



Руководство по настройке
Конфигурация надежности
Ethernet-коммутаторы ЦОД
серия QSW-8113



Оглавление

1. НАСТРОЙКА REUP	14
1.1. Обзор	14
1.2. Приложения	14
1.2.1. Связь в сети с двумя uplink-каналами	14
1.2.1.1. Сценарий	14
1.2.1.2. Развертывание	15
1.3. Функции	15
1.3.1. Базовые концепты	15
1.4. Обзор	16
1.4.1. Двухканальное резервирование REUP	16
1.4.1.1. Принцип работы	17
1.4.2. Связанная конфигурация	18
1.4.2.1. Включение двухканального резервирования на интерфейсе	18
1.4.3. Preemption mode (приоритетный режим) и время задержки REUP	19
1.4.3.1. Принцип работы	19
1.4.3.2. Связанная конфигурация	19
1.4.4. Обновление MAC-адреса	19
1.4.4.1. Принцип работы	19
1.4.4.2. Связанная конфигурация	20
1.4.5. Баланс нагрузки VLAN	21
1.4.5.1. Принцип работы	21
1.4.5.2. Связанная конфигурация	22
1.4.6. Отслеживание состояния канала	22
1.4.6.1. Принцип работы	23
1.4.6.2. Связанная конфигурация	23
1.5. Конфигурация	24
1.5.1. Настройка основных функций REUP	26
1.5.2. Эффект конфигурации	26
1.5.2.1. Шаги настройки	26
1.5.2.2. Проверка	26
1.5.2.3. Связанные команды	27
1.5.2.4. Пример конфигурации	27
1.5.2.5. Распространенные ошибки	29
1.5.3. Настройка preemption mode и функции задержки REUP	29
1.5.3.1. Эффект конфигурации	29
1.5.3.2. Примечания	29



1.5.3.3. Шаги настройки	29
1.5.3.4. Проверка	29
1.5.3.5. Связанные команды	29
1.5.3.6. Пример конфигурации	30
1.5.3.7. Распространенные ошибки	31
1.5.4. Настройка обновления MAC-адреса	31
1.5.4.1. Эффект конфигурации	31
1.5.4.2. Примечания	31
1.5.4.3. Шаги настройки	32
1.5.4.4. Проверка	32
1.5.4.5. Связанные команды	32
1.5.4.6. Пример конфигурации	34
1.5.4.7. Распространенные ошибки	36
1.5.5. Настройка балансировки нагрузки VLAN	36
1.5.5.1. Эффект конфигурации	36
1.5.5.2. Примечания	36
1.5.5.3. Шаги настройки	36
1.5.5.4. Проверка	36
1.5.5.5. Связанные команды	37
1.5.5.6. Пример конфигурации	37
1.5.5.7. Распространенные ошибки	38
1.5.6. Настройка отслеживания каналов	39
1.5.6.1. Эффект конфигурации	39
1.5.6.2. Примечания	39
1.5.6.3. Шаги настройки	39
1.5.6.4. Проверка	39
1.5.6.5. Связанные команды	39
1.5.6.6. Пример конфигурации	40
1.5.6.7. Распространенные ошибки	41
1.6. Мониторинг	41
1.6.1. Отображение	41
1.6.1.1. Отладка	42
2. НАСТРОЙКА RLDР	43
2.1. Обзор	43
2.2. Приложения	43
2.2.1. Обнаружение однонаправленной связи	43
2.2.1.1. Сценарий	43



2.2.1.2. Развертывание	43
2.2.2. Обнаружение двунаправленной пересылки	44
2.2.2.1. Сценарий	44
2.2.2.2. Развертывание	44
2.2.3. Обнаружение петель downlink-канала	44
2.2.3.1. Сценарий	44
2.2.3.2. Развертывание	45
2.3. Функции	45
2.3.1.1. Базовые концепты	45
2.3.2. Обзор	47
2.3.3. Развертывание обнаружения RLDP	47
2.3.3.1. Принцип работы	47
2.3.3.2. Связанная настройка	47
2.4. Конфигурация	48
2.4.1. Настройка основных функций RLDP	48
2.4.1.1. Эффект конфигурации	48
2.4.1.2. Примечания	48
2.4.1.3. Шаги настройки	49
2.4.1.4. Проверка	49
2.4.1.5. Связанные команды	50
2.4.1.6. Пример конфигурации	51
2.4.1.7. Распространенные ошибки	54
2.5. Мониторинг	54
2.5.1. Отображение	54
3. НАСТРОЙКА DLDP	55
3.1. Обзор	55
3.2. Приложения	55
3.2.1. Обнаружение внутрисетевого сегмента DLDP	55
3.2.1.1. Сценарий	55
3.2.1.2. Развертывание	55
3.2.2. Обнаружение межсетевого сегмента DLDP	56
3.2.2.1. Сценарий	56
3.2.2.2. Развертывание	56
3.3. Функции	56
3.3.1. Базовые концепты	56
3.3.2. Обзор	57
3.3.3. Обнаружение DLDP	57



3.3.3.1. Принцип работы	58
3.3.3.2. Связанная конфигурация	58
3.3.4. Привязка MAC-адреса	58
3.3.4.1. Принцип работы	58
3.3.4.2. Связанная конфигурация	58
3.3.5. Пассивное обнаружение DLDP	58
3.3.5.1. Принцип работы	58
3.3.5.2. Связанная конфигурация	58
3.4. Конфигурация	59
3.4.1. Включение обнаружения DLDP	59
3.4.1.1. Эффект конфигурации	59
3.4.1.2. Примечания	59
3.4.1.3. Шаги настройки	60
3.4.1.4. Проверка	60
3.4.1.5. Связанные команды	60
3.4.1.6. Пример конфигурации	62
3.4.1.7. Распространенные ошибки	63
3.5. Мониторинг	63
3.5.1. Очистка	63
3.5.2. Отображение	63
4. НАСТРОЙКА VRRP	64
4.1. Обзор	64
4.1.1. Протоколы и стандарты	64
4.2. Приложения	64
4.2.1. Избыточность маршрутизации	64
4.2.1.1. Сценарий	64
4.2.1.2. Развертывание	65
4.2.2. Балансировка нагрузки	65
4.2.2.1. Сценарий	65
4.2.2.2. Развертывание	66
4.3. Функции	66
4.3.1. Базовые концепты	66
4.3.2. Обзор	67
4.3.3. VRRP	67
4.3.3.1. Принцип работы	68
4.3.3.2. Связанная конфигурация	69
4.4. Конфигурация	72



4.4.1. Настройка IPv4 VRRP	74
4.4.1.1. Эффект конфигурации	74
4.4.1.2. Примечания	74
4.4.1.3. Шаги настройки	75
4.4.1.4. Проверка	76
4.4.1.5. Связанные команды	76
4.4.1.6. Пример конфигурации	83
4.4.1.7. Распространенные ошибки	86
4.4.1.8. Пример конфигурации	87
4.4.2. Настройка IPv6 VRRP	91
4.4.2.1. Эффект конфигурации	91
4.4.2.2. Примечания	91
4.4.2.3. Шаги настройки	91
4.4.2.4. Проверка	92
4.4.2.5. Связанные команды	92
4.4.2.6. Пример конфигурации	97
4.4.2.7. Пример конфигурации	99
4.4.3. Настройка VRRP-MSTP	102
4.4.3.1. Эффект конфигурации	102
4.4.3.2. Примечания	102
4.4.3.3. Шаги настройки	102
4.4.3.4. Проверка	103
4.4.3.5. Связанные команды	104
4.4.3.6. Пример конфигурации	110
4.5. Мониторинг	119
4.5.1. Отображение	119
4.5.1.1. Отладка	120
5. НАСТРОЙКА VRRP PLUS	121
5.1. Обзор	121
5.2. Приложения	121
5.2.1. Включение балансировки нагрузки в группе VRRP	121
5.2.1.1. Сценарий	121
5.2.1.2. Развертывание	123
5.3. Функции	124
5.3.1. Базовые понятия	124
5.3.2. Обзор	124
5.3.3. VRRP Plus	124



5.3.3.1. Основные принципы	124
5.4. Конфигурация	128
5.4.1. Настройка VRRP Plus	128
5.4.1.1. Эффект конфигурации	128
5.4.1.2. Примечания	128
5.4.1.3. Шаги настройки	129
5.4.1.4. Проверка	129
5.4.1.5. Связанные команды	129
5.4.1.6. Пример конфигурации	133
5.4.1.7. Распространенные ошибки	138
5.5. Мониторинг	138
5.5.1. Отображение	138
5.5.2. Отладка	139
6. НАСТРОЙКА BFD	140
6.1. Обзор	140
6.1.1. Протоколы и стандарты	140
6.2. Приложения	140
6.2.1. Поддержка BFD для OSPF	140
6.2.1.1. Сценарий	140
6.2.1.2. Развертывание	141
6.2.2. Поддержка BFD для статической маршрутизации	141
6.2.2.1. Сценарий	141
6.2.2.2. Развертывание	142
6.3. Функции	142
6.3.1. Базовые концепты	142
6.3.1.1. Обзор	145
6.3.2. Установление сеанса BFD	146
6.3.2.1. Принцип работы	146
6.3.3. Обнаружение сеанса BFD	147
6.3.3.1. Принцип работы	147
6.3.4. Поддержка BFD для приложений	148
6.3.4.1. Принцип работы	148
6.3.5. Защита BFD	151
6.3.5.1. Принцип работы	151
6.3.6. BFD Flapping Dampening	151
6.3.6.1. Принцип работы	151
6.4. Конфигурация	152



6.4.1. Настройка основных функций BFD	152
6.4.1.1. Эффект конфигурации	152
6.4.1.2. Примечания	152
6.4.1.3. Шаги настройки	153
6.4.1.4. Проверка	156
6.4.1.5. Пример конфигурации	157
6.4.1.6. Распространенные ошибки	159
6.4.2. Настройка защиты BFD	159
6.4.2.1. Эффект конфигурации	159
6.4.2.2. Примечания	159
6.4.2.3. Шаги настройки	160
6.4.2.4. Проверка	160
6.4.2.5. Пример конфигурации	160
6.4.3. Настройка BFD Flapping Dampening	160
6.4.3.1. Эффект конфигурации	160
6.4.3.2. Примечания	161
6.4.3.3. Шаги настройки	161
6.4.3.4. Проверка	161
6.4.3.5. Пример конфигурации	162
6.5. Мониторинг	162
6.5.1. Отображение	162
6.5.1.1. Отладка	162
7. НАСТРОЙКА ПОДАВЛЕНИЯ IP-СОБЫТИЙ	163
7.1. Обзор	163
7.1.1. Протоколы и стандарты	163
7.2. Приложение	163
7.2.1. Routed Port Flap Dampening	163
7.2.1.1. Сценарий	163
7.2.1.2. Развертывание	164
7.3. Функции	164
7.3.1. Базовые концепты	164
7.4. Обзор	165
7.4.1. Port Flap Suppression	165
7.4.1.1. Принцип работы	165
7.4.1.2. Связанная конфигурация	165
7.5. Конфигурация	165
7.5.1. Включение подавления IP-событий	166



7.5.1.1. Эффект конфигурации	166
7.5.1.2. Примечания	166
7.5.1.3. Шаги настройки	166
7.5.1.4. Проверка	166
7.5.1.5. Связанные команды	166
7.5.1.6. Пример конфигурации	167
7.5.1.7. Распространенные ошибки	168
7.6. Мониторинг	168
7.6.1. Очистка	168
7.6.2. Отображение	168
7.6.2.1. Отладка	168
8. НАСТРОЙКА VSU	169
8.1. Обзор	169
8.2. Приложения	170
8.2.1. Единое управление несколькими устройствами	170
8.2.1.1. Сценарий	170
8.2.1.2. Развертывание	171
8.2.2. Упрощение сетевой топологии	171
8.2.2.1. Сценарий	171
8.2.2.2. Развертывание	172
8.3. Функции	172
8.3.1.1. Базовые концепты	172
8.3.1.2. Обзор	174
8.3.2. Канал виртуальной коммутации (VSL)	175
8.3.2.1. Принцип работы	175
8.3.3. Топология	176
8.3.3.1. Принцип работы	176
8.3.4. Dual-Active Detection (DAD)	179
8.3.4.1. Принцип работы	179
8.3.5. Переадресация трафика VSU	181
8.3.5.1. Принцип работы	181
8.3.6. Управление системой	183
8.3.6.1. Принцип работы	183
8.3.7. Определение нахождения устройства быстрым миганием (Quick Blinking Location)	184
8.3.8. Восстановление устройства в режиме восстановления	184
8.3.9. Автоматическое восстановление без перезагрузки для устройства в режиме восстановления при неисправности Master-устройства	184



8.4. Конфигурация	185
8.4.1. Настройка VSU в автономном режиме	188
8.4.1.1. Эффект конфигурации	188
8.4.1.2. Шаги настройки	188
8.4.1.3. Проверка	193
8.4.1.4. Пример конфигурации	193
8.4.1.5. Распространенные ошибки	195
8.4.2. Настройка VSU в режиме VSU	196
8.4.2.1. Настройка атрибутов VSU	196
8.4.2.2. Эффект конфигурации	196
8.4.2.3. Примечания	196
8.4.2.4. Шаги настройки	196
8.4.2.5. Проверка	199
8.4.2.6. Пример конфигурации	199
8.4.2.7. Настройка VSL	201
8.4.2.8. Эффект конфигурации	201
8.4.2.9. Примечания	201
8.4.2.10. Шаги настройки	201
8.4.2.11. Проверка	202
8.4.2.12. Пример конфигурации	203
8.4.2.13. Настройка Dual-Active Detection	204
8.4.2.14. Эффект конфигурации	204
8.4.2.15. Примечания	204
8.4.2.16. Шаги настройки	204
8.4.2.17. Проверка	207
8.4.2.18. Пример конфигурации	208
8.4.2.19. Распространенные ошибки	210
8.4.2.20. Настройка балансировки трафика	210
8.4.2.21. Эффект конфигурации	210
8.4.2.22. Примечания	210
8.4.2.23. Шаги настройки	210
8.4.2.24. Проверка	211
8.4.2.25. Пример конфигурации	211
8.4.2.26. Изменение режима VSU на автономный режим	212
8.4.2.27. Эффект конфигурации	212
8.4.2.28. Шаги настройки	212
8.4.2.29. Пример конфигурации	213
8.4.3. Настройка Определения нахождения устройства быстрым миганием	214



8.4.3.1. Эффект конфигурации	214
8.4.3.2. Примечания	214
8.4.3.3. Шаги настройки	215
8.4.3.4. Проверка	215
8.4.3.5. Пример конфигурации	216
8.4.4. Настройка интерфейса MGMT	216
8.4.4.1. Эффект конфигурации	216
8.4.4.2. Примечания	216
8.4.4.3. Шаги настройки	216
8.4.4.4. Проверка	217
8.4.4.5. Пример конфигурации	217
8.4.5. Настройка восстановления устройства в режиме восстановления	218
8.4.5.1. Эффект конфигурации	218
8.4.5.2. Примечания	218
8.4.5.3. Шаги настройки	218
8.4.5.4. Проверка	219
8.4.5.5. Пример конфигурации	219
8.4.6. Настройка автоматического восстановления без перезагрузки в режиме восстановления	219
8.4.6.1. Эффект конфигурации	219
8.4.6.2. Шаги настройки	220
8.4.6.3. Проверка	220
8.4.6.4. Пример конфигурации	220
8.5. Мониторинг и обслуживание	221
8.5.1. Отображение	221
9. НАСТРОЙКА RNS	222
9.1. Обзор	222
9.2. Приложение	222
9.2.1. Тестирование и оценка эффективности службы	222
9.2.1.1. Сценарий	222
9.2.1.2. Развертывание	222
9.2.2. Обнаружение сетевых сбоев	223
9.2.2.1. Сценарий	223
9.2.2.2. Развертывание	223
9.3. Функции	224
9.3.1. Базовые концепты	224
9.3.2. Тест RNS	224
9.3.2.1. Принцип работы	224



9.3.2.2. Связанная конфигурация	225
9.3.3. Отслеживание поддержки RNS	226
9.3.3.1. Принцип работы	226
9.3.3.2. Связанная конфигурация	226
9.4. Конфигурация	227
9.4.1. Настройка основных функций RNS	229
9.4.1.1. Эффект конфигурации	229
9.4.1.2. Примечания	229
9.4.1.3. Шаги настройки	230
9.4.1.4. Проверка	230
9.4.1.5. Связанные команды	231
9.4.1.6. Пример конфигурации	234
9.4.2. Настройка эхо-теста ICMP	235
9.4.2.1. Эффект конфигурации	235
9.4.2.2. Примечания	235
9.4.2.3. Шаги настройки	235
9.4.2.4. Проверка	236
9.4.2.5. Связанные команды	236
9.4.3. Настройка теста DNS	240
9.4.3.1. Эффект конфигурации	240
9.4.3.2. Примечания	240
9.4.3.3. Шаги настройки	240
9.4.3.4. Проверка	240
9.4.3.5. Связанные команды	240
9.4.3.6. Пример конфигурации	243
9.4.3.7. Распространенные ошибки	244
9.4.4. Настройка поддержки отслеживания для RNS	244
9.4.4.1. Эффект конфигурации	244
9.4.4.2. Примечания	244
9.4.4.3. Шаги настройки	244
9.4.4.4. Проверка	245
9.4.4.5. Связанные команды	245
9.4.4.6. Пример конфигурации	247
9.4.4.7. Распространенные ошибки	250
9.5. Мониторинг	251
9.5.1. Отображение	251
9.5.1.1. Отладка	251



10. ОБЩАЯ ИНФОРМАЦИЯ	253
10.1. Гарантия и сервис	253
10.2. Техническая поддержка	253
10.3. Электронная версия документа	253



1. НАСТРОЙКА REUP

1.1. Обзор

Протокол Rapid Ethernet Uplink Protection Protocol (REUP) обеспечивает функцию быстрой защиты uplink-канала.

В сети с двумя uplink-каналами REUP используется для обеспечения нормальной связи между каналами, блокировки избыточных каналов, предотвращения образования петель и быстрого резервного копирования.

Upstream-интерфейсы REUP настраиваются парами. Если оба интерфейса в норме, интерфейс работает в резервном состоянии. Интерфейс в резервном состоянии не пересылает пакеты данных. Когда пересылающий интерфейс неисправен, резервный интерфейс немедленно переключается в состояние пересылки и обеспечивает передачу данных. Кроме того, REUP также отправляет пакеты обновления адреса upstream-устройствам, чтобы upstream-устройства могли немедленно обновить свои MAC-адреса. Эта функция REUP гарантирует, что потоки данных уровня 2 могут быть восстановлены в течение 50 мс после отказа канала.

REUP является взаимоисключающим со Spanning Tree Protocol (STP) на основе интерфейсов. В этом случае устройство запускает STP ниже и запускает REUP поверх, чтобы реализовать резервирование и защиту от сбоев для upstream-канала. REUP обеспечивает базовую избыточность канала, когда STP отключен, а также обеспечивает восстановление после сбоя быстрее на уровне миллисекунд, чем STP.

Протоколы и стандарты

REUP — это проприетарный протокол QTECH Network, и для справки не существует стандарта или протокола.

1.2. Приложения

Приложение	Описание
<u>Связь в сети с двумя uplink-каналами</u>	Пересылка пакетов в сети с двумя uplink-каналами

1.2.1. Связь в сети с двумя uplink-каналами

1.2.1.1. Сценарий

Для связи в сети с двойным uplink-каналом коммутатор доступа имеет два пути upstream-канала, как показано на Рисунке 1-1.



Рисунок 1-1. Сеть с двумя uplink-каналами

1.2.1.2. Развертывание

- Включите REUP на интерфейсе 1 и интерфейсе 2 коммутатора доступа D/E, чтобы реализовать быстрое переключение при сбое канала.
- Включите получение сообщения об обновлении MAC-адреса REUP на интерфейсах, подключенных к коммутаторам A/B/C, чтобы быстро очищать MAC-адреса на интерфейсах при сбое канала.

1.3. Функции

1.3.1. Базовые концепты

Пара REUP

Укажите интерфейс в качестве резервного интерфейса другого интерфейса для настройки пары REUP. Один интерфейс является активным интерфейсом, а другой интерфейс является резервным интерфейсом. Когда два интерфейса работают нормально, один интерфейс настраивается как интерфейс пересылки, тогда как другой интерфейс настраивается как резервный интерфейс. Вы можете определить интерфейс, который будет настроен в качестве резервного интерфейса. См. соответствующую информацию в разделе [Preemption mode \(приоритетный режим\) и время задержки REUP](#).



Сообщение об обновлении MAC-адреса

Сообщения об обновлении MAC-адреса относятся к пакетам FLUSH, отправляемым сетью QTECH на uplink-устройства через приватную многоадресную рассылку. Когда в uplink-устройство QTECH Network включает функцию получения сообщений об обновлении MAC-адреса и получает сообщения об обновлении MAC-адреса, устройство обновляет MAC-адреса соответствующих интерфейсов.

Группа обновления MAC-адреса

Несколько интерфейсов добавляются в группу. Если один интерфейс в группе получает сообщение об обновлении MAC-адреса, MAC-адреса других интерфейсов в группе будут обновлены. В этом случае группа называется группой обновления MAC-адреса.

Пакет обновления MAC-адреса

Пакеты, отправляемые для обновления MAC-адресов для поддержки uplink-устройств, называются пакетами обновления MAC-адресов.

Link Tracking Group (группа отслеживания канала)

Uplink- и downstream-интерфейсы устройства добавляются в группу. Если все upstream-интерфейсы в группе отключены, все downstream-интерфейсы в этой группе принудительно отключаются. В этом случае эта группа называется Link Tracking Group.

1.4. Обзор

Особенность	Описание
<u>Двухканальное резервирование REUP</u>	Когда канал неисправен, другой канал может быстро переключиться в состояние передачи
<u>Preemption mode (приоритетный режим) и время задержки REUP</u>	Когда оба канала работают нормально, preemption mode можно использовать для определения канала, используемого для пересылки данных, и времени задержки, используемого для определения времени ожидания перед переключением
<u>Обновление MAC-адреса</u>	Во время переключения каналов MAC-адрес интерфейса обновляется, чтобы ускорить сходимость пакетов
<u>Баланс нагрузки VLAN</u>	Когда два канала работают нормально, использование пропускной способности канала может быть максимальным
<u>Отслеживание состояния канала</u>	Когда upstream-канал неисправен, downstream-канал переключается

1.4.1. Двухканальное резервирование REUP

Когда активный канал неисправен, канал в резервном состоянии быстро переключается в состояние передачи и начинает пересылать данные, сводя к минимуму прерывание обслуживания, вызванное сбоем канала.

1.4.1.1. Принцип работы

Укажите интерфейс в качестве резервного интерфейса другого интерфейса для настройки пары REUP. Когда два интерфейса работают нормально, канал находится в состоянии пересылки (пересылка пакетов данных), а другой канал находится в резервном состоянии (не пересылает данные). Когда активный канал неисправен, канал в резервном состоянии быстро переключается в состояние передачи и начинает пересылку данных. Когда неисправный канал восстанавливается, он переходит в резервное состояние и не пересылает пакеты данных. Конечно, вы можете настроить preemption mode, чтобы указать, будет ли канал, восстановленный после сбоя, вытеснять канал, который в данный момент находится в состоянии передачи.

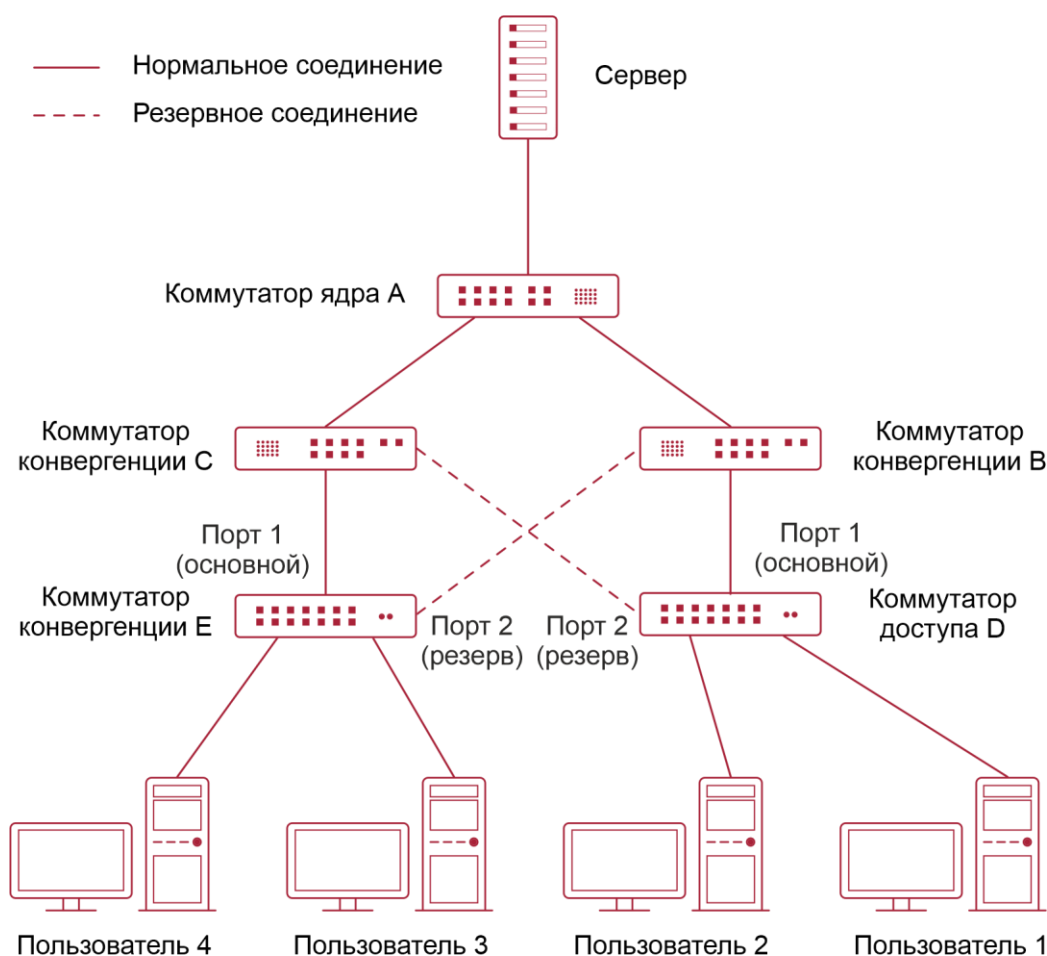


Рисунок 1-2. Топология с двумя нормальными каналами

Как показано на Рисунке 1-2, соедините интерфейсы 1 и 2 коммутатора D (E) с коммутаторами uplink B и C (C и B) и настройте REUP на интерфейсах 1 и 2. Когда каналы работают нормально, интерфейс 1 находится в состоянии пересылки и пересылает пакеты данных, а интерфейс 2 находится в резервном состоянии и не пересылает пакеты данных.



Рисунок 1-3. Топология с интерфейсом 1 коммутатора D (Е) неисправна

Как только интерфейс 1 выходит из строя, интерфейс 2 немедленно начинает пересылать пакеты данных и восстанавливает uplink-передачу коммутатора. В режиме без вытеснения, когда канал интерфейса 1 восстанавливается, интерфейс 1 находится в резервном состоянии и не пересылает пакеты данных, в то время как интерфейс 2 продолжает пересылать пакеты данных.

1.4.2. Связанная конфигурация

1.4.2.1. Включение двухканального резервирования на интерфейсе

По умолчанию двухканальное резервирование на интерфейсе отключено.

Вы можете запустить команду **switchport backup interface**, чтобы настроить физический интерфейс уровня 2 (или интерфейс AP уровня 2) в качестве интерфейса резервного копирования и включить функцию двухканального резервирования REUP.

Вы должны включить функцию двухканального резервирования REUP на интерфейсе. Функция включает переключение канала REUP только при неисправности интерфейса.

ПРИМЕЧАНИЕ: REUP, ERPS и RERP не используют общие интерфейсы.

ПРИМЕЧАНИЕ: устройства, поддерживающие REUP, должны отключить функцию storm control всех интерфейсов уровня 2.



1.4.3. Preemption mode (приоритетный режим) и время задержки REUP

1.4.3.1. Принцип работы

Вы можете определить, какая ссылка должна использоваться первой, настроив **preemption mode REUP**. Если для **preemption mode** сначала задана пропускная способность, REUP сначала выбирает канал с высокой пропускной способностью. Вы также можете установить **preemption mode**, чтобы сначала принудительно выбрать стабильную и надежную ссылку.

Чтобы избежать частого переключения активного/резервного канала, вызванного аномальными отказами, в REUP предусмотрена функция задержки прерывания. Когда два канала восстановлены, переключение каналов выполняется, когда неисправный канал становится стабильным после задержки (по умолчанию 35 секунд).

1.4.3.2. Связанная конфигурация

Настройка **preemption mode** и времени задержки REUP

По умолчанию **preemption mode** отключен, а время задержки составляет 35 секунд.

Вы можете запустить команду **switchport backup interface preemption mode**, чтобы настроить **preemption mode**.

Вы можете запустить команду **switchport backup interface preemption delay**, чтобы настроить время задержки.

Меньшая задержка означает более частое упреждающее переключение после восстановления неисправного канала.

ПРИМЕЧАНИЕ: REUP использует значение атрибута **Bandwidth** для интерфейса AP как фактическую пропускную способность интерфейса AP, которая равна значению атрибута **Speed** (количество подключенных интерфейсов-участников $\times$ количество интерфейсов-участников).

ПРИМЕЧАНИЕ: когда uplink-линия связи включает STP, время задержки **preemption mode REUP** превышает 35 секунд.

1.4.4. Обновление MAC-адреса

Во время переключения каналов MAC-адрес интерфейса обновляется, чтобы ускорить сходимость пакетов.

1.4.4.1. Принцип работы

Как показано на Рисунке 1-2, интерфейс 1 и интерфейс 2 коммутатора D (E) включены с двухканальным резервированием REUP. Интерфейс 1 работает как активный интерфейс. Во время обычной связи коммутатор A узнает MAC-адреса пользователей 1 и 2 (пользователей 3 и 4) из интерфейсов, подключенных к коммутатору B (C).

Когда интерфейс 1 коммутатора D (E) неисправен, интерфейс 2 быстро переключается в состояние пересылки и начинает пересылать пакеты данных. В этом случае коммутатор A не узнает MAC-адреса пользователей 1 и 2 (пользователей 3 и 4) на интерфейсах, подключающихся к коммутатору B (C). Пакеты данных, отправляемые сервером пользователям 1 и 2 (пользователям 3 и 4), пересылаются на коммутатор C (B) коммутатором A, в результате чего пакеты от сервера пользователям 1 и 2 (пользователям 3 и 4) теряются.

Чтобы избежать описанных выше проблем, вы можете включить функцию обновления MAC-адреса на коммутаторе D (E). Когда интерфейс 2 начинает пересылать пакеты, коммутатор D (E) отправляет сообщение об обновлении MAC-адреса на интерфейс 2.



Получив сообщение об обновлении MAC-адреса, коммутатор A обновляет MAC-адрес на интерфейсе коммутатора A. Таким образом, коммутатор A перенаправляет пакеты, отправленные сервером пользователям, на интерфейсы коммутатора B (C), чтобы ускорить конвергенцию пакетов.

Кроме того, импортируйте настройку группы обновления MAC-адресов, то есть классифицируйте несколько интерфейсов в одну группу. Когда интерфейс в этой группе получает сообщение об обновлении MAC-адреса, MAC-адреса на других интерфейсах в группе обновляются, чтобы уменьшить побочный эффект лавинной рассылки, вызванный обновлением MAC-адреса.

Чтобы быть совместимым с upstream-устройствами, не поддерживающими сообщения обновления MAC-адреса, коммутатор D (E) будет отправлять пакеты обновления MAC-адреса для пользователей 1 и 2 (пользователей 3 и 4) вверх, когда интерфейс 2 переключается в состояние пересылки. Таким образом, коммутатор A может обновить MAC-адреса пользователей 1 и 2 (пользователей 3 и 4) на соответствующие интерфейсы и восстановить передачу данных по downlink-каналу коммутатора A.

1.4.4.2. Связанная конфигурация

Включение отправки сообщений об обновлении MAC-адреса на интерфейсе

По умолчанию отправка сообщений об обновлении MAC-адреса на интерфейсе отключена.

Вы можете запустить команду **mac-address-table move update transit**, чтобы включить отставку обновлений MAC-адресов на все интерфейсы устройства.

Если отправка сообщений об обновлении MAC-адреса не включена, сообщения об обновлении MAC-адреса не будут отправляться при выполнении переключения двухканального резервирования REUP.

Включение получения сообщений об обновлении MAC-адреса на интерфейсе

По умолчанию получение сообщений об обновлении MAC-адреса на интерфейсе отключено.

Вы можете запустить команду **mac-address-table move update receive**, чтобы включить получение обновлений MAC-адреса на всех интерфейсах устройства.

Если получение сообщений об обновлении MAC-адреса не включено, устройство не может получать сообщения об обновлении MAC-адреса от downlink-устройств во время переключения двухканального резервирования REUP и не будет обновлять MAC-адреса.

Настройка VLAN для отправки сообщений об обновлении MAC-адреса

По умолчанию VLAN для отправки сообщений об обновлении MAC-адреса является VLAN по умолчанию, к которой принадлежит интерфейс.

Вы можете запустить команду **mac-address-table move update transit vlan**, чтобы настроить VLAN, в которой интерфейсы отправляют сообщения об обновлении MAC-адресов.

Если настроена VLAN, в которой интерфейсы отправляют сообщения обновления MAC-адреса, сообщения отправляются в настроенной VLAN; в противном случае сообщения отправляются в VLAN по умолчанию, к которой принадлежит интерфейс.

Настройка VLAN для получения сообщений об обновлении MAC-адреса

По умолчанию сообщения об обновлении MAC-адреса принимаются во всех VLAN.



Вы можете запустить команду **no mac-address-table move update receive vlan**, чтобы настроить VLAN, в которой интерфейсы не получают сообщения об обновлении MAC-адресов. Сообщения об обновлении MAC-адреса принимаются в оставшихся VLAN.

Если не настроена VLAN, в которой интерфейсы получают сообщения об обновлении MAC-адреса, сообщения об обновлении MAC-адреса принимаются во всех настроенных VLAN; в противном случае сообщения об обновлении MAC-адреса принимаются в остальных VLAN.

Настройка группы обновления MAC-адресов

По умолчанию группа обновления MAC-адресов отсутствует.

Вы можете запустить команду **mac-address-table update group**, чтобы добавить интерфейс в группу обновления MAC-адресов. По умолчанию интерфейс добавляется в первую группу обновлений.

Если группа обновления MAC-адресов не настроена, обновление MAC-адреса не будет выполняться при получении пакетов обновления MAC-адреса.

Настройка максимального количества пакетов обновления MAC-адреса, отправляемых в секунду

По умолчанию максимальное количество пакетов обновления MAC-адреса, отправляемых в секунду, равно 150.

Вы можете запустить команду **mac-address-table move update max-update-rate**, чтобы настроить максимальное количество пакетов обновления MAC-адреса, отправляемых в секунду.

Чем больше количество пакетов, тем больше процессорного времени используется для отправки пакетов и тем меньше downlink-пакетов теряется.

1.4.5. Баланс нагрузки VLAN

1.4.5.1. Принцип работы

Функция балансировки нагрузки VLAN позволяет REUP пересылать пакеты данных взаимоисключающих VLAN для двух интерфейсов, чтобы полностью использовать пропускную способность канала.

Как показано на Рисунке 1-4, настроить двухканальное резервирование REUP и включить балансировку нагрузки VLAN REUP на интерфейсе 1 и интерфейсе 2 коммутатора D, а также сопоставить VLAN 1 с экземпляром 1 и VLAN 2 с экземпляром 2. Данные VLAN 1 (экземпляр 1) передаются через интерфейс 1, а все остальные данные VLAN 2 (экземпляр 2) передаются через интерфейс 2. Выполните ту же операцию на коммутаторе E.

Когда интерфейс неисправен, другой интерфейс берет на себя передачу на всех VLAN. Когда неисправный интерфейс восстановлен и не становится неисправным в течение задержки preemption mode, передача VLAN переключается обратно на восстановленный интерфейс.

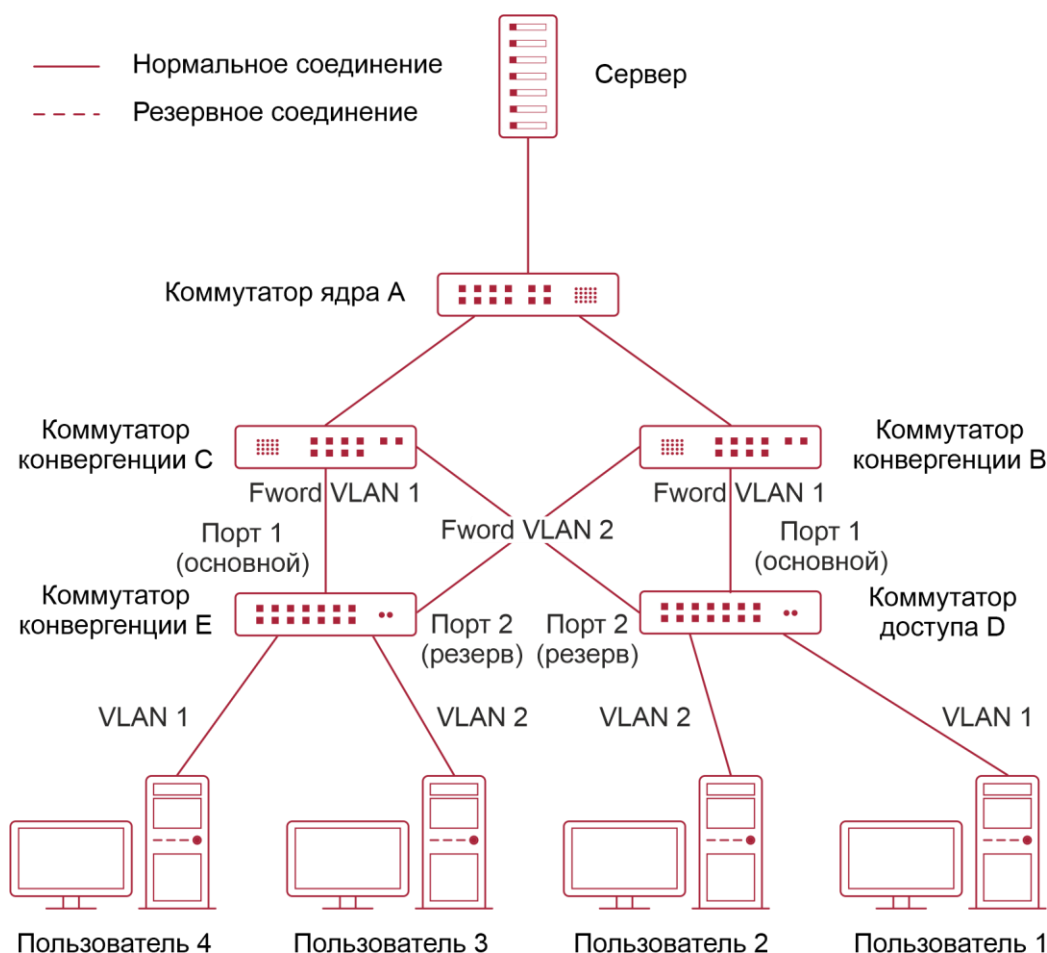


Рисунок 1-4. Топология с двумя нормальными каналами балансировки нагрузки

1.4.5.2. Связанная конфигурация

Включение балансировки нагрузки VLAN на интерфейсе

По умолчанию функция балансировки нагрузки VLAN на интерфейсе отключена.

Вы можете запустить команду **switchport backup interface prefer instance**, чтобы включить функцию балансировки нагрузки VLAN.

Если эта функция не включена, полоса пропускания канала не может быть полностью использована при пересылке пакетов, когда два канала работают нормально. Вы должны включить функцию балансировки нагрузки VLAN на порту, чтобы интерфейс мог участвовать в балансировке нагрузки VLAN.

ПРИМЕЧАНИЕ: отображение экземпляров баланса нагрузки REUP VLAN управляется модулем MSTP унифицированным образом. Дополнительные сведения о настройке экземпляров см. в описании в разделе Ethernet Switching/Настройка MSTP.

ПРИМЕЧАНИЕ: функцию балансировки нагрузки VLAN можно настроить только на trunk-, uplink- или гибридных интерфейсах.

1.4.6. Отслеживание состояния канала

Отслеживание канала означает, что при отказе upstream-канала службы переключаются на downstream-канал, чтобы резервный интерфейс мог продолжать пересылать пакеты.



1.4.6.1. Принцип работы

Отслеживание состояния канала обеспечивает функцию уведомления downlink-устройств для переключения канала, когда upstream-канал неисправен. Вы можете настроить uplink- и downstream-интерфейсы группы отслеживания состояния канала и привязать состояние канала нескольких downlink-интерфейсов к интерфейсам нескольких upstream-каналов для реализации синхронизации состояния канала. Когда все upstream-каналы в группе отслеживания неисправны, интерфейсы downstream-каналов принудительно отключаются, чтобы обеспечить переключение передачи downstream-каналов с активного канала на резервный.

Как показано на Рисунке 1-5, когда upstream-канал коммутатора В неисправен, отслеживание состояния канала быстро отключает downstream-интерфейс коммутатора В, так что передача по uplink-каналу коммутатора D переключается на коммутатор С.

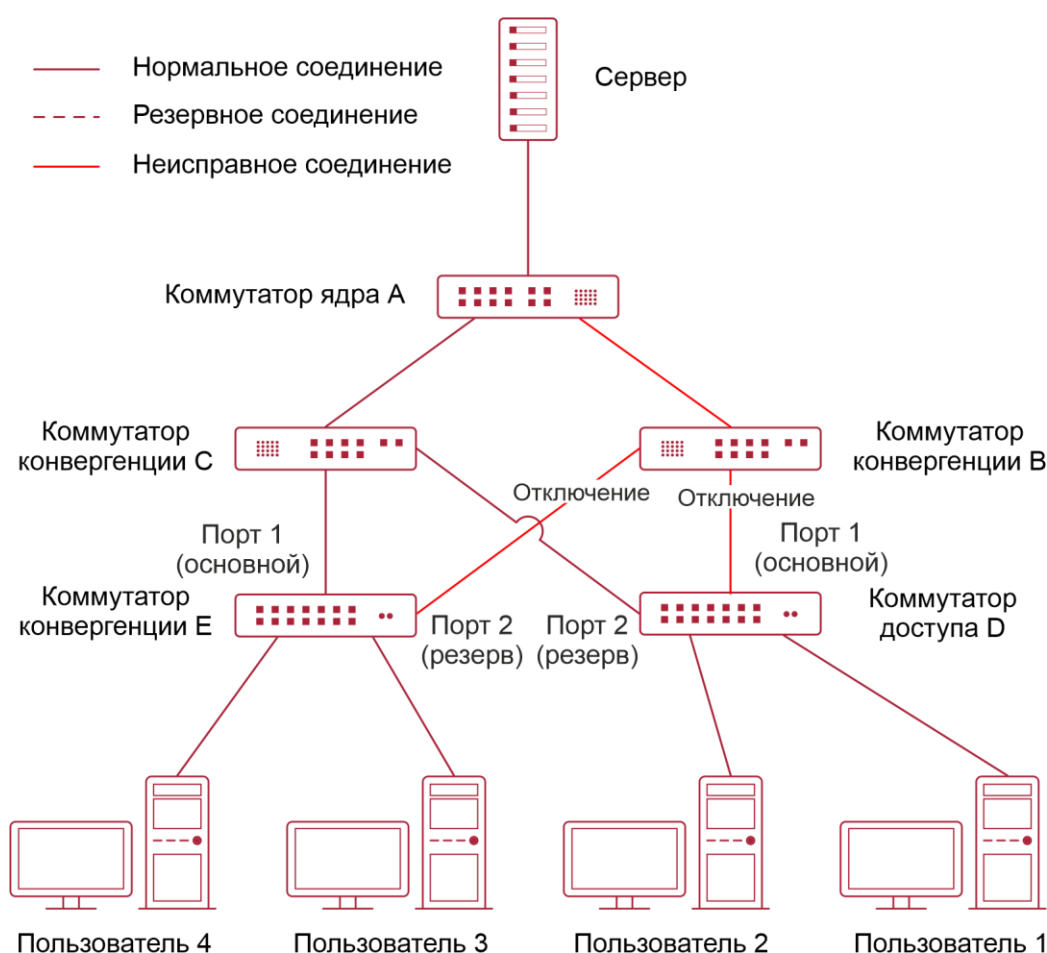


Рисунок 1-5. Топология, в которой upstream-канал активного канала неисправен

1.4.6.2. Связанная конфигурация

Включение отслеживания каналов

Отслеживание каналов по умолчанию отключено.

Вы можете запустить команду **link state track [number]**, чтобы включить Link Tracking Group. Значение **number** варьируется от 1 до 2. Первая Link Tracking Group включена по умолчанию (значение числа по умолчанию равно 1).



Если отслеживание каналов не включено, статус соответствующего upstream-интерфейса не может быть обнаружен, и переключение переадресации пакетов не может быть реализовано вовремя.

Включение функции задержки downlink-канала для Link Tracking Group

По умолчанию задержка downlink-канала для отслеживания канала равна 0 секунд.

Вы можете запустить команду **link state track number up-delay timer**, чтобы включить Link Tracking Group. Значение **number** варьируется от 1 до 2. Первая Link Tracking Group включена по умолчанию (значение числа по умолчанию равно 1). Значение таймера находится в диапазоне от 0 до 300 секунд, что по умолчанию равно 0 секунд.

Включив функцию увеличения задержки downlink-канала, вы можете избежать частых переключений downlink-канала, вызванных нестабильностью upstream-канала в Link Tracking Group. То есть, когда upstream-канал становится активным, downstream-канал становится активным после задержки.

Добавление интерфейса в Link Tracking Group

По умолчанию интерфейс не добавляется в Link Tracking Group.

Вы можете запустить команду **link state group [number] {upstream | downstream}** для настройки upstream- и downstream-интерфейсов Link Tracking Group. Значение **number** находится в диапазоне от 1 до 2. Интерфейс добавляется в первую Link Tracking Group по умолчанию (значение числа по умолчанию равно 1).

Если интерфейс не добавлен в группу отслеживания, статус соответствующего upstream-интерфейса не может быть обнаружен, и переключение переадресации пакетов не может быть реализовано вовремя.

1.5. Конфигурация

Конфигурация	Описание и команда	
Настройка основных функций REUP	(Обязательно) Используется для включения двухканального резервирования REUP	
	switchport backup interface	Включает двухканальное резервирование REUP
Настройка preemption mode и функции задержки REUP	(Опционально) Используется для определения preemption mode и времени задержки. Значения по умолчанию используются, если они не настроены	
	switchport backup interface preemption mode	Устанавливает preemption mode
	switchport backup interface preemption delay	Устанавливает время задержки для preemption mode



Конфигурация	Описание и команда	
<u>Настройка обновления MAC-адреса</u>	(Опционально) Используется для быстрого обновления MAC-адресов	
	mac-address-table update group	Устанавливает идентификатор группы обновления MAC-адреса коммутатора
	mac-address-table move update transit	Включает отправку сообщений об обновлении MAC-адреса
	mac-address-table move update transit vlan	Включает отправку идентификатора VLAN сообщений обновления MAC-адреса
	mac-address-table move update	Настраивает максимальное количество пакетов обновления MAC-адреса, отправляемых в секунду. Значение находится в диапазоне от 0 до 32 000. Значение по умолчанию — 150
	mac-address-table move update receive	Включает получение сообщений об обновлении MAC-адреса
	mac-address-table move update receive vlan	Настраивает диапазон VLAN для обработки сообщений об обновлении MAC-адреса
<u>Настройка балансировки нагрузки VLAN</u>	(Опционально) Используется для включения балансировки нагрузки VLAN	
	switchport backup interface prefer instance	Настраивает баланс нагрузки канала VLAN для REUP
<u>Настройка отслеживания</u>	(Опционально) Используется для включения отслеживания каналов	
	link state track up-delay	Включает задержку downlink-канала для группы отслеживания состояния канала



Конфигурация	Описание и команда	
	link state track	Включает отслеживание каналов группу состояний
<u>Настройка отслеживания</u>	link state group	Добавьте интерфейс в качестве upstream- или downstream-интерфейса указанной группы отслеживания состояния канала

1.5.1. Настройка основных функций REUP

1.5.2. Эффект конфигурации

- Когда канал неисправен, другой нормальный канал немедленно переключается в состояние пересылки для пересылки пакетов.
- Интерфейс принадлежит только одной паре REUP. Каждый активный канал имеет только один резервный канал. Резервный канал может использоваться как резервный увнвл только для одного активного канала. Активные и резервные каналы должны использовать разные интерфейсы.
- REUP поддерживает физические интерфейсы уровня 2 и интерфейсы AP, но не поддерживает интерфейсы участников AP.
- Активный и резервный интерфейсы могут быть разных типов и иметь разную скорость. Например, интерфейс AP может использоваться как активный интерфейс, тогда как физический интерфейс настроен как резервный интерфейс.
- Интерфейсы, сконфигурированные с REUP, не участвуют в расчете STP.
- Каждое устройство может быть настроено максимум на 16 пар REUP.
- Интерфейсы, успешно настроенные с помощью REUP, не могут изменить интерфейсы на интерфейсы уровня 3 или быть добавлены к AP.

1.5.2.1. Шаги настройки

Включение двухканального резервирования REUP

- Обязательный.
- Если нет особых требований, следует включить двухканальное резервирование REUP на интерфейсе принимающего коммутатора.

1.5.2.2. Проверка

Запустите команду **show interfaces switchport backup [detail]**, чтобы проверить, настроено ли двухканальное резервирование REUP.



1.5.2.3. Связанные команды

Включение двухканального резервирования REUP

Команда	<code>switchport backup interface <i>interface-id</i></code>
Описание параметров	<i>interface-id</i> : указывает идентификатор резервного интерфейса
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если интерфейс, в котором находится режим, является активным интерфейсом, интерфейс, соответствующий параметру <i>interface-id</i> , является резервным интерфейсом. Если активный канал неисправен, быстро восстанавливается передача по резервному каналу

1.5.2.4. Пример конфигурации

Включение двухканального резервирования REUP

Сценарий:

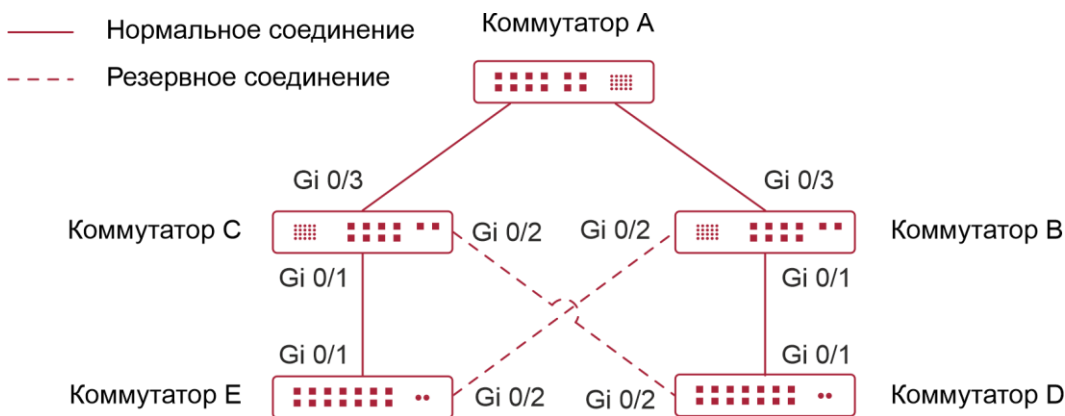


Рисунок 1-6. Двойная uplink-сеть

	Как показано на Рисунке 1-6, от коммутатора D к коммутатору А ведут два upstream-канала: коммутатор D > коммутатор В > коммутатор А и коммутатор D > коммутатор С > коммутатор А. Есть два upstream-канала от коммутатора Е к коммутатору А: коммутатор Е > коммутатор В > коммутатор А и коммутатор Е > коммутатор С > коммутатор А
Шаги настройки	Настройте двухканальное резервирование (интерфейс Gi0/1 — активный интерфейс, а Gi0/2 — резервный интерфейс) REUP на коммутаторе доступа D (Е)
D	SwitchD> enable



	<pre>SwitchD# configure terminal SwitchD(config)# interface GigabitEthernet 0/1 SwitchD(config-if-GigabitEthernet 0/1)# switchport mode trunk SwitchD(config-if-GigabitEthernet 0/1)#switchport backup interface GigabitEthernet 0/2 SwitchD(config-if-GigabitEthernet 0/1)# exit</pre>
E	<pre>SwitchE> enable SwitchE# configure terminal SwitchE(config)# interface GigabitEthernet 0/1 SwitchE(config-if-GigabitEthernet 0/1)# switchport mode trunk SwitchE(config-if-GigabitEthernet 0/1)#switchport backup interface GigabitEthernet 0/2 SwitchD(config-if-GigabitEthernet 0/1)# exit</pre>
Проверка	Проверьте информацию о двухканальном резервировании, настроенную для коммутатора D (E)
D	<pre>SwitchD#show interfaces switchport backup detail Switch Backup Interface Pairs: Active Interface Backup Interface State ----- Gi0/1 Gi0/2 Active Up/Backup Standby Interface Pair : Gi0/1, Gi0/2 Preemption Mode : off Preemption Delay : 35 seconds Bandwidth : Gi0/1(100000 Mbits), Gi0/2(100000 Mbits)</pre>
E	<pre>SwitchE#show interfaces switchport backup detail Switch Backup Interface Pairs: Active Interface Backup Interface State ----- Gi0/1 Gi0/2 Active Up/Backup Standby Interface Pair : Gi0/1, Gi0/2 Preemption Mode : off Preemption Delay : 35 seconds Bandwidth : Gi0/1(100000 Mbits), Gi0/2(100000 Mbits)</pre>



1.5.2.5. Распространенные ошибки

- Другие пары REUP настраиваются на сконфигурированном интерфейсе.
- Сконфигурированный интерфейс не является физическим интерфейсом уровня 2 или интерфейсом AP.

1.5.3. Настройка preemption mode и функции задержки REUP

1.5.3.1. Эффект конфигурации

Ограничьте preemption mode и время задержки для переключения канала REUP.

1.5.3.2. Примечания

Необходимо настроить двухканальное резервирование REUP.

1.5.3.3. Шаги настройки

- Опционально.
- Если активный канал должен всегда пересылать пакеты или пропускная способность канала должна использоваться для определения канала для пересылки пакетов, необходимо настроить соответствующий preemption mode и время задержки.

1.5.3.4. Проверка

Запустите команду **show interfaces switchport backup [detail]**, чтобы проверить, соответствуют ли preemption mode и время задержки конфигурациям.

1.5.3.5. Связанные команды

Настройка preemption mode REUP

Команда	<code>switchport backup interface interface-id preemption mode {forced bandwidth off}</code>
Описание параметров	<i>interface-id</i> : указывает идентификатор резервного интерфейса. mode : устанавливает preemption mode: forced : указывает форсированный режим. bandwidth : указывает режим полосы пропускания. off : указывает, что preemption mode выключен
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Preemption mode включает форсированный режим, режим полосы пропускания и выключение. В режиме полосы пропускания сначала выбирается интерфейс с высокой пропускной способностью для передачи данных; в форсированном режиме активный интерфейс выбирается первым для передачи данных; в выключенном режиме preemption (приоритезация) не выполняется. Режим по умолчанию выключен



Настройка времени задержки REUP

Команда	switchport backup interface <i>interface-id</i> preemption delay <i>delay-time</i>
Описание параметров	<i>interface-id</i> : указывает идентификатор резервного интерфейса. <i>delay-time</i> : указывает время задержки
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Задержка preemption указывает время задержки после восстановления неисправного канала до времени, когда снова выполняется переключение каналов

1.5.3.6. Пример конфигурации

Настройка **preemption mode** и времени задержки REUP

Сценарий	Как показано на Рисунке 1-6, от коммутатора D к коммутатору A ведут два upstream-канала: коммутатор D > коммутатор B > коммутатор A и коммутатор D > коммутатор C > коммутатор A. Есть два upstream-канала от коммутатора E к коммутатору A: коммутатор E > коммутатор B > коммутатор A и коммутатор E > коммутатор C > коммутатор A
Шаги настройки	Настройте preemption mode на полосу пропускания на коммутаторе доступа D (E) и время задержки на 40 секунд
D	<pre>SwitchD> enable SwitchD# configure terminal SwitchD(config)# interface GigabitEthernet 0/1 SwitchD(config-if-GigabitEthernet 0/1)#switchport backup interface gi 0/2 preemption mode bandwidth SwitchD(config-if-GigabitEthernet 0/1)#switchport backup interface gi 0/2 preemption delay 40 SwitchD(config-if-GigabitEthernet 0/1)# exit</pre>
E	<pre>SwitchE> enable SwitchE# configure terminal SwitchD(config)# interface GigabitEthernet 0/1 SwitchD(config-if-GigabitEthernet 0/1)#switchport backup interface gi 0/2 preemption mode bandwidth SwitchD(config-if-GigabitEthernet 0/1)#switchport backup interface gi 0/2 preemption delay 40</pre>



	SwitchD(config-if-GigabitEthernet 0/1)# exit
Проверка	Проверьте информацию о двухканальном резервировании, настроенную для коммутатора D (E)
D	<pre>SwitchD#show interfaces switchport backup detail Switch Backup Interface Pairs: Active Interface Backup Interface State ----- Gi0/1 Gi0/2 Active Up/Backup Standby Interface Pair : Gi0/1, Gi0/2 Preemption Mode : bandwidth Preemption Delay : 40 seconds Bandwidth : Gi0/1(100000 Mbits), Gi0/2(100000 Mbits)</pre>
E	<pre>SwitchE#show interfaces switchport backup detail Switch Backup Interface Pairs: Active Interface Backup Interface State ----- Gi0/1 Gi0/2 Active Up/Backup Standby Interface Pair : Gi0/1, Gi0/2 Preemption Mode : bandwidth Preemption Delay : 40 seconds Bandwidth : Gi0/1(100000 Mbits), Gi0/2(100000 Mbits)</pre>

1.5.3.7. Распространенные ошибки

Сконфигурированный интерфейс не является физическим интерфейсом уровня 2 или интерфейсом AP.

1.5.4. Настройка обновления MAC-адреса

1.5.4.1. Эффект конфигурации

Быстро удаляются и обновляются MAC-адреса интерфейса во время переключения каналов, чтобы ускорить конвергенцию пакетов.

1.5.4.2. Примечания

- Необходимо настроить двухканальное резервирование REUP.
- Для каждого устройства можно настроить до 8 групп обновления адресов. Каждая группа обновления адресов может иметь максимум 8 интерфейсов-участников, и интерфейс может принадлежать нескольким группам обновления адресов.



1.5.4.3. Шаги настройки

- Обязательный.
- Если нет особых требований, следует настроить функцию обновления MAC-адреса.

1.5.4.4. Проверка

Запустите команду **show mac-address-table update group [detail]**, чтобы просмотреть конфигурацию группы обновлений.

1.5.4.5. Связанные команды

Настройка идентификатора группы обновления MAC-адреса коммутатора

Команда	mac-address-table update group [group-num]
Описание параметров	<i>group-num</i> : указывает идентификатор группы обновления MAC-адреса
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Чтобы уменьшить большой флуд, вызванный обновлением MAC-адреса, который может повлиять на нормальную передачу данных коммутатора, мы добавляем настройку группы обновления MAC-адреса. Только после того, как все интерфейсы на пути коммутации будут добавлены в одну и ту же группу обновления MAC-адресов, downlink-передача данных может быть быстро восстановлена

Включение отправки сообщений об обновлении MAC-адреса

Команда	mac-address-table move update transit
Командный режим	Командный режим
Руководство по использованию	Чтобы уменьшить переключение каналов и потерю потоков данных downlink-канала, необходимо включить отставку сообщений об обновлении MAC-адреса на коммутаторе, выполняющем переключение

Включение отправки идентификатора VLAN в сообщениях об обновлении MAC-адреса

Команда	mac-address-table move update transit vlan vid
Описание параметров	<i>vid</i> : указывает идентификатор VLAN для отправки сообщений об обновлении MAC-адреса



Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После того, как отправка сообщений об обновлении MAC-адреса включена, сообщения об обновлении MAC-адреса могут быть отправлены на uplink-устройства во время переключения канала

Настройка максимального количества пакетов обновления MAC-адреса, отправляемых в секунду

Команда	<code>mac-address-table move update max-update-rate pkts-per-second</code>
Описание параметров	<i>pkts-per-second</i> : указывает максимальное количество пакетов обновления MAC-адреса, отправляемых в секунду. Значение находится в диапазоне от 0 до 32 000. Значение по умолчанию — 150
Командный режим	Режим конфигурации
Руководство по использованию	Во время переключения канала REUP отправляет пакеты обновления MAC-адреса с заданным количеством на uplink-устройства каждую секунду, чтобы восстановить downlink-передачу данных uplink-устройства

Включение получения сообщений об обновлении MAC-адреса

Команда	<code>mac-address-table move update receive</code>
Командный режим	Режим конфигурации
Руководство по использованию	Во время двухканального резервирования downlink-поток данных могут быть потеряны, поскольку таблица MAC-адресов uplink-коммутатора не обновляется в режиме реального времени. Чтобы уменьшить потери потоков данных уровня 2, вам необходимо обновить таблицу MAC-адресов uplink-коммутатора. В этом случае вам необходимо включить получение сообщений об обновлении MAC-адреса на uplink-коммутаторе

Настройка диапазона VLAN для обработки сообщений об обновлении MAC-адреса

Команда	<code>mac-address-table move update receive vlan vlan-range</code>
Описание параметров	<i>vlan-range</i> : указывает диапазон VLAN для обработки сообщений обновления MAC-адреса



Командный режим	Режим конфигурации
Руководство по использованию	Эта команда используется для отключения функции обработки сообщений об обновлении MAC-адреса в определенных VLAN. Для VLAN, отключенной с функцией обработки сообщений обновления MAC-адреса, пакеты обновления MAC-адреса могут использоваться для восстановления downlink-передачи uplink-устройств; однако производительность сходимости для сбоев канала будет снижена

1.5.4.6. Пример конфигурации

Настройка обновления MAC-адреса

Сценарий	Как показано на Рисунке 1-6, от коммутатора D к коммутатору A ведут два upstream-канала: коммутатор D > коммутатор B > коммутатор A и коммутатор D > коммутатор C > коммутатор A. Есть два upstream-канала от коммутатора E к коммутатору A: коммутатор E > коммутатор B > коммутатор A и коммутатор E > коммутатор C > коммутатор A
Шаги настройки	- Включите отправку сообщений об обновлении MAC-адреса на коммутаторе доступа D (E). - Включите получение пакетов обновления MAC-адреса на коммутаторе B (C). - Добавьте все интерфейсы на пути коммутации REUP в одну и ту же группу обновления MAC-адресов. - В среде Gi0/1 и Gi0/3 коммутатора B — это интерфейсы на пути коммутации uplink коммутатора D, а Gi0/3 и Gi0/2 — интерфейсы на пути коммутации uplink коммутатора E. Вы можете добавить интерфейсы Gi0/1, Gi0/2 и Gi0/3 в одну и ту же группу обновления адреса. Аналогичным образом можно получить конфигурацию коммутатора C. - Включите получение пакетов обновления MAC-адреса на коммутаторе A. - Добавьте все интерфейсы на пути коммутации REUP коммутатора A в одну и ту же группу обновления MAC-адресов
D	<pre>SwitchD> enable SwitchD# configure terminal SwitchD(config)# mac-address-table move update transit SwitchD(config)# exit</pre>
E	<pre>SwitchE> enable SwitchE# configure terminal SwitchE((config)# mac-address-table move update transit SwitchE(config)# exit</pre>



B	<pre>SwitchB# configure terminal SwitchB(config)# mac-address-table move update receive SwitchB(config)# interface range gigabitEthernet 0/1 -3 SwitchB(config-if-range)#switchport mode trunk SwitchB(config-if-range)# mac-address-table update group 1 SwitchB(config-if-range)# end</pre>
C	<pre>SwitchB# configure terminal SwitchB(config)# mac-address-table move update receive SwitchB(config)# interface range gigabitEthernet 0/1 -3 SwitchB(config-if-range)#switchport mode trunk SwitchB(config-if-range)# mac-address-table update group 1 SwitchB(config-if-range)# end</pre>
A	<pre>SwitchA# configure terminal SwitchA(config)# mac-address-table move update receive SwitchA(config)# interface range gigabitEthernet 0/1 -2 SwitchA(config-if-range)# switchport mode trunk SwitchA(config-if-range)# mac-address-table update group 1 SwitchA(config-if-range)# end</pre>
Проверка	Проверьте информацию о группах обновления адресов на коммутаторах D, E, C, B и A
D	<pre>SwitchD# show run incl mac-ad mac-address-table move update transit</pre>
E	<pre>SwitchE# show run incl mac-ad mac-address-table move update transit</pre>
B	<pre>SwitchB# show mac-address-table update group detail show mac-address-table update group detailMac-address-table Update Group:1 Received mac-address-table update message count:0 Group member Receive Count Last Receive Switch-ID Receive Time ----- Gi0/1 0 0000.0000.0000 Gi0/2 0 0000.0000.0000 Gi0/3 0 0000.0000.0000</pre>



C	<pre>SwitchC# show mac-address-table update group detail Mac-address-table Update Group:1 Received mac-address-table update message count:0 Group member Receive Count Last Receive Switch-ID Receive Time ----- Gi0/1 0 0000.0000.0000 Gi0/2 0 0000.0000.0000 Gi0/3 0 0000.0000.0000</pre>
A	<pre>SwitchA# show mac-address-table update group detail Mac-address-table Update Group:1 Received mac-address-table update message count:0 Group member Receive Count Last Receive Switch-ID Receive Time ----- Gi0/1 0 0000.0000.0000 Gi0/2 0 0000.0000.0000</pre>

1.5.4.7. Распространенные ошибки

Сконфигурированный интерфейс не является физическим интерфейсом уровня 2 или интерфейсом AP.

1.5.5. Настройка балансировки нагрузки VLAN

1.5.5.1. Эффект конфигурации

Максимально используйте пропускную способность канала.

1.5.5.2. Примечания

- Необходимо настроить двухканальное резервирование REUP.
- Интерфейс доступа не может совместно использоваться балансировкой нагрузки VLAN и STP.
- Для интерфейсов, успешно настроенных с балансировкой нагрузки VLAN, вы не можете изменять атрибуты интерфейсов, но можете изменять атрибуты VLAN интерфейсов.

1.5.5.3. Шаги настройки

- Если максимальное использование полосы пропускания не требуется, эта конфигурация не является обязательной.
- Если требуется балансировка нагрузки VLAN, необходимо выполнить соответствующую настройку.

1.5.5.4. Проверка

Запустите команду **show interfaces switchport backup [detail]**, чтобы проверить, настроен ли баланс нагрузки VLAN.



1.5.5.5. Связанные команды

Настройка балансировки нагрузки VLAN

Команда	switchport backup interface <i>interface-id</i> prefer instance <i>instance-range</i>
Описание параметров	<i>interface-id</i> : указывает идентификатор резервного интерфейса. <i>instance-range</i> : указывает диапазон загрузки экземпляра интерфейса резервного копирования
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Вы можете изменить сопоставление между экземплярами и VLAN, используя функцию сопоставления экземпляров MSTP

1.5.5.6. Пример конфигурации

Настройка балансировки нагрузки VLAN

Сценарий	Как показано на Рисунке 1-6, от коммутатора D к коммутатору A ведут два upstream-канала: коммутатор D > коммутатор B > коммутатор A и коммутатор D > коммутатор C > коммутатор A. Есть два upstream-канала от коммутатора E к коммутатору A: коммутатор E > коммутатор B > коммутатор A и коммутатор E > коммутатор C > коммутатор A
Шаги настройки	Настройте сопоставления экземпляров на коммутаторе D (E), чтобы сопоставить VLAN 1 с экземпляром 1, VLAN 2 с экземпляром 2, VLAN 3 с экземпляром 3 и VLAN 4 с экземпляром 4. Дополнительные сведения см. в Ethernet Switching/Руководстве по настройке MSTP. Настройте функцию балансировки нагрузки VLAN на коммутаторе D (E)
D	<pre>SwitchD> enable SwitchD# configure terminal SwitchD(config)# interface GigabitEthernet 0/1 SwitchD(config-if-GigabitEthernet 0/1)# switchport mode trunk SwitchD(config-if-GigabitEthernet 0/1)#switchport backup interface gi0/2 prefer instance 2 SwitchD(config-if-GigabitEthernet 0/1)# exit</pre>
E	<pre>SwitchE> enable SwitchE# configure terminal SwitchE(config)# interface GigabitEthernet 0/1 SwitchE(config-if-GigabitEthernet 0/1)# switchport mode trunk</pre>



	<pre>SwitchD(config-if-GigabitEthernet 0/1)#switchport backup interface gi0/2 prefer instance 4 SwitchD(config-if-GigabitEthernet 0/1)# exit</pre>
Проверка	Проверьте информацию о двухканальном резервировании, настроенную для коммутатора D (E)
D	<pre>SwitchD#show interfaces switchport backup detail Switch Backup Interface Pairs: Active Interface Backup Interface State ----- Gi0/1 Gi0/2 Active Up/Backup Up Instances Preferred on Active Interface: Instance 0-1,3-64 Mapping VLAN 1,3-4094 Instances Preferred on Backup Interface: Instance 2 Mapping VLAN 2 Interface Pair : Gi0/1, Gi0/2 Preemption Mode : balance Preemption Delay : 35 seconds Bandwidth : Gi0/1(800 kbits), Gi0/2(100000 kbits)</pre>
E	<pre>SwitchE#show interfaces switchport backup detail Switch Backup Interface Pairs: Active Interface Backup Interface State ----- Gi0/1 Gi0/2 Active Up/Backup Up Instances Preferred on Active Interface: Instance 0-3,5-64 Mapping VLAN 1-3,5-4094 Instances Preferred on Backup Interface: Instance 4 Mapping VLAN 4 Interface Pair : Gi0/1, Gi0/2 Preemption Mode : balance Preemption Delay : 35 seconds Bandwidth : Gi0/1(800 kbits), Gi0/2(100000 kbits)</pre>

1.5.5.7. Распространенные ошибки

Сопоставления между идентификаторами VLAN и экземплярами не настроены.



1.5.6. Настройка отслеживания каналов

1.5.6.1. Эффект конфигурации

Обнаружив, что upstream-канал отключен, принудительно отключите downstream-канал, чтобы можно было выполнить переключение каналов.

1.5.6.2. Примечания

- Необходимо настроить двухканальное резервирование REUP.
- Для функции отслеживания состояния связи каждый интерфейс принадлежит только к одной группе отслеживания состояния связи, и для каждого устройства можно настроить до 2 групп отслеживания состояния связи. Каждая группа отслеживания состояния канала может иметь 8 upstream-интерфейсов и 256 downstream-интерфейсов.

1.5.6.3. Шаги настройки

- Обязательный.
- Если нет особых требований, следует настроить функцию отслеживания upstream-канала.

1.5.6.4. Проверка

Запустите команду **show link state group**, чтобы просмотреть настроенную информацию об отслеживании каналов.

1.5.6.5. Связанные команды

Включение группы отслеживания состояния каналов

Команда	link state track [num]
Описание параметров	<i>num</i> : указывает идентификатор группы отслеживания состояния связи
Командный режим	Режим конфигурации
Руководство по использованию	Вы можете создать группу отслеживания каналов, а затем добавить интерфейс в указанную группу отслеживания

Включение задержки downlink-канала для группы отслеживания состояния канала

Команда	link state track num up-delay timer
Описание параметров	<i>num</i>: указывает идентификатор группы отслеживания состояния связи. <i>timer</i>: указывает время задержки downlink-канала, которое по умолчанию равно 0 секунд



Командный режим	Режим глобальной конфигурации
Руководство по использованию	Вы должны включить функцию задержки, чтобы downstream-канал мог работать после задержки

Добавление интерфейса в Link Tracking Group

Команда	link statigroup <i>num</i> {upstream downstream}
Описание параметров	<i>num</i>: указывает идентификатор состояния Link Tracking Group. upstream: добавляет интерфейс в качестве upstream-интерфейса группы отслеживания. downstream: добавляет интерфейс в качестве downstream-интерфейса группы отслеживания
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Вы можете создать Link Tracking Group, а затем добавить интерфейс в указанную группу отслеживания

1.5.6.6. Пример конфигурации

Настройка Link Tracking Group

Сценарий	Как показано на Рисунке 1-6, от коммутатора D к коммутатору A ведут два upstream-канала: коммутатор D > коммутатор B > коммутатор A и коммутатор D > коммутатор C > коммутатор A. Есть два upstream-канала от коммутатора E к коммутатору A, то есть коммутатор E > коммутатор B > коммутатор A и коммутатор E > коммутатор C > коммутатор A
Шаги настройки	- Создайте Link Tracking Group 1 на коммутаторе B (C). - На коммутаторе B (C) добавьте интерфейсы Gi0/1 и Gi0/2 в качестве downlink-интерфейсов Link Tracking Group и добавьте интерфейс Gi0/3 в качестве upstream-интерфейса Link Tracking Group
В	<pre>SwitchB> enable SwitchB# configure terminal SwitchB(config)# link state track 1 SwitchB(config)# interface GigabitEthernet 0/1 SwitchB(config-if-GigabitEthernet 0/1)#link state group 1 downstream SwitchB(config-if-GigabitEthernet 0/1)#exit SwitchB(config)# interface GigabitEthernet 0/2</pre>



	<pre>SwitchB(config-if-GigabitEthernet 0/2)# link state group 1 downstream SwitchB(config-if-GigabitEthernet 0/2)#exit SwitchB(config)# interface GigabitEthernet 0/3 SwitchB(config-if-GigabitEthernet 0/3)#link state group 1 upstream SwitchB(config-if-GigabitEthernet 0/3)#exit</pre>
C	<pre>SwitchC> enable SwitchC# configure terminal SwitchC(config)# link state track 1 SwitchC(config)# interface GigabitEthernet 0/1 SwitchC(config-if-GigabitEthernet 0/1)#link state group 1 downstream SwitchC(config-if-GigabitEthernet 0/1)#exit SwitchC(config)# interface GigabitEthernet 0/2 SwitchC(config-if-GigabitEthernet 0/2)# link state group 1 downstream SwitchC(config-if-GigabitEthernet 0/2)#exit SwitchC(config)# interface GigabitEthernet 0/3 SwitchC(config-if-GigabitEthernet 0/3)#link state group 1 upstream SwitchC(config-if-GigabitEthernet 0/3)#exit</pre>
Проверка	Проверьте информацию о Link Tracking Group, настроенную для коммутатора В (С)
B	<pre>SwitchB#show link state group Link State Group:1 Status: enabled, Down Upstream Interfaces :Gi0/3(Down) Downstream Interfaces : Gi0/2(Down)</pre>

1.5.6.7. Распространенные ошибки

Интерфейсы добавляются в Link Tracking Group, когда Link Tracking Group не включена.

1.6. Мониторинг

1.6.1. Отображение

Описание	Команда
Отображает информацию о двухканальном резервировании REUP	show interfaces[<i>interface-id</i>]switchport backup [detail]



Описание	Команда
Отображает конфигурации группы обновления MAC-адресов	show mac-address-table update group [detail]
Отображает статистику REUP об отправленных сообщениях об обновлении MAC-адреса	show mac-address-table move update
Отображает информацию о состоянии Link Tracking Group	show link state group

1.6.1.1. Отладка

ПРИМЕЧАНИЕ: системные ресурсы заняты при выводе отладочной информации. Поэтому отключайте отладку сразу после использования.

Описание	Команда
Включает всю отладку REUP	debug reup all
Отладка нормального рабочего процесса REUP	debug reup process
Отладка сообщения об обновлении MAC-адреса REUP	debug reup packet
Отладка пакетов обновления MAC-адреса REUP	debug reup macupdt
Отладка оперативного резервирования	debug reup ha
Отладка ошибки, возникающей при работе REUP	debug reup error
Отладки получили события	debug reup evnet
Отладка статистики при выполнении операций show	debug reup status



2. НАСТРОЙКА RLDP

2.1. Обзор

Протокол Rapid Link Detection Protocol (RLDP) обеспечивает быстрое обнаружение отказов однонаправленных каналов, отказов направленной переадресации и отказов loop downlink-канала Ethernet. При обнаружении сбоя соответствующие порты будут автоматически закрыты в соответствии с конфигурацией обработки сбоя, или пользователь будет уведомлен о необходимости вручную закрыть порты, чтобы избежать неправильной переадресации потока или петли Ethernet уровня 2.

2.2. Приложения

Приложение	Описание
Обнаружение однонаправленной связи	Обнаружение отказа однонаправленного канала
Обнаружение двунаправленной пересылки	Обнаружение сбоя двунаправленного канала
Обнаружение петель downlink-канала	Обнаружение петли канала

2.2.1. Обнаружение однонаправленной связи

2.2.1.1. Сценарий

Как показано на следующем рисунке, А подключен к В через оптоволокно. Эти две линии являются линиями Tx и Rx оптического волокна. Обнаружение однонаправленного соединения включено на А и В. Если какой-либо из Tx порта А, Rx порта В, Tx порта В и Rx порта А отказывает, однонаправленный сбой будет обнаружен и обработан в соответствии с RLDP. Если сбой устранен, администратор может вручную восстановить RLDP на А и В и возобновить обнаружение.



Рисунок 2-1.

А и В — коммутаторы уровня 2 или уровня 3.

Tx порта А устройства А подключен к Rx порта В устройства В.

Rx порта А устройства А подключен к Tx порта В устройства В.

2.2.1.2. Развертывание

- Глобальный RLDP включен.

- Настройте обнаружение однонаправленного соединения для портов А и В и определите метод обработки сбоев.

2.2.2. Обнаружение двунаправленной пересылки

2.2.2.1. Сценарий

Как показано на следующем рисунке, А подключен к В через оптическое волокно, а две линии являются линиями Tx и Rx оптоволоконного кабеля. Обнаружение однонаправленного соединения включено на А и В. Если Tx порта А, Rx порта В, Rx порта А и Tx порта В не работают, двунаправленный сбой будет обнаружен и обработан в соответствии с RLDP. Если сбой устранен, администратор может вручную восстановить RLDP на А и В и возобновить обнаружение.



Рисунок 2-2.

А и В — коммутаторы уровня 2 или уровня 3.

Tx порта А устройства А подключен к Rx порта В устройства В.

Rx порта А устройства А подключен к Tx порта В устройства В.

2.2.2.2. Развертывание

- Глобальный RLDP включен.
- Настройте BFD для портов А и В и определите метод обработки сбоев.

2.2.3. Обнаружение петель downlink-канала

2.2.3.1. Сценарий

Как показано на следующем рисунке, А, В и С соединены в петлю. Обнаружение петель downlink-канала включено на А, и петля обнаруживается и обрабатывается.

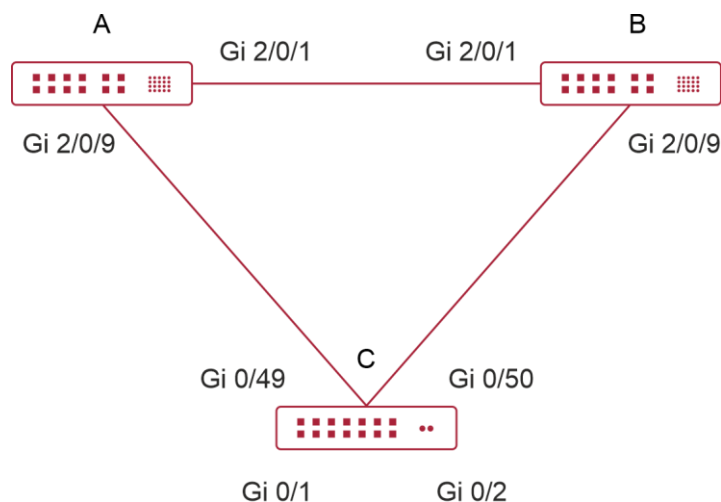


Рисунок 2-3.



А, В и С — коммутаторы уровня 2 или уровня 3.

2.2.3.2. Развертывание

- Глобальный RLDП включен на А.
- Настройте обнаружение петель downlink-канала на портах Gi 2/0/1 и Gi 2/0/9 А и определите метод устранения сбоев.

2.3. Функции

Большинство механизмов обнаружения каналов Ethernet обнаруживают возможность подключения через автоматическое согласование на физическом уровне. Однако в некоторых случаях устройства подключены на физическом уровне и работают нормально, но связь по каналу уровня 2 отключена или ненормальна. RLDП распознает соседнее устройство и обнаруживает отказ канала путем обмена с устройством пакетами Prob, Echo или Loop.

2.3.1.1. Базовые концепты

Сбой однонаправленной связи

Отказ однонаправленного соединения происходит в случае перекрестного оптического волокна, отсоединенного оптического волокна, обрыва цепи оптического волокна, одной разомкнутой линии в витой паре или однонаправленного обрыва цепи промежуточного устройства между двумя устройствами. В таких случаях один конец канала подключен, а другой отключен, поэтому поток перенаправляется неправильно или протокол защиты от петель (например, STP) дает сбой.

Сбой двунаправленной связи

Отказ двунаправленного соединения происходит в случае двух оптических волокон, двух разомкнутых линий в витой паре или двунаправленного обрыва промежуточного устройства между двумя устройствами. В таких случаях оба конца канала отключаются, поэтому поток перенаправляется неправильно.

Сбой петли

Downlink-устройство неправильно подключено, образуя петлю, что приводит к широковещательному шторму (broadcast storm).

Пакет RLDП

RLDП определяет три типа пакетов: пакеты Prob, пакеты Echo и пакеты Loop.

- Пакеты Prob представляют собой многоадресные пакеты уровня 2 для согласования соседей и обнаружения однонаправленных или двунаправленных каналов. Формат инкапсуляции по умолчанию — SNAP, который автоматически меняется на EthernetII, если сосед отправляет пакеты EthernetII.
- Пакеты Echo представляют собой одноадресные пакеты уровня 2 в ответ на пакеты Prob и используются для обнаружения однонаправленных или двунаправленных каналов. Формат инкапсуляции по умолчанию — SNAP, который автоматически меняется на EthernetII, если сосед отправляет пакеты EthernetII.
- Пакеты Loop представляют собой многоадресные пакеты уровня 2 для обнаружения петли downlink-канала. Их можно только получить. Формат инкапсуляции по умолчанию — SNAP.

Интервал обнаружения RLDП и максимальное время обнаружения

Интервал обнаружения и максимальное время обнаружения можно настроить для RLDП. Интервал обнаружения определяет период отправки пакетов Prob и Loop. Когда устройство получает пакет Prob, оно немедленно отвечает пакетом Echo. Интервал



обнаружения и максимальное время обнаружения определяют максимальное время обнаружения (равное интервалу обнаружения × максимальное время обнаружения + 1) для обнаружения однонаправленного или двунаправленного соединения. Если ни пакет Prob, ни пакет Echo от соседа не могут быть получены в течение максимального времени обнаружения, будет инициирована обработка однонаправленного или двунаправленного сбоя.

Согласование соседей RLDLP

При настройке обнаружения однонаправленного или двунаправленного соединения порт может узнавать устройство peer-end как своего соседа. Один порт может узнать одного соседа, который является переменным. Если согласование включено, обнаружение однонаправленного или двунаправленного соединения начинается после того, как порт находит соседний узел посредством согласования, которое завершается успешно, когда порт получает пакет Prob от соседнего узла. Однако, если RLDLP включен при сбое, порт не может узнать соседний узел, поэтому обнаружение не может начаться. В этом случае восстановите состояние канала перед включением RLDLP.

Обработка неисправного порта в соответствии с RLDLP

- Предупреждение: печатайте системный журнал только для указания отказавшего порта и типа отказа.
- Завершение работы SVI: распечатайте системный журнал, а затем запросите SVI в соответствии с Access VLAN или Native VLAN порта и выключите SVI, если порт является портом физического обмена или портом-членом AP уровня 2.
- Нарушение порта: распечатайте системный журнал и настройте неисправный порт как находящийся в состоянии нарушения, после чего порт физически перейдет в состояние Linkdown.
- Блокировка: распечатайте системный журнал и настройте состояние пересылки порта как Block, и порт не будет пересылать пакеты.

Восстановление неисправного порта по RLDLP

- Ручной сброс: вручную сбросить все неисправные порты до инициализированного состояния и перезапустить обнаружение соединения.
- Ручное или автоматическое аварийное восстановление: восстанавливает все неисправные порты до инициализированного состояния вручную или с периодичностью (30 секунд по умолчанию и настраивается) и перезапускает обнаружение каналов.
- Автоматическое восстановление: при обнаружении однонаправленного или двунаправленного канала, если обработка отказавших портов не указана как нарушение порта, восстанавливайте порты до инициализированного состояния на основе пакетов Prob и перезапускайте обнаружение канала.

Состояние порта в соответствии с RLDLP

- Нормальный: указывает состояние порта после включения обнаружения соединения.
- Ошибка: указывает состояние порта после обнаружения однонаправленного или двунаправленного сбоя соединения или сбоя петли.



2.3.2. Обзор

Особенность	Описание
<u>Развертывание обнаружения RLDLP</u>	Включает обнаружение однонаправленного или двунаправленного канала, или обнаружение петли downlink-канала на наличие сбоев и реализует обработку

2.3.3. Развертывание обнаружения RLDLP

RLDP обеспечивает обнаружение однонаправленного канала, обнаружение двунаправленной переадресации и обнаружение петли downlink-канала.

2.3.3.1. Принцип работы

Обнаружение однонаправленной связи

Когда эта функция включена, порт регулярно отправляет пакеты Prob и получает пакеты Echo от соседнего узла, а также получает пакеты Prob от соседнего узла и отвечает пакетами Echo. В течение максимального времени обнаружения, если порт получает пакеты Prob, но не пакеты Echo или ни один из них, будет инициирована обработка однонаправленного сбоя, и обнаружение прекратится.

Обнаружение двунаправленной пересылки

Когда эта функция включена, порт регулярно отправляет пакеты Prob и получает пакеты Echo от соседнего узла, а также получает пакеты Prob от соседнего узла и отвечает пакетами Echo. В течение максимального времени обнаружения, если порт не получает ни пакетов Prob, ни пакетов Echo от соседнего узла, будет инициирована обработка двунаправленного сбоя, и обнаружение прекратится.

Обнаружение петель downlink-канала

Когда эта функция включена, порт регулярно отправляет пакеты Loop. В следующих случаях отказ петли будет инициирован после того, как тот же или другой порт получит пакеты: в одном случае выходной и входной порты являются одним и тем же маршрутизируемым портом или портом-участниками AP уровня 3; в другом случае выходные и входные порты являются портами обмена или портами-участниками AP уровня 2 в одной и той же VLAN по умолчанию и в состоянии Forward. Будет реализовано исправление сбоя, и обнаружение прекратится.

2.3.3.2. Связанная настройка

Настройка обнаружения RLDLP

По умолчанию обнаружение RLDLP отключено.

Вы можете запустить глобальную команду **rldp enable** или интерфейсную команду **rldp port**, чтобы включить обнаружение RLDLP и указать тип обнаружения и обработку.

Вы можете запустить команду **rldp neighbor-negotiation** для согласования соседей, **rldp-detect-interval** для указания интервала обнаружения, **rldp-detect-max** для указания времени обнаружения или **rldp reset** для восстановления неисправного порта.



2.4. Конфигурация

Конфигурация	Описание и команда	
Настройка основных функций RLDP	(Обязательно) Он используется для включения обнаружения RLDP в режиме глобальной конфигурации	
	rldp enable	Включает глобальное обнаружение RLDP на всех портах
Настройка основных функций RLDP	(Обязательно) Он используется для указания в режиме конфигурации интерфейса типа обнаружения и обработки сбоев для интерфейса	
	rldp port	Включает обнаружение RLDP на порту и указывает тип обнаружения и обработку отказа
	(Опционально) Он используется для настройки интервала обнаружения, времени обнаружения и согласования соседей в режиме глобальной конфигурации	
	rldp detect-interval	Изменяет глобальные параметры RLDP на всех портах, такие как интервал обнаружения, максимальное время обнаружения и согласование соседей
	rldp detect-max	
	rldp neighbor-negotiation	
	(Опционально) Используется в привилегированном режиме	
	rldp reset	Восстанавливает все порты

2.4.1. Настройка основных функций RLDP

2.4.1.1. Эффект конфигурации

Включите обнаружение однонаправленного канала RLDP, обнаружение двунаправленной пересылки или обнаружение петли downlink-канала для обнаружения сбоев.

2.4.1.2. Примечания

- Обнаружение петель эффективно для всех портов-участников AP, если оно настроено на одном из портов. Обнаружение однонаправленного соединения и обнаружение двунаправленной пересылки эффективны только на порте-участнике AP.
- Обнаружение петли на физическом порту, добавленном к AP, должно быть настроено так же, как и на других портах-участниках. Есть три случая. Во-первых,



если обнаружение петель настроено не для вновь добавленного порта, а для существующих портов-участников, новый порт принимает конфигурацию и результаты обнаружения существующих портов. Во-вторых, если только что добавленный порт и существующие порты-участники имеют разную конфигурацию обнаружения петель, новый порт принимает конфигурацию и результаты обнаружения существующих портов.

- При настройке RLDП на порту AP вы можете настроить обработку сбоев только как «shutdown-port» («порт выключения»), на котором будут изменены другие настройки.
- Когда для порта настроен «shutdown-port», обнаружение RLDП не может быть восстановлено в случае сбоя. После устранения неполадок вы можете запустить команду **rldp reset** или **errdisable recovery**, чтобы восстановить порт и возобновить обнаружение.

2.4.1.3. Шаги настройки

Включение RLDП

- Обязательный.
- Включите обнаружение RLDП на всех портах в режиме глобальной конфигурации.

Включение согласования соседей

- Опционально.
- Включите функцию в режиме глобальной конфигурации, и обнаружение портов будет запущено при успешном согласовании соседей.

Настройка интервала обнаружения

- Опционально.
- Укажите интервал обнаружения в режиме глобальной конфигурации.

Настройка максимального времени обнаружения

- Опционально.
- Укажите максимальное время обнаружения в режиме глобальной конфигурации.

Настройка обнаружения в порту

- Обязательный.
- Настройте однонаправленное обнаружение RLDП, двунаправленное обнаружение RLDП или обнаружение петли нисходящей линии связи в режиме конфигурации интерфейса и укажите обработку сбоя.

Восстановление всех неисправных портов

- Опционально.
- Включите эту функцию в привилегированном режиме, чтобы восстановить все неисправные порты и возобновить обнаружение.

2.4.1.4. Проверка

Отображение информации о глобальном RLDП, порте и соседе.



2.4.1.5. Связанные команды

Включение глобального обнаружения RLDLP

Команда	rldp enable
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Включить глобальное обнаружение RLDLP

Включение обнаружения RLDLP на интерфейсе

Команда	rldp port { unidirection-detect bidirection-detect loop-detect } { warning shutdown-svi shutdown-port block }
Описание параметров	unidirection-detect: указывает на обнаружение однонаправленной связи. bidirection-detect: указывает на обнаружение двунаправленной переадресации. loop-detect: указывает на обнаружение петли downlink-канала. warning: указывает, что исправление отказа является предупреждением. shutdown-svi: указывает, что исправление сбоя закрывает SVI, на котором включен интерфейс. shutdown-port: указывает, что исправление сбоя является нарушением порта. block: указывает, что исправление сбоя отключает изучение и переадресацию порта
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Интерфейсы включают в себя порты коммутатора уровня 2, маршрутизируемые порты уровня 3, порты-участники AP уровня 2 и порты-участники AP уровня 3

Изменение глобальных параметров обнаружения RLDLP

Команда	rldp {detect-interval <i>interval</i> detect-max <i>num</i> neighbor-negotiation }
Описание параметров	detect-interval <i>interval</i>: указывает интервал обнаружения. detect-max <i>num</i>: указывает время обнаружения. neighbor-negotiation: указывает на согласование соседей



Командный режим	Режим глобальной конфигурации
Руководство по использованию	При необходимости измените все параметры RLDP на всех портах

Восстановление неисправного порта

Команда	rldp reset
Командный режим	Привилегированный режим
Руководство по использованию	Восстановит все неисправные порты в инициализированное состояние и возобновит обнаружение

Отображение информации о состоянии RLDP

Команда	show rldp [interface <i>interface-name</i>]
Описание параметров	<i>interface-name</i> : указывает интерфейс для отображения информации
Командный режим	Привилегированный режим, режим глобальной конфигурации или режим конфигурации интерфейса
Руководство по использованию	Отображение информации о состоянии RLDP

2.4.1.6. Пример конфигурации

Включение обнаружения RLDP в кольцевой топологии

Сценарий:

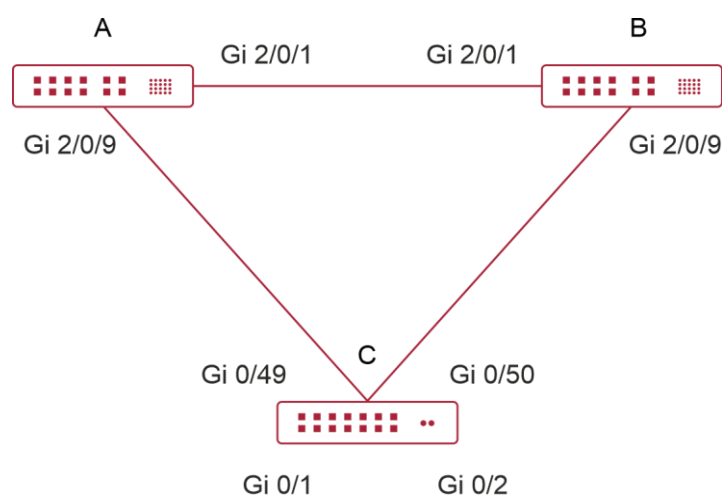


Рисунок 2-4.



	Как показано на Рисунке 2-4, секции агрегации и доступа находятся в кольцевой топологии. STP включен на всех устройствах для предотвращения образования петель и обеспечения защиты от избыточности. Чтобы избежать однонаправленного или двунаправленного сбоя соединения, приводящего к сбою STP, включено обнаружение однонаправленного и двунаправленного соединения RLDП между устройствами агрегации, а также между устройством агрегирования и устройством доступа. Чтобы избежать заикливания из-за неправильного downlink-соединения устройств агрегации, включите обнаружение петли downlink-канала RLDП на портах downlink-канала устройств агрегации и устройства доступа. Чтобы избежать образования петель из-за неправильного downlink-подключения устройства доступа, включите обнаружение петли downlink канала RLDП на downlink-портах устройства доступа
Шаги настройки	• SW А и SW В являются устройствами агрегации, а SW С является устройством доступа. Пользователи, подключены к SW С. SW А, SW В и SW С имеют кольцевую топологию, и на каждом из них включен STP. Для настройки STP обратитесь к соответствующему руководству по настройке. • Включите RLDП на SW А, включите обнаружение однонаправленных и двунаправленных каналов на двух портах и включите обнаружение петель на downlink-порту. • Включите RLDП на SW В, включите обнаружение однонаправленных и двунаправленных каналов на двух портах и включите обнаружение петель на downlink-порту. • Включите RLDП на SW С, включите обнаружение однонаправленного и двунаправленного соединения на двух uplink-портах и включите обнаружение петель на двух downlink-портах
А	<pre> A#configure terminal A(config)#rldp enable A(config)#interface GigabitEthernet 2/0/1 A(config-if-GigabitEthernet 2/0/1)#rldp port unidirection-detect shutdown-port A(config-if-GigabitEthernet 2/0/1)#rldp port bidirection-detect shutdown-port A(config-if-GigabitEthernet 2/0/1)# exit A(config)#interface GigabitEthernet 2/0/9 A(config-if-GigabitEthernet 2/0/1)#rldp port unidirection-detect shutdown-port A(config-if-GigabitEthernet 2/0/1)#rldp port bidirection-detect shutdown-port A(config-if-GigabitEthernet 2/0/1)#rldp port loop-detect shutdown-port A(config-if-GigabitEthernet 2/0/1)#exit </pre>



В	Примените конфигурацию на SW A
С	<pre> C#configure terminal C(config)#rldp enable C(config)#interface GigabitEthernet 0/49 C(config-if-GigabitEthernet 0/49)#rldp port unidirection-detect shutdown-port C(config-if-GigabitEthernet 0/49)#rldp port bidirection-detect shutdown-port C(config-if-GigabitEthernet 0/49)# exit C(config)#interface GigabitEthernet 0/50 C(config-if-GigabitEthernet 0/50)#rldp port unidirection-detect shutdown-port C(config-if-GigabitEthernet 0/50)#rldp port bidirection-detect shutdown-port C(config-if-GigabitEthernet 0/50)#exit C(config)#interface GigabitEthernet 0/1 C(config-if-GigabitEthernet 0/1)# rldp port loop-detect shutdown-port C(config-if-GigabitEthernet 0/1)#exit C(config)#interface GigabitEthernet 0/2 C(config-if-GigabitEthernet 0/2)# rldp port loop-detect shutdown-port C(config-if-GigabitEthernet 0/2)#exit </pre>
Проверка	Проверьте информацию RLDP на SW A, SW B и SW C. Возьмем, к примеру, SW A
А	<pre> A#show rldp rldp state : enable rldp hello interval: 3 rldp max hello : 2 rldp local bridge : 08c6.b322.33aa ----- Interface GigabitEthernet 2/0/1 port state : normal neighbor bridge : 08c6.b300.51b1 neighbor port : GigabitEthernet 2/0/1 unidirection detect information: action: shutdown-port state : normal bidirection detect information: </pre>



	<pre> action: shutdown-port state : normal Interface GigabitEthernet 2/0/9 port state : normal neighbor bridge : 08c6.b300.41b0 neighbor port : GigabitEthernet 0/49 unidirection detect information: action: shutdown-port state : normal bidirection detect information: action: shutdown-port state : normal loop detect information: action: shutdown-port state : normal </pre>
--	--

2.4.1.7. Распространенные ошибки

- Одновременно включаются функции RLDP и аутентификация приватного многоадресного адреса или TPP.
- Согласование соседей не включено при настройке обнаружения однонаправленного или двунаправленного канала. RLDP должен быть включен на соседнем устройстве, иначе будет обнаружен однонаправленный или двунаправленный сбой.
- Если обнаружение RLDP настроено на реализацию после согласования соседнего узла при настройке обнаружения однонаправленного или двунаправленного соединения, обнаружение невозможно реализовать, так как из-за сбоя соединения соседние узлы не могут быть изучены. В этой ситуации вам предлагается сначала восстановить состояние канала.
- Вам предлагается не указывать обработку сбоя как Shutdown SVI для маршрутизируемого порта.
- Вам предлагается не указывать обработку отказа как Block для порта, на котором включен протокол защиты от петель, например, STP.

2.5. Мониторинг

2.5.1. Отображение

Описание	Команда
Отображает состояние RLDP	show rldp [interface interface-name]



3. НАСТРОЙКА DLDP

3.1. Обзор

Data Link Detection Protocol (DLDP) — это протокол, используемый для быстрого обнаружения неисправных каналов Ethernet.

Типичный механизм обнаружения соединения Ethernet определяет возможность подключения физического соединения посредством автосогласования на физическом уровне. Такой механизм имеет ограничения при обнаружении исключений передачи данных уровня 3, несмотря на нормальные физические соединения.

DLDP предоставляет надежную информацию об обнаружении каналов уровня 3. После обнаружения неисправного канала DLDP отключает логическое состояние портов 3-го уровня, чтобы реализовать быструю конвергенцию протоколов 3-го уровня.

3.2. Приложения

Приложение	Описание
<u>Обнаружение внутрисетевого сегмента DLDP</u>	Исходный IP-адрес обнаруженного порта и обнаруженный IP-адрес находятся в одном сегменте сети
<u>Обнаружение межсетевого сегмента DLDP</u>	Исходный IP-адрес обнаруженного порта и обнаруженный IP-адрес находятся в разных сегментах сети

3.2.1. Обнаружение внутрисетевого сегмента DLDP

3.2.1.1. Сценарий

В этом разделе описывается базовый сценарий приложения DLDP, в котором исходный IP-адрес обнаруженного порта и обнаруженный IP-адрес находятся в одном сегменте сети.

На Рисунке 3-1 порт Gi 0/1 уровня 3 на устройстве A и порт Gi 0/2 уровня 3 на устройстве C находятся в одном сегменте сети. Чтобы обнаружить подключение канала уровня 3 от Gi 0/1 до Gi 0/2, включите обнаружение DLDP на Gi 0/1 или Gi 0/2.



Рисунок 3-1.

Gi 0/1 и Gi 0/2 — это порты уровня 3 в одном сегменте сети.

B — это сеть в том же сегменте сети, что и Gi 0/1 и Gi 0/2.

3.2.1.2. Развертывание

Включите DLDP на Gi 0/1 или Gi 0/2.



3.2.2. Обнаружение межсетевого сегмента DLDP

3.2.2.1. Сценарий

В этом разделе описывается сценарий приложения DLDP, в котором исходный IP-адрес обнаруженного порта и обнаруженный IP-адрес находятся в разных сегментах сети.

На Рисунке 3-2 порт Gi 0/1 уровня 3 на устройстве A и порт Gi 0/4 уровня 3 на устройстве D находятся в разных сегментах сети. Чтобы определить подключение канала уровня 3 от Gi 0/1 к Gi 0/4, включите DLDP на Gi 0/1 и настройте IP-адрес следующего перехода DLDP (IP-адрес порта Gi 0/2 на устройстве B).



Рисунок 3-2.

Gi 0/1 и Gi 0/4 — это порты уровня 3 в разных сегментах сети.

3.2.2.2. Развертывание

Включите DLDP на Gi 0/1 и настройте IP-адрес следующего перехода DLDP.

3.3. Функции

3.3.1. Базовые концепты

Интервал обнаружения DLDP и время повторной передачи

Интервал обнаружения: указывает интервал, с которым передаются пакеты обнаружения DLDP (ICMP echo).

Время повторной передачи: укажите максимальное количество раз, когда пакеты обнаружения DLDP могут быть повторно переданы в случае сбоя обнаружения DLDP.

Когда сетевое устройство не получает ответный пакет от peer end в течение интервала обнаружения, умноженного на время повторной передачи, устройство определяет, что произошел сбой канала 3-го уровня, и отключает логическое состояние своего порта 3-го уровня. (несмотря на нормальное физическое соединение). Когда соединение канала уровня 3 восстанавливается, устройство восстанавливает свой порт уровня 3 в логическое состояние Up.

Режимы обнаружения DLDP

Активный режим и пассивный режим — это два режима обнаружения DLDP.

Активный режим (по умолчанию): пакеты обнаружения ICMP отправляются активно.

Пассивный режим: пакеты обнаружения ICMP принимаются пассивно.

DLDP next hop

Next hop: указывает следующий узел, подключенный к обнаруженному IP-адресу при обнаружении межсетевого сегмента DLDP.

В некоторых случаях DLDP необходимо определять доступность IP в сегментах сети, не подключенных напрямую. Вам необходимо настроить IP-адрес next hop для



обнаруженного порта, чтобы позволить DLDP получить MAC-адрес next hop через пакет ARP перед отправкой правильного пакета ICMP.

В этой ситуации нужно избегать возврата ответного пакета по другому каналу; в противном случае DLDP ошибочно оценит, что обнаруженный порт не получает ответ ICMP.

Время восстановления DLDP

Время восстановления: укажите время, необходимое DLDP для получения последовательных пакетов ответа (ответ ICMP), прежде чем он сможет определить восстановление после сбоя канала.

В некоторых случаях обнаружение канала может быть нестабильным. Например, канал только периодически пингуется. В этом случае DLDP неоднократно меняет статус канала между Up и Down, что может еще больше дестабилизировать кольцевую сеть.

Время восстановления указывает время, необходимое DLDP для получения последовательных ответных пакетов, прежде чем DLDP сможет перевести канал из состояния Down в состояние Up. Время восстановления по умолчанию — три раза, что указывает на то, что канал должен быть успешно пропингован три раза, прежде чем он будет установлена в состояние Up. Настройка времени восстановления снижает чувствительность обнаружения каналов, но повышает стабильность. Связанные параметры настраиваются в зависимости от состояния сети.

MAC-адрес, привязанный к DLDP

Привязанный MAC-адрес: указывает MAC-адрес, связанный с обнаруженным IP-адресом.

В сложной сетевой среде DLDP может получить недопустимый MAC-адрес, если по обнаруженному каналу передаются аномальные пакеты ARP (вызывающие спуфинг ARP), что приведет к аномальному обнаружению DLDP.

Чтобы решить эту проблему, вы можете привязать обнаруженный IP-адрес (или IP-адрес next hop) к статическому MAC-адресу, чтобы избежать сбоя DLDP в случае спуфинга ARP.

3.3.2. Обзор

Особенность	Описание
Обнаружение DLDP	Обнаруживает соединение канала уровня 3. Когда канал уровня 3 неисправен, DLDP отключает порт уровня 3
Привязка MAC-адреса	Привязывает обнаруженный IP-адрес к MAC-адресу обнаруженного устройства, чтобы избежать исключений DLDP, в противном случае вызванных спуфингом ARP
Пассивное обнаружение DLDP	Когда оба конца обнаруженного канала включены с DLDP, вы можете настроить один конец в пассивном режиме для экономии пропускной способности и ресурсов ЦП

3.3.3. Обнаружение DLDP

DLDP обнаруживает подключение к каналу уровня 3. Когда канал уровня 3 неисправен, DLDP отключает соответствующий порт уровня 3.



3.3.3.1. Принцип работы

После включения обнаружения DLDP оно отправляет пакет ARP для получения MAC-адреса и исходящего порта обнаруженного устройства или устройства next-hop. Затем DLDP периодически отправляет эхо-пакеты IPv4 ICMP на MAC-адрес и исходящий порт для обнаружения соединения. Если DLDP не получает ответный пакет IPv4 ICMP от обнаруженного устройства в течение определенного периода времени, DLDP определяет, что канал неисправен, и устанавливает для порта уровня 3 значение Down.

3.3.3.2. Связанная конфигурация

- Включение обнаружения DLDP

По умолчанию обнаружение DLDP на портах отключено.

Запустите команду **dldp** с указанным обнаруженным IP-адресом, чтобы включить обнаружение DLDP.

Вы можете настроить IP-адрес next-hop, MAC-адрес обнаруженного устройства, интервал передачи, время повторной передачи и время восстановления в зависимости от фактической среды.

3.3.4. Привязка MAC-адреса

Функция привязки MAC-адреса используется для привязки обнаруженного IP-адреса (или IP-адреса next-hop) к MAC-адресу обнаруженного устройства (или устройства next-hop), чтобы избежать исключений DLDP, в противном случае вызванных спуфингом ARP.

3.3.4.1. Принцип работы

Вы можете привязать обнаруженный IP-адрес (или IP-адрес next-hop) к статическому MAC-адресу, чтобы избежать сбоя DLDP в случае спуфинга ARP.

3.3.4.2. Связанная конфигурация

По умолчанию MAC-адреса не привязаны к обнаружению DLDP.

Привяжите MAC-адрес обнаруженного устройства при запуске команды **dldp**, чтобы включить обнаружение DLDP. Если указан IP-адрес next-hop, привяжите MAC-адрес устройства next-hop.

После включения обнаружения DLDP оно отправляет пакеты ARP и ICMP с фиксированным IP-адресом получателя и фиксированным MAC-адресом получателя. Если исходный IP-адрес и MAC-адрес в полученном пакете не соответствуют связанным IP-адресу и MAC-адресу, DLDP не будет обрабатывать пакет.

3.3.5. Пассивное обнаружение DLDP

Когда оба конца обнаруженного канала включены с DLDP, вы можете настроить один конец в пассивном режиме для экономии пропускной способности и ресурсов ЦП.

3.3.5.1. Принцип работы

После того, как устройство на локальном конце отправит эхо-пакет ICMP, peer-устройство определяет возможность подключения к каналу в соответствии со временем приема пакета, используя определенные параметры обнаружения, которые совпадают с параметрами на локальном конце, тем самым экономя полосу пропускания и ресурсы ЦП.

3.3.5.2. Связанная конфигурация

По умолчанию пассивное обнаружение DLDP отключено.



Запустите команду **dldp passive**, чтобы включить пассивное обнаружение DLDP.

После включения пассивного обнаружения DLDP оно будет возвращать ответный пакет ICMP при получении эхо-пакета ICMP вместо активной отправки эхо-пакетов ICMP на peer end. Если DLDP не получает эхо-пакет ICMP в течение определенного периода, он определяет, что связь с портом peer ненормальна.

3.4. Конфигурация

Конфигурация	Описание и команда	
<u>Включение обнаружения DLDP</u>	(Обязательно) Используется для включения обнаружения DLDP в режиме конфигурации интерфейса	
	dldp	Включает обнаружение DLDP
	(Обязательно) Используется для включения пассивного обнаружения DLDP в режиме конфигурации интерфейса	
	dldp passive	Включает пассивное обнаружение DLDP
	(Опционально) Он используется для настройки интервала обнаружения, времени повторной передачи и времени восстановления обнаружения DLDP в режиме глобальной конфигурации	
	dldp interval	Глобально изменяет параметры DLDP, чтобы применить изменения к обнаружению DLDP на всех портах
	dldp retry	
	dldp resume	

3.4.1. Включение обнаружения DLDP

3.4.1.1. Эффект конфигурации

Обнаружение соединения канала уровня 3. Когда канал уровня 3 неисправен, DLDP отключает порт уровня 3.

3.4.1.2. Примечания

- DLDP поддерживает настройку нескольких IP-адресов на порту уровня 3. DLDP устанавливает для порта значение Down, когда ни один из IP-адресов не получает ответа ICMP. Если один IP-адрес возобновляет связь, DLDP снова устанавливает для порта значение Up.
- DLDP использует первый IP-адрес порта уровня 3 в качестве исходного IP-адреса пакетов обнаружения.



3.4.1.3. Шаги настройки

Включение обнаружения DLDP

- Обязательный.
- Когда вы включаете обнаружение DLDP в режиме конфигурации интерфейса, вы можете настроить IP-адрес *next-hop*, MAC-адрес, интервал передачи, время повторной передачи и время восстановления в зависимости от фактической среды.

Настройка режима обнаружения DLDP

- Опционально.
- Вы можете настроить активное или пассивное обнаружение DLDP в режиме конфигурации интерфейса в зависимости от фактической среды.
- Если необходимо включить обнаружение DLDP на обоих концах канала уровня 3, вы можете настроить пассивное обнаружение DLDP на одном конце, чтобы сэкономить ресурсы полосы пропускания и ЦП.

Глобальная настройка параметров DLDP

- Опционально.
- Вы можете изменить параметры обнаружения DLDP на всех портах в режиме глобальной конфигурации в зависимости от требований. Параметры включают интервал передачи пакетов, время повторной передачи пакетов и время восстановления.

3.4.1.4. Проверка

Отображение информации об устройстве DLDP, включая состояние и статистику обнаружения DLDP на всех портах.

3.4.1.5. Связанные команды

Включение обнаружения DLDP

Команда	<code>dldp ip-address [next-hop-ip] [mac-address mac-addr] [interval tick] [retry retry-num] [resume resume-num]</code>
Описание параметров	<i>ip-address</i>: указывает обнаруженный IP-адрес. <i>next-hop-ip</i>: указывает IP-адрес <i>next-hop</i>. <i>mac-addr</i>: указывает MAC-адрес обнаруженного устройства, которое необходимо связать. Если указан IP-адрес <i>next-hop</i>, привяжите MAC-адрес устройства <i>next-hop</i>. <i>tick</i>: указывает интервал, с которым передаются пакеты обнаружения. Диапазон значений составляет от 5 до 6000 тиков (1 тик = 10 мс). Значение по умолчанию — 100 тиков (1 с). <i>retry-num</i>: диапазон значений от 1 до 3600. Значение по умолчанию — 4. <i>resume-num</i>: указывает время восстановления. Диапазон значений от 1 до 200. Значение по умолчанию — 3
Командный режим	Режим конфигурации интерфейса



Руководство по использованию	Порт, для которого необходимо включить обнаружение DLDP, должен быть портом уровня 3, например, портом маршрутизатора, портом L3AP и портом SVI
------------------------------	---

Настройка режима обнаружения DLDP

Команда	dldp passive
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Перед настройкой режима обнаружения DLDP необходимо включить обнаружение DLDP

Глобальное изменение параметров обнаружения DLDP

Команда	dldp { interval tick retry retry-num resume resume-num }
Описание параметров	<i>tick</i>: указывает интервал, с которым передаются пакеты обнаружения. Диапазон значений составляет от 5 до 6000 тиков (1 тик = 10 мс). Значение по умолчанию — 100 тиков (1 с). <i>retry-num</i>: указывает интервал повторной передачи пакетов обнаружения. Значение колеблется от 4 до 3600. Значение по умолчанию — 4. <i>resume-num</i>: указывает время восстановления. Диапазон значений от 1 до 200. Значение по умолчанию — 3
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Используйте эту команду для быстрого изменения параметров обнаружения DLDP на всех портах при изменении фактической среды

Отображение состояния DLDP

Команда	show dldp statistic [interface interface-name]
Описание параметров	<i>interface-name</i> : указывает порт уровня 3, на котором будет отображаться статус DLDP
Командный режим	Привилегированный режим, режим глобальной конфигурации и режим конфигурации интерфейса
Руководство по использованию	Используйте эту команду для отображения состояния DLDP на определенном порту. Вы также можете использовать эту команду для отображения статуса DLDP на всех портах



3.4.1.6. Пример конфигурации

Включение обнаружения DLDP на портах уровня 3 на устройстве А и устройстве В в сети уровня 3

Сценарий:

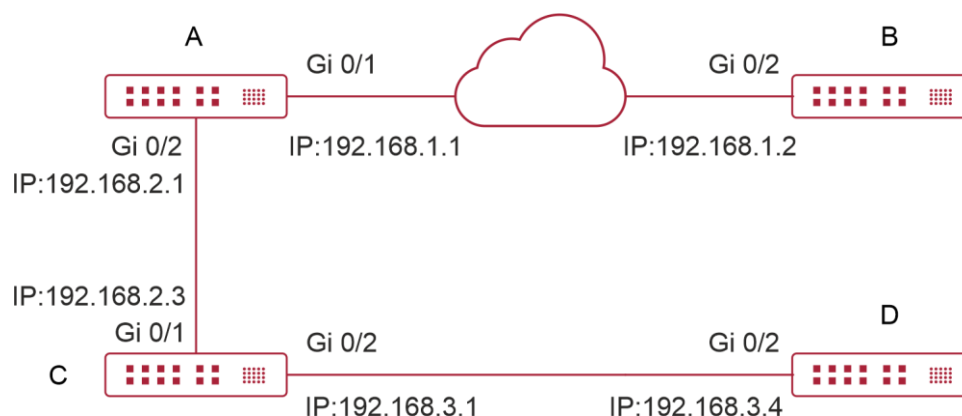


Рисунок 3-3.

Проверка	- Включите обнаружение DLDP на портах маршрутизатора Gi 0/1 и Gi 0/2 на устройстве A, чтобы обнаружить подключение канала уровня 3 между устройством A и устройством B, а также между устройством A и устройством D. - Чтобы управлять портом маршрутизатора Gi 0/2 устройства B, включите пассивное обнаружение DLDP на порту							
A	A#configure terminal A(config)#interface GigabitEthernet 0/1 A(config-if-GigabitEthernet 0/1)#dldp 192.168.1.2 A(config-if-GigabitEthernet 0/1)# exit A(config)#interface GigabitEthernet 0/2 A(config-if-GigabitEthernet 0/1)#dldp 192.168.3.4 192.168.2.3							
B	B#configure terminal B(config)#interface GigabitEthernet 0/2 B(config-if-GigabitEthernet 0/1)#dldp 192.168.1.1 B(config-if-GigabitEthernet 0/1)#dldp passive							
Проверка	Отобразите состояние DLDP на устройстве A и устройстве B, чтобы проверить, включено ли обнаружение DLDP и работает ли оно нормально							
A	A# show dldp <table><thead><tr><th>Interface</th><th>Type</th><th>Ip Next-hop</th><th>Interval</th><th>Retry</th><th>Resume</th><th>State</th></tr></thead></table>	Interface	Type	Ip Next-hop	Interval	Retry	Resume	State
Interface	Type	Ip Next-hop	Interval	Retry	Resume	State		



	Gi0/1	Active	192.168.1.2	100	4	3	Up
	Gi0/1	Active	192.168.3.4 192.168.2.3	100	4	3	Up
B	B# show dldp						
	Interface	Type	Ip Next-hop	Interval	Retry	Resume	State

	Gi0/2	Passive	192.168.1.1	100	4	3	Up

3.4.1.7. Распространенные ошибки

- Недостижимый одноадресный маршрут IPv4 ошибочно принимается за сбой обнаружения DLDP.
- Обнаружение DLDP завершается ошибкой, так как реер-устройство не поддерживает ответы ARP/ICMP.
- IP-адрес next-hop не настроен при обнаружении DLDP в межсетевом сегменте.

3.5. Мониторинг

3.5.1. Очистка

Описание	Команда
Очищает статистику DLDP	clear dldp [interface <i>interface-name</i> [<i>ip-address</i>]]

3.5.2. Отображение

Описание	Команда
Отображает статус DLDP	show dldp [interface <i>interface-name</i>]
Отображает статистику DLDP по состояниям портов Up/Down	show dldp statistic



4. НАСТРОЙКА VRRP

4.1. Обзор

Протокол резервирования виртуального маршрутизатора (VRRP) — отказоустойчивый протокол маршрутизации.

VRRP использует схему Master-backup (главный-резервный), чтобы обеспечить перенос функций с Master-маршрутизатора на резервный в случае сбоя Master-маршрутизатора, не влияя на внутреннюю и внешнюю передачу данных или изменяя конфигурацию локальной сети (LAN). Группа VRRP сопоставляет несколько маршрутизаторов с одним виртуальным маршрутизатором. VRRP гарантирует, что только один маршрутизатор в данный момент от имени виртуального маршрутизатора передает пакеты, который является выбранным Master. В случае отказа Master-маршрутизатора его заменит один из резервных маршрутизаторов. При VRRP кажется, что хост в локальной сети использует только один маршрутизатор, и маршрутизация остается работоспособной даже в случае сбоя first-hop маршрутизатора.

- VRRP применим к сценариям LAN, которые требуют избыточности выходов маршрутизации.

4.1.1. Протоколы и стандарты

- RFC2338: протокол резервирования виртуального маршрутизатора.
- RFC3768: протокол резервирования виртуального маршрутизатора (VRRP).
- RFC5798: протокол избыточности виртуального маршрутизатора (VRRP) версии 3 для IPv4 и IPv6.

4.2. Приложения

Приложение	Описание
<u>Избыточность маршрутизации</u>	Настройте маршрутизаторы в локальной сети как одну группу VRRP, чтобы обеспечить простую избыточность маршрутизации
<u>Балансировка нагрузки</u>	Настройте маршрутизаторы в локальной сети как несколько групп VRRP для балансировки нагрузки трафика

4.2.1. Избыточность маршрутизации

4.2.1.1. Сценарий

Настройте маршрутизаторы в LAN как одну группу VRRP, где хосты используют виртуальный IP-адрес этой группы в качестве адреса шлюза по умолчанию.

- Пакеты от хоста 1, хоста 2 и хоста 3 в другие сети перенаправляются выбранным Master-маршрутизатором (маршрутизатор А на Рисунке 4-1).

Если маршрутизатор А выходит из строя, Master будет переизбран между маршрутизатором В и маршрутизатором С для пересылки пакетов, обеспечивая простую избыточность маршрутизации.

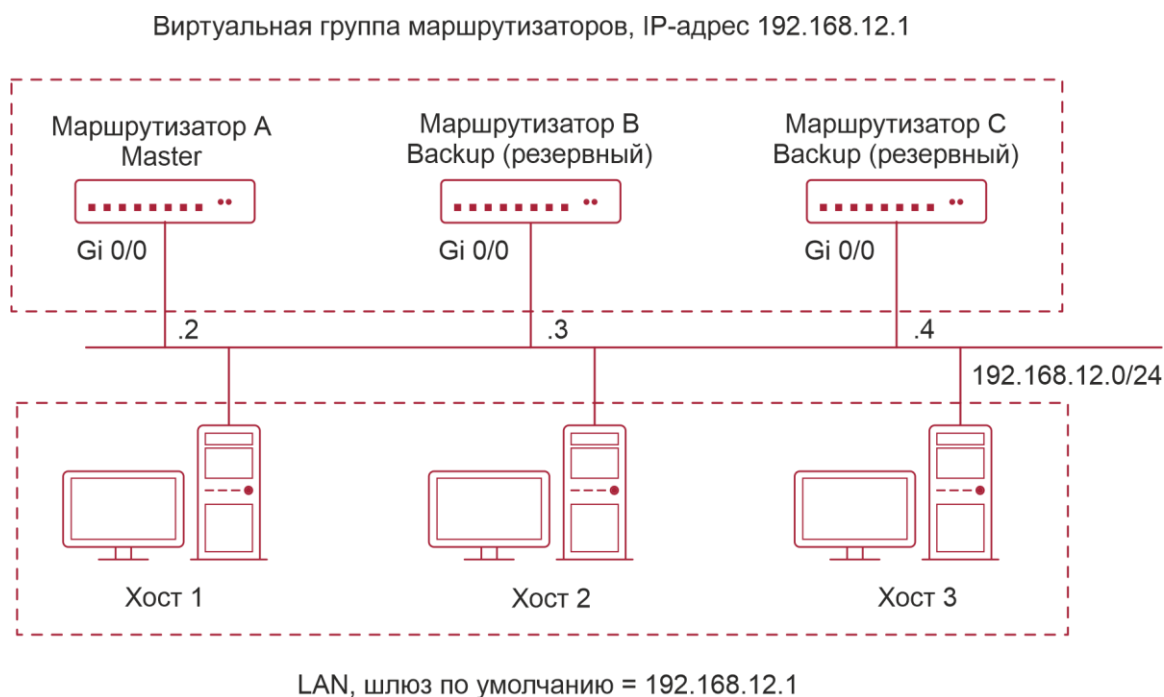


Рисунок 4-1.

4.2.1.2. Развертывание

- Маршрутизатор A, маршрутизатор B и маршрутизатор C подключены к LAN через интерфейсы Ethernet.
- На маршрутизаторе A, маршрутизаторе B и маршрутизаторе C VRRP настроен на интерфейсах Ethernet, подключенных к LAN.
- Эти интерфейсы Ethernet находятся в одной группе VRRP с виртуальным IP-адресом 192.168.12.1.
- Адрес шлюза для хоста 1, хоста 2 и хоста 3 — это IP-адрес группы VRRP, а именно 192.168.12.1.

4.2.2. Балансировка нагрузки

4.2.2.1. Сценарий

Настройте маршрутизаторы в LAN как несколько групп VRRP. Хосты в LAN используют виртуальные IP-адреса групп в качестве своих шлюзов, и каждый маршрутизатор выполняет резервное копирование для других маршрутизаторов в другой группе.

- Пакеты от хоста 1 и хоста 2 в другие сети с адресом шлюза по умолчанию в качестве виртуального IP-адреса виртуального маршрутизатора 1 перенаправляются Master-устройством виртуального маршрутизатора 1 (маршрутизатором A на Рисунке 4-2).
- Пакеты от хоста 3 и хоста 4 в другие сети с адресом шлюза по умолчанию в качестве виртуального IP-адреса виртуального маршрутизатора 2 перенаправляются Master-устройством виртуального маршрутизатора 2 (маршрутизатор B на Рисунке 4-2).

- Избыточность маршрутизации достигается на маршрутизаторе А и маршрутизаторе В, а трафик LAN используется совместно для достижения балансировки нагрузки.

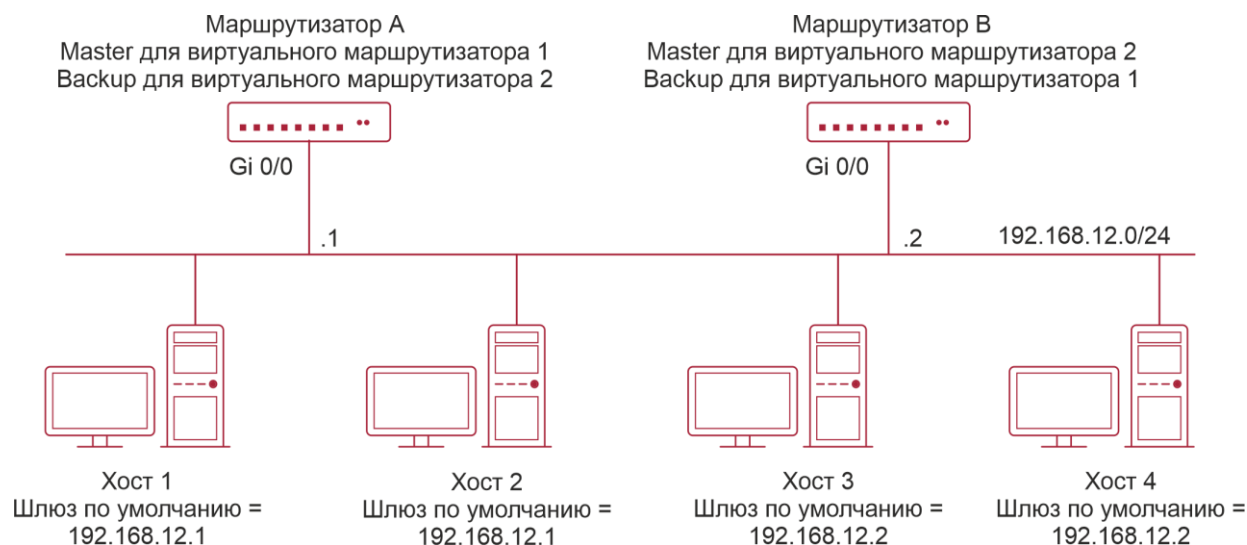


Рисунок 4-2.

4.2.2.2. Развертывание

- Маршрутизатор А и маршрутизатор В подключены к LAN через интерфейсы Ethernet.
- На маршрутизаторе А и маршрутизаторе В два виртуальных маршрутизатора настроены на интерфейсах Ethernet, подключенных к LAN.

Маршрутизатор А принимает IP-адрес 192.168.12.1 интерфейса Ethernet Gi0/0 в качестве IP-адреса виртуального маршрутизатора 1. Таким образом, для виртуального маршрутизатора 1 маршрутизатор А становится Master, а маршрутизатор В — Backup.

- Маршрутизатор В принимает IP-адрес 192.168.12.2 интерфейса Ethernet Gi0/0 в качестве IP-адреса виртуального маршрутизатора 2. Таким образом, для виртуального маршрутизатора 2 маршрутизатор В становится Master, а маршрутизатор А становится Backup.
- В LAN узел 1 и узел 2 используют IP-адрес 192.168.12.1 виртуального маршрутизатора 1 в качестве адреса шлюза по умолчанию, а узел 3 и узел 4 используют IP-адрес 192.168.12.2 виртуального маршрутизатора 2 в качестве адреса шлюза по умолчанию.

4.3. Функции

4.3.1. Базовые концепты

Виртуальный маршрутизатор

Виртуальный маршрутизатор, также называемый группой VRRP, считается шлюзом по умолчанию для хостов в LAN. Группа VRRP содержит идентификатор виртуального маршрутизатора (VRID) и набор виртуальных IP-адресов.



Виртуальный IP-адрес

Указывает IP-адрес виртуального маршрутизатора. Виртуальный маршрутизатор может быть настроен с одним или несколькими IP-адресами.

Владелец IP-адреса

Если группа VRRP имеет виртуальный IP-адрес, как у интерфейса Ethernet на одном реальном маршрутизаторе, маршрутизатор считается владельцем виртуального IP-адреса. В таком случае приоритет маршрутизатора равен 255. Если доступен собственный интерфейс Ethernet, VRRP группа автоматически перейдет в состояние Master. Владелец IP-адреса получает и обрабатывает пакеты с IP-адресом назначения как у виртуального маршрутизатора.

Виртуальный MAC-адрес

Виртуальный MAC-адрес группы VRRP — это MAC-адрес IEEE 802, отформатированный как 00-00-5E-00-01-{VRID} с назначенными первыми пятью октетами, а последние два — групповым VRID. Группа VRRP отвечает на запрос протокола разрешения адресов (ARP) своим виртуальным MAC-адресом вместо реального MAC-адреса.

Master-маршрутизатор

В группе VRRP только Master-маршрутизатор отвечает на запросы ARP и пересылает IP-пакеты. Если реальный маршрутизатор является владельцем IP-адреса, он становится Master-маршрутизатором.

Backup-маршрутизатор

В группе VRRP Backup-маршрутизаторы контролируют только состояние Master-маршрутизатора, но не отвечают на запросы ARP и не пересылают IP-пакеты. В случае сбоя Master Backup-маршрутизаторы воспользуются шансом побороться за позицию.

Preemption Mode

Если группа VRRP работает в preemption mode, Backup-маршрутизатор с более высоким приоритетом заменит Master-маршрутизатор с более низким приоритетом.

4.3.2. Обзор

Особенность	Описание
<u>VRRP</u>	VRRP обеспечивает избыточность шлюзов терминалов по умолчанию в среде с множественным доступом (например, Ethernet). Это позволяет Backup-маршрутизатору пересылать пакеты, когда Master-маршрутизатор не работает, обеспечивая прозрачное переключение маршрутизации и повышая качество сетевых услуг

4.3.3. VRRP

В случае, если Master-маршрутизатор неисправен, VRRP обеспечивает перенос функций с Master-маршрутизатора на Backup, не влияя на внутреннюю и внешнюю передачу данных или изменяя конфигурацию локальной сети.



4.3.3.1. Принцип работы

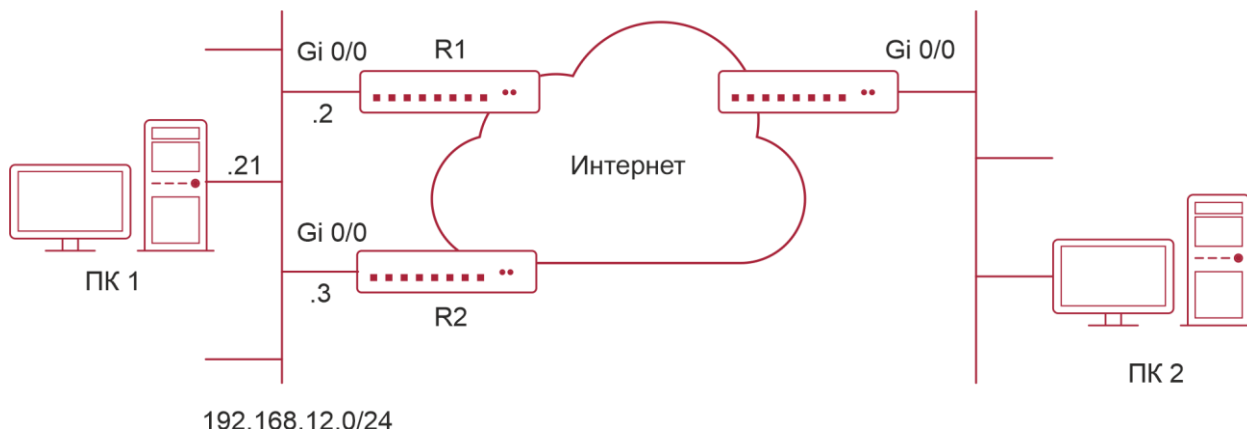


Рисунок 4-3. Принцип работы VRRP

Режим работы VRRP

Протоколы RFC2338, RFC3768 и RFC5798 определяют формат и механизм работы пакетов VRRP. Многоадресные пакеты VRRP периодически отправляются Master-маршрутизатором с указанными адресами назначения для объявления нормальной работы или для выбора Master. VRRP позволяет маршрутизатору в локальной сети автоматически заменять Master, который пересылает IP-пакеты в случае сбоя последнего. Это помогает обеспечить оперативное резервирование и отказоустойчивость маршрутизации на основе IP, а также обеспечить непрерывность и надежность связи для хостов в локальной сети. Группа VRRP обеспечивает избыточность за счет нескольких реальных маршрутизаторов. Однако только один маршрутизатор действует как Master для пересылки пакетов, а остальные являются Backup-маршрутизаторами. Переключение маршрутизатора в группе VRRP полностью прозрачно для хостов в локальной сети.

Master Election Process (процесс выбора Master)

Стандарты RFC определяют процесс выбора Master следующим образом:

- VRRP обеспечивает простой механизм выбора Master. Сначала сравните приоритеты VRRP, настроенные на интерфейсах маршрутизаторов в группе VRRP. Маршрутизатор с наивысшим приоритетом выбирается в качестве Master. Если эти приоритеты равны, сравните первичные IP-адреса этих маршрутизаторов. Маршрутизатор с наибольшим IP-адресом выбирается в качестве Master.
- После выбора Master-маршрутизатора другие маршрутизаторы становятся Backup-маршрутизаторами (и переходят в резервное состояние) и контролируют состояние Master-маршрутизатора с помощью пакетов VRRP, которые отправляет Master-маршрутизатор. Если Master-маршрутизатор работает, он регулярно отправляет многоадресные пакеты VRRP, известные как Advertisement-пакеты (объявления), чтобы уведомить Backup-маршрутизаторы о своем состоянии. Если Backup-маршрутизаторы не получают такие пакеты в течение установленного периода времени, все они перейдут в состояние Master. В таком случае предыдущий этап выбора Master повторяется. Таким образом, маршрутизатор с наивысшим приоритетом будет выбран в качестве нового Master, что обеспечит резервирование VRRP.

После выбора Master-маршрутизатора группы VRRP он отвечает за пересылку пакетов для хостов в локальной сети.



Коммуникационный процесс

Процесс связи VRRP можно пояснить на Рисунке 4-3. Маршрутизаторы R1 и R2 подключены к сегменту LAN 192.168.12.0/24 через интерфейсы Ethernet Gi0/0 с поддержкой VRRP. Хосты в локальной сети используют виртуальный IP-адрес группы VRRP в качестве адреса шлюза по умолчанию. Только виртуальный маршрутизатор распознается хостами. Однако Master-маршрутизатор в группе неизвестен. Например, когда ПК 1 планирует взаимодействовать с ПК 2, ПК 1 отправляет пакеты на шлюз по умолчанию с виртуальным IP-адресом; Master-маршрутизатор в группе получает пакеты и пересылает их на ПК 2. В этом процессе ПК 1 воспринимает только виртуальный маршрутизатор, а не R1 или R2. Master-маршрутизатор в группе выбирается между R1 и R2. В случае сбоя Master он будет автоматически заменен другим маршрутизатором.

4.3.3.2. Связанная конфигурация

Включение VRRP

По умолчанию VRRP отключен на интерфейсе.

В режиме конфигурации интерфейса запустите команду **vrrp group ip ipaddress [secondary]** or **vrrp group ipv6 ipv6-address**, чтобы установить VRID и виртуальный IP-адрес для включения VRRP.

VRRP должен быть включен на интерфейсе.

Настройка строки аутентификации IPv4 VRRP

По умолчанию VRRP находится в режиме без аутентификации.

Запустите команду **vrrp group authentication string**, чтобы установить строку аутентификации в режиме аутентификации MD5 или простой текстовый пароль в текстовом режиме для группы IPv4 VRRP. В режиме простой текстовой аутентификации пароль содержит не более 8 байт.

Члены группы VRRP могут общаться друг с другом, только если они находятся в одном и том же режиме аутентификации. В режиме простой текстовой аутентификации все маршрутизаторы в группе VRRP должны иметь одинаковый пароль аутентификации. Простой текстовый пароль аутентификации не может гарантировать безопасность, а только предотвращает/подсказывает неправильные конфигурации VRRP.

Настройка интервала объявления VRRP

По умолчанию интервал объявления Master-маршрутизатора составляет 1 секунду.

Запустите команду **vrrp [ipv6] group timers advertise { advertise-interval | csec centisecond-interval }** для изменения интервала.

Если таймер обучения VRRP не настроен, такой же интервал объявления должен быть установлен для группы VRRP, иначе маршрутизаторы в состоянии Backup будут отбрасывать полученные пакеты VRRP.

Настройка preemption mode VRRP

По умолчанию группа VRRP работает в preemption mode.

Чтобы включить preemption mode для группы VRRP, выполните команду **vrrp [ipv6] group preempt [delay seconds]**. Необязательный параметр **delay seconds** по умолчанию равен 0.

Если группа VRRP работает в preemption mode, маршрутизатор станет Master группы, когда обнаружит, что его приоритет выше, чем у текущего Master. Если группа VRRP работает в режиме без приоритета, маршрутизатор не станет Master, даже если обнаружит, что его приоритет выше, чем у текущего Master. Нет большого смысла настраивать preemption mode, когда группа VRRP использует IP-адрес



Ethernet-интерфейса, в этом случае группа имеет наивысший приоритет и автоматически становится Master в группе. Необязательный параметр **delay seconds** определяет задержку перед тем, как Backup-маршрутизатор VRRP объявит свой идентификатор Master.

Включение режима Асцепт (режима приема) IPv6 VRRP

По умолчанию режим Асцепт отключен для группы IPv6 VRRP.

Чтобы включить режим Асцепт, запустите команду **vrrp ipv6 group accept_mode**.

После включения режима Асцепт виртуальный маршрутизатор IPv6 VRRP в состоянии Master получает и обрабатывает пакеты с IP-адресом виртуального маршрутизатора в качестве пункта назначения; когда режим Асцепт отключен, виртуальный маршрутизатор отбрасывает такие пакеты, кроме пакетов Neighbor Advertisement (NA) и Neighbor Solicitation (NS). Кроме того, Master виртуальный маршрутизатор IPv6 VRRP в состоянии Owner получает и обрабатывает пакеты с IP-адресом виртуального маршрутизатора в качестве пункта назначения по умолчанию, независимо от того, настроен режим Асцепт или нет.

Настройка приоритета маршрутизатора VRRP

По умолчанию все приоритеты маршрутизатора в группе VRRP равны 100.

Чтобы настроить приоритет, запустите команду **vrrp [ipv6] group priority level**.

Если маршрутизатор в preemption mode владеет виртуальным IP-адресом группы и наивысшим приоритетом, он становится Master в группе, а другие маршрутизаторы с более низким приоритетом в группе становятся Backup (или мониторинговыми) маршрутизаторами.

Настройка отслеживаемого интерфейса VRRP

По умолчанию ни один интерфейс не отслеживается группой VRRP.

Чтобы настроить такой интерфейс, запустите команду **vrrp group track { interface-type interface-number | bfd interface-type interface-number ipv4-address } [priority]** or **vrrp ipv6 group track interface-type interface-number [priority]**.

После настройки интерфейса для мониторинга группы VRRP приоритет маршрутизатора будет динамически регулироваться в зависимости от состояния интерфейса. Как только интерфейс станет недоступным, приоритет маршрутизатора в группе будет снижен на заданное значение, а другой функциональный маршрутизатор с более высоким приоритетом в этой группе станет Master.

Настройка отслеживаемого IP-адреса VRRP

По умолчанию группа VRRP не отслеживает IP-адреса.

Чтобы настроить такой адрес, запустите команду **vrrp group track ip-address [interval interval-value] [timeout timeout-value] [retry retry-value] [priority]** or **vrrp ipv6 group track { ipv6-global-address | { ipv6-linklocal-address interface-type interface-number } } [interval interval-value] [timeout timeout-value] [retry retry-value] [priority]**.

После настройки IP-адреса для мониторинга группы VRRP приоритет маршрутизатора будет регулироваться динамически в зависимости от доступности адреса. Если адрес недоступен (команда **ping** не удалась), приоритет маршрутизатора в группе будет снижен на заданное значение, а другой маршрутизатор с более высоким приоритетом в этой группе станет Master.

Настройка таймера обучения VRRP

По умолчанию таймер обучения отключен для группы VRRP.

Чтобы включить его, запустите команду **vrrp [ipv6] group timers learn**.



После настройки таймера обучения Backup-маршрутизатор VRRP узнает интервал объявления пакетов NA от Master-устройства. На основании этого вместо локально установленного интервала Backup-маршрутизатор рассчитывает интервал для определения выхода из строя Master. Эта команда обеспечивает синхронизацию интервалов объявлений между Backup-маршрутизаторами и Master.

Настройка описания группы VRRP

По умолчанию для группы VRRP не настроено описание.

Чтобы настроить такую строку, запустите команду **vrrp [ipv6] group description text**.

Описание VRRP помогает различать группы VRRP. Описание имеет не более 80 байт, в противном случае предлагается неправильная конфигурация.

Настройка задержки VRRP

По умолчанию для группы VRRP не настроена задержка.

Чтобы включить её, запустите команду **vrrp delay { minimum min-seconds | reload reload-seconds }**. Два типа задержки варьируются от 0 до 60 секунд.

Команда настраивает задержку запуска группы VRRP на интерфейсе. Существует два типа задержки VRRP: задержка после запуска системы и задержка после возобновления работы интерфейса. Вы можете настроить их соответственно или одновременно. После того, как задержка настроена для группы VRRP на интерфейсе, группа VRRP запускается после задержки, а не сразу после запуска системы или возобновления работы интерфейса, что обеспечивает отсутствие preemption. Если интерфейс получает пакет VRRP во время задержки, задержка будет отменена, и VRRP будет запущен немедленно. Эта конфигурация будет эффективна как для групп VRRP IPv4, так и для групп IPv6 интерфейса.

Настройка версии IPv4 VRRP

По умолчанию IPv4 использует стандарт VRRPv2.

Чтобы указать версию для IPv4 VRRP, запустите команду **vrrp group version { 2 | 3 }**.

Когда значение параметра установлено на 2, принимается VRRPv2; когда значение параметра установлено на 3, принимается VRRPv3.

Указание Sub VLAN в Super VLAN для получения пакетов IPv4 VRRP

По умолчанию пакеты IPv4 VRRP отправляются на первый Up Sub VLAN-интерфейс в Super VLAN.

Чтобы указать первую Sub VLAN в состоянии Up Super VLAN для приема пакетов IPv4 VRRP, выполните команду **vrrp detection-vlan first-subvlan**; чтобы указать Sub VLAN, запустите команду **vrrp discovery-vlan subvlan-id**. Если VRRP и VRRP Plus включены одновременно на интерфейсе Super VLAN, пакеты VRRP отправляются на все Up-интерфейсы Sub VLAN в рамках Super VLAN.

Обе приведенные выше конфигурации уменьшают количество пакетов VRRP и позволяют избежать влияния на производительность маршрутизатора и использования пропускной способности сети. Тем не менее, маршрутизаторы, составляющие группу IPv4 VRRP, должны быть связаны между собой в пределах первого интерфейса UP Sub VLAN или указанной Sub VLAN Super VLAN.

Настройка поддержки BFD для IPv4 VRRP на интерфейсе

По умолчанию поддержка протокола двунаправленного обнаружения пересылки (BFD) для VRRP не включена на интерфейсе.

Чтобы включить его, запустите команду **vrrp group bfd ip-address**.

Для Backup-маршрутизатора запустите эту команду, чтобы сопоставить группу IPv4 VRRP с BFD, не заботясь о настроенном IP-адресе. Для Master, поскольку первичный IP-адрес Backup-маршрутизатора неизвестен, IP-адрес маршрутизатора может быть указан только администратором.

Чтобы включить поддержку BFD, убедитесь, что параметры сеанса IP и BFD настроены на целевом интерфейсе.

После того, как поддержка BFD включена для указанной группы IPv4 VRRP, при отказе Master-маршрутизатора Backup-маршрутизатор может обнаружить это в течение одной секунды.

Настройка глобального IPv4 VRRP BFD

По умолчанию VRRP не использует глобальный режим IPv4 VRRP BFD при определении состояния Master.

Чтобы включить глобальный IPv4 VRRP BFD, выполните команду **vrrp bfd interface-type interface-number ip-address**.

После включения глобального IPv4 VRRP BFD несколько групп IPv4 VRRP могут совместно использовать сеансы BFD, обеспечивая быстрое обнаружение и аварийное переключение Master-backup.

Чтобы включить поддержку BFD, убедитесь, что параметры сеанса IP и BFD настроены на целевом интерфейсе.

4.4. Конфигурация

Конфигурация	Описание и команда
<u>Настройка IPv4 VRRP</u>	(Обязательно) Он используется для включения IPv4 VRRP
	vrrp group ip ipaddress [secondary] Включает IPv4 VRRP
	(Опционально) Он используется для настройки параметров IPv4 VRRP
	vrrp group authentication string Настраивает строку аутентификации IPv4 VRRP
	vrrp group timers advertise { advertise-interval csec centisecond interval } Настраивает интервал объявления IPv4 VRRP
	vrrp group preempt [delay seconds] Настраивает preemption mode IPv4 VRRP
	vrrp group priority level Настраивает приоритет маршрутизатора IPv4 VRRP



Конфигурация	Описание и команда	
	vrrp group track { <i>interface-type interface-number</i> bfd <i>interface-type interface-number ipv4-address</i> } [<i>priority</i>]	Настраивает отслеживаемый интерфейс IPv4 VRRP
	vrrp group track <i>ip-address</i> [interval <i>interval-value</i>] [timeout <i>timeout value</i>] [retry <i>retry-value</i>] [<i>priority</i>]	Настраивает отслеживаемый IP-адрес IPv4 VRRP
	vrrp group timers learn	Настраивает таймер обучения IPv4 VRRP
	vrrp group description <i>text</i>	Настраивает описание группы IPv4 VRRP
<u>Настройка IPv4 VRRP</u>	vrrp delay { minimum <i>min-seconds</i> reload <i>reload-seconds</i> }	Настраивает задержку IPv4 VRRP
	vrrp group version { 2 3 }	Настраивает версию IPv4 VRRP
	vrrp detection-vlan { first-subvlan <i>subvlan-id</i> }	Указывает sub VLAN super VLAN для получения пакетов IPv4 VRRP
	vrrp group bfd <i>ip-address</i>	Настраивает поддержку BFD для IPv4 VRRP на интерфейсе
	vrrp bfd <i>interface-type interface-number ip-address</i>	Настраивает глобальный IPv4 VRRP BFD
<u>Настройка IPv6 VRRP</u>	(Обязательно) Он используется для включения IPv6 VRRP	
	vrrp group ipv6 <i>ipv6-address</i>	Включает IPv6 VRRP
	(Опционально) Он используется для настройки параметров IPv6 VRRP	
	vrrp ipv6 group timers advertise { <i>advertise-interval</i> csec <i>centisecond interval</i> }	Настраивает интервал объявления IPv6
	vrrp ipv6 group preempt [delay <i>seconds</i>]	Настраивает preemption mode IPv6 VRRP



Конфигурация	Описание и команда	
	vrrp ipv6 group accept_mode	Включает Accept mode для группы IPv6 VRRP
	vrrp ipv6 group priority level	Настраивает приоритет маршрутизатора IPv6 VRRP
	vrrp ipv6 group track interface-type interface-number [interface-priority]	Настраивает отслеживаемый интерфейс IPv6 VRRP
<u>Настройка IPv6 VRRP</u>	vrrp ipv6 group track { ipv6-global address { ipv6-linklocal-address interface-type interface-number } } [interval interval-value] [timeout timeout-value] [retry retry-value] [priority]	Настраивает отслеживаемый IP-адрес IPv6 VRRP
	vrrp ipv6 group timers learn	Настраивает таймер обучения IPv6 VRRP
	vrrp ipv6 group description text	Настраивает описание группы IPv6 VRRP
	vrrp delay { minimum min-seconds reload reload-seconds }	Настраивает задержку IPv6 VRRP
<u>Настройка VRRP-MSTP</u>	Конфигурация аналогична конфигурации IPv4 VRRP	

4.4.1. Настройка IPv4 VRRP

4.4.1.1. Эффект конфигурации

- Настройте группу VRRP на интерфейсе определенного сегмента локальной сети, задав VRID и виртуальный IP-адрес.
- Настройте несколько групп VRRP на интерфейсе, чтобы достичь балансировки нагрузки и предложить более стабильные и надежные сетевые услуги.
- Настройте отслеживаемые интерфейсы VRRP для мониторинга сбоев в режиме реального времени, изменения приоритетов интерфейсов и реализацию динамического переключения Master-backup.

4.4.1.2. Примечания

- Для достижения VRRP маршрутизаторы в группе VRRP должны быть настроены с одним и тем же виртуальным IPv4-адресом.



- Чтобы добиться взаимного резервирования между несколькими группами VRRP IPv4, настройте несколько групп VRRP IPv4 с одинаковой конфигурацией VRRP на другом интерфейсе и настройте для них разные приоритеты, чтобы они совместно действовали как Master и backup-группы.
- Включите VRRP на интерфейсах уровня 3.

4.4.1.3. Шаги настройки

Включение IPv4 VRRP

По умолчанию IPv4 VRRP отключен на интерфейсе. Вы можете включить его в зависимости от вашего требования.

Настройка строки аутентификации IPv4 VRRP

По умолчанию VRRP находится в режиме без аутентификации. Вы можете включить режим проверки подлинности с помощью простого текста в зависимости от ваших потребностей.

Настройка интервала объявления IPv4 VRRP

По умолчанию Master-маршрутизатор отправляет пакеты объявления каждую секунду. Вы можете изменить интервал в зависимости от ваших потребностей.

Настройка preemption mode IPv4 VRRP

По умолчанию группа VRRP работает в preemption mode с нулевой задержкой.

Настройка приоритета маршрутизатора IPv4 VRRP

Приоритет маршрутизатора по умолчанию для группы VRRP равен 100. Вы можете изменить приоритет в зависимости от ваших потребностей.

Настройка отслеживаемого интерфейса IPv4 VRRP

По умолчанию группа IPv4 VRRP не отслеживает интерфейс, а значение изменения приоритета равно 10. Чтобы обеспечить мониторинг сбоев посредством мониторинга интерфейса, настройте этот элемент.

Настройка таймера обучения IPv4 VRRP

По умолчанию таймер обучения отключен для группы VRRP. Включите эту функцию, если Backup-маршрутизаторам необходимо узнать интервал объявления Master-маршрутизатора.

Настройка группы IPv4 VRRP Описание

По умолчанию для группы VRRP не настроено описание. Чтобы четко различать группы VRRP, настройте описания.

Настройка задержки IPv4 VRRP

По умолчанию задержка IPv6 VRRP не настроена. Чтобы гарантировать эффективный режим без preemption mode, настройте задержку.

Настройка версии IPv4 VRRP

По умолчанию IPv4 использует стандарт VRRPv2. Для его изменения используйте соответствующую команду.

Указание Sub VLAN Super VLAN для получения пакетов IPv4 VRRP

По умолчанию пакеты IPv4 VRRP отправляются только на первый интерфейс UP Sub VLAN Super VLAN, но вы можете настроить конкретную Sub VLAN.



Настройка поддержки BFD для IPv4 VRRP на интерфейсе

По умолчанию поддержка BFD не настроена на интерфейсе. Для его настройки используйте соответствующую команду.

Настройка глобального IPv4 VRRP BFD

По умолчанию глобальный IPv4 VRRP BFD не включен. Для его реализации используйте соответствующую команду.

4.4.1.4. Проверка

Запустите команду **show vrrp**, чтобы проверить конфигурацию.

4.4.1.5. Связанные команды

Включение IPv4 VRRP

Команда	vrrp group ip ipaddress [secondary]
Описание параметров	<i>group</i> : указывает VRID группы VRRP, диапазон которых зависит от модели продукта. <i>ipaddress</i> : указывает IP-адрес группы VRRP. secondary : указывает дополнительный IP-адрес группы VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если виртуальный IP-адрес не указан, маршрутизаторы не могут присоединиться к группе VRRP. Если дополнительный IP-адрес не применяется, настроенный IP-адрес будет основным IP-адресом группы VRRP

Настройка строки аутентификации IPv4 VRRP

Команда	vrrp group authentication string
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>string</i> : указывает строку аутентификации группы VRRP (пароль в виде простого текста состоит не более чем из 8 байтов)
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	В группе VRRP для маршрутизаторов должен быть настроен один и тот же пароль аутентификации. Простой текстовый пароль аутентификации не может гарантировать безопасность, а только предотвращает/подсказывает неправильные конфигурации VRRP. Эта команда применима только к VRRPv2 вместо VRRPv3. Аутентификация отменена для пакетов VRRPv3 (IPv4 VRRP и IPv6 VRRP). Если для группы IPv4 VRRP выбран VRRPv2, команда действует; если выбран VRRPv3, команда неэффективна



Настройка интервала объявления IPv4 VRRP

Команда	vrrp group timers advertise { <i>advertise-interval</i> csec <i>centisecond-interval</i> }
Описание параметров	<i>group</i>: указывает VRID группы VRRP. <i>advertise-interval</i>: указывает интервал объявления группы VRRP (единица измерения: секунда). csec <i>centisecond-interval</i>: интервал, через который Master-маршрутизатор в резервной группе отправляет пакеты VRRP. Это целое число от 50 до 99. Единицей измерения является сантисекунда. Значение по умолчанию не указано. Команда эффективна только для пакетов VRRPv3. Если он настроен для пакетов VRRPv2, интервал по умолчанию равен одной секунде
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если маршрутизатор выбран в качестве Master в группе VRRP, он отправляет пакеты обновления VRRP с заданным интервалом, чтобы объявить о своем состоянии VRRP, приоритете и другой информации. Согласно стандартам RFC, если группа IPv4 VRRP использует VRRPv3 для отправки многоадресных пакетов, максимальный интервал объявления составляет 40 секунд. Таким образом, если задан интервал более 40 секунд, будет применяться этот максимальный интервал, несмотря на то, что конфигурация действует

Настройка preemption mode IPv4 VRRP

Команда	vrrp group preempt [<i>delay seconds</i>]
Описание параметров	<i>group</i>: указывает VRID группы VRRP. delay <i>seconds</i>: указывает задержку preemption для Master-маршрутизатора, чтобы заявить о своем статусе. Значение по умолчанию — 0 секунд
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если группа VRRP работает в preemption mode, маршрутизатор с более высоким приоритетом займет место Master с более низким приоритетом. Если группа VRRP работает в режиме без preemption mode, маршрутизатор с более высоким приоритетом, чем у Master, остается Backup. Нет большого смысла настраивать Preemption mode, когда группа VRRP использует IP-адрес Ethernet-интерфейса, в этом случае группа имеет наивысший приоритет и автоматически становится Master в группе



Настройка приоритета маршрутизатора IPv4 VRRP

Команда	vrrp group priority level
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>level</i> : указывает приоритет интерфейса в группе VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эта команда используется для ручной настройки приоритета маршрутизатора VRRP

Настройка отслеживаемого интерфейса IPv4 VRRP

Команда	vrrp group track { interface-type interface-number bfd interface-type interface-number ipv4-address } [priority]
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>interface-type interface-number</i> : указывает интерфейс для отслеживания. <i>bfd interface-type interface-number</i> : указанный соседний IP-адрес, отслеживаемый через BFD. <i>priority</i> : указывает масштаб изменения приоритета VRRP при изменении состояния контролируемого интерфейса. Значение по умолчанию — 10
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Отслеживаемый интерфейс должен быть маршрутизируемым логическим интерфейсом уровня 3 (например, Routed (маршрутизируемый) порт, интерфейс SVI, интерфейс Loopback или Tunnel (туннельный) интерфейс). Приоритет маршрутизатора, которому принадлежит виртуальный IP-адрес, связанный с группой VRRP, должен быть 255, и на нем нельзя настроить отслеживаемый интерфейс

Настройка отслеживаемого IP-адреса IPv4 VRRP

Команда	vrrp group track ipv4-address [interval interval-value] [timeout timeout-value] [retry retry-value] [priority]
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>ipv4-address</i> : указывает отслеживаемый IPv4-адрес.



	interval <i>interval-value</i>: указывает интервал проверки. Единица секунды. Если не настроено вручную, значение по умолчанию равно 3 секундам. timeout <i>timeout-value</i>: указывает тайм-аут датчика ожидания ответов. Если по истечении времени ожидания ответ не получен, считается, что место назначения недоступно. Единица секунды. Если не настроено вручную, значение по умолчанию равно 1 секунде. retry <i>retry-value</i>: указывает на повторные попытки датчика. Если тестовый пакет отправляется непрерывно в течение времени, равного значению повторной попытки, но ответа не получено, считается, что пункт назначения недоступен. Единица факт попытки. Если не настроено, значение по умолчанию равно 3 раза. priority: указывает масштаб изменения приоритета VRRP при изменении состояния контролируемого интерфейса. Значение по умолчанию — 10
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Для мониторинга хоста укажите его IPv4-адрес для группы IPv4 VRRP. Если группа VRRP владеет фактическим IP-адресом интерфейса Ethernet, приоритет группы равен 255, а отслеживаемый IP-адрес не может быть настроен

Настройка таймера обучения IPv4 VRRP

Команда	vrrp group timers learn
Описание параметров	<i>group</i> : указывает VRID группы VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Как только таймер обучения включен на маршрутизаторе VRRP, Backup-маршрутизатор узнает интервал объявления Master-устройства в течение таймера. На основании этого Backup-маршрутизатор рассчитывает интервал для определения отказа Master-маршрутизатора вместо использования локально настроенного интервала объявления. Эта команда обеспечивает синхронизацию с таймером обучения между главным и Backup-маршрутизаторами



Настройка описания группы IPv4 VRRP

Команда	vrrp group description text
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>text</i> : указывает описание группы VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Описание VRRP помогает различать группы VRRP. Описание имеет не более 80 байт, в противном случае предлагается неправильная конфигурация

Настройка задержки IPv4 VRRP

Команда	vrrp delay { minimum min-seconds reload reload-seconds }
Описание параметров	minimum min-seconds : указывает задержку VRRP после изменения состояния интерфейса. reload reload-seconds : указывает задержку VRRP после запуска системы
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После того, как задержка настроена для группы VRRP на интерфейсе, группа VRRP запускается после задержки, а не сразу после запуска системы или возобновления работы интерфейса, что обеспечивает отсутствие preemption. Если интерфейс получает пакет VRRP во время задержки, задержка будет отменена, и VRRP будет запущен немедленно. Два типа задержки имеют общий диапазон значений от 0 до 60 секунд. Эта конфигурация будет эффективна как для групп VRRP IPv4, так и для групп IPv6 интерфейса

Настройка версии IPv4 VRRP

Команда	vrrp group version { 2 3 }
Описание параметров	2 : указывает на VRRPv2. 3 : указывает на VRRPv3
Командный режим	Режим конфигурации интерфейса



Руководство по использованию	Учитывая совместимость между VRRPv2 и VRRPv3, укажите стандарт для IPv4 VRRP на основе фактического состояния сети. VRRPv2 разработан в RFC3768, а VRRPv3 описан в RFC5798. Эта команда применима только к IPv4 VRRP
------------------------------	---

Указание Sub VLAN Super VLAN для получения пакетов IPv4 VRRP

Команда	vrrp detection-vlan {first-subvlan subvlan-id}
Описание параметров	first-subvlan: отправляет пакеты IPv4 VRRP только на первый UP Sub VLAN-интерфейс в Super VLAN. subvlan-id: отправляет пакеты IPv4 VRRP в указанную Sub VLAN
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эта команда используется для указания Sub VLAN Super VLAN для получения пакетов IPv4 VRRP. Пакеты IPv4 VRRP отправляются в Super VLAN с использованием следующих трех методов. Пакеты отправляются на первый UP Sub VLAN-интерфейс в Super VLAN, или на указанный интерфейс Sub VLAN в Super VLAN, или на все интерфейсы Sub VLAN в Super VLAN. Если VRRP и VRRP Plus включены одновременно на интерфейсе Super VLAN, пакеты VRRP отправляются на все Up Sub VLAN-интерфейсы в рамках Super VLAN. Эта команда настраивается на интерфейсе VLAN и эффективна только для интерфейсов Super VLAN

Настройка поддержки BFD для IPv4 VRRP на интерфейсе

Команда	vrrp group bfd ip-address
Описание параметров	group: указывает VRID группы VRRP. ip-address: указывает IP-адрес интерфейса
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Для Ваксир-маршрутизатора запустите эту команду, чтобы сопоставить группу IPv4 VRRP с BFD, не заботясь о настроенном IP-адресе. Для Master, поскольку первичный IP-адрес Ваксир-маршрутизатора неизвестен, IP-адрес маршрутизатора может быть указан только администратором. Если настроен глобальный IPv4 VRRP BFD, эта конфигурация не может быть выполнена. Чтобы включить поддержку BFD, убедитесь, что параметры сеанса IP и BFD настроены на целевом интерфейсе



Настройка глобального IPv4 VRRP BFD

Команда	vrrp bfd interface-type interface-number ip-address
Описание параметров	<i>interface-type interface-number</i> : указывает тип интерфейса и идентификатор. <i>ip-address</i> : указывает IP-адрес интерфейса
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Если настроен глобальный IPv4 VRRP BFD, настроенная поддержка BFD будет удалена. Чтобы включить поддержку BFD, убедитесь, что параметры сеанса IP и BFD настроены на целевом интерфейсе. Глобальный сеанс IPv4 VRRP BFD применим только к группе IPv4 VRRP, состоящей из двух маршрутизаторов



4.4.1.6. Пример конфигурации

Настройка группы IPv4 VRRP и отслеживаемого интерфейса

Сценарий:

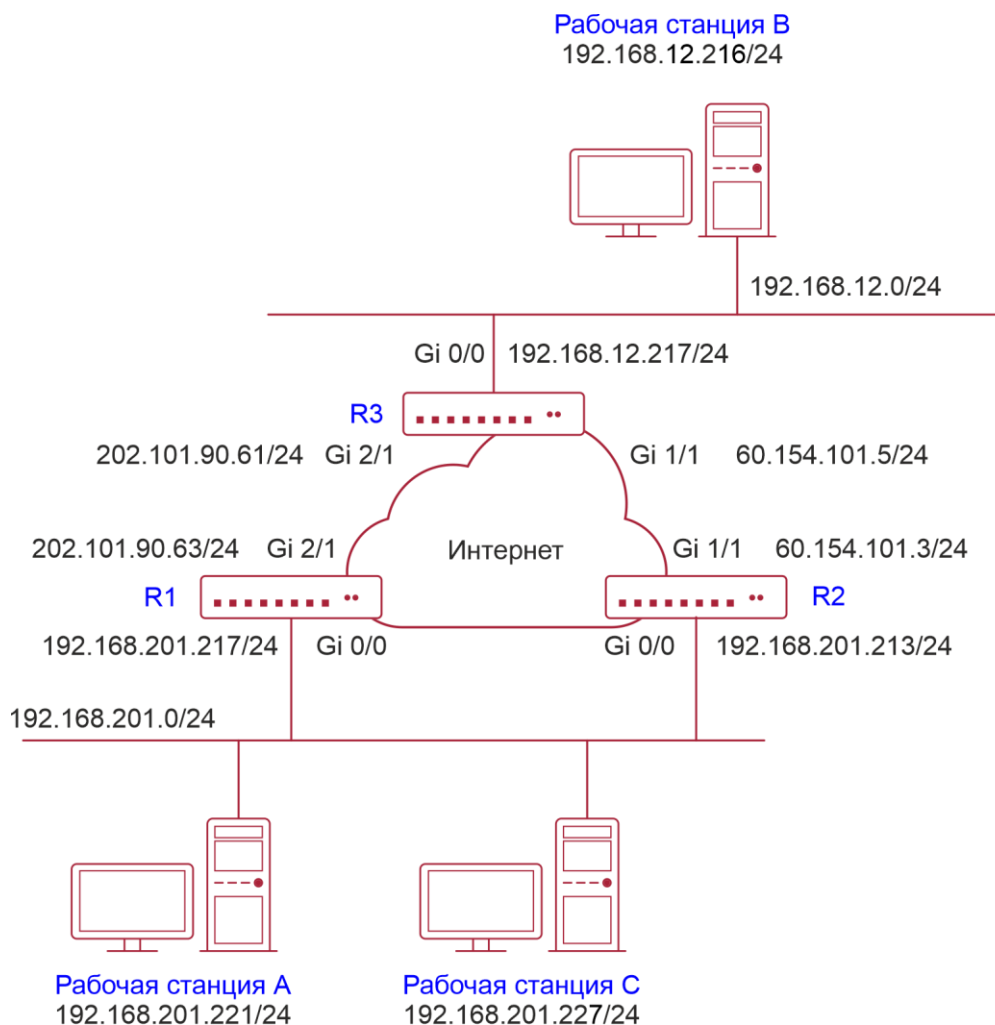


Рисунок 4-4.

Шаги настройки	- Кластер рабочей станции А и рабочей станции В (192.168.201.0/24) использует виртуальный IP-адрес 192.168.201.1 группы VRRP, состоящей из маршрутизаторов R1 и R2, в качестве адреса шлюза для связи с рабочей станцией В (192.168.12.0). /24). - GigabitEthernet 2/1 на R1 настроен как отслеживаемый интерфейс. - Нет VRRP, но на R3 настроена обычная функция маршрутизации
R3	<pre> R3#configure terminal R3(config)#interface GigabitEthernet 0/0 // Команда «no switchport» требуется только для коммутатора. R3(config-if-GigabitEthernet 0/0)#no switchport </pre>



	<pre> R3(config-if-GigabitEthernet 0/0)#ip address 192.168.12.217 255.255.255.0 R3(config-if-GigabitEthernet 0/0)#exit R3(config)#interface GigabitEthernet 1/1 // Команда «no switchport» требуется только для коммутатора. R3(config-if-GigabitEthernet 1/1)#no switchport R3(config-if-GigabitEthernet 1/1)#ip address 60.154.101.5 255.255.255.0 R3(config-if-GigabitEthernet 1/1)#exit R3(config)#interface GigabitEthernet 2/1 // Команда «no switchport» требуется только для коммутатора. R3(config-if-GigabitEthernet 2/1)#no switchport R3(config-if-GigabitEthernet 2/1)#ip address 202.101.90.61 255.255.255.0 R3(config-if-GigabitEthernet 2/1)#exit R3(config)#router ospf R3(config-router)#network 202.101.90.0 0.0.0.255 area 10 R3(config-router)#network 192.168.12.0 0.0.0.255 area 10 R3(config-router)#network 60.154.101.0 0.0.0.255 area 10 </pre>
R1	<pre> R1#configure terminal R1(config)#interface GigabitEthernet 0/0 R1(config-if-GigabitEthernet 0/0)#ip address 192.168.201.217 255.255.255.0 R1(config-if-GigabitEthernet 0/0)#vrrp 1 priority 120 R1(config-if-GigabitEthernet 0/0)#vrrp 1 timers advertise 3 R1(config-if-GigabitEthernet 0/0)#vrrp 1 ip 192.168.201.1 R1(config-if-GigabitEthernet 0/0)#vrrp 1 track GigabitEthernet 2/1 30 R1(config-if-GigabitEthernet 0/0)#exit R1(config)#interface GigabitEthernet 2/1 R1(config-if-GigabitEthernet 2/1)#ip address 202.101.90.63 255.255.255.0 R1(config-if-GigabitEthernet 2/1)#exit R1(config)#router ospf R1(config-router)#network 202.101.90.0 0.0.0.255 area 10 R1(config-router)#network 192.168.201.0 0.0.0.255 area 10 </pre>
R2	<pre> R2#configure terminal R2(config)#interface GigabitEthernet 0/0 R2(config-if-GigabitEthernet 0/0)#ip address 192.168.201.213 255.255.255.0 R2(config-if-GigabitEthernet 0/0)#vrrp 1 ip 192.168.201.1 R2(config-if-GigabitEthernet 0/0)#vrrp 1 timers advertise 3 </pre>



	<pre> R2(config-if-GigabitEthernet 0/0)#exit R2(config)#interface GigabitEthernet 1/1 // Команда «no switchport» требуется только для коммутатора. R2(config-if-GigabitEthernet 1/1)#no switchport R2(config-if-GigabitEthernet 1/1)#ip address 60.154.101.3 255.255.255.0 R2(config-if-GigabitEthernet 1/1)#exit R2(config)#router ospf R2(config-router)#network 60.154.101.0 0.0.0.255 area 10 R2(config-router)#network 192.168.201.0 0.0.0.255 area 10 </pre>
Проверка	Запустите команду show vrrp, чтобы проверить конфигурацию. - Проверьте, снижает ли R1, выступающий в роли Master, свой приоритет VRRP с 120 до 90, когда GigabitEthernet2/1, подключенный к глобальной сети (WAN), недоступен. Если да, R2 становится Master. - Проверьте, восстанавливает ли R1 свой приоритет VRRP с 30 до 120, когда GigabitEthernet 2/1, подключенный к глобальной сети, восстанавливается. Если да, R1 переизбирается Master
R1	<pre> R1#show vrrp GigabitEthernet 0/0 - Group 1 State is Master Virtual IP address is 192.168.201.1 configured Virtual MAC address is 0000.5e00.0101 Advertisement interval is 3 sec Preemption is enabled min delay is 0 sec Priority is 120 Master Router is 192.168.201.217 (local), priority is 120 Master Down interval is 10.59 sec Tracking state of 1 interface, 1 up up GigabitEthernet 2/1 priority decrement=30 </pre>
R2	<pre> R2#show vrrp GigabitEthernet 0/0 - Group 1 State is Backup Virtual IP address is 192.168.201.1 configured Virtual MAC address is 0000.5e00.0101 Advertisement interval is 3 sec </pre>



	Preemption is enabled min delay is 0 sec Priority is 100 Master Router is 192.168.201.217 , priority is 120 Master Down interval is 10.82 sec
--	--

4.4.1.7. Распространенные ошибки

- На маршрутизаторах в группе VRRP настраиваются разные виртуальные IP-адреса, в результате чего в группе имеется несколько Master-маршрутизаторов.
- На маршрутизаторах в группе VRRP настроены разные интервалы объявлений VRRP, а таймер обучения не настроен, что приводит к множеству Master-маршрутизаторов в группе.
- Различные версии VRRP настроены на маршрутизаторах в группе VRRP, в результате чего в группе имеется несколько Master-маршрутизаторов.
- Для VRRPv2 все интерфейсы Ethernet маршрутизаторов в группе VRRP находятся в режиме простой текстовой аутентификации, но несовместимы в строках аутентификации, что приводит к множеству Master-маршрутизаторов в группе.



4.4.1.8. Пример конфигурации

Настройка нескольких групп IPv4 VRRP

Сценарий:

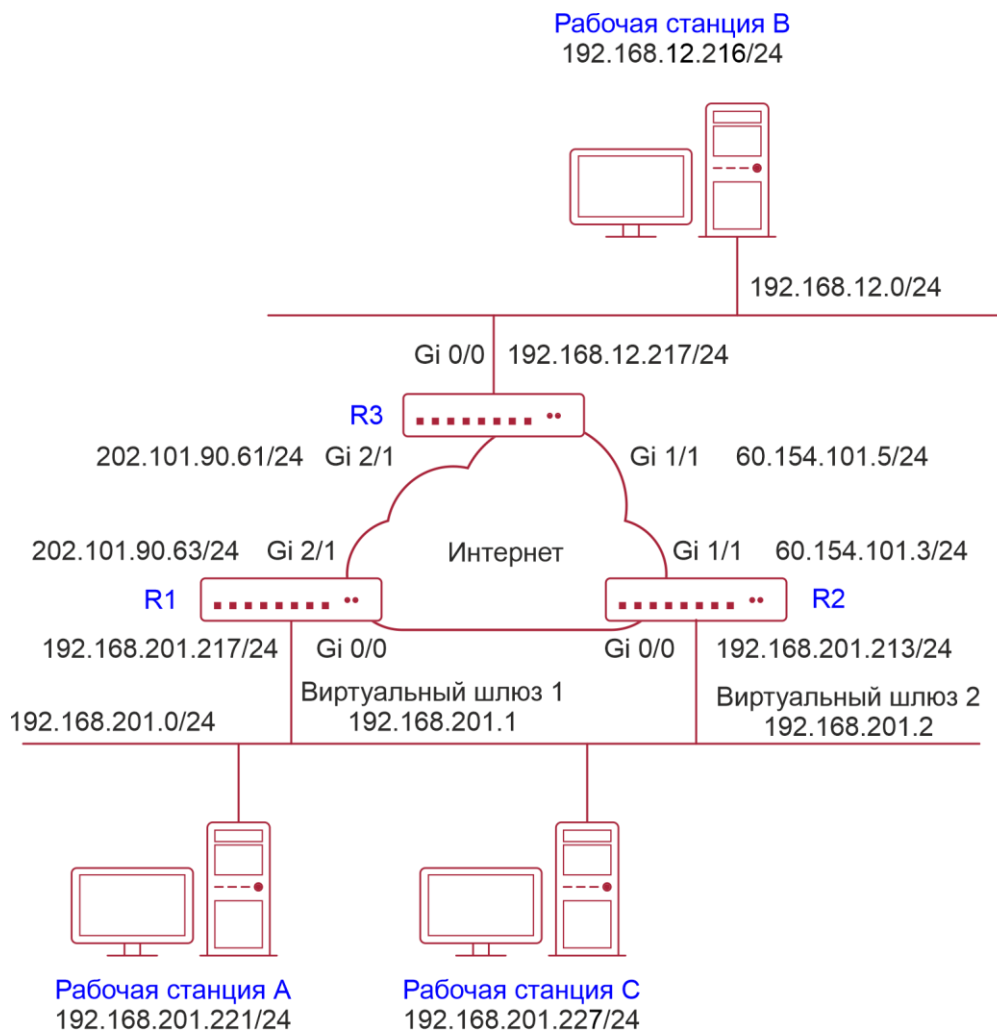


Рисунок 4-5.

Шаги настройки	- Кластер рабочих станций пользователей (192.168.201.0/24) использует Backup-группу, состоящую из маршрутизаторов R1 и R2. Шлюз для частичных рабочих станций (например, А) указывает на виртуальный IP-адрес 192.168.201.1 Backup-группы 1, а шлюз для других частичных рабочих станций (например, С) указывает на виртуальный IP-адрес 192.168.201.2 Backup-группы 2. Многоадресная маршрутизация IPv4 включена на всех маршрутизаторах. - R1 действует как Master-маршрутизатор в группе 2 и как Backup-маршрутизатор в группе 1. - R2 действует как Backup-маршрутизатор в группе 2 и как Master-маршрутизатор в группе 1
----------------	---



R3	<pre> R3#configure terminal R3(config)#interface GigabitEthernet 0/0 // Команда «no switchport» требуется только для коммутатора. R3(config-if-GigabitEthernet 0/0)#no switchport R3(config-if-GigabitEthernet 0/0)#ip address 192.168.12.217 255.255.255.0 R3(config-if-GigabitEthernet 0/0)#exit R3(config)#interface GigabitEthernet 1/1 // Команда «no switchport» требуется только для коммутатора. R3(config-if-GigabitEthernet 1/1)#no switchport R3(config-if-GigabitEthernet 1/1)#ip address 60.154.101.5 255.255.255.0 R3(config-if-GigabitEthernet 1/1)#exit R3(config)#interface GigabitEthernet 2/1 // Команда «no switchport» требуется только для коммутатора. R3(config-if-GigabitEthernet 2/1)#no switchport R3(config-if-GigabitEthernet 2/1)#ip address 202.101.90.61 255.255.255.0 R3(config-if-GigabitEthernet 2/1)#exit R3(config)#router ospf R3(config-router)#network 202.101.90.0 0.0.0.255 area 10 R3(config-router)#network 192.168.12.0 0.0.0.255 area 10 R3(config-router)#network 60.154.101.0 0.0.0.255 area 10 </pre>
R1	<pre> R1#configure terminal R1(config)#interface GigabitEthernet 0/0 R1(config-if-GigabitEthernet 0/0)#ip address 192.168.201.217 255.255.255.0 R1(config-if-GigabitEthernet 0/0)#vrrp 1 timers advertise 3 R1(config-if-GigabitEthernet 0/0)#vrrp 1 ip 192.168.201.1 R1(config-if-GigabitEthernet 0/0)#vrrp 2 priority 120 R1(config-if-GigabitEthernet 0/0)#vrrp 2 timers advertise 3 R1(config-if-GigabitEthernet 0/0)#vrrp 2 ip 192.168.201.2 R1(config-if-GigabitEthernet 0/0)#vrrp 2 track GigabitEthernet 2/1 30 R1(config-if-GigabitEthernet 0/0)#exit R1(config)#interface GigabitEthernet 2/1 R1(config-if-GigabitEthernet 2/1)#ip address 202.101.90.63 255.255.255.0 R1(config-if-GigabitEthernet 2/1)#exit R1(config)#router ospf R1(config-router)#network 202.101.90.0 0.0.0.255 area 10 </pre>



	R1(config-router)#network 192.168.201.0 0.0.0.255 area 10
R2	<pre> R2#configure terminal R2(config)#interface GigabitEthernet 0/0 R2(config-if-GigabitEthernet 0/0)#ip address 192.168.201.213 255.255.255.0 R2(config-if-GigabitEthernet 0/0)#vrrp 1 ip 192.168.201.1 R2(config-if-GigabitEthernet 0/0)#vrrp 1 timers advertise 3 R2(config-if-GigabitEthernet 0/0)#vrrp 1 priority 120 R2(config-if-GigabitEthernet 0/0)#vrrp 2 ip 192.168.201.2 R2(config-if-GigabitEthernet 0/0)#vrrp 2 timers advertise 3 R2(config-if-GigabitEthernet 0/0)#exit R2(config)#interface GigabitEthernet 1/1 R2(config-if-GigabitEthernet 1/1)#ip address 60.154.101.3 255.255.255.0 R2(config-if-GigabitEthernet 1/1)#exit R2(config)#router ospf R2(config-router)#network 60.154.101.0 0.0.0.255 area 10 R2(config-router)#network 192.168.201.0 0.0.0.255 area 10 </pre>
Проверка	Запустите команду show vrrp, чтобы проверить конфигурацию. - Проверьте, снижает ли маршрутизатор R1, выступающий в качестве Master-маршрутизатора в группе 2, приоритет группы VRRP с 30 до 90, когда обнаруживает, что интерфейс GigabitEthernet 2/1, подключенный к глобальной сети, недоступен. Если да, то R2 в группе 2 становится Master-маршрутизатором. - Проверьте, увеличивает ли R1 приоритет группы VRRP с 30 до 120, когда он обнаруживает, что интерфейс GigabitEthernet 2/1, подключенный к глобальной сети, снова становится доступным. Если да, то R1 снова становится Master-маршрутизатором в группе 2
R1	<pre> R1#show vrrp GigabitEthernet 0/0 - Group 1 State is Backup Virtual IP address is 192.168.201.1 configured Virtual MAC address is 0000.5e00.0101 Advertisement interval is 3 sec Preemption is enabled min delay is 0 sec Priority is 100 Master Router is 192.168.201.213 , priority is 120 </pre>



	Master Advertisement interval is 3 sec Master Down interval is 10.82 sec GigabitEthernet 0/0 - Group 2 State is Master Virtual IP address is 192.168.201.2 configured Virtual MAC address is 0000.5e00.0102 Advertisement interval is 3 sec Preemption is enabled min delay is 0 sec Priority is 120 Master Router is 192.168.201.217 (local), priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.59 sec Tracking state of 1 interface, 1 up: up GigabitEthernet 2/1 priority decrement=30
R2	R2#show vrrp GigabitEthernet 0/0 - Group 1 State is Master Virtual IP address is 192.168.201.1 configured Virtual MAC address is 0000.5e00.0101 Advertisement interval is 3 sec Preemption is enabled min delay is 0 sec Priority is 120 Master Router is 192.168.201.213 (local), priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.59 sec GigabitEthernet 0/0 - Group 2 State is Backup Virtual IP address is 192.168.201.2 configured Virtual MAC address is 0000.5e00.0102 Advertisement interval is 3 sec Preemption is enabled min delay is 0 sec Priority is 100 Master Router is 192.168.201.217 , priority is 120



	Master Advertisement interval is 3 sec Master Down interval is 10.82 sec
--	--

4.4.2. Настройка IPv6 VRRP

4.4.2.1. Эффект конфигурации

- Настройте группу IPv6 VRRP на интерфейсе определенного сегмента локальной сети, задав VRID и виртуальный IPv6-адрес.
- Настройте несколько групп IPv6 VRRP на интерфейсе для достижения баланса нагрузки и обеспечения более стабильных и надежных сетевых услуг.
- Настройте отслеживаемые интерфейсы VRRP для мониторинга сбоев в режиме реального времени, изменения приоритетов интерфейсов и динамического переключения на резервные копии.

4.4.2.2. Примечания

- Для достижения VRRP маршрутизаторы в группе VRRP должны быть настроены с одним и тем же виртуальным IPv6-адресом.
- Чтобы добиться взаимного резервного копирования для нескольких резервных групп IPv6 VRRP, вам необходимо настроить несколько групп IPv6 VRRP с идентичной конфигурацией VRRP на интерфейсе и настроить для них разные приоритеты, чтобы сделать маршрутизаторы Master и Backup взаимными.
- VRRP должен быть включен на интерфейсах уровня 3.

4.4.2.3. Шаги настройки

Включение IPv6 VRRP

По умолчанию IPv6 VRRP не включен на интерфейсе. Вы можете включить его в зависимости от вашего требования.

Настройка интервала объявления IPv6 VRRP

По умолчанию Master-маршрутизатор отправляет пакеты объявлений каждую секунду. Вы можете изменить интервал в зависимости от ваших потребностей.

Настройка preemption mode IPv6 VRRP

По умолчанию группа VRRP работает в preemption mode с нулевой задержкой.

Включение режима Ассепт для группы IPv6 VRRP

По умолчанию режим Ассепт отключен для группы IPv6 VRRP. Чтобы потребовать, чтобы группа IPv6 VRRP VRRP в состоянии Master получала и обрабатывала пакеты с IP-адресом назначения как у виртуального маршрутизатора, включите режим Ассепт.

Настройка приоритета маршрутизатора IPv6 VRRP

Приоритет маршрутизатора по умолчанию для группы VRRP равен 100. Вы можете изменить приоритет в зависимости от ваших потребностей.

Настройка отслеживаемого интерфейса IPv6 VRRP

По умолчанию отслеживаемый интерфейс не настроен. Вы можете изменить интервал в зависимости от ваших потребностей.



Настройка отслеживаемого IP-адреса IPv6 VRRP

По умолчанию отслеживаемый адрес IPv6 не настроен, а значение изменения приоритета равно 10. Вы можете настроить эту функцию по своему усмотрению.

Настраивает таймер обучения IPv6 VRRP

По умолчанию таймер обучения отключен для группы VRRP. Включите эту функцию, если Баскир-маршрутизаторам необходимо узнать интервал объявления Master-маршрутизатора.

Настройка описания группы IPv6 VRRP

По умолчанию для группы VRRP не настроено описание. Чтобы четко различать группы VRRP, настройте описания.

Настройка задержки IPv4 VRRP

По умолчанию задержка IPv6 VRRP не настроена. Чтобы гарантировать эффективный режим без preemption mode, настройте задержку.

4.4.2.4. Проверка

Запустите команду **show vrrp**, чтобы проверить конфигурацию.

4.4.2.5. Связанные команды

Включение IPv6 VRRP

Команда	vrrp group ipv6 ipv6-address
Описание параметров	<i>group</i> : указывает VRID группы VRRP, диапазон которых зависит от модели продукта. <i>ipv6-address</i> : указывает IPv6-адрес группы VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Группы VRRP IPv6 и группы VRRP IPv4 имеют общий диапазон VRID от 1 до 255. Один VRID применим к группе VRRP IPv4 и группе VRRP IPv6 одновременно. Первый настроенный адрес должен быть локальным адресом канала, который можно удалить только после других виртуальных адресов

Настройка интервала объявления IPv6 VRRP

Команда	vrrp ipv6 group timers advertise { advertise-interval csec centisecond-interval }
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>advertise-interval</i> : указывает интервал объявления группы VRRP (единица измерения: секунда). <i>csec centisecond-interval</i> : интервал, через который Master-маршрутизатор в резервной группе отправляет пакеты VRRP. Это целое число от 50 до 99. Единицей измерения является сантисекунда. Значение по умолчанию не указано. Команда



	эффективна только для пакетов VRRPv3. Если он настроен для пакетов VRRPv2, интервал по умолчанию равен одной секунде
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если маршрутизатор выбран в качестве Master в группе VRRP, он отправляет пакеты объявления VRRP с заданным интервалом, чтобы объявить о своем состоянии VRRP, приоритете и другой информации. Согласно стандартам RFC, если группа IPv6 VRRP использует VRRPv3 для отправки многоадресных пакетов, максимальный интервал объявления составляет 40 секунд. Таким образом, если задан интервал более 40 секунд, будет применяться этот максимальный интервал, несмотря на то, что конфигурация действует

Настройка preemption mode

Команда	vrrp ipv6 group preempt [delay seconds]
Описание параметров	<i>group</i> : указывает VRID группы VRRP. delay seconds : указывает задержку preemption для Master-маршрутизатора, чтобы заявить о своем статусе. Значение по умолчанию — 0 секунд
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если группа VRRP работает в preemption mode, маршрутизатор с более высоким приоритетом займет место Master с более низким приоритетом. Если группа VRRP работает в режиме без вытеснения, маршрутизатор с более высоким приоритетом, чем у ведущего, остается резервным. Нет большого смысла настраивать режим Preemption, когда группа VRRP использует IP-адрес Ethernet-интерфейса, в этом случае группа имеет наивысший приоритет и автоматически становится Master в группе

Включение режима Асепт для группы IPv6 VRRP

Команда	vrrp ipv6 group accept_mode
Описание параметров	<i>group</i> : указывает VRID группы VRRP
Командный режим	Режим конфигурации интерфейса



Руководство по использованию	По умолчанию группе IPv6 VRRP в состоянии Master не разрешено получать пакеты с IPv6-адресом назначения, как у группы VRRP. Однако он получает пакеты NA и NS независимо от того, настроен ли режим приема. Кроме того, владелец IP-адреса в состоянии Master получает и обрабатывает пакеты с целевым IPv6-адресом как с адресом группы VRRP, независимо от того, настроен режим Асепт или нет
------------------------------	---

Настройка приоритета маршрутизатора IPv6 VRRP

Команда	vrrp ipv6 group priority level
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>level</i> : указывает приоритет маршрутизатора VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эта команда используется для ручной настройки приоритета маршрутизатора VRRP

Настройка отслеживаемого интерфейса IPv6 VRRP

Команда	vrrp ipv6 group track interface-type interface-number [priority]
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>interface-type interface-number</i> : указывает интерфейс для отслеживания. <i>priority</i> : указывает масштаб изменения приоритета VRRP при изменении состояния контролируемого интерфейса. Значение по умолчанию — 10
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Отслеживаемый интерфейс должен быть маршрутизируемым логическим интерфейсом уровня 3 (например, Routed (маршрутизируемый) порт, интерфейс SVI, Loopback-интерфейс или Tunnel (туннельный) интерфейс. Приоритет маршрутизатора, которому принадлежит виртуальный IP-адрес, связанный с группой VRRP, должен быть 255, и на нем нельзя настроить отслеживаемый интерфейс



Настройка отслеживаемого IP-адреса IPv6 VRRP

Команда	vrrp ipv6 group track { <i>ipv6-global-address</i> <i>ipv6-linklocal-address</i> <i>interface-type</i> <i>interface-number</i> } [interval <i>interval-value</i>] [timeout <i>timeout-value</i>] [retry <i>retry-value</i>] [<i>priority</i>]
Описание параметров	<i>group</i>: указывает VRID группы VRRP. <i>ipv6-global-address</i>: указывает глобальный индивидуальный адрес IPv6. <i>ipv6-linklocal-address</i>: указывает локальный адрес канала IPv6. <i>interface-type interface-number</i>: указывает интерфейс для отслеживания. interval interval-value: указывает интервал проверки. Единица секунды. Если не настроено вручную, значение по умолчанию равно 3 секундам. timeout timeout-value: указывает тайм-аут датчика ожидания ответов. Если по истечении времени ожидания ответ не получен, считается, что место назначения недоступно. Единица секунды. Если не настроено вручную, значение по умолчанию равно 1 секунде. retry retry-value: указывает на повторные попытки датчика. Если тестовый пакет отправляется непрерывно в течение времени, равного значению повторной попытки, но ответа не получено, считается, что пункт назначения недоступен. Единица факт попытки. Если не настроено, значение по умолчанию равно 3 раза. <i>priority</i>: указывает масштаб изменения приоритета VRRP при изменении состояния контролируемого интерфейса. Значение по умолчанию — 10
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Для мониторинга хоста укажите его IPv6-адрес для группы IPv6 VRRP. Если отслеживаемый IP-адрес узла является локальным адресом канала, укажите сетевой интерфейс. Если группа VRRP владеет фактическим IP-адресом интерфейса Ethernet, приоритет группы равен 255, а отслеживаемый IP-адрес не может быть настроен

Настраивает таймер обучения IPv6 VRRP

Команда	vrrp ipv6 group timers learn
Описание параметров	<i>group</i> : указывает VRID группы VRRP
Командный режим	Режим конфигурации интерфейса



Руководство по использованию	Как только таймер обучения включен на маршрутизаторе VRRP, Backup-маршрутизатор узнает интервал объявления Master-устройства в течение таймера. На основании этого Backup-маршрутизатор рассчитывает интервал для определения отказа Master-маршрутизатора вместо использования локально настроенного интервала объявления. Эта команда обеспечивает синхронизацию с таймером обучения между Master и Backup-маршрутизаторами
------------------------------	---

Настройка описания группы IPv6 VRRP

Команда	vrrp ipv6 group description text
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>text</i> : указывает описание группы VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Описание VRRP помогает различать группы VRRP. Описание имеет не более 80 байт, в противном случае запрашивается неправильная конфигурация

Настройка задержки IPv4 VRRP

Команда	vrrp delay { minimum min-seconds reload reload-seconds }
Описание параметров	minimum min-seconds : указывает задержку VRRP после изменения состояния интерфейса. reload reload-seconds : указывает задержку VRRP после запуска системы
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После того, как задержка настроена для группы VRRP на интерфейсе, группа VRRP запускается после задержки, а не сразу после запуска системы или возобновления работы интерфейса, что обеспечивает отсутствие preemption. Если интерфейс получает пакет VRRP во время задержки, задержка будет отменена, и VRRP будет запущен немедленно. Два типа задержки имеют общий диапазон значений от 0 до 60 секунд. Эта конфигурация будет эффективна как для групп VRRP IPv4, так и для групп IPv6 интерфейса



4.4.2.6. Пример конфигурации

Настройка группы IPv6 VRRP и отслеживаемого интерфейса

Сценарий:

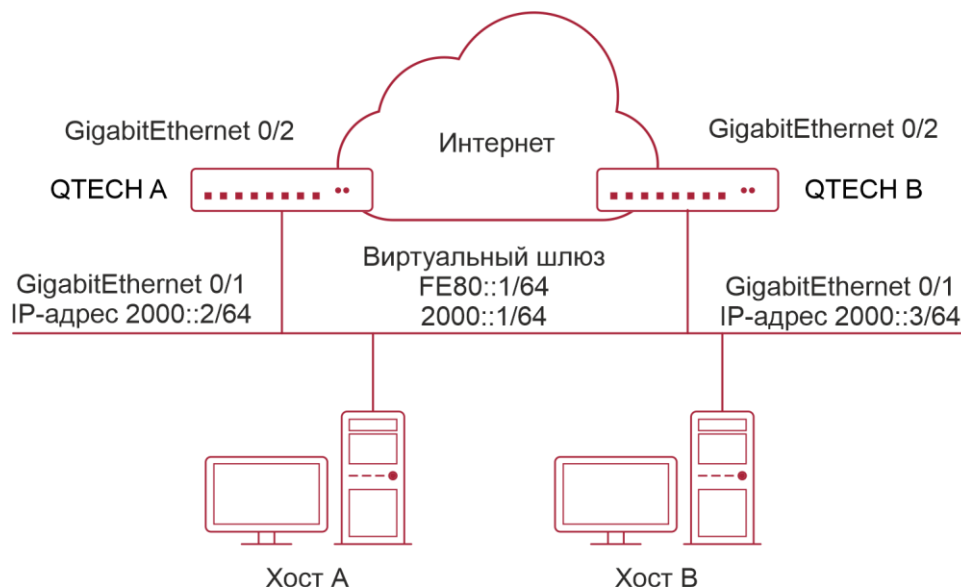


Рисунок 4-6.

Шаги настройки	- Узел А и узел В получают доступ к Интернет-ресурсам через шлюз по умолчанию 2000::1/64. - QTECH А и QTECH В принадлежат к группе IPv6 VRRP 1, и их виртуальные адреса 2000::1/64 и FE80::1 соответственно. - QTECH А отслеживает интерфейс GigabitEthernet 0/2, подключенный к Интернету. Когда GigabitEthernet 0/2 недоступен, QTECH А снижает свой приоритет, а QTECH В действует как шлюз
QTECH А	<pre> QTECH A#configure terminal QTECH A(config)#interface GigabitEthernet 0/1 QTECH A(config-if-GigabitEthernet 0/1)#no switchport QTECH A(config-if-GigabitEthernet 0/1)#ipv6 address 2000::2/64 QTECH A(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 FE80::1 QTECH A(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 2000::1 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 priority 120 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 timers advertise 3 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 track GigabitEthernet 0/2 50 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 accept_mode </pre>
QTECH В	<pre> QTECH B#configure terminal </pre>



	<pre> QTECH B(config)#interface GigabitEthernet 0/1 QTECH B(config-if-GigabitEthernet 0/1)#no switchport QTECH B(config-if-GigabitEthernet 0/1)#ipv6 address 2000::3/64 QTECH B(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 FE80::1 QTECH B(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 2000::1 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 priority 100 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 timers advertise 3 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 accept_mode </pre>
Проверка	Запустите команду show vrrp, чтобы проверить конфигурацию. - Проверьте, снижает ли QTECH A, выступающий в роли Master-маршрутизатора, приоритет группы VRRP со 120 до 70, когда обнаруживает, что интерфейс GigabitEthernet 0/2, подключенный к WAN, недоступен. Если да, QTECH B становится Master. - Проверьте, увеличивает ли QTECH A приоритет группы VRRP с 50 до 120, когда обнаруживает, что интерфейс GigabitEthernet 0/2, подключенный к WAN, снова становится доступным. Если да, QTECH A снова становится Master
QTECH A	<pre> QTECH A#show ipv6 vrrp 1 GigabitEthernet 0/1 - Group 1 State is Master Virtual IPv6 address is as follows: FE80::1 2000::1 Virtual MAC address is 0000.5e00.0201 Advertisement interval is 3 sec Accept_Mode is enabled Preemption is enabled min delay is 0 sec Priority is 120 Master Router is FE80::1234 (local), priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.59 sec Tracking state of 1 interface, 1 up: up GigabitEthernet 0/2 priority decrement=50 </pre>
QTECH B	<pre> QTECH B#show ipv6 vrrp 1 GigabitEthernet 0/1 - Group 1 </pre>



	State is Backup Virtual IPv6 address is as follow: FE80::1 2000::1 Virtual MAC address is 0000.5e00.0201 Advertisement interval is 3 sec Accept_Mode is enabled Preemption is enabled min delay is 0 sec Priority is 100 Master Router is FE80::1234, priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.82 sec
--	---

4.4.2.7. Пример конфигурации

Несколько резервных групп VRRP (под IPv6)

Сценарий:

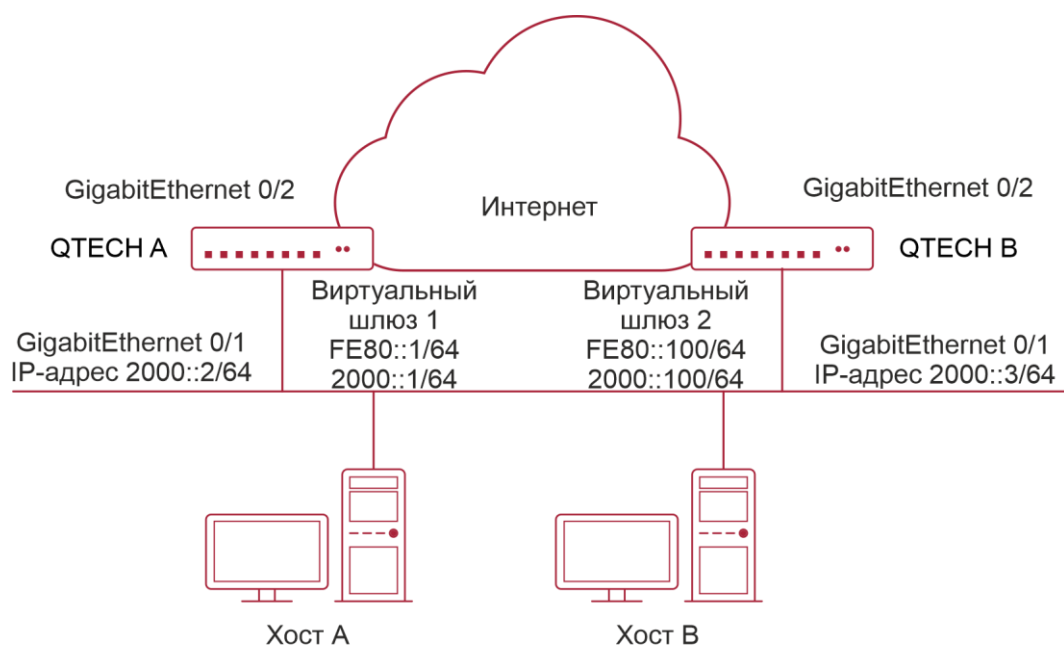


Рисунок 4-7.



Шаги настройки	• Хост А и Хост В получают доступ к Интернет-ресурсам через шлюзы 2000::1/64 и 2000::100/64 соответственно. • QTECH А и QTECH В принадлежат к группе IPv6 VRRP 1, и их виртуальные адреса 2000::1/64 и FE80::1 соответственно. • QTECH А и QTECH В принадлежат к резервной группе 2 виртуального маршрутизатора IPv6, и их виртуальные адреса — 2000::100/64 и FE80::100 соответственно. • QTECH А и QTECH В действуют как шлюзы и прямые потоки, являясь резервным маршрутизатором друг для друга
QTECH А	<pre> QTECH A#configure terminal QTECH A(config)#interface GigabitEthernet 0/1 QTECH A(config-if-GigabitEthernet 0/1)#no switchport QTECH A(config-if-GigabitEthernet 0/1)#ipv6 address 2000::2/64 QTECH A(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 FE80::1 QTECH A(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 2000::1 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 priority 120 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 timers advertise 3 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 accept_mode QTECH A(config-if-GigabitEthernet 0/1)#vrrp 2 ipv6 FE80::100 QTECH A(config-if-GigabitEthernet 0/1)# vrrp 2 ipv6 2000::100 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 2 priority 100 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 2 timers advertise 3 QTECH A(config-if-GigabitEthernet 0/1)#vrrp ipv6 2 accept_mode </pre>
QTECH В	<pre> QTECH B#configure terminal QTECH B(config)#interface GigabitEthernet 0/1 QTECH B(config-if-GigabitEthernet 0/1)#no switchport QTECH B(config-if-GigabitEthernet 0/1)#ipv6 address 2000::3/64 QTECH B(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 FE80::1 QTECH B(config-if-GigabitEthernet 0/1)#vrrp 1 ipv6 2000::1 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 priority 100 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 timers advertise 3 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 1 accept_mode QTECH B(config-if-GigabitEthernet 0/1)#vrrp 2 ipv6 FE80::100 QTECH B(config-if-GigabitEthernet 0/1)# vrrp 2 ipv6 2000::100 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 2 priority 120 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 2 timers advertise 3 QTECH B(config-if-GigabitEthernet 0/1)#vrrp ipv6 2 accept_mode </pre>



Проверка	Запустите команду show vrrp , чтобы проверить конфигурацию
QTECH A	<pre> QTECH A#show ipv6 vrrp GigabitEthernet 0/1 - Group 1 State is Master Virtual IPv6 address is as follows: FE80::1 2000::1 Virtual MAC address is 0000.5e00.0201 Advertisement interval is 3 sec Accept_Mode is enabled Preemption is enabled min delay is 0 sec Priority is 120 Master Router is FE80::1234 (local), priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.59 sec GigabitEthernet 0/1 - Group 2 State is Backup Virtual IPv6 address is as follows: FE80::100 2000::100 Virtual MAC address is 0000.5e00.0202 Advertisement interval is 3 sec Accept_Mode is enabled Preemption is enabled min delay is 0 sec Priority is 100 Master Router is FE80::5678, priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.82 sec </pre>
QTECH B	<pre> QTECH B#show ipv6 vrrp GigabitEthernet 0/1 - Group 1 State is Backup Virtual IPv6 address is as follow: FE80::1 </pre>



	2000::1 Virtual MAC address is 0000.5e00.0201 Advertisement interval is 3 sec Accept_Mode is enabled Preemption is enabled min delay is 0 sec Priority is 100 Master Router is FE80::1234, priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.82 sec GigabitEthernet 0/1 - Group 2 State is Master Virtual IPv6 address is as follows: FE80::100 2000::100 Virtual MAC address is 0000.5e00.0202 Advertisement interval is 3 sec Accept_Mode is enabled Preemption is enabled min delay is 0 sec Priority is 120 Master Router is FE80::5678(local), priority is 120 Master Advertisement interval is 3 sec Master Down interval is 10.59 sec
--	---

4.4.3. Настройка VRRP-MSTP

4.4.3.1. Эффект конфигурации

При одновременном применении MSTP и VRRP достигается резервное копирование на уровне канала и шлюза, а надежность сети значительно повышается.

4.4.3.2. Примечания

- Настройте маршрутизаторы в резервной группе VRRP с одним и тем же виртуальным IPv4-адресом.
- Включен VRRP на интерфейсе уровня 3.

4.4.3.3. Шаги настройки

Включение IPv4 VRRP

По умолчанию IPv4 VRRP не включен на интерфейсе. Чтобы включить IPv4 VRRP, настройте этот элемент.



Настройка строки аутентификации IPv4 VRRP

По умолчанию VRRP находится в режиме без аутентификации. Чтобы включить аутентификацию с помощью простого текстового пароля для VRRP, настройте этот элемент.

Настройка интервала объявления IPv4 VRRP

По умолчанию Master-маршрутизатор отправляет пакеты VRRP GWADV на интерфейсе в одну секунду. Чтобы вручную установить значение, настройте этот элемент.

Настройка preemption mode IPv4 VRRP

По умолчанию группы VRRP работают в preemption mode с нулевой задержкой.

Настройка приоритета маршрутизатора IPv4 VRRP

Приоритет маршрутизатора по умолчанию для группы VRRP равен 100. Вы можете изменить приоритет в зависимости от ваших потребностей.

Настройка отслеживаемого интерфейса IPv4 VRRP

По умолчанию группа IPv4 VRRP не отслеживает интерфейс. Чтобы обеспечить мониторинг отказов посредством мониторинга интерфейса, настройте этот элемент.

Настройка таймера обучения IPv4 VRRP

По умолчанию синхронизированное обучение не включено для Backup-группы VRRP. Чтобы Backup-маршрутизаторы могли получать пакеты VRRP GWADV от Master-маршрутизатора, настройте этот элемент.

Настройка описания группы IPv4 VRRP

По умолчанию для группы VRRP не настроено описание. Чтобы было удобно различать группы VRRP, настройте этот пункт.

Настройка задержки IPv4 VRRP

По умолчанию задержка VRRP для группы VRRP не настроена. Настройте задержку, чтобы гарантировать стабильный переход из режима без preemption в preemption mode.

Настройка версии IPv4 VRRP

По умолчанию для пакетов IPv4 VRRP используется стандарт VRRPv2. Чтобы изменить его вручную, настройте этот элемент.

Указание Sub VLAN Super VLAN для получения пакетов IPv4 VRRP

По умолчанию пакеты IPv4 VRRP отправляются только на первый интерфейс UP Sub VLAN в Super VLAN, но вы можете настроить определенный интерфейс Sub VLAN для отправки таких пакетов.

Настройка поддержки BFD для IPv4 VRRP на интерфейсе

По умолчанию связь между IPv4 VRRP и BFD не настроена на интерфейсе. Чтобы включить такую связь, настройте этот элемент.

Настройка глобального IPv4 VRRP BFD

По умолчанию глобальный IPv4 VRRP BFD не используется для определения того, активен ли Master-маршрутизатор. Чтобы включить это, настройте этот элемент.

4.4.3.4. Проверка

Запустите команду **show vrrp**, чтобы проверить конфигурацию.



4.4.3.5. Связанные команды

Включение IPv4 VRRP

Команда	vrrp group ip <i>ipaddress</i> [secondary]
Описание параметров	<i>group</i> : указывает VRID группы VRRP, диапазон которых зависит от модели продукта. <i>ipaddress</i> : IP-адрес группы VRRP. secondary : указывает вторичный IP-адрес VRRP-группы
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если виртуальный IP-адрес не указан, маршрутизаторы не могут присоединиться к группе VRRP. Если дополнительный IP-адрес не применяется, настроенный IP-адрес будет основным IP-адресом группы VRRP

Настройка строки аутентификации IPv4 VRRP

Команда	vrrp group authentication <i>string</i>
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>string</i> : указывает строку аутентификации группы VRRP (пароль в виде простого текста состоит не более чем из 8 байтов)
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	В группе VRRP для маршрутизаторов должен быть настроен один и тот же пароль аутентификации. Простой текстовый пароль аутентификации не может гарантировать безопасность, а только предотвращает/подсказывает неправильные конфигурации VRRP. Эта команда применима только к VRRPv2 вместо VRRPv3. Аутентификация отменена для пакетов VRRPv3. Если для группы IPv4 VRRP выбран VRRPv2, команда действует; если выбран VRRPv3, команда неэффективна

Настройка интервала объявления IPv4 VRRP

Команда	vrrp group timers advertise { <i>advertise-interval</i> csec <i>centisecond-interval</i> }
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>advertise-interval</i> : указывает интервал объявления группы VRRP (единица измерения: секунда).



	csec <i>centisecond-interval</i> : интервал, через который Master-маршрутизатор в резервной группе отправляет пакеты VRRP. Это целое число от 50 до 99. Единицей измерения является сантисекунда. Значение по умолчанию не указано. Команда эффективна только для пакетов VRRPv3. Если он настроен для пакетов VRRPv2, интервал по умолчанию равен одной секунде
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если маршрутизатор выбран в качестве Master в группе VRRP, он отправляет пакеты объявления VRRP с заданным интервалом, чтобы объявить о своем состоянии VRRP, приоритете и другой информации. Согласно стандартам RFC, если группа IPv4 VRRP использует VRRPv3 для отправки многоадресных пакетов, максимальный интервал объявления составляет 40 секунд. Таким образом, если задан интервал более 40 секунд, будет применяться этот максимальный интервал, несмотря на то, что конфигурация действует

Настройка preemption mode IPv4 VRRP

Команда	vrrp group preempt [<i>delay seconds</i>]
Описание параметров	<i>group</i> : указывает VRID группы VRRP. delay seconds : указывает задержку preemption для Master-маршрутизатора, чтобы заявить о своем статусе. Значение по умолчанию — 0 секунд
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Если группа VRRP работает в preemption mode, маршрутизатор с более высоким приоритетом займет место Master, у которого более низкий приоритет. Если группа VRRP работает в режиме без preemption, маршрутизатор с более высоким приоритетом, чем у Master, остается Backup. Нет большого смысла настраивать режим Preemption, когда группа VRRP использует IP-адрес Ethernet-интерфейса, в этом случае группа имеет наивысший приоритет и автоматически становится Master в группе

Настройка приоритета маршрутизатора IPv4 VRRP

Команда	vrrp group priority level
Описание параметров	<i>group</i> : указывает VRID группы VRRP. <i>level</i> : указывает приоритет интерфейса в группе VRRP



Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эта команда используется для ручной настройки приоритета группы VRRP

Настройка отслеживаемого интерфейса IPv4 VRRP

Команда	vrrp group track { <i>interface-type interface-number</i> bfd <i>interface-type interface-number ipv4-address</i> } [<i>priority</i>]
Описание параметров	<i>group</i>: указывает VRID группы VRRP. <i>interface-type interface-number</i>: указывает интерфейс для отслеживания. bfd <i>interface-type interface-number ipv4-address</i>: указанный соседний IP-адрес, отслеживаемый через BFD. <i>priority</i>: указывает масштаб изменения приоритета VRRP при изменении состояния контролируемого интерфейса. Значение по умолчанию — 10
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Отслеживаемый интерфейс должен быть маршрутизируемым логическим интерфейсом уровня 3 (например, маршрутизируемый (Routed) порт, интерфейс SVI, интерфейс Loopback или туннельный (Tunnel) интерфейс). Приоритет маршрутизатора, которому принадлежит виртуальный IP-адрес, связанный с группой VRRP, должен быть 255, и на нем нельзя настроить отслеживаемый интерфейс

Настройка отслеживаемого IP-адреса IPv4 VRRP

Команда	vrrp group track ipv4-address [interval <i>interval-value</i>] [timeout <i>timeout-value</i>] [retry <i>retry-value</i>] [<i>priority</i>]
Описание параметров	<i>group</i>: указывает VRID группы VRRP. <i>ipv4-address</i>: указывает отслеживаемый IPv4-адрес. interval <i>interval-value</i>: указывает интервал проверки. Единица секунды. Если не настроено вручную, значение по умолчанию равно 3 секундам. timeout <i>timeout-value</i>: указывает тайм-аут датчика ожидания ответов. Если по истечении времени ожидания ответ не получен, считается, что место назначения недоступно. Единица секунды. Если не настроено вручную, значение по умолчанию равно 1 секунде. retry <i>retry-value</i>: показывает повторные попытки датчика. Если тестовый пакет отправляется непрерывно в течение времени, равного



	значению повторной попытки, но ответа не получено, считается, что пункт назначения недоступен. Единица факт попытки. Если не настроено, значение по умолчанию равно 3 раза. <i>priority</i>: указывает масштаб изменения приоритета VRRP при изменении состояния контролируемого интерфейса. Значение по умолчанию — 10
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Для мониторинга хоста укажите его IPv4-адрес для группы IPv4 VRRP. Если группа VRRP владеет фактическим IP-адресом интерфейса Ethernet, приоритет группы равен 255, а отслеживаемый IP-адрес не может быть настроен

Настройка таймера обучения IPv4 VRRP

Команда	<code>vrrp group timers learn</code>
Описание параметров	<i>group</i> : указывает VRID группы VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Как только таймер обучения включен на маршрутизаторе VRRP, Backup-маршрутизатор узнает интервал объявления Master-устройства в течение таймера. На основании этого Backup-маршрутизатор рассчитывает интервал для определения отказа Master-маршрутизатора вместо использования локально настроенного интервала объявления. Эта команда обеспечивает синхронизацию с таймером обучения между Master и Backup-маршрутизаторами

Настройка описания группы IPv4 VRRP

Команда	<code>vrrp group description text</code>
Описание параметров	<i>group</i>: указывает VRID группы VRRP. <i>text</i>: указывает описание группы VRRP
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Описание VRRP помогает различать группы VRRP. А описание имеет не более 80 байт, в противном случае запрашивается неправильная конфигурация



Настройка задержки IPv4 VRRP

Команда	vrrp delay { minimum <i>min-seconds</i> reload <i>reload-seconds</i> }
Описание параметров	minimum <i>min-seconds</i> : указывает задержку VRRP после изменения состояния интерфейса. reload <i>reload-seconds</i> : указывает задержку VRRP после запуска системы
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После того, как задержка настроена для группы VRRP на интерфейсе, группа VRRP запускается после задержки, а не сразу после запуска системы или возобновления работы интерфейса, что обеспечивает отсутствие preemption. Если интерфейс получает пакет VRRP во время задержки, задержка будет отменена, и VRRP будет запущен немедленно. Два типа задержки имеют общий диапазон значений от 0 до 60 секунд. Эта конфигурация будет эффективна как для групп VRRP IPv4, так и для групп IPv6 интерфейса

Настройка версии IPv4 VRRP

Команда	vrrp group version { 2 3 }
Описание параметров	2 : указывает на VRRPv2. 3 : указывает на VRRPv3
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Учитывая совместимость между VRRPv2 и VRRPv3, укажите стандарт для IPv4 VRRP на основе фактического состояния сети. VRRPv2 разработан в RFC3768, а VRRPv3 описан в RFC5798. Эта команда применима только к IPv4 VRRP

Указание Sub VLAN Super VLAN для получения пакетов IPv4 VRRP

Команда	vrrp detection-vlan { first-subvlan <i>subvlan-id</i> }
Описание параметров	first-subvlan : отправляет пакеты IPv4 VRRP только на первый интерфейс UP Sub VLAN в Super VLAN. <i>subvlan-id</i> : отправляет пакеты IPv4 VRRP в указанную Sub VLAN
Командный режим	Режим конфигурации интерфейса



Руководство по использованию	Эта команда используется для указания Sub VLAN Super VLAN для получения пакетов IPv4 VRRP. Пакеты IPv4 VRRP отправляются в Super VLAN с использованием следующих трех методов. Пакеты отправляются на первый интерфейс UP Sub VLAN в Super VLAN, или на указанный интерфейс Sub VLAN в Super VLAN, или на все интерфейсы Sub VLAN в Super VLAN. Если на интерфейсе Super VLAN включены и VRRP, и VRRP PLUS, пакеты VRRP отправляются на все интерфейсы UP Sub VLAN интерфейса Super VLAN. Эта команда настраивается на интерфейсе VLAN и эффективна только для интерфейсов Super VLAN
------------------------------	---

Настройка поддержки BFD для IPv4 VRRP на интерфейсе

Команда	vrrp group bfd ip-address
Описание параметров	<i>group</i>: указывает VRID группы VRRP. <i>ip-address</i>: указывает IP-адрес интерфейса
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Для Backup-маршрутизатора запустите эту команду, чтобы сопоставить группу IPv4 VRRP с BFD, не заботясь о настроенном IP-адресе. Для Master, поскольку первичный IP-адрес Backup-маршрутизатора неизвестен, IP-адрес маршрутизатора может быть указан только администратором. Если настроен глобальный IPv4 VRRP BFD, эта конфигурация не может быть выполнена. Чтобы включить поддержку BFD, убедитесь, что параметры сеанса IP и BFD настроены на целевом интерфейсе

Настройка глобального IPv4 VRRP BFD

Команда	vrrp bfd interface-type interface-number ip-address
Описание параметров	<i>interface-type interface-number</i>: указывает тип интерфейса и идентификатор. <i>ip-address</i>: указывает IP-адрес интерфейса
Командный режим	Режим глобальной конфигурации



Руководство по использованию	Если настроен глобальный IPv4 VRRP BFD, настроенная поддержка BFD будет удалена. Чтобы включить поддержку BFD, убедитесь, что параметры сеанса IP и BFD настроены на целевом интерфейсе. Глобальный сеанс IPv4 VRRP BFD применим только к группе IPv4 VRRP, состоящей из двух маршрутизаторов
------------------------------	--

4.4.3.6. Пример конфигурации

Настройка VRRP+MSTP

Сценарий:

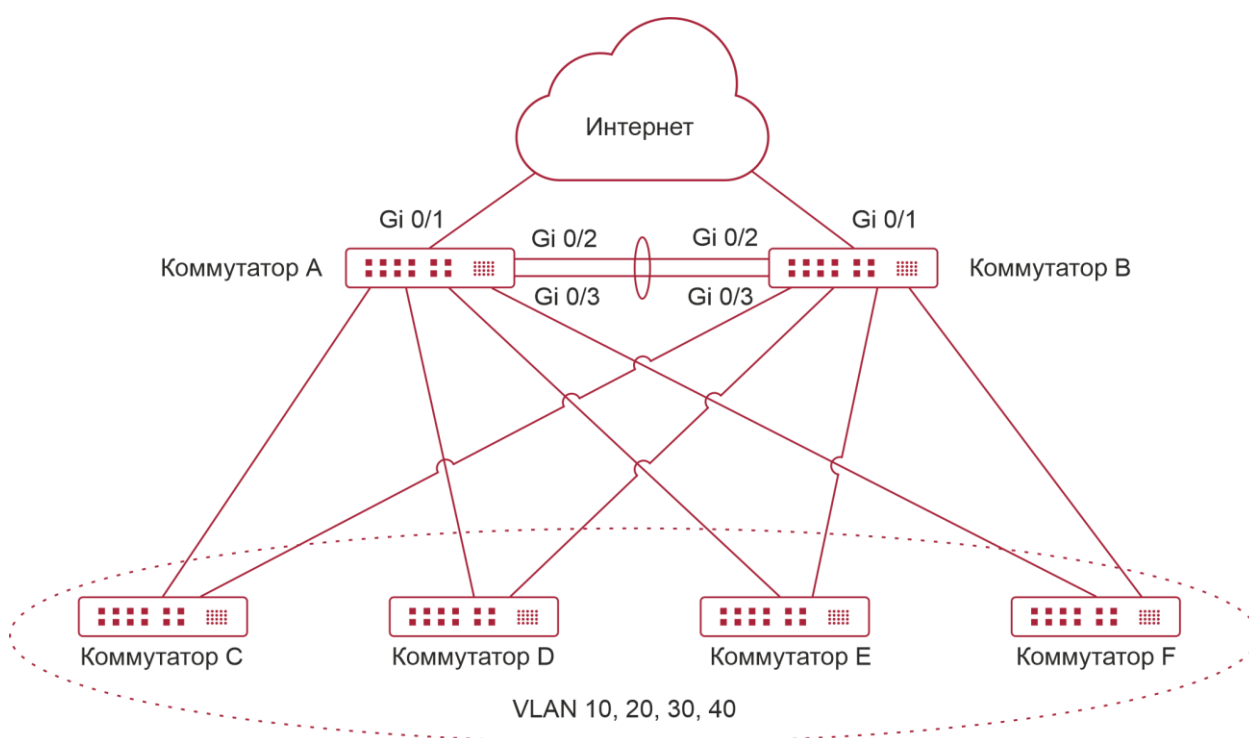


Рисунок 4-8.

Шаги настройки	Включите MSTP на маршрутизаторах (в данном примере это коммутаторы A, B, C, D, E и F). Настройте сопоставление VLAN-Instance (сопоставление VLAN 10 и VLAN 20 с экземпляром 1, VLAN 30 и VLAN 40 с экземпляром 2, а остальные VLAN с экземпляром 0) и настройте шлюзы (коммутатор A и коммутатор B в этом примере) в качестве root bridge соответствующих экземпляров. Добавьте SVI всех VLAN в соответствующие Backup-группы VRRP и настройте шлюзы в качестве главного и резервного маршрутизаторов для соответствующих групп резервного копирования. См. сведения о конфигурации в следующей таблице.
----------------	--



	Шлюз	VLAN ID	SVI	Backup-группа	Виртуальный IP-адрес	Состояние
	Switch A	10	192.168.10.2	VRRP 10	192.168.10.1	Master
	Switch B		192.168.10.3			Backup
	Switch A	20	192.168.20.2	VRRP 20	192.168.20.1	Master
	Switch B		192.168.20.3			Backup
	Switch A	30	192.168.30.2	VRRP 30	192.168.30.1	Backup
	Switch B		192.168.30.3			Master
	Switch A	40	192.168.40.2	VRRP 40	192.168.40.1	Backup
	Switch B		192.168.40.3			Master
	Настройте восходящий порт (порт Gi 0/1 коммутатора А и коммутатора В) Master-маршрутизаторов в качестве контролируемого интерфейса Master-маршрутизатора. Шаг 1. Создайте VLAN. Создайте VLAN 10, VLAN 20, VLAN 30 и VLAN 40 соответственно на коммутаторе А и коммутаторе В. Шаг 2. Настройте регионы MST. Сопоставьте VLAN 10 и VLAN 20 с экземпляром 1, VLAN 30 и VLAN 40 с экземпляром 2, а остальные VLAN с экземпляром 0. Шаг 3. Настройте коммутатор А в качестве root bridge для MST 0 и MST 1, а коммутатор В — в качестве root bridge для MST 2. Шаг 4. Включите MSTP. Шаг 5. Настройте SVI всех VLAN, добавьте SVI в соответствующие резервные группы и настройте виртуальные IP-адреса для групп. См. конфигурацию в таблице выше. Шаг 6. Настройте Master-маршрутизаторы и Backup-маршрутизаторы для всех групп. Шаг 7. Настройте uplink-порты Master-маршрутизаторов в качестве контролируемых портов групп VRRP. Внимание: Контролируемые порты должны быть портами уровня 3. Шаг 8. Настройте интернет-интерфейсы основных маршрутизаторов как интерфейсы AP					
Switch A	// Создать VLAN 10, VLAN 20, VLAN 30 и VLAN 40 на коммутаторе А. SwitchA#configure terminal Введите команды конфигурации, по одной в строке. Конец с CNTL/Z.					



```

SwitchA(config)#vlan range 10,20,30,40
SwitchA(config-vlan-range)#exit
// Сопоставьте VLAN 10 и VLAN 20 с экземпляром 1, VLAN 30 и VLAN 40 с
экземпляром 2, а остальные VLAN с экземпляром 0.
SwitchA(config)#spanning-tree mst configuration
SwitchA(config-mst)#instance 1 vlan 10,20
%Предупреждение: вы должны создать VLANы перед настройкой
отношения экземпляра-VLAN
SwitchA(config-mst)#instance 2 vlan 30,40
%Предупреждение: вы должны создать VLANы перед настройкой
отношения экземпляра-VLAN
SwitchA(config-mst)#exit
// На коммутаторе А настройте приоритет MST 0 и MST 1 на 4096, а
приоритет MST 2 на 8192.
SwitchA(config)#spanning-tree mst 0 priority 4096
SwitchA(config)#spanning-tree mst 1 priority 4096
SwitchA(config)#spanning-tree mst 2 priority 8192
// Включение MSTP
SwitchA(config)#spanning-tree
Enable spanning-tree.
// Настройте SVI всех VLAN, добавьте SVI в соответствующие резервные
группы и настройте виртуальные IP-адреса для групп.
SwitchA(config)#interface vlan 10
SwitchA(config-if-VLAN 10)#ip address 192.168.10.2 255.255.255.0
SwitchA(config-if-VLAN 10)#vrrp 10 ip 192.168.10.1
SwitchA(config-if-VLAN 10)#exit
SwitchA(config)#interface vlan 20
SwitchA(config-if-VLAN 20)#ip address 192.168.20.2 255.255.255.0
SwitchA(config-if-VLAN 20)#vrrp 20 ip 192.168.20.1
SwitchA(config-if-VLAN 20)#exit
SwitchA(config)#interface vlan 30
SwitchA(config-if-VLAN 30)#ip address 192.168.30.2 255.255.255.0
SwitchA(config-if-VLAN 30)#vrrp 30 ip 192.168.30.1
SwitchA(config-if-VLAN 30)#exit
SwitchA(config)#interface vlan 40
SwitchA(config-if-VLAN 40)#ip address 192.168.40.2 255.255.255.0
SwitchA(config-if-VLAN 40)#vrrp 40 ip 192.168.40.1
SwitchA(config-if-VLAN 40)#exit

```




	<pre>// Увеличьте приоритет VRRP 10 и VRRP 20 коммутатора А до 120. SwitchA(config)#interface vlan 10 SwitchA(config-if-VLAN 10)#vrrp 10 priority 120 SwitchA(config-if-VLAN 10)#exit SwitchA(config)#interface vlan 20 SwitchA(config-if-VLAN 20)#vrrp 20 priority 120 SwitchA(config-if-VLAN 20)#exit // Настройте порт Gi 0/1 коммутатора А как порт Route и его IP-адрес как 10.10.1.1/24. SwitchA(config)#interface gigabitEthernet 0/1 SwitchA(config-if-GigabitEthernet 0/1)#no switchport SwitchA(config-if-GigabitEthernet 0/1)#ip address 10.10.1.1 255.255.255.0 SwitchA(config-if-GigabitEthernet 0/1)#exit // Настройте порт Gi 0/1 коммутатора А в качестве контролируемого порта для VRRP 10 и VRRP 20 и уменьшите приоритет до 30. SwitchA(config)#interface vlan 10 SwitchA(config-if-VLAN 10)#vrrp 10 track gigabitEthernet 0/1 30 SwitchA(config-if-VLAN 10)#exit SwitchA(config)#interface vlan 20 SwitchA(config-if-VLAN 20)#vrrp 20 track gigabitEthernet 0/1 30 SwitchA(config-if-VLAN 20)#exit // Настройте порты Gi 0/2 и Gi 0/3 как порты AP, которые являются транковыми портами. SwitchA#configure terminal Введите команды конфигурации, по одной в строке. Конец с CNTL/Z. SwitchA(config)#interface range gigabitEthernet 0/2-3 SwitchA(config-if-range)#port-group 1 SwitchA(config)#interface aggregateport 1 SwitchA(config-if-AggregatePort 1)#switchport mode trunk</pre>
Switch B	<pre>// Создайте VLAN 10, VLAN 20, VLAN 30 и VLAN 40 на коммутаторе В. SwitchB#configure terminal Введите команды конфигурации, по одной в строке. Конец с CNTL/Z. SwitchB(config)#vlan range 10,20,30,40 SwitchB(config-vlan-range)#exit // Сопоставьте VLAN 10 и VLAN 20 с экземпляром 1, VLAN 30 и VLAN 40 с экземпляром 2, а остальные VLAN с экземпляром 0. SwitchB(config)#spanning-tree mst configuration</pre>



```

SwitchB(config-mst)#instance 1 vlan 10,20
%Предупреждение: вы должны создать VLANы перед настройкой
отношения экземпляра-VLAN

SwitchB(config-mst)#instance 2 vlan 30,40
%Предупреждение: вы должны создать VLANы перед настройкой
отношения экземпляра- VLAN

SwitchB(config-mst)#exit
// На коммутаторе В настройте приоритет MST 2 на 4096, а приоритет MST
0 и MST 1 на 8192.

SwitchB(config)#spanning-tree mst 2 priority 4096
SwitchB(config)#spanning-tree mst 0 priority 8192
SwitchB(config)#spanning-tree mst 1 priority 8192
// Включение MSTP
SwitchB(config)#spanning-tree
Enable spanning-tree.
// Настройте SVI всех VLAN, добавьте SVI в соответствующие резервные
группы и настройте виртуальные IP-адреса для групп.

SwitchB(config)#interface vlan 10
SwitchB(config-if-VLAN 10)#ip address 192.168.10.3 255.255.255.0
SwitchB(config-if-VLAN 10)#vrrp 10 ip 192.168.10.1
SwitchB(config-if-VLAN 10)#exit
SwitchB(config)#interface vlan 20
SwitchB(config-if-VLAN 20)#ip address 192.168.20.3 255.255.255.0
SwitchB(config-if-VLAN 20)#vrrp 20 ip 192.168.20.1
SwitchB(config-if-VLAN 20)#exit
SwitchB(config)#interface vlan 30
SwitchB(config-if-VLAN 30)#ip address 192.168.30.3 255.255.255.0
SwitchB(config-if-VLAN 30)#vrrp 30 ip 192.168.30.1
SwitchB(config-if-VLAN 30)#exit
SwitchB(config)#interface vlan 40
SwitchB(config-if-VLAN 40)#ip address 192.168.40.3 255.255.255.0
SwitchB(config-if-VLAN 40)#vrrp 40 ip 192.168.40.1
SwitchB(config-if-VLAN 40)#exit
// Увеличьте приоритет VRRP 30 и VRRP 40 коммутатора В до 120.
SwitchB(config)#interface vlan 30
SwitchB(config-if-VLAN 30)#vrrp 30 priority 120
SwitchB(config-if-VLAN 30)#exit
SwitchB(config)#interface vlan 40

```



	<pre>SwitchB(config-if-VLAN 40)#vrrp 40 priority 120 SwitchB(config-if-VLAN 40)#exit</pre> // Настройте порт Gi 0/1 коммутатора В как порт маршрута и его IP-адрес как 10.10.1.1/24. <pre>SwitchB(config)#interface gigabitEthernet 0/1 SwitchB(config-if-GigabitEthernet 0/1)#no switchport SwitchB(config-if-GigabitEthernet 0/1)#ip address 10.10.2.1 255.255.255.0 SwitchB(config-if-GigabitEthernet 0/1)#exit</pre> // Настройте порт Gi 0/1 коммутатора В в качестве контролируемого порта для VRRP 30 и VRRP 40, а приоритет интерфейса — 30. <pre>SwitchB(config)#interface vlan 30 SwitchB(config-if-VLAN 30)#vrrp 30 track gigabitEthernet 0/1 30 SwitchB(config-if-VLAN 30)#exit SwitchB(config)#interface vlan 40 SwitchB(config-if-VLAN 40)#vrrp 40 track gigabitEthernet 0/1 30 SwitchB(config-if-VLAN 40)#exit</pre> // Настройте порты Gi 0/2 и Gi 0/3 как порты AP, которые являются транковыми портами. <pre>SwitchB #configure terminal</pre> Введите команды конфигурации, по одной в строке. Конец с CNTL/Z. <pre>SwitchB (config)#interface range gigabitEthernet 0/2-3 SwitchB (config-if-range)#port-group 1 SwitchB (config)#interface aggregateport 1 SwitchB (config-if-AggregatePort 1)#switchport mode trunk</pre>
Switch A	Проверьте конфигурацию. <pre>SwitchA#show running-config</pre> <pre>! vlan 10 ! vlan 20 ! vlan 30 ! vlan 40</pre>



```
!  
spanning-tree  
spanning-tree mst configuration  
instance 0 vlan 1-9, 11-19, 21-29, 31-39, 41-4094  
instance 1 vlan 10, 20  
instance 2 vlan 30, 40  
spanning-tree mst 0 priority 4096  
spanning-tree mst 1 priority 4096  
spanning-tree mst 2 priority 8192  
interface GigabitEthernet 0/1  
no switchport  
no ip proxy-arp  
ip address 10.10.1.1 255.255.255.0  
!  
interface GigabitEthernet 0/2  
port-group 1  
!  
interface GigabitEthernet 0/3  
port-group 1  
!  
interface AggregatePort 1  
switchport mode trunk  
!  
interface VLAN 10  
no ip proxy-arp  
ip address 192.168.10.2 255.255.255.0  
vrrp 10 priority 120  
vrrp 10 ip 192.168.10.1  
vrrp 10 track GigabitEthernet 0/1 30  
!  
interface VLAN 20  
no ip proxy-arp  
ip address 192.168.20.2 255.255.255.0  
vrrp 20 priority 120  
vrrp 20 ip 192.168.20.1  
vrrp 20 track GigabitEthernet 0/1 30
```



	<pre> ! interface VLAN 30 no ip proxy-arp ip address 192.168.30.2 255.255.255.0 vrrp 30 ip 192.168.30.1 ! interface VLAN 40 no ip proxy-arp ip address 192.168.40.2 255.255.255.0 vrrp 40 ip 192.168.40.1 // Проверить статус VRRP. SwitchA#show vrrp brief Interface Grp Pri timer Own Pre State Master addr Group addr VLAN 10 10 120 3 - P Master 192.168.10.2 192.168.10.1 VLAN 20 20 120 3 - P Master 192.168.20.2 192.168.20.1 VLAN 30 30 100 3 - P Backup 192.168.30.3 192.168.30.1 VLAN 40 40 100 3 - P Backup 192.168.40.3 192.168.40.1 // Отключите upstream-канал коммутатора А и проверьте статус VRRP. SwitchA#show vrrp brief Interface Grp Pri timer Own Pre State Master addr Group addr VLAN 10 10 90 3 - P Backup 192.168.10.3 192.168.10.1 VLAN 20 20 90 3 - P Backup 192.168.20.3 192.168.20.1 VLAN 30 30 100 3 - P Backup 192.168.30.3 192.168.30.1 VLAN 40 40 100 3 - P Backup 192.168.40.3 192.168.40.1 </pre>
Switch B	<pre> // Проверить конфигурацию. SwitchB#show running-config ! vlan 10 ! vlan 20 ! vlan 30 ! vlan 40 ! spanning-tree </pre>



```

spanning-tree mst configuration
instance 0 vlan 1-9, 11-19, 21-29, 31-39, 41-4094
instance 1 vlan 10, 20
instance 2 vlan 30, 40
spanning-tree mst 0 priority 8192
spanning-tree mst 1 priority 8192
spanning-tree mst 2 priority 4096
interface GigabitEthernet 0/1
no switchport
no ip proxy-arp
ip address 10.10.2.1 255.255.255.0
!
interface GigabitEthernet 0/2
port-group 1!
interface GigabitEthernet 0/3
port-group 1
!
interface AggregatePort 1
switchport mode trunk
!
interface VLAN 10
no ip proxy-arp
ip address 192.168.10.3 255.255.255.0
vrrp 10 ip 192.168.10.1
!
interface VLAN 20
no ip proxy-arp
ip address 192.168.20.3 255.255.255.0
vrrp 20 ip 192.168.20.1
!
interface VLAN 30
no ip proxy-arp
ip address 192.168.30.3 255.255.255.0
vrrp 30 priority 120
vrrp 30 ip 192.168.30.1
vrrp 30 track GigabitEthernet 0/1 30

```



```

!
interface VLAN 40
no ip proxy-arp
ip address 192.168.40.3 255.255.255.0
vrrp 40 priority 120
vrrp 40 ip 192.168.40.1
vrrp 40 track GigabitEthernet 0/1 30
// Проверить статус VRRP.
SwitchB#show vrrp brief
Interface Grp Pri timer Own Pre State Master addr Group addr
VLAN 10 10 100 3 - P Backup 192.168.10.2 192.168.10.1
VLAN 20 20 100 3 - P Backup 192.168.20.2 192.168.20.1
VLAN 30 30 120 3 - P Master 192.168.30.3 192.168.30.1
VLAN 40 40 120 3 - P Master 192.168.40.3 192.168.40.1
// Отключите upstream-канал коммутатора В и проверьте статус VRRP.
SwitchB#show vrrp brief
Interface Grp Pri timer Own Pre State Master addr Group addr
VLAN 10 10 100 3 - P Master 192.168.10.3 192.168.10.1
VLAN 20 20 100 3 - P Master 192.168.20.3 192.168.20.1
VLAN 30 30 120 3 - P Master 192.168.30.3 192.168.30.1
VLAN 40 40 120 3 - P Master 192.168.40.3 192.168.40.1

```

4.5. Мониторинг

4.5.1. Отображение

Описание	Команда
Отображает краткую или подробную информацию о IPv4/IPv6 VRRP	show [ipv6] vrrp [brief group]
Отображает информацию о группе IPv4/IPv6 VRRP на указанном интерфейсе	show [ipv6] vrrp interface type number [brief]
Отображает статистику пакетов VRRP	show vrrp packet statistics [interface-type interface-number]



4.5.1.1. Отладка

ПРИМЕЧАНИЕ: системные ресурсы заняты при выводе отладочной информации. Поэтому отключайте отладку сразу после использования.

Описание	Команда
Отладка ошибок, событий, пакетов и статуса VRRP	debug [ipv6] vrrp
Отладка ошибок VRRP	debug [ipv6] vrrp errors
Отладка событий VRRP	debug [ipv6] vrrp events
Отладка пакетов VRRP	debug vrrp packets [acl acl-id [icmp protocol] interface type number [group]] debug ipv6 vrrp packets [acl acl-name [icmp protocol] interface type number [group]]
Отладка статусов VRRP	debug [ipv6] vrrp state



5. НАСТРОЙКА VRRP PLUS

5.1. Обзор

Virtual Router Redundancy Protocol Plus (VRRP Plus) является расширением VRRP. Он использует VRRP для реализации резервного копирования шлюза и балансировки нагрузки в локальной сети IEEE 802.3 (LAN).

Недостатком VRRP является то, что маршрутизатор в состоянии Backup не может пересылать пакеты. Чтобы использовать VRRP для реализации балансировки нагрузки, вам необходимо вручную настроить несколько групп VRRP и установить адреса шлюзов хостов в локальной сети на виртуальные IP-адреса разных групп VRRP. Это увеличивает нагрузку на администратора сети. VRRP Plus предназначен для решения этой проблемы.

В VRRP Plus автоматически реализуется балансировка нагрузки. То есть трафик разных хостов автоматически распределяется между участниками группы VRRP Plus, и нет необходимости настраивать несколько групп VRRP или устанавливать адреса шлюзов хостов в LAN на виртуальные IP-адреса разных групп VRRP. Это значительно снижает нагрузку на администратора сети.

5.2. Приложения

Приложение	Описание
<u>Включение балансировки нагрузки в группе VRRP</u>	Реализуйте балансировку нагрузки в группе VRRP без настройки нескольких групп или настройки разных шлюзов по умолчанию для хостов

5.2.1. Включение балансировки нагрузки в группе VRRP

5.2.1.1. Сценарий

Включите балансировку нагрузки в группе VRRP без настройки нескольких групп VRRP или настройки различных шлюзов по умолчанию для хостов.

Как показано на Рисунке 5-1, настройте данные следующим образом:

- Настройте группу VRRP, состоящую из маршрутизатора А и маршрутизатора В, и включите функцию VRRP Plus.
- Настройте шлюз по умолчанию для каждого хоста в качестве Master виртуального IP-адреса группы VRRP.

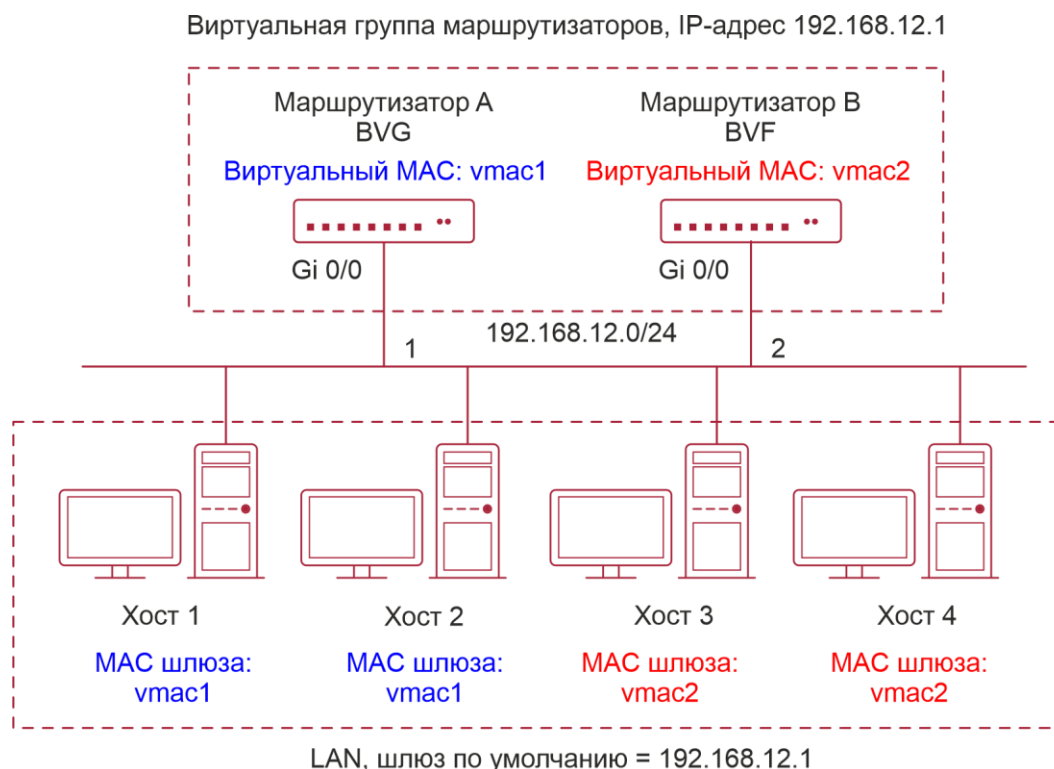


Рисунок 5-1. Топология приложения IPv4 VRRP Plus

1. Два устройства уровня 3 (L3), маршрутизатор A и маршрутизатор B, образуют группу VRRP Plus, а виртуальный IP-адрес группы — 192.168.12.1. Маршрутизатор A является Master-устройством VRRP и функционирует как балансировочный виртуальный шлюз (BVG). Маршрутизатор B является Ваксир-устройством VRRP и функционирует как балансировочный виртуальный сервер пересылки (BVF).
2. Хосты с 1 по 4 — это хосты в локальной сети с сетевым сегментом 192.168.12.0/24. Их адреса шлюза по умолчанию установлены на виртуальный IP-адрес 192.168.12.1 группы VRRP Plus.
3. Политика балансировки нагрузки настроена на устройстве для ответа на запросы ARP, отправленные с разных хостов. Например, когда Хост 1 и Хост 2 запрашивают ARP шлюза, MAC-адрес 0000.5e00.0101 возвращается на Хост 1 и Хост 2. Когда Хост 3 и Хост 4 запрашивают ARP шлюза, MAC-адрес 08c6.b316.0201 вернулся на Хост 3 и Хост 4. Таким образом, пакеты, которыми обмениваются хост 1/хост 2 и внешняя сеть, отправляются на маршрутизатор A, а пакеты, которыми обмениваются между хостом 3/хостом 4 и внешней сетью, отправляются на маршрутизатор B, тем самым реализуя балансировку нагрузки.

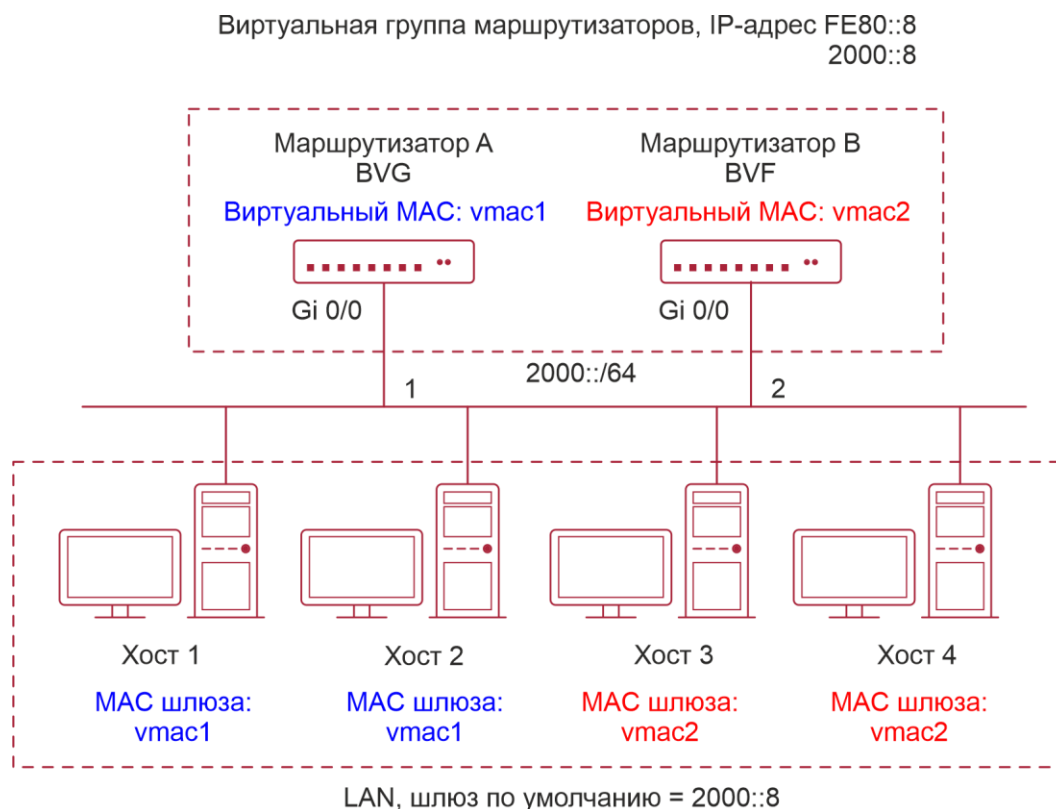


Рисунок 5-2. Топология приложения IPv6 VRRP Plus

1. Два устройства L3, маршрутизатор A и маршрутизатор B, образуют группу IPv6 VRRP Plus, а виртуальные IPv6-адреса — fe80::8 и 2000::8. Маршрутизатор A является Master-устройством IPv6 VRRP и функционирует как BVG. Маршрутизатор B является Backup-устройством IPv6 VRRP и функционирует как BVF.
2. Хосты с 1 по 4 являются хостами в локальной сети, и для всех их адресов шлюзов IPv6 установлено значение 2000::8.
3. Когда хост 1 и хост 2 запрашивают обнаружение соседей (ND) шлюза, MAC-адрес 0000.5e00.0201 возвращается на хост 1 и хост 2. Когда хост 3 и хост 4 запрашивают ND шлюза, MAC-адрес 08C6.B372.7701 возвращается на хост 3 и хост 4. Таким образом, пакеты, которыми обмениваются хост 1/хост 2 и внешняя сеть, отправляются на маршрутизатор A, а пакеты, которыми обмениваются между хостом 3/хостом 4 и внешней сетью, отправляются на маршрутизатор B, тем самым реализуя балансировку нагрузки.

5.2.1.2. Развертывание

Разверните VRRP Plus на маршрутизаторе A и маршрутизаторе B, чтобы реализовать балансировку нагрузки на локальном хосте.



5.3. Функции

5.3.1. Базовые понятия

BVG

BVG выделяет виртуальные MAC-адреса участникам группы VRRP Plus. Он отвечает на запросы шлюза ARP/ND в локальной сети и пересылает пакеты хостов в локальной сети.

BVF

BVF пересылает пакеты хостов в локальной сети. Если виртуальный MAC-адрес выделен в BVF, BVF участвует в пересылке пакетов; в противном случае BVF не участвует в пересылке пакетов.

5.3.2. Обзор

Особенность	Описание
<u>VRRP Plus</u>	Расширьте возможности VRRP и используйте VRRP для реализации резервного копирования шлюза и балансировки нагрузки в локальной сети IEEE 802.3

5.3.3. VRRP Plus

В VRRP Plus автоматически реализуется балансировка нагрузки. То есть трафик разных хостов автоматически распределяется между участниками группы VRRP Plus, и нет необходимости настраивать несколько групп VRRP или устанавливать адреса шлюзов хостов в LAN на виртуальные IPv4-адреса разных групп VRRP.

5.3.3.1. Основные принципы

Хосты в локальной сети используют IPv4/IPv6-адрес унифицированного шлюза (то есть виртуальный IP-адрес группы VRRP). Когда разные хосты запрашивают ARP шлюза, BVG отвечает разными виртуальными MAC-адресами. Таким образом, трафик разных хостов распределяется по разным участникам группы VRRP Plus, тем самым реализуя балансировку нагрузки.

Связь между VRRP Plus и VRRP

VRRP Plus опирается на VRRP и работает следующим образом:

Master-устройство в VRRP соответствует BVG в VRRP Plus, а Backup-устройство в VRRP соответствует BVF в VRRP Plus. Адреса шлюзов хостов в локальной сети устанавливаются на виртуальный IPv4-адрес VRRP.

Правила распределения MAC-адресов BVG и BVF

BVG выделяет виртуальные MAC-адреса для BVF. Для группы IPv4 VRRP Plus BVG напрямую использует виртуальный MAC-адрес VRRP для обеспечения совместимости между IPv4 VRRP Plus и VRRP. То есть виртуальный MAC-адрес, используемый BVG, — 00-00-5E-00-01-{VRID}, где VRID — номер группы VRRP. Виртуальный MAC-адрес, используемый BVF, — 08-C6-B3-16-{MemberID}-{VRID}, где MemberID — это идентификатор участника BVF в группе VRRP Plus. В настоящее время в группе VRRP Plus может быть до четырех участников. BVG использует идентификатор участника 01, а другие BVF используют идентификаторы участников от 02 до 04. Для группы IPv6 VRRP Plus BVG напрямую использует виртуальный MAC-адрес IPv6 VRRP для обеспечения



совместимости между IPv6 VRRP Plus и VRRP. То есть виртуальный MAC-адрес, используемый BVG, — 00-00-5E-00-02-{VRID}, где VRID — это номер группы IPv6 VRRP.

Политика балансировки нагрузки VRRP Plus

BVG отвечает на запросы шлюза ARP/NS, отправленные с хостов в локальной сети. В зависимости от конкретной политики балансировки нагрузки BVG отвечает узлам с разными виртуальными MAC-адресами. Существует три типа политик балансировки нагрузки:

- Host-dependent policy (Политика, зависящая от хоста): указанный виртуальный MAC-адрес используется для ответа на запросы, отправленные указанным хостом.
- Round-robin policy (Политика циклического перебора): виртуальные MAC-адреса в резервной группе используются циклически для ответа на запросы шлюза ARP/NS, отправляемые хостами.
- Weighted policy (Взвешенная политика): ответы на запросы ARP/NA предоставляются в зависимости от возможностей пересылки каждого устройства.

При изменении режима балансировки нагрузки балансировка нагрузки всегда выполняется в новом режиме балансировки нагрузки. Например, если ранее использовался режим опроса ответа, а позже используется взвешенный режим, то балансировка нагрузки реализуется в взвешенном режиме вне зависимости от более ранних ответов устройства. Если используется взвешенная политика и общий вес виртуальных маршрутизаторов в группе VRRP Plus равен 0, запросы ARP/NS не отвечают.

Прокси-виртуального MAC-адреса

Когда устройство с виртуальным MAC-адресом выходит из строя в Backup-группе, трафик хостов, использующих этот виртуальный MAC-адрес в качестве MAC-адреса шлюза, будет прерван.

BVG в Backup-группе VRRP Plus может быстро обнаружить неисправность и автоматически выделить виртуальный MAC-адрес неисправного BVF другому устройству в Backup-группе. Новое устройство действует как прокси неисправного устройства для пересылки пакетов виртуального MAC-адреса. Кроме того, это прокси-устройство берет на себя трафик исходных хостов, чтобы предотвратить прерывание трафика. Виртуальный MAC-адрес, назначенный устройству в Backup-группе, может называться Master виртуальным MAC-адресом, а виртуальный MAC-адрес, используемый этим устройством от имени другого устройства, называется прокси-виртуальным MAC-адресом.

Время перенаправления и время ожидания прокси-виртуального MAC-адреса

VRRP Plus обеспечивает функцию прокси для виртуального MAC-адреса, чтобы другое устройство могло заменить неисправное устройство с виртуальным MAC-адресом для пересылки пакетов. Если BVF восстанавливается после сбоя, его роль пересылки восстанавливается, и BVF продолжает пересылать пакеты виртуального MAC-адреса, выделенного для этого BVF. Если неисправный BVF не восстановлен, Backup-группа прекращает перенаправление трафика на этот виртуальный MAC-адрес. То есть при повторном получении запросов ARP этот виртуальный MAC-адрес больше не отвечает. По прошествии достаточно длительного периода времени считается, что хосты, которые используют MAC-адрес в качестве MAC-адреса шлюза, уже обновляют запись таблицы ARP/ND с адресом шлюза, и трафик уже перехватывается другими устройствами. В это время этот виртуальный MAC-адрес может быть удален, и пакеты, отправленные на этот виртуальный MAC-адрес, отбрасываются.

VRRP Plus поддерживает настройку времени перенаправления и времени ожидания Backup-группы. Когда устройство неисправно, Backup-группа выделяет виртуальный

MAC-адрес неисправного устройства другому устройству. В течение времени перенаправления Backup-группа продолжает использовать этот виртуальный MAC-адрес для ответа на запросы ARP/NS. По истечении времени перенаправления Backup-группа больше не использует этот виртуальный MAC-адрес для ответа на запросы. По истечении тайм-аута Backup-группа удаляет этот виртуальный MAC-адрес и прекращает использовать этот виртуальный MAC-адрес для переадресации прокси-сервера. На Рисунке 5-3 показаны изменения роли виртуального MAC-адреса в зависимости от времени перенаправления и времени ожидания.

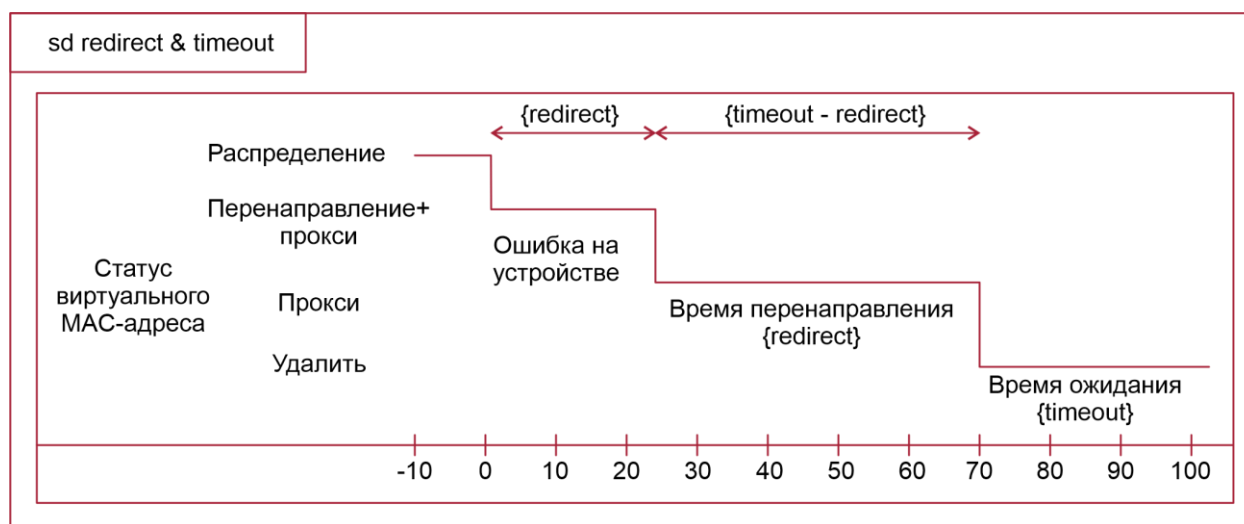


Рисунок 5-3. Изменения роли виртуального MAC-адреса в зависимости от времени перенаправления и времени ожидания

Переадресация на основе веса

VRRP Plus поддерживает конфигурацию веса Backup-группы. Для разных устройств настроены разные веса. Таким образом, больше трафика распределяется на устройство с большим весом, а меньше трафика распределяется на устройство с меньшим весом, тем самым полностью используя производительность пересылки различных устройств. Когда вес BVF в Backup-группе меньше нижнего порога, BVF автоматически выходит из роли пересылки. Когда вес восстанавливается и превышает верхний порог, BVF автоматически применяет роль пересылки. Роль пересылки может быть восстановлена, когда существует один или несколько оставшихся виртуальных MAC-адресов или прокси-виртуальных MAC-адресов.

Ассоциация VRRP Plus с BFD

VRRP Plus поддерживает ассоциацию с обнаружением двунаправленной пересылки (BFD) для настройки веса на основе состояния канала. Каждое устройство в Backup-группе может связать свой вес со статусом связи. Когда связь ненормальна или прервана, устройство автоматически уменьшает свой вес. Когда вес слишком мал, устройство автоматически выходит из роли пересылки. Если Backup-группа в настоящее время использует взвешенную политику балансировки нагрузки, трафик можно распределить на основе нового веса. Когда соответствующий канал восстанавливается, устройство может автоматически восстановить свой первоначальный вес и роль пересылки. Если Backup-группа в настоящее время использует взвешенную политику балансировки нагрузки, трафик можно распределять на основе восстановленного веса. IPv6 VRRP Plus в настоящее время не поддерживает ассоциацию с BFD.



Захват (Seizure) переадресации на основе веса

VRRP Plus поддерживает функцию захвата роли пересылки. В VRRP Plus в балансировке нагрузки могут участвовать не более четырех устройств. То есть Ваксир-группа VRRP Plus создает не более четырех виртуальных MAC-адресов. Если в группу VRRP Plus добавлено более четырех устройств, только четыре устройства участвуют в пересылке пакетов. Остальные устройства только прослушивают состояние других устройств и не участвуют в пересылке пакетов. Только когда устройство, участвующее в пересылке пакетов, неисправно, другое устройство, которое изначально не участвовало в пересылке пакетов, займет место неисправного устройства для пересылки пакетов. Предположим, что в Ваксир-группе VRRP Plus уже есть четыре устройства, и все эти устройства участвуют в пересылке пакетов; пятое устройство добавлено в группу VRRP Plus, и возможности пересылки этого устройства сильны, или исходная роль пересылки сталкивается с отказом канала и, следовательно, ухудшением производительности пересылки. В этом случае, если включен режим захвата, пятое устройство может захватить роль пересылки у устройства с меньшим весом (то есть с меньшими возможностями пересылки). Большой вес настраивается для устройства с более сильными возможностями переадресации. Когда вес устройства в состоянии прослушивания оказывается больше веса пересылающего устройства, устройство в состоянии прослушивания автоматически захватывает роль пересылки у пересылающего устройства. То есть устройство с более сильными возможностями пересылки пересылает пакеты, тогда как устройство с более низкими возможностями пересылки находится в состоянии прослушивания. Это может свести к минимуму трату ресурсов.

BVG в Ваксир-группе отвечает за выделение виртуальных MAC-адресов. Таким образом, роль BVG не может быть занята, и может быть занята только роль пересылки BVF. Если устройство BVG неисправно, VRRP повторно выбирает новое Master-устройство, которое берет на себя роль BVG.

Факторы, влияющие на политику пересылки

1. После настройки VRRP Plus на запросы ARP/NS, полученные от хостов, можно отвечать на основе различных политик балансировки нагрузки для реализации балансировки нагрузки между этими хостами. Однако балансировка нагрузки не может быть реализована для хостов, которые узнали адреса виртуальных шлюзов VRRP до настройки VRRP Plus. Следовательно, если VRRP Plus настроен после изменения состояния VRRP на Master, реальная балансировка нагрузки не может быть реализована до устаревания ARP/ND, полученных хостами. Балансировка нагрузки реализуется только после того, как ARP/ND шлюза записываются по возрасту хостов, и хосты запрашивают новые адреса шлюза.
2. Периодическая отправка gratuitous ARP на интерфейс также влияет на функцию балансировки нагрузки VRRP Plus. Когда VRRP Plus включен, функция отправки gratuitous ARP виртуальных IP-адресов VRRP будет отключена. Когда виртуальный IP-адрес перекрывается с фактическим IP-адресом, gratuitous ARP этого адреса больше не отправляются.
3. Когда между хостом и локальным устройством возникает конфликт адресов, модуль ARP/NA будет передавать пакеты gratuitous ARP/NA с этим адресом. Если возникает конфликт виртуального адреса VRRP Plus, отправка пакета gratuitous ARP/NA приведет к повторному изучению MAC-адреса шлюза хоста, что негативно повлияет на функцию балансировки нагрузки VRRP Plus. Поэтому функция балансировки нагрузки VRRP Plus в настоящее время не поддерживается в этом сценарии.



5.4. Конфигурация

Элемент конфигурации	Описание и команда	
<u>Настройка VRRP Plus</u>	(Обязательно) Он используется для включения функции VRRP Plus	
	vrrp balance	Включает функцию VRRP Plus Backup-группы VRRP с указанным идентификатором группы в режиме конфигурации интерфейса
	(Опционально) Он используется для настройки параметров Backup-группы VRRP Plus	
	vrrp load-balancing	Настраивает политику балансировки нагрузки VRRP Plus в режиме конфигурации интерфейса
	vrrp timers redirect	Настраивает время перенаправления и время ожидания прокси-виртуального MAC-адреса в Backup-группе VRRP Plus в режиме конфигурации интерфейса
	vrrp weighting	Настраивает вес, а также верхний и нижний пороги Backup-группы VRRP Plus в режиме конфигурации интерфейса
	vrrp forwarder preempt	Настраивает функцию захвата переадресации Backup-группы VRRP Plus в режиме конфигурации интерфейса
	vrrp weighting track	Настраивает объект отслеживания веса Backup-группы VRRP Plus в режиме настройки интерфейса

5.4.1. Настройка VRRP Plus

5.4.1.1. Эффект конфигурации

- Включите функцию VRRP Plus. (По умолчанию эта функция отключена)
- Настройте объект отслеживания веса Backup-группы VRRP Plus.

5.4.1.2. Примечания

Чтобы включить функцию VRRP Plus, необходимо настроить виртуальный IP-адрес VRRP для соответствующей Backup-группы.



5.4.1.3. Шаги настройки

Включение VRRP Plus на интерфейсе

По умолчанию VRRP Plus включен. Выполните эту настройку, если требуется VRRP Plus.

Настройка политики балансировки нагрузки VRRP Plus

После включения VRRP Plus по умолчанию используется циклическая (round-robin) политика балансировки нагрузки.

Настройка времени перенаправления и времени ожидания прокси-виртуального MAC-адреса в Backup-группе VRRP Plus

После включения VRRP Plus время перенаправления устанавливается на 300 с, а тайм-аут по умолчанию — на 14 400 с.

Настройка веса, а также верхнего и нижнего порогов Backup-группы VRRP Plus

После включения VRRP Plus вес Backup-группы устанавливается равным 100, нижний порог равен 1, а верхний порог равен 100 по умолчанию.

Настройка функции захвата пересылки (Forwarding Seizure) Backup-группы VRRP Plus

После включения VRRP Plus функция захвата пересылки включена по умолчанию.

Настройка объекта отслеживания веса Backup-группы VRRP Plus

Объект отслеживания веса Backup-группы VRRP Plus по умолчанию отключен. Выполните эту настройку, если требуется функция отслеживания.

5.4.1.4. Проверка

Запустите команду **show [ipv6] group vrrp balance**, чтобы отобразить конфигурацию Backup-группы VRRP. Если у Backup-группы есть задачи пересылки пакетов, в столбце серверов пересылки отображается «local», а также отображается виртуальный MAC-адрес, назначенный этой Backup-группе.

5.4.1.5. Связанные команды

Включение VRRP Plus на интерфейсе

Команда	vrrp [ipv6] group balance
Описание параметров	ipv6 : указывает, что эта конфигурация применяется к IPv6. По умолчанию конфигурация применяется к IPv4. group : указывает идентификатор группы VRRP. Диапазон значений идентификатора группы зависит от модели продукта
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	VRRP Plus можно включить только после настройки группы VRRP



Настройка политики балансировки нагрузки Backup-группы VRRP Plus

Команда	<code>vrrp [ipv6] group load-balancing {host-dependent round-robin weighted }</code>
Описание параметров	ipv6: указывает, что эта конфигурация применяется к IPv6. По умолчанию конфигурация применяется к IPv4. group: указывает идентификатор группы VRRP. host-dependent: указывает политику балансировки нагрузки, зависящую от хоста. round-robin: указывает политику балансировки нагрузки round-robin (циклического перебора). weighted: указывает взвешенную политику балансировки нагрузки
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После включения VRRP Plus политика балансировки нагрузки, зависящая от хоста, используется по умолчанию. Политика балансировки нагрузки всей Backup-группы определяется политикой, настроенной на BVG. Если вы хотите использовать одну и ту же политику балансировки нагрузки после изменения роли устройства BVG, настройте одну и ту же политику на всех устройствах в Backup-группе

Настройка времени перенаправления и времени ожидания прокси-виртуального MAC-адреса в Backup-группе VRRP Plus

Команда	<code>vrrp [ipv6] group timers redirect redirect timeout</code>
Описание параметров	ipv6: указывает, что эта конфигурация применяется к IPv6. По умолчанию конфигурация применяется к IPv4. group: указывает идентификатор группы VRRP. redirect: указывает время перенаправления. Диапазон значений от 0 до 3600 с. Значение по умолчанию — 300 с, то есть 5 минут. timeout: указывает время ожидания. Значение варьируется от (перенаправление + 600) до 64 800 с. Значение по умолчанию — 14 400, то есть 4 часа
Командный режим	Режим конфигурации интерфейса



Руководство по использованию	После включения VRRP Plus время перенаправления устанавливается на 300 с, а тайм-аут по умолчанию — на 14 400 с. Когда устройство неисправно, Backup-группа выделяет виртуальный MAC-адрес неисправного устройства другому устройству. В течение времени перенаправления Backup-группа продолжает использовать этот виртуальный MAC-адрес для ответа на запросы ARP/NS. По истечении времени перенаправления Backup-группа больше не использует этот виртуальный MAC-адрес для ответа на запросы. По истечении времени ожидания Backup-группа удаляет этот виртуальный MAC-адрес
------------------------------	--

Настройка веса, а также верхнего и нижнего порогов Backup-группы VRRP Plus

Команда	<code>vrrp [ipv6] group weighting maximum [lower lower] [upper upper]</code>
Описание параметров	ipv6: указывает, что эта конфигурация применяется к IPv6. По умолчанию конфигурация применяется к IPv4. maximum: указывает вес Backup-группы. Значение находится в диапазоне от 2 до 254. Значение по умолчанию — 100. lower lower: указывает нижний порог Backup-группы. Значение варьируется от 1 до (максимум - 1). Значение по умолчанию — 1. upper upper: указывает верхний порог Backup-группы. Значение варьируется от нижнего до максимального. Значение по умолчанию — 100
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После включения VRRP Plus вес, а также верхний и нижний пороги Backup-группы VRRP Plus настраиваются по умолчанию. Вы можете использовать эту команду для настройки разных весов для разных устройств, чтобы больше трафика распределялось на устройство с большим весом и меньше трафика распределялось на устройство с меньшим весом. Когда вес BVF в Backup-группе ниже нижнего порога, BVF автоматически выходит из роли пересылки. Когда вес восстанавливается и становится выше верхнего порога, роль пересылки BVF автоматически восстанавливается

Настройка функции захвата пересылки (Forwarding Seizure) Backup-группы VRRP Plus

Команда	<code>vrrp [ipv6] group forwarder preempt</code>
Описание параметров	ipv6: указывает, что эта конфигурация применяется к IPv6. По умолчанию конфигурация применяется к IPv4. group: указывает идентификатор группы VRRP



Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После включения VRRP Plus функция захвата пересылки включена по умолчанию. VRRP Plus поддерживает настройку функции захвата пересылки Backup-группы. Когда вес устройства в состоянии прослушивания оказывается больше веса пересылающего устройства, устройство в состоянии прослушивания автоматически захватывает роль пересылки у пересылающего устройства. То есть устройство с более сильными возможностями пересылки пересылает пакеты, тогда как устройство с более низкими возможностями пересылки находится в состоянии прослушивания

Настройка объекта отслеживания веса Backup-группы VRRP Plus

Команда	<code>vrrp [ipv6] group weighting track object-number [decrement value]</code>
Описание параметров	ipv6: указывает, что эта конфигурация применяется к IPv6. По умолчанию конфигурация применяется к IPv4. group: указывает идентификатор группы VRRP. object-number: указывает номер отслеживаемого объекта. Диапазон значений от 1 до 700. decrement value: указывает вес уменьшения, когда объект отслеживания в нерабочем состоянии. Значение уменьшения находится в диапазоне от 1 до 255. Значение по умолчанию — 10
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	После включения VRRP Plus объекты отслеживания по умолчанию не настраиваются. После настройки объекта отслеживания значение уменьшения по умолчанию равно 10



5.4.1.6. Пример конфигурации

Включение балансировки нагрузки в группе IPv4 VRRP

Сценарий:

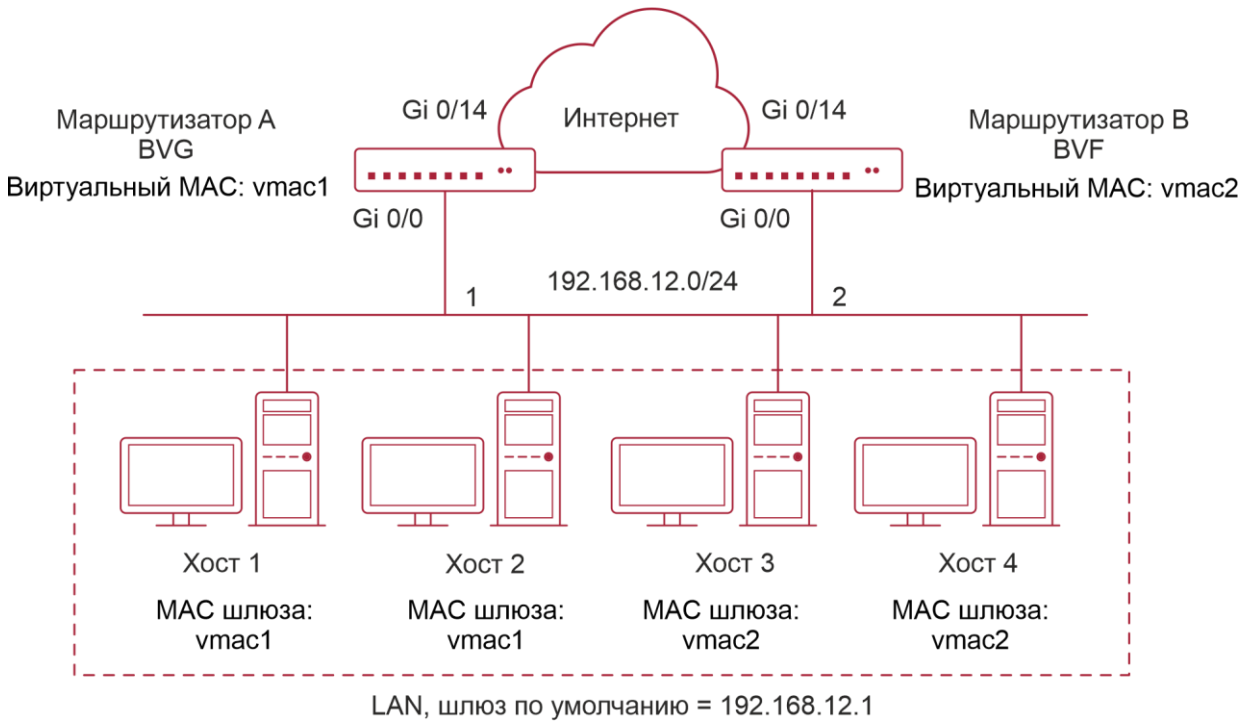


Рисунок 5-4.

Шаги настройки	• Настройте объект отслеживания соответственно на маршрутизаторе А и маршрутизаторе В для отслеживания состояния локального интерфейса GigabitEthernet 0/14. • Настройте группу VRRP и включите VRRP Plus соответственно на маршрутизаторе А и маршрутизаторе В. Настройте локальные IP-адреса так, чтобы маршрутизатор А стал устройством BVG (Master), а маршрутизатор В стал устройством BVF (Backup). • Настройте взвешенную политику балансировки нагрузки. Настройте объект отслеживания веса и установите значение уменьшения на 100. • Сохраните настройки по умолчанию для веса, верхнего и нижнего порогов, времени перенаправления, тайм-аута и захвата пересылки Backup-группы. • Установите адреса шлюзов по умолчанию для хостов 1 и 4 в локальной сети на виртуальный IP-адрес VRRP, то есть 192.168.12.1
Маршрутизатор А	<pre>QTECH A#config QTECH A(config)#track 1 interface GigabitEthernet0/14 line-protocol</pre>



	<pre> QTECH A(config)#interface GigabitEthernet0/0 // На коммутаторе используется "no switchport". QTECH A(config-if-GigabitEthernet 0/0)#no switchport QTECH A(config-if-GigabitEthernet 0/0)#ip address 192.168.12.3 255.255.255.0 QTECH A(config-if-GigabitEthernet 0/0)#vrrp 1 ip 192.168.12.1 QTECH A(config-if-GigabitEthernet 0/0)#vrrp 1 balance QTECH A(config-if-GigabitEthernet 0/0)#vrrp 1 load-balancing weighted QTECH A(config-if-GigabitEthernet 0/0)#vrrp 1 weighting track 1 decrement 100 </pre>
Маршрутизатор В	<pre> QTECH B#config QTECH B(config)#track 1 interface GigabitEthernet0/14 line-protocol QTECH B(config)#interface GigabitEthernet0/0 QTECH B(config-if-GigabitEthernet 0/0)#no switchport QTECH B(config-if-GigabitEthernet 0/0)#ip address 192.168.12.2 255.255.255.0 QTECH B(config-if-GigabitEthernet 0/0)#vrrp 1 ip 192.168.12.1 QTECH B(config-if-GigabitEthernet 0/0)#vrrp 1 balance QTECH B(config-if-GigabitEthernet 0/0)#vrrp 1 load-balancing weighted QTECH B(config-if-GigabitEthernet 0/0)#vrrp 1 weighting track 1 decrement 100 </pre>
Проверка	Запустите команду show vrrp balance, чтобы отобразить конфигурацию группы VRRP Plus. Если у Backup-группы есть задачи пересылки пакетов, в столбце серверов пересылки отображается «local», а также отображается виртуальный MAC-адрес, назначенный этой Backup-группе
Маршрутизатор А	<pre> QTECH A# show vrrp balance interface GigabitEthernet0/0 State is BVG Virtual IP address is 192.168.12.1 Hello time 1 sec, hold time 3 sec Load balancing: weighted Redirect time 300 sec, forwarder time-out 14400 sec Weighting 100 (configured 100), thresholds: lower 1, upper 100 Track object 1, state: up, decrement weight: 100 There are 2 forwarders Forwarder 1 (local) MAC address: 0000.5e00.0101 Owner ID is 0000.0001.0006 </pre>



	Preemption disabled (BVG cannot be preempted) Forwarder 2 MAC address: 08c6.b316.0201 Owner ID is 08c6.b322.33a3 Preemption enabled
Маршрутизатор В	QTECH B# show vrrp balance interface GigabitEthernet0/0 State is BVF Virtual IP address is 192.168.12.1 Hello time 1 sec, hold time 3 sec Load balancing: weighted Redirect time 300 sec, forwarder time-out 14400 sec Weighting 100 (configured 100), thresholds: lower 1, upper 100 Track object 1, state: up, decrement weight: 100 There are 2 forwarders Forwarder 1 MAC address: 0000.5e00.0101 Owner ID is 0000.0001.0006 Preemption disabled (BVG cannot be preempted) Forwarder 2 (local) MAC address: 08c6.b316.0201 Owner ID is 08c6.b322.33a3 Preemption enabled



Включение балансировки нагрузки в группе IPv6 VRRP

Сценарий:

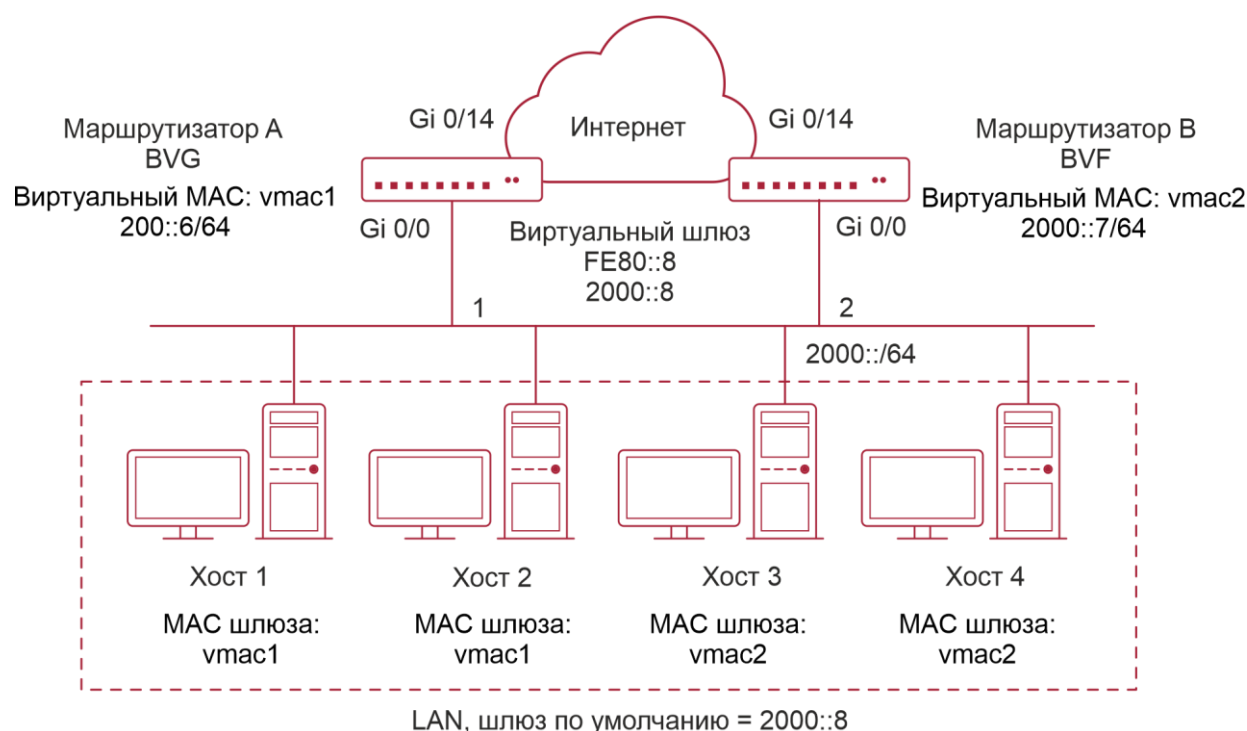


Рисунок 5-5.

Шаги настройки	- Настройте группу IPv6 VRRP и включите IPv6 VRRP Plus соответственно на маршрутизаторе А и маршрутизаторе В. Настройте приоритеты таким образом, чтобы маршрутизатор А стал устройством BVG (Master), а маршрутизатор В стал устройством BVF (Backup). - Настройте взвешенную политику балансировки нагрузки для Backup-группы IPv6 VRRP Plus. - Сохраните настройки по умолчанию для веса, верхнего и нижнего порогов, времени перенаправления, тайм-аута и захвата пересылки Backup-группы. - Установите адреса шлюзов по умолчанию для хостов с 1 по 4 в локальной сети на 2000::8
Маршрутизатор А	<pre> RouterA#config RouterA(config)#interface GigabitEthernet0/0 // На коммутаторе используется "no switchport". RouterA(config-if-GigabitEthernet 0/0)#no switchport RouterA(config-if-GigabitEthernet 0/0)#ipv6 address 2000::6/64 RouterA(config-if-GigabitEthernet 0/0)#vrrp 1 ipv6 fe80::8 </pre>



	<pre>RouterA(config-if-GigabitEthernet 0/0)#vrrp 1 ipv6 2000::8 RouterA(config-if-GigabitEthernet 0/0)#vrrp ipv6 1 priority 120 RouterA(config-if-GigabitEthernet 0/0)#vrrp ipv6 1 load-balancing weighted RouterA(config-if-GigabitEthernet 0/0)#vrrp ipv6 1 balance</pre>
Маршрутизатор В	<pre>RouterB#config RouterB(config)#interface GigabitEthernet0/0 RouterB(config-if-GigabitEthernet 0/0)#no switchport RouterB(config-if-GigabitEthernet 0/0)# ipv6 address 2000::7/64 RouterB(config-if-GigabitEthernet 0/0)#vrrp 1 ipv6 fe80::8 RouterB(config-if-GigabitEthernet 0/0)#vrrp 1 ipv6 2000::8 RouterB(config-if-GigabitEthernet 0/0)#vrrp ipv6 1 load-balancing weighted RouterB(config-if-GigabitEthernet 0/0)#vrrp ipv6 1 balance</pre>
Проверка	Запустите команду show ipv6 vrrp balance, чтобы отобразить конфигурацию группы VRRP Plus. Если у Backup-группы есть задачи пересылки пакетов, в столбце серверов пересылки отображается «local», а также отображается виртуальный MAC-адрес, назначенный этой Backup-группе
Маршрутизатор А	<pre>RouterA# show ipv6 vrrp balance interface GigabitEthernet0/0 GigabitEthernet 0/0 - Group 1 State is BVG Virtual IPv6 address is as follows: FE80::8 2000::8 Hello time 1 sec, hold time 3 sec Load balancing: weighted Redirect time 300 sec, forwarder time-out 14400 sec Weighting 100 (configured 100), thresholds: lower 1, upper 100 There are 2 forwarders Forwarder 1 (local) MAC address: 0000.5e00.0201 Owner ID is 08c6.b3fb.96f3 Preemption disabled (BVG cannot be preempted) Forwarder 2 MAC address:</pre>



	08C6.B372.7701 Owner ID is 08c6.b3fb.6c42 Preemption enabled
Маршрутизатор В	RouterB# show ipv6 vrrp balance interface GigabitEthernet0/0 GigabitEthernet 0/0 - Group 1 State is BVF Virtual IPv6 address is as follows: FE80::8 2000::8 Hello time 1 sec, hold time 3 sec Load balancing: weighted Redirect time 300 sec, forwarder time-out 14400 sec Weighting 150 (configured 150), thresholds: lower 1, upper 100 There are 2 forwarders Forwarder 1 MAC address: 0000.5e00.0201 Owner ID is 08c6.b3fb.96f3 Preemption disabled (BVG cannot be preempted) Forwarder 2 (local) MAC address: 08C6.B372.7701 Owner ID is 08c6.b3fb.6c42 Preemption enabled

5.4.1.7. Распространенные ошибки

VRRP Plus не действует, поскольку виртуальный IP-адрес VRRP не настроен для связанной группы.

5.5. Мониторинг

5.5.1. Отображение

Описание	Команда
Отображает краткую или подробную конфигурацию VRRP Plus	show [ipv6] vrrp balance



Описание	Команда
Отображает действия группы VRRP Plus на указанном интерфейсе	show [ipv6] vrrp balance interface

5.5.2. Отладка

ПРИМЕЧАНИЕ: системные ресурсы заняты при выводе отладочной информации. Поэтому отключайте отладку сразу после использования.

Описание	Команда
Отладка функции VRRP Plus	debug [ipv6] vrrp balance
Отладка ошибок	debug[ipv6] vrrp balance error
Отладка событий группы VRRP Plus	debug [ipv6] vrrp balance event
Отладка сообщений между модулем VRRP и модулем отслеживания	debug[ipv6] vrrp balance messages
Отладка пакетов VRRP Plus	debug[ipv6] vrrp balance packets
Отладка статуса группы VRRP Plus	debug [ipv6] vrrp balance state
Отладка таймеров группы VRRP Plus	debug vrrp balance timer



6. НАСТРОЙКА BFD

6.1. Обзор

Сбои связи прерывают работу сети и, таким образом, влияют на услуги. Поэтому очень важно быстро обнаруживать сбои связи на каналах с соседними устройствами, чтобы обеспечить своевременные действия и доступность услуг. Обнаружение двунаправленной пересылки (BFD) обеспечивает метод быстрого обнаружения возможности соединения пути пересылки между двумя соседними маршрутизаторами с не полной нагрузкой. Он может быстро обнаруживать сбои на пути двунаправленной пересылки между двумя маршрутизаторами для протоколов верхнего уровня, таких как протоколы маршрутизации и многопротокольная коммутация по меткам (MPLS). В результате для поддержания производительности существующей сети используется резервный путь пересылки.

6.1.1. Протоколы и стандарты

- draft-ietf-bfd-base-09: обнаружение двунаправленной пересылки.
- draft-ietf-bfd-generic-05: общее применение BFD.
- draft-ietf-bfd-mib-06: информационная база управления обнаружением двунаправленной пересылки.
- draft-ietf-bfd-v4v6-1hop-09: BFD для IPv4 и IPv6 (Single Hop).
- draft-ietf-bfd-multihop-07: BFD для IPv4 и IPv6 (Multi-hop).
- draft-ietf-bfd-mpls-07: BFD для LSP MPLS.

ПРИМЕЧАНИЕ: в настоящее время draft-ietf-bfd-mib-06 и draft-ietf-bfd-multihop-07 не поддерживаются.

6.2. Приложения

Приложение	Описание
Поддержка BFD для OSPF	OSPF использует BFD для быстрого определения статуса соседа
Поддержка BFD для статической маршрутизации	Статическая маршрутизация использует BFD для быстрого определения доступности маршрута next-hop

6.2.1. Поддержка BFD для OSPF

6.2.1.1. Сценарий

Протокол Open Shortest Path First (OSPF) динамически обнаруживает соседний узел с помощью пакетов приветствия (hello packets). После включения BFD устанавливается сеанс BFD с соседом, находящимся в полной близости, для определения статуса соседа. Когда соседний узел выходит из строя, OSPF немедленно выполняет конвергенцию сети. Время сходимости можно сократить со 120 секунд (по умолчанию в нешироковещательной сети приветственные пакеты OSPF передаются с интервалом в 30 секунд, а время отказа соседнего узла в четыре раза превышает интервал, то есть 120 секунд) до 1 секунды.

В качестве примера приведен следующий рисунок. Маршрутизатор А и маршрутизатор В подключены через коммутатор уровня 2, OSPF настроен на маршрутизаторах для

установления маршрутов, а поддержка BFD для OSPF включена на интерфейсах маршрутизатора А и маршрутизатора В. Когда связь между маршрутизатором В и коммутатором уровня 2 выходит из строя, BFD может быстро обнаружить неисправность и сообщить об этом в OSPF, чтобы инициировать быструю конвергенцию OSPF.

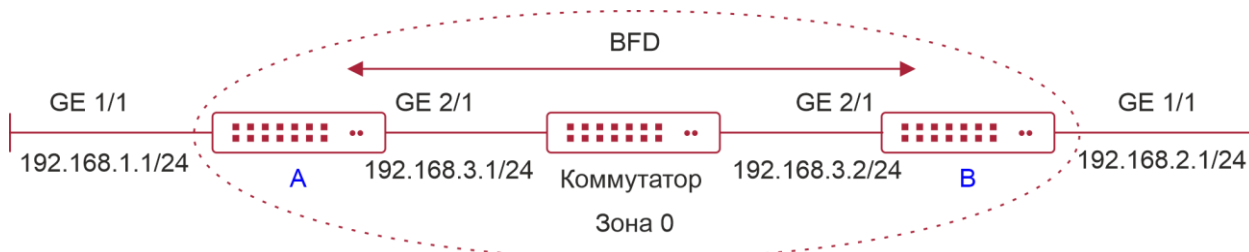


Рисунок 6-1.

А и В являются маршрутизаторами.

Коммутатор — это коммутатор уровня 2.

А и В подключены через коммутатор уровня 2.

6.2.1.2. Развертывание

- Настройте IP-адреса для взаимосвязанных интерфейсов маршрутизатора А и маршрутизатора В.
- Запустите OSPF на маршрутизаторе А и маршрутизаторе В.
- Установите параметры BFD на взаимосвязанных интерфейсах маршрутизатора А и маршрутизатора В.
- Включите поддержку BFD для OSPF на маршрутизаторе А и маршрутизаторе В.

6.2.2. Поддержка BFD для статической маршрутизации

6.2.2.1. Сценарий

Поддержка BFD для статической маршрутизации предотвращает выбор маршрутизаторами ошибочного статического маршрута в качестве пути пересылки и обеспечивает быстрое отключение маршрутизации при отказе, используя доступный резервный путь пересылки.

В отличие от протоколов динамической маршрутизации, статическая маршрутизация не имеет механизма обнаружения соседей (ND). Когда BFD поддерживает статическую маршрутизацию, доступность статического маршрута для next-hop зависит от состояния сеанса BFD. В случае сбоя сеанса BFD next-hop статического маршрута считается недостижимым и не будет добавлен в базу информации о маршрутизации (RIB).

В качестве примера приведем следующий рисунок. Маршрутизатор А и маршрутизатор В подключены через коммутатор уровня 2, на маршрутизаторах настроена статическая маршрутизация для установления путей пересылки, а поддержка статической маршрутизации BFD включена на интерфейсах маршрутизатора А и маршрутизатора В. Когда связь между маршрутизатором В и коммутатором уровня 2 неисправна, BFD может быстро обнаружить неисправность и сообщить о ней статической маршрутизации, чтобы заставить систему удалить статический маршрут из RIB, тем самым предотвращая ошибки маршрутизации.



Рисунок 6-2.

А и В являются маршрутизаторами.

Коммутатор — это коммутатор уровня 2.

А и В подключены через коммутатор уровня 2.

6.2.2.2. Развертывание

- Настройте IP-адреса для взаимосвязанных интерфейсов маршрутизатора А и маршрутизатора В.
- Настройте статическую маршрутизацию на маршрутизаторе А и маршрутизаторе В.
- Установите параметры BFD для взаимосвязанных интерфейсов маршрутизатора А и маршрутизатора В.
- Включите поддержку BFD для статической маршрутизации на маршрутизаторе А и маршрутизаторе В.

6.3. Функции

6.3.1. Базовые концепты

Формат пакета

Пакеты обнаружения, передаваемые BFD, представляют собой пакеты протокола пользовательских датаграмм (UDP), которые подразделяются на управляющие пакеты и эхо-пакеты. Эхо-пакеты касаются только локальной системы сеанса BFD. Поэтому их форматы не указаны. BFD определяет формат только управляющих пакетов. В настоящее время существует две версии (версия 0 и версия 1) формата управляющих пакетов. Версия 1 используется по умолчанию для установления сеанса BFD. Если устройство получает пакеты версии 0 от реер-системы, оно автоматически переключается на версию 0.

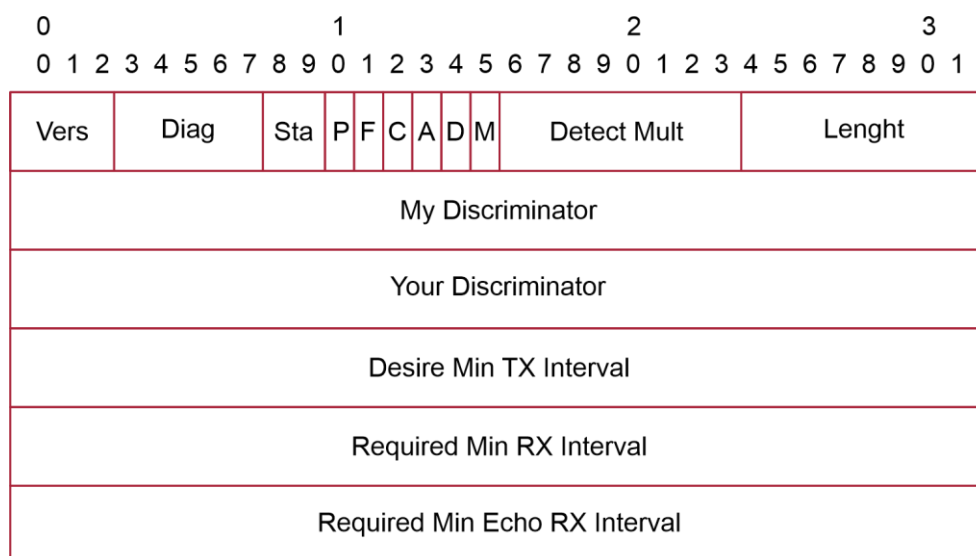


Рисунок 6-3.

Поле	Описание
Vers	Указывает номер версии протокола BFD, который в настоящее время равен 1
Diag	Указывает причину последнего изменения состояния сеанса локальной системы, в том числе: 0 – Нет диагностики. 1 – Время обнаружения контроля истекло. 2 – Функция эха (Echo) не удалась. 3 – Сосед сообщил об отключении сеанса. 4 – Сброс плоскости переадресации. 5 – Путь не в рабочем состоянии. 6 – Объединенный путь не в рабочем состоянии. 7 – Административно не в рабочем состоянии
Sta	Указывает состояние локального сеанса BFD, включая: 0 – AdminDown. 1 – Down. 2 – Init. 3 – Up
P	Указывает, что передатчик в сеансе BFD добавляет этот бит в запрос проверки при изменении параметра, ожидая ответа реер-узла



Поле	Описание
F	Указывает бит, который должен быть установлен в ответном пакете для ответа на бит P
C	Указывает, что control plane независима. Если установлено, изменения control plane не влияют на обнаружение BFD. Например, если control plane является OSPF, при перезапуске OSPF или плавном перезапуске (GR) BFD может продолжать определять состояние канала
A	Указывает на наличие аутентификации. Если установлено, сеанс должен быть аутентифицирован
D	Указывает на запрос по требованию. Если этот параметр установлен, передатчик желает обнаруживать соединения в режиме по запросу
M	Указывает многоточечный (multipoint) бит, который будет использоваться в расширениях «точка-многоточка» (point-to-multipoint). В настоящее время он должен быть установлен на 0
Detect Mult	Указывает множитель времени ожидания обнаружения. Он используется детектором для расчета времени ожидания обнаружения
Length	Указывает длину пакета
My Discriminator	Указывает дискриминатор local end, подключенного сеансом BFD
Your Discriminator	Указывает дискриминатор remote end, подключенного сеансом BFD
Desired Min Tx Interval	Указывает минимальный интервал передачи пакетов BFD, поддерживаемый local end
Required Min RX Interval	Указывает минимальный интервал получения пакетов BFD, поддерживаемый local end
Required Min Echo RX Interval	Указывает минимальный интервал приема эхо-пакетов, поддерживаемый local end. Он устанавливается равным 0, если local end не поддерживает функцию эха
Auth Type	(Опционально) Указывает тип аутентификации, в том числе: • Простой пароль (Simple Password) • Keyed MD5 • Meticulous Keyed MD5



Поле	Описание
	- Keyed SHA1 - Meticulous Keyed SHA1
Auth Length	Указывает длину данных аутентификации
Authentication Data	Указывает область данных аутентификации

Статус сеанса

Сеанс BFD может находиться в любом из четырех основных состояний: Down, Init, Up и AdminDown.

1. Down: указывает, что сеанс находится в состоянии Down или только что установлен.
2. Init: указывает, что локальная система связалась с peer-системой и хочет перевести сеанс в состояние Up.
3. Up: указывает, что сеанс успешно согласован.
4. AdminDown: указывает, что сеанс находится в состоянии AdminDown.

BFD переносит state machine на основе состояния локального сеанса и полученных пакетов BFD от peer end.

State machine BFD устанавливается и отключается с использованием механизма three-way handshake (трехстороннего рукопожатия), чтобы гарантировать, что оба конца знают об изменении статуса.

Интервал передачи и время обнаружения

Оба конца согласовывают параметры BFD во время установления сеанса BFD, чтобы определить интервал передачи и время обнаружения.

После установления сеанса BFD обе стороны могут динамически согласовывать параметры BFD (например, минимальный интервал передачи и минимальный интервал приема). После того, как протоколы на обоих концах передают соответствующие пакеты согласования, они принимают новый интервал передачи и время обнаружения, не влияя на текущее состояние сеанса.

6.3.1.1. Обзор

Особенность	Описание
Установление сеанса BFD	Устанавливает сеанс BFD
Обнаружение сеанса BFD	Быстро обнаруживает двунаправленный путь пересылки
Поддержка BFD для приложений	Быстро объявляет результат обнаружения BFD



Особенность	Описание
<u>Защита BFD</u>	Защищает BFD от атак для стабильности
<u>BFD Flapping</u> <u>Dampening</u>	Защищает стабильность связанных приложений в случае нестабильности линии

6.3.2. Установление сеанса BFD

Обнаружение BFD начинается с установления сеанса BFD.

6.3.2.1. Принцип работы

Процесс установления сеанса

Сам BFD не может обнаружить соседей. Ему нужен протокол верхнего уровня, чтобы указать соседа для установления сеанса.

Как показано на следующем рисунке, два маршрутизатора с OSPF и BFD подключены через коммутатор уровня 2.

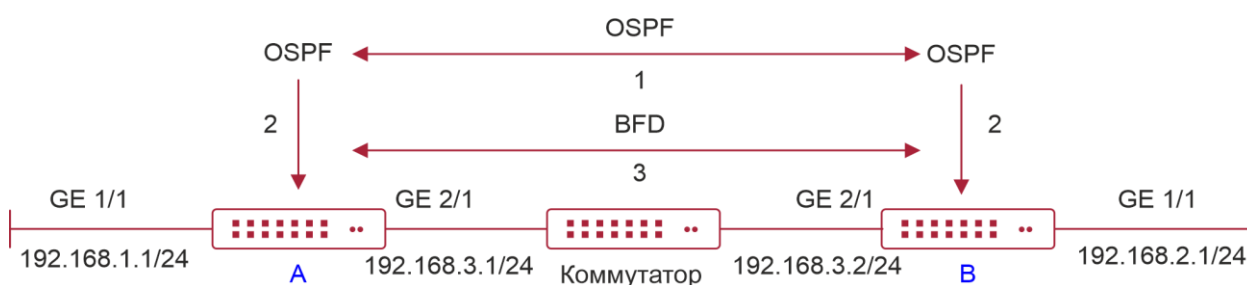


Рисунок 6-4.

Процесс установления сеанса BFD:

- OSPF обнаруживает соседа и устанавливает с ним соединение.
- OSPF инструктирует BFD установить сеанс с соседним устройством.
- BFD устанавливает сеанс с соседом.

Режим установления сеанса BFD

Протокол BFD указывает, что сеанс BFD может быть установлен в двух режимах:

- Активный режим

Перед установлением сеанса BFD активно передает управляющий пакет для установления сеанса BFD независимо от того, получает ли он управляющий пакет для установления сеанса BFD от peer end.

- Пассивный режим

BFD не передает активно управляющий пакет для установления сеанса BFD до того, как сеанс будет установлен, а ждет, пока не получит управляющий пакет для установления сеанса BFD от peer end.

ПРИМЕЧАНИЕ: пассивный режим в настоящее время не поддерживается.



Согласование параметров сеанса BFD

Оба конца согласовывают параметры сеанса BFD во время установления сеанса BFD, чтобы определить интервал передачи и время обнаружения. Обратите внимание на следующие моменты:

- Параметры сеанса BFD (включая **Desired Min Tx Interval**, **Required Min RX Interval**, и **Detect Mult**) должны быть установлены для интерфейсов на обоих концах. В противном случае сеанс BFD не может быть установлен.
- Интерфейсы на обоих концах согласовывают параметры сеанса BFD и определяют сеанс на основе параметров во время установления сеанса BFD.
- После установления сеанса BFD обе стороны могут динамически согласовывать параметры BFD (например, минимальный интервал передачи и минимальный интервал приема). После того, как протоколы на обоих концах передают соответствующие пакеты согласования, они принимают новый интервал передачи и время обнаружения, не влияя на текущее состояние сеанса.

6.3.3. Обнаружение сеанса BFD

Обнаружение канала начинается после установления сеанса BFD. BFD периодически передает пакеты управления BFD. Если ему не удастся получить пакеты BFD от peer end в течение времени обнаружения, он считает, что сеанс отключен, и уведомляет связанное приложение об ускорении конвергенции.

6.3.3.1. Принцип работы

Процесс обнаружения

Как показано на следующем рисунке, два маршрутизатора с OSPF и BFD подключены через коммутатор уровня 2.

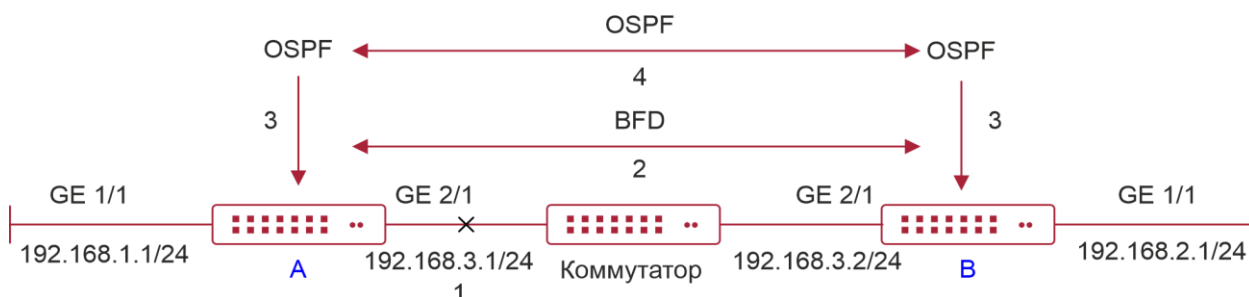


Рисунок 6-5.

Процедура обработки после отключения сеанса BFD:

- Сбой связи между маршрутизатором A и коммутатором.
- Сеанс BFD между маршрутизатором A и маршрутизатором B отключен.
- BFD уведомляет локальный OSPF о том, что путь пересылки к соседнему узлу неисправен.
- OSPF указывает соседа в статус «Не исправен» (Down). Если резервный путь пересылки доступен, он запускает конвергенцию протокола, чтобы включить альтернативный путь пересылки.

Режим обнаружения

BFD поддерживает следующие режимы обнаружения:



- Асинхронный режим

В асинхронном режиме системы периодически передают друг другу пакеты управления BFD. Если системе не удастся получить пакеты управления BFD от peer end в течение времени обнаружения, она объявляет, что сеанс отключен.

- Режим запроса

В режиме запроса предполагается, что каждая система имеет независимый метод подтверждения своей связи с другими системами. После установления сеанса BFD система прекращает передачу управляющих пакетов BFD, если только ей не требуется явная проверка подключения. В таком случае система передает пакет управления BFD с последовательностью кадров. Если системе не удастся получить возвращенный пакет в течение времени обнаружения, она объявляет, что сеанс отключен. Если он получает ответ от peer end, путь пересылки доступен.

- Режим эха

В эхо-режиме локальная система периодически передает эхо-пакеты BFD, а удаленная система получает и возвращает пакеты по пути пересылки. Если локальная система не может получить несколько последовательных эхо-пакетов в течение времени обнаружения, она объявляет, что сеанс отключен. Функцию эха можно использовать вместе с двумя предыдущими режимами обнаружения. Функция обнаружения эхо-пакетов не требует участия control plane удаленной системы. Пакеты возвращаются через forwarding plane удаленной системы, что снижает задержку и обеспечивает более быстрое обнаружение ошибок по сравнению с передачей контрольных пакетов. Включение функции эха в асинхронном режиме может значительно уменьшить передачу управляющих пакетов, поскольку обнаружение выполняется функцией эха. Включение функции эха в режиме запроса может полностью отменить передачу управляющих пакетов после установления сеанса. Функция эха должна быть включена на обоих концах сеанса BFD. В противном случае функция эха не действует.

ПРИМЕЧАНИЕ: режим запроса не поддерживается и в настоящее время не может быть настроен.

ПРИМЕЧАНИЕ: только сеанс BFD версии 1 поддерживает эхо-режим BFD.

ПРИМЕЧАНИЕ: эхо-режим не поддерживается для сеанса IPv6 BFD с локальным адресом канала в качестве исходного или целевого адреса.

6.3.4. Поддержка BFD для приложений

Благодаря поддержке BFD связанные приложения могут использовать быстрое обнаружение ошибок BFD для повышения производительности конвергенции протоколов. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.

6.3.4.1. Принцип работы

После включения поддержки BFD для определенного приложения сеанс BFD устанавливается на основе конфигурации BFD. При возникновении сбоя канала BFD может быстро идентифицировать сбой и уведомить связанное приложение для обработки, тем самым улучшая его конвергенцию. В настоящее время BFD поддерживает следующие приложения:

Поддержка BFD для RIP

После того как поддержка BFD для протокола информации о маршрутизации (RIP) включена, RIP может использовать обнаружение ошибок BFD, которое быстрее, чем механизм ND RIP, для улучшения конвергенции протоколов. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.



ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для RIP см. в IP Routing Configuration разделе Настройка RIP.

Поддержка BFD для OSPF

После включения поддержки BFD для OSPF OSPF может использовать обнаружение ошибок BFD, которое быстрее, чем механизм ND OSPF, для улучшения конвергенции протоколов. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для OSPF см. в IP Routing Configuration разделе Настройка OSPF.

Поддержка BFD для OSPFv3

После включения поддержки BFD для OSPFv3 OSPFv3 может использовать обнаружение ошибок BFD, которое быстрее, чем механизм ND OSPFv3, для улучшения конвергенции протоколов. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для OSPFv3 см. в IP Routing Configuration разделе Настройка OSPFv3.

Поддержка BFD для BGP

После включения поддержки BFD для протокола пограничного шлюза (BGP) BGP может использовать обнаружение ошибок BFD, которое быстрее, чем механизм ND BGP, для улучшения конвергенции протоколов. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для BGP см. в IP Routing Configuration разделе Настройка BGP.

Поддержка BFD для IS-IS

Протокол промежуточной системы к промежуточной системе (IS-IS) динамически обнаруживает соседа с помощью Hello-пакетов. После включения BFD IS-IS использует BFD для установления сеанса BFD с соседом, находящимся в состоянии Up, и определения состояния соседа. Когда сосед BFD выходит из строя, IS-IS немедленно выполняет конвергенцию сети. Время сходимости можно сократить с 30 секунд (по умолчанию в сети «точка-точка» Hello-пакеты IS-IS передаются с интервалом в 10 секунд, а время отказа соседа утроено интервалом, то есть 30 секунд) до 1 секунды.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для IS-IS см. в IP Routing Configuration разделе Настройка IS-IS.

Поддержка BFD для статической маршрутизации

После включения поддержки статической маршрутизации BFD BFD не позволяет маршрутизаторам выбирать недоступный статический маршрут в качестве пути пересылки во время маршрутизации и позволяет маршрутизаторам быстро переключаться на доступный резервный путь пересылки.

В отличие от протоколов динамической маршрутизации, статическая маршрутизация не имеет механизма ND. Таким образом, после настройки поддержки BFD для статической маршрутизации доступность статического маршрута для next-hop зависит от состояния сеанса BFD. Если сеанс BFD обнаруживает ошибку, next-hop статического маршрута недоступен, и статический маршрут не добавляется в RIB.

Если удаленная система удаляет сеанс BFD во время установления сеанса BFD, сеанс BFD становится отключенным (состояние Down). В этом случае система гарантирует, что поведение пересылки статической маршрутизации не будет затронуто.



ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для статической маршрутизации см. в IP Routing Configuration разделе Настройка статической маршрутизации.

Поддержка BFD для PBR

После настройки поддержки BFD для PBR BFD не позволяет маршрутизаторам выбирать недоступный маршрут политики в качестве пути пересылки во время маршрутизации и позволяет маршрутизаторам быстро переключаться на доступный резервный путь пересылки.

Поддержка BFD для PBR эквивалентна статической маршрутизации. BFD отслеживает и обнаруживает путь пересылки к указанному соседу. При сбое сеанса BFD BFD уведомляет PBR о том, что следующий переход недоступен. В этом случае маршрут политики к next hop не вступает в силу.

Если удаленная система удаляет сеанс BFD во время установления сеанса BFD, сеанс BFD становится отключенным. В этом случае система гарантирует, что поведение пересылки PBR не будет затронуто.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для PBR см. в IP Routing Configuration разделе Настройка PBR.

Поддержка BFD для VRRP

Поддержка BFD для протокола резервирования виртуальных маршрутизаторов (VRRP) может заменить механизм ND VRRP для быстрого определения рабочего состояния активных и резервных маршрутизаторов. Когда происходит сбой, он ускоряет переключение между активными и резервными маршрутизаторами и повышает производительность сети. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.

VRRP также может использовать BFD для отслеживания указанного соседа. Если путь пересылки к соседу выходит из строя во время сеанса BFD, он автоматически снижает приоритет VRRP до определенной степени, чтобы инициировать переключение между активными и резервными маршрутизаторами. Эта конфигурация вступает в силу только тогда, когда протокол динамической маршрутизации или другие приложения уведомляют BFD о необходимости установить сеанс с соседним устройством.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для VRRP см. в разделе [Настройка VRRP](#).

Поддержка BFD для VRRP Plus

Поддержка BFD для VRRP Plus может заменить обнаружение BVF, проводимое балансирующим виртуальным шлюзом (BVG) VRRP Plus, для быстрого определения рабочего состояния балансирующих виртуальных функций (BVF). Когда возникает ошибка, она ускоряет переключение объекта пересылки и повышает производительность сети. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.

VRRP Plus основан на протоколе VRRP. Таким образом, для поддержки BFD не требуется дополнительной настройки, и необходимо только включить VRRP на устройствах на обоих концах, а сеанс BFD правильно связан.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для VRRP Plus см. в разделе [Настройка VRRP Plus](#).

Поддержка BFD для MPLS

Поддержка BFD для коммутации по меткам с несколькими протоколами (MPLS) означает, что пути с коммутацией по меткам (LSP) используют BFD для быстрого определения статуса соседа. Поддерживаются следующие режимы обнаружения:



1. BFD обнаруживает статические LSP.
2. BFD обнаруживает LSP, созданные протоколом распространения меток (LDP).
3. BFD может обнаруживать обратные ссылки LSP с помощью протоколов IP.

Поддержка BFD для интерфейсов уровня 3

BFD поддерживает изменение состояния интерфейсов уровня 3. В режиме конфигурации интерфейса используйте команду **bfd bind peer-ip**, чтобы определить прямой адрес указанного интерфейса уровня 3. После выполнения этой команды CLI создается сеанс BFD, и состояние интерфейса уровня 3 может быть изменено в зависимости от результата обнаружения сеанса BFD, например, BFD Down или BFD Up. Эта функция часто используется в различных типах быстрого перенаправления (FRR), которые используют BFD для определения состояния интерфейса для реализации быстрого переключения FRR.

ПРИМЕЧАНИЕ: в поддержке BFD для интерфейсов уровня 3 поддерживается только коммутация LDP FRR.

Поддержка BFD для портов-участников L3AP

После включения поддержки BFD для портов-членов точки доступа уровня 3 (L3AP) BFD может быстро обнаруживать сбой, возникающий на канале порта-участника, чтобы трафик на этом канале быстро распределялся по другим действующим каналам-участникам. Как правило, время обнаружения неисправности может быть сокращено в пределах 1 секунды.

ПРИМЕЧАНИЕ: дополнительные сведения о поддержке BFD для портов-членов L3AP см. в Ethernet Switching разделе Настройка AP.

6.3.5. Защита BFD

Защита BFD используется для защиты BFD от нестабильности сеансов, вызванного атаками (например, атакой большого количества пакетов ping на устройства).

6.3.5.1. Принцип работы

Протокол BFD очень чувствителен. Если устройство с поддержкой BFD подвергается атаке (например, атаке с помощью большого количества пакетов ping) и сеансы BFD становятся нестабильными, можно настроить защиту BFD для обеспечения защиты. Если на устройстве включены и BFD, и защита BFD, устройство отбрасывает пакет BFD от previous-hop, что влияет на установление сеанса BFD между устройством previous-hop и другими устройствами.

6.3.6. BFD Flapping Dampening

Сеанс BFD может часто переключаться между Down и Up из-за нестабильности соединения. В результате связанное приложение (например, статическая маршрутизация) может часто переключать пути пересылки, что влияет на работающие службы. BFD Flapping Dampening может решить эту проблему.

6.3.6.1. Принцип работы

Сеанс BFD может часто переключаться между Down и Up. Эта функция позволяет пользователям устанавливать задержку объявления об изменении статуса. После того, как сеанс BFD находится в рабочем состоянии в течение определенного периода времени, BFD уведомляет связанное приложение об открытии BFD. В противном случае BFD уведомляет связанное приложение о недоступности BFD.



6.4. Конфигурация

Конфигурация	Описание и команда	
Настройка основных функций BFD	(Обязательно) Используется для установления сеанса BFD	
	bfd interval	Устанавливает параметры BFD
	N/A	Настраивает поддержку BFD для приложений. ПРИМЕЧАНИЕ: команда настройки зависит от связанных приложений. Подробнее см. в их руководствах по настройке
	(Опционально) Он используется для настройки режима обнаружения BFD, медленного таймера и поддержки BFD для интерфейсов уровня 3	
	bfd echo	Настраивает режим эха BFD
	bfd slow-timer	Настраивает таймер замедления BFD
Настройка основных функций BFD	bfd bind peer-ip	Настраивает поддержку BFD для интерфейсов уровня 3
Настройка защиты BFD	(Опционально) Используется для защиты BFD от атак	
	bfd cpp	Включает защиту BFD
Настройка BFD Flapping Dampening	(Опционально) Используется для защиты связанных протоколов от нестабильности BFD	
	bfd up-dampening	Настраивает BFD Flapping Dampening

6.4.1. Настройка основных функций BFD

6.4.1.1. Эффект конфигурации

- Настройте поддержку BFD для приложений.
- Установите сеанс BFD.
- Сеанс BFD обнаруживает сбой канала.

6.4.1.2. Примечания

- Обратите внимание на следующие моменты при настройке параметров сеанса BFD:
1. Рекомендуется, чтобы настройки параметров были одинаковыми на обоих концах сеанса BFD, чтобы протоколы приложений, связанные с BFD, вступали в силу

одновременно, и предотвращалось возникновение односторонней переадресации из-за разного времени подавления на обоих концах.

2. Учитывайте разницу в пропускной способности разных интерфейсов при настройке параметров. Если для минимального интервала передачи и минимального интервала приема установлены очень маленькие значения, передача данных может быть затруднена из-за очень большой занятости полосы пропускания BFD.
 - Обратите внимание на следующие моменты при настройке поддержки BFD для приложений:
 1. Убедитесь, что он включен для соседей сеанса BFD. В противном случае сеанс BFD не может быть установлен. Если протокол динамической маршрутизации или другое приложение требует, чтобы BFD установил сеанс с соседним устройством, сеанс BFD также может быть установлен.
 2. Если интерфейс, указанный в сеансе BFD, отличается от фактического исходящего интерфейса пакетов BFD из-за IP-маршрутизации или если интерфейс, указанный при создании сеанса BFD, отличается от фактического входящего интерфейса пакетов BFD, сеанс BFD установить невозможно.
 - Обратите внимание на следующие моменты при настройке режима обнаружения BFD:
 1. В процессе, когда forwarding plane реер-устройства возвращает эхо-пакеты, переданные локальным концом, на локальный конец, эхо-пакеты могут быть потеряны из-за перегрузки реер-устройства, что приводит к сбою обнаружения сеанса. В этом случае настройте политики качества обслуживания (QoS), чтобы обеспечить предпочтительную обработку эхо-пакетов, или отключите функцию эха.
 2. Функция обнаружения эха BFD не поддерживает обнаружение multi-hop. Убедитесь, что функция эха отключена при настройке multi-hop.
 3. Режим эха вступает в силу только после включения этого режима на обоих концах сеанса BFD.
 4. Перед включением эхо-режима BFD запустите команду **no ip redirects** на соседях сеанса BFD, чтобы отключить функцию перенаправления пакетов ICMP, и запустите команду **no ip deny land**, чтобы отключить функцию распределенного отказа в обслуживании (DDoS) (предотвратить наземную атаку).

6.4.1.3. Шаги настройки

Настройка параметров BFD

- Обязательный.
- Параметры BFD должны быть установлены на выходе сеанса BFD маршрутизаторов на обоих концах, обнаруженных BFD, если не предъявляются особые требования.
- Учитывайте разницу в пропускной способности разных интерфейсов при настройке параметров. Если для минимального интервала передачи и минимального интервала приема установлены очень маленькие значения, передача данных может быть затруднена из-за очень большой занятости полосы пропускания BFD.

Команда	bfd interval <i>milliseconds</i> min_rx <i>milliseconds</i> multiplier <i>interval-multiplier</i>
Описание параметров	interval <i>milliseconds</i> : указывает минимальный интервал передачи в миллисекундах.



	min_rx milliseconds: указывает минимальный интервал приема в миллисекундах. multiplier interval-multiplier: указывает множитель времени ожидания обнаружения
По умолчанию	Параметр сеанса BFD не настроен
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Функция быстрой переадресации должна быть включена до включения функции BFD на маршрутизаторах

Включение режима эха BFD

- (Опционально) По умолчанию порты работают в асинхронном режиме. Если сеанс BFD должен выполняться в эхо-режиме, этот эхо-режим необходимо настроить.
- Завершите настройку портов коммутаторов или маршрутизаторов.
- Сеанс выполняется в асинхронном режиме, пока любой из маршрутизаторов на обоих концах настроен для работы в асинхронном режиме. Если маршрутизаторы на обоих концах по умолчанию настроены на работу в эхо-режиме, сеанс BFD в конечном итоге выполняется в эхо-режиме.

Команда	bfd echo
По умолчанию	Режим эха BFD отключен
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эту команду нельзя настроить для агрегированных портов. По умолчанию при установке параметров сеанса BFD система автоматически включает режим эха. Минимальный интервал TX и минимальный интервал RX эхо-пакетов принимают параметры сеанса interval milliseconds и min_rx milliseconds. Перед включением эхо-режима BFD запустите команду no ip redirects на соседях сеанса BFD, чтобы отключить функцию перенаправления пакетов ICMP, и запустите команду no ip deny land, чтобы отключить функцию распределенного отказа в обслуживании (DDoS) (предотвратить наземную атаку)

Настройка медленного таймера BFD

- (Опционально) Таймер замедления по умолчанию составляет 3000 миллисекунд. Значение может быть изменено по мере необходимости.
- Настройте эту функцию в режиме глобальной конфигурации коммутаторов или маршрутизаторов.



- В эхо-режиме BFD или построении сеанса для управления пакетами используется медленный таймер. Если значение увеличивается, время, необходимое для согласования и установления сеанса BFD, увеличивается, а время, необходимое для передачи медленных пакетов BFD в эхо-режиме, увеличивается.

Команда	bfd slow-timer [<i>milliseconds</i>]
Описание параметров	<i>milliseconds</i> : указывает на медленный таймер BFD с единицей измерения в миллисекундах. Диапазон значений составляет от 1000 до 30 000, а значение по умолчанию 3000 принимается, если оно не установлено
По умолчанию	Интервал передачи медленных управляющих пакетов составляет 3000 миллисекунд
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Эта команда используется для указания медленного таймера в эхо-режиме

Настройка поддержки BFD для интерфейсов уровня 3

- (Опционально) В настоящее время эта функция используется только тогда, когда MPLS LDP используется для FRR.
- Настройте эту функцию на интерфейсах коммутаторов или маршрутизаторов.

Команда	bfd bind peer-ip <i>src-address</i> [<i>source-ip</i> <i>dst-address</i>] process-pst
Описание параметров	<i>src-address</i> : указывает peer IP-адрес интерфейса. <i>dst-address</i> : указывает локальный IP-адрес интерфейса
По умолчанию	Поддержка BFD для интерфейсов уровня 3 по умолчанию не настроена
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эта команда используется для включения поддержки BFD для интерфейсов уровня 3, чтобы быстро обнаруживать возможность подключения интерфейсов уровня 3

Настройка поддержки BFD для приложений

- Обязательный.
- Эта функция отключена по умолчанию.
- Команда настройки зависит от связанных приложений. Подробнее см. в их руководствах по настройке.
- Эта функция должна быть настроена на обоих концах, чтобы можно было установить сеанс BFD.



- В режиме конфигурации маршрутизации RIP запустите команду **bfd all interfaces**, чтобы включить поддержку BFD для RIP на всех интерфейсах. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка RIP.
- В режиме конфигурации маршрутизации OSPF запустите команду **bfd all interfaces**, чтобы включить поддержку BFD для OSPF на всех интерфейсах. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка OSPF.
- В режиме конфигурации маршрутизации OSPFv3 запустите команду **bfd all interfaces**, чтобы включить поддержку BFD для OSPFv3 на всех интерфейсах. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка OSPFv3.
- В режиме настройки маршрутизации BGP запустите команду **bfd** для перехода к соседнему адресу, чтобы включить поддержку BFD для BGP. Дополнительные сведения см. в разделе Настройка BGP.
- В режиме конфигурации маршрутизации IS-IS запустите команду **bfd all interfaces**, чтобы включить поддержку BFD для IS-IS на всех интерфейсах. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка IS-IS.
- В режиме глобальной конфигурации запустите команду **ip route static bfd [vrf vrf-name] interface-type interface-number gateway [source ip-address]**, чтобы включить поддержку BFD для статической маршрутизации. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка статической маршрутизации.
- В режиме глобальной конфигурации запустите команду **ipv6 route static bfd [vrf vrf-name] interface-type interface-number gateway [source ip-address]**, чтобы включить поддержку BFD для статической маршрутизации IPv6. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка статической маршрутизации.
- Запустите команду **set ip next-hop verify-availability next-hop-address bfd [vrf vrf-name] interface-type interface-number gateway**, чтобы включить поддержку BFD для PBR. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка PBR.
- Запустите команду **set ipv6 next-hop verify-availability next-hop-address bfd [vrf vrf-name] interface-type interface-number gateway**, чтобы включить поддержку BFD для IPv6 PBR. Дополнительные сведения см. в разделе IP Routing Configuration/Настройка PBR.
- Запустите команду **vrrp bfd interface-type interface-number ip-address**, чтобы включить поддержку BFD для VRRP. Дополнительные сведения см. в разделе [Настройка VRRP](#).
- VRRP Plus основан на протоколе VRRP. Поэтому для поддержки BFD для VRRP Plus не требуется дополнительная настройка. Только VRRP должен быть включен на устройствах на обоих концах, и сеанс BFD правильно связан.

6.4.1.4. Проверка

Команда проверки зависит от связанных приложений. Подробнее см. в их руководствах по настройке.



6.4.1.5. Пример конфигурации

Настройка поддержки BFD для OSPF

Сценарий:

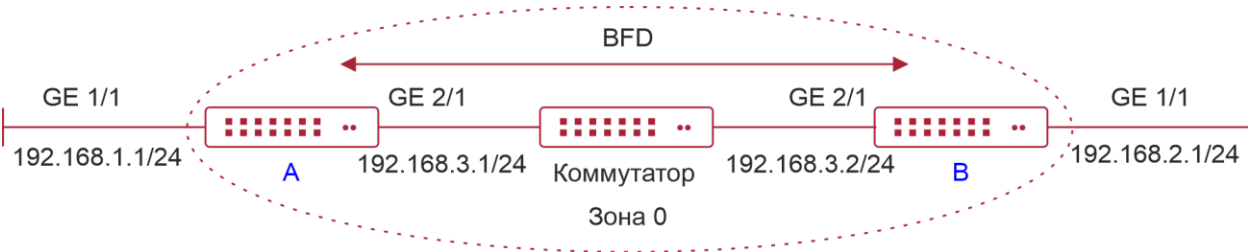


Рисунок 6-6.

Шаги настройки	• Настройте IP-адреса для взаимосвязанных интерфейсов маршрутизатора A и маршрутизатора B. • Запустите OSPF на маршрутизаторе A и маршрутизаторе B. • Установите параметры BFD для взаимосвязанных интерфейсов маршрутизатора A и маршрутизатора B. • Включите поддержку BFD для OSPF на маршрутизаторе A и маршрутизаторе B
A	<pre>A#configure terminal A(config)#interface GigabitEthernet2/1 A(config-if-GigabitEthernet2/1)# no switchport // На маршрутизаторах конфигурация не требуется. A(config-if-GigabitEthernet2/1)#ip address 192.168.3.1 255.255.255.0 A(config-if-GigabitEthernet2/1)#bfd interval 200 min_rx 200 multiplier 5 A(config-if-GigabitEthernet2/1)# exit A(config)#interface GigabitEthernet1/1 A(config-if-GigabitEthernet1/1)# no switchport // На маршрутизаторах конфигурация не требуется. A(config-if-GigabitEthernet1/1)#ip address 192.168.1.1 255.255.255.0 A(config-if-GigabitEthernet1/1)# exit A(config)# router ospf 123 A(config-router)# log-adj-changes detail A(config-router)# network 192.168.3.0.0.0.255 area 0 A(config-router)# network 192.168.1.0.0.0.255 area 0 A(config-router)# bfd all-interfaces A(config-router)# end</pre>



B	<pre> B#configure terminal B(config)#interface GigabitEthernet2/1 B(config-if-GigabitEthernet2/1)# no switchport // На маршрутизаторах конфигурация не требуется. B(config-if-GigabitEthernet2/1)#ip address 192.168.3.2 255.255.255.0 B(config-if-GigabitEthernet2/1)#bfd interval 200 min_rx 200 multiplier 5 B(config-if-GigabitEthernet2/1)# exit B(config)#interface GigabitEthernet1/1 B(config-if-GigabitEthernet1/1)# no switchport // На маршрутизаторах конфигурация не требуется. B(config-if-GigabitEthernet1/1)#ip address 192.168.2.1 255.255.255.0 B(config-if-GigabitEthernet1/1)# exit B(config)# router ospf 123 B(config-router)# log-adj-changes detail B(config-router)# network 192.168.3.0.0.0.255 area 0 B(config-router)# network 192.168.2.0.0.0.255 area 0 B(config-router)# bfd all-interfaces B(config-router)# end </pre>
Проверка	Показать подтверждение
A	<pre> A# show bfd neighbors details OurAddr NeighAddr LD/RD RH/RS Holdown(mult) State Int 192.168.3.1 192.168.3.2 1/2 Up 532 (3) Up Ge2/1 Local Diag: 0, Demand mode: 0, Poll bit: 0 MinTxInt: 200000, MinRxInt: 200000, Multiplier: 5 Received MinRxInt: 50000, Received Multiplier: 3 Holdown (hits): 600(22), Hello (hits): 200(84453) Rx Count: 49824, Rx Interval (ms) min/max/avg: 208/440/332 Tx Count: 84488, Tx Interval (ms) min/max/avg: 152/248/196 Registered protocols: OSPF Uptime: 02:18:49 Last packet: Version: 1 - Diagnostic: 0 I Hear You bit: 1 - Demand bit: 0 Poll bit: 0 - Final bit: 0 Multiplier: 3 - Length: 24 My Discr.: 2 - Your Discr.: 1 </pre>



	Min tx interval: 50000 - Min rx interval: 50000 Min Echo interval: 0
B	B# show bfd neighbors details OurAddr NeighAddr LD/RD RH/RS Holdown(mult) State Int 192.168.3.2 192.168.3.1 2/1 Up 532 (5) Up Ge2/1 Local Diag: 0, Demand mode: 0, Poll bit: 0 MinTxInt: 50000, MinRxInt: 50000, Multiplier: 3 Received MinRxInt: 200000, Received Multiplier: 5 Holdown (hits): 600(22), Hello (hits): 200(84453) Rx Count: 49824, Rx Interval (ms) min/max/avg: 209/440/332 last: 66 ms ago Tx Count: 84488, Tx Interval (ms) min/max/avg: 153/249/197 last: 190 ms ago Registered protocols: OSPF Uptime: 02:18:49 Last packet: Version: 1 - Diagnostic: 0 I Hear You bit: 1 - Demand bit: 0 Poll bit: 0 - Final bit: 0 Multiplier: 5 - Length: 24 My Discr.: 1 - Your Discr.: 2 Min tx interval: 200000 - Min rx interval: 200000 Min Echo interval: 0

6.4.1.6. Распространенные ошибки

- Параметры BFD не устанавливаются для интерфейсов устройств на одном конце.
- Поддержка BFD для приложений отключена.
- Поддержка BFD для приложений включена только на одном конце.

6.4.2. Настройка защиты BFD

6.4.2.1. Эффект конфигурации

Если устройство с поддержкой BFD подвергается атаке (например, атаке с использованием большого количества ring-пакетов) и сеанс BFD соответственно нестабилен, для обеспечения защиты можно включить защиту BFD.

6.4.2.2. Примечания

- Основные функции BFD должны быть настроены.
- Если на устройстве включены и BFD, и защита BFD, устройство отбрасывает пакет BFD от previous-hop, что влияет на установление сеанса BFD между previous-hop и другими устройствами.
- Эта функция и ограничения применимы только к коммутаторам.



6.4.2.3. Шаги настройки

Включение защиты BFD

- Опционально.
- Настройте эту функцию в режиме глобальной конфигурации на коммутаторах или маршрутизаторах.
- Функция защиты BFD повышает приоритет обработки пакетов BFD и обеспечивает нормальную работу служб BFD в сценарии, в котором устройства подвергаются атаке.

Команда	bfd cpp
По умолчанию	Функция защиты BFD включена по умолчанию
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Включите функцию защиты BFD, чтобы обеспечить защиту, если устройство сталкивается с нестабильностью BFD из-за атак

6.4.2.4. Проверка

Запустите команду **show running-config**, чтобы проверить конфигурацию интерфейса.

6.4.2.5. Пример конфигурации

Включение защиты BFD

Шаги настройки	• Настройте эту функцию на коммутаторе в сети, где существуют атаки. • Настройте функцию защиты BFD
	<pre>QTECH #configure terminal QTECH (config)# bfd cpp QTECH (config)# end</pre>

6.4.3. Настройка BFD Flapping Dampening

6.4.3.1. Эффект конфигурации

- Сеанс BFD может часто переключаться между Down и Up из-за нестабильности соединения. В результате соответствующее приложение (такое как статическая маршрутизация) может часто переключать пути пересылки, что влияет на работающие службы.
- Пользователи могут установить задержку для объявления об изменении статуса, после чего BFD уведомляет связанное приложение о BFD Up. После того, как сеанс BFD находится в рабочем состоянии (Up) в течение определенного периода времени, BFD уведомляет связанное приложение об BFD Up. В противном случае BFD уведомляет об отключении BFD (Down). Цель состоит в том, чтобы



уменьшить колебания связанных протоколов, вызванные нестабильными соединениями.

6.4.3.2. Примечания

- Основные функции BFD должны быть настроены.
- Если сеанс BFD не часто переключается между состояниями Down и Up, включение BFD Flapping Dampening приведет к задержке уведомления связанного приложения BFD Up.

6.4.3.3. Шаги настройки

Настройка BFD Flapping Dampening

- (Опционально) По умолчанию BFD Flapping Dampening отключено на портах. Если сеанс BFD часто переключается между Down и Up, рекомендуется включить эту функцию.
- Настройте эту функцию на портах коммутаторов или маршрутизаторов.
- При включенном BFD Flapping Dampening облегчается то, что связанные приложения, такие как перерасчет маршрута, обрабатывают количество объявлений из-за частого изменения состояния BFD. Чем больше настроенное время, тем больше требуемое время стабильности BFD. BFD уведомляет модуль приложения о BFD Up только после того, как время стабильности достигает настроенного времени.

Команда	bfd up-dampening [milliseconds]
Описание параметров	<i>milliseconds</i> : указывает задержку для объявления об изменении состояния, после которой BFD уведомляет связанное приложение BFD Up с единицей измерения в миллисекундах. Диапазон значений от 0 до 300 000. Значение 0 указывает, что BFD немедленно уведомляет прикладной уровень, когда сеанс переключается с Down на Up, и значение по умолчанию равно 0
По умолчанию	Функция BFD Flapping Dampening по умолчанию отключена
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эту функцию необходимо включать только тогда, когда связь нестабильна. Если сеанс BFD не часто переключается между состояниями Down и Up, включение BFD Flapping Dampening приведет к задержке уведомления связанного приложения BFD Up

6.4.3.4. Проверка

Запустите команду **show running-config**, чтобы проверить конфигурацию интерфейса.



6.4.3.5. Пример конфигурации

Настройка BFD Flapping Dampening с задержкой объявления, равной 60 000 миллисекунд

Шаги настройки	- Настройте эту функцию в среде, где BFD часто переключается из-за нестабильности канала. - Установите задержку для объявления об изменении статуса на 60 000 миллисекунд
	<pre>QTECH #configure terminal QTECH (config)# interface fastEthernet 0/2 QTECH (config)# bfd up-dampening 60000 QTECH (config)# end</pre>

6.5. Мониторинг

6.5.1. Отображение

Описание	Команда
Отображает информацию о сеансе BFD	<code>show bfd neighbors [vrf vrf-name] [client {ap bgp ospf rip vrrp static-route pbr vrrp-balance ldp-lsp static-lsp backward-lsp-with-ip pst }][ipv4 ip-address ipv6 ip-address][details]</code>

6.5.1.1. Отладка

ПРИМЕЧАНИЕ: системные ресурсы заняты при выводе отладочной информации. Поэтому отключайте отладку сразу после использования.

Описание	Команда
Отладка событий BFD	<code>debug bfd event[interface interface-type interface-number ipv4ip-address ipv6 ipv6-address]</code>
Отладка пакетов BFD	<code>debug bfd packet[interface interface-type interface-number ipv4ip-address ipv6 ipv6-address]</code>



7. НАСТРОЙКА ПОДАВЛЕНИЯ IP-СОБЫТИЙ

7.1. Обзор

Когда порт уровня 3 на устройстве уровня 3 часто переключается между Up и Down из-за ручного включения/отключения или других внешних причин, таблица маршрутизации на устройстве будет постоянно нестабильна. Если настроен протокол маршрутизации, этот протокол может распространять нестабильность на всю сеть, вызывая повторные обновления и перерасчет соседних маршрутов, что приводит к нерациональному использованию пропускной способности сети и дестабилизации сети. Многократные обновления маршрутов и перерасчет на устройствах потребляют много ресурсов ЦП, что влияет на нормальную работу клиентских сетей.

Подавление IP-событий обнаруживает аномальные значения переключений Up/Down и автоматически подавляет частые изменения состояния порта, который предотвращает распространение отказов single-point-канала протоколом маршрутизации. Когда порт восстановится, он будет автоматически отключен, что уменьшит количество сетевых сбоев и потребление ресурсов процессора при одновременном повышении стабильности сети.

7.1.1. Протоколы и стандарты

RFC2439: подавление нестабильности маршрута BGP (BGP Route Flap Dampening).

ПРИМЕЧАНИЕ: по своей сути, алгоритм подавления IP-событий (IP Event Dampening), аналогичен алгоритму, используемому в BGP Route Flap Dampening.

7.2. Приложение

Приложение	Описание
Routed Port Flap Dampening	Отслеживает изменение состояния порта уровня 3 на маршрутизаторе и подавляет нестабильность портов

7.2.1. Routed Port Flap Dampening

7.2.1.1. Сценарий

В сети, в которой работает протокол маршрутизации, когда порт на маршрутизаторе, подключенном к другому маршрутизатору, часто переключается Up и Down, соседние маршруты будут неоднократно обновляться и пересчитываться. Протокол маршрутизации может распространять нестабильность на всю сеть, вызывая нестабильность сети. На подключенных маршрутизаторах можно включить подавление IP-событий, чтобы отслеживать изменения состояния портов и подавлять нестабильность портов, тем самым уменьшая нестабильность сети и потребление ресурсов ЦП при одновременном повышении стабильности сети.



Рисунок 7-1.

А и В являются маршрутизаторами.

7.2.1.2. Развертывание

Настройте подавление IP-событий на порту GE0/1 на маршрутизаторе А и порту GE0/1 на маршрутизаторе В соответственно.

ПРИМЕЧАНИЕ: подинтерфейсы и виртуальные шаблоны интерфейсов на маршрутизаторах не поддерживают функцию подавления.

7.3. Функции

7.3.1. Базовые концепты

Штраф

Порт, который переходит в состояние Up или Down, получает штраф за каждое изменение состояния, но штраф уменьшается экспоненциально, когда порт стабилен. Таким образом, можно разумно отслеживать поведение порта и управлять им.

Порог подавления

Когда совокупный штраф для порта превышает порог подавления, считается, что порт нестабилен и будет заблокирован.

Период полураспада

Период полураспада — это период, необходимый для уменьшения штрафа до половины исходного значения, когда порт стабилен. Он определяет скорость, с которой штраф затухает экспоненциально. Чем короче период полураспада, тем быстрее спадает штраф и тем быстрее обнаруживается стабильность порта, но снижается чувствительность обнаружения нестабильности (flap).

Порог повторного использования

Когда порт долго стабилен и его штраф уменьшается до определенной степени (ниже порога подавления), порт считается стабильным и не блокируется.

Максимальное время подавления

Когда порт продолжает оставаться нестабильным и достигает очень большого штрафа, порт нельзя будет использовать в течение длительного времени. Чтобы избежать этой проблемы, определяется максимальное время блокировки, чтобы всегда поддерживать продолжительность блокировки порта ниже определенного значения, независимо от того, как долго порт находится в состоянии нестабильности.



7.4. Обзор

Особенность	Описание
Port Flap Suppression	Настройте критерии и параметры flap suppression на портах, чтобы коммутаторы могли выявлять и подавлять часто переключающиеся порты, что обеспечивает стабильность маршрута и предотвращает распространение нестабильности маршрута

7.4.1. Port Flap Suppression

7.4.1.1. Принцип работы

Порту, настроенному с подавлением IP-событий, назначается штраф. Порт получает штраф 1000 каждый раз, когда он выходит из строя, но штраф уменьшается со временем. Если порт снова выходит из строя, штраф соответственно увеличивается. Когда совокупный штраф превышает порог подавления, порт будет подавлен. Для затронутого протокола верхнего уровня подавляемый порт всегда находится в состоянии Down независимо от фактического состояния порта. Когда штраф уменьшится до порога повторного использования, порт не будет подавлен, и протокол верхнего уровня сможет определить фактическое состояние порта.

Если порт уровня 3 не настроен с подавлением IP-событий или не подавляется им, протокол маршрутизации или другой протокол, связанный с состоянием порта, по-прежнему работает нормально. Когда порт подавлен, протокол верхнего уровня считает, что порт отключен. Любое изменение состояния порта до того, как он будет восстановлен, не влияет на таблицу маршрутизации, а также на расчет и объявление маршрута, выполняемые протоколом маршрутизации верхнего уровня.

7.4.1.2. Связанная конфигурация

Настройка подавления IP-событий

- По умолчанию подавление IP-событий отключено на портах уровня 3.
- Запустите команду **dampening** [*half-life-period* [*reuse-threshold suppress-threshold max-suppress* [**restart** [*restart-penalty*]]]], чтобы включить или отключить подавление IP-событий на портах уровня 3.

7.5. Конфигурация

Конфигурация	Описание и команда	
Включение подавления IP-событий	(Обязательно) Он используется для подавления нестабильности портов уровня 3	
	dampening	Настраивает подавление IP-событий



7.5.1. Включение подавления IP-событий

7.5.1.1. Эффект конфигурации

Если порт, для которого настроено подавление IP-событий, продолжает оставаться нестабильным до тех пор, пока не будет превышен предварительно заданный порог, для порта устанавливается значение Down.

7.5.1.2. Примечания

Когда порт уровня 3 на коммутаторе преобразуется в порт уровня 2 (например, из маршрутизируемого порта в порт коммутатора), конфигурация подавления IP-событий на порте будет удалена.

7.5.1.3. Шаги настройки

Настройка подавления IP-событий

- Обязательный.
- Выполните настройку в режиме настройки интерфейса уровня 3.
- Вы можете указать период полураспада, порог повторного использования, порог подавления, максимальное время подавления и начальный штраф. Если вы не зададите эти параметры, будут использоваться их значения по умолчанию.

7.5.1.4. Проверка

Используйте любую из следующих команд, чтобы проверить, действует ли конфигурация:

- **show running-config**
- **show interfaces [interface-id] dampening**, которое используется для проверки конфигурации подавления IP-событий на указанном порту.

7.5.1.5. Связанные команды

Включение подавления IP-событий на порту

Команда	dampening [<i>half-life-period</i> [<i>reuse-threshold</i> <i>suppress-threshold</i> <i>max-suppress</i> [restart [<i>restart-penalty</i>]]]]
Описание параметров	<i>half-life-period</i>: указывает период полураспада. Диапазон значений: <1–30>; значение по умолчанию: 5 секунд. <i>reuse-threshold</i>: указывает порог повторного использования. Диапазон значений: <1–20 000>; значение по умолчанию: 1000. <i>suppress-threshold</i>: указывает порог подавления. Диапазон значений: <1–20 000>; значение по умолчанию: 2000. <i>max-suppress</i>: указывает максимальное время подавления. Диапазон значений: <1–255>; значение по умолчанию: четырехкратный период полураспада. restart <i>restart-penalty</i>: указывает первоначальный штраф. Диапазон значений: <1–20 000>; значение по умолчанию: 2000
Командный режим	Режим конфигурации интерфейса



Руководство по использованию	Подавление IP-событий может повлиять на прямые маршруты, маршруты хоста, статические маршруты, динамические маршруты и VRRP. Когда порт подавляется на основе настроенных критериев, затронутые модули определяют, что порт отключен, и поэтому удаляют соответствующие маршруты. Пакет данных не будет передаваться через порт. При повторном запуске команды dampening для порта, настроенного с подавлением IP-событий, информация о подавлении порта будет очищена, но flap-счетчик будет сохранен, если вы не используете команду очистки счетчиков для сброса счетчиков на порту. Если параметр max-suppress установлен на очень маленькое значение, что делает максимальный штраф меньше порога подавления, порт никогда не будет подавлен. При возникновении такой ошибки конфигурации будет напечатано следующее сообщение, указывающее на ошибку конфигурации: % Maximum penalty (10) is less than suppress penalty (2000). Increase maximum suppress time Если доступной системной памяти недостаточно для выполнения команды dampening, будет напечатано следующее сообщение, указывающее на сбой конфигурации: % No memory, configure dampening fail!
------------------------------	---

7.5.1.6. Пример конфигурации

Настройка подавления IP-событий на портах уровня 3

Сценарий:



Рисунок 7-2.

Шаги настройки	Включите подавление IP-событий на порте GigabitEthernet 0/1 на маршрутизаторе A и на порту GigabitEthernet 0/1 на маршрутизаторе B соответственно и установите half-time-period на 30 секунд, reuse-threshold на 1500, suppress-threshold на 10 000 и max-suppress до 120 секунд
A	<pre>QTECH (config)#interface GigabitEthernet 0/1 QTECH (config-if-GigabitEthernet 0/1)#dampening 30 1500 10000 100</pre>
B	<pre>QTECH (config)#interface GigabitEthernet 0/1 QTECH (config-if-GigabitEthernet 0/1)#dampening 30 1500 10000 100</pre>



Проверка	Запустите команду show interfaces dampening , чтобы проверить конфигурацию подавления IP-событий на соответствующих портах
	<pre>QTECH #show interfaces dampening GigabitEthernet 0/1 Flaps Penalty Supp ReuseTm HalfL ReuseV SuppV MaxSTmMaxP Restart 0 0 FALSE 0 30 1500 1000 100 15119 0</pre>

7.5.1.7. Распространенные ошибки

Порт на коммутаторе уровня 3 не преобразуется в маршрутизируемый порт с помощью команды **no swiport** до настройки подавления IP-событий.

7.6. Мониторинг

7.6.1. Очистка

Описание	Команда
Очищает счетчики интерфейса	clear counters

ПРИМЕЧАНИЕ: дополнительные сведения о команде **clear counter** см. в соответствующей главе, посвященной команде «Интерфейс».

7.6.2. Отображение

Описание	Команда
Отображает счетчики на заблокированных портах	show dampening interface
Отображает конфигурацию подавления IP-событий на портах	show interfaces dampening

7.6.2.1. Отладка

ПРИМЕЧАНИЕ: системные ресурсы заняты при выводе отладочной информации. Поэтому отключайте отладку сразу после использования.

Описание	Команда
Включает отладку подавления IP-событий	debug dampening interface



8. НАСТРОЙКА VSU

8.1. Обзор

Чтобы повысить надежность сетей, два устройства на уровне ядра и уровне конвергенции традиционных сетей конфигурируются с двумя ядрами для обеспечения резервирования. Устройства доступа и конвергенции соответственно подключаются к ядрам двумя каналами. На следующем рисунке показана типичная традиционная сетевая архитектура. Резервированная сетевая архитектура увеличивает сложность проектирования и эксплуатации сети. В то же время большое количество резервных каналов снижает использование сетевых ресурсов и окупаемость инвестиций.

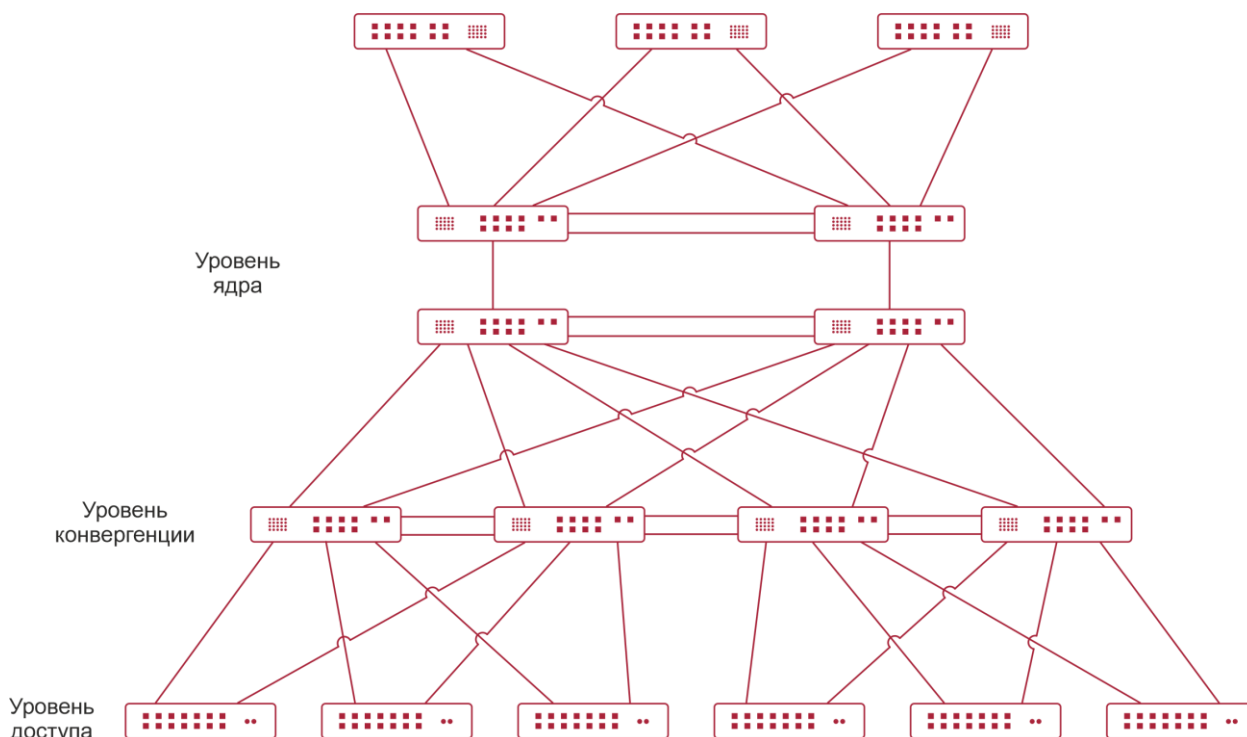


Рисунок 8-1. Традиционная сетевая архитектура

Virtual Switching Unit (VSU) — это своего рода технология виртуализации сетевых систем, которая поддерживает объединение нескольких устройств в одно виртуализированное устройство. Как показано на Рисунке 8-2, устройства доступа, агрегации и уровня ядра могут соответственно образовывать VSU, а затем эти VSU соединяются друг с другом, образуя сквозную сеть VSU. По сравнению с традиционной сетью эта сеть предоставляет:

- Упрощение топологии сети.
- Сокращение затрат на управление и обслуживание сети.
- Сокращение времени восстановления приложений и времени прерывания обслуживания.
- Повышение эффективности использования сетевых ресурсов.

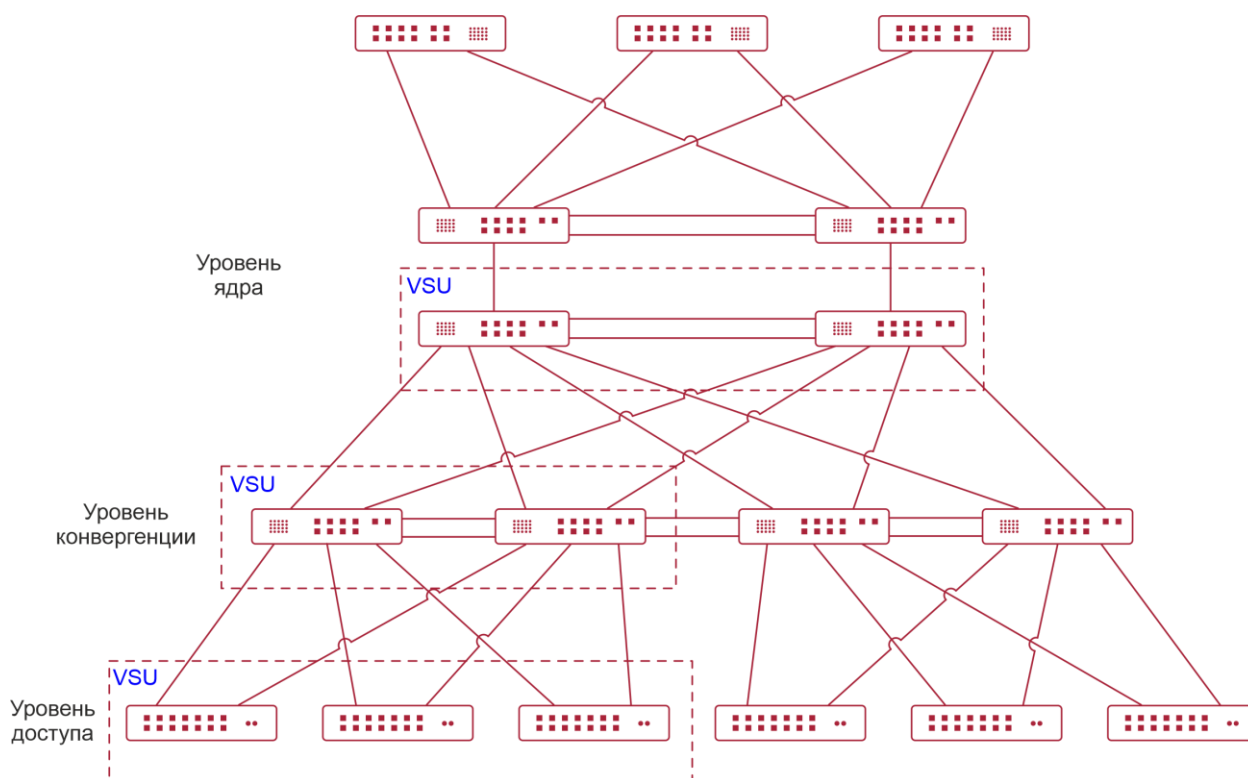


Рисунок 8-2. End-to-End сеть VSU

8.2. Приложения

Приложение	Описание
<u>Единое управление несколькими устройствами</u>	Использует несколько физических устройств в качестве логического устройства для унифицированного управления
<u>Упрощение сетевой топологии</u>	Использует VSU в качестве логического устройства для упрощения сетевой топологии

8.2.1. Единое управление несколькими устройствами

8.2.1.1. Сценарий

Когда несколько физических устройств образуют систему VSU, физические устройства можно рассматривать как логическое устройство. Все конфигурации управляются на глобальном Master-устройстве.

Как показано на Рисунке 8-3, четыре устройства (пронумерованные 1, 2, 3 и 4 слева направо) образуют систему VSU. Устройство 1 является глобальным Master-устройством, устройство 2 — глобальным slave-устройством, а устройства 3 и 4 — глобальными устройствами-кандидатами.

- Все устройства настраиваются просто на глобальном Master-устройстве.

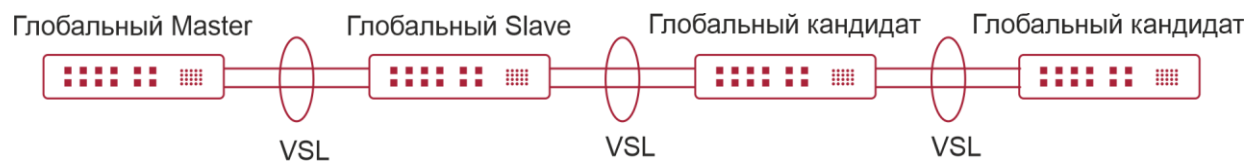


Рисунок 8-3.

Подробнее о VSL см. описание в разделе [Настройка VSL](#).

8.2.1.2. Развертывание

- Глобальное Master-устройство управляет всей системой VSU, запускает протоколы control-plane и участвует в пересылке данных.
- Глобальное slave-устройство участвует в пересылке данных, не запускает протоколы control-plane, работает как резервное и берет на себя работу глобального Master-устройства в случае сбоя.
- Глобальные устройства-кандидаты участвуют в пересылке данных и не запускают протоколы control-plane. Когда глобальное slave-устройство неисправно, глобальное устройство-кандидат может взять на себя работу глобального slave-устройства. В этом случае, когда глобальное Master- and slave-устройства неисправны, система VSU перезапустится.

8.2.2. Упрощение сетевой топологии

8.2.2.1. Сценарий

В традиционных сетях, как показано на Рисунке 8-4, необходимо добавить резервные устройства и линии для повышения надежности сети; однако также необходимо ввести множество алгоритмов для предотвращения петель, которые усложняют работу сети. В системе VSU все устройства рассматриваются как логическое устройство. Различные устройства дублируют друг друга, и нет необходимости вводить алгоритм предотвращения образования петель, что может упростить сеть.

- Два коммутатора агрегации образуют систему VSU. Нет необходимости настраивать алгоритм предотвращения образования петель. Два коммутатора взаимно резервированы.
- Коммутатор доступа подключается к коммутаторам агрегации через uplink AP.
- Когда коммутатор в системе VSU неисправен, другой канал все еще работает.

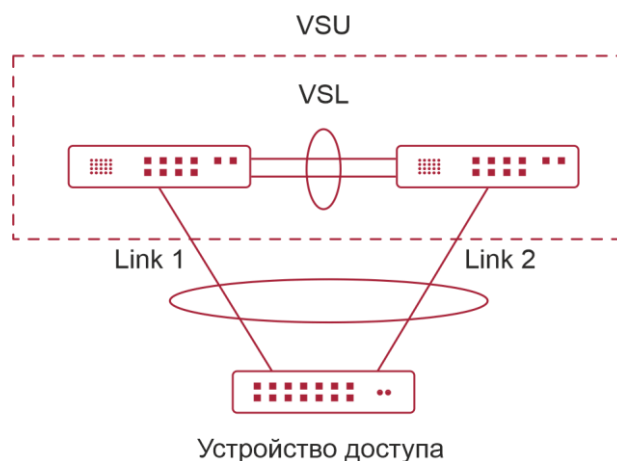


Рисунок 8-4.

8.2.2.2. Развертывание

- Глобальное Master-устройство управляет всей системой VSU, запускает протоколы control-plane и участвует в пересылке данных.
- Глобальное slave-устройство участвует в пересылке данных, не использует протоколы control-plane, работает как резервное и берет на себя работу глобального Master-устройства, когда глобальное Master-устройство неисправно.
- Коммутатор доступа ориентирован на пользователей и разрешает доступ с устройств пользователей.

8.3. Функции

8.3.1.1. Базовые концепты

Система VSU

Система VSU представляет собой единый логический объект, состоящий из двух или нескольких устройств в традиционной сетевой архитектуре. Например, система VSU уровня агрегации, показанная на следующем рисунке, может рассматриваться как единое устройство, взаимодействующее с уровнем ядра и уровнем доступа.

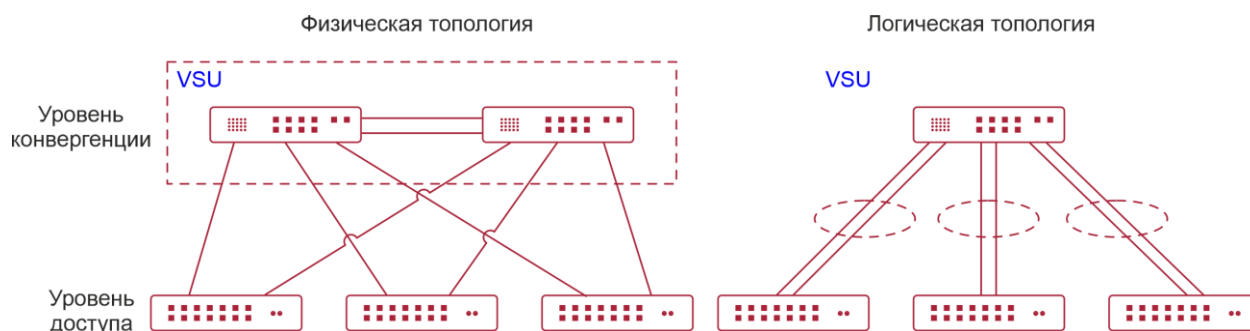


Рисунок 8-5. Уровень агрегации VSU

В приведенной выше сетевой структуре VSU устройства-члены образуют логический объект через внутренние каналы, а устройства уровня доступа подключаются к VSU через



агрегированные каналы. Таким образом, между уровнями доступа и агрегации нет петли уровня 2.

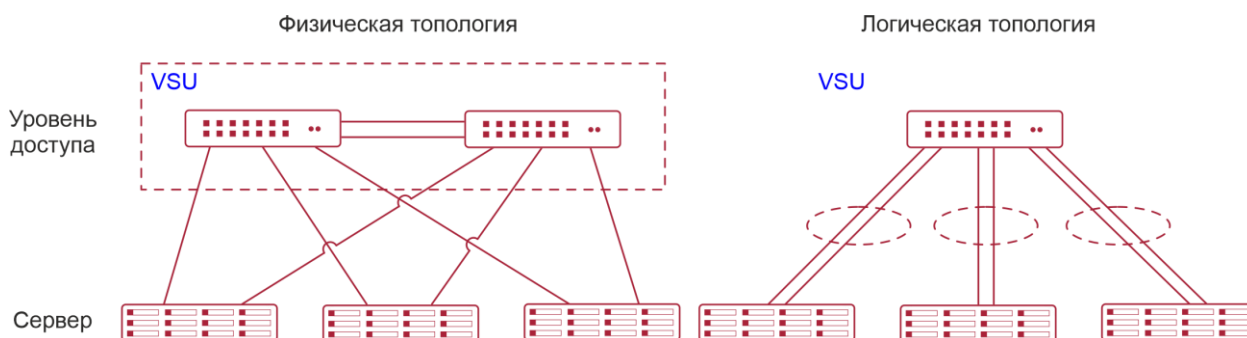


Рисунок 8-6. Уровень доступа VSU

Помимо устройств уровня ядра и уровня агрегации, устройства уровня доступа также могут формировать систему VSU. Сервер, которому требуется высокая доступность, может использовать несколько сетевых карт для формирования агрегированного порта (AP) для подключения устройств уровня доступа. Поскольку точка доступа может подключаться только к одному и тому же устройству доступа, возрастает риск отказа одного устройства. В этом случае для решения проблемы можно использовать VSU. В режиме VSU сервер использует несколько сетевых карт и привязывает их к AP для подключения различных устройств-членов одной и той же группы VSU. Таким образом можно предотвратить отказ одной точки и прерывание сети, вызванное отказом одного канала.

Идентификатор домена VSU

Домен VSU имеет только один идентификатор. Только устройства с одинаковыми идентификаторами домена могут формировать систему VSU.

Идентификатор устройства участника

Каждое устройство-участник в системе VSU имеет уникальный идентификатор, а именно идентификатор коммутатора. Идентификаторы коммутатора можно использовать для управления устройствами или настройки интерфейсов на устройствах-участниках. При добавлении устройства в систему VSU необходимо настроить идентификатор устройства и убедиться, что идентификатор уникален в той же системе VSU. В случае конфликта идентификаторов система VSU резервирует одно устройство в соответствии с приоритетом.

Роль устройства участника

Система VSU состоит из нескольких устройств. При установке системы VSU необходимо выбрать глобальное Master-устройство и глобальное slave-устройство. Все остальные устройства являются глобальными устройствами-кандидатами. Глобальное Master-устройство выбирается из нескольких устройств на основе протокола выбора. Все остальные устройства являются глобальными slave-устройствами в режиме горячего резерва 1:N. Когда поддерживается режим горячего резерва 1:1, одно устройство является глобальным Master-устройством, одно устройство является глобальным slave-устройством, а все остальные устройства являются глобальными устройствами-кандидатами.

Глобальное Master-устройство отвечает за управление всей системой VSU, выполнение протоколов control-plane и участие в пересылке данных. Другие устройства, включая глобальные slave-устройства и устройства-кандидаты, участвуют в пересылке данных, но



не используют протоколы control-plane. Все полученные потоки данных control-plane направляются на глобальное Master-устройство для обработки.

Глобальное slave-устройство также получает статусы глобального Master-устройства в режиме реального времени и обеспечивает избыточность 1:1 или 1:N с глобальным Master-устройством. Если глобальное Master-устройство выйдет из строя, глобальное slave-устройство возьмет на себя функции Master-устройства и будет управлять всей системой VSU.

Ниже приведен метод выбора Master-устройства системы VSU:

1. Правила выбора Master-устройства системы VSU включают (переходит к следующему правилу, если предыдущее правило не помогает при выборе Master-устройства): а) Выбор устройства с высшим приоритетом из запущенных в данный момент в качестве Master-устройства (все устройства не являются Master-устройствами во время запуска). б) Выбор устройства с наименьшим номером (No.) устройства. в) Выбор устройства с наименьшим MAC-адресом в качестве Master-устройства.
2. В режиме «горячего» резерва 1:N выбирается устройство с наиболее похожей конфигурацией с Master-устройством в качестве slave-устройства, чтобы предотвратить использование двух активных устройств. Порядок выбора следующий: ближайший/наивысший приоритет/наименьший MAC-адрес.
3. Система VSU поддерживает «горячее» добавление вспомогательного устройства. Даже если «горячо» добавленное устройство имеет более высокий приоритет, чем Master-устройство, система VSU не выполняет переключение активный/резервный.
4. Порядок запуска устройства-участника может повлиять на выбор Master-устройства. Устройство-член может не присоединиться к системе VSU, потому что оно запускается слишком медленно. В этом случае устройство будет добавлено в систему VSU в «горячем» режиме. Даже если устройство имеет более высокий приоритет, чем Master-устройство, система VSU не выполняет переключение активный/резервный.

8.3.1.2. Обзор

Особенность	Описание
<u>Канал виртуальной коммутации (VSL)</u>	В системе VSU для подключения всех устройств используется виртуальный канал
<u>Топология</u>	Описывает внутреннюю топологию системы VSU
<u>Dual-Active Detection (DAD)</u>	Предотвращает сосуществование двух Master-коммутаторов в домене VSU
<u>Управление системой</u>	Описывает возможные соединения между внешними устройствами и устройствами VSU
<u>Определение нахождения устройства быстрым миганием (Quick Blinking Location)</u>	Управляет устройствами в системе VSU



8.3.2. Канал виртуальной коммутации (VSL)

8.3.2.1. Принцип работы

VSL

Система VSU представляет собой сетевой объект, состоящий из нескольких устройств. Эти устройства должны обмениваться управляющей информацией и частью потоков данных. VSL — это специальный канал, используемый для передачи управляющей информации и потоков данных между устройствами системы VSU. Например, VSL может быть установлен между двумя устройствами через интерфейсы 10 Gigabit Ethernet. На Рисунке 8-7 показано положение VSL в системе VSU.

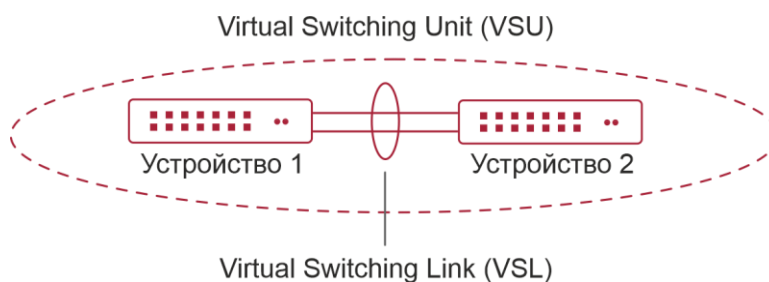


Рисунок 8-7. VSL

VSL существует в виде групп AP. Потоки данных, передаваемые через VSL, распределяют нагрузку между участниками порта агрегации в соответствии с алгоритмом балансировки трафика.

VSL-трафик

Потоки управления, передаваемые через VSL между устройствами, включают:

1. Пакеты протоколов, полученные устройствами-участниками: эти пакеты протоколов необходимо пересылать через VSL на глобальное Master-устройство для обработки.
2. Пакеты протокола, обрабатываемые глобальным мастер-устройством: эти пакеты протокола должны быть перенаправлены через VSL на интерфейсы других устройств-участников, а затем отправлены на реер-устройства с помощью этих интерфейсов.

Потоки данных, передаваемые через VSL между устройствами, включают:

1. Поток данных, поступающие во VLAN.
2. Потоки данных, которые необходимо пересылать между устройствами и передавать через VSL.

Кроме того, внутренние пакеты управления системы VSU также передаются через VSL. Пакеты управления включают в себя информацию о протоколе, переключаемую оперативным резервированием, и информацию о конфигурации, доставляемую хостом другим устройствам-участникам.

ПРИМЕЧАНИЕ: с точки зрения функции анализатора коммутируемых портов (SPAN) интерфейс, связанный с VSL, не может рассматриваться как исходный порт или порт назначения SPAN.



Сбой VSL

Если определенный канал-участник, подключенный к группе AP VSL, не работает, VSU автоматически изменит конфигурацию порта агрегации VSL, чтобы предотвратить передачу трафика через неисправный канал-участник.

Если все каналы связи отключены от группы AP VSL, топология VSU изменится. Если исходная топология VSU является топологией кольца, кольцо преобразуется в линию. Для получения дополнительной информации см. раздел [Топология](#).

Обнаружение ошибок на интерфейсе VSL

Когда на интерфейсе VSL обнаруживается большое количество последовательных ошибок, интерфейс должен быть отключен и переключен на другой интерфейс VSL. Метод обнаружения следующий:

Если на интерфейсе VSL обнаружены ошибочные кадры, выполнется исправление ошибочных кадров. По умолчанию система определяет интерфейс VSL каждые 5 секунд. Если количество кадров с ошибками больше, чем значение *num* по сравнению с обнаруженным в последний раз, предполагается, что кадры с ошибками обнаружены один раз. Если ошибочные кадры обнаруживаются последовательно в течение значения раз, предполагается, что интерфейс неисправен. Если при обнаружении ошибочных кадров доступно несколько каналов VSL, VSL будет переключен. Последний VSL не будет переключаться, чтобы предотвратить разделение топологии.

Различные пользовательские сценарии имеют разные требования к *num* и *times*. Значение по умолчанию для *num* равно 3, а для *times* — 10. Если у пользователей есть строгие требования к сценариям, выберите меньшие значения для *num* и *times*; если наоборот, выберите большие значения.

8.3.3. Топология

Система VSU поддерживает линейную и кольцевую топологию. Устройства подключаются через VSL для формирования линии, которая называется линейной топологией.

8.3.3.1. Принцип работы

Топология

Линейная топология проста. Она использует очень мало портов и кабелей. Два устройства соединены только одним каналом связи. Следовательно, VSL имеет низкую надежность.

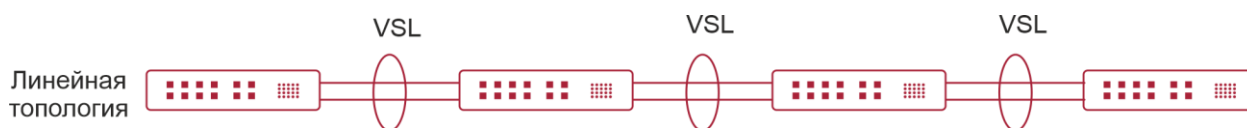


Рисунок 8-8. Линейная топология

За исключением линейной топологии, устройства также могут образовывать кольцевую топологию, как показано на Рисунке 8-9. В кольцевой топологии два канала связи между устройствами могут дублировать друг друга и выполнять резервирование каналов для повышения надежности системы VSU.

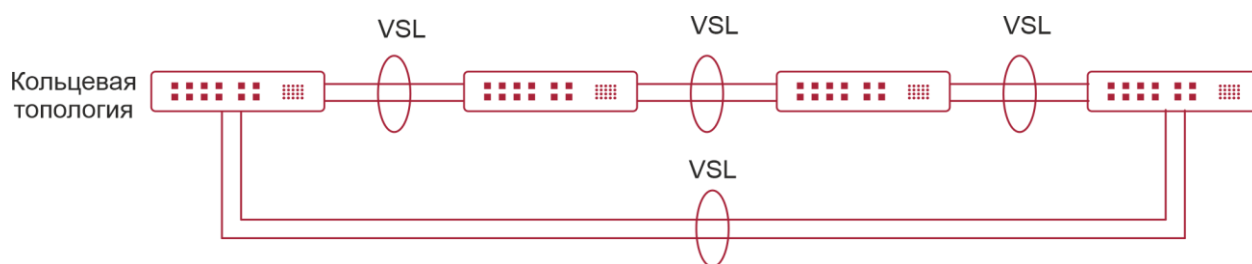


Рисунок 8-9. Кольцевая топология

ПРИМЕЧАНИЕ: рекомендуется выбрать кольцевую топологию для системы VSU, чтобы на нормальную работу всей системы VSU не повлияло ни одно неисправное устройство или VSL.

ПРИМЕЧАНИЕ: помимо выбора сети с кольцевой топологией, рекомендуется настроить несколько VSL для каждого участника VSL, чтобы повысить надежность одного VSL. Рекомендуется как минимум два канала, и можно настроить максимум четыре канала. Разумная конфигурация включает более двух VSL, пересекающих разные карты.

Конвергенция топологии

Перед созданием VSU устройства-участники должны обнаружить соседей с помощью протоколов обнаружения топологии и проверить устройства в системе VSU, чтобы подтвердить диапазон домена управления. Затем выбирается глобальное Master-устройство для управления всей системой VSU, а глобальное slave-устройство выбирается для резервного копирования Master-устройства. Тогда вся топология VSU сходится. Поскольку время запуска различается для разных устройств, время первой сходимости топологии также различно.

Преобразование кольцевой и линейной топологии

В кольцевой топологии, если канал VSL отключен, кольцевая топология преобразуется в линейную топологию. Вся система VSU по-прежнему будет нормально работать без отключения сети. Чтобы предотвратить отказ других каналов и узлов VSL, рекомендуется определить местонахождение сбоев VSL и восстановить доступность VSL. После восстановления канала VSL топология линии преобразуется в топологию кольца.

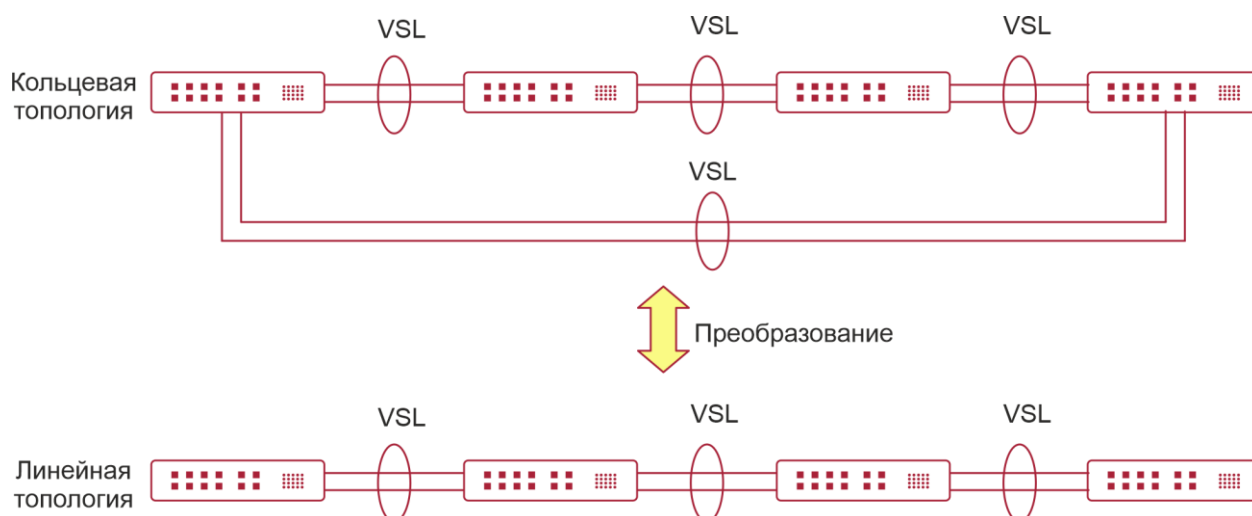


Рисунок 8-10. Тип преобразования соединения. Кольцо в линию и линия в кольцо



Разделение топологии

В линейной топологии, если канал VSL отключен, линейная топология будет разделена, как показано на Рисунке 8-11. Группа VSU разделена на две группы. В этом случае в сети могут существовать два устройства с абсолютно одинаковыми конфигурациями, что приведет к ненормальной работе сети. Следовательно, для решения проблемы разделения топологии необходимо развернуть функцию многоактивного обнаружения (MAD).

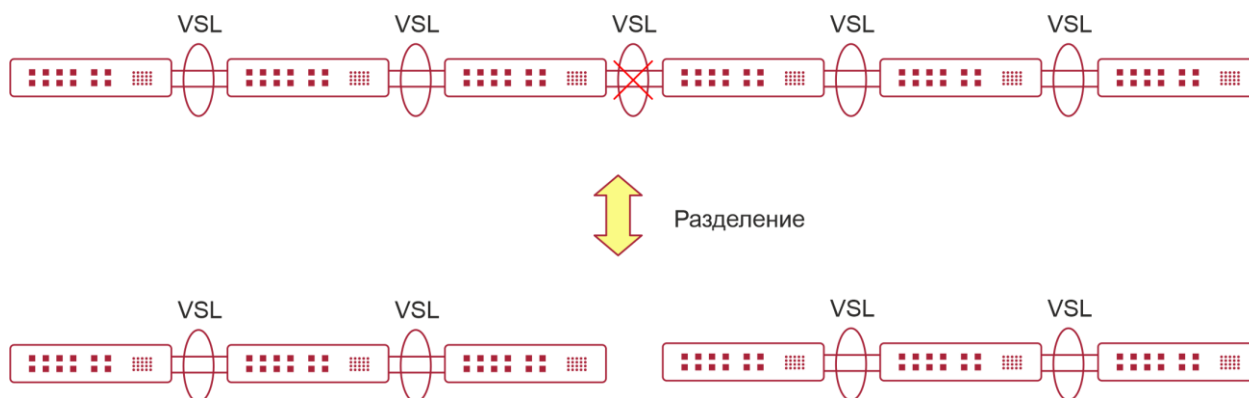


Рисунок 8-11. Разделение топологии

Объединение топологии

Если две группы VSU соединены через канал VSL, топология линии будет объединена. Во время объединения топологий перезапустите одну группу VSU, а затем добавьте другую группу VSU в «горячем» режиме.

Принцип комбинирования топологий: минимизация влияния на сервисы при комбинировании топологий. Правила следующие (Выбор происходит последовательно, по одному из правил. Если по первому правилу не получается выбрать оптимальную топологию, продолжается выбор по следующему пункту):

- Используйте приоритет устройства в качестве первого критерия для оценки комбинирования топологии. Зарезервируйте группу VSU, содержащую устройство с наивысшим приоритетом.
- Если предыдущий пункт не может помочь принять решение, выберите группу VSU с меньшим идентификатором коммутатора (один из двух глобальных Master-коммутаторов).
- Если предыдущий пункт не может помочь принять решение, зарезервируйте группу VSU с меньшим MAC-адресом (адрес глобального Master-коммутатора).

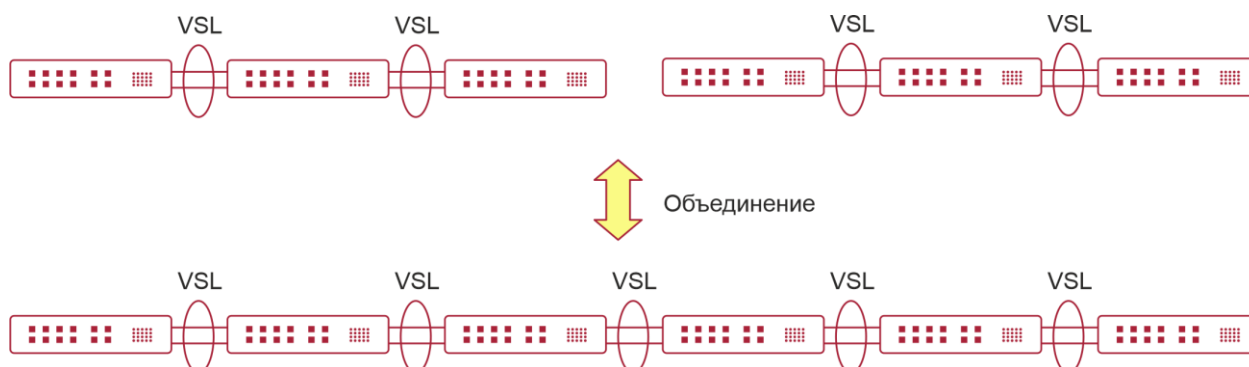


Рисунок 8-12. Объединение топологий

ПРИМЕЧАНИЕ: при объединении топологии двух групп VSU необходимо выбрать одну из двух групп VSU. Группа VSU, которая не прошла отбор, автоматически перезапустится и добавится в другую группу VSU.

8.3.4. Dual-Active Detection (DAD)

8.3.4.1. Принцип работы

Когда VSL отключен, slave-устройство переключается в режим Master. Если исходное Master-устройство все еще работает, ряд проблем, включая конфликт IP-адресов в локальной сети, будет вызван из-за наличия двух Master-устройств и их полностью одинаковых конфигураций. В этом случае система VSU должна обнаружить два устройства и принять меры по восстановлению. Система VSU предоставляет два следующих метода выполнения MAD:

- Обнаружение двунаправленной пересылки (BFD)
- Обнаружение на основе AP

Правила MAD

1. Выберите группу VSU с наивысшим приоритетом.
2. Если предыдущий пункт не может помочь принять решение, выберите группу VSU с большим количеством физических устройств.
3. Если предыдущий пункт не может помочь принять решение, выберите группу VSU с более высоким уровнем Состояния (Состояние: общая пропускная способность всех физических интерфейсов (кроме интерфейсов управления и VSL) в состоянии UP в топологии).
4. Если предыдущий пункт не может помочь принять решение, выбирается группа VSU с меньшим идентификатором коммутатора (один из двух глобальных Master-коммутаторов).
5. Если предыдущий пункт не может помочь принять решение, выбирается группа VSU с меньшим MAC-адресом (адрес двух глобальных Master-коммутаторов).
6. Если предыдущий пункт не может помочь принять решение, выбирается группа VSU с большим временем запуска (время запуска глобальных Master-коммутаторов).

BFD

Система VSU поддерживает BFD для обнаружения нескольких Master-устройств. Рисунок 8-13 показывает топологию. Канал добавлен для двух устройств по краям специально для MAD. Когда канал VSL отключен между глобальным Master- и

slave-устройствами, два Master-устройства существуют одновременно. Если функция BFD установлена, два Master-устройства будут отправлять пакеты BFD друг другу через канал BFD. Таким образом, в текущей системе обнаруживаются одни и те же устройства. Наконец, выключите систему VSU Master-устройства в соответствии с некоторыми правилами и войдите в состояние восстановления, чтобы избежать сбоев в работе сети.

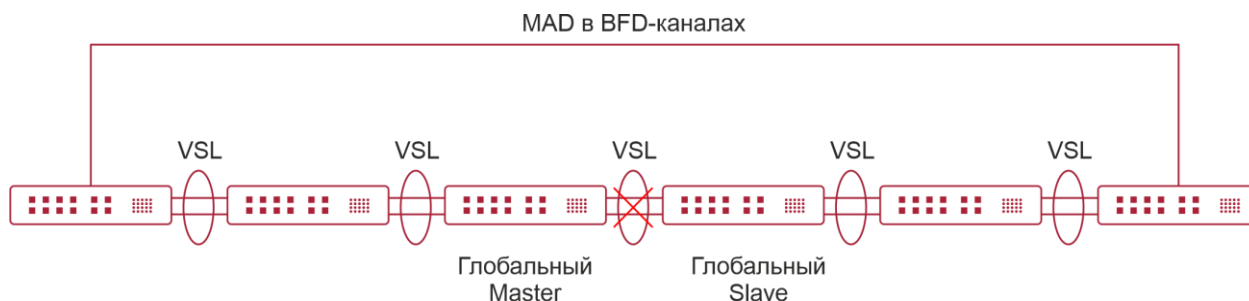


Рисунок 8-13. BFD

ПРИМЕЧАНИЕ: при наличии пары каналов BFD рекомендуется развертывать каналы обнаружения на двух концах топологии.

ПРИМЕЧАНИЕ: вам необходимо использовать расширение BFD, и вы не можете настроить порт Dual-Active Detection с помощью существующих конфигураций и команд BFD.

MAD

Система VSU также поддерживает механизм dual-active detection MAD. Рисунок 8-14 показывает топологию. И система VSU, и upstream-устройство должны поддерживать функцию MAD. Когда канал VSL отключен, два Master-устройства существуют одновременно. Два Master-устройства соответственно отправляют пакеты MAD на порты-участники MAD-AP, а затем пакеты MAD пересылаются друг другу через upstream-устройство. Как показано на Рисунке 8-14, MAD-AP имеет четыре порта-участника. Каждый порт-участник подключен к другому устройству системы VSU. Когда происходит разделение топологии, все четыре порта-участника отправляют и получают пакеты MAD. Таким образом, в текущей системе обнаруживаются одни и те же устройства. Наконец, выключите систему VSU Master-устройства в соответствии с некоторыми правилами и войдите в состояние восстановления, чтобы избежать сбоев в работе сети.

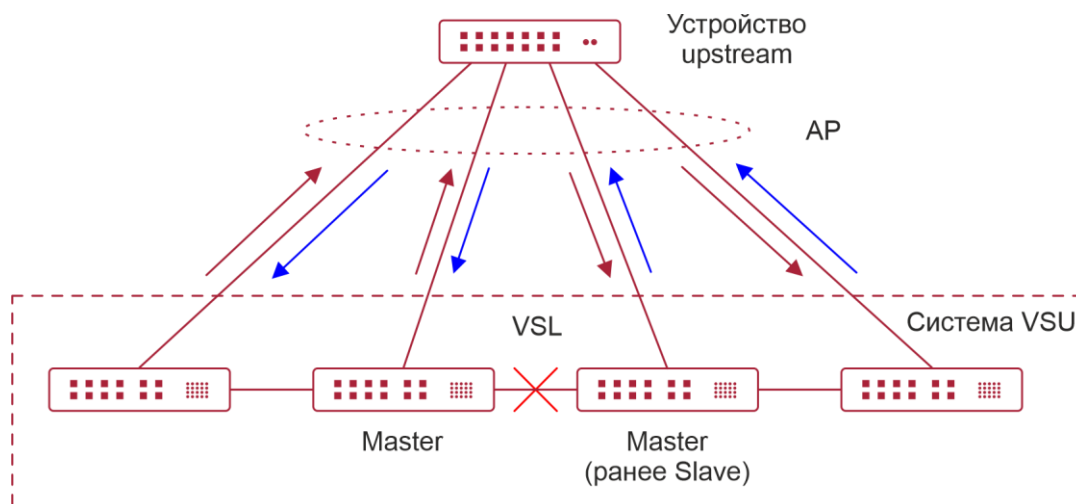


Рисунок 8-14. MAD на основе upstream- и downstream-устройств

ПРИМЕЧАНИЕ: в приведенной выше топологии восходящее устройство должно быть устройством QTECH и поддерживать функцию пересылки пакетов MAD.

8.3.5. Переадресация трафика VSU

8.3.5.1. Принцип работы

Группа между устройств AP

AP связывает несколько физических каналов вместе, образуя логический канал. Система VSU поддерживает точку доступа на устройствах-участниках.

Как показано на Рисунке 8-15, два устройства образуют группу VSU. Устройство внешнего доступа Switch A подключено к VSU в виде AP. С точки зрения коммутатора A нет никакой разницы между AP на Рисунке 8-15 и общей группой AP.

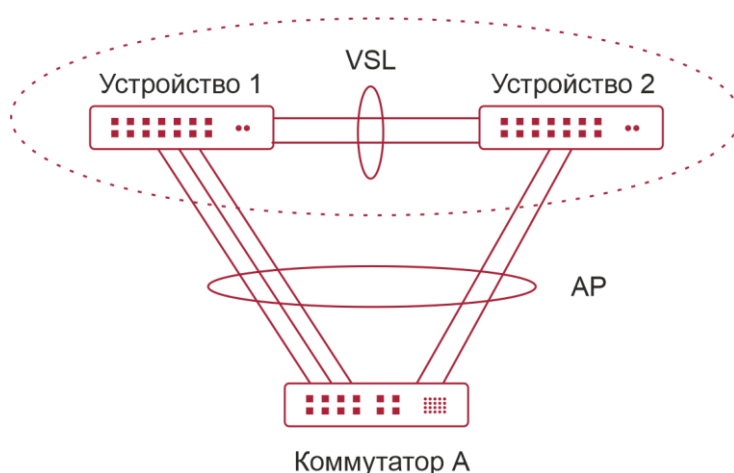


Рисунок 8-15. Порт агрегации между устройствами



Поиск неисправностей

Рекомендуется настроить AP для нескольких устройств с физической связью между периферийным устройством и каждым устройством VSU. С одной стороны, полоса пропускания VSL может быть зарезервирована (укажите приоритет AP, принадлежащий тому же шасси, что и выход, для передачи трафика между AP и между шасси и предотвращения передачи ненужного трафика по каналу VSL). С другой стороны, надежность сети может быть повышена (если какое-то шасси неисправно, порты-участники обычных устройств могут работать нормально).

В следующих разделах описаны возможные неисправности AP, работающего на разных устройствах, и их последствия.

- Ошибка одиночного канала

Если один канал AP, работающий на разных устройствах, неисправен, а другие каналы работают нормально, AP на разных устройствах перераспределит трафик для оставшихся нормальных каналов.

- Сбой соединения всех портов-участников AP между устройствами на глобальном Master-устройстве

Если каналы всех портов-участников AP между устройствами на глобальном Master-устройстве не работают, только порты-участники других устройств-участников продолжают нормально работать. Что касается потока данных, передаваемого через AP в систему VSU, если выход для пересылки потока данных находится на глобальном Master-устройстве, система будет пересылать поток данных на соответствующий выход на глобальном Master-устройстве по каналу VSL.

Протоколы control-plane все еще работают на глобальном Master-устройстве. Следовательно, пакеты протокола, поступающие в систему VSU, необходимо пересылать на глобальное Master-устройство по каналу VSL для вычисления протокола.

- Сбой всех каналов других устройств-участников

Если все каналы AP на нескольких устройствах и одного устройства A не работают, только порты-участники других устройств-участников продолжают нормально работать. Что касается потока данных, передаваемого через AP в систему VSU, если выход для пересылки потока данных находится на устройстве-участнике A, система будет пересылать поток данных на соответствующий выход на устройстве-участнике A через VSL.

- Глобальная ошибка Master-устройства

Если глобальное Master-устройство неисправно, выполняется «горячее» резервное переключение, чтобы переключить исходное slave-устройство на Master-устройство. При этом порты участников на других устройствах-участниках продолжают работать. Сбой соединения обнаружен на реер-устройстве, подключенном к VSU через этот AP. Следовательно, алгоритм балансировки трафика необходимо скорректировать, чтобы распределять поток данных по обычным каналам.

- Ошибка устройства-участника

Если устройство-участник неисправно, канал участника AP, подключенный к этому устройству-участнику, отключается. Однако другие каналы участников по-прежнему работают нормально. Сбой соединения обнаружен на реер-устройстве, подключенном к VSU через этот AP. Следовательно, алгоритм балансировки трафика необходимо скорректировать, чтобы выделить пути пересылки потоков данных для обычных каналов.

Балансировка трафика

В системе VSU трафик может иметь несколько выходов. AP и ECMP имеют собственные алгоритмы балансировки трафика, например, с использованием MAC-адресов назначения



или источника. Дополнительные сведения см. в разделе Ethernet Switching/Настройка AP. В этом руководстве по настройке можно подробно настроить локальную переадресацию в первую очередь (LFF). Пакеты, полученные устройством, сначала пересылаются на это устройство. Таким образом, пакеты могут пересылаться на другие устройства без использования VSL.

8.3.6. Управление системой

8.3.6.1. Принцип работы

Доступ к консоли

Консоль Master-устройства системы VSU одновременно управляет несколькими устройствами в системе. Консоли подчиненных устройств и устройств-кандидатов не поддерживают ввод командной строки. Однако вы можете настроить систему VSU на Master-устройстве для указанного устройства-участника и войти в консоль Master-устройства через последовательный порт slave-устройства. Сеанс можно использовать для перенаправления на Master-консоль устройства.

Присвоение имени слоту

С точки зрения устройства шасси, в режиме VSU слот именуется номером устройства (Switch ID). Поэтому номер слота из одномерного превращается в двумерный. Например, кабельный зажим 1/1 указывает на слот с номером 1 слота 1 на устройстве-участнике.

Присвоение имени интерфейсу

В режиме работы VSU номер слота может встречаться в нескольких устройствах. Поэтому интерфейс именуется по номеру устройства (Switch ID).

Например, интерфейс `gigabitEthernet 1/0/1` указывает на гигабитный порт 1 в слоте 0 устройства, идентификатор которого равен 1; `interface gigabitEthernet 2/0/2` указывает на гигабитный порт 2 в слоте 0 устройства, идентификатор которого равен 2.

Доступ к файловой системе

В рабочем режиме VSU вы можете получить доступ к файловой системе на других устройствах-участниках с Master-устройства. Подробный метод доступа такой же, как и для локальной файловой системы. Единственное отличие состоит в том, что используются разные префиксы URL.

Обновление системы

Как правило, система VSU требует согласованности версий основных номеров версий программ устройств-участников. Однако устройств-участников так много, что выполнение обновления по одному в автономном режиме занимает слишком много времени и энергии, а также легко допустить ошибку. Коммутаторы QTECH предоставляют непревзойденное решение для обновления системы, которое поможет вам выполнить обновление системы, используя два следующих метода:

- При создании системы VSU: система автоматически выровняет номера версий основных программ всех устройств-участников. Как только основные версии программы обнаружат несоответствие, будет выбрана основная программа Master-устройства для синхронизации со всеми устройствами-участниками.
- После установки системы VSU: основная версия программы будет автоматически синхронизирована со всеми устройствами-участниками с помощью файла, загружаемого по TFTP.

SYSLOG (системный лог)

Все устройства-участники системы VSU могут отображать SYSLOG. SYSLOG, сгенерированный Master-устройством, отображается на консоли Master-устройства в том

же формате, что и в автономном режиме. SYSLOG, созданный другими устройствами-участниками, также отображается на консоли Master-устройства, но формат сообщения отличается от формата сообщения в автономном режиме, поскольку добавляется информация о номере устройства.

Например, информация SYSLOG, сгенерированная в автономном режиме, выглядит так: "%VSU-5-DTM_TOPO_CVG:Node discovery done. Topology converged.". Информация SYSLOG, сгенерированная устройством-участником с номером 3: "%VSU-5-DTM_TOPO_CVG:(3) Node discovery done. Topology converged.".

8.3.7. Определение нахождения устройства быстрым миганием (Quick Blinking Location)

В сетевом кабельном пространстве аппаратная, в которой расположены коммутаторы, и операционная консоль часто находятся в разных местах. Если в среде много устройств, сетевые администраторы не могут легко определить местонахождение определенных устройств.

Определение нахождения устройства быстрым миганием позволяет сетевым администраторам находить устройства с помощью быстрого мигания. Включив эту функцию для устройства на консоли, вы можете легко найти соответствующее устройство в аппаратной.

ПРИМЕЧАНИЕ: если включено Определение нахождения устройства быстрым миганием, светодиодный индикатор состояния не может отображать исходное состояние до тех пор, пока не будет отключено Определение нахождения устройства быстрым миганием.

8.3.8. Восстановление устройства в режиме восстановления

По умолчанию устройство в режиме восстановления автоматически перезагружается и снова добавляется в топологию VSU после восстановления канала VSL.

Команда включения автоматического перезапуска без восстановления предназначена для отключения функции автоматического перезапуска и восстановления в режиме восстановления. После восстановления канала VSL администратор может снова включить функцию автоматического перезапуска и восстановления в режиме восстановления или вручную ввести команду перезапуска для устройств в режиме восстановления.

ПРИМЕЧАНИЕ: после отключения функции автоматического перезапуска и восстановления в режиме восстановления устройство остается в режиме восстановления до тех пор, пока администратор снова не включит эту функцию или не перезагрузит устройство вручную.

8.3.9. Автоматическое восстановление без перезагрузки для устройства в режиме восстановления при неисправности Master-устройства

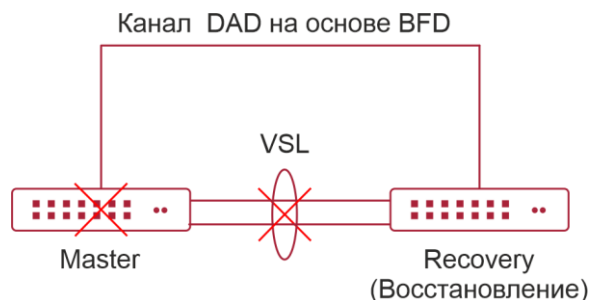




Рисунок 8-16. Автоматическое восстановление без перезагрузки в режиме восстановления

По умолчанию устройство в режиме восстановления не выполняет никаких операций, когда Master-устройство выходит из строя. Устройство остается в режиме восстановления до тех пор, пока не восстановится канал VSL.

Администратору предоставляется команда конфигурации для включения функции автоматического восстановления без перезагрузки в режиме восстановления. В этом случае устройство в режиме восстановления может автоматически стать Master-устройством без перезапуска и взять на себя функции исходного Master-устройства.

ПРИМЕЧАНИЕ: после включения функции автоматического восстановления без перезагрузки в режиме восстановления устройство в режиме восстановления становится Master-устройством при обнаружении неисправности Master-устройства.

8.4. Конфигурация

Конфигурация	Конфигурация и команда	
<u>Настройка VSU в автономном режиме</u>	(Обязательно) Используется для настройки VSU в автономном режиме	
	switch domain	virtual Настраивает идентификатор домена
	switch	Настраивает идентификатор коммутатора
	switch priority	Настраивает приоритет коммутатора
	vsl-port	Вход в режим конфигурации интерфейса VSL
<u>Настройка VSU в автономном режиме</u>	port-member interface	Настраивает интерфейс участника VSL
	switch mode	convert virtual Изменяет автономный режим на режим VSU
	(Опционально) Используется для настройки атрибутов устройства в режиме VSU	
	switch description	Настраивает описание устройства
	switch crc	Настраивает проверку ошибок на VSL



Конфигурация		Конфигурация и команда	
<u>Настройка VSU в режиме VSU</u>	<u>Настройка атрибутов</u>	(Опционально) Используется для настройки атрибутов устройства в режиме VSU	
		switch domain	Изменяет идентификатор домена
		switch renumber	Изменяет идентификатор коммутатора
		switch description	Настраивает описание устройства
<u>Настройка VSU в режиме VSU</u>	<u>Настройка атрибутов</u>	switch crc	Настраивает проверку ошибок на VSL
		(Опционально) Используется для настройки VSL	
		vsl-port	Вход в режим конфигурации интерфейса VSL
	<u>Настройка VSL</u>	port-member interface	Настраивает интерфейс участника VSL
		(Обязательно) Используется для настройки DAD	
		dual-active detection	Настраивает DAD
		dual-active interface bfd	Настраивает интерфейс BFD DAD
		dual-active interface	Настраивает AP как интерфейс DAD
		dual-active exclude interface	Настраивает исключенный интерфейс
	<u>Настройка балансировки трафика</u>	(Опционально) Используется для настройки балансировки трафика в режиме VSU	
		switch virtual aggregateport-lff enable	Настраивает режим AP LFF



Конфигурация		Конфигурация и команда	
		switch virtual ecmp-lff enable	Настраивает режим ECMP LFF
	<u>Настройка восстановления устройства в режиме восстановления</u>	(Опционально) Используется для настройки восстановления устройства в режиме восстановления	
		no recovery auto-restart enable	Отключает функцию автоматического перезапуска и восстановления в режиме восстановления
<u>Настройка VSU в режиме VSU</u>	<u>Настройка автоматического восстановления без перезагрузки в режиме восстановления</u>	(Опционально) Используется для настройки функции автоматического восстановления без перезагрузки в режиме восстановления	
		dual-active auto-recovery enable	Включает функцию автоматического восстановления без перезагрузки в режиме восстановления
	<u>Изменение режима VSU на автономный режим</u>	(Опционально) Используется для изменения режима VSU на автономный режим	
		switch convert mode standalone	Изменяет режим VSU на автономный режим
<u>Настройка Определения нахождения устройства быстрым миганием</u>		(Опционально) Используется для быстрого поиска устройства	
		led-blink	Включает Определение устройства быстрым миганием
<u>Настройка интерфейса MGMT</u>		(Опционально) Он используется для настройки устройства для создания одного интерфейса MGMT для каждого шасси или создания только одного интерфейса MGMT в режиме VSU	
		mgmt_mode unique	Настраивает систему для создания только одного интерфейса MGMT в режиме VSU



Конфигурация	Конфигурация и команда	
	no mgmt_mode	Настраивает систему для создания одного интерфейса MGMT для каждого шасси в режиме VSU

8.4.1. Настройка VSU в автономном режиме

8.4.1.1. Эффект конфигурации

Запустите коммутатор в автономном режиме, чтобы установить соответствующие параметры VSU для создания системы VSU.

8.4.1.2. Шаги настройки

Настройка атрибутов VSU

- По умолчанию коммутатор запускается в автономном режиме. Вам необходимо установить один и тот же идентификатор домена на двух шасси установленной системы VSU. Идентификатор домена должен быть уникальным в пределах локальной сети (LAN). Кроме того, вам необходимо установить идентификатор каждого шасси в VSU.
- Запустите команду **switch virtual domain domain_id**, чтобы настроить идентификатор домена. Эта команда является обязательной.
- Запустите команду **switch switch_id**, чтобы настроить идентификатор устройства в VSU. Эта команда является обязательной. Для устройств с одинаковыми приоритетами в системе VSU в качестве глобального Master-устройства выбирается устройство с наименьшим идентификатором устройства.
- Запустите команду **switch switch_id priority priority_num**, чтобы настроить приоритет устройства. Эта команда является обязательной.
- Диапазон значений от 1 до 255. Чем больше значение, тем выше приоритет.
- Запустите команду **switch switch_id description switch1**, чтобы настроить псевдоним устройства. Эта команда опциональна. Имя по умолчанию — QTECH. Для легкой идентификации устройств в сетевой среде этот пункт можно выбрать для установки псевдонима устройства.
- Допускается не более 32 символов.

Команда	switch virtual domain number
Описание параметров	<i>number</i> : указывает идентификатор домена VSU
По умолчанию	Идентификатор домена по умолчанию — 100
Командный режим	Режим конфигурации домена



Руководство по использованию	Только два устройства с одинаковым идентификатором домена могут формировать VSU. Идентификатор домена должен быть уникальным в локальной сети
------------------------------	---

Команда	switch <i>switch_id</i>
Описание параметров	<i>switch_id</i> : указывает идентификатор коммутатора в системе VSU. Значение варьируется в зависимости от продуктов
По умолчанию	Идентификатор устройства по умолчанию — 1
Командный режим	Режим конфигурации домена
Руководство по использованию	Идентификатор устройства идентифицирует каждого участника виртуального устройства. В режиме VSU формат имени интерфейса меняется на «коммутатор/слот/порт» с «слот/порт», где «коммутатор» — это идентификатор устройства. Если одно из шасси активно или если роль только что запущенного шасси неясна и оба имеют одинаковый приоритет, то шасси с меньшим идентификатором выбирается в качестве активного. Эту команду можно использовать только для изменения идентификатора устройства в автономном режиме. В режиме VSU запустите команду switch renumber, чтобы изменить идентификатор устройства. Измененный идентификатор устройства вступает в силу только после перезагрузки устройства, независимо от того, в автономном режиме или в режиме VSU

Команда	switch <i>switch_id</i> priority <i>priority_num</i>
Описание параметров	<i>switch_id</i>: указывает идентификатор коммутатора, для которого необходимо настроить приоритет. <i>priority_num</i>: указывает приоритет коммутатора в диапазоне от 1 до 255
По умолчанию	Приоритет устройства по умолчанию равен 100
Командный режим	Режим конфигурации домена
Руководство по использованию	Большее значение означает более высокий приоритет. Устройство с наивысшим приоритетом выбирается в качестве Master-устройства.



	Вы можете запустить эту команду в автономном режиме или в режиме VSU. Измененный приоритет вступает в силу только после перезагрузки устройства. Эта команда не используется для изменения значения <i>switch_id</i>. В автономном режиме, если для параметра <i>switch_id</i> установлено значение 1, запуск команды switch 2 priority 200 не работает. Вы можете сначала установить для <i>switch_id</i> значение 2, а затем запустить команду switch 2 priority 200. В режиме VSU <i>switch_id</i> указывает идентификатор работающего в данный момент коммутатора. Если идентификатор не существует, конфигурация не вступает в силу
--	--

Команда	switch <i>switch_id</i> description <i>dev-name</i>
Описание параметров	<i>switch_id</i> : указывает идентификатор устройства. <i>dev-name</i> : указывает описание устройства, не более 32 символов
Командный режим	Режим конфигурации домена
Руководство по использованию	Эта команда настраивается на устройстве в автономном режиме или в режиме VSU и вступает в силу сразу после настройки

Настройка VSL

- Чтобы установить систему VSU, необходимо решить, какие порты настроены как порты-участники VSL.
- Запустите команду **vsl-port**, чтобы войти в режим конфигурации интерфейса VSL. Эта команда является обязательной.
- Запустите команду **port-member interface** *interface-name* [**copper** | **fiber**] для добавления интерфейса VSL. Эта команда является обязательной.

Когда устройство входит в режим конфигурации интерфейса VSL, интерфейс VSL можно настроить или удалить.

Команда	vsl-port
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Вы можете запустить эту команду в автономном режиме или в режиме VSU



Команда	port-member interface <i>interface-name</i> [copper fiber]
Описание параметров	<i>interface-name</i>: указывает двумерное имя интерфейса, например Tengigabitethernet 1/1 и Tengigabitethernet 1/3. copper: указывает атрибут медного интерфейса. fiber: указывает атрибут оптического интерфейса
Командный режим	Режим конфигурации интерфейса VSL
Руководство по использованию	Добавьте интерфейс участника канала VSL. <i>interface-name</i> указывает двумерное имя интерфейса в автономном режиме. Двумерный интерфейс может быть 10-гигабитным интерфейсом или гигабитным интерфейсом. (Гигабитный интерфейс может быть опто-медным интерфейсом. Если тип носителя не указан, по умолчанию используется гигабитный медный интерфейс.) Для опто-медного интерфейса необходимо указать атрибут оптического или медного интерфейса. Интерфейс VSL для устройства шасси должен быть 10-гигабитным интерфейсом. Вы можете запустить эту команду в режиме VSU или в автономном режиме. Команда может вступить в силу после сохранения конфигурации команды и перезапуска устройства, на котором находится интерфейс участника VSL

ПРИМЕЧАНИЕ: в автономном режиме конфигурации VSL не могут вступить в силу немедленно, пока устройство не перейдет в режим VSU и не перезапустится.

Настройка проверки ошибок

- Запустите команду **switch crc**, чтобы настроить проверку ошибок. Эта команда не является обязательной. Запустите эту команду, чтобы изменить метод по умолчанию для проверки кадров ошибок.
- Если на интерфейсе VSL обнаружены ошибочные кадры, выполняется исправление ошибочных кадров. По умолчанию система проверяет интерфейсы VSL каждые 5 секунд. Если количество кадров с ошибками больше 3 по сравнению с количеством, обнаруженным в последний раз, предполагается, что кадры с ошибками обнаружены один раз. Если ошибочные кадры обнаруживаются последовательно 10 раз, предполагается, что интерфейс неисправен. Если при обнаружении ошибочных кадров доступно несколько каналов VSL, VSL будет переключен. Последний VSL не будет переключаться, чтобы предотвратить разделение топологии.

Команда	switch crc errors <i>error_num</i> times <i>time_num</i>
Описание параметров	<i>error_num</i>: настраивает увеличение количества ошибочных кадров между двумя обнаружениями. Когда количество ошибочных кадров превышает прирост, предполагается, что ошибочные кадры обнаруживаются один раз.



	<i>time_num</i> : настраивает количество раз, после которого необходимо выполнить действие (действие может отображать подсказку или отключать интерфейс)
По умолчанию	Значение ошибок по умолчанию равно 3; значение раз по умолчанию равно 10
Командный режим	Режим конфигурации домена
Руководство по использованию	По умолчанию система проверяет интерфейсы VSL каждые 5 секунд. Если количество кадров с ошибками больше 3 по сравнению с количеством, обнаруженным в последний раз, предполагается, что кадры с ошибками обнаружены один раз. Если ошибочные кадры обнаруживаются последовательно 10 раз, предполагается, что интерфейс неисправен. Действие по умолчанию для ненормального интерфейса — отображение запроса журнала. Вы можете установить действие на отключение интерфейса. Если интерфейс отключен, вы должны восстановить его, отключив и подключив его

ПРИМЕЧАНИЕ: разные продукты предъявляют разные требования к проверке кадра ошибки и разной обработке для интерфейсов VSL. В версии 11.0 проверка ошибок настраивается.

Изменение автономного режима на режим VSU

- Используйте команду **switch convert mode virtual**, чтобы изменить автономный режим на режим VSU.
- В автономном режиме программное обеспечение выполнит следующие действия после запуска команды **switch convert mode virtual**.
- Делается резервная копия файла глобальной конфигурации *config.text* в автономном режиме как *standalone.text* для последующего использования.
- Очищается содержимое конфигурационного файла *config.text*.
- Записывается соответствующая конфигурация VSU в специальный файл конфигурации *config_vsu.dat*.
- При наличии на коммутаторе файла *virtual_switch.text* система предложит перезаписать содержимое файла *virtual_switch.text* в файл *config.text* (файл *virtual_switch.text* является резервным файлом для файла *config.text* когда коммутатор переходит из режима VSU в автономный режим). Затем можно нажать Да или Нет. Наконец коммутатор перезагружается в режиме VSU и считывает параметры VSU в файле *config_vsu.dat*.

Команда	switch convert mode virtual
По умолчанию	Коммутатор по умолчанию находится в автономном режиме
Командный режим	Привилегированный режим EXEC



Руководство по использованию	Измените автономный режим на режим VSU
------------------------------	--

8.4.1.3. Проверка

Запустите команду **show switch virtual config [switch_id]**, чтобы проверить конфигурацию VSU текущего коммутатора в автономном режиме.

Команда	show switch virtual config [switch_id]
Описание параметров	<i>switch_id</i> : указывает идентификатор коммутатора. После указания этого параметра отображается только конфигурация VSU указанного устройства
Командный режим	Привилегированный режим EXEC
Руководство по использованию	Используйте эту команду для отображения конфигурации VSU в автономном режиме или в режиме VSU

ПРИМЕЧАНИЕ: соответствующие конфигурации VSU устанавливаются для одного физического коммутатора и хранятся в специальном файле конфигурации *config_vsu.dat*. Таким образом, вы можете просмотреть текущие конфигурации VSU, запустив команду **show switch virtual config**, а не команду **show running config**.

В автономном режиме рабочая информация VSU пуста. Когда вы вводите такие команды, как **show switch virtual**, система сообщит вам, что коммутатор находится в автономном режиме и информация о работе VSU отсутствует.

8.4.1.4. Пример конфигурации

Настройка VSU в автономном режиме

Сценарий:

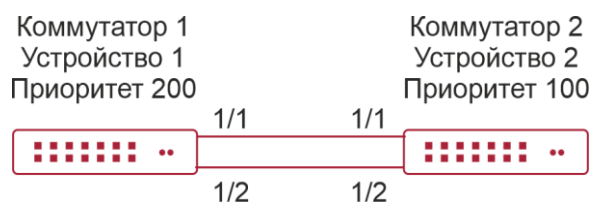


Рисунок 8-17.

Коммутатор 1 и коммутатор 2 образуют систему VSU. Идентификатор домена — 100. Шасси слева настроено как Chassis 1 с приоритетом 200, псевдонимом Switch 1 и интерфейсами VSL 1/1 и 1/2. Шасси справа настроено как Chassis 2 с приоритетом 100, псевдонимом к Switch 2 и интерфейсами VSL 1/1 и 1/2.



Шаги настройки	1. Выполните следующую настройку на коммутаторе 1: • Настройте атрибуты VSU и интерфейсы VSL. • Измените автономный режим на режим VSU. 2. Выполните следующую настройку на коммутаторе 2: • Настройте атрибуты VSU и интерфейсы VSL. • Измените автономный режим на режим VSU
Коммутатор 1	<pre>QTECH # configure terminal QTECH (config)# switch virtual domain 100 QTECH (config-vs-domain)#switch 1 QTECH (config-vs-domain)#switch 1 priority 200 QTECH (config-vs-domain)#witch 1 description switch-1 QTECH (config-vs-domain)# switch crc errors 10 times 20 QTECH (config-vs-domain))#exit QTECH (config)#vsl-port QTECH (config-vsl-port)#port-member interface Tengigabitethernet 1/1 QTECH (config-vsl-port)#port-member interface Tengigabitethernet 1/2 QTECH (config)#exit QTECH #switch convert mode virtual</pre>
Коммутатор 2	<pre>QTECH # configure terminal QTECH (config)# switch virtual domain 100 QTECH (config-vs-domain)# switch 2 QTECH (config-vs-domain)# switch 2 priority 200 QTECH (config-vs-domain)# switch 2 description switch-2 QTECH (config-vs-domain)# switch crc errors 10 times 20 QTECH (config-vs-domain))#exit QTECH (config)#vsl-port QTECH (config-vsl-port)#port-member interface Tengigabitethernet 1/1 QTECH (config-vsl-port)#port-member interface Tengigabitethernet 1/2 QTECH (config-vsl-port)#exit QTECH #switch convert mode virtual</pre>
Проверка	Запустите команду show switch virtual config , чтобы просмотреть атрибуты VSU Switch 1 и Switch 2
Коммутатор 1	<pre>QTECH #show switch virtual config switch_id: 1 (mac: 0x1201aeda0M)</pre>



	<pre> ! switch virtual domain 100 ! switch 1 switch 1 priority 100 ! switch convert mode virtual ! port-member interface Tengigabitethernet 1/1 ! port-member interface Tengigabitethernet 1/2 ! switch crc errors 10 times 20 ! </pre>
Коммутатор 2	<pre> QTECH #show switch virtual config switch_id: 2 (mac: 0x1201aeda0E) ! switch virtual domain 100 ! switch 2 switch 2 priority 100 ! switch convert mode virtual ! port-member interface Tengigabitethernet 1/1 ! port-member interface Tengigabitethernet 1/2 ! switch crc errors 10 times 20 ! </pre>

8.4.1.5. Распространенные ошибки

ПРИМЕЧАНИЕ: интерфейс VSL устройства шасси должен иметь скорость 10 Гбит/с или выше.



8.4.2. Настройка VSU в режиме VSU

8.4.2.1. Настройка атрибутов VSU

8.4.2.2. Эффект конфигурации

Во время работы системы VSU вы можете изменять такие параметры, как идентификатор домена, идентификатор коммутатора и приоритет Master-устройства или slave-устройства. Однако вы можете войти в консоль Master-устройства VSU только для изменения этих параметров, но не можете войти в режим глобальной конфигурации с консоли slave-устройства.

8.4.2.3. Примечания

Среди приведенных выше команд все команды конфигурации вступают в силу только после перезапуска коммутатора, за исключением команды **switch sw_id description switch1**, которая может вступить в силу немедленно.

8.4.2.4. Шаги настройки

Вход в режим настройки домена

- Опционально.
- Запустите эту команду, чтобы войти в режим конфигурации домена. Коммутаторы с одинаковым идентификатором домена образуют систему VSU. Вы можете изменить или настроить идентификатор домена, приоритет коммутатора и идентификатор коммутатора только после входа в режим конфигурации домена в режиме VSU.

Команда	switch virtual domain domain_id
Описание параметров	<i>domain_id</i> : указывает идентификатор виртуального домена системы VSU
По умолчанию	Идентификатор домена по умолчанию — 100
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Только два устройства с одинаковым идентификатором домена могут образовывать систему VSU. Идентификатор домена должен быть уникальным в локальной сети

Изменение идентификатора домена

- Опционально.
- Чтобы изменить значение *domain_id* для устройства, вы можете настроить этот элемент на консоли Master-устройства системы VSU.

Команда	switch switch_id domain new_domain_id
Описание параметров	<i>switch_id</i> : указывает идентификатор работающего в данный момент коммутатора в режиме VSU в диапазоне от 1 до 8.



	<i>new_domain_id</i> : указывает измененный идентификатор домена в диапазоне от 1 до 255
По умолчанию	Идентификатор домена по умолчанию — 100
Командный режим	Режим конфигурации домена
Руководство по использованию	Выполняйте эту команду только в режиме VSU. Кроме того, настройка может вступить в силу только после перезагрузки устройства

Изменение идентификатора коммутатора

- Опционально.
- Чтобы изменить значение *switch_id* для устройства, вы можете настроить этот элемент на консоли Master-устройства системы VSU.

Команда	switch switch_id renumber new_switch_id [force]
Описание параметров	<i>switch_id</i> : указывает идентификатор коммутатора. В системе VSU идентификатор коммутатора находится в диапазоне от 1 до 16 для box-коммутаторов. <i>new_switch_id</i> : указывает измененный идентификатор переключателя
Командный режим	Режим конфигурации домена
Руководство по использованию	Выполняйте эту команду только в режиме VSU. Кроме того, настройка может вступить в силу только после перезагрузки устройства

Изменение приоритета переключения

- Опционально.
- Чтобы изменить приоритет устройства, вы можете настроить этот элемент на консоли Master-устройства системы VSU.
- Большее значение означает более высокий приоритет. Выберите устройство с наивысшим приоритетом в качестве Master-устройства.

Команда	switch switch_id priority priority_num
Описание параметров	<i>switch_id</i> : указывает идентификатор коммутатора, для которого необходимо настроить приоритет. <i>priority_num</i> : указывает приоритет переключателя в диапазоне от 1 до 255 для box-коммутаторов
По умолчанию	Приоритет по умолчанию равен 100



Командный режим	Режим конфигурации домена
Руководство по использованию	Большее значение означает более высокий приоритет. Выберите устройство с наивысшим приоритетом в качестве Master-устройства. Вы можете запустить эту команду в автономном режиме или в режиме VSU. Измененный приоритет вступает в силу только после перезагрузки устройства. Эта команда не используется для изменения значения <i>switch_id</i>. В автономном режиме, если для параметра <i>switch_id</i> установлено значение 1, запуск команды switch 2 priority 200 не работает. Вы можете сначала установить для <i>switch_id</i> значение 2, а затем запустить команду switch 2 priority 200. В режиме VSU <i>switch_id</i> указывает идентификатор работающего в данный момент коммутатора. Если идентификатор не существует, конфигурация не вступает в силу

Настройка описания устройства

- Опционально.
- Чтобы настроить описание устройства, вы можете настроить этот элемент на консоли Master-устройства системы VSU.
- Запустите команду **switch switch_id description switch1**, чтобы настроить описание устройства. Допускается не более 32 символов.

Команда	switch switch_id description dev-name
Описание параметров	<i>switch_id</i>: указывает идентификатор коммутатора, для которого необходимо настроить приоритет. <i>dev-name</i>: указывает имя устройства
Командный режим	Режим конфигурации домена
Руководство по использованию	Вы можете запустить эту команду в автономном режиме или в режиме VSU. Конфигурация вступает в силу немедленно в режиме VSU

Настройка проверки кадров ошибок

- Опционально.
- Запустите команду **switch crc errors error_num times time_num**, чтобы настроить условия для запуска проверки кадра ошибки.

Команда	switch crc errors error_num times time_num
Описание параметров	<i>error_num</i> : настраивает увеличение количества ошибочных кадров между двумя обнаружениями. Когда количество ошибочных кадров превышает прирост, предполагается, что ошибочные кадры обнаруживаются один раз.



	<i>time_num</i> : настраивает количество раз, после которого необходимо выполнить действие (действие может отображать подсказку или отключать интерфейс)
По умолчанию	Значение ошибок по умолчанию равно 3; значение раз по умолчанию равно 10
Командный режим	Режим конфигурации домена
Уровень по умолчанию	14

Сохранение файла конфигурации

Запустите команду **exit**, чтобы выйти из режима настройки виртуального устройства, и запустите команду **write**, чтобы сохранить настройки в файле *config_vsu.dat*.

8.4.2.5. Проверка

Используйте команду **show switch virtual [topology | config]** для отображения текущей информации о работе VSU, топологии или параметров конфигурации.

Команда	show switch virtual [topology config]
Описание параметров	topology : указывает информацию о топологии. config : указывает конфигурации VSU
Командный режим	Привилегированный режим EXEC
Руководство по использованию	Просмотрите идентификатор домена и идентификатор устройства, статус и роль каждого устройства

8.4.2.6. Пример конфигурации

Настройка атрибутов VSU

Сценарий:



Рисунок 8-18.

Коммутатор 1 и Коммутатор 2 образуют систему VSU. Измените идентификатор шасси Коммутатора 2 на 3 и его приоритет на 150. Предположим, что Коммутатор 1 является глобальным Master-коммутатором, и выполните настройку глобального Master-коммутатора.



Шаги настройки	Измените конфигурации Коммутатора 2
Коммутатор 1	<pre> QTECH #config QTECH (config)# switch virtual domain 100 QTECH (config-vs-domain)# switch 2 renumber 3 QTECH (config-vs-domain)# switch 2 priority 150 QTECH (config-vs-domain)# switch 2 description switch-3 </pre>
Проверка	Запустите команду show switch virtual config для проверки
Коммутатор 1	<pre> QTECH #show switch virtual config switch_id: 1 (mac: 0x1201aeda0M) ! switch virtual domain 100 ! switch 1 switch 1 priority 100 ! switch convert mode virtual ! port-member interface Tengigabitethernet 1/1 ! port-member interface Tengigabitethernet 1/2 ! </pre>
	<pre> switch_id: 3 (mac: 0x1201aeda0E) ! switch virtual domain 100 ! switch 3 switch 3 priority 150 ! switch convert mode virtual ! port-member interface Tengigabitethernet 1/1 ! port-member interface Tengigabitethernet 1/2 </pre>



	!
	switch 3 description switch-3
	!

8.4.2.7. Настройка VSL

8.4.2.8. Эффект конфигурации

Когда коммутаторы формируют систему VSU или, когда система VSU работает, вы можете переключаться между общими интерфейсами и интерфейсами VSL. Однако вы можете войти только в консоль Master-устройства системы VSU для внесения изменений, но не можете войти в режим глобальной конфигурации с консоли slave-устройства.

8.4.2.9. Примечания

- Вы можете войти в консоль системы VSU, используя последовательный порт или telnet, чтобы добавить или удалить конфигурации интерфейсов участников VSL.
- Чтобы предотвратить неправильные подключения в действующих сценариях, AP VSL использует динамическое согласование. Сначала необходимо настроить пул интерфейсов VSL, а затем добавить пул интерфейсов VSL к тому же AP после успешного согласования. Интерфейсы, подключающиеся к одному и тому же устройству, находятся в пределах одного AP.

8.4.2.10. Шаги настройки

Вход в режим настройки интерфейса VSL

- Запустите команду **vsl-port**, чтобы войти в режим конфигурации VSL-PORT. Эта команда не является обязательной.
- Когда устройство входит в режим конфигурации VSL-PORT, интерфейс VSL можно настроить или удалить.

Команда	vsl-port
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Вы можете запустить эту команду в автономном режиме или в режиме VSU

Настройка интерфейса участника VSL

- Запустите команду **port-member interface** *interface-name* [**copper** | **fiber**] для добавления интерфейса VSL. Эта команда не является обязательной.
- Запустите команду **port-member interface**, чтобы настроить интерфейс участника VSL.

Команда	port-member interface <i>interface-name</i> [copper fiber]
Описание параметров	<i>interface-name</i> : указывает двумерное имя интерфейса, например GigabitEthernet 0/1 и GigabitEthernet 0/3.



	copper: указывает атрибут медного интерфейса. fiber: указывает атрибут оптического интерфейса
Командный режим	Режим конфигурации интерфейса VSL
Руководство по использованию	Вы можете запустить эту команду в режиме VSU или в автономном режиме. Команда может вступить в силу после сохранения конфигурации команды и перезапуска устройства, на котором находится интерфейс участника VSL

Во время работы системы VSU настроенные каналы участника VSL вступают в силу немедленно. Интерфейсы VSL необходимо настроить для всех устройств.

Для box-устройств интерфейсы VSL могут быть оптическими и медными интерфейсами Gigabit или выше.

Интерфейсы Split 40G «один-четыре» нельзя настроить как интерфейсы VSL.

ПРИМЕЧАНИЕ: для порта 40G (независимо от того, выполняется ли разбиение интерфейса) его интерфейсы-участники (а именно, четыре интерфейса 10G) не могут быть возведены до интерфейсов участника VSL.

ПРИМЕЧАНИЕ: если интерфейс настроен как интерфейс NLB reflex, этот интерфейс можно переключить на интерфейс участника VSL только после удаления конфигурации интерфейса NLB reflex.

ПРИМЕЧАНИЕ: для предотвращения образования петли, которая может возникнуть при выходе интерфейса-участника VSL из AP VSL, система автоматически переводит интерфейс-участник в состояние выключения при выполнении команды, чтобы заставить интерфейс-участник VSL выйти из AP VSL. После того, как интерфейс участника VSL выходит из AP VSL, вы можете повторно подключить канал и выполнить команду **no shutdown**, чтобы снова включить этот интерфейс. При настройке интерфейса VSL система сначала выключит его. Если конфигурация не удалась, и вы хотите использовать его как общий интерфейс, вы можете запустить команду **no shutdown**, чтобы снова включить этот интерфейс. Добавьте номер интерфейса участника, который должен быть трехмерным номером интерфейса. Например, в режиме конфигурации VSL-PORT, если вы запустите команду **port-member interface** TenGigabitEthernet 1/1/1, это указывает на то, что вы настраиваете глобальный трехмерный интерфейс 1/1/1 как интерфейс VSL.

ПРИМЕЧАНИЕ: если при изменении интерфейса VSL на общий интерфейс происходит разделение топологии VSU, интерфейс VSL нельзя удалить. Вы можете сначала отключить физический интерфейс, а затем удалить интерфейс VSL.

8.4.2.11. Проверка

Используйте команду **show switch virtual link [port]**, чтобы отобразить текущую информацию о работе канала VSL в режиме VSU.

Команда	show switch virtual link [port]
Описание параметров	port: отображает информацию о состоянии интерфейсов участников VSL



Командный режим	Привилегированный режим EXEC
-----------------	------------------------------

8.4.2.12. Пример конфигурации

Настройка VSL

Сценарий:

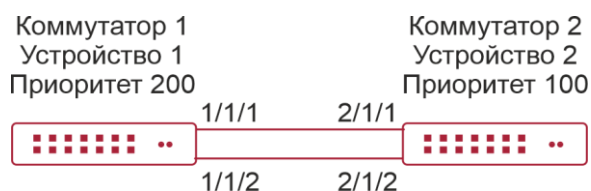


Рисунок 8-19.

Шаги настройки	Добавьте интерфейс 1/1/3 в качестве интерфейса VSL для Коммутатора 1 и удалите интерфейс 1/1/2 из интерфейса VSL
Коммутатор 1	<pre>QTECH #config QTECH (config)# vsl-port QTECH (config-vsl-port)# port-member interface Tengigabitethernet 1/1/3 QTECH (config-vsl-port)# no port-member interface Tengigabitethernet 1/1/2</pre>
Проверка	Запустите команду show switch virtual config , чтобы просмотреть VSL. Предположим, что Коммутатор 1 является глобальным Master-коммутатором, и запустите команду на глобальном Master-коммутаторе
Коммутатор 1	<pre>QTECH #show switch virtual config switch_id: 1 (mac: 0x1201aeda0M) ! switch virtual domain 100 ! switch 1 switch 1 priority 100 ! switch convert mode virtual ! port-member interface Tengigabitethernet 1/1 !</pre>



	<pre>port-member interface Tengigabitethernet 1/3 !</pre>
	<pre>switch_id: 3 (mac: 0x1201aeda0E) ! switch virtual domain 100 ! switch 3 switch 3 priority 150 ! switch convert mode virtual ! port-member interface Tengigabitethernet 1/1 ! port-member interface Tengigabitethernet 1/2 ! switch 3 description switch-3 !</pre>

8.4.2.13. Настройка Dual-Active Detection

8.4.2.14. Эффект конфигурации

Настройте соответствующий механизм обнаружения, чтобы предотвратить создание Dual-Active.

8.4.2.15. Примечания

- DAD можно настроить только в режиме VSU. Вам не разрешено настраивать механизм DAD в автономном режиме.
- Все конфигурации DAD вступают в силу сразу же после настройки на Master или slave-устройствах в режиме глобальной конфигурации с помощью команды **show running-config**.
- Информацию о конфигурации обнаруженного BFD можно отобразить только с помощью команды отображения dual-active detection, а не команды отображения BFD.

8.4.2.16. Шаги настройки

Настройка BFD DAD

- BFD DAD требует установления прямой связи между двумя коммутаторами. Интерфейсы на двух концах должны быть интерфейсами физической маршрутизации. Следующая конфигурация должна быть выполнена на обоих шасси.
- Войдите в режим конфигурации интерфейса DAD-интерфейса и настройте DAD-интерфейс как интерфейс маршрутизации.



- После выхода из режима настройки интерфейса выполните команду **switch virtual domain *domain_id***, чтобы войти в режим настройки домена.
- В режиме домена запустите команду **dual-active detection bfd**, чтобы включить BFD. Эта команда является необязательной и может использоваться, когда BFD DAD нужно настроить.
- В режиме конфигурации домена, запустите команду **dual-active bfd interface *interface-name***, чтобы настроить BFD DAD-интерфейс. Эта команда не является обязательной и может использоваться для настройки интерфейса BFD DAD при настройке BFD DAD.
- Удалите интерфейс BFD DAD. Если интерфейс BFD DAD недоступен, обнаружение BFD использовать нельзя.

Команда	switch virtual domain <i>domain_id</i>
Описание параметров	<i>domain_id</i> : указывает идентификатор домена
По умолчанию	Идентификатор домена по умолчанию — 100
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Только два устройства с одинаковым идентификатором домена могут образовывать систему VSU. Идентификатор домена должен быть уникальным в локальной сети

Команда	dual-active detection { aggregateport bfd }
Описание параметров	aggregateport : указывает режим обнаружения AP. bfd : указывает режим обнаружения BFD
По умолчанию	DAD отключен
Командный режим	Режим конфигурации домена
Руководство по использованию	Настройте эту команду только в режиме VSU

Команда	dual-active bfd interface <i>interface-name</i>
Описание параметров	<i>interface-name</i> : указывает тип интерфейса и идентификатор

Командный режим	Режим конфигурации домена
Руководство по использованию	Интерфейс BFD DAD должен быть интерфейсом маршрутизации и находиться на разных коммутаторах

Интерфейсы обнаружения BFD должны быть напрямую подключены к физическим портам маршрутизации. Два порта должны быть на разных устройствах.

Тип интерфейса не ограничен. Канал dual-active detection используется только для передачи пакетов BFD с небольшим объемом трафика. Поэтому рекомендуется использовать гигабитный интерфейс или интерфейс 100 Мбит/с в качестве интерфейса dual-active detection.

После преобразования интерфейса маршрутизации уровня 3, настроенного с двумя Master-устройствами, в интерфейс коммутатора уровня 2 (выполните команду **switchport** под этим интерфейсом), dual-active detection BFD будет автоматически удалено.

Рекомендуется напрямую подключать интерфейсы обнаружения BFD только к Master и slave-устройствам.

ПРИМЕЧАНИЕ: когда система VSU обнаруживает dual-active-конфликт и переводит другую группу VSU в состояние восстановления, решить проблему можно только путем устранения неисправности VSL, но не путем непосредственного восстановления группы VSU в состоянии восстановления; в противном случае в сети может возникнуть dual-active-конфликт.

Настройка DAD на основе AP

- Чтобы настроить DAD на основе AP, необходимо сначала настроить агрегированный порт (AP), а затем указать порт AP в качестве интерфейса DAD.
- Запустите команду **port-group ap-num**, чтобы добавить физический интерфейс участника к AP.
- После входа в режим конфигурации домена запустите команду **dual-active detection aggregateport**, чтобы включить режим обнаружения AP. Эта команда не является обязательной. Вы можете запустить эту команду, когда необходимо настроить обнаружение AP.
- Запустите команду **dual-active interface interface-name**, чтобы настроить AP в качестве DAD-интерфейса. Эта команда не является обязательной. Можно запустить эту команду, чтобы настроить AP в качестве интерфейса DAD, когда необходимо настроить обнаружение AP.
- Запустите команду **dad relay enable**, чтобы включить ретрансляцию пакетов dual-active detection для upstream- и downstream-интерфейсов. Эта команда не является обязательной. Вы можете запустить эту команду для ретрансляции пакетов DAD (пакетов dual-active detection), когда DAD на основе AP настроен.
- Отключение DAD на основе AP отключит DAD.
- Удалите обнаруженный интерфейс. Если интерфейс DAD на основе AP недоступен, DAD на основе AP использовать нельзя.
- Ретрансляция пакетов DAD на основе AP отключена по умолчанию.



Команда	dual-active detection { aggregateport bfd }
Описание параметров	aggregateport : указывает режим обнаружения AP. bfd : указывает режим обнаружения BFD
По умолчанию	DAD отключен
Командный режим	Режим конфигурации домена
Руководство по использованию	Настройте эту команду только в режиме VSU

Команда	dual-active interface <i>interface-name</i>
Описание параметров	<i>interface-name</i> : указывает тип интерфейса и идентификатор интерфейса. Должен быть указан интерфейс DAD на основе AP
Командный режим	Режим конфигурации домена
Руководство по использованию	Можно настроить только один интерфейс DAD на основе AP. Этот интерфейс должен быть создан до того, как вы настроите AP в качестве интерфейса DAD. Последующие настроенные интерфейсы DAD перезапишут предыдущие

Команда	dad relay enable
По умолчанию	Ретрансляция пакетов DAD на основе AP отключена по умолчанию
Командный режим	Режим конфигурации интерфейса
Руководство по использованию	Эта команда может быть выполнена только на AP

ПРИМЕЧАНИЕ: рекомендуется распределить физические интерфейсы, которые добавляются к интерфейсу обнаружения на основе AP, на разные устройства.

8.4.2.17. Проверка

Используйте **show switch virtual dual-active { aggregateport | bfd | summary }** для отображения текущей конфигурации DAD.



Команда	show switch virtual dual-active { aggregateport bfd summary }
Описание параметров	aggregateport: отображает информацию о DAD на AP. bfd: отображает информацию DAD на основе BFD. summary: отображает сводку DAD
Командный режим	Привилегированный режим EXEC

8.4.2.18. Пример конфигурации

Настройка BFD DAD

Сценарий:

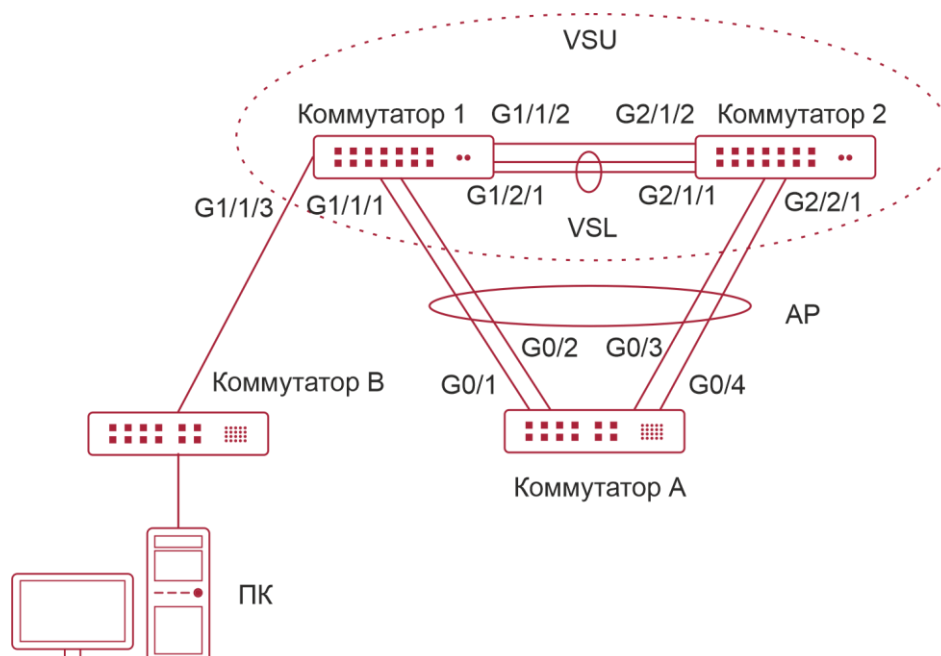


Рисунок 8-20.

Коммутатор 1 и Коммутатор 2 образуют систему VSU (идентификатор домена равен 1). Приоритеты Коммутатора 1 и Коммутатора 2 равны 200 и 150 соответственно. Связи между Te1/3/1 и Te1/3/2 Коммутатора 1 и Te2/3/1 и Te2/3/2 Коммутатора 2 устанавливаются соответственно для формирования VSL между Коммутатором 1 и Коммутатором 2. G0/1, интерфейсы G0/2, G0/3 и G0/4 Коммутатора А подключены к G1/1/1 и G1/2/1 Коммутатора 1 и G2/1/1 и G2/2/1 Коммутатора 2 для формирования группы AP, включающая четыре канала-участника. Идентификатор группы AP — 1. Все члены группы AP 1 являются гигабитными оптическими интерфейсами. G1/1/2 и G2/1/2 являются интерфейсами маршрутизации.

G1/1/2 и G2/1/2 — это пара интерфейсов BFD DAD.



Шаги настройки	- Настройте G1/1/2 и G2/1/2 в качестве интерфейсов маршрутизации. - Включите BFD DAD. - Настройте G1/1/2 и G2/1/2 как интерфейсы BFD DAD. Поскольку Коммутатор 1 и Коммутатор 2 находятся в системе VSU, предыдущую настройку можно выполнить либо на Коммутаторе 1, либо на Коммутаторе 2. В следующем примере настраиваются функции на Коммутаторе 1
Коммутатор 1	<pre>QTECH (config)# interface GigabitEthernet 1/1/2 QTECH (config-if-GigabitEthernet 1/1/2)# no switchport QTECH (config)# interface GigabitEthernet 2/1/2 QTECH (config-if-GigabitEthernet 2/1/2)# no switchport QTECH (config-if)# switch virtual domain 1 QTECH (c config-vs-domain)# dual-active detection bfd QTECH (config-vs-domain)# dual-active bfd interface GigabitEthernet 1/1/2 QTECH (config-vs-domain)# dual-active bfd interface GigabitEthernet 2/1/2</pre>
Коммутатор A	<pre>QTECH # configure terminal</pre> Введите команды конфигурации, по одной в строке. Оканчивайте с CNTL/Z. <pre>QTECH (config)# interface aggregateport 1 QTECH (config-if-aggregateport 1)# interface range GigabitEthernet 0/1-4 QTECH (config-if-aggregateport 1)# port-group 1 QTECH (config)# interface vlan 1 QTECH (config-if-vlan 1)#ip address 1.1.1.2 255.255.255.0 QTECH (config-if-vlan 1)#exit QTECH (config)#interface aggregateport 1 QTECH (config-if-AggregatePort 1)# dad relay enable QTECH (config-if-AggregatePort 1)# exit</pre>
Проверка	- Просмотрите конфигурацию DAD. - Просмотрите конфигурацию BFD DAD
Коммутатор 1	<pre>QTECH # show switch virtual dual-active summary</pre> BFD dual-active detection enabled: No Aggregateport dual-active detection enabled: Yes Interfaces excluded from shutdown in recovery mode:



	In dual-active recovery mode: NO QTECH # show switch virtual dual-active bfd BFD dual-active detection enabled: Yes BFD dual-active interface configured: GigabitEthernet 1/1/2: UP GigabitEthernet 2/1/2: UP
--	--

8.4.2.19. Распространенные ошибки

- Интерфейс BFD DAD не является интерфейсом маршрутизации.
- Ни BFD DAD, ни DAD на основе AP не включены и не активированы.

8.4.2.20. Настройка балансировки трафика

8.4.2.21. Эффект конфигурации

В системе VSU, если выходы распределены по нескольким устройствам, можно настроить Local Forward First (LFF).

8.4.2.22. Примечания

Конфигурация по умолчанию — LFF.

8.4.2.23. Шаги настройки

Настройка режима AP LFF

- В режиме конфигурации домена запустите команду **switch virtual aggregateport-lff enable**, чтобы включить режим AP LFF. Эта команда не является обязательной.
- Порты-участники AP могут быть распределены по двум шасси системы VSU. Вы можете настроить, будет ли исходящий трафик AP перенаправляться сначала через локальные порты-участники, исходя из фактических условий трафика.
- Если эта функция отключена, трафик перенаправляется на основе правил конфигурации AP. Дополнительные сведения см. в разделе Ethernet Switching/Настройка AP.

Команда	switch virtual aggregateport-lff enable
По умолчанию	Эта функция включена по умолчанию
Командный режим	Режим конфигурации домена
Руководство по использованию	Включите AP LFF в режиме VSU

Настройка режима ECMP LFF

- В режиме конфигурации домена запустите команду **switch virtual ecmp-lff enable**, чтобы включить режим ECMP LFF. Эта команда не является обязательной.



- Выход маршрутизации Equal-Cost MultiPath (ECMP) может быть распределен на два шасси системы VSU. Вы можете настроить, будет ли исходящий трафик ECMP перенаправляться сначала через локальные порты-участники, исходя из фактических условий трафика.
- Если эта функция отключена, трафик перенаправляется на основе правил конфигурации ECMP. Дополнительные сведения см. в разделе Ethernet Switching/Настройка AP.

Команда	switch virtual ecmp-lff enable
По умолчанию	Эта функция включена по умолчанию
Командный режим	Режим конфигурации домена
Руководство по использованию	Включите ECMP LFF в режиме VSU

ПРИМЕЧАНИЕ: в режиме VSU режим LFF AP между шасси и режим ECMP LFF отключены по умолчанию.

ПРИМЕЧАНИЕ: чтобы развернуть систему VSU для коммутаторов уровня 3, рекомендуется настроить балансировку нагрузки AP на основе IP (src-ip, dst-ip и src-dst-ip).

8.4.2.24. Проверка

Используйте команду **show switch virtual balance**, чтобы отобразить текущий режим балансировки трафика системы VSU.

Команда	show switch virtual balance
Командный режим	Привилегированный режим EXEC
Руководство по использованию	Используйте эту команду для отображения конфигурации режима балансировки трафика в режиме VSU

8.4.2.25. Пример конфигурации

Настройка LFF

Сценарий:

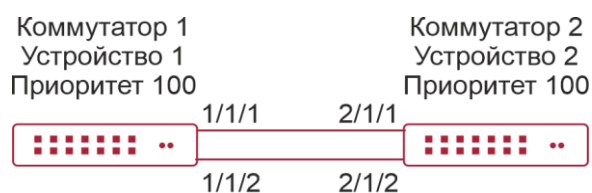


Рисунок 8-21.



	На Рисунке 8-21 Коммутатор 1 и Коммутатор 2 образуют систему VSU. Предполагается, что Коммутатор 1 является глобальным Master-коммутатором, и настройка выполняется на Коммутаторе 1
Шаги настройки	Настройте AP LFF
Коммутатор 1	<pre>QTECH #config QTECH (config)# switch virtual domain 100 QTECH (config-vs-domain)# switch virtual aggregateport-lff enable</pre>
Проверка	Запустите команду show switch virtual balance для проверки
Коммутатор 1	<pre>QTECH #show switch virtual balance Aggregate port LFF: enable Ecmp lff enable</pre>

8.4.2.26. Изменение режима VSU на автономный режим

8.4.2.27. Эффект конфигурации

Развести систему VSU на отдельные устройства, которые могут работать в автономном режиме.

8.4.2.28. Шаги настройки

- Запустите команду **switch convert mode standalone** [*switch_id*], чтобы изменить режим VSU на автономный режим. Эта команда не является обязательной.
- После выполнения этой команды система предложит вам следующее: Восстановить ли сохраненный файл конфигурации в автономном режиме? Если **yes**, файл конфигурации будет восстановлен; если **no**, конфигурация будет очищена.

Команда	switch convert mode standalone [<i>switch_id</i>]
Описание параметров	<i>switch_id</i> : указывает идентификатор коммутатора
По умолчанию	Коммутатор по умолчанию находится в автономном режиме
Командный режим	Привилегированный режим EXEC
Руководство по использованию	После запуска команды switch convert mode standalone Master-коммутатор создает резервную копию файлов глобальной конфигурации всех VSD в режиме VSU с идентификатором <i>vsd.virtual_switch.text.vsd ID</i> . Затем Master-коммутатор очищает файлы глобальной конфигурации <i>config.text</i> всех VSD в режиме VSU и



	спрашивает, перезаписать ли файлы глобальной конфигурации <i>config.text</i> на <i>vsd.standalone.text.vsd</i> ID. Если вы выберете yes, содержимое <i>vsd.standalone.text.vsd</i> ID перезапишет глобальный файл конфигурации <i>config.text</i> всех VSD; в противном случае Master-коммутатор не восстанавливает <i>config.text</i>. Наконец, перезапустите коммутатор. Эту команду можно использовать в автономном режиме или в режиме VSU. Если команда выполняется в автономном режиме, то переключение режима происходит на текущем коммутаторе. Если команда содержит параметр <i>sw_id</i> и выполняется в режиме VSU, то переключение режима выполняется на коммутаторе с идентификатором, указанным в <i>sw_id</i>. Если команда не содержит параметр <i>sw_id</i>, то переключение режима выполняется на Master-коммутаторе. Рекомендуется переключить режим slave-коммутатора, а затем Master-коммутатора
--	---

8.4.2.29. Пример конфигурации

Изменение режима VSU на автономный режим

Сценарий:

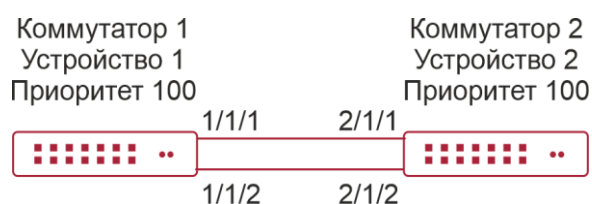


Рисунок 8-22.

На Рисунке 8-22 предполагается, что Коммутатор 1 и Коммутатор 2 образуют систему VSU, а Коммутатор 1 является глобальным Master-коммутатором.

Шаги настройки	Измените режим Коммутатора 1 на автономный режим. Измените режим Коммутатора 2 на автономный режим
Коммутатор 1	QTECH # switch convert mode standalone 1 QTECH # switch convert mode standalone 2
Проверка	Запустите команду show switch virtual config , чтобы отобразить состояние коммутатора
Коммутатор 1	QTECH #show switch virtual config switch_id: 1 (mac: 0x1201aeda0M) ! switch virtual domain 100



	<pre> ! switch 1 switch 1 priority 100 ! switch convert mode standalone ! port-member interface Tengigabitethernet 1/1 ! port-member interface Tengigabitethernet 1/3 ! </pre>
	<pre> switch_id: 2 (mac: 0x1201aeda0E) ! switch virtual domain 100 ! switch 2 switch 2 priority 150 ! switch convert mode standalone ! port-member interface Tengigabitethernet 1/1 ! port-member interface Tengigabitethernet 1/2 ! switch 2 description switch-2 ! </pre>

8.4.3. Настройка Определения нахождения устройства быстрым миганием

8.4.3.1. Эффект конфигурации

Включите Определение нахождения устройства быстрым миганием, чтобы светодиодный индикатор состояния коммутатора быстро мигал.

8.4.3.2. Примечания

Если не отключить Определение нахождения устройства быстрым миганием, система автоматически отключит функцию через 30 минут после ее включения.



8.4.3.3. Шаги настройки

Включение/отключение Определения нахождения устройства быстрым миганием

- Обязательный. Используйте эту функцию на коммутаторе, который необходимо найти.
- В привилегированном режиме EXEC запустите команду **led-blink**, чтобы включить Определение нахождения устройства быстрым миганием.

Команда	led-blink { enable disable } [device device_id]
Описание параметров	enable: включает Определение нахождения устройства быстрым миганием. disable: отключает Определение нахождения устройства быстрым миганием. <i>device_id:</i> указывает идентификатор устройства
По умолчанию	Определение нахождения устройства быстрым миганием отключено по умолчанию
Командный режим	Привилегированный режим EXEC
Руководство по использованию	Запустите эту команду без параметра <i>device_id</i>, чтобы включить или отключить быстрый мигающий поиск в автономном режиме. В режиме VSU вы можете установить параметр <i>device_id</i>, чтобы включить или отключить эту функцию для указанного устройства. Если вы игнорируете параметр <i>device_id</i>, вы можете включить или отключить эту функцию для всех устройств в системе VSU. Если вы не отключите эту функцию, система автоматически отключит функцию через 30 минут после ее включения. Эта конфигурация не может быть сохранена. Определение нахождения устройства быстрым миганием будет отключено при перезапуске

8.4.3.4. Проверка

Проверьте, быстро ли мигает светодиод состояния коммутатора.



8.4.3.5. Пример конфигурации

Включение Определения нахождения устройства быстрым миганием для двух устройств VSU

Сценарий	Предположим, что Коммутатор 1 и Коммутатор 2 образуют систему VSU, а Коммутатор 1 является глобальным Master-устройством
Шаги настройки	- Введите команду led-blink enable device 2 на консоли Коммутатора 1, чтобы включить Определение нахождения устройства быстрым миганием. - Введите команду led-blink disable device 2 на консоли Коммутатора 1, чтобы отключить Определение нахождения устройства быстрым миганием. (в оригинале стояло enable)
Проверка	Если включена функция Определения нахождения устройства быстрым миганием, проверьте, быстро ли мигает светодиод состояния Коммутатора 2

8.4.4. Настройка интерфейса MGMT

8.4.4.1. Эффект конфигурации

Настройте устройство для создания одного интерфейса MGMT для каждого шасси или создайте только один интерфейс MGMT для системы в режиме VSU.

8.4.4.2. Примечания

После настройки этой команды запустите команду **write**, чтобы сохранить конфигурацию. Конфигурация вступает в силу только в режиме VSU и только после перезагрузки устройства.

8.4.4.3. Шаги настройки

Настройка устройства для создания только одного интерфейса MGMT для системы в режиме VSU

- Опционально. Настраивайте эту функцию только тогда, когда устройству необходимо создать только один интерфейс MGMT для системы. По умолчанию для каждого шасси создается один интерфейс MGMT.
- В режиме глобальной конфигурации запустите команду **mgmt_mode**, чтобы настроить интерфейс MGMT.

Команда	mgmt_mode unique
Командный режим	Режим глобальной конфигурации
Руководство по использованию	В режиме VSU настройте систему на создание только одного интерфейса MGMT. Конфигурацию необходимо сохранить, и она вступает в силу только после перезагрузки устройства



Настройка устройства для создания одного интерфейса MGMT для каждого шасси

Опционально. После настройки устройства для создания только одного интерфейса MGMT используйте эту команду для восстановления конфигурации по умолчанию, то есть для создания одного интерфейса MGMT для каждого шасси.

Команда	no mgmt_mode
Командный режим	Режим глобальной конфигурации
Руководство по использованию	В режиме VSU настройте систему на создание одного интерфейса MGMT для каждого шасси. Конфигурацию необходимо сохранить, и она вступает в силу только после перезагрузки устройства

8.4.4.4. Проверка

После настройки запустите команду **write**, чтобы сохранить конфигурацию. После выполнения команды **reload** для перезагрузки устройства запустите команду **show interface**, чтобы отобразить количество интерфейсов MGMT.

8.4.4.5. Пример конфигурации

Настройка устройства для создания только одного интерфейса MGMT для системы в режиме VSU

Сценарий:

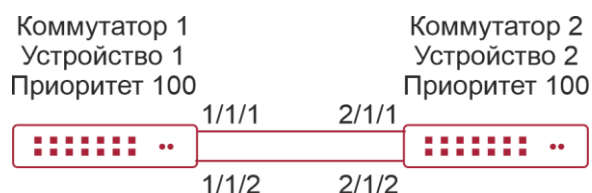


Рисунок 8-23.

Коммутатор 1 и Коммутатор 2 образуют VSU, а Коммутатор 1 является глобальным активным устройством.

Шаги настройки	Запустите команду, чтобы настроить устройство для создания только одного интерфейса MGMT для системы, запустите команду write , чтобы сохранить конфигурацию, и запустите команду reload , чтобы перезапустить устройство, чтобы конфигурация вступила в силу
	<pre>QTECH (config)#mgmt_mode unique Please write and reload system to take config effect! QTECH (config)#exit QTECH #write</pre>



	Building configuration... [OK] QTECH #reload Reload system?(Y/N)y
Проверка	После перезагрузки устройства запустите команду show interfaces inc mgmt . Система создает только один интерфейс MGMT даже при наличии нескольких шасси в режиме VSU. QTECH #show interfaces inc Mgmt ===== Mgmt 0 ===== Mgmt 0 is UP , line protocol is UP Hardware is Mgmt, address is 1414.3344.5519 (bia 1414.3344.5519)

8.4.5. Настройка восстановления устройства в режиме восстановления

8.4.5.1. Эффект конфигурации

Отключает функцию автоматического перезапуска и восстановления в режиме восстановления.

8.4.5.2. Примечания

Если функция автоматического перезапуска и восстановления отключена, ее необходимо снова включить или вручную перезапустить устройства, находящиеся в режиме восстановления.

8.4.5.3. Шаги настройки

Включение/отключение функции автоматического перезапуска в режиме восстановления

- Обязательный. Включите или отключите функцию на устройстве по мере необходимости.
- В режиме конфигурации config-vs-domain запустите команду **[no] recovery auto-restart enable**, чтобы включить или отключить функцию автоматического перезапуска.

Команда	recovery auto-restart enable
По умолчанию	Функция автоматического перезапуска и восстановления в режиме восстановления включена по умолчанию
Командный режим	Режим конфигурации config-vs-domain



Руководство по использованию	Эта команда может быть выполнена только в режиме VSU. После настройки команды ее необходимо сохранить, чтобы она вступила в силу немедленно
------------------------------	---

8.4.5.4. Проверка

Запустите команду **show run**, чтобы отобразить конфигурации.

8.4.5.5. Пример конфигурации

Отключение автоматического перезапуска и восстановления в режиме восстановления в режиме VSU

Сценарий:

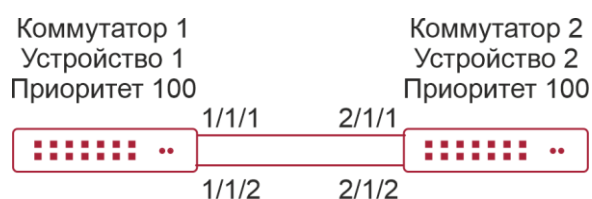


Рисунок 8-24.

Коммутатор 1 и Коммутатор 2 образуют VSU, а Коммутатор 1 является глобальным активным устройством. Функция dual-active detection включена.

Шаги настройки	- Запустите команду switch virtual domain 100 на Коммутаторе 1, чтобы войти в режим конфигурации config-vs-domain. - Запустите команду no recovery auto-restart enable на Коммутаторе 1, чтобы отключить функцию автоматического перезапуска и восстановления
Проверка	Отключите канал VSL. После завершения dual-active detection Коммутатор 2 переходит в режим восстановления. Повторно подключите канал VSL. Коммутатор 2 не перезагружается. Запустите команду recovery auto-restart enable на Коммутаторе 2, чтобы включить функцию автоматического перезапуска и восстановления. Коммутатор 2 автоматически сбрасывается

8.4.6. Настройка автоматического восстановления без перезагрузки в режиме восстановления

8.4.6.1. Эффект конфигурации

Включает функцию автоматического восстановления без перезагрузки в режиме восстановления.



8.4.6.2. Шаги настройки

Включение/отключение функции автоматического восстановления без перезагрузки в режиме восстановления

- Обязательный. Включите или отключите функцию на устройстве по мере необходимости.
- В режиме конфигурации config-vs-domain запустите команду **[no] dual-active auto-recovery enable**, чтобы включить или отключить функцию автоматического восстановления без перезапуска.

Команда	dual-active auto-recovery enable
По умолчанию	Функция автоматического восстановления без перезагрузки в режиме восстановления по умолчанию отключена
Командный режим	Режим конфигурации config-vs-domain
Руководство по использованию	Эта команда может быть выполнена только в режиме VSU. После настройки команды ее необходимо сохранить для немедленной проверки

8.4.6.3. Проверка

Запустите команду **show run**, чтобы отобразить конфигурации.

8.4.6.4. Пример конфигурации

Включение автоматического восстановления без перезагрузки в режиме восстановления в режиме VSU

Сценарий:

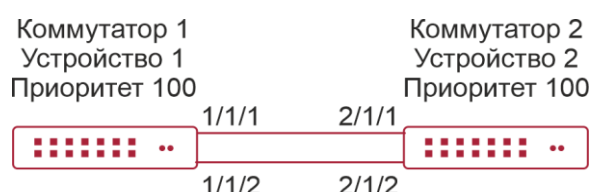


Рисунок 8-25.

Коммутатор 1 и Коммутатор 2 образуют VSU, а Коммутатор 1 является глобальным активным устройством. Функция DAD включена.

Шаги настройки	• Запустите команду switch virtual domain 100 на Коммутаторе 1, чтобы войти в режим конфигурации config-vs-domain. • Используйте команду dual-active auto-recovery enable на Коммутаторе 1, чтобы включить функцию автоматического восстановления без перезагрузки в режиме восстановления
----------------	---



Проверка	Отключите все каналы VSL. После завершения DAD Коммутатор 2 переходит в режим восстановления. Выключите питание Коммутатора 1. Убедитесь, что Коммутатор 2 автоматически становится Master-устройством без перезапуска
----------	--

8.5. Мониторинг и обслуживание

8.5.1. Отображение

Описание	Команда
Отображает текущую работу VSU, топологию или конфигурацию	show switch virtual [topology config role]
Отображает текущую конфигурацию dual-active	show switch virtual dual-active { bfd aggregateport summary }
Отображает текущую информацию о работе VSL в режиме VSU	show switch virtual link [port]
Перенаправляет на консоль Master-коммутатора или любого другого коммутатора	session { device <i>switch_id</i> Master }
Отображает текущий идентификатор коммутатора	show switch id



9. НАСТРОЙКА RNS

9.1. Обзор

Служба надежной сети (RNS) тестирует определенные службы, предоставляемые реер-устройством, для мониторинга доступности службы, целостности соединения end-to-end и качества службы. Использование результаты тестов RNS, вы можете:

- Своевременно изучить производительность сети и принять соответствующие меры для решения связанных с ней проблем с производительностью.
- Диагностировать и локализовать сбои в сети.

9.2. Приложение

9.2.1. Тестирование и оценка эффективности службы

9.2.1.1. Сценарий

Как показано на следующем рисунке, компания собирается развернуть систему видео-конференц-связи между штаб-квартирой и филиалами и выполнила соответствующие настройки качества обслуживания (QoS). Перед официальным развертыванием необходимо проверить, могут ли услуги предоставляться в обычном режиме в условиях существующего давления компании на службу. Система видео-конференц-связи чувствительна к задержке протокола пользовательских датаграмм (UDP) и джиттеру передачи UDP в сети. Традиционный инструмент проверки связи может тестировать производительность интернет-протокола управляющих сообщений (ICMP), но не может эффективно оценить производительность передачи UDP и не может удовлетворить требования к измерению джиттера.

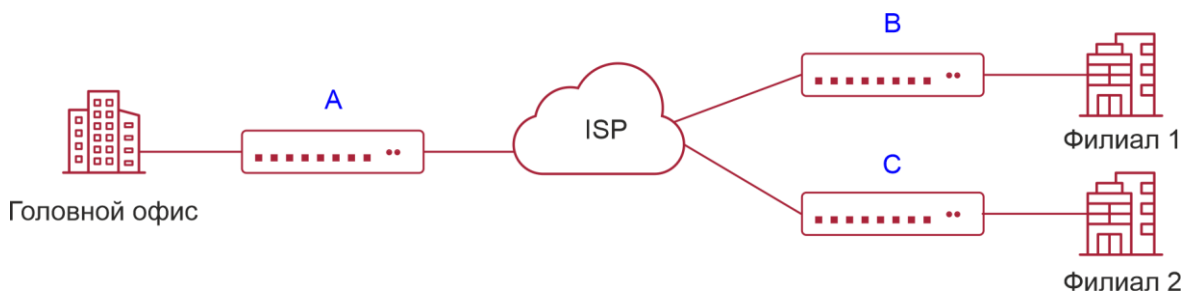


Рисунок 9-1.

A, B и C — коммутаторы.

9.2.1.2. Развертывание

- Настройте RNS на выходном коммутационном устройстве или коммутаторе каждой ветви, чтобы проверить джиттер и задержку UDP.
- На коммутаторе A укажите IP-адрес и UDP-порт выходного коммутационного устройства или коммутатора в штаб-квартире, после чего UDP-пакеты могут отправляться автоматически. В зависимости от конфигурации выходное коммутационное устройство или коммутатор в штаб-квартире могут автоматически отвечать на пакеты UDP. Выходное коммутационное устройство или коммутатор филиала обрабатывает отправленные и полученные пакеты и

вычисляет джиттер UDP. Чтобы узнать производительность в разные периоды времени, вам также необходимо настроить функции планирования, такие как периодический запуск/остановка и повторный запуск, для RNS.

9.2.2. Обнаружение сетевых сбоев

9.2.2.1. Сценарий

В сети кампуса, как показано на Рисунке 9-2, Студент 1 сообщает об ошибке доступа к веб-серверу, Студент 3 сообщает об ошибке доступа в Интернет, а Студент 6 сообщает об ошибке отправки/получения электронной почты.

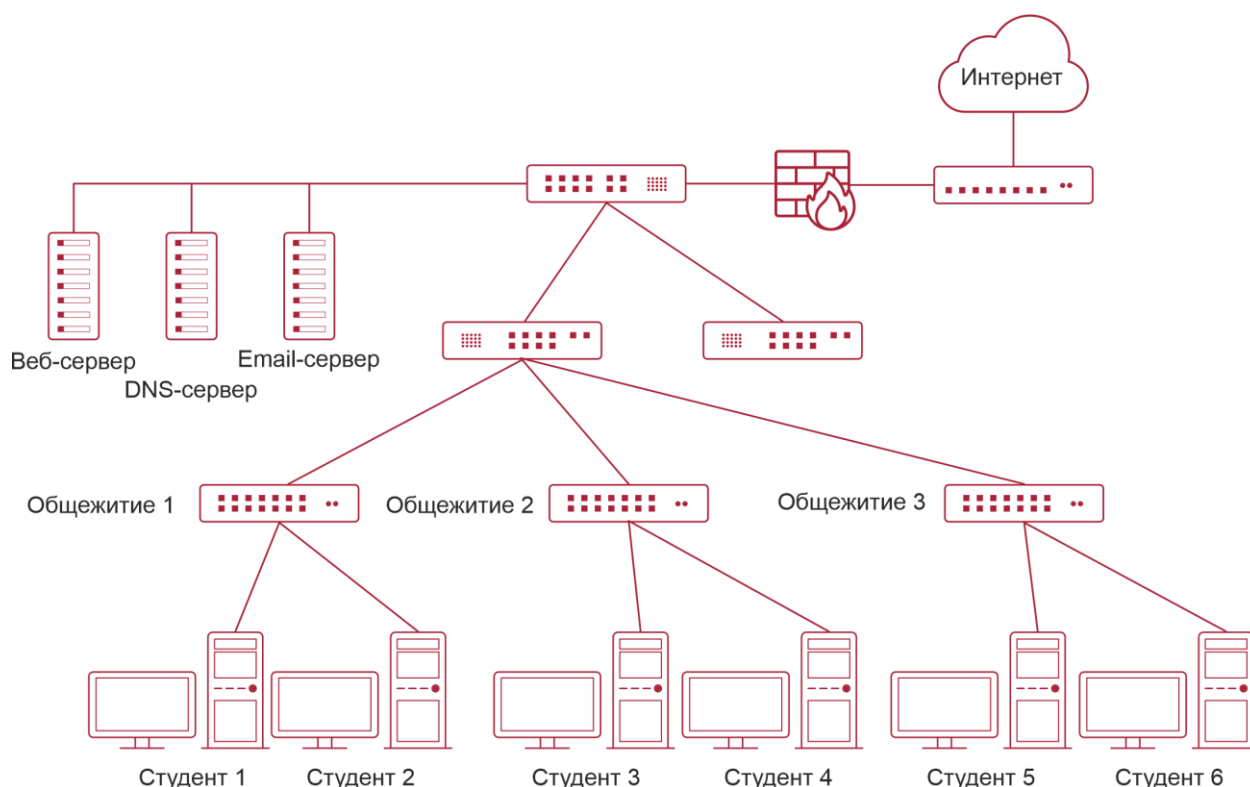


Рисунок 9-2.

9.2.2.2. Развертывание

- Администратор напрямую включает функцию DNS на коммутаторе доступа в общежитии, чтобы проверить, неисправен ли сервер службы доменных имен (DNS). В случае сбоя DNS автоматически запускается эхо-пакет ICMP для проверки доступности веб-сервера.
- При возникновении ошибки администратору нужно только запустить тест, а последующие тесты могут запускаться автоматически. Затем администратор может проверить результаты теста, чтобы найти неисправность, что значительно снижает нагрузку на администратора.



9.3. Функции

9.3.1. Базовые концепты

Экземпляр RNS

Экземпляр RNS можно рассматривать как процесс RNS. Перед выполнением RNS необходимо создать экземпляр RNS. В экземпляре RNS необходимо настроить параметры RNS, такие как тип теста, адрес назначения теста и частота тестирования. Идентификатор экземпляра является глобально уникальным.

Особенность	Описание
Тест RNS	Отслеживает сетевое подключение, доступность услуг, целостность подключения end-to-end и качество обслуживания
Отслеживание поддержки RNS	Отслеживает результаты теста и уведомляет соответствующий модуль о результатах

9.3.2. Тест RNS

Отслеживает сетевое подключение, доступность услуг, целостность подключения end-to-end и качество обслуживания. Например, проверьте, нормально ли работает функция DNS устройства. В настоящее время RNS поддерживает следующие типы тестов: ICMP-эхо, DNS и TCP.

9.3.2.1. Принцип работы

Эхо-тест ICMP

ICMP-эхо — это базовая функция RNS, реализованная в соответствии с RFC 2925. Пакет ICMP отправляется для проверки доступности пункта назначения, а также для расчета времени отклика сети и коэффициента потери пакетов.

Пакет эхо-запроса ICMP отправляется на IP-адрес назначения на основе заданного времени и частоты тестирования. После получения пакета эхо-запроса ICMP с IP-адреса назначения возвращается пакет эхо-ответа ICMP. С помощью эхо-теста ICMP вычисляется время отклика и скорость потери пакетов на основе информации, относящейся к полученному пакету эхо-ответа ICMP, например, время получения и количество пакетов. Таким образом, отражаются текущая производительность и состояние сети. Результаты эхо-тестирования ICMP и записи истории будут записаны, и вы можете использовать командную строку для их отображения.

ПРИМЕЧАНИЕ: необходимым условием успешного прохождения эхо-теста ICMP является то, что целевые устройства могут правильно отвечать на пакеты эхо-запроса ICMP.

DNS-тест

В тесте DNS имитируется DNS-клиент, который отправляет запрос на разрешение имени домена на указанный DNS-сервер. Вы можете определить, доступен ли DNS-сервер и скорость разрешения доменного имени, проверив результат разрешения доменного имени и время, необходимое для разрешения доменного имени. В тесте DNS имитируется процесс разрешения доменного имени, и сопоставление между разрешенным доменным именем и IP-адресом не сохраняется. Результаты тестирования DNS и записи истории будут записаны в тестовой группе. Вы можете использовать командную строку для проверки результатов тестирования и записей истории.



Процедура настройки теста экземпляра RNS

1. Создайте экземпляр и настройте тест на основе типа теста.
2. Запустите экземпляр.
3. Используйте экземпляр RNS для создания пакета определенного типа теста и отправки пакета на peer end.
4. После получения тестового пакета peer end возвращает ответный пакет соответствующего типа.
5. Экземпляр RNS вычисляет скорость потери пакетов и время приема-передачи в зависимости от того, получен ли ответный пакет, и времени получения ответного пакета.
6. Используйте команду **show** или **debug**, чтобы проверить результат теста.

ПРИМЕЧАНИЕ: выше описаны общие процедуры для тестов экземпляра RNS. Дополнительные сведения о настройке см. в следующих разделах.

9.3.2.2. Связанная конфигурация

Настройка интервала повторения теста

По умолчанию интервал повторения теста составляет 60 секунд.

В режиме настройки RNS запустите команду **frequency millisecond**, чтобы настроить интервал повторения теста.

Настройте частоту на основе следующей формулы, чтобы обеспечить правильный расчет теста.

$(\text{frequency milliseconds}) > (\text{timeout milliseconds}) \geq (\text{threshold milliseconds})$

Настройка времени ожидания (тайм-аута) теста

Тайм-аут по умолчанию зависит от типа теста. Вы можете запустить команду **show ip rns configuration**, чтобы отобразить время тайм-аута типа теста.

В режиме конфигурации RNS запустите команду **timeout milliseconds**, чтобы настроить время тайм-аута экземпляра.

Настройте время тайм-аута на основе формулы. Подробнее см. в «Руководстве по использованию» команды **frequency**.

Настройте пороговое значение времени тестирования.

Настройка порогового значения теста

По умолчанию пороговое значение теста равно 5000 мс.

В режиме конфигурации RNS запустите команду **threshold milliseconds**, чтобы настроить пороговое значение теста экземпляра.

Настройте порог на основе формулы. Подробнее см. в «Руководстве по использованию» команды **frequency**.

Настройка тега для теста

Конфигурация по умолчанию недоступна.

В режиме настройки RNS запустите команду **tag text**, чтобы настроить тестовый тег.

Вы можете запустить команду **tag**, чтобы указать тег для идентификации теста.

Настройка размера полезной нагрузки протокола

Размер полезной нагрузки протокола по умолчанию зависит от типа теста. По умолчанию размер полезной нагрузки протокола является минимальным или подходящим размером для пакетов протокола соответствующего типа теста.



В режиме конфигурации RNS запустите команду **request-data-size bytes**, чтобы настроить размер полезной нагрузки протокола.

Выполните эту настройку в режиме настройки IP RNS.

Настройка поля TOS тестового пакета

По умолчанию TOS равен 0.

В режиме настройки RNS запустите команду **tos number**, чтобы настроить поле TOS в заголовке IPv4 тестовых пакетов RNS.

Настройка VRF

В режиме конфигурации RNS запустите команду **vrf vrf-name** для виртуальной маршрутизации и пересылки (VRF) для экземпляра RNS.

9.3.3. Отслеживание поддержки RNS

Объекты, которые можно отслеживать, включают: результат тестирования экземпляра RNS, состояние списка RNS, состояние канала на интерфейсе и состояние списка отслеживания. При изменении состояния отслеживания срабатывает действие других модулей.

9.3.3.1. Принцип работы

Результат теста экземпляра RNS отслеживается следующим образом:

- Настройте объект отслеживания для отслеживания результатов теста экземпляра RNS.
- Когда результат теста экземпляра RNS изменяется, модуль RNS отправляет сообщение об изменении состояния модулю отслеживания.
- Модуль отслеживания получает результат теста. После заданной задержки, если результат проверки не изменился, статус объекта отслеживания изменяется, и модуль объекта отслеживания уведомляется об изменении. Если результат теста восстанавливается в течение периода, статус объекта отслеживания не изменяется и соответствующий модуль не уведомляется.

9.3.3.2. Связанная конфигурация

Настройка объекта отслеживания для отслеживания статуса канала интерфейса

По умолчанию функция отслеживания статуса канала интерфейса отключена.

Запустите команду **track interface line-protocol**, чтобы настроить объект отслеживания, который используется для отслеживания состояния канала интерфейса.

Если статус канала интерфейса — UP, статус объекта отслеживания — UP. Если статус канала интерфейса - DOWN, статус объекта отслеживания также будет DOWN.

Настройка объекта отслеживания для отслеживания результатов тестирования экземпляра RNS

По умолчанию функция отслеживания результата теста экземпляра RNS отключена.

Запустите команду **track rns**, чтобы настроить объект отслеживания, который используется для отслеживания результатов тестирования экземпляра RNS. Идентификатор экземпляра RNS находится в диапазоне от 1 до 500.

Если тест RNS прошел успешно, объект отслеживания находится в состоянии Up. Если проверка RNS не удалась, объект отслеживания находится в состоянии Down.



Настройка объекта отслеживания для отслеживания статуса Track List (списка отслеживания)

По умолчанию функция отслеживания состояния списка отслеживания отключена.

Запустите команду **track list**, чтобы настроить объект отслеживания, который используется для отслеживания состояния списка отслеживания. Результатом может быть результат операции AND или OR для статуса всех участников.

Если результатом этого объекта отслеживания является результат операции OR для статуса всех участников, то при успешном выполнении всех тестов RNS объект отслеживания находится в состоянии Up. Если один тест RNS не пройден, объект отслеживания находится в состоянии Down. Если результатом этого объекта отслеживания является результат операции OR для статуса всех участников, когда все тесты RNS не пройдены, объект отслеживания находится в состоянии Down. Если один тест RNS прошел успешно, объект отслеживания находится в состоянии Up.

Настройка участника списка отслеживания

По умолчанию для списка отслеживания не настроен ни один участник.

Запустите команду **object**, чтобы настроить участника списка отслеживания. Статус участника может быть таким же, как у соответствующего объекта отслеживания, или отличаться от него.

Настройка задержки для уведомления об изменении статуса объекта отслеживания

По умолчанию задержка уведомления об изменении состояния объекта отслеживания равна 0.

Запустите команду **delay**, чтобы настроить задержку уведомления отслеживания, включая задержку уведомления об изменении состояния объекта отслеживания с UP на DOWN и задержку уведомления об изменении состояния объекта отслеживания с DOWN на UP. Задержка колеблется от 0 до 180. Единицей измерения является секунда.

Более длительная задержка указывает на то, что требуется больше времени, прежде чем модуль, связанный с объектом отслеживания, будет уведомлен о статусе. Более короткая задержка указывает на то, что требуется меньше времени, прежде чем модуль, связанный с объектом отслеживания, будет уведомлен о статусе.

9.4. Конфигурация

Элемент конфигурации	Описание и команда	
<u>Настройка основных функций RNS</u>	(Обязательно) Используется для настройки основных функциональных параметров RNS	
	ip rns	Поддерживает подробную настройку и краткую настройку. Подробная конфигурация: объект операции RNS определяется и используется в качестве идентификатора конфигурации для последующих тестов и параметров.



Элемент конфигурации	Описание и команда	
		Краткая конфигурация: последующая настройка не требуется, и тесты можно запустить в один шаг. В настоящее время эхо-тесты ICMP, DNS и TCP можно запустить за один шаг
<u>Настройка основных функций RNS</u>	ip rns reaction-configuration	Настраивает упреждающий пороговый мониторинг и механизм запуска теста RNS
	ip rns reaction-trigger	Запускает другой тип теста RNS в состоянии ожидания, когда порог мониторинга превышает ожидаемое значение во время теста RNS
	ip rns schedule	Настраивает метод планирования, время начала и время жизни теста RNS
	ip rns restart	Перезапускает тест RNS
	ip rns reset	Очищает все конфигурации IP RNS
<u>Настройка эхо-теста ICMP</u>	(Опционально) Используется для реализации эхо-теста ICMP	
	icmp-echo	Создает экземпляр эхо-теста ICMP
	request-data-size	Настраивает размер полезной нагрузки протокола
	frequency	Настраивает интервал повторения теста
	tag	Настраивает тег
	threshold	Настраивает пороговое значение времени тестирования
	timeout	Настраивает время ожидания теста
	tos	Настраивает поле TOS в заголовке IPv4 тестовых пакетов
<u>Настройка теста DNS</u>	(Опционально) Используется для реализации теста DNS	
	dns	Создает тестовый экземпляр DNS



Элемент конфигурации	Описание и команда	
	frequency	Настраивает интервал повторения теста
	tag	Настраивает тег
<u>Настройка теста DNS</u>	threshold	Настраивает пороговое значение времени тестирования
	timeout	Настраивает время ожидания теста
	tos	Настраивает поле TOS в заголовке IPv4 тестовых пакетов
<u>Настройка поддержки отслеживания для RNS</u>	(Опционально) Используется для настройки поддержки отслеживания для других тестовых модулей	
	track rns	Настраивает объект отслеживания для отслеживания результатов теста экземпляра RNS
	track interface line-protocol	Настраивает объект отслеживания для отслеживания состояния канала интерфейса
	track list	Настраивает объект отслеживания для отслеживания состояния списка отслеживания
	object	Настраивает объект-участник для объекта списка отслеживания
	delay	Настраивает задержку для уведомления об изменении статуса объекта отслеживания

9.4.1. Настройка основных функций RNS

9.4.1.1. Эффект конфигурации

9.4.1.1.1. Подробная конфигурация: настраивает экземпляр RNS для завершения базовой настройки экземпляра RNS

Краткая конфигурация: настройка и запуск экземпляра RNS за один раз. (Опционально)

9.4.1.2. Примечания

- В режиме подробной конфигурации, если вы не настроите тип теста после входа в режим IP RNS, выполнив команду, экземпляр RNS не будет создан.



- В режиме подробной конфигурации после настройки экземпляра RNS необходимо выполнить команду **ip rns schedule** для настройки политики запуска; в противном случае тест не будет реализован.

9.4.1.3. Шаги настройки

Определение объекта операции RNS

- Обязательный.
- Если не требуется иное, определите объект операции RNS на каждом коммутаторе.
- Краткая конфигурация не является обязательной.

Настройка упреждающего порогового мониторинга и механизма запуска для теста RNS

- Выполните эту настройку, если требуется настроить упреждающий пороговый мониторинг и механизм запуска теста.
- Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.

Включение экземпляра RNS для запуска другого экземпляра RNS

- Выполните эту настройку, если требуется инициировать другой тест RNS в состоянии ожидания, когда пороговое значение мониторинга превышает ожидаемое во время теста RNS.
- Если параметры расписания не настроены для активированного экземпляра RNS, применяются параметры расписания по умолчанию.
- Если не требуется иное, примените эту конфигурацию к каждому коммутатору.

Настройка параметров расписания экземпляра RNS

- Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.
- В случае краткой настройки эта команда уже настроена с использованием значений по умолчанию, и ручная настройка не требуется.

Перезапуск экземпляра RNS

Выполните эту настройку или напрямую запустите команду **ip rns schedule X start-time now**, если требуется перезапустить экземпляр IP RNS в состоянии ожидания.

Очистка конфигураций всех экземпляров RNS

Выполните эту настройку, если требуется очистить конфигурации всех экземпляров IP RNS, например, когда настроено много экземпляров, но конфигурации признаны неправильными.

9.4.1.4. Проверка

Запустите команду **show ip rns configuration**, чтобы отобразить конфигурации экземпляров RNS.



9.4.1.5. Связанные команды

Определение объекта операции IP RNS

Команда	<code>ip rns operation-number</code>
Описание параметров	<i>operation-number</i> : указывает идентификатор экземпляра RNS. Диапазон значений от 1 до 500
Командный режим	Режим глобальной конфигурации
Руководство по использованию	В настоящее время RNS поддерживает только тесты, связанные с IPv4, но не тесты, связанные с IPv6. Можно настроить не более 500 тестов, в зависимости от производительности устройств. Функция тестирования — это только дополнительная функция. Когда настроено большое количество тестов, которые потребляют много системных ресурсов, функция тестирования может быть временно отключена, чтобы обеспечить нормальную работу основных служб, таких как переадресация маршрута. Подробная настройка (выполнение обязательных пунктов ip rns operation-number): запустите эту команду и войдите в режим настройки IP-RNS. В этом режиме вы можете определить различные типы тестов. Если тип теста не настроен, тест RNS не создается. После настройки теста RNS необходимо запустить команду ip rns schedule, чтобы настроить параметры его расписания; в противном случае тест не может быть проведен. После настройки типа теста RNS можно запустить команду ip rns для входа в режим типа теста. Чтобы изменить тип экземпляра RNS, необходимо сначала удалить экземпляр RNS, выполнив команду no ip rns в режиме глобальной конфигурации

Настройка механизма упреждающего порогового мониторинга и запуска теста

Команда	<code>ip rns reaction-configuration operation-number react monitored-element [action-type option][threshold-type {average [number-of-measurements] consecutive [occurrences] immediate never xofy [x-value y-value] }] [threshold-value upper-threshold lower-threshold]</code>
Описание параметров	<i>operation-number</i>: указывает идентификатор экземпляра RNS. Диапазон значений от 1 до 500. <i>monitored-element</i>: указывает контролируемый элемент. action-type option: указывает действие, предпринятое после запуска теста. average [number-of-measurements]: указывает, что последующие связанные действия запускаются, если среднее число измерений (<i>number-of-measurements</i>) отслеживаемого элемента превышает пороговое значение.



	consecutive [<i>occurrences</i>]: указывает, что тест запускается, если количество последовательных вхождений (<i>occurrences</i>) отслеживаемого элемента превышает пороговое значение. Значение по умолчанию для вхождений — 5. Диапазон значений — от 1 до 16. immediate: указывает, что тест запускается сразу после того, как контролируемый элемент превышает пороговое значение. never: указывает, что тест никогда не запускается. xofy [<i>x-value y-value</i>]: указывает, что результаты тестов X превышают пороговое значение в последних тестах Y. Значения X и Y по умолчанию равны 5. Значение X или Y находится в диапазоне от 1 до 16. threshold-value <i>upper-threshold lower-threshold</i>: укажите верхний и нижний пороги. • Когда <i>monitored-element</i> в состоянии rtt, пороговые значения — это время. Диапазон значений от 0 до 60 000 мс. • Обратите внимание, что вам не нужно настраивать threshold-value когда react установлено на timeout																		
Командный режим	Режим глобальной конфигурации																		
Руководство по использованию	Вы можете настроить несколько пороговых значений для одного теста RNS, чтобы отслеживать разные элементы. В следующей таблице показано сопоставление между типами тестов и контролируемыми элементами. <table><tr><td>monitored-element</td><td>icmp-echo</td><td>dns</td></tr><tr><td>timeout</td><td>0</td><td>0</td></tr><tr><td>rtt</td><td>0</td><td>0</td></tr></table> В следующей таблице перечислены пороговые значения по умолчанию для каждого отслеживаемого элемента. <table><tr><td>Monitored Element</td><td>Upper Threshold</td><td>Lower Threshold</td></tr><tr><td>timeout</td><td>-</td><td>-</td></tr><tr><td>rtt</td><td>5000 мс</td><td>0 мс</td></tr></table>	monitored-element	icmp-echo	dns	timeout	0	0	rtt	0	0	Monitored Element	Upper Threshold	Lower Threshold	timeout	-	-	rtt	5000 мс	0 мс
monitored-element	icmp-echo	dns																	
timeout	0	0																	
rtt	0	0																	
Monitored Element	Upper Threshold	Lower Threshold																	
timeout	-	-																	
rtt	5000 мс	0 мс																	

Включение экземпляра RNS для триггера (запуска) другого экземпляра RNS

Команда	ip rns reaction-trigger operation-number target-operation
Описание параметров	operation-number: указывает номер исходного экземпляра RNS, запускающего действие. Диапазон значений от 1 до 500. target-operation: указывает номер запущенного целевого экземпляра RNS. Диапазон значений от 1 до 500



Командный режим	Режим глобальной конфигурации
Руководство по использованию	Функция триггера обычно используется в сценарии диагностики неисправностей сети. В обычном сценарии вам не нужно настраивать функцию триггера

Настройка параметров расписания экземпляра RNS

Команда	ip rns schedule <i>operation-number</i> [life { forever <i>seconds</i> }] [start-time { <i>hh:mm[:ss]</i> [<i>month day</i> <i>daymonth</i>] pending now after <i>hh:mm:ss</i> }] [recurring]
Описание параметров	<i>operation-number</i>: указывает номер операции RNS. Диапазон значений от 1 до 500. life forever: указывает, что время жизни операции RNS действует вечно. life seconds: указывает время работы экземпляра RNS в секундах. <i>hh:mm[:ss]</i>: указывает время запуска экземпляра RNS в 24-часовом формате. <i>month</i>: указывает начальный месяц экземпляра RNS. Значение по умолчанию — текущий месяц. <i>day</i>: указывает дату начала экземпляра RNS. Значением по умолчанию является текущая дата. pending: указывает, что время запуска экземпляра RNS не определено, что является значением по умолчанию. now: указывает, что время начала операции сейчас, то есть операция начинается сейчас. after hh:mm:ss: указывает, что экземпляр RNS запускается после задержки чч:мм:сс. recurring: указывает, запускается ли экземпляр RNS каждый день в одно и то же время
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Если параметры расписания экземпляра RNS были настроены с помощью команды ip rns schedule, параметры нельзя изменить во время работы. Чтобы изменить конфигурацию, вам нужно запустить команду no ip rns schedule, чтобы удалить параметры расписания. life { seconds } указывает время работы экземпляра RNS. То есть тест останавливается через промежуток времени в секундах



Перезапуск теста RNS с помощью команды `ip rns restart`

Команда	<code>ip rns restart operation-number</code>
Описание параметров	<i>operation-number</i> : указывает номер экземпляра RNS. Диапазон значений от 1 до 500
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Эта команда перезапускает тест RNS, для которого настроена политика планирования и который находится в состоянии ожидания. Эта команда недопустима для теста RNS, для которого не настроена политика планирования

Очистка конфигураций всех экземпляров IP RNS с помощью команды `ip rns reset`

Команда	<code>ip rns reset</code>
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Эта команда очищает конфигурации всех экземпляров IP RNS. Используется только в крайних случаях, например, когда настроено много RNS-тестов, но конфигурации оказываются неверными

9.4.1.6. Пример конфигурации

Настройка основных функций RNS

Сценарий:



Рисунок 9-3.

Шаги настройки	• Настройте экземпляр 1 на Коммутаторе А. • Настройте метод планирования, время начала и время жизни экземпляра 1. • Настройте упреждающий пороговый мониторинг и механизм запуска экземпляра 1. • Активировать экземпляр 2 в состоянии ожидания, когда пороговое значение мониторинга экземпляра 1 превышает ожидаемое
----------------	--



Коммутатор A	<pre> A# configure terminal A(config)# ip rns 1 A(config-ip-rns)#icmp-echo 10.1.1.1 A(config-ip-rns-icmp-echo)#exit A(config)ip rns schedule 1 start-time now life forever A(config)ip rns reaction-configuration 1 react timeout threshold-type immediate action-type trigger A(config)ip rns reaction-trigger 1 2 </pre>
Проверка	Запустите команду show ip rns configuration, чтобы отобразить конфигурации экземпляра. <pre> Router#show ip rns configuration 1 Entry number: 1 Tag: QTECH 555 Type of operation to perform: icmp-echo Operation timeout (milliseconds): 5000 Operation frequency (milliseconds): 60000 Threshold (milliseconds): 5000 Recurring (Starting Everyday): FALSE Life (seconds): 3500 Next Scheduled Start Time:Start Time already passed Target address/Source address: 2.2.2.3/0.0.0.0 Request size (ARR data portion): 36 </pre>

9.4.2. Настройка эхо-теста ICMP

9.4.2.1. Эффект конфигурации

Создает экземпляр эхо-теста ICMP.

9.4.2.2. Примечания

Основные функции RNS должны быть настроены.

9.4.2.3. Шаги настройки

Создание экземпляра эхо-теста ICMP

- Обязательный.
- Если не требуется иное, создайте экземпляры эхо-теста ICMP на каждом коммутаторе.

Настройка общих необязательных параметров теста

- Обязательно, если требуется изменить общие необязательные параметры теста, например, интервал повторения, тег, порог времени, таймаут и TOS.

- Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.

Настройка размера полезной нагрузки протокола

- Выполните эту настройку, если требуется изменить размер полезной нагрузки протокола теста.
- Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.

9.4.2.4. Проверка

Запустите команду конфигурации **show ip rns**, чтобы отобразить конфигурации экземпляра.

9.4.2.5. Связанные команды

Создание экземпляра эхо-теста ICMP

Команда	icmp-echo { oob { <i>destination-ip-address</i> <i>destination-hostname</i> [name-server <i>ip-address</i>] } [source-ipaddr <i>ip-address</i>] via <i>type num</i> next-hop <i>ip-address</i> } { { <i>destination-ip-address</i> <i>destination-hostname</i> [name-server <i>ip-address</i>] } [source-ipaddr <i>ip-address</i>] [out-interface <i>type num</i> [next-hop <i>ip-address</i>]] }
Описание параметров	oob: указывает на тест на интерфейсе MGMT. <i>destination-ip-address</i>: указывает IP-адрес назначения. <i>destination-hostname</i>: указывает имя хоста назначения. name-server <i>ip-address</i>: указывает DNS-сервер при настройке имени хоста назначения. По умолчанию DNS-сервер настроен с использованием команды ip name-server и используется для разрешения адреса. source-ipaddr <i>ip-address</i>: указывает исходный IP-адрес. out-interface <i>type num</i>: определяет исходящий интерфейс (не интерфейс MGMT) тестового пакета. via <i>type num</i>: указывает интерфейс MGMT в качестве исходящего интерфейса тестового пакета. next-hop <i>A.B.C.D</i>: указывает IP-адрес next-hop
Командный режим	Режим конфигурации IP RNS (config-ip-rns)
Руководство по использованию	После запуска эхо-теста ICMP система отправляет пакет эхо-запроса ICMP, чтобы проверить, подключено ли устройство к целевому хосту. После создания тестового экземпляра ICMP-Echo система переходит в режим эха IP RNS ICMP. По умолчанию размер полезной нагрузки протокола пакета эхо-запроса ICMP составляет 36 байт. Вы можете запустить команду request-data-size , чтобы изменить размер пакета. Перед настройкой параметров необходимо настроить тип теста RNS (например, ICMP-эхо и DNS). Чтобы изменить тип экземпляра RNS,



необходимо удалить экземпляр RNS, выполнив команду **no ip rns** в режиме глобальной конфигурации

Настройка размера полезной нагрузки протокола экземпляра RNS

Команда	request-data-size bytes
Описание параметров	<i>bytes</i> : указывает байты тестового пакета. Минимальные и максимальные байты зависят от типа теста. Вам необходимо настроить этот параметр на основе командной строки в соответствующем тестовом режиме
Командный режим	Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	Эта команда используется для заполнения некоторых байтов в тестовом пакете, чтобы для теста можно было использовать большие пакеты

Настройка интервала повторения теста

Команда	frequency milliseconds
Описание параметров	<i>milliseconds</i> : указывает интервал отправки пакетов в мс. Значение по умолчанию — 60 000 мс. Значение варьируется от 10 до 604 800 000. Максимальное значение — одна неделя
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	После запуска экземпляра RNS периодически проводятся тесты. Вы можете запустить команду frequency , чтобы указать интервал повторения. Вам необходимо настроить частоту на основе следующей формулы, чтобы обеспечить правильный расчет теста. (frequency milliseconds) > (timeout milliseconds) >= (threshold milliseconds)

Настройка тега для экземпляра RNS

Команда	tag text
Описание параметров	<i>text</i> : устанавливает тег теста. Значение представляет собой строку длиной до 79 символов
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)



Руководство по использованию	Эта команда задает тег для теста, который часто используется для обозначения функции теста
------------------------------	--

Настройка временного порога для экземпляра RNS

Команда	threshold <i>milliseconds</i>
Описание параметров	<i>milliseconds</i> : указывает временной порог для теста. Значение находится в диапазоне от 0 до 60 000 в миллисекундах. Значение по умолчанию — 5000
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	Настройте пороговое значение на основе следующей формулы, чтобы обеспечить правильный расчет теста. $(\text{frequency } \textit{milliseconds}) > (\text{timeout } \textit{milliseconds}) \geq (\text{threshold } \textit{milliseconds})$

Настройка времени ожидания (тайм-аута) для экземпляра RNS

Команда	timeout <i>millisecond</i>
Описание параметров	<i>millisecond</i> : указывает время ожидания теста. Значение варьируется от 10 до 604 800 000. Единица измерения — мс. Тайм-аут по умолчанию зависит от типа теста
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	Настройте время ожидания на основе следующей формулы, чтобы обеспечить правильный расчет теста. $(\text{frequency } \textit{milliseconds}) > (\text{timeout } \textit{milliseconds}) \geq (\text{threshold } \textit{milliseconds})$

Настройка поля TOS в заголовке пакета IPv4 теста IP RNS

Команда	tos <i>number</i>
Описание параметров	<i>number</i> : устанавливает поле TOS в заголовке IPv4 тестовых пакетов. Значение находится в диапазоне от 0 до 255. Значение по умолчанию — 0
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)



Руководство по использованию	TOS — это 8-битное поле в заголовке пакета IPv4. Установив TOS, вы можете управлять приоритетом тестового пакета. Для разных полей TOS приоритеты обработки на промежуточных маршрутизаторах различаются
------------------------------	--

Настройка VRF теста RNS

Команда	<code>vrf vrf-name</code>
Описание параметров	<code>vrf-name</code> : указывает имя VRF
Командный режим	Режим конфигурации IP RNS DNS (<code>config-ip-rns-dns</code>) Режим конфигурации эха IP RNS ICMP (<code>config-ip-rns-icmp-echo</code>)
Руководство по использованию	Эта команда указывает VRF тестового пакета

Пример конфигурации



Рисунок 9-4.

	Настройте экземпляр RNS 1 и связанные параметры на Коммутаторе А
Коммутатор А	<code>A# configure terminal</code> <code>A(config)# ip rns 1</code> <code>A(config-ip-rns)#icmp-echo 10.2.2.2</code> <code>A(config-ip-rns-icmp-echo)#exit</code> <code>A(config)#ip rns schedule 1 start-time now life forever</code>
	Запустите команду конфигурации show ip rns , чтобы отобразить конфигурации экземпляра
Коммутатор А	<code>A#show ip rns configuration 1</code> Entry number: 1 Tag: Type of operation to perform: icmp-echo Operation timeout (milliseconds): 5000



	Operation frequency (milliseconds): 60000 Threshold (milliseconds): 5000 Recurring (Starting Everyday): FALSE Life (seconds): forever Next Scheduled Start Time: Start Time already passed Target address/Source address: 10.2.2.2/0.0.0.0 Request size (ARR data portion): 36
--	--

9.4.3. Настройка теста DNS

9.4.3.1. Эффект конфигурации

Создает тестовый экземпляр DNS.

9.4.3.2. Примечания

Основные функции RNS должны быть настроены.

9.4.3.3. Шаги настройки

Создание тестового экземпляра DNS

- Обязательный.
- Если не требуется иное, создайте тестовые экземпляры DNS на каждом коммутаторе.

Настройка общих необязательных параметров теста

- Обязательный, если требуется изменить общие необязательные параметры теста, например, интервал повторения, тег, порог времени, тайм-аут и TOS.
- Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.

9.4.3.4. Проверка

Запустите команду **show ip rns configuration**, чтобы отобразить конфигурации экземпляра.

9.4.3.5. Связанные команды

Создание тестового экземпляра DNS

Команда	dns { oob destination-hostname name-server ip-address }
Описание параметров	oob : указывает на тест на интерфейсе MGMT. destination-hostname : указывает имя хоста назначения. name-server ip-address : указывает IP-адрес DNS
Командный режим	Режим конфигурации IP RNS (config-ip-rns)



Руководство по использованию	После запуска теста DNS система отправляет пакет запроса анализа DNS, чтобы проверить, подключено ли устройство к целевому хосту. После создания тестового экземпляра DNS система переходит в режим IP RNS DNS. Перед настройкой параметров необходимо настроить тип теста RNS. Чтобы изменить тип экземпляра RNS, необходимо удалить экземпляры RNS, выполнив команду no ip rns в режиме глобальной конфигурации
------------------------------	--

Настройка интервала повторения теста

Команда	frequency <i>milliseconds</i>
Описание параметров	<i>milliseconds</i> : указывает интервал отправки пакетов в мс. Значение по умолчанию — 60 000 мс. Значение варьируется от 10 до 604 800 000. Максимальное значение — одна неделя
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	После запуска экземпляра RNS периодически проводятся тесты. Вы можете запустить команду frequency, чтобы указать интервал повторения. Вам необходимо настроить частоту на основе следующей формулы, чтобы обеспечить правильный расчет теста. (frequency milliseconds) > (timeout milliseconds) >= (threshold milliseconds)

Настройка тега для экземпляра RNS

Команда	tag <i>text</i>
Описание параметров	<i>text</i> : устанавливает тестовый тег. Значение представляет собой строку длиной до 79 символов
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	Эта команда задает тег для теста, который часто используется для обозначения функции теста

Настройка временного порога для экземпляра RNS

Команда	threshold <i>milliseconds</i>
Описание параметров	<i>milliseconds</i> : указывает временной порог для теста. Значение находится в диапазоне от 0 до 60 000 в миллисекундах. Значение по умолчанию — 5000



Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	Настройте пороговое значение на основе следующей формулы, чтобы обеспечить правильный расчет теста. $(\text{frequency milliseconds}) > (\text{timeout milliseconds}) \geq (\text{threshold milliseconds})$

Настройка временного порога для экземпляра RNS

Команда	<code>timeout millisecond</code>
Описание параметров	<i>millisecond</i> : указывает время ожидания теста. Значение варьируется от 10 до 604 800 000. Единица измерения — мс. Тайм-аут по умолчанию зависит от типа теста
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	Настройте время ожидания (тайм-аут) на основе следующей формулы, чтобы обеспечить правильный расчет теста. $(\text{frequency milliseconds}) > (\text{timeout milliseconds}) \geq (\text{threshold milliseconds})$

Настраивает поле TOS в заголовке IPv4 тестовых пакетов.

Команда	<code>tos number</code>
Описание параметров	<i>number</i> : устанавливает поле TOS в заголовке IPv4 тестовых пакетов. Значение находится в диапазоне от 0 до 255. Значение по умолчанию — 0
Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	TOS — это 8-битное поле в заголовке пакета IPv4. Установив TOS, вы можете управлять приоритетом тестового пакета. Для разных полей TOS приоритеты обработки на промежуточных маршрутизаторах различаются

Настройка VRF теста RNS

Команда	<code>vrf vrf-name</code>
Описание параметров	<i>vrf-name</i> : указывает имя VRF



Командный режим	Режим конфигурации IP RNS DNS (config-ip-rns-dns) Режим конфигурации эха IP RNS ICMP (config-ip-rns-icmp-echo)
Руководство по использованию	Эта команда указывает VRF тестового пакета

9.4.3.6. Пример конфигурации

Сценарий:



Рисунок 9-5.

Шаги настройки	Настройте экземпляр RNS 1 и связанные параметры на Коммутаторе А
Коммутатор А	A# configure terminal A(config)# ip rns 1 A(config-ip-rns)# dns www.QTECH.com name-server 10.2.2.2 A(config-ip-rns-dns)#exit A(config)ip rns schedule 1 start-time now life forever
Проверка	Запустите команду конфигурации show ip rns , чтобы отобразить конфигурации экземпляра
Коммутатор А	A#show ip rns configuration 1 Entry number: 1 Tag: Type of operation to perform: dns Operation timeout (milliseconds): 5000 Operation frequency (milliseconds): 60000 Threshold (milliseconds): 5000 Recurring (Starting Everyday): FALSE Life (seconds): forever Next Scheduled Start Time:Start Time already passed Target host name: www.QTECH.com Name Server: 10.2.2.2



9.4.3.7. Распространенные ошибки

Неверный IP-адрес DNS.

9.4.4. Настройка поддержки отслеживания для RNS

9.4.4.1. Эффект конфигурации

- Настройте функцию отслеживания для отслеживания результатов тестирования экземпляра RNS.
- Настройте функцию отслеживания для отслеживания состояния канала интерфейса.
- Настройте функцию отслеживания для отслеживания состояния списка отслеживания.
- Настройте функцию отслеживания для отслеживания состояния списка RNS.

9.4.4.2. Примечания

- Чтобы настроить функцию отслеживания для отслеживания результатов тестирования экземпляра RNS, необходимо настроить соответствующий экземпляр RNS.
- Чтобы настроить функцию отслеживания для отслеживания состояния канала интерфейса, необходимо настроить соответствующий интерфейс.
- Чтобы настроить функцию отслеживания для отслеживания состояния списка отслеживания, необходимо настроить участников для связанного списка отслеживания.
- Чтобы настроить функцию отслеживания для отслеживания состояния списка RNS, необходимо настроить участников для связанного списка RNS.

9.4.4.3. Шаги настройки

Настройка объекта трека

- Выполните эту операцию, если требуется создать объект отслеживания.
- Для создания объекта отслеживания доступны следующие четыре метода:
 1. Создайте объект отслеживания для отслеживания результатов тестирования экземпляра RNS: Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.
 2. Создайте объект отслеживания для отслеживания состояния канала интерфейса: выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.
 3. Создайте объект отслеживания для отслеживания состояния списка отслеживания: выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.
 4. Создайте объект отслеживания для отслеживания состояния списка RNS: выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.

Настройка задержки уведомления объекта отслеживания

- Выполните эту настройку, если требуется отложить уведомление об изменении статуса объекта отслеживания.
- Задержка уведомления об изменении состояния объекта пути отслеживания включает в себя задержку уведомления об изменении состояния объекта



отслеживания с UP на DOWN и задержку уведомления об изменении состояния объекта пути с DOWN на UP. Вы можете настроить либо задержку, либо обе задержки.

- Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.

Настройка элемента отслеживания

- Выполните эту настройку, если требуется настроить объект отслеживания для отслеживания статуса списка отслеживания.
- При настройке элемента отслеживания вы можете установить статус условий выполнения элемента на UP или DOWN.
- Выполните эту настройку на каждом коммутационном устройстве, если не требуется иное.

9.4.4.4. Проверка

Наблюдайте за состоянием объекта отслеживания, когда состояние объекта отслеживания (например, результаты тестирования экземпляра RNS, состояние канала интерфейса или состояние списка отслеживания) изменяется.

- После предустановленной задержки запустите команду **show track**, чтобы проверить, изменяется ли состояние текущего отслеживания.

9.4.4.5. Связанные команды

Настройка объекта отслеживания для отслеживания статуса канала интерфейса

Команда	track <i>object-number</i> interface <i>interface-type</i> <i>interface-number</i> line-protocol
Описание параметров	<i>object-number</i> : указывает номер объекта отслеживания. Диапазон значений от 1 до 700. <i>interface-type interface-number</i> : указывает тип и номер интерфейса
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Запустите эту команду, чтобы настроить объект отслеживания для отслеживания состояния канала интерфейса. Когда статус канала интерфейса — UP, статус соответствующего объекта отслеживания — UP

Настройка объекта отслеживания для отслеживания результатов теста RNS

Команда	track <i>object-number</i> rns <i>entry-number</i>
Описание параметров	<i>object-number</i> : указывает номер объекта отслеживания. Диапазон значений от 1 до 700.



	<i>entry-number</i> : указывает номер экземпляра RNS. Диапазон значений от 1 до 500
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Запустите эту команду, чтобы настроить объект отслеживания для отслеживания результатов теста RNS. Если проверка прошла успешно, объект отслеживания находится в состоянии Up

Настройка объекта отслеживания для отслеживания статуса списка отслеживания

Команда	track object-number listboolean { and or }
Описание параметров	<i>object-number</i> : указывает номер объекта отслеживания. Диапазон значений от 1 до 700
Командный режим	Режим глобальной конфигурации
Руководство по использованию	Запустите эту команду, чтобы настроить объект отслеживания для отслеживания состояния списка отслеживания. Результатом может быть результат операции AND или OR для всех статусов участников

Настройка участника отслеживания

Команда	object object-number [not]
Описание параметров	<i>object-number</i> : указывает номер объекта отслеживания. Диапазон значений от 1 до 700
Командный режим	Режим конфигурации отслеживания
Руководство по использованию	Запустите эту команду, чтобы настроить участника для списка отслеживания. Количество участников списка отслеживания, которые можно настроить, ограничено только емкостью объектов отслеживания

Настройка задержки уведомления объекта отслеживания

Команда	delay { up seconds [down seconds] [up seconds] down seconds }
Описание параметров	up seconds : определяет задержку для уведомления об изменении состояния объекта отслеживания с DOWN на UP. Значение находится в диапазоне от 0 до 180. Единицей измерения является секунда. Значение по умолчанию — 0.



	down seconds: задает задержку для уведомления об изменении состояния объекта отслеживания с UP на DOWN. Значение находится в диапазоне от 0 до 180. Единицей измерения является секунда. Значение по умолчанию — 0
Командный режим	Режим конфигурации отслеживания
Руководство по использованию	Когда статус объекта отслеживания часто меняется, статус клиента, использующего этот объект отслеживания, также будет часто меняться. Использование этой команды может задержать уведомление об изменении статуса объекта отслеживания. Например, если статус объекта отслеживания изменяется с UP на DOWN и настроена задержка down 10, статус DOWN объекта отслеживания уведомляется через 10 секунд. Если в течение этого времени статус объектов отслеживания снова меняется на UP, уведомление не отправляется. Для клиента, который использует этот объект отслеживания, статус объекта отслеживания всегда UP

Отображение статистики объекта отслеживания

Команда	show track [object-number]
Описание параметров	<i>object-number:</i> указывает номер объекта отслеживания. Диапазон значений от 1 до 700. По умолчанию все объекты отслеживаются
Командный режим	Привилегированный режим EXEC
Руководство по использованию	Запустите эту команду, чтобы отобразить статистику объектов отслеживания

9.4.4.6. Пример конфигурации

Настройка отслеживания объекта 3 для отслеживания состояния канала интерфейса FastEthernet 1/0

Шаги настройки	- Настройте объект отслеживания для отслеживания состояния канала интерфейса. - Настройте задержку уведомления об изменении состояния с UP на DOWN
	<pre> QTECH # configure terminal QTECH (config)# track 3 interface FastEthernet 1/0 line-protocol QTECH (config-track)# delay down 10 QTECH (config-track)# exit </pre>



Проверка	Измените статус канала интерфейса FastEthernet 1/0 на DOWN. • Немедленно проверьте статус объекта отслеживания и убедитесь, что статус все еще UP. • Через 10 секунд проверьте состояние объекта отслеживания и убедитесь, что состояние изменилось на DOWN
	<pre>QTECH # show track 3 Track 3 Interface FastEthernet 1/0 The state is Up, delayed Down (5 secs remaining) 1 change, current state last: 300 secs Delay up 0 secs, down 10 secs</pre>

Настройка объекта отслеживания 3 (когда статус объекта отслеживания 1 является UP, а статус объекта отслеживания 2 — DOWN, статус объекта отслеживания 3 — UP.)

Шаги настройки	• Настройка объекта отслеживания 1 и объекта отслеживания 2. • Настройте объект отслеживания 3, и его элементы включают объект отслеживания 1 и объект отслеживания 2
	<pre>QTECH # config QTECH (config)#track 1 interface gigabitEthernet 0/0 line-protocol QTECH (config-track)#delay up 20 down 40 QTECH (config-track)#exit QTECH (config)# QTECH (config)#track 2 interface gigabitEthernet 0/1 line-protocol QTECH (config-track)#delay down 30 QTECH (config-track)#exit QTECH (config)# track 3 list Boolean and QTECH (config-track)#object 1 QTECH (config-track)#object 2 not QTECH (config-track)# exit</pre>
Проверка	При изменении состояния объектов отслеживания 1 и 2 проверьте состояние объекта отслеживания 3. • Когда статус объекта отслеживания 1 изменится с DOWN на UP, а статус объекта отслеживания 2 останется DOWN, убедитесь, что статус объекта отслеживания 3 изменился с DOWN на UP. • Когда статус объекта отслеживания 1 остается UP, а статус объекта отслеживания 2 меняется с DOWN на UP, убедитесь, что статус объекта отслеживания 3 меняется с UP на DOWN



	QTECH # show track 3 Track 3 List boolean and Object 1 Object 2 not The state is Down 1 change,current state last:10 secs Delay up 0 secs,down 0 secs
--	---

Настройка объекта отслеживания 5 для отслеживания результатов тестирования экземпляра RNS 7

Шаги настройки	• Настройте тест RNS. • Настройте объект отслеживания для отслеживания результатов теста RNS. • Настройте задержку уведомления об изменении результата теста с успешного на неуспешный и задержку уведомления об изменении результата теста с неуспешного на успешный
	QTECH # configure terminal QTECH (config)#ip rns 7 QTECH (config-ip-rns)#icmp-echo 2.2.2.2 QTECH (config-ip-rns-icmp-echo)#exit QTECH (config)#ip rns schedule 7 start-time now life forever QTECH (config)# track 5 rns 7 QTECH (config-track)# delay up 20 down 30 QTECH (config-track)# exit
Проверка	Пусть результат проверки экземпляра RNS 7 изменится с успешного на неуспешный. • Когда результат проверки изменится на неуспешный, немедленно проверьте состояние объекта отслеживания 7 и убедитесь, что состояние все еще UP. • Через 30 секунд проверьте статус объекта отслеживания и убедитесь, что статус изменился на DOWN
	QTECH # show track 5 Track 5 Reliable Network Service 7 The state is Down



	2 change, current state last: 10 secs Delay up 20 secs, down 30 secs
--	---

Настройка объекта отслеживания 5 для отслеживания результатов тестирования списка RNS (состоящего из экземпляров RNS 1, 2–5 и 8)

Шаги настройки	- Настройте и запустите тест RNS (см. Настройка RNS). - Настройте объект отслеживания для отслеживания результатов теста списка RNS. - Настройте задержку уведомления об изменении результата теста с UP на DOWN и задержку уведомления об изменении результата теста с DOWN на UP
	<pre>QTECH (config)# track 5 rns-list 1,2-5,8 and QTECH (config-track)# delay up 20 down 30 QTECH (config-track)# exit</pre>
Проверка	Пусть результат проверки одного из экземпляров RNS 1, 2-5 и 8 изменится с успешного на неуспешный. - Когда результат проверки изменится на неуспешный, немедленно проверьте состояние объекта отслеживания 7 и убедитесь, что состояние все еще UP. - Через 30 секунд проверьте статус объекта отслеживания и убедитесь, что статус изменился на DOWN
	<pre>QTECH # show track 5 Track 5 rns-list 1,2-5,8 and The state is Down 2 change, current state last: 10 secs Delay up 20 secs, down 30 secs</pre>

9.4.4.7. Распространенные ошибки

- Объект отслеживания для отслеживания теста RNS настроен, но тест RNS не настроен.
- Объект отслеживания для отслеживания статуса канала интерфейса настроен, но соответствующий интерфейс не настроен.
- Объект отслеживания для отслеживания состояния списка отслеживания, но ни один участник списка RNS не настроен.
- Объект отслеживания настроен для отслеживания списка RNS, но тест RNS не настроен.



9.5. Мониторинг

9.5.1. Отображение

Описание	Команда
Отображает конфигурации одного или нескольких экземпляров RNS	show ip rns configuration [<i>operation-number</i>]
Отображает подробную статистику по одному или нескольким экземплярам RNS	show ip rns collection-statistics [<i>operation-number</i>]
Отображает текущее состояние RNS	show ip rns operational-state [<i>operation-number</i>]
Отображает информацию об упреждающем пороговом мониторинге одного или нескольких Экземпляров RNS	show ip rns reaction-configuration [<i>operation-number</i>]
Отображает информацию о тесте, запущенном одним или несколькими экземплярами RNS	show ip rns reaction-trigger [<i>operation-number</i>]
Отображает краткую статистику по одному или нескольким экземплярам RNS	show ip rns statistics [<i>operation-number</i>]
Отображает краткую статистику по одному или нескольким объектам отслеживания	show track [<i>object-number</i>]
Отображает краткую статистику о клиенте отслеживания	show track client

9.5.1.1. Отладка

ПРИМЕЧАНИЕ: системные ресурсы заняты при выводе отладочной информации. Поэтому отключайте отладку сразу после использования.

Описание	Команда
Отладка модуля отслеживания	debug track { all proc-event rdnd-event client }



Описание	Команда
Отладка модуля RNS	<code>debug rns { all interface lib rdnd-event restart rns_id [0, 500] }</code>



10. ОБЩАЯ ИНФОРМАЦИЯ

10.1. Гарантия и сервис

Процедура и необходимые действия по вопросам гарантии описаны на сайте QTECH в разделе «Поддержка» -> «[Гарантийное обслуживание](#)».

Ознакомиться с информацией по вопросам тестирования оборудования можно на сайте QTECH в разделе «Поддержка» -> «[Взять оборудование на тест](#)».

Вы можете написать напрямую в службу сервиса по электронной почте sc@qtech.ru.

10.2. Техническая поддержка

Если вам необходимо содействие в вопросах, касающихся нашего оборудования, то можете воспользоваться разделом технической поддержки пользователей QTECH на нашем сайте www.qtech.ru/support/.

Телефон Технической поддержки +7 (495) 269-08-81

Центральный офис +7 (495) 477-81-18

10.3. Электронная версия документа

Дата публикации 25.07.2025



https://files.qtech.ru/upload/switchers/QSW-8113/QSW-8113_reliability_config_guide.pdf