

OSNOVO

cable transmission

РУКОВОДСТВО ПО ЭКСПЛУАТАЦИИ

Управление коммутатором с помощью
интерфейса командной строки (CLI)

Прежде чем приступить к эксплуатации изделия,
внимательно прочтите настоящее руководство

www.osnovo.ru

Оглавление

1. CLI Command line interface (Введение).....	26
1.1 Accessing CLI (Доступ к CLI).....	26
1.1.1 Users access CLI via Console port (Подключение через порт Console)	26
1.1.2 Users access CLI via TELNET (Подключение через TELNET).....	27
1.2 CLI Mode introduction (CLI режимы).....	28
1.2.1 CLI Role of mode (CLI роли режимов).....	28
1.2.2 Identification of CLI mode (Идентификация CLI режимов)	29
1.2.3 Classification of CLI modes (Классификация CLI режимов).....	29
1.3 Introduction to command syntax (Синтаксис команд).....	31
1.3.1 Command composition (Композиция команд).....	31
1.3.2 Parameter type (Типы параметров).....	32
1.3.3 Command syntax rules (Правила синтаксиса).....	32
1.3.4 Command abbreviation (Сокращения).....	33
1.3.5 Grammar help (Помощь)	33
1.3.6 Command line error message (Сообщения об ошибках).....	34
1.4 Command line shortcuts (Короткие команды)	35
1.4.1 Row edit shortcut (Короткие команды ввода).....	35
1.4.2 Show command shortcuts (Короткие команды show)	36
1.5 Historical command (История команд).....	36
2. System management configuration (Конфигурация управление системой)	36
2.1 System security configuration (Конфигурация системы безопасности)	37
2.1.1 Multi user management control (Многопользовательский режим) 37	
2.1.2 Enable password control (Управление аутентификацией)	39
2.1.3 TELNET service control (Управление TELNET сервисом).....	40
2.1.4 SNMP Service control (Управление SNMP сервисом)	40
2.1.5 HTTP Service control (Управление HTTP сервисом)	41
2.1.6 HTTPS Service control (Управление HTTPS сервисом).....	42

2.1.7 SSH Service control (Управление SSH сервисом)	43
2.2 System maintenance and commissioning (Обслуживание системы)....	44
2.2.1 Configure the system clock (Настройка системных часов)	44
2.2.2 Configure terminal timeout attribute (Настройка времени бездействия)	45
2.2.3 System reset (Сброс системы)	45
2.2.4 View system information (Просмотр информации о системе)	46
2.2.5 Network connectivity debugging (Устранение сетевых неполадок)	46
2.2.6 Traceroute debugging (Устранение неполадок маршрутизации) ..	47
2.2.7 Telnet Client (Telnet клиент)	47
2.2.8 SSH Client (SSH клиент)	48
2.3 Profile management (Управление профилем).....	48
2.3.1 View configuration information (Просмотр конфигурации).....	49
2.3.2 Save configuration (Сохранение конфигурации)	50
2.3.3 Delete configuration file (Удаление файла конфигурации)	50
2.3.4 Download configuration file (Загрузка файла конфигурации).....	51
2.4 Software version upgrade (Обновление ПО)	53
2.4.1 Software version upgrade command (Команды обновления ПО) ..	53
2.4.2 Software upgrade process (Процесс обновления ПО).....	54
3. Configuration of ports (Конфигурация портов).....	55
3.1 General configuration (Основные настройки портов)	56
3.1.1 Opening and closing of ports (Открытие и закрытие портов)	56
3.1.2 Port rate configuration (Настройка скорости портов)	56
3.1.3 Display port information (Информация о порте)	57
3.2 Configure mirror (Зеркалирование портов)	57
3.2.1 Configure listening port and monitored port (Настройка портов) ...	57
3.2.2 Display configuration of mirror (Просмотр конфигурации).....	58
3.3 Configure storm-control (Подавление «шторма»)	58
3.3.1 Default configuration (Конфигурация по умолчанию).....	59
3.3.2 Broadcast suppression configuration (Broadcast подавление)	59

3.3.3 Multicast suppression configuration (Multicast подавление)	59
3.3.4 DLF suppression configuration (Multicast подавление).....	59
3.3.5 Suppression rate configuration (Уровни подавления)	60
3.3.6 Display storm-control configuration (Просмотр конфигурации)	60
3.4 Configure storm-constraint (Настройка storm-constraint).....	60
3.5 Configure flow-control (Настройка flow-control)	63
3.5.1 Default configuration (Конфигурация по умолчанию).....	63
3.5.2 Set port flow control (Включение flow control)	63
3.5.3 Close port flow control (Отключение flow control)	63
3.5.4 Display flow control information (Просмотр конфигурации).....	64
3.6 Configure port bandwidth (Настройка скорости портов).....	64
3.6.1 Default configuration (Конфигурация по умолчанию).....	64
3.6.2 Set port bandwidth (Включение port bandwidth)	64
3.6.3 Cancel port bandwidth (Отключение port bandwidth)	65
3.6.4 Display bandwidth control of port configuration (Просмотр конфигурации)	65
3.7 Configure trunk (Конфигурация trunk)	65
3.7.1 LACP protocol configuration (Настройка протокола LACP).....	66
3.7.2 Trunk Group Configuration (Конфигурация Trunk группы).....	68
3.7.3 TRUNK Member port configuration (Конфигурация портов)	68
3.7.4 TRUNK Load balancing configuration (Распределение нагрузки) ..	69
3.7.5 Trunk Display (Просмотр конфигурации)	70
3.8 Extra large frames (Jumbo frames)	70
3.8.1 Configure Jumbo frames (Настройка Jumbo frames)	70
3.8.2 Display Jumbo frames (Просмотр конфигурации)	70
3.9 Configure redundant ports (Резервирование портов)	71
3.9.1 Configuration of redundant ports (Настройка резервных портов).71	
3.9.2 Display redundant ports (Просмотр конфигурации)	72
3.10 UDLD configuration (Конфигурация UDLD)	72
3.11 LLDP configuration (Конфигурация LLDP)	73

3.11.1 LLDP configuration (Настройка LLDP)	74
3.11.2 Display LLDP (Просмотр конфигурации)	75
4. Configure port based security (Конфигурация безопасности портов)	75
4.1 Introduction (Введение).....	76
4.2 MACBinding configuration (Настройка MACBinding)	76
4.3 MACFilter configuration (Настройка MACFilter).....	77
4.4 Port learning restriction configuration (Настройка Port learning restriction)	78
4.5 Protection port configuration (Настройка Protection port)	79
4.5.1 Introduction to protection port (Защита порта)	79
4.5.2 Protection port configuration (Настройка Protection port).....	79
5. Configure port IP and MAC binding (Конфигурация привязки IP и MAC) ..	81
5.1 Introduction (Введение).....	81
5.2 IP and MAC binding configuration (Настройка IP и MAC binding)	82
6. Port loop detection (Определение петлевых соединений)	84
6.1 Protocol principle (Принцип работы протокола).....	84
6.1.1 Recovery mode (Режимы восстановления)	85
6.1.2 Protocol security (Безопасность протокола).....	85
6.2 Configuration introduction (Конфигурация протокола).....	86
6.2.1 Global configuration (Основные настройки).....	86
6.2.2 Port configuration (Настройки портов)	86
7. VLAN Configure (Конфигурация VLAN)	87
7.1 VLAN introduction (Введение)	87
7.1.1 VLAN ID (Идентификация VLAN)	88
7.1.2 VLAN Port member type (Типы портов-членов VLAN)	89
7.1.3 Default port VLAN (Первичный port VLAN)	89
7.1.4 VLAN mode of port (Режимы портов VLAN).....	89
7.1.5 VLAN RELAY (Порт VLAN RELAY)	90
7.1.6 Forwarding of data flow in VLAN (Обмен данными VLAN)	91
7.2 VLAN configuration (Настройка VLAN)	93

7.2.1 Creating and deleting VLAN (Создание и удаление VLAN)	93
7.2.2 Configure VLAN mode of port (Настройка портов VLAN)	94
7.2.3 VLAN configuration access mode (Настройка в режиме access) ...	94
7.2.4 VLAN configuration trunk mode (Настройка в режиме trunk)	95
7.2.5 VLAN configuration hybrid mode (Настройка в режиме trunk)	96
7.2.6 View VLAN information (Просмотр конфигурации).....	97
7.3 VLAN configuration example (Пример настройки VLAN)	97
7.3.1 Port based VLAN (VLAN основанный на портах)	97
7.3.2 802.1Q based VLAN (VLAN основанный на протоколе 802.1Q)	99
7.4 Mac, IP subnet, protocol VLAN(Протокол VLAN на Mac, IP).....	100
7.4.1 Introduction Mac, IP subnet, protocol VLAN (Введение)	100
7.4.2 Mac, IP subnet, protocol VLAN configuration (Конфигурация).....	101
7.5 Voice VLAN (VLAN для голосового трафика)	103
7.5.1 Voice VLAN Introduction (Введение)	103
7.5.2 Voice VLAN configuration (Конфигурация Voice VLAN)	104
7.6 VLAN mapping (VLAN отображение)	106
7.6.1 VLAN mapping Introduction (Введение).....	106
7.6.2 VLAN mapping configuration (Конфигурация VLAN mapping).....	106
7.7 QinQ (QinQ технология двойного тегирования)	107
7.7.1 QinQ (Введение)	107
7.7.2 QinQ configuration (Конфигурация QinQ)	108
8. QoS Configure (Конфигурация QoS).....	110
8.1 QoS introduction (Введение)	110
8.1.1 QoS based on cos (QoS основанный на cos)	112
8.1.2 QoS based on DSCP (QoS основанный на DSCP)	112
8.1.3 QoS based on Mac (QoS основанный на Mac).....	112
8.1.4 Policy based QoS (QoS основанный на политике)	113
8.2 QoS configuration (Конфигурация QoS)	113
8.2.1 QoS default configuration (QoS конфигурация по умолчанию) ...	113
8.2.2 Configure scheduling mode (Конфигурация режимов)	114

8.2.3 Configure queue weights (Установка весовых коэффициентов) ..	114
8.2.4 Configure mapping relationship between DSCP and qosprofile (Конфигурация взаимосвязи DSCP - qosprofile)	115
8.2.5 Configure port QoS based on DSCP (Конфигурация порта DSCP)	115
8.2.6 Configure port user priority (Конфигурация пользователя COS)	116
8.2.7 QoS configuration example (Пример конфигурации QoS)	116
8.2.8 Policy QoS configuration example (Пример конфигурации политики QoS).....	116
9. MSTP Configure (Конфигурация QoS)	117
9.1 MSTP introduction (Введение).....	117
9.1.1 Multi spanning tree domain (MST домен)	117
9.1.2 IST, CIST, CST	118
9.1.3 Intra domain operation (Операции домена).....	119
9.1.4 Inter domain operation (Операции между доменами).....	119
9.1.5 Skip count (Пропуск счета).....	120
9.1.6 Boundary port (Граничный порт).....	121
9.1.7 Interoperability MSTP and 802.1d STP (Совместимость)	121
9.1.8 Port role (Роль порта)	122
9.1.9 802.1D spanning tree Introduction (Введение)	125
9.2 MSTP configuration (Конфигурация MSTP)	127
9.2.1 Default configuration (Конфигурация по умолчанию).....	127
9.2.2 General configuration (Команды основной конфигурации)	127
9.2.3 Domain configuration (Конфигурация домена)	130
9.2.4 Instance configuration (Конфигурация Instance).....	130
9.2.5 Port configuration (Конфигурация порта)	131
9.2.6 PORTFAST Related configuration (Конфигурация PORTFAST)	134
9.2.7 Root Guard Related configuration (Конфигурация Root Guard)....	135
9.2.8 MSTP Configuration example (Пример конфигурации MSTP).....	136
10. EAPS Configure (Конфигурация EAPS)	138
10.1 EAPS introduction (Введение).....	138

10.2 EAPS Basic concepts (Основная концепция EAPS).....	138
10.3 EAPS Protocol introduction (EAPS протокол введение)	139
10.3.1 Link-Down Give an alarm (Link-Down оповещение)	139
10.3.2 Loop check (Проверка петли)	140
10.3.3 Ring recovery (Восстановление кольца)	140
10.3.4 EAPS compatible with extreme (EAPS совместимость).....	141
10.4 EAPS configuration (EAPS конфигурация)	141
10.5 EAPS Brief introduction of command (Команды EAPS).....	142
10.5.1 EAPS Configuration command (Команды конфигурации)	143
10.5.2 Single ring configuration (Пример организации кольца)	144
10.5.3 Cross ring data forwarding configuration (Пример организации пересекающихся колец).....	149
11. ERPS Configure (Конфигурация ERPS)	153
11.1 ERPS introduction (Введение)	153
11.1.1 ERPS Loop (ERPS петля)	153
11.1.2 ERPS node (ERPS нод)	154
11.1.3 Links and channels (Соединения и каналы)	154
11.1.4 ERPS VLAN (ERPS VLAN).....	155
11.2 ERPS working principle (Принципы работы ERPS).....	155
11.2.1 Normal state (Нормальное состояние)	155
11.2.2 Link failure (Обрыв соединения).....	156
11.2.3 Link recovery (Восстановление соединения)	157
11.3 ERPS Technical features (Возможности ERPS)	158
11.3.1 ERPS load balancing (Распределение нагрузки).....	158
11.3.2 Good safety (Безопасность)	159
11.3.3 Multi ring intersection and tangent (Пересечение колец)	160
11.4 ERPS Protocol command (Команды ERPS).....	160
11.5 ERPS Typical application (Применение ERPS)	163
11.5.1 Single ring example (Пример одиночного кольца)	163
11.5.2 Multi ring example (Пример нескольких колец)	166

11.5.3 Multi instance load balancing example (Пример распределения нагрузки)	172
12. AAA Configure (Конфигурация AAA)	181
12.1 802.1x introduction (Введение)	182
12.1.1 802.1x Equipment composition (Подключение оборудования) ..	182
12.1.2 Protocol package (Пакеты протокола)	184
12.1.3 Protocol flow interaction (Взаимодействие протоколов)	185
12.1.4 802.1x port status (Статусы 802.1x порта)	187
12.2 RADIUS introduction (Введение)	188
12.2.1 Protocol package (Пакеты протокола)	189
12.2.2 Protocol flow interaction (Взаимодействие протокола)	190
12.2.3 User authentication method (Методы аутентификации)	192
12.3 Configure 802.1x (Конфигурация 802.1x)	193
12.3.1 802.1xDefault configuration (Конфигурация по умолчанию)	193
12.3.2 Turn 802.1x on and off (Включение и отключение 802.1x)	194
12.3.3 Configure 802.1x port status (Выбор статуса портов)	194
12.3.4 Configure re authentication (Конфигурация реаутентификации)	195
12.3.5 Configure number of port access hosts (Управление хостами) ..	196
12.3.6 Configure interval and number of retransmissions (Интервалы передачи)	196
12.3.7 Configure port as transport port (Транспортный порт)	197
12.3.8 Configure 802.1x client version number (Номер верси клиента) ..	197
12.3.9 Configure check the client version number (Проверка номера версии клиента)	198
12.3.10 Configure authentication method (Выбор метода аутентификации)	198
12.3.11 Configure to check the timing package of the client (Проверка тайминга пакета клиента)	198
12.3.12 Display 802.1x information (Информация 802.1x)	199
12.4 Configure RADIUS (Конфигурация RADIUS)	199
12.4.1 RADIUS Default configuration (Конфигурация по умолчанию) ..	199

12.4.2 IP address of the authentication server (IP адрес сервера аутентификации)	200
12.4.3 Configure shared key (Конфигурация общего ключа)	200
12.4.4 Start and close billing (включение/отключение биллинга)	200
12.4.5 Configure radius port and attribute information (порт radius и информация атрибута)	201
12.4.6 Configure radius roaming function (Функция роуминга)	201
12.4.7 Display Radius information (Информация Radius)	202
12.4.8 Configuration example (Пример конфигурации)	202
12.4 TACACS+ Introduction (Протокол TACACS+)	202
13. GMRP Configure (Конфигурация GMRP)	204
13.1 GMRP Introduction (Протокол GMRP)	204
13.2 GMRP configurtion (Протокол GMRP)	205
13.2.1 GMRP settings (Установки GMRP)	205
13.2.2 View GMRP information (Просмотр GMRP информации)	206
13.2.3 GMRP Configuration example (Пример конфигурации)	206
14. IGMP SNOOPING Configure (IGMP SNOOPING)	207
14.1 IGMP SNOOPING Introduction (Протокол GMRP)	207
14.1.1 IGMP SNOOPING Processing (IGMP SNOOPING процессы)	208
14.1.2 Dynamic Multicast Layer 2 (Динамический Multicast)	209
14.1.3 Join a group (Включение в группу)	209
14.1.4 Left a group (Выход из группы)	211
14.1.5 IGMP Interrogator (Следование IGMP)	212
14.1.6 IGMP snooping Multicast filtering (Фильтрация IGMP snooping)	213
14.2 IGMP SNOOPING configuration (GMRP конфигурация)	213
14.2.1 IGMP snooping default configuration (IGMP начальная конфигурация)	213
14.2.2 Open and close IGMP snooping (Включение и отключение)	213
14.2.3 Configure lifetime (Установка lifetime)	214
14.2.4 Configure fast-leave (Настройка fast-leave)	214
14.2.5 Configure MROUTER (Настройка MROUTER)	215

14.2.6 Configure IGMP snooping query port (Настройка порта query) ..	215
14.2.7 Configure IGMP snooping query function (Настройка query)	215
14.2.8 Configure IGMP snooping multicast filtering (Настройка фильтрации)	215
14.2.9 Display information (Просмотр информации)	216
14.2.10 IGMP SNOOPING Configuration example (Пример конфигурации)	216
15. MVR Configuration (Конфигурация MVR)	217
15.1 MVR introduction (Введение)	217
15.2 MVR Configure (Конфигурация MVR)	218
15.3 MVR Configuration example (Пример конфигурации)	218
16. DHCP SNOOPING Configure (DHCP SNOOPING)	220
16.1 DHCP SNOOPING introduction (Введение)	220
16.1.1 DHCP SNOOPING Processing (Процесс DHCP SNOOPING)	221
16.1.2 DHCP SNOOPING Binding table (Таблица привязки)	221
16.1.3 Physical port of the linked server (Физический порт сервера) ...	222
16.2 DHCP SNOOPING Configuration (Конфигурация)	223
16.2.1 DHCP SNOOPING Default configuration (Настройки по умолчанию)	223
16.2.2 Global on and off (Глобальное включение и отключение)	223
16.2.3 Interface on and off (Включение и отключение для интерфейса)	223
16.2.4 Interface on and off OPTION82 (Включение и отключение OPTION82)	224
16.2.5 Display information (Просмотр информации)	224
16.3 DHCP SNOOPING Configuration example (Пример конфигурации) ..	225
17. DHCP CLIENT Configuration (DHCP Клиент)	226
17.1 DHCP CLIENT Configuration (Конфигурация клиента)	227
18. DHCP RELAY Configuration (DHCP RELAY)	227
18.1 DHCP RELAY introduction (Введение)	227
18.2 DHCP RELAY Configuration (Конфигурация DHCP RELAY)	229

18.2.1 DHCP relay of interface (Включение DHCP relay)	229
18.2.2 Display information (Просмотр информации)	229
18.2.3 DHCP RELAY configuration example (Пример конфигурации)...	229
19. DHCP SERVER Configuration (DHCP Сервер).....	231
19.1 DHCP SERVER introduction (Введение).....	231
19.2 DHCP SERVER configuration (Конфигурация DHCP SERVER)	231
19.2.1 Start global DHCP server (Включение global DHCP server).....	232
19.2.2 Start interface to receive DHCP server message (Включение приема сообщений DHCP server интерфейсом)	232
19.2.3 Configure address pool (Настройка пула адресов).....	232
19.2.4 Configure address pool range (Настройка диапазона пула адресов).....	233
19.2.5 Configure address pool subnet mask (Настройка пула адресов маски подсети)	233
19.2.6 Configure address pool lease (Настройка времени использования пула адресов)	233
19.2.7 Configure address pool default gateway (Настройка шлюза пула адресов).....	234
19.2.8 Configure address pool DNS server (Настройка DNS сервера пула адресов).....	234
19.2.9 Manual address exclude from address pool (Исключение адресов из пула)	235
19.2.10 Configure option82 (Подключение функции option82).....	235
19.2.11 Clear assigned address table entries (Очистка таблицы адресов)	236
19.2.12 Clear conflicting address table entries detected (Очистка таблицы конфликтующих адресов)	236
19.2.13 DHCP SERVER configuration example (Пример конфигурации)	236
20. Configure ACL (Конфигурация ACL).....	238
20.1 ACL resource library Introduction (Введение библиотека).....	238
20.2 ACL Filter Introduction (Введение ACL фильтр).....	241
20.3 ACL Repository configuration (ACL хранилище).....	243

20.4 Time period based ACL (Период времени ACL)	245
20.5 ACL Filter configuration (Конфигурация ACL фильтра)	248
20.6 ACL Filter configuration example (Пример конфигурации)	248
21. TCP/IP Basic configuration (Конфигурация TCP/IP)	249
21.1 Configure VLAN interface (VLAN интерфейс настройка)	250
21.2 Configure ARP (ARP конфигурация)	252
21.2.1 Configure static ARP (Конфигурация статического ARP)	252
21.2.2 View ARP information (Просмотр информации ARP)	253
21.3 Configure static routing (Конфигурация статической маршрутизации)	253
21.4 TCP/IP Basic configuration example (Пример конфигурации)	256
21.4.1 L3 layer interface (Интерфейс уровня L3)	257
21.4.2 Static routing (Статическая маршрутизация)	257
21.4.3 ARP (Протокол ARP)	257
22. Configure SNMP (Конфигурация SNMP)	258
22.1 SNMP introduction (Введение)	258
22.2 SNMP configuration (Конфигурация SNMP)	259
22.3 SNMP configuration example (Пример конфигурации)	261
23. Configure RMON (Конфигурация RMON)	261
23.1 RMON introduction (Введение)	261
23.2 RMON Configuration (Конфигурация RMON)	262
23.3 RMON Cconfiguration example (Пример конфигурации)	264
24. Cluster configuration (Конфигурация кластера)	265
24.1 Cluster management introduction (Введение)	265
24.1.1 Cluster definition (Определение кластера)	265
24.1.2 Cluster role (Роль кластера)	267
24.1.3 NDP introduction (Протокол NDP)	268
24.1.4 NTPD introduction (Протокол NTPD)	268
24.1.5 Cluster management (Управление кластером)	270
24.1.6 Manage VLAN (Управление VLAN)	273

24.2 Cluster configuration introduction (Введение)	274
24.3 Configuration management device (Конфигурация управляющего устройства).....	275
24.3.1 Enable NDP function (Включение функции NDP)	275
24.3.2 Configure NDP parameters (Конфигурация параметров NDP) ...	275
24.3.3 Enable NTDP function (Включение функции NTDP).....	276
24.3.4 Configure NTDP parameters (Конфигурация параметров NTDP)	276
24.3.5 Configure NTDP information (Конфигурация информации NTDP)	277
24.3.6 Enable cluster function (Функция создание кластера)	278
24.3.7 Establish cluster (Создание кластера)	278
24.3.8 Configure member interaction (Взаимодействие членов кластера)	281
24.3.9 Configure cluster member management (Настройка управления членами кластера)	282
24.4 Configure member devices (Конфигурация устройств-членов).....	282
24.4.1 Enable NDP function (Включение функции NDP)	282
24.4.2 Enable NTDP function (Включение функции NTDP).....	282
24.4.3 Configure manual collection NTDP information (Включение функции NDP)	282
24.5 Configure access cluster members (Конфигурация доступа устройств- членов)	283
24.6 Display cluster configuration (Просмотр конфигурации)	283
24.7 Cluster management example (Пример управления кластером)	284
25. SNTP Configuration (Конфигурация SNTP).....	287
25.1 SNTP introduction (Введение)	287
25.2 Configure SNTP (Конфигурация SNTP)	288
25.2.1 Default SNTP settings (Начальные установки)	288
25.2.2 Configure SNTP server (Конфигурация SNTP сервера)	288
25.2.3 Configure the interval of SNTP synchronization (Конфигурация времени синхронизации).....	289

25.2.4 Configure local time zone (Конфигурация часового пояса)	290
25.2.5 SNTP information display (Просмотр конфигурации).....	290
26. IGMP Configuration (Конфигурация IGMP).....	290
26.1 IGMP introduction (Введение)	290
26.2 IGMP configuration (Конфигурация IGMP).....	292
26.2.1 Start IGMP function of interface (Включение функции IGMP)	293
26.2.2 Configure the group filter access control list for the interface (Конфигурация фильтра группового списка доступа)	293
26.2.3 Configure the access control list filtered by the interface leaving the group (Выход из группы фильтра группового списка доступа)	294
26.2.4 Number of specific group queries for the configuration interface (Конфигурация количества запросов группы интерфейса)	295
26.2.5 Configure the specific group query interval of the interface (Конфигурация интервалов запросов группы интерфейса)	295
26.2.6 Configure the non querier timer time of the interface (Конфигурация таймера отсутствия запросов)	296
26.2.7 Configure the query timer interval of the interface (Конфигурация интервалов отправки запросов)	296
26.2.8 Configure the maximum response time of the interface (Конфигурация максимального времени ответа интерфейса)	297
26.2.9 Configure interface parameters (Конфигурация параметров интерфейса)	298
26.2.10 Configure protocol version of the interface (Конфигурация версии протокола интерфейса).....	298
26.3 IGMP configuration example (Пример конфигурации)	299
27. PIM-SM Configuration (Конфигурация PIM-SM)	300
27.1 PIM-SM introduction (Введение).....	301
27.2 PIM-SM configuration (Конфигурация PIM-SM)	304
27.2.1 Start multicast routing (Запуск функции multicast routing).....	305
27.2.2 Configure multicast routing table capacity (Конфигурация размера таблицы маршрутизации).....	305
27.2.3 Configure TTL value of multicast interface (Конфигурация значения TTL multicast интерфейса)	305

27.2.4 Start interface PIM SM function (Запуск функции PIM SM на интерфейсе)	306
27.2.5 Configure passive mode of interface (Конфигурация пассивного режима интерфейса).....	306
27.2.6 Configure interface priority (Конфигурация приоритета интерфейса)	307
27.2.7 Configure interface Hello message (Конфигурация сообщения Hello message).....	308
27.2.8 Configure interface Hello timer interval (Конфигурация интервала Hello таймера)	308
27.2.9 Configure the holding time of neighbors on the interface (Конфигурация holding time интерфейса).....	309
27.2.10 Configure neighbor list filtering of the interface (Конфигурация списка фильтрации интерфейса).....	309
27.2.11 Configure the source address of unicast registration message (Конфигурация адреса источника unicast registration message)	310
27.2.12 Configure the number limit of registered message (Конфигурация предельного количества зарегистрированных сообщений)	311
27.2.13 Check RP when configuring registration (Проверка RP при конфигурации регистрации сообщений)	311
27.2.14 Configure registration inhibit timer (Конфигурация таймера подавления регистрации)	312
27.2.15 Configure and register Kat timer (Конфигурация Kat таймера регистрации)	312
27.2.16 Registration address configuration (Конфигурация регистрации адреса)	312
27.2.17 Configure the checksum of the registered message in Cisco mode (Конфигурация проверки контрольной суммы в режиме Cisco).....	313
27.2.18 Configure static RP address (Конфигурация статического RP адреса)	314
27.2.19 Configure candidate RP (Конфигурация кандидата RP).....	314
27.2.20 Configure ignore RP set priority (Конфигурация приоритета отклонения RP).....	315
27.2.21 Configure c-rp-adv message in Cisco mode (Конфигурация c-rp-adv сообщения в режиме Cisco).....	315

27.2.22 Configure candidate BSR (Конфигурация кандидата BSR).....	316
27.2.23 Configure JP timer interval (Конфигурация JP таймера).....	316
27.2.24 Configure SPT switching (Конфигурация переключения SPT)	317
27.2.25 Configure SSM (Конфигурация SSM)	318
27.2.26 Configure multicast security (Конфигурация multicast безопасности)	318
27.3 PIM-SM configuration example (Пример конфигурации)	319
28. RIP Configuration (Конфигурация RIP).....	322
28.1 RIP introduction (Введение).....	322
28.2 RIP configuration (Конфигурация RIP)	324
28.2.1 Start RIP and enter RIP configuration mode (Запуск и режим конфигурации RIP)	324
28.2.2 Enable RIP interface (Запуск RIP интерфейса)	325
28.2.3 Configure unicast message transmission (Конфигурация передачи unicast сообщений)	326
28.2.4 Configure the working state of the interface (Конфигурация рабочего статуса интерфейса)	326
28.2.5 Configure default routing weights (Конфигурация веса маршрута)	327
28.2.6 Configure management distance (Конфигурация дистанции управления)	327
28.2.7 Configure timer (Конфигурация таймера)	328
28.2.8 Configuration version (Версия конфигурации).....	328
28.2.9 Introducing external routing (Введение внешней маршрутизации)	329
28.2.10 Configure routing filtering (Конфигурация фильтра маршрутизации)	329
28.2.11 Configure additional routing weight (Конфигурация дополнительного веса маршрута)	330
28.2.12 Configure RIP version of interface (Конфигурация версии RIP интерфейса)	331
28.2.13 Configure the transceiver status of the interface (Конфигурация статуса интерфейса).....	332

28.2.14 Configure horizontal segmentation (Конфигурация горизонтальной сегментации)	333
28.2.15 Message authentication (Аутентификация сообщений)	334
28.2.16 Configure interface weight (Конфигурация веса интерфейса)	335
28.2.17 Display information (Просмотр конфигурации)	335
28.3 RIP configuration example (Пример конфигурации)	336
29. RIPng Configuration (Конфигурация RIPng)	337
29.1 RIPng introduction (Введение)	338
29.2 RIPng configuration (Конфигурация RIPng)	338
29.2.1 Start RIPng and enter RIPng configuration mode (Запуск и режим конфигурации RIP)	339
29.2.2 Enable RIPng interface (Запуск RIPng интерфейса)	339
29.2.3 Configure specific neighbors to send updates (Конфигурация соседних устройств для отправки обновлений)	340
29.2.4 Configure the working state of the interface (Конфигурация рабочего статуса интерфейса)	340
29.2.5 Configure default routing weights (Конфигурация веса маршрута)	341
29.2.6 Configure timer (Конфигурация таймера)	341
29.2.7 Introducing external routing (Введение внешней маршрутизации)	342
29.2.8 Configure routing filtering (Конфигурация фильтра маршрутизации)	342
29.2.9 Configure additional routing weight (Конфигурация дополнительного веса маршрута)	343
29.2.10 Configure horizontal segmentation (Конфигурация горизонтальной сегментации)	344
29.2.11 Configure interface weight (Конфигурация веса интерфейса)	345
29.2.12 Display information (Просмотр конфигурации)	345
29.3 RIPng configuration example (Пример конфигурации)	346
30. OSPF Configuration (Конфигурация OSPF)	347
30.1 OSPF introduction (Введение)	347
30.2 OSPF configuration (Конфигурация OSPF)	349

30.2.1 Start OSPF and enter OSPF mode (Запуск протокола OSPF и режима конфигурации)	350
30.2.2 Enable interface (Включение интерфейса)	351
30.2.3 Specify host (Определение хоста)	351
30.2.4 Configure router ID (Конфигурация ID роутера)	352
30.2.5 Configure adjacency points (Конфигурация смежных точек)	353
30.2.6 Prohibit the interface from sending messages (Запрет на отправку сообщений)	354
30.2.7 Configure SPF calculation time (Вычисление времени SPF)	354
30.2.8 Configure management distance (Конфигурация дистанции управления)	355
30.2.9 Introducing external routing (Введение внешней маршрутизации)	356
30.2.10 Configure the network type of the interface (Конфигурация типа сетевого интерфейса)	357
30.2.11 Configure Hello message sending interval (Конфигурация интервалов отправки Hello сообщений)	359
30.2.12 Configure neighbor router expiration time (Конфигурация времени существования соседнего роутера)	359
30.2.13 Configure retransmission time (Конфигурация интервала ретрансляции)	360
30.2.14 Configure interface delay (Конфигурация интервала задержки интерфейса)	360
30.2.15 Configure the priority of interface in Dr election (Конфигурация приоритета интерфейса при выборе Dr)	361
30.2.16 Configure the cost of sending messages on the interface (Конфигурация важности сообщений при отправке интерфейсом) ..	362
30.2.17 MTU value for DD message sent by interface (Значение MTU value для сообщений DD отправленных интерфейсом)	363
30.2.18 Configure interface message authentication (Аутентификация сообщений интерфейса)	363
30.2.19 Configure regional virtual link (Конфигурация регионального виртуального соединения)	364

30.2.20 Configure regional routing aggregation (Конфигурация агрегации региональной маршрутизации)	366
30.2.21 Configure regional message authentication (Конфигурация аутентификации региональных сообщений)	366
30.2.22 Stub area configuration (Конфигурация Stub области)	367
30.2.23 Configure NSSA area (Конфигурация NSSA области)	367
30.2.24 Configure external route aggregation (Конфигурация агрегации внешних маршрутов)	368
30.2.25 Configure default weights for external routes (Конфигурация весов внешних маршрутов)	368
30.2.26 Display information (Просмотр конфигурации)	369
30.3 OSPF configuration example (Пример конфигурации)	370
31. OSPFv3 Configuration (Конфигурация OSPFv3)	371
31.1 OSPFv3 introduction (Введение)	372
31.2 OSPFv3 configuration (Конфигурация OSPFv3)	374
31.2.1 Start OSPFv3 (Запуск протокола OSPFv3)	374
31.2.2 Configure router ID (Конфигурация ID роутера)	375
31.2.3 Prohibit the interface from sending messages (Запрет на отправку сообщений)	376
31.2.4 Configure SPF calculation time (Вычисление времени SPF)	376
31.2.5 Configure management distance (Конфигурация дистанции управления)	377
31.2.6 Introducing external routing (Введение внешней маршрутизации)	378
31.2.7 Configure the network type of the interface (Конфигурация типа сетевого интерфейса)	379
31.2.8 Configure Hello message sending interval (Конфигурация интервалов отправки Hello сообщений)	381
31.2.9 Configure neighbor router expiration time (Конфигурация времени существования соседнего роутера)	381
31.2.10 Configure retransmission time (Конфигурация интервала ретрансляции)	382
31.2.11 Configure interface delay (Конфигурация интервала задержки интерфейса)	382

31.2.12 Configure the priority of interface in Dr election (Конфигурация приоритета интерфейса при выборе Dr)	383
31.2.13 Configure the cost of sending messages on the interface (Конфигурация важности сообщений при отправке интерфейсом) ..	384
31.2.14 MTU value for DD message sent by interface (Значение MTU value для сообщений DD отправленных интерфейсом)	385
31.2.15 Configure regional virtual link (Конфигурация регионального виртуального соединения).....	385
31.2.16 Configure regional routing aggregation (Конфигурация агрегации региональной маршрутизации)	386
31.2.17 Stub area configuration (Конфигурация Stub области)	387
31.2.18 Configure default weights for external routes (Конфигурация весов внешних маршрутов)	388
31.2.19 Display information (Просмотр конфигурации).....	388
31.3 OSPFv3 configuration example (Пример конфигурации).....	389
32. BGP Configuration (Конфигурация BGP)	390
32.1 BGP introduction (Введение)	390
32.2 BGP configuration (Конфигурация BGP).....	394
32.2.1 Start BGP and enter BGP configuration mode (Запуск протокола BGP)	394
32.2.2 Inject routing information into BGP protocol (Ввод маршрутную информацию в протокол BGP)	394
32.2.3 Configure BGP timer (Конфигурация таймера BGP).....	395
32.2.4 Configure BGP and IGP synchronization (Синхронизация протоколов BGP и IGP)	395
32.2.5 Configure interaction between BGP and IGP (Конфигурация взаимодействия между протоколами BGP и IGP)	396
32.2.6 Selection of BGP optimal path (Выбор оптимального пути BGP)	396
32.2.7 Configure BGP routing aggregation (Агрегация маршрутизации BGP)	397
32.2.8 Configure BGP routing reflector (Конфигурация BGP routing reflector)	398
32.2.9 Configure BGP as Federation (Конфигурация BGP Federation) ..	399

32.2.10 Configure BGP management distance (Конфигурация дистанции управления BGP).....	400
32.2.11 Configure BGP routing update mechanism (Конфигурация механизма обновления BGP маршрутизации).....	400
32.2.12 Configure BGP local autonomous system (Конфигурация BGP локальной автономной системы)	401
32.3 BGP configuration example (Пример конфигурации).....	402
33. VRRP Configuration (Конфигурация VRRP)	404
33.1 VRRP introduction (Введение).....	404
33.1.1 VRRP summary (Принцип работы протокола VRRP)	404
33.1.2 VRRP term (VRRP пояснения и терминология)	407
33.1.3 VRRP Protocol interaction (Взаимодействие протокола VRRP)	408
33.1.4 Election of virtual master router (Выбор виртуального мастера-роутера)	411
33.1.5 Status of virtual router (Статус виртуального роутера)	412
33.1.6 VRRP track (VRRP трекинг).....	414
33.2 VRRP configuration (Конфигурация VRRP)	416
33.2.1 Create and delete virtual routers (Создание и удаление виртуальных роутеров)	416
33.2.2 Configure the virtual IP address of virtual router (Конфигурация виртуального IP адреса виртуального роутера)	417
33.2.3 Configure parameters of virtual router (Конфигурация параметров виртуального роутера).....	417
33.2.4 Configure VRRP tracking (Конфигурация VRRP трекинга)	419
33.2.5 Start and shut down the virtual router (Включение и отключение виртуального роутера).....	420
33.2.6 View VRRP information (Просмотр конфигурации)	421
33.3 VRRP Configuration example (Пример конфигурации).....	422
34. VLLP Configuration (Конфигурация VLLP).....	424
34.1 VLLP introduction (Введение)	424
34.2 VLLP configuration (Конфигурация VLLP).....	428
34.2.1 Create VLLP device on layer 3 interface (создание VLLP оборудования на layer 3 интерфейсе).....	429

34.2.2 Enable VLLP device (Включение VLLP оборудования).....	429
34.2.3 Create VLLP port on layer 2 interface (Создание VLLP порта на layer 2 интерфейсе).....	429
34.2.4 Configure VLLP device priority (Конфигурация приоритета VLLP оборудования).....	430
34.2.5 Configure VLLP device query timer interval (Конфигурация интервала таймера запросов).....	430
34.2.6 Configure attached VLANs (Конфигурация дополнительных VLAN).....	431
34.2.7 Configure VLLP port priority (Конфигурация приоритета VLLP портов).....	431
34.2.8 View VLLP information (Просмотр конфигурации)	431
34.3 VLLP Configuration example (Пример конфигурации).....	432
35. Configure policy routing (Политика маршрутизации)	434
35.1 Policy routing introduction (Введение).....	434
35.2 Policy routing configuration (Конфигурация политики маршрутизации).....	434
35.2.1 Create a new policy route (Создание новой политики маршрутизации).....	435
35.2.2 Insert a policy route (Введение политики маршрутизации)	435
35.2.3 Delete a policy route (Удаление политики маршрутизации).....	436
35.2.4 Move a policy route (Перемещение политики маршрутизации)...	436
35.2.5 View policy routing information (Просмотр информации)	436
35.3 Policy routing configuration example (Пример конфигурации)	436
36. Configure system log (Логи системы).....	437
36.1 System log introduction (Введение).....	437
36.1.1 Format of log information (Формат логов).....	438
36.1.2 Storage of logs (Хранение логов).....	440
36.1.3 Display of logs (Просмотр информации)	441
36.1.4 Debugging tool (Инструментарий отладки)	441
36.2 System log configuration (Конфигурация логов системы)	442

36.2.1 Configure terminal real-time display switch (Конфигурация просмотра в режиме реального времени).....	442
36.2.2 Set log level (Установка уровня логов)	443
36.2.3 View log information (Просмотр логов)	443
36.2.4 Configure debugging switch (Конфигурация инструментов отладки)	443
36.2.5 View debugging information (Просмотр информации отладки)	445
36.3 SYSLOG configuration (Конфигурация функции SYSLOG).....	446
36.3.1 SYSLOG introduction (Введение)	446
36.3.2 SYSLOG configuration (Конфигурация SYSLOG)	447
36.3.3 SYSLOG configuration example (Пример конфигурации SYSLOG)	447
37. IPv6 Basic configuration (Базовая конфигурация IPv6).....	448
37.1 IPv6 introduction (Введение)	448
37.1.1 IPv6 Protocol features (Возможности протокола IPv6).....	448
37.1.2 IPv6 Address introduction (IPv6 адреса)	450
37.1.3 IPv6 neighbor discovery introduction (Введение IPv6 поиск соседних устройств).....	454
37.1.4 IPv6 PMTU discovery (Поиск IPv6 PMTU)	457
37.1.5 Protocol specification (Спецификации протокола)	458
37.2 IPv6 basic configuration tasks (Конфигурация IPv6)	459
37.2.1 Configure IPv6 unicast address (Конфигурация unicast адреса)	459
37.3 Configure IPv6 Neighbor Discovery Protocol (Конфигурация NDP) ...	460
37.3.1 Configure relevant parameters of RA message (Конфигурация параметров RA сообщения)	460
37.3.2 Configure the number of neighbor request messages (Конфигурация количества запросов)	464
37.3.3 IPv6 Static routing configuration (IPv6 статическая конфигурация)	464
37.4 IPv6 Display and maintenance (Просмотр конфигурации)	465
38. MLD SNOOPING configuration (Конфигурация MLD SNOOPING)	465
38.1 MLD SNOOPING introduction (Введение).....	466

38.1.1 MLD snooping process (Процесс MLD snooping)	466
38.1.2 Layer 2 Dynamic Multicast (Динамический Multicast)	467
38.1.3 Join a group (Включение в группу)	468
38.1.4 Leave a group (Выход из группы)	470
38.2 MLD SNOOPING Configuration (Конфигурация MLD SNOOPING)	471
38.2.1 MLD SNOOPING default configuration (Конфигурация MLD SNOOPING по умолчанию)	471
38.2.2 On and off MLD SNOOPING (Включение и отключение MLD SNOOPING)	471
38.2.3 Configure lifetime (Конфигурация lifetime).....	472
38.2.4 Configure fast-leave (Конфигурация fast-leave).....	472
38.2.5 Configure MROUTER (Конфигурация MROUTER).....	472
38.2.6 Configure VLAN querier (Конфигурация VLAN querier)	473
38.2.7 Display information (Просмотр конфигурации)	473
38.3 MLD SNOOPING Configuration example (Пример конфигурации).....	474
39. POE configuration (Конфигурация POE)	475
39.1 POE introduction (Введение)	475
39.2 POE configuration (Конфигурация PoE).....	476
39.2.1 Manual PoE configuration (Ручная конфигурация PoE)	476
39.2.2 PoE Policy configuration (Конфигурация политики PoE).....	476
39.2.3 PDQuery configuration (Конфигурация запросов PD).....	478

1. CLI Command line interface (Введение)

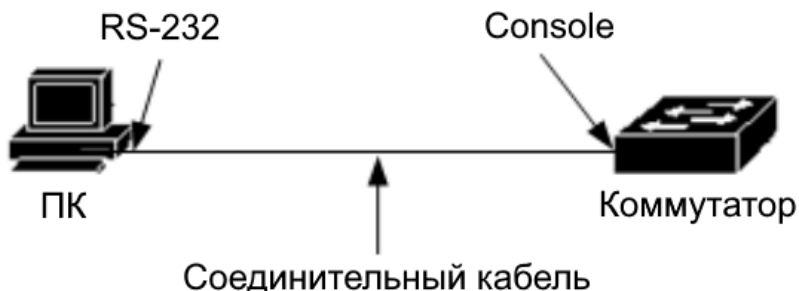
Управление коммутатором через COM-порт (RS-232) с помощью командной строки (CLI) предоставляет пользователю широкие возможности по настройке и управлению режимами работы коммутатора, или в случае, если по каким-либо причинам управление через WEB-интерфейс недоступно.

1.1 Accessing CLI (Доступ к CLI)

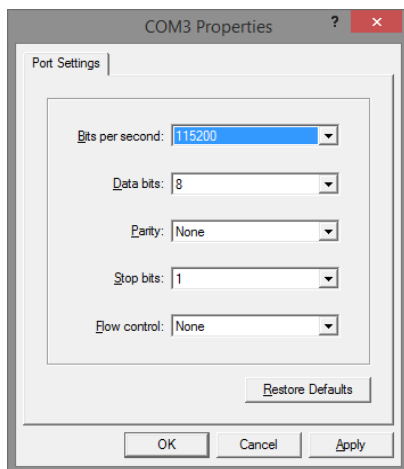
Процедура подготовки к управлению с помощью командной строки описана в руководстве по эксплуатации коммутатором (п.8).

1.1.1 Users access CLI via Console port (Подключение через порт Console)

- Соедините порт Console коммутатора с COM-портом компьютера с помощью кабеля (см. рис. ниже);



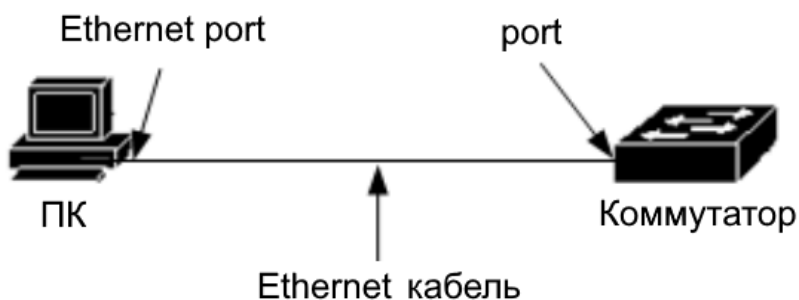
- Запустите HyperTerminal на ПК;
- Задайте имя для нового консольного подключения;
- Выберите COM-порт, к которому подключен коммутатор;
- Настройте COM-порт следующим образом:
 - ✓ Скорость передачи данных (Baud Rate) – 115200;
 - ✓ Биты данных (Data bits) – 8;
 - ✓ Четность (Parity) – нет (no);
 - ✓ Стоп биты (Stop bits) – 1;
 - ✓ Управление потоком (flow control) – нет (no).



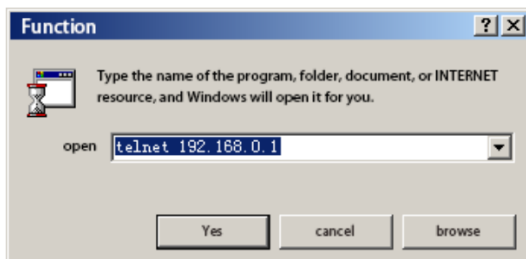
- Включите коммутатор, система предложит войти в интерфейс CLI (управление через командную строку).

1.1.2 Users access CLI via TELNET (Подключение через TELNET)

- Соедините порт коммутатора с Ethernet портом компьютера с помощью кабеля (см. рис. ниже);



- Чтобы получить доступ к CLI коммутатора через Telnet, ПК и коммутатор должны находиться в одной сети. Установите IP адрес сетевого адаптера ПК (например 192.168.0.3) IP адрес коммутатора **192.168.0.1** (по умолчанию);
- Если ПК соединен с коммутатором откроется Telnet интерфейс (встроен в командную строку CMD семейства операционных систем Microsoft Windows см. рис. ниже);



- Если пароль не установлен, то Telnet интерфейс предоставит доступ к CLI коммутатора. В противном случае потребуется аутентификация (по умолчанию: Login: **admin** Password: **admin**).

Внимание !

- IP адрес порта коммутатора зависит от VLAN layer 3 интерфейса. Перед тем как подключиться к коммутатору, должен быть установлен IP адрес VLAN интерфейса. По умолчанию IP адрес vlan1 192.168.0.1. Иной IP адрес VLAN интерфейса может быть установлен через порт console.

- Подключить ПК к коммутатору можно через один из портов непосредственно кабелем Ethernet, или через локальную сеть (если имеется возможность соединения ПК и VLAN коммутатора).

1.2 CLI Mode introduction (CLI режимы)

1.2.1 CLI Role of mode (CLI роли режимов)

Основные функции режимов CLI:

- Упрощение классификации пользователей и предотвращение неавторизованного использования командной строки.

Пользователи разделяются на два уровня: обычный и привилегированный. Обычные пользователи могут только просматривать статусы некоторых операций коммутатора с использованием команд просмотра.

Привилегированные пользователи (в дополнение к перечисленным выше) имеют возможность конфигурировать настройки коммутатора.

- Упрощение управлением коммутатора.

Т.к. коммутатор имеет множество конфигураций, то для пользователей удобно объединять их в различные режимы для использования. Сходные по функциональному назначению команды объединены в режимы. Например связанные с управлением VLAN команды и команды управления интерфейсом объединены в соответствующие режимы.

1.2.2 Identification of CLI mode (Идентификация CLI режимов)

При использовании командной строки пользователь может узнать текущий режим CLI, который указан в CLI prompt. CLI prompt состоит из двух частей, одна идентифицирует хост, а другая идентифицирует режим.

CLI prompt хост использует имя хоста системы. Имя хоста системы (по умолчанию *switch*) может быть изменено. Таким образом, CLI prompt начинает работу с именем хоста «switch».

Изменения части CLI prompt, касающейся режима CLI не предусмотрены. Каждый режим привязан к соответствующей данному режиму строке. Существуют фиксированные строки и изменяемые строки. Например строка режима VLAN конфигурации фиксированная, а строка режима интерфейса конфигурации изменяемая.

Пример:

- CLI prompt switch # определяет привилегированный режим (switch определяет хост, # определяет режим).

- CLI prompt switch (config-ge1 / 1) # определяет режим интерфейса конфигурации и конфигурирует GE1 / 1 порт (switch определяет хост, (config-ge1 / 1), # определяет режим).

- CLI prompt switch (config-vlan2) # определяет режим интерфейса конфигурации vlan2 (switch определяет хост, (config-vlan2), # определяет режим).

1.2.3 Classification of CLI modes (Классификация CLI режимов)

Режимы CLI разделены на четыре группы: обычный (normal mode), привилегированный (privileged mode), глобальной конфигурации

(global configuration mode) и конфигурационный (configuration sub mode, включает несколько cli подрежимов).

Простым пользователям доступен только обычный режим, привилегированным пользователям доступны все cli режимы.

При подключении через console и telnet сначала доступен только нормальный режим. После ввода соответствующей команды возможен переход в привилегированный режим при условии успешной аутентификации, (*простой пользователь не имеет возможности перехода в другие режимы*). После входа в привилегированный режим возможно перейти в режим глобальной конфигурации. Для перехода из режима глобальной конфигурации в другие подрежимы (sub mode) следует ввести соответствующие команды.

Таблица основных cli режимов коммутатора.

Режим	Описание	Prompt	Команда входа	Команда выхода
Normal mode	Просмотр текущей информации о коммутаторе.	Switch>	Режим доступен автоматически при первом входе.	exit или quit для выхода из telnet terminal.
Privileged mode	Просмотр текущей информации о коммутаторе. Команды debugging, version upgrade, configuration maintenance.	Switch#	enable в режиме normal mode.	disable возврат в normal mode. exit или quit для выхода из telnet terminal.
whole situation mode of configuration	Основные команды недоступные в режимах configuration sub mode, такие как configuration static routing.	Switch(config)#	Перейти в режим команд из privileged mode.	exit, quit, или end возврат в privileged mode.
Interface configuration mode	Команды для конфигурации интерфейса портов и VLAN.	Port: Switch(config-ge1/1)# VLAN port: Switch(config-vlan1)#	< if name > в режиме global configuration mode.	exit или quit для выхода из global configuration mode. end для выхода из privileged mode.

VLAN mode of configuration	Команды для конфигурации VLAN, такие как создать и удалить VLAN.	Switch(config-vlan)#	vlan database order в режиме global configuration mode.	exit или quit для выхода из global configuration mode. end для выхода из privileged mode.
MSTP configuration mode	Команды для конфигурации MSTP, такие как создать и удалить MSTP интерфейсы.	Switch(config-mst)#	Войти в spanning tree MST в режиме global configuration mode.	exit или quit для выхода из global configuration mode. end для выхода из privileged mode.
Terminal configuration mode	Команды для конфигурации console и telnet, такие как timeout интерфейса.	Switch(config-line)#	Войти в строку команд vty в режиме global configuration mode.	exit или quit для выхода из global configuration mode. end для выхода из privileged mode.

1.3 Introduction to command syntax (Синтаксис команд)

1.3.1 Command composition (Композиция команд)

CLI команды состоят из ключевых слов объединенных с параметрами. Сначала располагается ключевое слово, далее следующее ключевое слово или параметр. Ключевые слова и параметры могут чередоваться. Команда может содержать только ключевое слово (без параметра).

Пример:

- Команда < write > состоит только из ключевого слова;
- Команда < show version > состоит из двух ключевых слов;
- Команда VLAN < VLAN ID > состоит из ключевого слова и параметра;
- Команда instance < instance ID > VLAN < VLAN ID > состоит из двух ключевых слов и двух параметров, которые чередуются.

1.3.2 Parameter type (Типы параметров)

Параметры CLI команд разделяются на два типа: обязательные и опциональные. При вводе команды обязательные параметры должны быть тоже введены, ввод опциональных параметров не является необходимым.

Пример:

- Для команды `VLAN < VLAN ID >` параметр является обязательным;
- Для команды `< show interface [if name] >` параметр является опциональным.

1.3.3 Command syntax rules (Правила синтаксиса)

При описании команд в тексте данной инструкции используются следующие правила:

- 1) Ключевые слова описываются словами, совпадающими по смыслу.

- 2) Параметры заключаются в символы `< >`

Пример: `VLAN < VLAN ID >`

- 3) Опциональный параметр заключается в прямоугольные скобки `[ ]`

Пример: `VLAN [< VLAN ID >]` в этом же случае символы `< >` могут быть опущены: `VLAN [VLAN ID]` соответственно параметр `VLAN ID` может не вводиться. Если параметр обязательный, то в скобки он не заключается.

- 4) Если необходимо выбрать одно из нескольких ключевых слов или параметров, то для выделения используются фигурные скобки `{ }`.

Для разделения ключевых слов и параметров используется символ `|` с обязательными пробелами до и после `|`

Пример: `spanning-tree mst link-type {point-to-point | shared}`

Выбор одного из параметров `no arp {<ip-address> | <ip-prefix>}`

Ключевые слова и параметры `Show spanning-tree mst {none | instance <0-15>}ng`

- 5) Если необходимо выбрать одно ключевое слово или параметр, то для выделения используются прямоугольные скобки `[ ]`. Для разделения ключевых слов и параметров используется символ `|` с обязательными пробелами до и после `|`

Пример: debug ip tcp [recv | send]

Ключевые слова recv и send могут быть выбраны или не выбраны
show ip route [<ip-address> | <ip-prefix>] show interface [<if-name> | switchport]

- 6) Если необходимо выбрать группу из нескольких ключевых слов или параметров, следует добавить символ «*» после группы (Group) ключевых слов или параметров.

Пример: ping <ip-address> [-n <count> | -l <size> | -r <count> | -s <count> | -j <count> <ip-address>* | -k <count> <ip-address>* | -w <timeout>]*

-j <count> <ip-address>* --- Несколько IP адресов могут быть добавлены последовательно

-k <count> <ip-address>* --- Несколько IP адресов могут быть добавлены последовательно

- 7) Если параметр представлен одним и более словами, следует разделить слова символом «-» и использовать строчные буквы

Пример (правильно): <vlan-id>, <if-name>, <router-id>, <count>

Пример (не правильно): <1-255>, <A.B.C.D>, <WORD>, <IFNAME>

1.3.4 Command abbreviation (Сокращения)

При вводе команд в CLI интерфейс, допустимо использование сокращенной записи ключевых слов. CLI поддерживает использование префиксов, соответствующих ключевым словам. В зависимости от уникальности префикса, CLI определяет соответствующую ключевому слову команду. Эта функция создает дополнительное удобство для пользователей CLI, т.к. сокращает число вводимых символов необходимых для завершения команды.

Пример: команда show version ---- sh ver

1.3.5 Grammar help (Помощь)

Интерфейс CLI поддерживает функцию помощи (по синтаксису) для команд и параметров на каждом уровне.

- 1) При вводе в cli символа ? будет показано первое ключевое слово и описание всех связанных с ним команд.
Пример: switch (config) #?
- 2) Введите первую часть команды, далее пробел и ?. Будут выведены все ключевые слова или параметры с описаниями.
Пример: switch #show?
- 3) Введите часть ключевого слова и символ ? Будут выведены все ключевые слова с описаниями относящиеся к введенному префиксу.
Пример: switch #show ver?
- 4) Введите первую часть команды, далее пробел и tab. Будут выведены все ключевые слова для следующего уровня (кроме параметров, в этом случае информация не выводится).
- 5) Введите часть ключевого слова, далее tab. Если только одно ключевое слово соответствует префиксу, то он будет дополнен. Если префиксу соответствует несколько ключевых слов, то они все будут выведены.

1.3.6 Command line error message (Сообщения об ошибках)

Сообщение об ошибке появляется при вводе команд с неверным синтаксисом. Типовые сообщения об ошибках приведены в таблице ниже:

Таблица типовых сообщений об ошибках.

Сообщение	Причина ошибки
Invalid input or Unrecognized command	Не найдено ключевое слово. Введен некорректный параметр. Введено слишком много ключевых слов или параметров.
Incomplete command	Неполная команда, не введено ключевое слово или параметр.
Ambiguous command	Неполная команда, несколько ключевых слов соответствует префиксу.

1.4 Command line shortcuts (Короткие команды)

1.4.1 Row edit shortcut (Короткие команды ввода)

CLI интерфейс поддерживает ввод коротких команд, которые упрощают работу с командной строкой. Короткие команды представляют собой определенные сочетания клавиш, использование которых делает процесс ввода более быстрым. Сочетания клавиш для ввода коротких команд приведены в таблице ниже:

Таблица коротких команд.

Клавиши	Назначение
Ctrl+p или ↑	Предыдущая команда
Ctrl+n или ↓	Следующая
Ctrl+u	Удалить всю строку
Ctrl+a	Перевод курсора в начало строки
Ctrl+f или →	Перевод курсора на одну позицию вправо
Ctrl+b или ←	Перевод курсора на одну позицию влево
Ctrl+d	Удалить последнюю позицию
Ctrl+h	Удалить предпоследнюю позицию
Ctrl+k	Удалить все позиции после курсора
Ctrl+w	Удалить все позиции перед курсором
Ctrl+e	Перевод курсора на конец строки
Ctrl+c или Ctrl+z	Прервать выполнение команды. В режиме глобальной конфигурации или режиме sub mode, CLI возвращается в привилегированный режим, если CLI находится в обычном или привилегированном режиме, режим CLI остается неизменным.
Tab	Закончить неполное ключевое слово. Если префиксу соответствует несколько ключевых слов, то все они выводятся, если префиксу не соответствует ни одно ключевое слово, то выводится сообщение об ошибке.

В некоторых случаях клавиши ↑ ↓ → ← не имеют функционального назначения.

1.4.2 Show command shortcuts (Короткие команды show)

Команды, которые начинаются со слова show, предназначены для вывода (показа) какой-либо информации. В некоторых случаях из-за большого объема выводимая информация не помещается в рамки одного окна. Для решения этой проблемы предусмотрена функция которая разделяет выводимую информацию на несколько окон. После вывода на экран одного окна, система ожидает команды для вывода следующего окна. Ниже представлена таблица коротких команд для управления этими возможностями.

Таблица коротких команд show.

Клавиши	Назначение
Пробел	Показать следующее окно
Enter	Показать следующую строку
Ctrl+c	Прервать выполнение команды и выйти в режим cli mode.
Другие клавиши	То же что Ctrl+c

1.5 Historical command (История команд)

CLI интерфейс поддерживает функцию хранения истории введенных команд. В памяти сохраняется до 20и последних введенных команд. Для просмотра истории (show history) следует ввести Ctrl+P, Ctrl+n или кнопки ↑, ↓ для выбора команды. Эта функция упрощает пользователю повторно вводить необходимые команды.

2. System management configuration (Конфигурация управление системой)

Перед переходом к изучению дальнейших возможностей по конфигурированию коммутатора следует провести настройку некоторых основных параметров управления системой. В этом разделе

описываются основные настройки системы управления коммутатором:

- System security configuration (Конфигурация системы безопасности)
- System maintenance and commissioning (Обслуживание и эксплуатация)
- Monitoring system (Мониторинг системы)
- Profile management (Управление профилем)
- Software version upgrade (Обновление ПО)

2.1 System security configuration (Конфигурация системы безопасности)

Для предотвращения допуска неавторизованных лиц к управлению коммутатором предусмотрен ряд настроек которые включают следующие:

- Multi user management control (Управление многопользовательским режимом)
- Enable Password control (Защита паролем)
- TELNET service control (Управление TELNET)
- SNMP service control (Управление SNMP)
- HTTP service control (Управление HTTP)

2.1.1 Multi user management control (Многопользовательский режим)

Многопользовательский режим повышает безопасность системы и обеспечивает возможность управления и обслуживания коммутатора несколькими пользователями одновременно. Многопользовательский режим обеспечивает безопасность путем индивидуальной аутентификации пользователей (присвоение каждому пользователю своего имени и пароля для входа в систему). При настройке многопользовательского режима устанавливаются определенные права для каждого пользователя.

Многопользовательский режим предусматривает два уровня прав для пользователей: обычный и привилегированный. Обычные пользователи могут использовать только обычный режим командной строки CLI, команды display для просмотра информации о коммутаторе.

Привилегированные пользователи могут использовать все режимы командной строки CLI, все команды управления коммутатором.

Многопользовательский режим применяется только для управления через telnet и недоступен для управления через console. При управлении через console для входа в систему не требуется аутентификация, пользователь может непосредственно войти в командную строку CLI. При входе в систему через telnet требуется ввести имя пользователя и пароль. Доступ к командной строке CLI возможен только после успешной аутентификации.

По умолчанию *имя пользователя и пароль* **admin**. Пользователь admin является администратором (привилегированным пользователем), его настройки не могут быть изменены на обычного пользователя, и он не может быть удален.

Таблица команд многопользовательского режима.

Команда	Описание	CLI ржим
username <user-name> password <key> {normal privilege}	Добавление пользователя, изменение пароля, прав пользователя (если пользователь уже существует). 1й параметр - user name, 2й параметр - пароль, опционально - права пользователя (normal – обычный, privilege – привилегированный).	Global configuration mode
no username [user-name]	Удалить пользователя.	Global configuration mode
show running-config	Просмотр текущей конфигурации системы, конфигурации многопользовательского режима.	Privileged mode

2.1.2 Enable password control (Управление аутентификацией)

Управление аутентификацией применяется для перевода управления коммутатора из обычного режима в привилегированный режим. Перед включением аутентификации пользователю доступна для просмотра только информация о коммутаторе. После включения и прохождения аутентификации пользователь получает доступ к настройкам коммутатора.

Возможность включения аутентификации не имеет привязки к конкретному пользователю. Любой пользователь, который входит через console или telnet в привилегированный режим управления должен ввести корректный пароль. После успешной аутентификации пользователь также может оставаться в обычном режиме управления.

При вводе команды в обычном режиме пользователю необходимо ввести пароль. После успешной аутентификации система перейдет в привилегированный режим управления. В противном случае система останется в обычном режиме управления.

При включенной аутентификации система по умолчанию не будет требовать ввод пароля для входа в обычный режим.

Таблица команд управления режимом аутентификации.

Команда	Описание	CLI режим
enable password <key>	Установка пароля для входа в систему.	Global configuration mode
no enable password	Удаление пароля, «пустой» пароль.	Global configuration mode
show running-config	Просмотр текущей конфигурации, в т.ч. пароля.	Global configuration mode
enable	Включение режима аутентификации. Переход в привилегированный режим.	Global configuration mode

Для обеспечения безопасности системы администратору следует включить режим аутентификации и установить пароль.

2.1.3 TELNET service control (Управление TELNET сервисом)

В случаях когда отсутствует необходимость дистанционного управления коммутатором и управление осуществляется локально через console, для обеспечения безопасности системы и предотвращения неавторизованного дистанционного управления, администратор может отключить управление по telnet (включено по умолчанию). Команды приведены в таблице ниже:

Таблица команд управления telnet сервисом

Команда	Описание	CLI ржим
security-manage telnet enable	Включить telnet	Global configuration mode
security-manage telnet disable	Отключить telnet	Global configuration mode
security-manage telnet number <1-100>	Диапазон параметра от 1 до 100, по умолчанию 5.	Global configuration mode
security-manage telnet access-group <1-99> (Note: subject to the actual product)	Выбор группы ACL, вкл управление IP адресом. Если группа ACL нестандартная или отсутствует, управление IP адресом будет невозможно.	Global configuration mode
no security-manage telnet access-group	Отключить управление IP адресом.	Global configuration mode
show security-manage	Просмотр конфигурации управления сервисом.	Privileged mode

2.1.4 SNMP Service control (Управление SNMP сервисом)

Управление сервисом SNMP и IP адресом доступа к коммутатору через ACL может быть включено или отключено (on / off). Команды приведены в таблице ниже:

Таблица команд управления SNMP сервисом

Команда	Описание	CLI ржим
security-manage snmp enable	Включить SNMP	Global configuration mode
security-manage snmp disable	Отключить SNMP	Global configuration mode
security-manage snmp access-group <1-99> (Note: subject to the actual product)	Выбор группы ACL, вкл управление IP адресом. Если группа ACL нестандартная или отсутствует, управление IP адресом будет невозможно.	Global configuration mode
no security-manage snmp access-group	Отключить управление IP адресом.	Global configuration mode
show security-manage	Просмотр конфигурации управления сервисом.	Privileged mode

2.1.5 HTTP Service control (Управление HTTP сервисом)

Управление сервисом HTTP и IP адресом доступа к коммутатору через ACL может быть включено или отключено (on / off).

Команды приведены в таблице ниже:

Таблица команд управления HTTP сервисом

Команда	Описание	CLI ржим
security-manage http enable	Включить HTTP	Global configuration mode
security-manage http disable	Отключить HTTP	Global configuration mode

security-manage http access-group <1-99> (Note: subject to the actual product)	Выбор группы ACL, вкл управление IP адресом. Если группа ACL нестандартная или отсутствует, управление IP адресом будет невозможно.	Global configuration mode
no security-manage http access-group	Отключить управление IP адресом.	Global configuration mode
show security-manage	Просмотр конфигурации управления сервисом.	Privileged mode

2.1.6 HTTPS Service control (Управление HTTPS сервисом)

Управление сервисом HTTPS и IP адресом доступа к коммутатору через ACL может быть включено или отключено (on / off). Команды приведены в таблице ниже:

Таблица команд управления HTTPS сервисом

Команда	Описание	CLI ржим
security-manage https enable	Включить HTTPS	Global configuration mode
security-manage https disable	Отключить HTTP	Global configuration mode
security-manage https access-group <1-99> (Note: subject to the actual product)	Выбор группы ACL, вкл управление IP адресом. Если группа ACL нестандартная или отсутствует, управление IP адресом будет невозможно.	Global configuration mode
no security-manage https access-group	Отключить управление IP адресом.	Global configuration mode
show security-manage	Просмотр конфигурации управления сервисом.	Privileged mode

2.1.7 SSH Service control (Управление SSH сервисом)

Многие программы сетевого сервиса такие как FTP, pop и Telnet не являются безопасными т.к. передают данные и пароли открытым текстом, что делает возможным их перехват злоумышленниками. Методы обеспечения безопасности этих сервисов имеют ряд уязвимостей для атак типа «middleman» (ретрансляция данных через промежуточный сервер). При использовании сервиса SSH производится шифрование передаваемых данных, что делает атаки типа «middleman» не эффективными, кроме того предотвращается DNS spoofing и IP spoofing. Еще одним преимуществом сервиса SSH является сжатие передаваемых данных и повышение скорости передачи.

Сервис SSH не заменяет Telnet, но обеспечивает безопасный канал передачи данных для FTP, pop и PPP. Команды приведены в таблице ниже:

Таблица команд управления SSH сервисом

Команда	Описание	CLI режим
security-manage ssh enable	Включить ssh	Global configuration mode
security-manage ssh disable	Отключить ssh	Global configuration mode
security-manage ssh access-group <1-99>	Выбор группы ACL, вкл управление IP адресом. Если группа ACL нестандартная или отсутствует, управление IP адресом будет невозможно.	Global configuration mode
no security-manage ssh access-group	Отключить управление IP адресом.	Global configuration mode
ip sshd auth-retries <times>	Установить количество попыток аутентификации.	Global configuration mode
ip sshd silence-period <seconds>	Установить время ожидания после исчерпания всех попыток аутентификации.	Global configuration mode
show ip sshd	Просмотр конфигурации настроек sshd	Privileged mode
show security-manage	Просмотр конфигурации управления сервисом.	Privileged mode

2.2 System maintenance and commissioning (Обслуживание системы)

Подготовка к эксплуатации и дальнейшее обслуживание коммутатора включает в себя следующие пункты:

- Configure the system clock (Настройка системных часов)
- Configure terminal timeout attribute (Настройка времени бездействия)
- System reset (Сброс системы)
- View system information (Просмотр системной информации)
- Network connectivity debugging (Устранение сетевых неполадок)
- Traceroute debugging (Устранение неполадок маршрутизации)

2.2.1 Configure the system clock (Настройка системных часов)

В коммутаторе предусмотрены системные часы, которые синхронизированы с реальным временем. С помощью команд можно устанавливать и просматривать текущее время. Для удобства эксплуатации коммутатора системные часы обеспечены встроенным питанием, которое делает не нужной установку времени после длительного отключения коммутатора от энергоснабжения. После первого включения коммутатора следует проверить точность установки системных часов и при необходимости откорректировать.

Таблица команд управления системными часами

Команда	Описание	CLI ржим
set date-time <year> <month> <day> <hour> <minute> <second>	Установка текущего времени и даты в последовательности: год, месяц, день, час, минуты, секунды (вводятся значения параметров).	Privileged mode
show date-time	Просмотр текущей даты и времени системных часов.	Normal mode. privileged mode

2.2.2 Configure terminal timeout attribute (Настройка времени бездействия)

Для обеспечения безопасности в коммутаторе предусмотрена функция автоматического выхода из режима управления при бездействии, по истечении определенного интервала времени, который может быть изменен пользователем.

Процессы выхода из режима управления через console и telnet различаются. При управлении через console по истечении временного интервала, командная строка CLI переходит в обычный режим. При управлении через telnet по истечении временного интервала, соединение telnet прерывается, коммутатор выходит из режима управления. По умолчанию временной интервал бездействия составляет 10 мин, пользователь может изменить его или отключить эту функцию.

Таблица команд настройки времени бездействия

Команда	Описание	CLI ржим
exec-timeout <minutes> [seconds]	Установка времени бездействия. При выборе значения параметра 0, функция будет отключена.	Global configuration mode
no exec-timeout	Установка времени бездействия 10 мин (заводская установка, по умолчанию).	Global configuration mode
show running-config	Просмотр текущей конфигурации системы, в т.ч. установленного времени бездействия.	Privileged mode

2.2.3 System reset (Сброс системы)

Таблица команды сброса системы

Команда	Описание	CLI ржим
reset	Сброс системы	Privileged mode

2.2.4 View system information (Просмотр информации о системе)

Для просмотра информации о системе и текущего состояния коммутатора предусмотрено множество различных команд. В таблице ниже приведены несколько основных команд необходимых для просмотра информации о системе.

Таблица команд для просмотра информации о системе

Команда	Описание	CLI ржим
show version	Показывает номер версии системы и время её выхода.	Normal mode privileged mode
show snmp system information	Показывает основную информацию о системе, время непрерывной работы с момента запуска.	Normal mode privileged mode
show history	Показывает список последних команд введенных в командную строку.	Normal mode privileged mode

2.2.5 Network connectivity debugging (Устранение сетевых неполадок)

Для поиска причины неполадок подключения коммутатора к другому оборудованию необходимо проверить соединение между устройствами путем введения команды ping.

Если соединение установлено, то будет получен ответ от устройства (для выполнения команды ping необходимо знать IP адреса коммутатора и устройства, при этом адреса должны принадлежать одной подсети). Если ответ от устройства не получен, следует установить причину отсутствия соединения.

Таблица команды ping

Команда	Описание	CLI ржим
ping <ip-address> [-n <count> -l <size> -w <timeout>]*	Команда может использоваться как с опциональными параметрами, так и без них. Ctrl + C - прерывание выполнения команды.	Privileged mode

2.2.6 Traceroute debugging (Устранение неполадок маршрутизации)

Для поиска причины неполадок маршрутизации, при подключении через промежуточные сетевые устройства, следует ввести команду trace route и IP адрес конечного устройства. В процессе выполнения команды будут показаны IP адреса промежуточных устройств (при этом адреса должны принадлежать одной подсети).

Коммутатор поддерживает несколько опций команды trace route, что дает пользователю дополнительные возможности для более точного определения причин возможных неполадок.

Таблица команды trace route

Команда	Описание	CLI ржим
trace-route <ip-address> [-h <maximum-hops> -j <count> <ip-address>* -w <timeout>]*	Команда может использоваться как с опциональными параметрами, так и без них. Ctrl + C - прерывание выполнения команды.	Privileged mode

2.2.7 Telnet Client (Telnet клиент)

Коммутатор предоставляет пользователю возможность дистанционного подключения к другим сетевым устройствам с помощью команды telnet client.

Таблица команды telnet client

Команда	Описание	CLI ржим
telnet <ip-address>	IP адрес подключаемого сетевого устройства.	Privileged mode

2.2.8 SSH Client (SSH клиент)

Коммутатор предоставляет пользователю возможность дистанционного подключения к другим сетевым устройствам с помощью команды ssh client.

Таблица команды ssh client

Команда	Описание	CLI ржим
ssh <ip-address> [user-name]	IP адрес подключаемого сетевого устройства, имя пользователя.	Privileged mode

2.3 Profile management (Управление профилем)

Существует два типа конфигурации: текущая (current configuration) и стартовая (initial configuration). Текущая конфигурация создается с учетом требований эксплуатации коммутатора и хранится в памяти системы. Стартовая конфигурация используется при первичном включении коммутатора и хранится на внутренней flash памяти коммутатора (configuration file). При использовании соответствующих команд стартовая конфигурация может быть изменена. Команда Save (сохранить) позволяет сохранить внесенные изменения и при последующем запуске будет загружена уже измененная (текущая) конфигурация.

Файлы конфигурации имеют тот же формат что и текст командной строки (прост и интуитивно понятен). Формат файла конфигурации имеет следующий вид:

- Файл состоит из текстовых команд.
- Может быть сохранена только измененная конфигурация, неизменная конфигурация не сохраняется

- Формат команд аналогичен применяемому в командной строке. Команды объединены в сегменты, сегменты разделяются знаком "!". В режиме global configuration команды со сходными значениями объединены в секции, разделяются знаком "!".
- В режиме configuration sub mode, перед командой должен стоять пробел, в режиме global configuration mode пробел перед командой должен отсутствовать.
- Файл конфигурации заканчивается командой "end".

Команды управления файлом конфигурации включают следующие:

- View configuration information (Просмотр конфигурации)
- Save configuration (Сохранение конфигурации)
- Delete configuration profile (Удаление конфигурации)
- Download from configuration file (Загрузка конфигурации)

2.3.1 View configuration information (Просмотр конфигурации)

Пользователь имеет возможность просмотра текущей (current configuration) и стартовой (initial configuration) конфигурации системы. Файл конфигурации хранится на внутренней flash памяти коммутатора, если файл конфигурации отсутствует, то при включении коммутатора загружаются заводские настройки (по умолчанию). В этом случае при вводе команды для просмотра стартовой конфигурации, система выдаст сообщение об отсутствии файла.

Таблица команд просмотра конфигурации

Команда	Описание	CLI ржим
show running-config	Просмотр текущей конфигурации системы.	Privileged mode
show startup-config	Просмотр стартовой конфигурации системы.	Privileged mode

2.3.2 Save configuration (Сохранение конфигурации)

После внесения изменений в текущую конфигурацию системы пользователь имеет возможность сохранить эти изменения для того чтобы при последующем включении коммутатора загрузилась последняя сохраненная версия конфигурации. В противном случае при включении коммутатора будет загружена конфигурация без учета внесенных изменений.

Таблица команд сохранения конфигурации

Команда	Описание	CLI ржим
write	Сохранить текущую конфигурацию	Privileged mode

Внимание ! Если пользователь не сохранит последние внесенные в текущую конфигурацию изменения, то при выключении коммутатора изменения настроек будут потеряны. При последующем включении коммутатора будет загружена последняя сохраненная конфигурация.

2.3.3 Delete configuration file (Удаление файла конфигурации)

Если пользователю необходимо чтобы при включении были загружены заводские настройки системы (настройки по умолчанию), то следует удалить файл конфигурации. Удаление файла конфигурации не приводит к изменениям в текущей конфигурации до выключения коммутатора. Для возвращения к заводским настройкам после удаления файла следует перезагрузить коммутатор.

Таблица команд удаления файла конфигурации

Команда	Описание	CLI ржим
delete startup-config	Удаление файла конфигурации	Privileged mode

Внимание !

Будьте внимательны, удаление файла конфигурации приведет к потере текущей конфигурации после перезагрузки коммутатора.

2.3.4 Download configuration file (Загрузка файла конфигурации)

В случае повреждения или удаления файла конфигурации или необходимости вернуться к заводской конфигурации коммутатора предусмотрена возможность загрузки заводского файла с ПК в память коммутатора.

Загрузка заводского файла конфигурации не приводит к изменениям в текущей конфигурации до выключения коммутатора. Для возвращения к заводским настройкам после загрузки файла следует перезагрузить коммутатор.

Загрузка и скачивание файла конфигурации также доступна через web, подробно процедура описана в полной инструкции по эксплуатации коммутатора.

Таблица команд загрузки файла конфигурации

Команда	Описание	CLI ржим
upload configure <ip-address> <file-name>	Скачивание файла конфигурации на ПК. 1й параметр IP адрес ПК 2й параметр имя файла.	Privileged mode
download configure <ip-address> <file-name>	Загрузка файла конфигурации в коммутатор. 1й параметр IP адрес ПК 2й параметр имя файла.	Privileged mode

Загрузите файл конфигурации в коммутатор, используйте протокол TFTP. Запустите TFTP client на коммутаторе и TFTP на ПК. Далее действуйте по пунктам ниже:

Шаг 1: Сформируйте сетевое окружение.

Шаг 2: Сохраните файл конфигурации TFTP software на ПК.

Шаг 3: Сохраните файл конфигурации на коммутаторе.

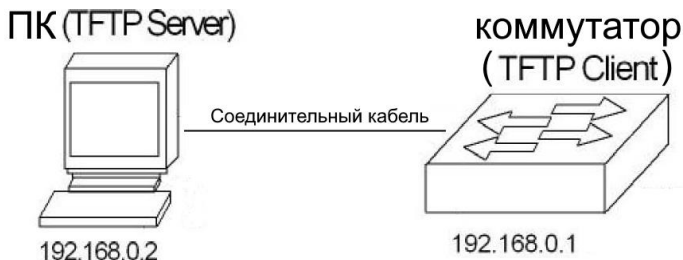
Шаг 4: Запустите команду configuration file upload на коммутаторе для скачивания файла конфигурации на ПК.

Шаг 5: При необходимости запустите команду configuration file download command на коммутаторе для загрузки файла конфигурации в коммутатор.

Шаг 6: Перезагрузите коммутатор для обновления конфигурации.

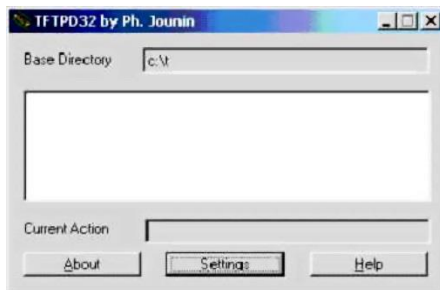
Пример: Загрузка и скачивание файла конфигурации на коммутатор с VLAN и IP адресом.

Шаг 1: Сетевое окружение.

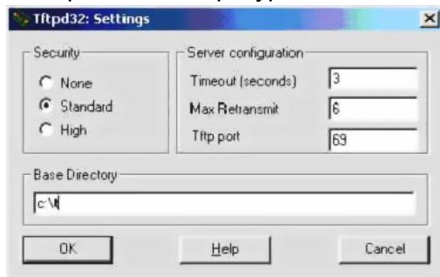


Соединить ПК с конфигурационным портом коммутатора сетевым кабелем. Установить TFTP server на ПК и установить IP адрес сетевого адаптера (например 192 168.0.2) в одной подсети с IP адресом коммутатора (по умолчанию 192.168.0.1), проверить наличие соединения.

Шаг 2: Запустите TFTP сервер на ПК и сконфигурируйте его параметры.



Откройте директорию с файлом конфигурации, далее нажмите [settings].



Введите путь в base directory и нажмите [OK] для завершения.

Шаг 3: Введите команду write на коммутаторе для Сохранения файла текущей конфигурации.

Шаг 4: Запустите команду configuration file upload на коммутаторе для скачивания файла конфигурации на ПК (192 168.0.2).

Шаг 5: Для загрузки файла конфигурации с ПК (192 168.0.2) в коммутатор запустите команду configuration file download command на коммутаторе.

Шаг 6: Для обновления конфигурации вводом команды switch#reset перезагрузите коммутатор.

2.4 Software version upgrade (Обновление ПО)

Коммутатор поддерживает online обновление внутреннего программного обеспечения. Операция выполняется с помощью возможностей TFTP.

2.4.1 Software version upgrade command (Команды обновления ПО)

В режиме global configuration mode команда обновления файла ПО:

download switch <ip-address> <file-name>

Где < IP address > IP адрес ПК с запущенным TFTP server и < file name > имя файла сохраненного на TFTP server.

В режиме global configuration mode команда обновления файла kernel:

download kernel <ip-address> <file-name>

Где < IP address > IP адрес ПК с запущенным TFTP server и < file name > имя файла kernel сохраненного на TFTP server.

В режиме global configuration mode команда обновления файла patch:

download patch <ip-address> <file-name>

Где < IP address > IP адрес ПК с запущенным TFTP server и < file name > имя файла patch сохраненного на TFTP server.

В режиме global configuration mode команда обновления файла uboot:

download uboot <ip-address> <file-name>

Где < IP address > IP адрес ПК с запущенным TFTP server и < file name > имя файла uboot сохраненного на TFTP server.

Внимание!

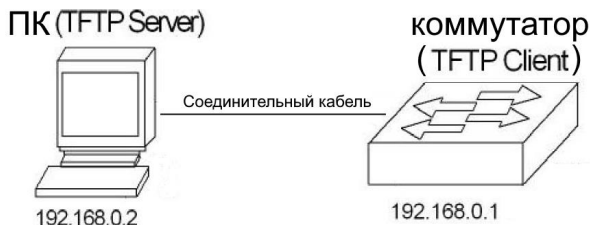
Не выключайте питание коммутатора во время процесса обновления, в противном случае произойдет повреждение файлов и коммутатор не запустится. После окончания загрузки файлов следует перезагрузить коммутатор. Процесс обновления ПО занимает несколько минут.

Обновление программного обеспечения также доступно через web, подробно процедура описана в полной инструкции по эксплуатации коммутатора.

2.4.2 Software upgrade process (Процесс обновления ПО)

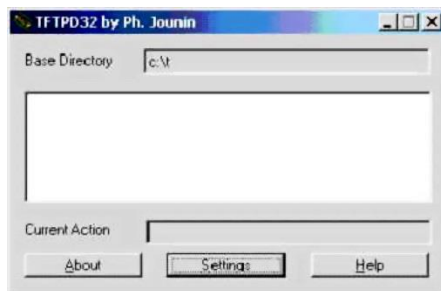
Процесс обновления ПО коммутатора с VLAN и IP адресом по шагам:

Шаг 1: Сетевое окружение.

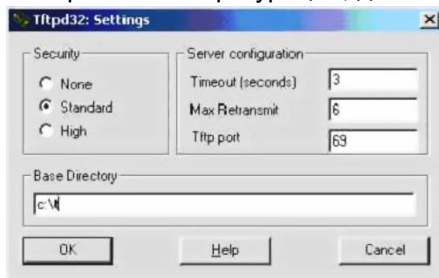


Соединить ПК с конфигурационным портом коммутатора сетевым кабелем. Установить TFTP server на ПК и установить IP адрес сетевого адаптера (например 192 168.0.2) в одной подсети с IP адресом коммутатора (по умолчанию 192.168.0.1), проверить наличие соединения.

Шаг 2: Запустите TFTP сервер на ПК и сконфигурируйте его параметры.



Откройте директорию с файлом конфигурации, далее нажмите [settings].



Введите путь в base directory и нажмите [OK] для завершения.

Шаг 3: Введите команду switch# на коммутаторе: Switch# download switch 192.168.0.2 switch.img далее нажмите enter и дождитесь завершения процесса. *Не выключайте питание коммутатора до окончания процесса!*

По окончании процесса коммутатор предложит провести перезагрузку: Do you wish to reset? [Y/N], в зависимости от необходимости следует выбрать (Y-да / N-нет). Изменения вступят в силу только после перезагрузки коммутатора.

Шаг 4: Для перезагрузки коммутатора введите команду switch#reset.

3. Configuration of ports (Конфигурация портов)

Данный раздел описывает процедуры по конфигурации портов коммутатора и включает следующие возможности:

- General configuration of ports (Основные настройки)
 - Configure (Прочие настройки)
 - Configure TORM-CONTROL (Настройка TORM-CONTROL)
 - Configure LOW-CONTROL (Настройка LOW-CONTROL)
- Configure port bandwidth (Настройка пропускной способности)
 - Configure truck (Настройка truck)

3.1 General configuration (Основные настройки портов)

Путем соответствующих настроек администратор имеет возможность контролировать доступ пользователей к портам коммутатора. Если пользователю не разрешен доступ к порту, то администратор закрывает порт. В данном разделе представлены следующие возможности:

- Opening and closing of ports (Открытие и закрытие портов)
- Port rate configuration (Настройка скорости Port rate)
- Display port information (Информация о порте)

3.1.1 Opening and closing of ports (Открытие и закрытие портов)

Порты коммутатора открыты по умолчанию. Если пользователю не разрешен доступ к порту, то администратор закрывает порт. В режиме interface configuration mode для изменения статуса порта применяются следующие команды:

no shutdown – открыть порт, например открыть port GE1/1:

Switch(config-ge1/1)#no shutdown

Shutdown – закрыть порт, например закрыть port GE1/1:

Switch(config-ge1/1)#shutdown

3.1.2 Port rate configuration (Настройка скорости портов)

Режим работы портов по умолчанию автоматический. В режиме interface configuration mode для изменения режима работы порта (port rate) применяются следующие команды:

speed {autonegotiate |full-1000 |full-100 |full-10 |half-100 |half-10 }

autonegotiate – автоматический (или адаптивный)

full-1000 - Full duplex Gigabit

full-100 - Full duplex 100M

full-10 - Full duplex 10M

half-100 - Half duplex 100M

half-10 - Half duplex 10M

Например: скорость порта 1/1 full duplex 100M

Switch(config-ge1/1)# speed full-100

3.1.3 Display port information (Информация о порте)

Команда вывода информации об одном или нескольких портах в режимах normal mode или privileged mode:

`show interface [if-name]`

Например: показать информацию о port GE1/1

Switch# `show interface ge1/1`

Например: показать информацию о всех портах

Switch# `show interface`

3.2 Cofigure mirror (Зеркалирование портов)

Функция зеркалирования портов (Port mirroring) предназначена для мониторинга трафика принимаемых и отправляемых пакетов одним или несколькими портами. Зеркалированный порт может отслеживать принимаемые и отправляемые данные другими портами. В данном разделе рассматривается конфигурация параметров зеркалирования:

- Configure the listening port and monitored port of mirror (Настройка конфигурации listening портов и monitored портов)
- Show mirror configuration (Просмотр конфигурации)

3.2.1 Configure listening port and monitored port (Настройка портов)

Для настройки, в режиме interface configuration mode администратор выбирает listening порт и monitored порт, например GE1/1 и GE1/2 и вводит команду:

Switch(config-ge1/1)# `mirror interface ge1/2 direction both`

Таким образом порт GE1/1 переходит в режим listening, а порт GE1/2 в режим monitored, при этом порт GE1/1 отслеживает как отправляемые, так и принимаемые портом GE1/2 пакеты.

Для настройки параметров monitored портов применяются следующие команды:

Switch(config-ge1/1)#`mirror interface <if-name> direction {both | receive | transmit}`

Таким образом порт GE1/1 становится listening, < if name > становится monitored, а {both | receive | transmit} определяют параметры мониторинга: receive – отслеживание принимаемых пакетов; transmit – отслеживание отправляемых пакетов; both – отслеживание как отправляемых, так и принимаемых пакетов.

Для удаления monitored порта в режиме interface configuration mode, применяется следующая команда: Switch(config-ge1/1)#no mirror interface <if-name>, где < if name > удаляемый порт, например:

```
Switch(config-ge1/1)# no mirror interface ge1/2
```

Таким образом порт GE1/1 перестает отслеживать трафик пакетов порта GE1/2. При удалении всех monitored портов будет удален и listening порт.

3.2.2 Display configuration of mirror (Просмотр конфигурации)

Для просмотра конфигурации зеркалирования в обычном или привилегированных режимах следует ввести команду:

```
Switch# show mirror
```

Внимание ! Один и тот же порт не может быть одновременно назначен listening и monitored. Может быть назначено несколько портов monitored и только один порт listening.

3.3 Configure storm-control (Подавление «шторма»)

Функция Storm Control обеспечивает защиту от влияния широковещательных (Multicast) пакетов и DLF пакетов на передаваемый/получаемый портами коммутатора трафик, предотвращает перегрузку сети циркулирующими пакетами данных. Данную функцию поддерживают все порты коммутатора. В данном разделе представлены следующие возможности по настройке Storm Control:

- Default configuration (Конфигурация по умолчанию)
- Broadcast suppression configuration (Подавление Broadcast)
- Multicast suppression configuration (Подавление Multicast)
- DLF suppression configuration (Подавление DLF)
- Displays the storm-control configuration (Просмотр конфигурации)

3.3.1 Default configuration (Конфигурация по умолчанию)

Коммутатор поддерживает установки broadcast, multicast и DLF отдельно для каждого порта. Каждая установка имеет три уровня подавления. По умолчанию на портах включено подавление broadcast (уровень подавления 64K) для предотвращения сетевого шторма. Подавление DLF пакетов и multicast пакетов отключено.

3.3.2 Broadcast suppression configuration (Broadcast подавление)

Для включения подавления broadcast на порту в режиме interface configuration введите команду:

```
storm-control broadcast
```

Для отключения подавления broadcast на порту в режиме interface configuration введите команду:

```
no storm-control broadcast
```

3.3.3 Multicast suppression configuration (Multicast подавление)

Для включения подавления multicast на порту в режиме interface configuration введите команду:

```
storm-control multicast
```

Для отключения подавления multicast на порту в режиме interface configuration введите команду:

```
no storm-control multicast
```

3.3.4 DLF suppression configuration (Multicast подавление)

Для включения подавления DLF на порту в режиме interface configuration введите команду:

```
storm-control dlf
```

Для отключения подавления DLF на порту в режиме interface configuration введите команду:

```
no storm-control dlf
```

3.3.5 Suppression rate configuration (Уровни подавления)

Для выбора уровня подавления на порту в режиме interface configuration введите команду:

```
storm-control ratelimit <1-1024000>
```

Где параметр изменяется в пределах 1-1024000 (по умолчанию 64К).

3.3.6 Display storm-control configuration (Просмотр конфигурации)

Для просмотра конфигурации storm-control в обычном или привилегированных режимах следует ввести команду:

```
show storm-control
```

3.4 Configure storm-constraint (Настройка storm-constraint)

Управление порогом трафика на портах коммутатора применяется для предотвращения возникновения сетевого шторма в сети Ethernet. Порт, на котором включена функция будет регулярно определять трафик поступающих unicast, multicast и broadcast сообщений. Если количество сообщений определенного типа превысит пороговое значение, пользователь может заблокировать, закрыть порт или сохранить это событие в логах.

Блокировка порта (Block mode) – при превышении порога определенным типом сообщений порт блокирует отправку сообщений, но продолжает вести учет трафика, при уменьшении трафика порт возобновляет отправку сообщений данного типа.

Закрытие порта (Shutdown mode) – при превышении порога порт закрывается и прекращает отправку сообщений любого типа, пользователь может открыть порт с помощью команды или отключить порог учета трафика в настройках.

Примечание: Трафик сообщений одного типа может быть ограничен как описанным выше методом, так и функцией storm-control, но эти функции не могут быть применены одновременно. В противном случае эффект подавления может быть неопределенным. Например, установление порога для сообщений multicast не может быть использовано совместно с функцией multicast suppression.

Таблица команд конфигурации storm-constraint

Команда	Описание	CLI режим
storm-constrain (broadcast multicast unicast) min-rate <1-1488100> max-rate <1-1488100>	Включение функции и выбор broadcast, multicast, unicast сообщений. Уровень подавления.	Interface configuration mode
no storm-constrain (broadcast multicast unicast all)	Отключение storm control.	Interface configuration mode
storm-constrain action (block shutdown)	Выбор действия storm control (блокировка или закрытие порта). По умолчанию отключено.	Interface configuration mode
no storm-constrain action	Отключение действия (блокировки или закрытия порта).	Interface configuration mode
storm-constrain enable (trap)	Включение логирования (при включенном storm control).	Interface configuration mode
no storm-constrain enable (log)	Отключение логирования (при включенном storm control).	Interface configuration mode
storm-constrain interval <6-180>	Выбор интервала отслеживания storm control. По умолчанию установлен интервал 5 сек.	Interface configuration mode
no storm-constrain interval	Установка интервала отслеживания storm control на значение по умолчанию.	Interface configuration mode
no storm-constrain	Отключение storm control.	Interface configuration mode
show storm-constrain	Просмотр установок storm control для всех интерфейсов.	Privileged mode
show storm-constrain interface IFNAME	Просмотр установок storm control для выбранного интерфейса.	Privileged mode

Таблица значений при просмотре конфигурации storm-constraint

	Описание
interface	Имя интерфейса
type	Тип сообщения: (1) broadcast; (2) multicast; (3) unicast.
rate	Уровень: Min – низкий порог; Max – высокий порог.
action	Действие: (1) block - блокировка; (2) shutdown – закрытие.
punish-status	Статус порта: (1) block – заблокирован (превышен установленный порог); (2) normal – обычная работа; (3) shutdown – закрыт (превышен установленный порог).
log	Логирование: Logs on / off (включено/отключено).
interval	Интервал отслеживания storm control. По умолчанию установлен интервал 5 сек.
last-punish-time	Время последнего срабатывания функции storm control.

При работе функции storm-constraint action следует учесть следующие особенности.

Если выбран режим блокировки порта (Block mode) при превышении максимального порога определенным количеством сообщений, порт блокирует отправку сообщений, но продолжает вести учет трафика, при уменьшении трафика до минимального порогового значения порт возобновляет отправку сообщений. Пользователь может сконфигурировать значения максимального и минимального порогов.

Если выбран режим закрытия порта (Shutdown mode) при превышении порога порт закрывается и прекращает отправку сообщений любого типа. Автоматическое возобновление работы порта не предусмотрено, пользователь должен отключить функцию shutdown для данного порта.

Когда трафик порта превышает максимальное пороговое значение или опускается к минимальному пороговому значению, фиксируется Log / trap информация.

3.5 Configure flow-control (Настройка flow-control)

Функция Flow-control предназначена для предотвращения потери пакетов в случае если порт заблокирован. В режиме half duplex mode, функция flow control реализуется с использованием технологии back pressure, которая понижает скорость передачи данных источником. В режиме full duplex mode, функция flow control соответствует стандарту ieee802.1 3х, заблокированный порт отправляет источнику данных пакет "pause" для временной приостановки отправки данных.

В данном разделе представлены следующие возможности по настройке функции flow-control:

- Default configuration (Конфигурация по умолчанию)
- Set port flow control (Включение функции flow-control)
- Close port flow control (Отключение функции flow-control)
- Display flow control information (Просмотр конфигурации)

3.5.1 Default configuration (Конфигурация по умолчанию)

Коммутатор поддерживает установку функции flow-control отправляемых и принимаемых пакетов отдельно для каждого порта. По умолчанию на портах функция flow-control отключена.

3.5.2 Set port flow control (Включение flow control)

Для включения функции flow-control на порту в режиме interface configuration mode введите команду:

```
flowcontrol
```

3.5.3 Close port flow control (Отключение flow control)

Для отключения функции flow-control на порту в режиме interface configuration mode введите команду:

```
no flowcontrol
```

3.5.4 Display flow control information (Просмотр конфигурации)

Для просмотра конфигурации flow-control на всех портах в обычном или привилегированных режимах следует ввести команду:

```
show flowcontrol
```

Для просмотра конфигурации flow-control на выбранном порту в обычном или привилегированных режимах следует ввести команду:

```
show flowcontrol interface <if-name>
```

где < if name > номер (или имя) выбранного порта.

3.6 Configure port bandwidth (Настройка скорости портов)

Функция Port bandwidth используется для управления скоростью принимаемого и отправляемого портом трафика.

В данном разделе представлены следующие возможности по настройке функции port bandwidth:

- Default configuration (Конфигурация по умолчанию)
- Set the port sending or receiving bandwidth control (Включение функции port bandwidth)
- Cancel port send or receive bandwidth control (Отключение функции port bandwidth)
- Displays the bandwidth control of the port configuration (Просмотр конфигурации)

3.6.1 Default configuration (Конфигурация по умолчанию)

Коммутатор поддерживает установку функции port bandwidth отдельно для каждого порта. По умолчанию на портах функция port bandwidth отключена.

3.6.2 Set port bandwidth (Включение port bandwidth)

Для включения функции port bandwidth на порту в режиме interface configuration mode введите команду:

`portrate {egress | ingress} <rate>`

Egress означает управление скоростью передачи отправляемых пакетов.

Ingress означает управление скоростью передачи принимаемых пакетов.

`< rate >` определяет скорость передачи в диапазоне 1-1024000 Кбит.

3.6.3 Cancel port bandwidth (Отключение port bandwidth)

Для отключения функции port bandwidth на порту в режиме interface configuration mode введите команду:

`no portrate {egress | ingress}`

Egress означает отключение управления скоростью передачи отправляемых пакетов.

Ingress означает отключение управления скоростью передачи принимаемых пакетов.

3.6.4 Display bandwidth control of port configuration (Просмотр конфигурации)

Для просмотра конфигурации port bandwidth на портах в обычном или привилегированных режимах следует ввести команду:

`show portrate interface <if-name>`

где `< if name >` номер (или имя) выбранного порта.

3.7 Configure trunk (Конфигурация trunk)

Trunk позволяет агрегировать несколько портов коммутатора в один логический порт. Trunk может быть использован для увеличения пропускной способности, создания резервных соединений, распределения нагрузки. Когда trunk группа используется как логический порт, коммутатор выбирает порт из группы портов для отправки пакетов согласно политики агрегирования, установленной пользователем. Порт и политика агрегирования trunk группы конфигурируется программой, но отправление потока данных выполняется через физический порт.

Все порты в trunk группе должны быть настроены на одинаковую скорость в режиме full duplex. Коммутатор поддерживает до 32-х trunk групп, каждая группа может включать до 8-и членов trunk. Каждый порт может принадлежать только одной trunk группе.

LACP это протокол, основанный на стандарте IEEE802 3aD. LACP протокол предназначен для взаимодействия с устройством на противоположном конце линии через lacpdu (link aggregation control protocol data unit).

Интерфейс агрегированной группы включает LACP протокол. Интерфейс уведомляет устройство на противоположном конце линии о приоритете своей системы LACP, системе Mac, приоритете LACP протокола порта, номере порта и операционном ключе путем отправки lacpdu. После получения lacpdu, устройство на противоположном конце сравнивает информацию с информацией полученной другими интерфейсами и устанавливает соединение с соответствующим интерфейсом.

Операционный ключ это автоматически генерированная комбинация агрегационного контроля согласно конфигурации портов-членов (скорость порта, режим duplex, статус, VLAN, VLAN ID, тип соединения). В агрегационной группе порты-члены в выбранном состоянии имеют один и тот же операционный ключ.

В данном разделе представлены следующие возможности по настройке trunk конфигурации:

- LACP protocol configuration (Конфигурация протокола LACP)
- TRUNK Group configuration (Конфигурация TRUNK группы)
- TRUNK Member port configuration (Конфигурация порта-члена)
- TRUNK Load balancing policy configuration (Конфигурация распределения нагрузки)
- TRUNK Display (Просмотр конфигурации)

3.7.1 LACP protocol configuration (Настройка протокола LACP)

Таблица команд конфигурации LACP

Команда	Описание	CLI режим
lacp system-priority <1-65535>	Установка приоритета LACP	Global configuration mode

no lacp system-priority	Отключение приоритета (по умолчанию 32768)	Global configuration mode
lacp max-active-link-number <1-8>	Верхний предел LACP для агрегированного порта	Global configuration mode
no lacp max-active-link-number	Восстановление верхнего предела LACP для агрегированного порта (по умолчанию 8)	Global configuration mode
lacp port-priority <1-65535>	Установка LACP приоритета порта	Interface configuration mode
no lacp port-priority	Восстановление LACP приоритета порта (по умолчанию 32768)	Interface configuration mode
lacp timeout (short long)	Установка LACP задержки на порту	Interface configuration mode
show lacp summary	Простой просмотр всех LACP агрегаций	Privileged mode
show lacp detail	Просмотр всех LACP агрегаций	Privileged mode
show lacp <1-8>	Просмотр деталей LACP агрегаций	Privileged mode
show lacp port IFNAME	Просмотр деталей LACP порта	Privileged mode
show lacp system-id	Просмотр LACP статуса системы	Privileged mode
show lacp counter <1-8>	Просмотр статистики LACP агрегированного порта	Privileged mode
show lacp counter	Просмотр всех LACP агрегированных портов	Privileged mode
clear lacp <1-8> counters	Обнуление статистики LACP агрегированного порта	Privileged mode
clear lacp counters	Обнуление статистики всех LACP агрегированных портов	Privileged mode

3.7.2 Trunk Group Configuration (Конфигурация Trunk группы)

В данном разделе представлены команды для конфигурации trunk группы (агрегирование) в режиме global configuration mode.

Создать trunk группу:

```
trunk <trunk-id>
```

где параметр (номер группы) < trunk ID > выбирается в пределах 1-32. Имя интерфейса (trunk группы) trunk + ID.

Например, имя интерфейса trunk группы с номером ID1 - trunk1. Для входа в интерфейс в режиме configuration mode используется команда "interface trunk + ID number" (interface trunk 1), далее можно продолжить конфигурирование группы.

Для создания LACP trunk группы в режиме global configuration mode введите команду:

```
trunk <1-32> dynamic
```

Для уделения LACP trunk группы в режиме global configuration mode введите команду:

```
no trunk <trunk-id>
```

При удалении trunk группы убедитесь, что в ней отсутствуют порты-члены.

3.7.3 TRUNK Member port configuration (Конфигурация портов)

Для добавления порта-члена в trunk группу в режиме interface configuration mode введите команду:

```
trunk interface IFNAME (passive)
```

где < if name > номер (имя) добавляемого в trunk группу порта. Интерфейс должен быть уровня layer-2. Trunk группа может содержать до 8-и портов (интерфейсов) уровня layer-2. В случае если trunk группа является статической LACP trunk группой, добавленный порт (интерфейс) по умолчанию будет в активном состоянии, которое можно изменить на пассивное.

Для удаления всех портов-членов из trunk группы в режиме interface configuration mode введите команду:

```
no trunk interface
```

Для удаления одного из портов-членов из trunk группы в режиме interface configuration mode введите команду:

```
no trunk interface <if-name>
```

Для удаления нескольких портов из trunk группы используйте предыдущую команду несколько раз.

3.7.4 TRUNK Load balancing configuration (Распределение нагрузки)

Для включения функции оптимального распределения нагрузки (трафика) в trunk группе в режиме interface configuration mode введите команду:

```
trunk load-balance {dst-mac | dst-ip | src-dst-mac | src-dst-ip | src-mac | src-ip}
```

Для оптимизации распределения трафика могут быть использованы следующие методы:

- DST-MAC метод, основанный на MAC адресе назначения;
- DST-IP метод, основанный на IP адресе назначения;
- SRC-DST-MAC метод, основанный как на исходном MAC, так и на MAC адресе назначения;
- SRC-DST-IP метод, основанный как на исходном IP, так и на IP адресе назначения.
- SRC-MAC метод, основанный на исходном MAC адресе;
- SRC-IP метод, основанный на исходном IP адресе.

Для возврата к установкам по умолчанию функции оптимального распределения нагрузки (трафика) в trunk группе в режиме interface configuration mode введите команду:

```
no trunk load-balance
```

SRC MAC метод, основанный на исходном MAC адресе используется по умолчанию.

3.7.5 Trunk Display (Просмотр конфигурации)

Для просмотра конфигурации всех trunk групп в обычном или привилегированном режиме введите команду:

```
show trunk
```

Для просмотра конфигурации выбранной trunk группы в обычном или привилегированном режиме введите команду:

```
show trunk <trunk-id>
```

где < trunk ID > ID номер trunk группы.

3.8 Extra large frames (Jumbo frames)

Коммутатор имеет возможность работать с кадрами увеличенного размера (extra large или Jumbo frame). Для реализации этой возможности необходимо настроить порты коммутатора таким образом, чтобы они могли принимать Jumbo frame.

3.8.1 Configure Jumbo frames (Настройка Jumbo frames)

Для того чтобы порт мог принимать Jumbo frame в режиме port configuration mode введите команду:

```
Switch(config-ge1/1)#jumbo frame 2000
```

3.8.2 Display Jumbo frames (Просмотр конфигурации)

Для просмотра конфигурации Jumbo frame порта в обычном или привилегированном режиме введите команду:

```
Switch#show jumbo frame ge1/1
```

где < ge1/1 > номер порта, 2000 – размер Jumbo frame (bytes).

3.9 Configure redundant ports (Резервирование портов)

В некоторых случаях для обеспечения стабильности сетевых подключений возникает необходимость создать резервные подключения (дублирование подключений). В этом случае одно подключение обеспечивается двумя портами коммутатора. Если порт передающий данные отключается, то для передачи данных система подключает резервный порт.

Активный порт в группе резервирования называется *Conversely*, неактивный порт группы называется *disable*.

В данном разделе представлены следующие возможности по настройке *redundant* (резервных) портов:

- Configuration of redundant ports (Конфигурация резервных портов)
- Display of redundant ports (Просмотр конфигурации)

3.9.1 Configuration of redundant ports (Настройка резервных портов)

Коммутатор поддерживает 8 групп резервирования, каждая из которых включает в себя 2 порта (первичный и вторичный). Каждый порт может входить только в одну группу.

При конфигурировании портов в группе следует учесть следующие пункты:

1. Если одновременно подключены оба порта группы, *primary port* будет находиться в состоянии *active state*, а *secondary port* будет находиться в состоянии *disable state*;
2. Если подключен только один порт группы, то он будет находиться в состоянии *active state*, а второй порт в состоянии *disable state*;
3. Если произойдет отключение соединения порта в состоянии *active state*, второй порт будет пытаться установить соединение (перейти в состояние *active state*);

Настройка принудительного переключения (*force switch*). При включении *force switch*, если соединение восстанавливается на *primary port*, когда *secondary port* находится в состоянии *active*, а *primary port* отключен (*disabled*), пользователю следует выбрать между текущим подключением и исходным. Если данная настройка отключена, то будет восстановлено исходное подключение.

Таблица команд конфигурации резервных портов

Команда	Описание	CLI режим
redundant-port <1-8> primary-port IFNAME secondary-port IFNAME [force-switch]	Группа резервирования < 1-8 > номер группы Ifname имя первичного порта, ifname имя вторичного порта, включение force switch	Global configuration mode
redundant-port <1-8> force-switch	Включение force switch на портах группы резервирования.	Global configuration mode
no redundant-port <1-8>	Удаление группы резервирования.	Global configuration mode
no redundant-port <1-8> force-switch	Отключение force switch на портах группы резервирования.	Global configuration mode

3.9.2 Display redundant ports (Просмотр конфигурации)

Для просмотра конфигурации всех групп резервирования (redundant) портов в обычном или привилегированном режиме введите команду:

```
show redundant-port
```

3.10 UDLD configuration (Конфигурация UDLD)

UDLD (unidirectional link detection) – протокол уровня layer-2, используется для мониторинга физической конфигурации соединений Ethernet по витой паре и оптоволоконному кабелю. Протокол UDLD может определять ненаправленные соединения, соединения в одном направлении, отключать интерфейсы и отправлять тревожные оповещения. Ненаправленные соединения могут создавать проблемы при развертывании spanning trees и маршрутизации.

Внимание: необходимо чтобы протокол UDLD поддерживался устройствами на обоих концах линии.

Протокол UDLD имеет два режима работы: Normal mode (по умолчанию) и aggressive mode.

- Normal (обычный) mode: отключение порта при длительном разрыве соединения.
- Aggressive (агрессивный) mode: протокол UDLD определяет ненаправленное подключение путем попытки установления соединения отправкой сообщения (длиной 8сек). Если ответ UDLD не получен, то порт переводится в состояние errdisable state.

Таблица команд конфигурации протокола UDLD

Команда	Описание	CLI режим
udld enable	Включение протокола UDLD.	Global configuration mode
udld message time <time>	Установление интервала отправки UDLD сообщений.	Global configuration mode
udld port	Включение протокола UDLD на порту.	Interface configuration mode
udld aggressive	Включение режима aggressive. Обычный режим вкл по умолчанию.	Interface configuration mode
show udld <ifname>	Просмотр информации UDLD интерфейса.	Privileged mode

3.11 LLDP configuration (Конфигурация LLDP)

LLDP (Link Layer Discovery Protocol) – стандартный протокол link layer discovery method, используется для организации управления адресами, идентификации оборудования и другой информации о локальном оборудовании с различными типами TLV (type / length / value), инкапсуляции lldpdu (Link Layer Discovery Protocol Data Unit) и ближайшего подключенного оборудования. После получения этой информации, ближайшие устройства сохраняют её в стандартной базе MIB (management information base) для установления соединения и коммуникационного статуса.

В данном разделе представлены следующие возможности по настройке LLDP протокола:

- LLDP configuration (Конфигурация протокола LLDP)
- LLDP Display (Просмотр конфигурации)

3.11.1 LLDP configuration (Настройка LLDP)

Существует четыре режима работы порта с LLDP:

TXRX: отправление и прием LLDP сообщений.

TX: только отправление LLDP сообщений.

Rx: только прием LLDP сообщений.

Disable: отправление и прием LLDP сообщений отключен.

Когда рабочий режим LLDP на порту изменяется, порт инициирует процедуру protocol state machine. Для предотвращения частой смены режима работы порта, приводящей к длительной инициализации, предусмотрена настройка времени задержки инициализации.

Таблица команд конфигурации протокола LLDP

Команда	Описание	CLI режим
lldp global enable	Включение протокола LLDP.	Global configuration mode
lldp hold-multiplier <num>	Включение LLDP TTL multiple.	Global configuration mode
lldp timer [<reinit-delay><time>][< tx-delay><time>][< tx-interval ><time>]	Настройка LLDP интервалов времени задержки.	Global configuration mode
lldp enable	Включение LLDP (интерфейса).	Interface configuration mode
lldp admin-status{ disable rx tx rxtx}	Настройка режима работы порта LLDP.	Interface configuration mode
lldp check-change-interval <time>	Настройка времени обновления интерфейса.	Interface configuration mode

lldp management-address <A.B.C.D>	Настройка LLDP управления адресом интерфейса.	Interface configuration mode
lldp tlv-enable{ dot1-tlv dot3-tlv med-tlv }	Настройка интерфейса LLDP с расширенными возможностями.	Interface configuration mode

3.11.2 Display LLDP (Просмотр конфигурации)

Таблица команд просмотра конфигурации протокола LLDP

Команда	Описание	CLI режим
show lldp configuration [ifname]	Просмотр общей конфигурации LLDP	Privileged mode
show lldp local-information [ifname]	Просмотр локальной конфигурации LLDP	Privileged mode
show lldp neighbor-information [ifname]	Просмотр LLDP информации о ближайших подключениях	Privileged mode
show lldp statistics [ifname]	Просмотр статистики LLDP сообщений	Privileged mode
show lldp status [ifname]	Просмотр информации о статусе LLDP	Privileged mode

4. Configure port based security (Конфигурация безопасности портов)

Данный раздел описывает процедуры по конфигурации безопасности портов коммутатора на основе MAC и включает следующие возможности:

- MAC Binding configuration (Привязка MAC)
- MAC Filter configuration (Фильтрация MAC)
- Port learning restriction configuration (Ограничения на портах)
- Protection port configuration (Защита портов)

4.1 Introduction (Введение)

Безопасность портов на основе MAC может быть обеспечена четырьмя функциями: MAC binding (привязка MAC), MAC filtering (фильтрация MAC), port learning control (ограничения на портах) и port protection (защита портов). Эти функции повышают возможности по обеспечению безопасности уровня layer 2 при эксплуатации коммутатора.

Привязка MAC к порту и ограничение выделенных MAC адресов для подключения порта к сети (к одному порту может быть привязано несколько MAC) соответствует стандарту 802.1x. Эта функция может быть использована для подключения устройств не поддерживающих 802.1x таких как принтеры, файловые серверы и т.д.

Фильтрация MAC адресов ограничивает доступ к сети выделенных MAC. В основном используется для ограничения доступа к сети нелегальных устройств. Выделенный MAC адрес не может получать пакеты данных от портов коммутатора. Каждый порт имеет возможность фильтровать несколько MAC адресов. Данная функция может применяться совместно с функцией ACL для отражения вирусной сетевой атаки с использованием поддельных MAC адресов.

Функция Port learning control может ограничивать количество MAC адресов, динамически связанных с портом. Когда количество MAC адресов превышает заданное значение, порт перестает принимать пакеты поступающие с нового MAC адреса.

Функции MAC binding и 802.1x могут быть одновременно настроены на одном порту. Функции Mac filtering и port learning restrictions могут быть одновременно настроены на одном порту. Функции Mac binding, 802.1x, MAC filtering и port learning restrictions не могут быть настроены на одном порту одновременно.

4.2 MACBinding configuration (Настройка MACBinding)

Функция MAC binding поддерживает ручную и автоматическую привязку MAC адресов. При использовании ручной привязки адресов к портам пользователь последовательно вводит MAC адреса используя команды. При автоматической привязке производится считывание

портом данных в layer 2 (L2) непосредственно из таблицы MAC адресов после ввода соответствующей команды.

Таблица команд конфигурации MAC binding

Команда	Описание	CLI режим
switchport-security mac-bind НННН.НННН.НННН vlan <1-4094> qosprofile {qp0 qp1 qp2 qp3 qp4 qp5 qp6 qp7}	Ручная привязка MAC к порту с указанием приоритета.	Interface configuration mode
switchport-security mac-bind auto-conversion number <1-8191> qosprofile {qp0 qp1 qp2 qp3 qp4 qp5 qp6 qp7}	Автоматическая привязка MAC к порту с указанием приоритета.	Interface configuration mode
switchport-security mac-bind auto-conversion vlan <1-4094> qosprofile {qp0 qp1 qp2 qp3 qp4 qp5 qp6 qp7}	Автоматическая привязка MAC к определенному VLAN с указанием приоритета.	Interface configuration mode
show port-security mac-bind [IFNAME]	Просмотр конфигурации MAC binding.	Privileged mode

Внимание: Возможные причины сбоев при выполнении привязки:

Порт имеет конфигурацию 802.1x

На порту включена функция MAC filtering или learning restrictions;

MAC адрес привязан к другим портам или имеет конфигурацию MAC filtering;

Таблица MAC адресов L2 коммутатора заполнена.

4.3 MACFilter configuration (Настройка MACFilter)

Функция MAC filtering поддерживает ручную и автоматическую привязку MAC адресов. При использовании ручной привязки фильтруемых адресов к портам пользователь последовательно вводит MAC адреса используя команды. При автоматической привязке производится считывание портом данных в layer 2 (L2) непосредственно из таблицы MAC адресов после ввода соответствующей команды.

Таблица команд конфигурации MAC filtering

Command	Description	CLI mode
switch port-security mac-filter HHHH.NHHH.NHHH vlan <1-4094>	Ручная привязка MAC к порту.	Interface configuration mode
switch port-security mac-filter auto-conversion number <1-8191>	Автоматическая привязка MAC к порту.	Interface configuration mode
switch port-security mac-filter HHHH.NHHH.NHHH vlan <1-4094>	Автоматическая привязка MAC к определенному VLAN.	Interface configuration mode
show port-security mac-filter [IFNAME]	Просмотр конфигурации MAC filter.	Privileged mode

Внимание: Возможные причины сбоев при выполнении привязки:

Порт имеет конфигурацию MAC binding или 802.1x

MAC адрес привязан к другим портам или имеет конфигурацию MAC binding;

Таблица MAC адресов L2 коммутатора заполнена.

4.4 Port learning restriction configuration (Настройка Port learning restriction)

Функция Port learning restriction ограничивает количество MAC адресов связанных с одним портом коммутатора. Максимально возможное количество связанных MAC адресов с портом 8191. Если количество связанных с портом MAC адресов превысит установленное значение, порт останавливает обмен пакетами с устройствами с новыми адресами.

Таблица команд конфигурации Port learning restriction

Команда	Описание	CLI режим
switchport port-security learn-limit <0-8191>	Установка количества MAC адресов связанных с портом.	Interface configuration mode

no switchport port-security learn-limit	Отмена ограничения количества MAC адресов связанных с портом.	Interface configuration mode
show port-security learn-limit [IFNAME]	Просмотр конфигурации port learning.	Privileged mode

Пример: Порт GE1/5 связан максимум с 7ю MAC адресами

```
Switch(config)interface ge1/5
```

```
Switch(config-ge1/5)switchport port-security learn-limit 7
```

Внимание: Возможные причины сбоев при выполнении настройки:

Порт имеет конфигурацию MAC binding или 802.1x

4.5 Protection port configuration (Настройка Protection port)

4.5.1 Introduction to protection port (Защита порта)

Для реализации two-layer уровня изоляции между пакетами данных, порты коммутатора могут быть разделены между различными VLAN. При настройке портов в режим protection ports, порты относящиеся к одному VLAN становятся изолированными, что повышает безопасность и расширяет возможности по настройке сети.

Каждый порт может быть настроен как protected port. В этом режиме порты не могут обмениваться данными друг с другом, но могут поддерживать коммуникацию с незащищенными портами. Существует два режима работы:

- Один порт работает в незащищенном режиме, остальные порты в режиме protection port (изолированы);
- Несколько портов работают в режиме protection ports, остальные порты в небезопасном режиме.

4.5.2 Protection port configuration (Настройка Protection port)

Для того чтобы перевести порт (например GE1/1) в статус protected port в режиме port configuration mode введите команду:

```
Switch(config-ge1/1)#switchport port-security protect
```

Для просмотра портов со статусом protected port в обычном или привилегированном режиме введите команду:

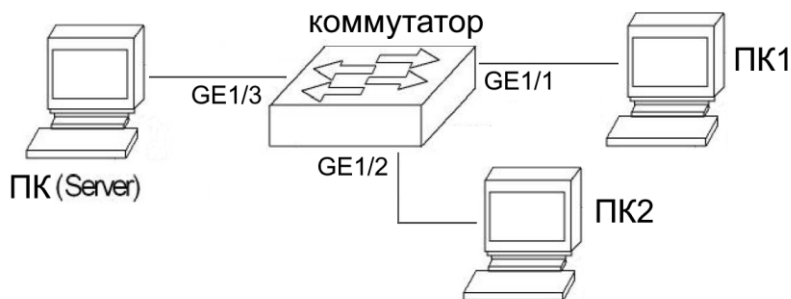
```
Switch#show port-security protect
```

После ввода команды будут показаны порты в режиме protected:

Port	Port protected
-----	-----
ge1/1	ON

Пример:

К портам коммутатора GE1/1, GE1/2 и GE1/3 соответственно подключены ПК1, ПК2 (находятся в одном VLAN) и сервер. ПК1, ПК2 должны поддерживать коммуникацию только с сервером (но не между собой).



Команды настройки конфигурации описанной выше:

```
Switch#config terminal
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport port-security protect
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#switchport port-security protect
Switch(config-ge1/2)#exit
Switch(config)#exit
Switch#show port-security protect
Port    Port protected
-----  -----
ge1/1    ON
ge1/2    ON
```

5. Configure port IP and MAC binding (Конфигурация привязки IP и MAC)

Данный раздел описывает процедуры по конфигурации привязки к портам коммутатора MAC и IP адресов и включает следующие пункты:

- IP and MAC Binding configuration (Привязка IP и MAC)
- Configuration example (Пример конфигурации)
- Configuration troubleshooting (Возможные сбои)

5.1 Introduction (Введение)

Конфигурация привязки IP и MAC адресов к портам коммутатора уровня layer 2 является мерой статической защиты против ARP. Путем отправки ARP сообщений с фальшивыми MAC адресами происходит замена в таблице локального ARP кэша на подставные MAC адреса. Одновременно происходит отправка данных атакующим злоумышленникам. Настройка привязки статических IP и MAC адресов к портам коммутатора дает возможность эффективно фильтровать ARP пакеты данных при атаке.

В добавление к возможности предотвращения ARP spoofing, привязка IP и MAC адресов обеспечивает соединение точка-точка между IP и MAC адресами (один IP может обмениваться с одним MAC и один MAC может обмениваться с одним IP). При изменении схемы подключения, соединение устройств в сети станет невозможным. Динамическая защита реализуется в соответствии со стандартом 802.1x anti ARP Spoofing и протоколом DHCP snooping.

В коммутаторе реализованы четыре различных функции: IP MAC binding, ACL, 802.1x anti ARP Spoofing and DHCP snooping. При настройке конфигурации следует учитывать их совместимость.

Таблица совместимости функций при настройке конфигурации

Функция	IP MAC binding	ACL	802.1x	DHCP SNOOPING
IP MACbinding	-	X	O	O
ACL	X	-	X	X
802.1x	O	X	-	X
DHCP SNOOPING	O	X	X	-

CFP является ограниченным ресурсом, в среднем к каждому порту коммутатора возможно привязать не более 16-и IP MAC адресов. Обычно для сети со множеством доступных хостов достаточно контролировать несколько IP и MAC адресов, для чего подходит функция static IP MAC. Для предотвращения сбоев при передаче трафика следует избегать исчерпания ресурсов CFP.

В зависимости от ситуации используются протоколы 802.1x или DHCP snooping. При использовании static IP address configuration и протокола 802.1x (для доступа к сети), для достижения эффективности необходимо применять 801.1x anti ARP Spoofing. Если используются динамические IP адреса, следует применять протокол DHCP snooping.

5.2 IP and MAC binding configuration (Настройка IP и MAC binding)

Команды подключения IP и MAC binding:

```
Switch#configure terminal
Switch(config)#interface ge1/1
Switch(config-ge1/1)#ip mac-bind A.B.C.D MAC
```

Команды отключения IP and MAC binding:

```
Switch#configure terminal
Switch(config)#interface ge1/1
Switch(config-ge1/1)#no ip mac-bind A.B.C.D MAC
```

Просмотр таблицы конфигурации всех портов:

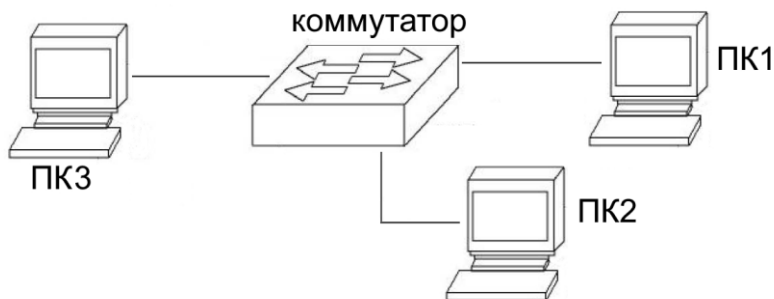
```
show ip mac-bind
```

Просмотр таблицы конфигурации интерфейса:

```
show ip mac-bind IFNAME
```

Пример:

К портам коммутатора подключены ПК1, ПК2 и ПК3. Требуется привязать IP и MAC адреса к портам коммутатора для защиты от ARP атаки.



```
Switch#configure terminal
Switch(config)#interface ge1/1
Switch(config-ge1/1)#ip mac-bind 192.168.1.100 0011.5b34.42ad
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#ip mac-bind 192.168.1.101 0011.6452.135d
Switch(config-ge1/2)#interface ge1/3
Switch(config-ge1/3)#ip mac-bind 192.168.1.102 0011.804d.a246
Switch(config-ge1/3)#end
Switch#show ip mac-bind
[ge1/1] sum: 1
      MAC                IP
      0011.5b34.42ad     192.168.1.100
[ge1/2] sum: 1
      MAC                IP
      0011.6452.135d     192.168.1.101
[ge1/3] sum: 1
      MAC                IP
      0011.804d.a246     192.168.1.102
Switch#show ip mac-bind ge1/1
[ge1/1] sum: 1
      MAC                IP
      0011.5b34.42ad     192.168.1.100
Switch#show running-config
!
spanning-tree mst configuration
!
Interface vlan1
 ip address 10.10.10.1/24
!
interface ge1/1
 ip mac-bind 192.168.1.100 0011.5b34.42ad
!
interface ge1/2
```

```
ip mac-bind 192.168.1.101 0011.6452.135d
!  
interface ge1/3  
ip mac-bind 192.168.1.102 0011.804d.a246  
!  
line vty  
!  
end
```

Внимание: Возможные причины сбоев при выполнении привязки:

Ресурсы системы CFP превышены;

Функция ACL filtering включена на интерфейсе;

Интерфейс имеет уровень layer 3 или включена функция trunk.

6. Port loop detection (Определение петлевых соединений)

Появление петлевых соединений на портах коммутатора приводит к возникновению сетевого шторма из-за направления всех пакетов данных с MAC адреса источника на порт и вызывает сбой передачи данных. Данный раздел включает следующие пункты:

- Protocol principle (Принцип протокола)
- Configuration (Конфигурация)

6.1 Protocol principle (Принцип работы протокола)

Протокол Ethernet loopback detection (ELD) может определять петлевые соединения по взаимодействию пакетов данных и блокировать порт в случае выявления петель. Работа протокола ELD основана на вычислении для выбранного порта.

Если на порту коммутатора активирован протокол ELD, включается таймер, когда время таймера истекает, отправляется пакет определения петель. Если отправленный пакет принимается в течении цикла таймера, то это означает, что порт попал в петлевое соединение, петля блокируется и таблица FDB порта очищается.

Если порт является членом нескольких VLAN, то пакет определения петель автоматически отправляется всем VLAN которым порт принадлежит.

6.1.1 Recovery mode (Режимы восстановления)

При обнаружении петлевого соединения порт блокируется. Для восстановления его работы пользователь может выбрать автоматический или ручной метод, предусмотренный ELD протоколом.

- Автоматический режим: когда порт заблокирован ELD протокол активирует таймер. Когда время таймера истекает, происходит разблокировка, на порту заново запускается таймер определения петель.
- Ручной режим: когда порт заблокирован, ELD протокол не активирует таймер для восстановления работы порта. Пользователь должен ввести команду для разблокировки.

6.1.2 Protocol security (Безопасность протокола)

ELD протокол является уязвимым с точки зрения сетевых атак. Злоумышленник может отправить пакеты (согласно формату ELD протокола) на порт на котором активирован ELD протокол, в результате чего порт будет заблокирован при отсутствии сетевых петель.

Для предотвращения ошибок и атак подобного рода ELD протокол предлагает два возможных решения.

- ELD протокол не зависит от прочего оборудования, таким образом пакет данных может быть зашифрован и отправлен совместно с ключом. Злоумышленник не сможет замаскировать пакет протокола без ключа.
- Предотвращение атак путем захвата отраженных пакетов данных. Для предотвращения атак формат пакетов данных принятых коммутатором за один цикл может быть специально сконфигурирован пользователем.

6.2 Configuration introduction (Конфигурация протокола)

6.2.1 Global configuration (Основные настройки)

Таблица команд основных настроек протокола

Команда	Описание	CLI режим
loop-detection detection-time <1-65535>	Установка времени определения петли (по умолчанию 5 сек). Время должно быть в два раза меньше чем время восстановления работы порта.	Global configuration mode
loop-detection resume-time <10-65535>	Установка времени восстановления работы порта (в автоматическом режиме). Время должно быть в два раза больше чем время определения петли (по умолчанию 600 сек).	Global configuration mode
loop-detection protocol-safety	Включение проверки безопасности протокола (по умолчанию выкл).	Global configuration mode
loop-detection respond-packets	Установка количества пакетов которое должно быть принято за определенный период времени (при включенной проверке безопасности протокола, по умолчанию 10).	Global configuration mode

6.2.2 Port configuration (Настройки портов)

Таблица команд настройки портов

Команда	Описание	CLI режим
Loop-detection enable	Включение ELD протокола на порту.	Interface configuration mode
Loop-detection resume	Возобновление проверки возникновения петель на порту.	Interface configuration mode
loop-detection resume-mode {automation manual}	Выбор режима восстановления работы порта (по умолчанию режим автоматический).	Interface configuration mode

loopback-detection shutdown-mode {no-shutdown shutdown}	Выбор действия при обнаружении петли (вкл/откл порт).	Interface configuration mode
show loop- detection [ifname]	Просмотр конфигурации протокола и интерфейса.	Interface configuration mode

7. VLAN Configure (Конфигурация VLAN)

7.1 VLAN introduction (Введение)

VLAN значительно расширяет возможности масштабирования физической сети. Обычно физическая сеть имеет маленький масштаб, т.е. сеть может объединять ограниченное количество устройств. Использование VLAN разделения сети позволяет увеличить число подключаемых устройств в несколько раз. VLAN имеет такие же функции и атрибуты как и физическая сеть. VLAN дает следующие преимущества:

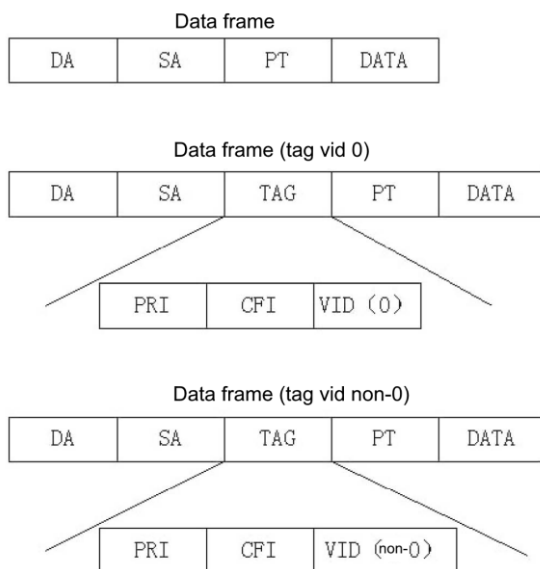
- *VLAN эффективно контролирует сетевой трафик.* В физической сети все пакеты данных передаются на все сетевые устройства вне зависимости от необходимости, что увеличивает нагрузку на сеть и оборудование. VLAN при необходимости объединяет устройства в логическую сеть. VLAN является вещательным доменом, пакеты данных передаются в пределах VLAN. Благодаря делению на VLAN, трафик в сети эффективно распределяется.
- *VLAN повышает безопасность сети.* Устройства подключенные к VLAN могут обмениваться данными между собой в рамках VLAN используя возможности функций уровня layer 2. Если возникает необходимость коммуникации с другим VLAN, обмен данными происходит с использованием функций layer 3. Если между VLAN не установлено соединение с использованием layer 3, то коммуникация не возможна, что повышает безопасность каждого VLAN. Для обмена данными между VLAN должен быть создан канал коммуникации на основе layer 3.

- VLAN дает возможность упрощенного перемещения сетевых устройств. В физических сетях (при перемещении оборудования в другую сеть) требуется изменять сетевые настройки устройств. Т.к. VLAN является логической сетью, которая разделяет устройства не по физическому расположению, при перемещении устройств их принадлежность к VLAN не изменяется. Таким образом, для мобильных устройств отсутствует необходимость изменять сетевые настройки при перемещении.

7.1.1 VLAN ID (Идентификация VLAN)

Каждый VLAN имеет свой идентификационный номер VLAN ID. Диапазон VLAN ID находится в пределах 0 - 4095, где 0 и 4095 не используется (только 1 - 4094). VLAN ID однозначно идентифицирует VLAN. Коммутатор поддерживает 4094 VLAN. По умолчанию коммутатор имеет `vlan1`, который не может быть удален, при создании новых VLAN следует использовать номера от 2 до 4094.

Существует три формата кадров данных (data frames) передаваемых VLAN: data frames без tag, data frames с tags и vid 0 и data frames с tags с vid non-0. См. Форматы на рис. Ниже:



Внутри коммутатора все data frames отмечены. Если коммутатор принимает data frame без тэга, то он сам добавляет tag к data frame, и выбирает VLAN ID для заполнения vid. При получении тегированной data frame с vid 0 коммутатор выбирает VLAN ID для заполнения tagged vid. Если data frame отмечена с vid non-0, то коммутатор оставляет её без изменений.

7.1.2 VLAN Port member type (Типы портов-членов VLAN)

Коммутатор поддерживает VLAN созданный на основе протокола 802.1Q. Существует два типа портов-членов VLAN: не тегированные (untagged) и тегированные (tagged). VLAN может содержать порты-члены обоих типов (untagged и tagged).

VLAN может как не содержать портов-членов, так и содержать один или несколько портов. Порт-член VLAN может быть как не тегированным, так и тегированным.

Порт может быть тегированным или не тегированным портом-членом одного или нескольких VLAN. Если порт является тегированным портом-членом двух и более VLAN, то он называется VLAN relay port. Также порт может являться не тегированным портом-членом одного или нескольких VLAN и одновременно тегированным членом одного или нескольких других VLAN.

7.1.3 Default port VLAN (Первичный port VLAN)

По умолчанию порт относится только к одному VLAN. Первичный VLAN используется для определения пакетов данных без tag с tag и vid 0 на порту. Первичный VLAN называют port vid или PVID. По умолчанию первичный port VLAN имеет номер 1.

7.1.4 VLAN mode of port (Режимы портов VLAN)

Существует три режима работы портов-членов VLAN: access mode (режим доступа), trunk mode (режим trunk) и hybrid mode (гибридный режим). При настройке VLAN, в начале пользователь

должен выбрать режим портов VLAN.

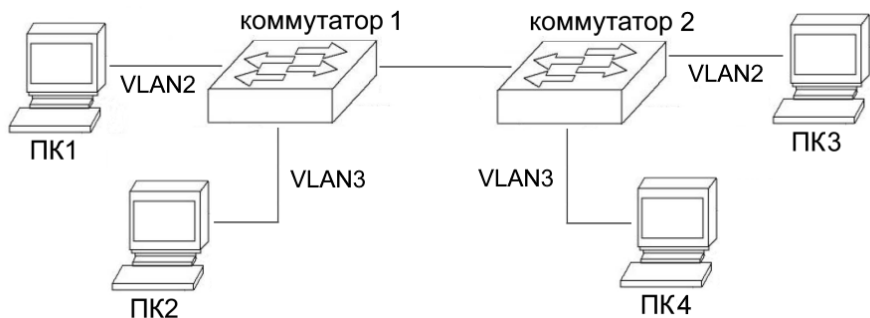
- access mode, порт напрямую доступен пользователю. Данный порт может быть только не тегированным (untagged) членом одного VLAN. Первичный VLAN определяет пользователь. Если порт является не тегированным членом только одного VLAN, то его режим работы определяется как режим доступа (access mode).
- trunk mode (relay port), порт напрямую связан с коммутатором. Данный порт может быть тегированным членом одного или нескольких VLAN, но не может быть не тегированным членом какого-либо VLAN. Порт является членом первичного VLAN 1 и его принадлежность *не может* быть изменена.
- hybrid mode (relay port), порт напрямую связан с коммутатором. Данный порт может быть тегированным членом одного или нескольких VLAN и / или не тегированным членом одного или нескольких VLAN. Порт является членом первичного VLAN 1 и его принадлежность *может* быть изменена.

На практике, пользователь может выбирать режимы портов VLAN в зависимости от текущей необходимости.

7.1.5 VLAN RELAY (Порт VLAN RELAY)

Если порт является тегированным портом-членом двух и более VLAN, то он называется VLAN relay port. Если два коммутатора соединены через VLAN trunk порты, то два и более общих VLAN могут быть разделены между коммутаторами.

На рисунке ниже приведен пример VLAN relay. Два коммутатора соединены через VLAN relay порты VLAN 2 и VLAN 3. Каждый коммутатор разделен на два VLAN (VLAN 2 и VLAN 3) соответственно. Каждый VLAN имеет по одному пользователю. ПК 1 может обмениваться данными с ПК 3, ПК 2 может обмениваться данными с ПК 4, а ПК 1 и ПК 3 не могут обмениваться данными с ПК 2 и ПК 4.



7.1.6 Forwarding of data flow in VLAN (Обмен данными VLAN)

Когда на порт коммутатора приходит пакет данных выполняются следующие шаги (согласно требованиям уровня layer 2):

- ✓ Определение VLAN которому принадлежит пакет;
- ✓ Определение типа пакета (broadcast, multicast, unicast packet);
- ✓ Определение выходного порта в соответствии с различием пакетов данных (0, 1 или более выходных портов). Если выходной порт отсутствует, пакет отбрасывается;
- ✓ Определение метки пакета в соответствии с режимом работы выходного порта (порта-члена VLAN);
- ✓ Отправка пакета через выходной порт.

Определение принадлежности пакета к VLAN:

Если принимаемый пакет тегирован и значение vid в теге не нулевое, то VLAN которому принадлежит пакет, имеет в теге значение vid.

Если принимаемый пакет не тегирован или тегирован, а значение vid в теге нулевое, то VLAN которому принадлежит пакет, является первичным VLAN.

Определение типа пакета данных:

Если MAC адрес назначения получаемого пакета FF: FF: FF: FF: FF: FF, то пакет является broadcast пакетом.

Если получаемый пакет не broadcast типа и 40й бит MAC адреса назначения 1, то пакет является multicast пакетом.

Если получаемый пакет не broadcast и не multicast типа, то пакет является unicast пакетом.

Определение выходного порта пакета:

Если получаемый пакет broadcast типа, все порты-члены VLAN которому принадлежит пакет, являются выходными портами.

Если получаемый пакет multicast типа, сначала происходит обращение к таблице multicast forwarding table (уровень layer 2) согласно MAC адресу назначения и VLAN. Если найден соответствующий multicast вход, общий порт (и действие) выходного порта-члена VLAN, то он является выходным портом пакета. Если общий порт не отсутствует, то пакет сбрасывается. Если не найден соответствующий multicast вход (уровень layer 2) в таблице, то выходные порты определяются согласно режиму forwarding mode уровня layer-2 для multicast forwarding. При нерегистрированном режиме multicast forwarding mode, пакеты multicast определяются как broadcast пакеты, и все порты-члены VLAN являются выходными. При режиме registered forwarding mode, выходной порт отсутствует и пакет сбрасывается.

Если получаемый пакет unicast типа, сначала происходит обращение к таблице forwarding table (уровень layer 2) согласно MAC адресу назначения и VLAN. Если найден соответствующий вход, общий порт (и действие) выходного порта-члена VLAN, то он является выходным портом пакета. Если соответствующий вход не найден в таблице (уровень layer 2) пакеты определяются как broadcast пакеты, и все порты-члены VLAN являются выходными, если выходной порт отсутствует то пакет сбрасывается.

Отправка пакетов:

После определения выходных портов для полученного пакета, происходит его отправка.

Если выходной порт-член VLAN не тегированный, то пакет который принадлежит этому VLAN будет отправлен через выходной порт без тега.

Если выходной порт-член VLAN тегированный, то пакет который принадлежит этому VLAN будет помечен и отправлен через выходной порт. Значение vid в теге соответствует значению VLAN, которому этот пакет принадлежит.

7.2 VLAN configuration (Настройка VLAN)

Данный раздел содержит детальное описание настройки конфигурации VLAN и включает следующие пункты:

- Creating and deleting (Создание и удаление VLAN)
- VLAN mode of VLANs configuration port (Настройка портов VLAN)
- VLAN configuration in access mode (Режим access)
- VLAN configuration in trunk mode (Режим trunk)
- VLAN configuration in hybrid mode (Режим hybrid)
- VLAN subnet configuration (Настройка подсети)
- View VLAN information (Просмотр конфигурации VLAN)

7.2.1 Creating and deleting VLAN (Создание и удаление VLAN)

Для создания и удаления VLAN пользователь должен в режиме global configuration mode войти в режим настройки VLAN configuration. По умолчанию коммутатор имеет VLAN1, который нельзя удалить.

Таблица команд для создания и удаления VLAN:

Команда	Описание	CLI режим
vlan <vlan-id>	Создать VLAN. Команда не будет выполнена, если такой же VLAN уже существует. Диапазон значений параметра от 2 до 4094. VLAN1 существует по умолчанию.	VLAN configuration mode
no vlan <vlan-id>	Удалить VLAN. Команда не будет выполнена, если такой VLAN отсутствует. Диапазон значений параметра от 2 до 4094. VLAN1 не может быть удален.	VLAN configuration mode

7.2.2 Configure VLAN mode of port (Настройка портов VLAN)

Перед тем как начать настройку портов VLAN пользователь должен установить режим работы порта. По умолчанию на портах установлен режим access mode.

Таблица команд настройки режимов портов:

Команда	Описание	CLI режим
switchport mode access	Установка режима access. Порт становится не тегированным членом vlan1 по умолчанию.	Interface configuration mode
switchport mode trunk	Установка режима trunk. Порт становится тегированным членом vlan1 по умолчанию.	Interface configuration mode
no switchport trunk	Отмена режима trunk. Порт переходит в режим access.	Interface configuration mode
switchport mode hybrid	Установка режима hybrid. Порт становится не тегированным членом vlan1 по умолчанию.	Interface configuration mode
no switchport hybrid	Отмена режима hybrid. Порт переходит в режим access.	Interface configuration mode

7.2.3 VLAN configuration access mode (Настройка в режиме access)

Перед тем как начать настройку VLAN configuration, пользователь должен убедиться, что порт находится в режиме access mode. В этом режиме порт является не тегированным членом vlan1.

Таблица команд настройки VLAN в режиме access:

Команда	Описание	CLI режим
switchport access vlan <vlan-id>	Привязка к VLAN не тегированного порта. Диапазон значений параметра от 2 до 4094.	Interface configuration mode

no switchport access vlan	Отмена привязки порта к VLAN. Порт становится не тегированным членом vlan1 по умолчанию.	Interface configuration mode
---------------------------	---	------------------------------

7.2.4 VLAN configuration trunk mode (Настройка в режиме trunk)

Перед тем как начать настройку VLAN configuration, пользователь должен убедиться, что порт находится в режиме trunk mode. В этом режиме порт является тегированным членом vlan1.

Таблица команд настройки VLAN в режиме trunk:

Команда	Описание	CLI режим
switchport trunk native vlan <vlan-id>	Привязка порта PVID к VLAN. Диапазон значений параметра от 2 до 4094.	Interface configuration mode
switchport trunk allowed vlan all	Привязка тегированного порта ко всем существующим VLAN. Тегированный порт также станет привязан к VLAN созданным в дальнейшем.	Interface configuration mode
switchport trunk allowed vlan none	Исключение тегированного порта из всех VLAN кроме vlan1.	Interface configuration mode
switchport trunk allowed vlan add <vlan-list>	Привязка тегированного порта к одному или нескольким VLAN. Параметр < VLAN list > может быть указан в виде <1>, <2-4> или <1,3,5>.	Interface configuration mode
switchport trunk allowed vlan remove <vlan-list>	Исключение тегированного порта из одного или нескольких VLAN. Параметр < VLAN list > может быть указан в виде <1>, <2-4> или <1,3,5>.	Interface configuration mode

7.2.5 VLAN configuration hybrid mode (Настройка в режиме trunk)

Перед тем как начать настройку VLAN configuration, пользователь должен убедиться, что порт находится в режиме hybrid mode. В этом режиме порт является не тегированным членом vlan1.

Таблица команд настройки VLAN в режиме hybrid:

Команда	Описание	CLI режим
switchport hybrid native vlan <vlan-id>	Привязка к VLAN не тегированного порта. Диапазон значений параметра от 2 до 4094.	Interface configuration mode
no switchport hybrid native vlan	Отмена привязки порта к VLAN. Порт становится членом vlan1 по умолчанию.	Interface configuration mode式
switchport hybrid allowed vlan all	Привязка тегированного порта ко всем VLAN (кроме vlan1). Тегированный порт также станет привязан к VLAN созданным в дальнейшем.	Interface configuration mode
switchport hybrid allowed vlan none	Исключение тегированного порта из всех VLAN кроме vlan1. Порт становится членом vlan1 по умолчанию.	Interface configuration mode
switchport hybrid allowed vlan add <vlan-list> egress-tagged enable	Привязка тегированного порта к одному или нескольким VLAN. Параметр < VLAN list > может быть указан в виде <1>, <2-4> или <1,3,5>.	Interface configuration mode
switchport hybrid allowed vlan add <vlan-list> egress-tagged disable	Привязка не тегированного порта к одному или нескольким VLAN. Параметр < VLAN list > может быть указан в виде <1>, <2-4> или <1,3,5>.	Interface configuration mode

switchport hybrid allowed vlan remove <vlan-list>	Исключение порта из одного или нескольких VLAN. Параметр < VLAN list > может быть указан в виде <1>, <2-4> или <1,3,5>. Порт становится членом vlan1 по умолчанию.	Interface configuration mode
--	--	------------------------------

7.2.6 View VLAN information (Просмотр конфигурации)

Таблица команд просмотра конфигурации настройки VLAN:

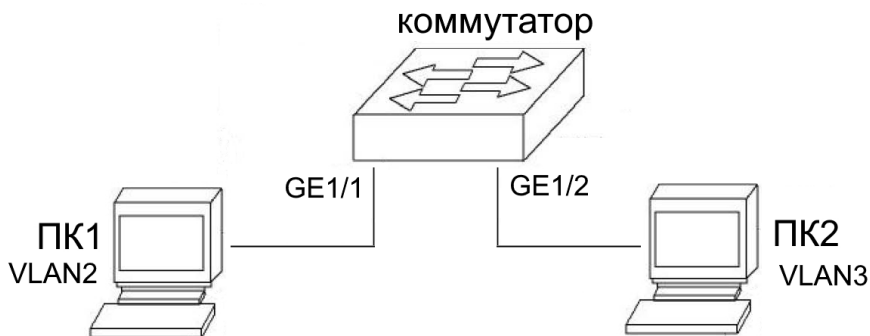
Команда	Описание	CLI режим
show vlan [vlan-id]	Просмотр конфигурации одного из существующих VLAN. Диапазон значений параметра от 1 до 4094. Если параметр отсутствует, то будет показана конфигурация всех VLAN.	Normal mode privileged mode
show interface	Просмотр информации всех портов связанных с VLAN. (VLAN mode, первичный VLAN, и.т.д.).	Normal mode privileged mode
show running-config	Просмотр текущей конфигурации в т.ч. VLAN configuration.	privileged mode

7.3 VLAN configuration example (Пример настройки VLAN)

7.3.1 Port based VLAN (VLAN основанный на портах)

Пример:

Два пользователя ПК1 и ПК2 должны находиться в разных VLAN (с различными задачами и сетевым окружением). ПК1 принадлежит VLAN2 и подключен к порту коммутатора GE1/1. ПК2 принадлежит VLAN3 и подключен к порту коммутатора GE1/2.



Команды настройки конфигурации описанной выше:

Создание VLAN

```

Switch#config t
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#vlan 3
Assign ports to VLANsSwitch#config t
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode access
Switch(config-ge1/1)#switchport access vlan 2
Switch(config-ge1/1)#exit
Switch(config)#interface ge1/2
Switch(config-ge1/2)#switchport mode access
Switch(config-ge1/2)#switchport access vlan 3
  
```

После настройки данной конфигурации ПК из разных VLAN не смогут обмениваться данными. Для установления коммуникации между разными VLAN, используются возможности маршрутизации 3го уровня (3-layer).

Если обмен данными между ПК одного VLAN отсутствует, следует проверить настройки:

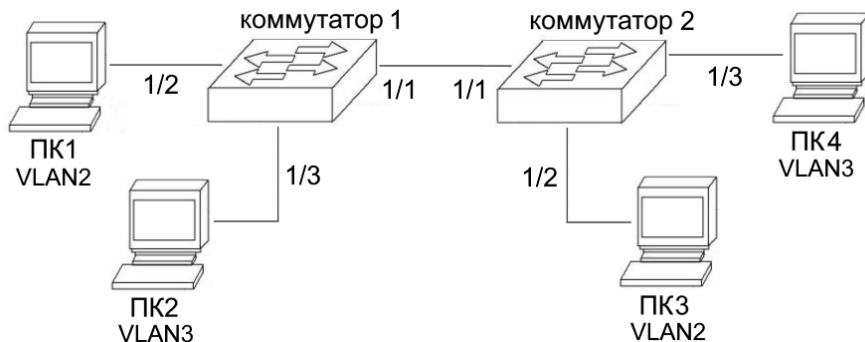
show vlan – просмотр портов-членов всех VLAN.

show vlan <vlan-id> - просмотр портов-членов определенного VLAN.

7.3.2 802.1Q based VLAN (VLAN основанный на протоколе 802.1Q)

Пример:

Пользователь	Принадлежит VLAN	port	Принадлежит коммутатору	cascade port
ПК 1	2	1/2	коммутатор 1	1/1
ПК 2	3	1/3	коммутатор 1	1/1
ПК 3	2	1/2	коммутатор 2	1/1
ПК 4	3	1/3	коммутатор 2	1/1



Команды настройки конфигурации описанной выше (для двух коммутаторов):

Коммутатор 1

```
Switch#config t
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#vlan 3
Switch#config t
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode trunk
Switch(config-ge1/1)#switchport trunk allowed vlan add 2
Switch(config-ge1/1)#switchport trunk allowed vlan add 3
Switch(config)#interface ge1/2
Switch(config-ge1/2)#switchport mode access
Switch(config-ge1/2)#switchport access vlan 2
Switch(config-ge1/2)#exit
Switch(config)#interface ge1/3
```

```
Switch(config-ge1/3)#switchport mode access  
Switch(config-ge1/3)#switchport access vlan 3
```

Коммутатор 1

```
Switch#config t  
Switch(config)#vlan database  
Switch(config-vlan)#vlan 2  
Switch(config-vlan)#vlan 3  
Switch#config t  
Switch(config)#interface ge1/1  
Switch(config-ge1/1)#switchport mode trunk  
Switch(config-ge1/1)#switchport trunk allowed vlan add 2  
Switch(config-ge1/1)#switchport trunk allowed vlan add 3  
Switch(config)#interface ge1/2  
Switch(config-ge1/2)#switchport mode access  
Switch(config-ge1/2)#switchport access vlan 2  
Switch(config-ge1/2)#exit  
Switch(config)#interface ge1/3  
Switch(config-ge1/3)#switchport mode access  
Switch(config-ge1/3)#switchport access vlan 3
```

Если ПК относящиеся к одному VLAN не обмениваются данными, то следует проверить настройку:

- Если ПК соединен с портом принадлежащем соответствующему VLAN, включите режим access mode, для подключения к VLAN.
- Cascade port 1/1 включен в каждый VLAN и должен находиться в режиме trunk mode.

7.4 Mac, IP subnet, protocol VLAN(Протокол VLAN на Mac, IP)

7.4.1 Introduction Mac, IP subnet, protocol VLAN (Введение)

VLAN основанный на MAC адресе разделяется согласно MAC адресу источника сообщения. После получения не тегированного сообщения (untagged или tag 0) от порта, устройство определяет VLAN которому принадлежит сообщение согласно MAC адресу источника и автоматически направляет сообщение определенному VLAN;

VLAN основанный на IP подсети разделяется согласно IP адресу источника и маски подсети (subnet mask). После получения не тегированного сообщения (untagged) от порта, устройство определяет VLAN которому принадлежит сообщение согласно IP адресу источника сообщения, и автоматически направляет сообщение определенному VLAN. Данная возможность используется для передачи сообщений отправленных определенным сегментом сети или IP адреса соответствующему VLAN;

VLAN основанный на протоколе - разным VLAN ID соответствуют пакеты получаемые портом согласно типу протокола. Протоколы используемые для разделения VLANs основаны на IP, IPv6, IPX и т.д.

7.4.2 Mac, IP subnet, protocol VLAN configuration (Конфигурация)

Перед тем как начать конфигурацию VLAN основанных на MAC, IP подсети, протоколе следует создать VLAN.

Команда	Описание	CLI режим
mac-vlan mac WORD vlan <1-4094>	Создание VLAN основанного на MAC адресе источника.	VLAN configuration mode
no mac-vlan mac WORD	Удаление VLAN основанного на MAC адресе источника.	VLAN configuration mode
no mac-vlan	Удаление всех VLAN основанных на MAC адресе источника.	VLAN configuration mode
mac-vlan enable	Включение функции интерфейса mac-vlan.	Interface configuration mode
mac-vlan disable	Отключение функции интерфейса mac-vlan.	Interface configuration mode
show mac-vlan	Просмотр всех VLAN основанных на MAC адресе источника.	Privileged mode

ip-subnet-vlan ip A.B.C.D A.B.C.D vlan <1-4094>	Создание VLAN основанного на IP подсети источника.	VLAN configuration mode
no ip-subnet-vlan ip A.B.C.D A.B.C.D	Удаление VLAN основанного на IP подсети источника.	VLAN configuration mode
no ip-subnet-vlan	Удаление всех VLAN основанных на IP подсети источника.	VLAN configuration mode
ip-subnet-vlan enable	Включение функции интерфейса VLAN основанного на IP подсети источника.	Interface configuration mode
ip-subnet-vlan disable	Отключение функции интерфейса VLAN основанного на IP подсети источника.	Interface configuration mode
show ip-subnet-vlan	Просмотр всех VLAN основанных на IP подсети источника.	Privileged mode
protocol-vlan ether-type (ip ipv6 ipx ... <0-65535>) vlan <1-4094>	Создание VLAN основанного на протоколе.	Interface configuration mode
no protocol-vlan ether-type (ip ipv6 ipx ... <0-65535>)	Удаление VLAN основанного на протоколе.	Interface configuration mode
no protocol-vlan	Удаление всех VLAN основанных на протоколе.	Interface configuration mode
show protocol-vlan	Просмотр всех VLAN основанных на протоколе.	Privileged mode
show vlan-partition interface IFNAME	Просмотр статуса интерфейса VLAN основанных на MAC и IP подсети.	Privileged mode

7.5 Voice VLAN (VLAN для голосового трафика)

7.5.1 Voice VLAN Introduction (Введение)

Voice VLAN это специальный VLAN для выделения голосового трафика IP телефонии. Путем выделения voice VLAN и добавления порта соединяющего голосовое оборудование, параметры QoS (качества обслуживания) могут быть сконфигурированы для повышения приоритета передаваемого голосового трафика и улучшения качества звука.

Оборудование может определять поток голосовых данных на основе MAC адреса источника (oui field) в трафике приходящем на порт. Сообщения поступающие с MAC адреса источника (oui адрес голосового оборудования) определяются системой и направляются в voice VLAN для передачи.

Пользователь может заранее установить oui адрес или использовать oui адрес выделенный по умолчанию как показано ниже:

Номер	Oui адрес	Производитель
1	0001-e300-0000	Siemens phone
2	0003-6b00-0000	Cisco phone
3	0004-0d00-0000	Avaya phone
4	00d0-1e00-0000	Pingtel phone
5	0060-b900-0000	Philips/NEC phone
6	00e0-7500-0000	Polycom phone
7	00e0-bb00-0000	3com phone

Добавьте порт доступа для IP в voice VLAN. Определите соответствующий oui адрес путем идентификации MAC источника сообщений. Система сделает запрос ACL правил и установит приоритет сообщений.

7.5.2 Voice VLAN configuration (Конфигурация Voice VLAN)

Перед тем как начать конфигурацию Voice VLAN следует создать соответствующий VLAN.

Команда	Описание	CLI режим
voice-vlan oui WORD mask WORD	Создание пользовательского oui.	Global configuration mode
voice-vlan oui WORD mask WORD description WORD	Создание пользовательского oui (с именем).	Global configuration mode
no voice-vlan oui WORD mask WORD	Удаление пользовательского oui (по oui адресу и маске).	Global configuration mode
no voice-vlan oui description WORD	Удаление пользовательского oui (по имени).	Global configuration mode
no voice-vlan oui	Удаление всех пользовательских oui.	Global configuration mode
no voice-vlan default-oui WORD mask WORD	Удаление oui по умолчанию (по oui адресу и маске).	Global configuration mode
no voice-vlan default-oui description WORD	Удаление oui по умолчанию (по имени)	Global configuration mode
no voice-vlan default-oui	Удаление всех oui по умолчанию.	Global configuration mode
voice-vlan default-oui resume	Восстановление всех oui по умолчанию.	Global configuration mode
show voice-vlan oui	Просмотр всех существующих oui.	Privileged mode

voice vlan <1-4094> (enable disable)	Включение voice VLAN интерфейса.	Interface configuration mode
voice vlan qos remark cos <0-7> dscp <0-63>	Конфигурация по приоритету QoS. Значение cos 6, значение DSCP 46.	Interface configuration mode
no voice vlan qos	Восстановление QoS приоритета по умолчанию.	Interface configuration mode
no voice vlan	Удаление voice VLAN интерфейса.	Interface configuration mode
show voice-vlan	Просмотр конфигурации voice VLAN интерфейса.	Privileged mode

Пример:

Сконфигурировать поток голосовых данных (0009.ca00.0000) к порту коммутатора GE1/1 и от порта GE1/2 с тегом tag 2.

Команды настройки конфигурации описанной выше:

```
Switch#con t
Switch(config)#vlan da
Switch(config-vlan)#vlan 2
Switch(config-vlan)#exit
Switch(config)#int ge1/1
Switch(config-ge1/1)#sw mod hy
Switch(config-ge1/1)#sw hybrid allowed vlan add 2 egress-tagged disable
Switch(config-ge1/1)#voice vlan 2 en
Switch(config-ge1/1)#voice vlan 2 enable
Switch(config-ge1/1)#int ge1/2
Switch(config-ge1/2)#sw mod tr
Switch(config-ge1/2)#sw trunk allowed vlan add 2
Switch(config-ge1/2)#exit
Switch(config)#voice-vlan oui 0009.ca00.0000 mask ffff.ff00.0000
Switch(config)#
```

7.6 VLAN mapping (VLAN отображение)

7.6.1 VLAN mapping Introduction (Введение)

Функция VLAN mapping (отображения) может изменять тег сообщения (VLAN tag) т.е. изменяет тег отображающий VLAN ID на другой VLAN ID.

7.6.2 VLAN mapping configuration (Конфигурация VLAN mapping)

Перед тем как начать конфигурацию VLAN mapping, следует создать соответствующий VLAN.

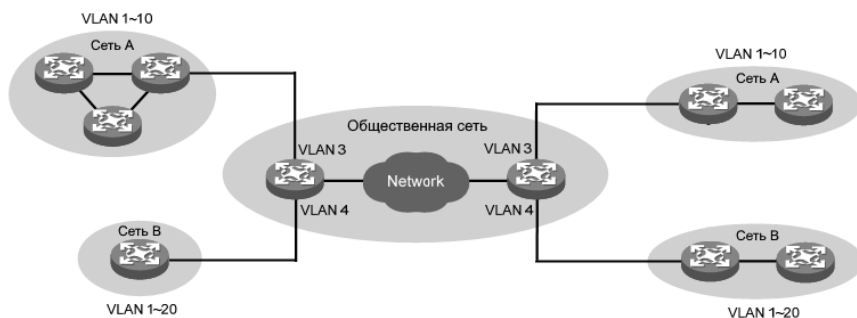
Команда	Описание	CLI режим
vlan-mapping vlan <1-4094> map-vlan <1-4094>	Конфигурация VLAN mapping взаимосвязей портов.	Interface configuration mode
no vlan-mapping vlan <1-4094>	Отмена VLAN mapping взаимосвязей порта.	Interface configuration mode
no vlan-mapping	Отмена VLAN mapping всех взаимосвязей порта.	Interface configuration mode
vlan-mapping enable	Включение VLAN mapping взаимосвязей портов	Interface configuration mode
vlan-mapping disable	Закрытие VLAN mapping взаимосвязей порта.	Interface configuration mode
show vlan-mapping	Просмотр всех VLAN mappings конфигураций.	Privileged mode

7.7 QinQ (QinQ технология двойного тегирования)

7.7.1 QinQ (Введение)

Возможности QinQ обеспечиваются функциями второго уровня (2-layer) технологии VPN. Происходит инкапсуляция исходящего тега VLAN tag для пакетов данных частной сети пользователя, таким образом данные получают VLAN tag от оператора публичной сети (public network). В публичной сети данные пересылаются согласно исходящим VLAN tag, и связаны с элементом таблицы MAC адресов VLAN где хранится исходящий тег, одновременно тег частной сети VLAN tag пересылается как часть данных сообщения.

QinQ дает возможность пользователю использовать один VLAN для обслуживания сети с несколькими VLAN. Как показано на рисунке ниже: частный VLAN сети пользователя А (VLAN 1 - 10), и частный VLAN сети пользователя В (VLAN 1 - 20). VLAN назначенный пользователем для сети А - VLAN 3, VLAN назначенный для сети В - VLAN 4. В случае когда сообщение с VLAN tag пользователя сети попадает в сеть оператора, VLAN tag с VLAN ID 3 будет инкапсулирован вне сообщения или когда сообщение с VLAN tag 4 оператора попадает в сеть с тегом ID VLAN В, сообщения разных пользователей и разных сетей разделяются при передаче через общественную сеть. Даже если диапазоны VLAN двух частных сетей перекрываются при передаче через общественную сеть, конфликта не происходит.



QinQ позволяет создать в сети максимальное количество VLAN 4094х4094, что вполне обеспечивает потребности по количеству VLAN, и решает следующие основные проблемы.

- Помогает снизить дефицит VLAN ID в публичных сетях.
- Пользователи могут создавать собственные (частные) VLAN ID, которые не будут конфликтовать с общественными VLAN ID.
- Предоставляет относительно простые 2-layer VPN решения для сетей на предприятиях и частных пользователей.

Существует разделение на два типа QinQ: основной (basic) QinQ и гибкий (flexible) QinQ.

- Basic QinQ: основной QinQ базируется на режиме работы порта. При включении на порту функции basic QinQ, когда порт получает сообщение, коммутатор отмечает тег VLANa по умолчанию для сообщения. Если полученное сообщение имеет VLAN tag, то оно становится дважды тегированным; если полученное сообщение не имеет VLAN tag, то оно получает тег порта VLAN tag по умолчанию.
- Flexible QinQ: гибкий QinQ базируется на комбинации режимов работы порта и VLAN. В добавлении к возможностям basic QinQ, сообщения приходящие на порт могут подвергаться различным действиям в соответствии с разными VLAN, и получать разные исходящие VLAN tags с внутренними VLAN ID.

7.7.2 QinQ configuration (Конфигурация QinQ)

Таблица команд конфигурации QinQ

Команда	Описание	CLI режим
qinq tpid WORD	Значение в VLAN (0xtpid 8100 по умолчанию)	Interface configuration mode
no qinq tpid	Восстановление TPID порта по умолчанию	Interface configuration mode
qinq uplink	Установка режима порта uplink port	Interface configuration mode

no qinq uplink	Отключение режима порта uplink port	Interface configuration mode
qinq customer	Установка режима порта customer port	Interface configuration mode
no qinq customer	Отключение режима порта customer port	Interface configuration mode
qinq outer-vid <1-4094> inner-vid VLAN_ID	Конфигурация конверсии VLAN интерфейса	Interface configuration mode
no qinq outer-vid <1-4094> [inner-vid VLAN_ID]	Отключение конверсии VLAN интерфейса	Interface configuration mode
show qinq	Просмотр статуса всех QinQ конфигураций	Privileged mode

Пример:

Порт GE1/1 коммутатора В подключен к общественной сети, порт GE1/2 подключен к ПК и телефонному серверу. VLAN 600 используется для ПК, VLAN 500 используется для IP телефонии. Сообщения VLAN 100 и VLAN 200 проходят через общественную сеть. Требуется обеспечить прохождение через общественную сеть данных пользователей ПК через VLAN 200 и данных IP телефонии через VLAN 100 (см. рис. ниже).



Команды настройки конфигурации коммутатора B:

```
Switch#con t
Switch(config)#vlan da
Switch(config-vlan)#vlan 100
Switch(config-vlan)#vlan 200
Switch(config-vlan)#exit
Switch(config)#int ge1/1
Switch(config-ge1/1)#sw mod tr
Switch(config-ge1/1)#switchport trunk allowed vlan add 100
Switch(config-ge1/1)#switchport trunk allowed vlan add 200
Switch(config-ge1/1)#qinq uplink
Switch(config-ge1/1)#int ge1/2
Switch(config-ge1/2)#switchport mode hybrid
Switch(config-ge1/2)#switchport hybrid allowed vlan add 100 egress-
tagged dis
Switch(config-ge1/2)#switchport hybrid allowed vlan add 200 egress-
tagged dis
Switch(config-ge1/2)#qinq customer
Switch(config-ge1/2)#qinq outer-vid 100 inner-vid 500
Switch(config-ge1/2)#qinq outer-vid 200 inner-vid 600
Switch(config-ge1/2)#
Switch#show qinq
```

ifname	tpid	dtag-mode	outer-vid	inner-vid
ge1/1	0x8100	uplink	-	-
ge1/2	0x8100	customer	100	500
ge1/2	0x8100	customer	200	600

```
Switch#
```

8. QoS Configure (Конфигурация QoS)

8.1 QoS introduction (Введение)

Функция QoS позволяет устанавливать приоритет трафика через коммутатор для наиболее важных данных, оптимизировать пропускную способность сети и тем самым делать работу сети более предсказуемой и стабильной.

В коммутаторе очередность отправки выходных пакетов данных определяется в соответствии с информацией о приоритете поступающей на вход коммутатора.

Функция QoS реализуется на основе cos (802.1p), DSCP (DiffServ) и Mac адресах. DSCP QoS может быть сконфигурирован на одном физическом порту, физический порт по умолчанию запускает (802.1p). QoS на основе Mac адреса конфигурируется с помощью функции MAC binding. Каждый физический порт может иметь только одну конфигурацию функции QoS.

Коммутатор поддерживает восемь уровней приоритета очередности от 0 до 7. Уровень 7 имеет наивысший приоритет, уровень 0 имеет низший приоритет. Также предусмотрено четыре режима приоритета очередности: SP, RR, WRR и wdr.

SP – строгий режим. Всегда отправляет пакеты с уровнем приоритета 7 в первую очередь. Пакеты уровня 6 отправляются только после отправки всех пакетов уровня 7. Т.е. пакеты отправляются строго с учетом уровня приоритета от 7 до 0. Последними будут отправлены пакеты с уровнем приоритета 0.

RR – избирательный режим. При отправке пакетов коммутатор выбирает и отправляет пакеты с учетом понижения уровня приоритета, при этом последовательно отправляется по одному пакету каждого уровня.

WRR – взвешенно-избирательный режим. При отправке пакетов коммутатор выбирает и отправляет пакеты с учетом понижения уровня приоритета, учитывая весовой коэффициент пакета. При этом последовательно отправляется по одному пакету каждого уровня с наибольшим весовым коэффициентом.

Wdr – взвешенно-отложено-избирательный режим. В случае если имеется 4 пакета с весовым коэффициентом 3, то за один раз может быть отправлено 5 пакетов, а в следующий раз может быть отправлено 3 пакета. Такие гибкие возможности наиболее подходят для сетей с высокой нагрузкой.

Qosprofile это атрибут для конфигурации взаимосвязи между протоколом 802.1p и приоритетом очередности (priority queue). Этот атрибут не может быть изменен пользователем. Взаимосвязь между протоколом 802.1p и приоритетом очередности представлена в таблице ниже.

QosProfile	802.1p(CoS) value	Priority queue
Qp0	0	0
Qp1	1	1
Qp2	2	2
Qp3	3	3
Qp4	4	4
Qp5	5	5
Qp6	6	6
Qp7	7	7

8.1.1 QoS based on cos (QoS основанный на cos)

QoS основанный на Cos включен на порту по умолчанию. Возможность замены значения приоритета тега VLAN tag пакета данных приходящих на порт определяется очередностью пакета согласно взаимосвязи между пользовательским значением cos и очередностью. Если пакет не имеет VLAN tag или его значение 0, то коммутатор присвоит пакету vid порта (установленный пользователем) и установит приоритет по умолчанию, далее определит приоритет пакета на выходе в соответствии с установками по умолчанию.

8.1.2 QoS based on DSCP (QoS основанный на DSCP)

Если на порту включен QoS основанный на DSCP, замена полученного значения DSCP IP пакета приходящего на порт и определяющего очередность пакета на выходе, происходит в соответствии с пользовательскими установками значения DSCP и установленным приоритетом (queue).

8.1.3 QoS based on Mac (QoS основанный на Mac)

Когда на порт коммутатора поступает пакет данных, коммутатор производит поиск в таблице forwarding table (layer-2) MAC адреса назначения vid и VLAN tag пакета. Если данные найдены, то

очередность пакета на выходе будет определяться в соответствии с взаимосвязями установленными между приоритетом очередности с таблицей назначения.

8.1.4 Policy based QoS (QoS основанный на политике)

Политика QoS включает в себя классы, правила и действия. Класс используется для идентификации потока. Пользователи могут определять правила с помощью команд для классификации пакетов. Действия используются для определения QoS сообщения соответствующего правилу. Если на порту включен QoS основанный на политике, коммутатор классифицирует пакеты данных поступающие на порт. Для пакетов данных соответствующих классификации коммутатор произведет действия в согласно существующей политике. Для прочих пакетов данных коммутатор установит очередность на выходе в соответствии с приоритетом определенным заданной взаимосвязью.

8.2 QoS configuration (Конфигурация QoS)

8.2.1 QoS default configuration (QoS конфигурация по умолчанию)

Таблица настроек QoS по умолчанию

Параметр	Значение	Статус
Number of queues	8	N
Scheduling mode	WRR	Y
Enable SP scheduling mode	disable	Y
Enable RR scheduling mode	disable	Y
Whether to enable wdr scheduling mode	disable	Y
Queue weight	QP0[1],QP1[2],QP2[4],QP3[8] QP4[16],QP5[32],QP6[64] QP7[127]	Y
Mapping relationship between COS and qosprofile	COS0~ [qp0] COS1~ [qp1] COS2~ [qp2]	N

	COS3~ [qp3] COS4~ [qp4] COS5~ [qp5] COS6~ [qp6] COS7~ [qp7]	
Mapping relationship between DSCP and qosprofile	DSCP0~DSCP7[qp0] DSCP8~DSCP15[qp1] DSCP16~DSCP23[qp2] DSCP24~DSCP31[qp3] DSCP32~DSCP39[qp4] DSCP40~DSCP47[qp5] DSCP48~DSCP55[qp6] DSCP56~DSCP63[qp7]	Y
Does the interface enable DSCP based QoS	disable	Y
Does the interface enable cos based QoS	enable	N
Interface user priority (COS value)	0	Y

8.2.2 Configure scheduling mode (Конфигурация режимов)

По умолчанию на коммутаторе включен режим WRR. Режимы SP, RR и wdrp могут быть установлены с помощью команд представленных в таблице ниже:

Команда	Описание	CLI режим
qos sched {rr sp wrr wdrp}	Установка режима QoS	Interface configuration mode

8.2.3 Configure queue weights (Установка весовых коэффициентов)

Весовые коэффициенты определяют количество отправляемых пакетов согласно приоритету очередности в избирательном режиме. При установке весовых коэффициентов следует учитывать что значение при низком приоритете не должно превышать высокоприоритетных значений очередности.

Таблица команд установки весовых коэффициентов

Команда	Описание	CLI режим
qos qosprofile {qp0 qp1 qp2 qp3 qp4 qp5 qp6 qp7} weight<1-127>	Установка коэффициента для каждого приоритета очереди.	Interface configuration mode
no qos qosprofile {qp0 qp1 qp2 qp3 qp4 qp5 qp6 qp7} weight	Отключение коэффициента (по умолчанию).	Interface configuration mode

8.2.4 Configure mapping relationship between DSCP and qosprofile (Конфигурация взаимосвязи DSCP - qosprofile)

Таблица команд конфигурации взаимосвязи между DSCP и qosprofile

Команда	Описание	CLI режим
qos dsc-map-qp <0-63> qosprofile {qp0 qp1 qp2 qp3 qp4 qp5 qp6 qp7}	Установка взаимосвязи между DSCP и qosprofile.	Interface configuration mode
no qos qosprofile {qp0 qp1 qp2 qp3 qp4 qp5 qp6 qp7} weight	Отключение взаимосвязи между DSCP и qosprofile (по умолчанию).	Interface configuration mode

8.2.5 Configure port QoS based on DSCP (Конфигурация порта DSCP)

Функция QoS может быть сконфигурирована только на физическом порту (не trunk или layer 3 интерфейсе).

Таблица команд конфигурации QoS основанной на DSCP.

Команда	Описание	CLI режим
qos dscp-based	Включение функции QoS на порту основанной на DSCP.	Interface configuration mode
no qos dscp-based	Отключение функции QoS на порту основанной на DSCP.	Interface configuration mode

8.2.6 Configure port user priority (Конфигурация пользователя COS)

Таблица команд конфигурации пользователя COS.

Команда	Описание	CLI режим
qos user-priority <0-7>	Установка приоритета пользователя (COS value)	Interface configuration mode
no qos user-priority	Отключение приоритета пользователя (COS value по умолчанию)	Interface configuration mode

8.2.7 QoS configuration example (Пример конфигурации QoS)

Установить на порту GE1/3 пользовательский приоритет (COS value) 3, и QoS на основе cos (по умолчанию):

```
Switch#configure terminal
Switch#(config)#interface ge1/3
Switch#(config-ge1/3)#qos user-priority 3
Switch#(config-ge1/3)#end
```

Установить на порту GE1/3 QoS основанный на DSCP, DSCP значение 3 приоритет очередности 2:

```
Switch#configure terminal
Switch#(config)#interface ge1/3
Switch#(config-ge1/3)#qos dscp-map-qp 3 qosprofile qp2
Switch#(config-ge1/3)#qos dscp-based
Switch#(config-ge1/3)#end
```

8.2.8 Policy QoS configuration example (Пример конфигурации политики QoS)

Сконфигурировать ACL для получения данных от Mac1, Mac2 и Mac3 (ACL правила могут быть изменены).

```
access-list 700 permit host 0000.0000.1111 vid any ip any any
access-list 701 permit host 0000.0000.2222 vid any ip any any
access-list 702 permit host 0000.0000.3333 vid any ip any any
```

Сконфигурировать QoS классы в соответствии потокам данных Mac1, Mac2 и Mac3 (правило cos или DSCP могут быть изменены).

```
qos class 10 match acl 700
qos class 11 match acl 701
qos class 12 match acl 702
```

Сконфигурировать политику QoS изменить отметку приоритет 802.1p потоков данных Mac1, Mac2 и Mac3 (стратегия может быть изменена)

```
qos policy 10 class 10 remark cos 7
qos policy 10 class 11 remark cos 5
qos policy 10 class 12 remark cos 3
Issue QoS policy to port
interface ge1/2
qos apply-policy 10
```

9. MSTP Configure (Конфигурация QoS)

9.1 MSTP introduction (Введение)

Коммутатор поддерживает IEEE802.1d, IEEE802.1w, IEEE802.1s STP стандарты и протоколы.

MSTP протокол использует RSTP для быстрого конвергирования, таким образом множество VLAN агрегируются в разветвленную сеть с независимой топологией. Такая архитектура обеспечивает множественную маршрутизацию потоков данных, распределяет нагрузку, снижает заикливание и поддерживает большое количество VLAN в сети.

9.1.1 Multi spanning tree domain (MST домен)

Для корректной работы протокола multi spanning tree (MST) должны быть последовательно проведены соответствующие MST настройки коммутатора. Настройки коммутатора связанные с MST являются MST доменом.

Конфигурация MST определяет домен к которому относится каждый коммутатор. Конфигурация включает имя домена, номер пересмотра (revision number), MST instance и VLAN assignment mapping. На основе этой информации генерируется уникальный элемент (digest) конфигурации MST, соответствующий элементу summaries домена (эти элементы должны быть одинаковыми). Информация о настройках MST выводится по команде show spanning tree MST config.

Домен может иметь один или несколько членов с одной и той же конфигурацией MST. Каждый член должен обладать возможностью поддерживать RSTP BPDUs. Количество MST доменов в сети не ограничено, но каждый домен поддерживает до 16и instances. Каждому spanning tree instance может соответствовать один VLAN.

9.1.2 IST, CIST, CST

Internal spanning tree (IST) разворачивается в MST домене. В каждом MST домене MSTP обслуживает множество создаваемых instances. Instance 0 – специальный instance домена, который называется ist. Все остальные MST instances имеют номера от 1 до 15.

Для отправления и получения BPDUs используется ist. Вся другая информация о spanning tree instance в сжатом виде содержится в MSTI BPDUs. Т.к. MSTI BPDUs содержит информацию о всех instances, то она должна быть обработана коммутатором поддерживающим multiple spanning tree instances. Номера BPDUs означают упрощение.

Все MST instances в домене используют один таймер протокола, но каждый MST instance имеет свои собственные технологические параметры, такие как ID маршрутного коммутатора, root path cost и т.д. По умолчанию все VLAN относятся к ist.

Common и internal spanning tree (CIST) являются сборником всех ist в каждом MST домене и общим spanning tree (соединяют MST домен и single spanning tree).

Spanning tree вычисляемая внутри домена, является subtree относящейся к CST, которая содержит все домены коммутатора. CIST формируется по результатам вычислений spanning tree протоколами 802.1w и 802.1d. CIST в MST домене такая же как CST вне домена.

Common spanning tree (CST) это spanning tree разворачиваемая между MST доменами.

9.1.3 Intra domain operation (Операции домена)

Ist соединяет все MSTP коммутаторы в домене. Когда ist создает маршрут, то ist становится ist master (мастером), коммутатором с низшим мостом ID и высшим соединением с маршрутом CST в домене. Если в сети только один домен, ist master является CST маршрутом. Если CST маршрут проходит вне домена, MSTP коммутатор находящийся на границе домена выбирается как ist master.

При запуске MSTP коммутатора, он отправляет запрос BPDU на назначение его как маршрут CST и ist master и path cost к маршруту CST, ist master устанавливается на 0. Одновременно коммутатор инициализирует все MST instances и запрашивает возможность быть их маршрутизатором. Если коммутатор получает информацию MST маршрутизации, то он становится главенствующим по данным хранящимся в текущем порту (low bridge ID, low path cost и т.д.), требование стать ist master отменяется.

В процессе инициализации, домен может получить несколько суб-доменов, каждый со своим собственным ist master. Когда коммутатор получает более высокоприоритетное ist сообщение, он покидает старый суб-домен и вступает в новый суб-домен, который может содержать реальный ist master. Таким образом, все суб-домены сокращаются, кроме тех которые содержат реальные ist master.

Для корректного функционирования все коммутаторы MST домена должны опознавать один и тот же ist master. Коммутаторы в любых двух доменах синхронизируют роли порта одного из их MST instances только в том случае если они сходятся в публичном ist master.

9.1.4 Inter domain operation (Операции между доменами)

Если в сети присутствует несколько доменов или коммутаторов ранних моделей с поддержкой протокола 802.1d, MSTP создает и обслуживает CST, которые включают MST домены и коммутаторы

ранних моделей с STP во всех сетях. MST instances включаются в ist на границе домена чтобы стать CST.

Ist соединяет все коммутаторы в MSTP домене и имеет вид subtree в CST (охватывая все коммутаторы и домены). Маршрутизатор subtree становится ist master. MST домен принимает вид виртуального коммутатора, примыкающего к коммутатору STP и MST домену.

Только CST instances отправляют и получают BPDU, MST instances добавляют информацию о своем spanning tree в BPDU, взаимодействуют с соседними коммутаторами и рассчитывают финальную топологию spanning tree. По этому параметры spanning tree относятся к передаче BPDU (hello time, forward time, Max age and Max hops) и конфигурируются только в CST instances, но не оказывают влияния на все MST instances. Параметры относящиеся к топологии spanning tree (switch priority, port VLAN cost, port VLAN priority) могут быть сконфигурированы в CST instances и MST instances.

Коммутаторы MSTP используют версию 3 RSTP BPDU или 802.1d BPDU для коммуникации с коммутаторами поддерживающими протокол 802.1d. Коммутаторы MSTP используют MSTP BPDU для коммуникации с коммутаторами MSTP.

9.1.5 Skip count (Пропуск счета)

В BPDU конфигурирующим расчет топологии spanning tree, ist и MST instances не используют сообщения age и maximum age information. Вместо этого используется path cost маршрута и механизм hop count эквивалентный IP TTL.

Можно установить наибольшее число hops в домене и применить его для ist и всех MST instances в этом домене. Имплементация расчета hop count такая же как для message age (определяется после triggering a reconfiguration). Instance root коммутатор всегда отправляет BPDU (or-m-record) с cost 0 и hop count максимального размера. Когда коммутатор получает BPDU, он вычитает оставшиеся hops по одному и возвращает оставшиеся hops их сгенерировавшему BPDU. Когда счет достигает 0, коммутатор отбрасывает BPDU и обновляет информацию порта.

В домене сообщения message age и maximum age information в части RSTP BPDU остаются такими же, оставшаяся часть такая же, и то же значение передается на специальный порт на границе домена.

9.1.6 Boundary port (Граничный порт)

Граничный порт – это домен spanning tree который соединяет MST домен с отдельным RSTP running spanning tree доменом, или отдельным 801.1d spanning tree доменом, или другим MST доменом с отличающейся конфигурацией. Граничный порт соединен с LAN. Назначенный коммутатор LAN это любой отдельный spanning tree коммутатор или коммутатор с другой конфигурацией MST домена.

В граничном порту роль MST порта не важна, и их статус устанавливается таким же, как и у ist порта (когда ist port отправляет данные, MST порт на границе тоже отправляет данные). Ist порт на границе может иметь роль отличную от backup port.

При граничном соединении (shared boundary), MST порт ожидает истечения времени задержки отправки данных в заблокированном статусе в режиме learning state. MST порт ожидает истечение другой временной задержки для перехода в режим отправки данных.

Если граничный порт подключен в режиме точка-точка и является ist root port (маршрутным), MST порт переключится в режим отправки данных, как только ist port изменит свой режим на отправление.

Если граничный порт переходит в режим отправления в instance, то он отправляет все instances и топология изменяется. Если граничный порт с маршрутом ist root или выделенный порт получит уведомление об изменении топологии, то ist instance и все MST instances на активном порту MSTP коммутатора получают изменение топологии.

9.1.7 Interoperability MSTP and 802.1d STP (Совместимость)

Коммутатор MSTP поддерживает встроенный протокол механизма миграции, который используется координации

взаимодействия с 802.1d. Если коммутатор получает 802.1d конфигурированный BPDU от порта, то он отправляет 802.1d BPDU на этот порт. Когда граничный порт домена получает 802.1d BPDU, MSTP BPDU или RSTP BPDU от другого домена, MSTP коммутатор может принять его.

Если коммутатор перестает получать 802.1d BPDU, то он не возвращается в режим MSTP автоматически, потому что он не может определить находится ли передающий коммутатор в соединении пока не определен его статус (назначение). Когда подключенный (к первому) коммутатор становится членом домена, первый коммутатор может продолжить определять назначение роли порта и начать процесс миграции протокола (force negotiation with the neighbor switch).

Если все коммутаторы поддерживают RSTP, то они все могут поддерживать MSTP BPDU и RSTP BPDU. Таким образом, MSTP коммутаторы могут отправлять любую версию конфигурации (version 0), TCN BPDU или версию 3 MSTP BPDU на граничный порт. Для граничного порта соединенного с LAN, назначенный коммутатор может иметь отдельное spanning tree или это может быть коммутатор с другой MST конфигурацией.

9.1.8 Port role (Роль порта)

MSTP использует быстрый алгоритм конвергации RSTP. Ниже описываются основные моменты взаимодействия MSTP port roles и быстрой конвергенции в сочетании с RSTP.

RSTP обеспечивает быструю конвергенцию ролей порта (port roles) и определения активной топологии. RSTP основывается на протоколе IEEE802 1D STP, выбирает высокий приоритет для коммутатора, как маршрутного. Когда RSTP определяет роль для порта:

Root port (маршрутный порт) – обеспечивает оптимальный путь (path cost) когда коммутатор отправляет данные на маршрутный коммутатор.

Designated port (назначенный порт) – связан и соединен с назначенным коммутатором. Когда LAN отправляет пакеты данных маршрутному коммутатору, создается lowest path cost (путь низшего порядка). Указанный порт (specified port) связан с указанным

коммутатором, который соединяется с LAN.

Alternate port (альтернативный порт) - обеспечивает альтернативный путь к порту маршрутного коммутатора.

Backup port (резервный порт) – обеспечивает запасной путь от выделенного порта к spanning tree. Резервный порт существует только когда два порта соединены вместе (точка-точка), или когда два или более коммутаторов объединены в сегмент LAN (shared).

Disable port (отключенный порт) – роль порта в работе spanning tree отсутствует.

Master port (мастер порт) – находится в маршруте домена или кратчайший путь в общем маршруте, т.е. порт соединяет домен с общим маршрутом.

Роль маршрутного порта или определенного порта участвует в работе активной топологии. Роль замещающего или резервного порта не участвует в работе активной топологии.

В масштабе всей сети со стабильной топологией и фиксированными ролями портов, RSTP обеспечивает немедленный переход маршрутных и определенных портов в статус отправления данных, когда все замещенные и резервные порты находятся в статусе сброса. Статус порта (Port status) определяет состояние отправки или обучения (learning) порта.

Быстрая конвергенция.

RSTP обеспечивает быстрое восстановление работоспособности в следующих случаях: сбой работы коммутатора, сбой работы порта или сбой работы LAN. Обеспечивается быстрое восстановление крайних портов (edge ports), новых маршрутных портов и соединений точка-точка.

Edge Ports (крайние порты) – при данной конфигурации порта, крайний порт немедленно переходит в статус отправки (forwarding state). Порт может быть открыт как связанный порт (boundary port) только в случае если порт соединен с отдельным терминалом или с коммутатором который не участвует в расчете spanning tree.

Root Ports (маршрутные порты) – Если протокол RSTP выбирает новый маршрутный порт, то старый маршрутный порт блокируется а новый немедленно изменяет свой статус на отправление (forwarding state).

Point to point links (соединение точка-точка) – когда порты соединяются как точка-точка локальный порт становится назначенным

(designated port), происходит соединение со связанным с ним портом посредством proposal agreement handshake для определения быстрой конвергенции по петлевой технологии (fast convergence loop free topology).

Изменение топологии. Ниже описываются различия между протоколами RSTP и 802.1d при развертывании технологии spanning tree topology.

Detection (детектирование) – в протоколе 802.1d любые изменения статусов заблокировано (blocking) и отправление (forwarding) вызывают изменение топологии. В протоколе RSTP только изменение статуса от блокировки к передаче вызывает изменение топологии (увеличение количества соединений). Изменение статуса крайнего порта не вызывает изменения топологии. Когда RSTP коммутатор обнаруживает изменение топологии, он передает информацию всем не крайним портам, кроме порта принимающего TC информацию.

В отличие от протокола 802.1d, который использует TCN BPDU, протокол RSTP эту информацию не использует. Для взаимодействия с протоколом 802.1d RSTP коммутатор генерирует и обрабатывает TCP BPDU.

Когда RSTP коммутатор получает сообщение с TCN от 802.1d коммутатора на определенный порт, он откликается на BPDU с 802.1d и устанавливает флаговый TCA бит. Если TC во время отсчета таймера (того же что и таймер смены топологии в 802.1d) активен, соединен с маршрутным портом 802.1d коммутатора и получает конфигурацию BPDU с TCA, TC будет перезапущен (restart \ reset). Это действие требуется для поддержки 802.1d коммутаторов. RSTP BPDU не имеет флагового TCA бита.

Propagation (распространение) – когда RSTP коммутатор получает TC сообщение от других коммутаторов через определенный или маршрутный порт, то он передает всем не крайним портам, определенным портам и маршрутным портам (кроме принимающего порта receiving port). Все такие порты коммутатора запускают TC во время отсчета таймера и передают полученную (learned) информацию.

Protocol migration (миграция протокола) – для обеспечения обратной совместимости с 802.1d коммутаторами, RSTP выборочно отправляет 802.1d BPDU конфигурацию и TCN BPDU основанную на каждом порту. При инициализации запускается таймер задержки (минимум определяется в течение периода когда отправляется RSTP

BPDU), и отправляется RSTP. Во время отсчета таймера, коммутатор обрабатывает все BPDU полученные от порта и игнорирует тип протокола.

После прерывания таймера миграции, если коммутатор получает 802.1d BPDU, он предполагает что он соединен с 802.1d коммутатором и начинает использовать 802.1d протокол BPDU. Если RSTP коммутатор использует 802.1d BPDU на порту и получает RSTP BPDU после прерывания таймера, порт перезапустит таймер и начнет использовать RSTP BPDU.

9.1.9 802.1D spanning tree Introduction (Введение)

Протокол spanning tree основывается на следующих пунктах:

- Уникальный групповой адрес (group address) (01-80-c2-00-00-00) определяет все коммутаторы принадлежащие LAN. Этот адрес может быть опознан всеми коммутаторами.
- Каждый коммутатор имеет уникальную идентификацию.
- Каждый порт каждого коммутатора имеет уникальный идентификатор. Для управления конфигурацией spanning tree необходима координация относительного приоритета каждого коммутатора, приоритета каждого порта каждого коммутатора и учета значимости пути каждого порта.

Маршрутный коммутатор имеет наивысший приоритет. Каждый порт коммутатора имеет значимость пути (root path cost), которая является суммой значимостей путей каждого сегмента сети от коммутатора к маршрутному коммутатору. Порт с минимальной значимостью пути называется маршрутным (root port). Если несколько портов имеют одинаковую значимость пути (root path cost), порт с наивысшим приоритетом становится маршрутным портом.

В каждом LAN существует назначенный коммутатор (designated switch), который принадлежит коммутатору с наименьшей значимостью пути в LAN. Порт соединяющий LAN и назначенный коммутатор называется назначенным портом LAN. Если более двух портов определенного коммутатора соединены с LAN, то порт с наивысшим приоритетом выбирается назначенным.

Необходимые элементы решений spanning tree:

- Определение маршрутного коммутатора. Сначала все коммутаторы считаются маршрутными. Коммутатор отправляет конфигурацию BPDU в LAN, свой root_ ID, Bridge_ ID (одинаковые значения). Когда коммутатор получает конфигурацию BPDU от другого коммутатора и если он находит путь в принятой конфигурации BPDU_ поле ID превышает число маршрутов коммутатора, значение ID параметра, то кадр будет сброшен. В противном случае маршрут коммутатора будет обновлен (ID, root path cost root_ path_ Cost и пр.), коммутатор продолжит отправлять конфигурацию BPDU с обновленными данными.

- Определение маршрутного порта. Порт с минимальным значением root path cost в коммутаторе назначается маршрутным. Если несколько портов имеют одинаковое минимальное значение root path cost, то порт с наивысшим приоритетом назначается маршрутным. Если два и более порта имеют одинаковое минимальное значение и одинаковый приоритет, то порт с меньшим номером назначается маршрутным.

- Назначенный коммутатор LAN. Сначала все коммутаторы считаются назначенными в LAN. Когда коммутатор получает BPDU от другого коммутатора (из того же LAN) с меньшим значением root path cost, коммутатор перестает считать себя назначенным. Если два и более коммутаторов в LAN имеют одинаковое значение root path cost, коммутатор с наивысшим приоритетом становится назначенным.

- Если назначенный коммутатор получает конфигурацию BPDU от другого коммутатора LAN который считает себя назначенным, первый назначенный коммутатор отправляет в ответ конфигурацию BPDU для подтверждения своего назначения.

- Определение определенного порта (specified port). Порт определенного коммутатора соединенный с LAN называется определенным. Если определенный коммутатор имеет два и более портов соединенных с LAN, порт с наименьшей идентификацией назначается определенным. Кроме маршрутного и определенного портов, все остальные порты будут переведены в статус заблокированных.

Т.о. после определения маршрутного коммутатора, маршрутного порта коммутатора, назначенного коммутатора, назначенного порта каждого LAN определяется топология spanning tree.

9.2 MSTP configuration (Конфигурация MSTP)

9.2.1 Default configuration (Конфигурация по умолчанию)

Таблица значений параметров настроек MSTP по умолчанию

Параметр команды	Значение
spanning-tree mst enable(start mstp) включение-отключение MSTP	Off (отключено)
Spanning-tree mst priority(Switch CIST priority) Приоритет коммутатора	32768
spanning-tree mst hello-time(Switch cist hello-time) Время отклика коммутатора cist hello-time	2 сек.
spanning-tree mst forward-time(Switch cist forward-time) Время отправки cist forward-time	15 сек.
spanning-tree mst max-age(Switch cist max-age) Максимальное время cist max-age	20 сек.
spanning-tree mst max-hops(Switch cist max-hops) Максимальное время cist max-hops	20 сек.
instance 1 priority (Priority instance) Приоритет instance 1	32768
spanning-tree mst instance 1 priority(Port instance priority) Приоритет порта instance 1	128
spanning-tree mst instance 1 path-cost(Port instance path-cost) Значение path-cost instance 1 порта	20000000
spanning-tree mst priority (Port cist priority) Приоритет порта	128
spanning-tree mst path-cost (Port cist path-cost) Значение path-cost порта	20000000

9.2.2 General configuration (Команды основной конфигурации)

Включить MSTP (при запуске коммутатора отключен по умолчанию).

Команды запуска MSTP:

```
Switch#configure terminal
Switch(config)#spanning-tree mst enable
```

Команды отключения MSTP:

```
Switch#configure terminal
Switch(config)#no spanning-tree mst
```

Установка параметра Max age.

Max age это конфигурация всех instances. Max age это время (в секундах), в которое коммутатор ожидает получения информации о конфигурации spanning tree, перед началом распознавания. Пределы диапазона от 6 до 40 сек, значение по умолчанию 20 сек.

Команды конфигурации:

```
Switch#configure terminal
Switch(config)#spanning-tree mst max-age <seconds>
```

Установка Max hops

Max hops это количество отсчетов в установленных в домене до сброса BPDU. Пределы изменения параметра от 1 до 40, значение по умолчанию 20.

Команды конфигурации:

```
Switch#configure terminal
Switch(config)#spanning-tree mst max-hops <hop-count>
```

Установка времени отправки (forward time)

Конфигурация forward time для всех instances. Forward time это время (в секундах), которое порт ожидает от момента сброса до получения и от получения до отправления данных. Пределы диапазона от 4 до 30 сек, значение по умолчанию 15 сек. Согласно сгенерированному протоколом числу forward time должно соответствовать следующему условию: $2 * (\text{forward time} - 1) \geq \text{max age}$.

Команды конфигурации:

```
Switch#configure terminal
Switch(config)#spanning-tree mst forward-time <seconds>
```

Установка времени Hello time

Hello time является конфигурацией всех instances. Hello time это временной интервал (в секундах) генерации информации о конфигурации маршрутным коммутатором. Пределы диапазона от 1 до 10 сек, значение по умолчанию 2 сек. Согласно сгенерированному протоколом числу forward time должно соответствовать следующему

условию: $2 * (\text{Hello time} + 1) = \text{max age}$.

Команды конфигурации:

```
Switch#configure terminal
```

```
Switch(config)# spanning-tree mst hello-time <seconds>
```

Установка приоритета моста CIST Bridge

Значение параметра по умолчанию 32768, диапазон значений параметра < 0-61440 >, приоритет CIST является множественным значением 4096.

Команды конфигурации:

```
Switch#configure terminal
```

```
Switch(config)#spanning-tree mst priority <priority>
```

Установка конфигурации совместимой с Cisco

Сетевой коммутатор поддерживает протокол MSTP основанный на 802.1s, длина каждого MSTI сообщения составляет 16 бит. Длина каждого MSTI сообщения BPDU Cisco составляет 26 бит. Для обеспечения совместимости с коммутаторами Cisco, совместимый коммутатор должен быть запущен во время конфигурирования настроек сети.

При запуске конфигурации совместимой с Cisco, при определении домена имя домена, номер версии должны совпадать. По умолчанию функция отключена.

Команды включения функции совместимости Cisco:

```
Switch#configure terminal
```

```
Switch(config)#spanning-tree mst cisco-interoperability enable
```

Команды отключения функции совместимости Cisco:

```
Switch#configure terminal
```

```
Switch(config)#spanning-tree mst cisco-interoperability disable
```

Сброс протокола проверки задачи (Reset protocol check task)

Для обеспечения совместимости протоколов 802.1d и STP, система может автоматически определять протокол другой системы. Протокол исполняемый портом определяется согласно протоколу исполняемому на другой стороне.

В некоторых случаях требуется произвести сброс протокола. Например, после установления соединения система исполняет протокол STP на одном порту, через некоторое время устройство на другой стороне исполняющее протокол STP заменяется хостом. В этот момент требуется изменить конфигурацию порта на fast port, но данный порт

исполняет STP протокол, и выполнение задачи согласования протоколов останавливается. Требуется сбросить задачу согласования протоколов и позволить начать согласование протокола с хостом. Команды сброса задачи определения протокола (Reset protocol reconnaissance task для устройства в целом):

```
Switch#clear spanning-tree detected protocols
```

Команды сброса задачи определения протокола (Reset protocol reconnaissance task для порта):

```
Switch#clear spanning-tree detected protocols interface <if-name>
```

9.2.3 Domain configuration (Конфигурация домена)

Если два и более устройства входят в один домен, то для взаимодействия они должны иметь единый VLAN instance mapping relationship, одной версии и одно имя домена.

Домен включает в себя один или несколько членов с одной конфигурацией MST и каждый член поддерживает возможности RSTP BPDU. Отсутствует ограничение по числу членов в сети, но каждый домен может поддерживать до 16 instances.

Конфигурация instances описана в соответствующем разделе. Ниже представлена конфигурация имени домена и версии конфигурации.

Команды конфигурации имени домена (domain name)

```
Switch#configure terminal:
```

```
Switch(config)#spanning-tree mst configuration
```

```
Switch(config-mst)#region <region-name>
```

Команды конфигурации номера версии (revision number)

```
Switch#configure terminal:
```

```
Switch(config)#spanning-tree mst configuration
```

```
Switch(config-mst)# revision <revision-num>
```

9.2.4 Instance configuration (Конфигурация Instance)

Система поддерживает 16 instances, диапазон номеров (instance ID) от 0 до 15. VLAN может относиться только к одному spanning tree

instance. По умолчанию существует только один instance 0, и все VLAN принадлежат этому instance.

Команды конфигурации instance:

```
Switch#configure terminal
Switch(config)#spanning-tree mst configuration
Switch(config-mst)#instance <instance-id> vlan <vlan-id>
```

Конфигурация приоритета MSTI Bridge

По умолчанию значение параметра 32768, диапазон значений параметра < 0-61440 >, приоритет MSTI является множественным значением 4096.

Команды конфигурации приоритета MSTI Bridge:

```
Switch#configure terminal
Switch(config)#spanning-tree mst configuration
Switch(config-mst)#instance <instance-id> priority <priority>
```

9.2.5 Port configuration (Конфигурация порта)

Ниже описывается информация о конфигурации порта относящаяся к MSTP (представлена простая конфигурация).

Команды конфигурации порта для включения в instance:

```
Switch#configure terminal
Switch(config)#interface <if-name>
Switch(config)#spanning-tree mst instance <instance-id>
```

Конфигурация приоритета порта CIST. Диапазон значений параметра от 0 до 240, по умолчанию значение 128, значение приоритета порта CIST может быть множественным до 16.

Команды конфигурации:

```
Switch#configure terminal
Switch(config)#interface <if-name>
Switch(config)#spanning-tree mst priority <priority>
```

Конфигурация приоритета порта MSTI. Диапазон значений параметра от 0 до 240, по умолчанию значение 128, значение приоритета порта MSTI может быть множественным до 16.

Команды конфигурации:

```
Switch#configure terminal
Switch(config)#interface <if-name>
```

```
Switch(config)#spanning-tree mst instance <instance-id> priority  
<priority>
```

Конфигурация path cost порта CIST

По умолчанию значение параметра 200000000, диапазон значений параметра 1-2000000000.

Таблица зависимости значений path cost mapping от пропускной (bandwidth)

bandwidth (bps)	Path cost
100,000(100K)	200000000
1,000,000(1M)	20000000
10,000,000(10M)	2000000
100,000,000(100M)	200000
1,000,000,000(1G)	20000
10,000,000,000(10G)	2000
100,000,000,000(100G)	200
1,000,000,000,000(1T)	20
>10000000000000	2

Команды конфигурации path-cost:

```
Switch#configure terminal  
Switch(config)#interface <if-name>  
Switch(config)#spanning-tree mst path <path-cost>
```

Конфигурация path cost порта MSTI

По умолчанию значение параметра 200000000, диапазон значений параметра 1-2000000000. Взаимосвязь пропускной способности (bandwidth) и значения представлены в таблице выше.

Команды конфигурации path-cost порта MSTI:

```
Switch#configure terminal  
Switch(config)#interface <if-name>  
Switch(config)#spanning-tree mst instance <instance-id> path-cost <path-cost>
```

Конфигурация номера версии отправляемого пакета протокола

Отправить пакет протокола MSTP с конфигурацией по умолчанию.

Диапазон значений параметра конфигурации 0-3, mapping relationship 0-

stp, 2-rstp, 3-mstp.

Команды конфигурации:

```
Switch#configure terminal
```

```
Switch(config)#interface <if-name>
```

```
Switch(config)# spanning-tree mst force-version <version-id>
```

Конфигурация типа соединения

Если порт соединен с другим портом в режиме точка-точка, локальный порт становится назначенным, RSTP устанавливает взаимосвязь методом быстрой миграции по соглашению (proposal agreement process), соединенный порт становится маршрутным, определяется ацикличная топология. Ниже описывается установление взаимосвязи путем процесса предложение-согласие (proposal agreement).

Когда порт коммутатора получает proposal сообщение и порт выбирается как новый маршрутный, RSTP заставляет остальные порты синхронизироваться с информацией нового маршрутного порта.

Если все остальные порты синхронизированы с первостепенной маршрутной информацией полученной от маршрутного порта, коммутатор считается синхронизированным.

Когда протокол RSTP заставляет синхронизироваться с новой маршрутной информацией, определенный порт в статусе отправки и не сконфигурирован как крайний порт, то он изменит свой статус на заблокированный. В основном, когда RSTP заставляет порт синхронизироваться с новыми маршрутными сообщениями, порт переходит в заблокированный статус.

Когда соглашение отправляется на соответствующий порт коммутатора, требуется подтверждение, что вся информация была отправлена. Когда коммутаторы соединены как точка-точка по соглашению с ролями портов, RSTP незамедлительно переводит порт в статус отправления.

В случае shared соединения, статус порта должен быть определен путем вычислительного процесса протокола 802.1d. По умолчанию установлен тип соединения точка-точка.

Команды установки соединения точка-точка (point-to-point):

```
Switch#configure terminal
```

```
Switch(config)#interface <if-name>
```

```
Switch(config)#spanning-tree mst link-type point-to-point
```

Команды установки соединения shared:

```
Switch#configure terminal
Switch(config)#interface <if-name>
Switch(config-ge1/2)#spanning-tree mst link-type shared
```

9.2.6 PORTFAST Related configuration (Конфигурация PORTFAST)

Port Fast

Функция Port fast моментально изменяет статус access или trunk портов с заблокированного на передающий (forwarding), минуя статусы listening и learning. Функция port fast может быть использована для соединения удаленной рабочей станции с сервером, который позволяет устройствам соединяться с сетью немедленно без ожидания развертывания spanning tree.

Команды установки fast port:

```
Switch#configure terminal
Switch(config)#interface <if-name>
Switch(config)#spanning-tree mst portfast
```

BPDU Filtering

Функция фильтрации BPDU filtering может быть включена на коммутаторе (global) или на каждом порту коммутатора, но некоторые параметры будут отличаться.

В режиме global layer, можно использовать spanning tree MST portfast BPDU команду для включения BPDU фильтрации на порту portfast BPDU filter.

В режиме port layer, можно использовать spanning tree MST portfast BPDU фильтр для открытия BPDU фильтра на каждом порту.

Эта функция закрывает port fast порт от получения или отправления BPDU.

Команды конфигурации BPDU фильтрации (global configuration mode):

```
Switch#configure terminal
Switch(config)# spanning-tree mst portfast bpdu-filter
```

Команды конфигурации BPDU фильтрации (interface configuration mode):

```
Switch#configure terminal
Switch(config)#interface <if-name>
Switch(config)#spanning-tree mst portfast bpdu-filter enable
```

BPDU Guard

Защита BPDU может быть включена на коммутаторе (globally) или на каждом порту коммутатора, но некоторые параметры будут отличаться. В режиме global layer, можно использовать spanning tree MST portfast BPDU для включения функции BPDU guard на порту portfast BPDU. В режиме port layer, можно включить функцию BPDU guard на любом порту.

Когда порт с включенной функцией BPDU guard получает BPDU, spanning tree отключает порт. В действующей конфигурации, port fast порт не получает BPDU. Получение BPDU port fast портом указывает на неправильную конфигурацию, как подключение неавторизованного оборудования, BPDU guard переходит в статус отключения ошибки (error disabled state).

Статус отключения ошибки (Error disabled) означает что когда порт с включенной функцией BPDU guard получает BPDU, система запускает таймер отключения ошибки. Функция Error disable перезапустит порт после истечения установленного времени.

Команды конфигурации (global configuration mode):

```
Switch#configure terminal
```

```
Switch(config)# spanning-tree mst portfast bpdu-guard
```

Команды конфигурации (interface configuration mode):

```
Switch#configure terminal
```

```
Switch(config)#interface <if-name>
```

```
Switch(config)#spanning-tree mst portfast bpdu-guard enable
```

Команды включения error-disable:

```
Start the error disable mechanism
```

```
Switch#configure terminal
```

```
Switch(config)#spanning-tree mst errdisable-timeout enable
```

Команды конфигурации error disable timeout:

```
Switch#configure terminal
```

```
Switch(config)#spanning-tree mst errdisable-timeout interval <seconds>
```

9.2.7 Root Guard Related configuration (Конфигурация Root Guard)

Современные сети уровня layer 2 могут содержать множество соединенных и связанных коммутаторов. В случае такой топологии, spanning tree может проводить собственную реконфигурацию и

выбирать customer коммутатор как маршрутный (root). Для обхода этого следует сконфигурировать порт с помощью root guard SP коммутатора соединенного с коммутатором в customer сети. Если вычисления spanning tree приводят к тому, что порт customer сети выбирается маршрутным, функция root guard переводит порт в статус root inconsistent (blocked) для предотвращения назначения customer коммутатора маршрутным или путем маршрута.

Если коммутатор не принадлежащий SP сети становится маршрутным, порт блокируется (root inconsistent STAT) и spanning tree выбирает новый маршрутный коммутатор. Соответственно коммутатор не становится маршрутным и маршрутный путь будет отсутствовать.

Если коммутатор работает в режиме MST, функция root guard заставляет порт стать определенным (specified port). Если граничный порт (boundary port) заблокирован в ist instance функцией root guard, то этот порт блокируется и во всех MST instances. Граничный порт соединен с LAN. Определенный коммутатор должен поддерживать 802.1d или быть сконфигурированным в разных MST доменах.

Когда порт открыт, функция root guard применяется ко всем VLAN, которым принадлежит порт. VLANы могут быть агрегированы в MST instance.

Команды включения root guard:

```
Switch#configure terminal
Switch(config)#interface <if-name>
Switch(config)#spanning-tree mst guard root
```

9.2.8 MSTP Configuration example (Пример конфигурации MSTP)

Конфигурация:

Три коммутатора соединены в по схеме кольцо. Требуется включить протокол spanning tree на каждом коммутаторе для предотвращения нарушения кольца. Конфигурация каждого коммутатора представлена отдельно.

Конфигурация коммутатора 1:

```
Switch>en
Switch#configure terminal
Switch(config)#spanning mst enable
```

Конфигурация коммутатора 2:

```
Switch>en
```

```
Switch#configure terminal
```

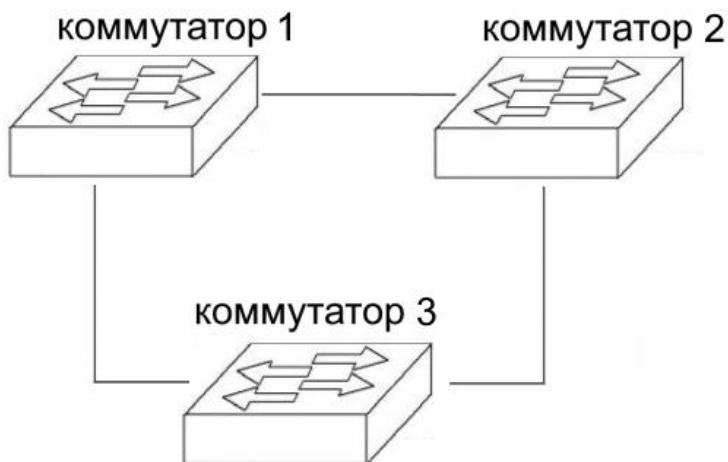
```
Switch(config)#spanning mst enable
```

Конфигурация коммутатора 3:

```
Switch>en
```

```
Switch#configure terminal
```

```
Switch(config)#spanning mst enable
```



Внимание:

Проверьте какой коммутатор выбран как маршрутный (root bridge):

Выполните просмотр spanning tree MST, убедитесь, что значение параметра cistroot соответствует наименьшему MAC адресу из трех, что результат выбора пути корректный.

```
Switch#show spanning-tree mst
```

Для просмотра статуса порта коммутатора в spanning tree выполните команду show spanning tree MST interface GE1/1 и убедитесь что значение параметра порта GE1/1 в instance 0.

```
Switch#show spanning-tree mst interface ge1/1
```

10. EAPS Configure (Конфигурация EAPS)

10.1 EAPS introduction (Введение)

Ethernet automatic protection switching (EAPS) - это протокол автоматической защиты переключений. EAPS использует Ethernet стандарты и технологию VLAN для обеспечения работы петлевой топологии (loop topology) и механизма восстановления петлевых соединений (loop recovery mechanism). В случае разрыва кольцевого соединения, EAPS дает возможность восстановить соединение и передачу данных в течение 1 сек. Количество используемых нодов (nodes) EAPS не ограничено, также не ограничено время восстановления EAPS. EAPS не полагается на работу других устройств, т.е. кольцо EAPS ring может включать в себя устройства, которые не поддерживают протокол EAPS.

10.2 EAPS Basic concepts (Основная концепция EAPS)

1. EAPS домен: В сети домен EAPS организован в отдельное кольцо. Это серии нодов устройств, которые формируют отдельную петлю. EAPS домен включает мастер нод (master node) и один или более транзитных нодов (transit nodes).
2. Мастер нод это коммутатор исполняющий EAPS или EAPS нод устройство. EAPS домен имеет только один мастер нод.
3. Транзитный нод это коммутатор исполняющий EAPS или EAPS нод устройство, не являющийся мастер нодом в EAPS домене.
4. Первичный (Primary) порт это порт, соединяющий EAPS нод устройство в EAPS домене. Нод устройства имеет только один primary порт EAPS домена соединенный с кольцом.
5. Вторичный (Secondary) порт это порт соединяющий EAPS нод устройства в EAPS домене. Нод устройства имеют только один secondary порт в EAPS домене соединенный с кольцом.
6. Управляющий (Control) VLAN, это VLAN отвечающий за протокол передачи пакетов в EAPS домене. EAPS домен может иметь только один управляющий VLAN.

7. Защищенный (Protected) VLAN, это VLAN который передает служебные данные в EAPS домене. EAPS домен должен содержать один или более защищенных VLAN.

10.3 EAPS Protocol introduction (EAPS протокол введение)

EAPS домен содержит мастер нод и один или несколько транзитных нодов. Каждый EAPS нод включает один и тот же управляющий LAN и несколько защищенных VLAN. Каждый EAPS нод содержит первичный (primary) порт и вторичный (secondary) порт в EAPS домене. Эти два порта принадлежат управляющему VLAN и всем защищенным VLAN этого кольца. EAPS кольцо формируется всеми нодами EAPS домена через первичный и вторичный порты каждого нод устройства.

В обычных условиях, когда все первичные и вторичные порты в EAPS домене поддерживают соединение, блокировка вторичного порта мастер нода вызывает разрыв петли передачи служебных данных в EAPS домене. Когда в EAPS домене происходит сбой, открывается вторичный порт мастер нода (меняет статус на отправление forwarding), что позволяет отправлять служебные данные и восстановить нормальную отправку.

Транзитный нод не имеет отличий в процессах от первичного и вторичного портов.

Ниже будет описано два вида проверки сбоев и восстановления петлевых соединений протоколом EAPS.

10.3.1 Link-Down Give an alarm (Link-Down оповещение)

Когда транзитный нод обнаруживает, на что его первичном или вторичном порту происходит обрыв соединения, то он отправляет пакет протокола link-down от управляющего VLAN мастер ноду через другой link up порт.

Когда мастер нод получает пакет link-down протокола, мастер нод незамедлительно переходит в статус failed state, открывает вторичный порт (устанавливает статус вторичного порта forwarding), обновляет свои layer 2 и 3 forwarding, отправляет ring-down-flush-fdb для

уведомления EAPS домена о другом направлении, обновляет таблицу forwarding table и запоминает layer 2 и 3 forwarding.

Когда мастер нод обнаруживает, что на локальном первичном порту обрывается соединение (link down), то он действует как при получении пакета протокола link-down.

Когда мастер нод обнаруживает, что на локальном вторичном порту обрывается соединение (link down), он незамедлительно переходит в статус failed state, обновляет сой layer 2 и layer 3 forwarding, отправляет пакет протокола ring-down-flush-fdb, и уведомляет EAPS домен о других направлениях, обновляет таблицу forwarding table и запоминает layer 2 и layer 3 forwarding.

10.3.2 Loop check (Проверка петли)

Мастер нод регулярно отправляет протокол состояния (health protocol package) с первичного порта. Если кольцо в исправном состоянии, мастер нод может получать health protocol package на своем вторичном порту. В этот момент мастер нод перезапускает таймер (fail period timer), и статус мастера нода остается исправным.

Если до истечения таймера, health protocol package не получен, мастер нод изменит статус на failed state, откроет вторичный порт (изменит статус вторичного порта на forwarding), обновит layer 2 и layer 3 forwarding, отправит уведомление ring-dsown-flush-fdb EAPS домену о других направлениях, обновит таблицу forwarding table и запомнит layer 2 и layer 3 forwarding.

10.3.3 Ring recovery (Восстановление кольца)

Вне зависимости от исправности кольца мастер нод будет отправлять health protocol package со своего первичного порта. Если мастер нод находится в статусе failed state, а health protocol packet получен от вторичного порта, кольцо вернется в статус исправного (complete state). В это время мастер нод изменит статус вторичного порта на заблокирован (blocking), обновит layer 2 и 3 forwarding, отправит ring-up-flush-fdb package для уведомления других устройств о необходимости обновить их layer 2 и 3 forwarding и запомнит layer 2 и 3 forwarding.

Когда порт транзитного ноды восстанавливает соединение (link up) после обрыва, мастер нод определяет что кольцо исправно, вторичный порт мастера ноды может находиться в статусе forwarding state. В этом случае образуется временное кольцо. Когда один порт транзитного ноды имеет статус link up и другой порт меняет статус с link down на link up, транзитный нод переходит в статус предотправления (pre forwarding state). В этом статусе порт не может отправлять служебные данные и разрывать возможные петли. Когда мастер нод восстанавливается и отправляет ring-up-flush-fdb, транзитный нод после получения пакета протокола изменит статус на link-up, порт перейдет в статус forwarding state и восстановит нормальную отправку служебных данных.

Если транзитный нод не сможет получить ring-up-flush-fdb protocol packet, он переведет порт в статус forwarding status после двойного истечения fail time.

10.3.4 EAPS compatible with extreme (EAPS совместимость)

Extreme company является первой компанией чья продукция поддерживает протокол EAPS. Протокол EAPS также поддерживается коммутаторами стандарта rfc3619. Существует некоторая разница между EAPS используемым устройствами extreme company и протоколом распознавания rfc3619. Протокол EAPS который поддерживает коммутатор полностью совместим с устройствами extreme device, совместимость обеспечивается по умолчанию.

10.4 EAPS configuration (EAPS конфигурация)

Коммутатор может поддерживать до 16и EAPS доменов.

Базовая конфигурация протокола EAPS включает следующие основные элементы: control VLAN (управляющий VLAN), node mode (режим ноды), primary port (первичный порт), secondary port (вторичный порт), protected VLAN (защищенный VLAN), hello time и fail time (время приветствия и время сбоя). Hello time и fail time сконфигурированы по умолчанию (Hello time 1 сек и fail time 3 сек).

ВНИМАНИЕ !

1. Первичный порт должен принадлежать управляющему VLAN EAPS домена и быть trunk членом всех защищенных VLAN.
2. Протокол EAPS не может одновременно работать с MSTP. Если запущен MSTP или сконфигурирован MSTP instance, протокол EAPS не может быть запущен.
3. Когда VLAN запускает протокол vltp, то он не может быть сконфигурирован как управляющий VLAN или защищенный VLAN EAPS.
4. Управляющий VLAN EAPS может содержать только первичный и вторичный порты, и может находиться только в режиме trunk mode VLAN.
5. Если VLAN сконфигурирован как управляющий VLAN EAPS домена и домен запущен, то VLAN не может быть удален, и его порты члены не могут быть удалены или изменены. Управляющий VLAN не может быть сконфигурирован с помощью интерфейсов layer 3.
6. Первичный порт и вторичный порт в защищенном VLAN могут находиться только в режиме trunk mode. Остальные порты члены не имеют ограничений.
7. Только по одному порту может быть сконфигурировано как первичный и вторичный порты в одном EAPS домене.
8. Только один VLAN может быть управляющим VLAN или защищенным VLAN в EAPS домене.
9. Управляющий VLAN всеми нодами в EAPS домене должен быть одним и тем же.

10.5 EAPS Brief introduction of command (Команды EAPS)

Для создания EAPS домена, следует убедиться что VLAN и порты соответствуют необходимым требованиям.

Существует ряд рекомендаций для конфигурации EAPS домена. Во-первых создается EAPS домен. Перед запуском EAPS домена, сконфигурируйте другие параметры основываясь на рекомендациях

выше, в противном случае запуск может не осуществиться. Если требуется установить время Hello time большее чем текущее fail time, сначала увеличте fail time, иначе требуемая конфигурация не установится.

Когда запущен EAPS домен, управляющий VLAN, режим, первичный и вторичный порты не могут быть изменены. Защищенный VLAN, fail timer, hello time, совместимость с устройствами extreme могут быть изменены. Первичный и вторичный порты поддерживают LACP транковые группы (trunk groups).

10.5.1 EAPS Configuration command (Команды конфигурации)

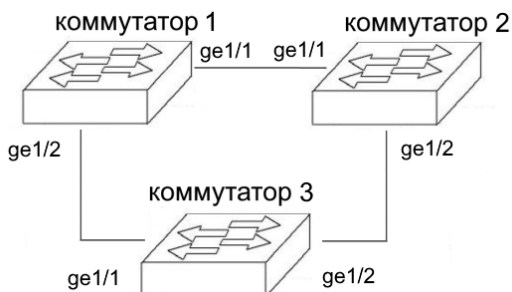
Таблица команд конфигурации EAPS

Команда	Описание	CLI режим
eaps create <ring-id>	Создать EAPS домен.	Global configuration mode
eaps control-vlan <ring-id> <vlan-id>	Создать управляющий VLAN EAPS домена.	Global configuration mode
eaps protected-vlan <ring-id> <vlan-id>	Добавить защищенный VLAN в EAPS домен.	Global configuration mode
eaps mode <ring-id> <master transit>	Выбрать режим нода EAPS домена.	Global configuration mode
eaps primary-port <ring-id> <ifname>	Создать первичный порт EAPS домена.	Global configuration mode
eaps secondary-port <ring-id> <ifname>	Создать вторичный порт EAPS домена.	Global configuration mode
eaps data-span <ring-id>	Создать EAPS кольцо отправления данных.	Global configuration mode
eaps fail-time <ring-id> <secs>	Установка времени таймера fail timer EAPS домена (по умолчанию 3с).	Global configuration mode

eaps hello-time <ring-id> <secs>	Установка времени периодичности отправки EAPS доменом health packets (по умолчанию 1с). Время Hello timer должно быть меньше чем fail time.	Global configuration mode
eaps extreme-interopability <ring-id> <enable disable>	Включение\отключение совместимости с устройствами extreme (по умолчанию вкл).	Global configuration mode
eaps enable <ring-id>	Запуск EAPS домена.	Global configuration mode
eaps disable <ring-id>	Отключение EAPS домена.	Global configuration mode
show eaps	Просмотр конфигурации запущенного EAPS домена.	Normal mode / privileged mode
Show eaps <ring-id>	Просмотр деталей конфигурации запущенного EAPS домена.	Normal mode / privileged mode

10.5.2 Single ring configuration (Пример организации кольца)

Есть три коммутатора: коммутатор1, коммутатор2 и коммутатор3. VLAN 1 не формирует петлю пока трафик направляется в соответствии с EAPS протоколом. Одновременно гарантируется что будет установлен режим ожидания standby link когда одно из соединений между коммутаторами 1-3 разорвется. Согласно рекомендациям, коммутатор1 может быть установлен в мастер режим, а коммутаторы 2,3 в транзитный режим. Включить управляющи VLAN VLAN 2 для передачи пакетов протоколов.



Конфигурация коммутатора1:

Коммутатор1 сконфигурирован как мастер кольца 1 EAPS домена. VLAN2 - управляющий VLAN, VLAN1 – защищенный VLAN, первичный порт GE1/1, вторичный порт GE1/2. Остальные параметры по умолчанию.

```
Switch#configure terminal
#Add VLANs 2 and 3
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-3
Switch(config-vlan)#exit
```

#Configure GE1/1 trunk член VLAN1 и VLAN2.

```
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode trunk
Switch(config-ge1/1)#switchport trunk native vlan 3
Switch(config-ge1/1)#switchport trunk allowed vlan add 1,2
```

#Configure GE1/2 trunk член VLAN1 и VLAN2.

```
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#switchport mode trunk
Switch(config-ge1/2)#switchport trunk native vlan 3
Switch(config-ge1/2)#switchport trunk allowed vlan add 1,2
```

```
Switch(config-ge1/2)#exit
Switch(config)#exit
Switch#show vlan
```

VLAN	Name	State	Member ports ([u]-Untagged, [t]-Tagged)
1	vlan1	active	[t]ge1/1 [t]ge1/2 [u]ge1/3 [u]ge1/4 [u]ge1/5 [u]ge1/6 [u]ge1/7 [u]ge1/8 [u]ge1/9 [u]ge1/10 [u]ge1/11 [u]ge1/12
2	vlan2	active	[t]ge1/1 [t]ge1/2
3	vlan3	active	[u]ge1/1 [u]ge1/2

```
Switch#configure terminal
```

Create an EAPS Domain ring 1

```
Switch(config)#eaps create 1
```

```
#Configure VLAN 2 as the control VLAN
Switch(config)#eaps control-vlan 1 2

#Configure VLAN 1 as a protected VLAN
Switch(config)#eaps protected-vlan 1 1

#Configure Switch1 as the master node
Switch(config)#eaps mode 1 master

#Configure GE1/1 as primary port
Switch(config)#eaps primary-port 1 ge1/1

#Configure GE1/2 as secondary -port
Switch(config)#eaps secondary-port 1 ge1/2

#start-upEAPS Domain ring1
Switch(config)#eaps enable 1
```

Конфигурация коммутатора2:

Коммутатор2 сконфигурирован как transport кольца1 EAPS домена, VLAN 2 – управляющий VLAN, VLAN 1 – защищенный VLAN, первичный порт GE1/1, вторичный порт GE1/2. Остальные параметры по умолчанию.

```
Switch#configure terminal
#Add VLANs 2 and 3
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-3
Switch(config-vlan)#exit
```

```
#Configure GE1/1 trunk член VLAN1 и VLAN 2
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode trunk
Switch(config-ge1/1)#switchport trunk native vlan 3
Switch(config-ge1/1)#switchport trunk allowed vlan add 1,2

#Configure GE1/2 trunk член VLAN1 и VLAN2
Switch(config-ge1/1)#interface ge1/2
```

```
Switch(config-ge1/2)#switchport mode trunk
Switch(config-ge1/2)#switchport trunk native vlan 3
Switch(config-ge1/2)#switchport trunk allowed vlan add 1,2
```

```
Switch(config-ge1/2)#exit
Switch(config)#exit
Switch#show vlan
```

VLAN	Name	State	Member ports ([u]-Untagged, [t]-Tagged)
1	vlan1	active	[t]ge1/1 [t]ge1/2 [u]ge1/3 [u]ge1/4 [u]ge1/5 [u]ge1/6 [u]ge1/7 [u]ge1/8 [u]ge1/9 [u]ge1/10
2	vlan2	active	[t]ge1/1 [t]ge1/2
3	vlan3	active	[u]ge1/1 [u]ge1/2

```
Switch#configure terminal
#Create an EAPs domain ring 1
Switch(config)#eaps create
```

```
#Configure VLAN 2 as the control VLAN
Switch(config)#eaps control-vlan 1 2
```

```
#Configure VLAN 1 as a protected VLAN
Switch(config)#eaps protected-vlan 1
```

```
#Configure switch as a transit node
Switch(config)#eaps mode 1 transit
```

```
#Configure GE1 / 1 as primary port
Switch(config)#eaps primary-port 1 ge1/1
```

```
#Configure GE1 / 2 as secondary port
Switch(config)#eaps secondary-port 1 ge1/2
```

```
#Start EAPs domain ring 1
Switch(config)#eaps enable 1
```

Конфигурация коммутатора3:

Коммутатор3 сконфигурирован как transport кольца1 EAPS домена, VLAN 2 – управляющий VLAN, VLAN 1 – защищенный VLAN, первичный порт GE1/1, вторичный порт GE1/2. Остальные параметры по умолчанию.

```
Switch#configure terminal
# Add VLANs 2 and 3
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-3
Switch(config-vlan)#exit
```

Configure GE1/1 trunk член VLAN1 и VLAN2

```
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode trunk
Switch(config-ge1/1)#switchport trunk native vlan 3
Switch(config-ge1/1)#switchport trunk allowed vlan add 1,2
```

Configure GE1/2 trunk член VLAN 1 и VLAN2

```
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#switchport mode trunk
Switch(config-ge1/2)#switchport trunk native vlan 3
Switch(config-ge1/2)#switchport trunk allowed vlan add 1,2
```

```
Switch(config-ge1/2)#exit
Switch(config)#exit
Switch#show vlan
```

VLAN	Name	State	Member ports ([u]-Untagged, [t]-Tagged)
1	vlan1	active	[t]ge1/1 [t]ge1/2 [u]ge1/3 [u]ge1/4 [u]ge1/5 [u]ge1/6 [u]ge1/7 [u]ge1/8 [u]ge1/9 [u]ge1/10 [u]ge1/11 [u]ge1/12
2	vlan2	active	[t]ge1/1 [t]ge1/2
3	vlan3	active	[u]ge1/1 [u]ge1/2

```
Switch#configure terminal
#Create an EAPs domain ring 1
Switch(config)#eaps create 1
```

```
#Configure VLAN 2 as the control VLAN
Switch(config)#eaps control-vlan 1

#Configure VLAN 1 as a protected VLAN
Switch(config)#eaps protected-vlan 1 1

#Configure switch3 as a transit node
Switch(config)#eaps mode 1 transit

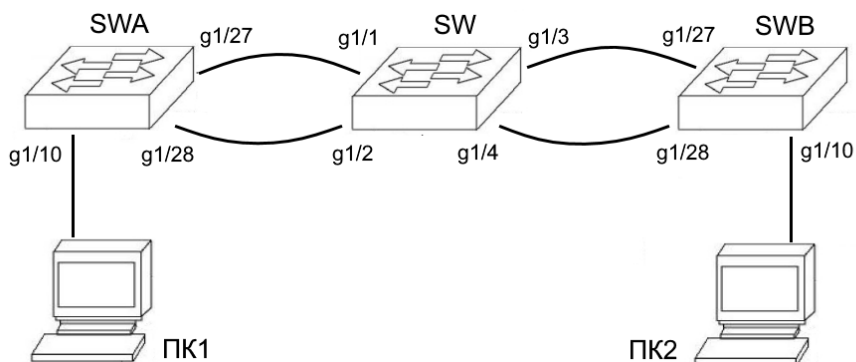
#Configure GE1/1 as primary portSwitch(config)#eaps primary-port 1 ge1/1

#Configure GE1 / 2 secondary -port
Switch(config)#eaps secondary-port 1 ge1/2

#start-upEAPS Domain ring 1
Switch(config)#eaps enable 1
```

10.5.3 Cross ring data forwarding configuration (Пример организации пересекающихся колец)

Есть три коммутатора: SWA, SW и SWB, которые обеспечивают взаимодействие vlan1 и vlan2 по кольцевой топологии по протоколу EAPS согласно схеме ниже.



SWA ring1 управляет VLAN111 и защищает VLAN 1 и 2.

Команды конфигурации представлены ниже:

```
vlan database
```

```
vlan 2
```

```
vlan 111
```

```
interface ge1/10
```

```
switchport access vlan 2
```

```
interface ge1/27
```

```
switchport mode trunk
```

```
switchport trunk allowed vlan add 2
```

```
switchport trunk allowed vlan add 111
```

```
interface ge1/28
```

```
switchport mode trunk
```

```
switchport trunk allowed vlan add 2
```

```
switchport trunk allowed vlan add 111
```

```
eaps create 1
```

```
eaps mode 1 Transit
```

```
eaps primary-port 1 ge1/27
```

```
eaps secondary-port 1 ge1/28
```

```
eaps control-vlan 1 111
```

```
eaps protected-vlan 1 1
```

```
eaps protected-vlan 1 2
```

```
eaps enable 1
```

SW ring 1 соединено с SWA для управления VLAN 111 и защиты VLAN 1 и 2. Ring 2 соединено с SWB, управляет VLAN 222 и защищает VLAN 3333 (необходимо добавить virtual VLAN и interface). Если нужно реализовать cross ring forwarding отправки данных ring 1 и ring 2, необходимо сконфигурировать команду EAPS data span.

Команды конфигурации представлены ниже:

```
vlan database
```

```
vlan 2
```

```
vlan 111
```

```
vlan 222
```

```
vlan 3333
```

```

interface ge1/1
switchport mode trunk
switchport trunk allowed vlan add 2
switchport trunk allowed vlan add 111

interface ge1/2
switchport mode trunk
switchport trunk allowed vlan add 2
switchport trunk allowed vlan add 111

interface ge1/3
switchport mode trunk
switchport trunk allowed vlan add 2
switchport trunk allowed vlan add 222
switchport trunk allowed vlan add 3333    ###Add virtual vlan 3333

interface ge1/4
switchport mode trunk
switchport trunk allowed vlan add 2
switchport trunk allowed vlan add 222
switchport trunk allowed vlan add 3333    ###Add virtual vlan 3333

```

```

eaps create 1
eaps mode 1 Master
eaps primary-port 1 ge1/1
eaps secondary-port 1 ge1/2
eaps control-vlan 1 111
eaps protected-vlan 1 1
eaps protected-vlan 1 2
eaps data-span 1
eaps enable 1

```

```

eaps create 2
eaps mode 2 Transit
eaps primary-port 2 ge1/3
eaps secondary-port 2 ge1/4
eaps control-vlan 2 222
eaps protected-vlan 2 3333    ###Here is virtual protection vlan
eaps data-span 2

```

```
eaps enable 2
```

SWB ring 2 соединено с SW ring 2 для управления VLAN 222 и защиты VLAN 1 и 2.

Команды конфигурации представлены ниже:

```
vlan database
```

```
vlan 2
```

```
vlan 222
```

```
interface ge1/10
```

```
switchport access vlan 2
```

```
interface ge1/27
```

```
switchport mode trunk
```

```
switchport trunk allowed vlan add 2
```

```
switchport trunk allowed vlan add 222
```

```
interface ge1/28
```

```
switchport mode trunk
```

```
switchport trunk allowed vlan add 2
```

```
switchport trunk allowed vlan add 222
```

```
eaps create 2
```

```
eaps mode 2 Master
```

```
eaps primary-port 2 ge1/27
```

```
eaps secondary-port 2 ge1/28
```

```
eaps control-vlan 2 222
```

```
eaps protected-vlan 2 1
```

```
eaps protected-vlan 2 2
```

```
eaps enable 2
```

После того как vlan1 и vlan2 получили конфигурации для коммуникации друг с другом, режим EAPS нода может быть изменен при необходимости.

11. ERPS Configure (Конфигурация ERPS)

11.1 ERPS introduction (Введение)

Ethernet ring protection switching (ERPS) - это протокол защиты сетевых колец g.8032 разработанный ITU. Это протокол уровня соединений (link layer) используемый в кольцевых сетях Ethernet. В кольцевой сети (Ethernet ring network) протокол предотвращает возникновение сетевых штормов вызванных заикливанием данных, и может быстро восстановить коммуникацию между сетевыми нодами при разрыве соединения в сетевом кольце (Ethernet ring). ERPS протокол реализует механизм защиты сетевых колец Ethernet, который быстро восстанавливает передачу данных по сети в случае сбоя кольцевого соединения, что гарантирует высокую стабильность работы коммутаторов при использовании кольцевой топологии сети.

11.1.1 ERPS Loop (ERPS петля)

В принципах работы протокола ERPS заложена минимизация кольца. Каждое кольцо должно быть наименьшим, существует разделение на главное кольцо (main ring) и подкольца (sub ring). Главное кольцо имеет закрытый тип, а подкольца могут быть как закрытого, так и открытого типа. Необходимые конфигурации могут быть установлены с помощью команд.

Каждое ERPS кольцо (первичное или вторичное) имеет пять статусов:

1. idle state (бездействие): все физические связи кольца имеют соединение.
2. Protection status (защита): одна или несколько физических связей кольца теряют соединение.
3. Manual switch status (ручной): ручное изменение статуса кольца.
4. Forced switch state (принудительный): принудительное изменение статуса кольца.
5. Pending status (ожидающий): промежуточный (ожидающий) статус.

11.1.2 ERPS node (ERPS нод)

Оборудование и коммутаторы уровня layer-2 входящие в ERPS кольцо называются нодом. В ERPS кольцо могут входить не более двух портов каждого нода, один порт - RPL port и другой - обычный порт.

Роли нодов делятся на два типа:

- (1) Пересекающиеся (intersecting nodes) – ноды принадлежащие нескольким кольцам в пересекающихся ERP Скольцах;
- (2) Непересекающиеся (Non intersecting nodes) – ноды принадлежащие одному ERPS кольцу в пересекающихся ERPS кольцах.

Режимы нодов определяются ERPS протоколом включающим RPL владельца нода (owner node), RPL соседний нод (neighbor node) и обычный (ordinary ring node).

RPL owner node – назначается пользователем и может быть только один в ERPS кольце. RPL порт заблокирован для предотвращения образования петель в ERPS кольце. Когда нод RPL owner получает уведомление о сбое, что другие ноды или соединения ERPS кольца нарушены, то автоматически разблокируется RPL порт, который возобновит получение и отправку трафика, для обеспечения непрерывности передачи данных.

RPL neighbor node – нод напрямую связан с RPL портом RPL owner node. В обычных условиях RPL порт RPL owner node и RPL порт RPL neighbor node заблокированы для предотвращения образования петель. Когда ERPS кольцо дает сбой, RPL порт RPL owner нода и RPL порт RPL neighbor нода разблокируются.

Ordinary ring node – Все ноды в ERPS кольце кроме RPL owner node и RPL neighbor node являются обычными нодами. RPL порт ordinary ring нода не отличается от обычного ordinary ring порта. Ring порт ordinary ring нода отвечает за мониторинг статуса соединения (link status) связанного с ним напрямую ERPS protocol и уведомляет другие ноды об изменении статуса соединения вовремя.

11.1.3 Links and channels (Соединения и каналы)

RPL (ring protection link) – защитное соединение (соединение RPL порта RPL owner нода), каждое ERPS кольцо имеет только одно RPL. Когда Ethernet кольцо находится в статусе idle state, RPL link

заблокирован и не отправляет сообщения и данные для предотвращения образования петель.

Sub ring link – подкольцевое соединение в пересекающихся кольцах (intersecting ring), принадлежит подкольцу (sub ring) и управляется подкольцом (sub ring).

Raps (ring auto protection switch) – кольцо авто защиты, виртуальный канал. В пересекающемся кольце, путь для передачи сообщений протокола кольца sub ring между пересекающимися нодами, но не принадлежащий sub ring.

11.1.4 ERPS VLAN (ERPS VLAN)

Существует два типа VLAN в ERPS протоколе:

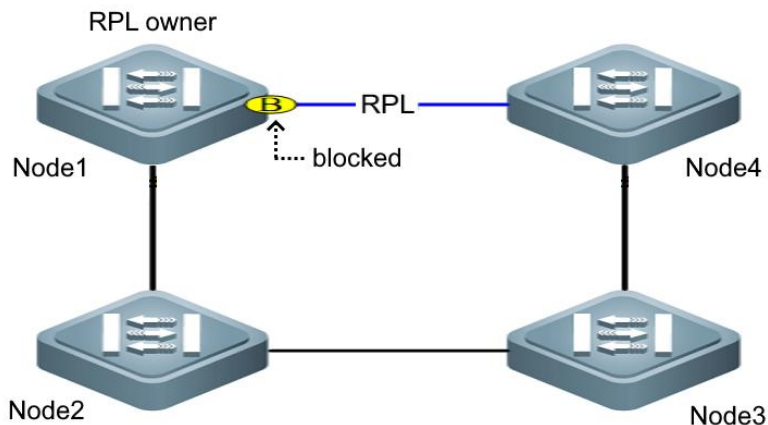
- Raps VLAN – используется для передачи сообщений ERPS протокола. Порты устройств соединенные с ERPS кольцом принадлежат raps VLAN, только порты соединенные с ERPS кольцом могут принадлежать VLAN. Разные кольца должны иметь разные raps VLAN. Не допускается конфигурация IP адресов в интерфейсе raps VLAN.

- Data VLAN – VLAN противоположный raps VLAN, data VLAN используется для передачи сообщений с данными. Data VLAN может включать как ERPS ring порты, так и не ERPS ring порты.

11.2 ERPS working principle (Принципы работы ERPS)

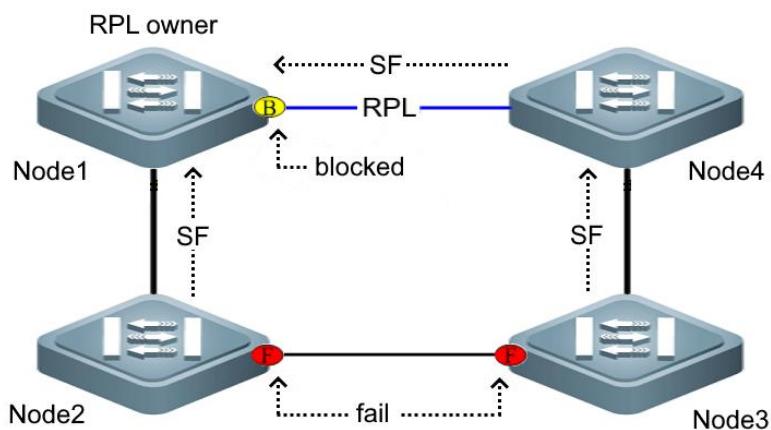
11.2.1 Normal state (Нормальное состояние)

- Все ноды физически соединены по топологии кольцо.
- Протокол защиты (loop protection protocol) обеспечивает защиту от закливания блокировкой соединения RPL link. Как показано на рисунке ниже, соединение между node1 и node4 является RPL link.
- Для каждого соединения между нодами выполняется процедура определения обрыва (Fault detection).

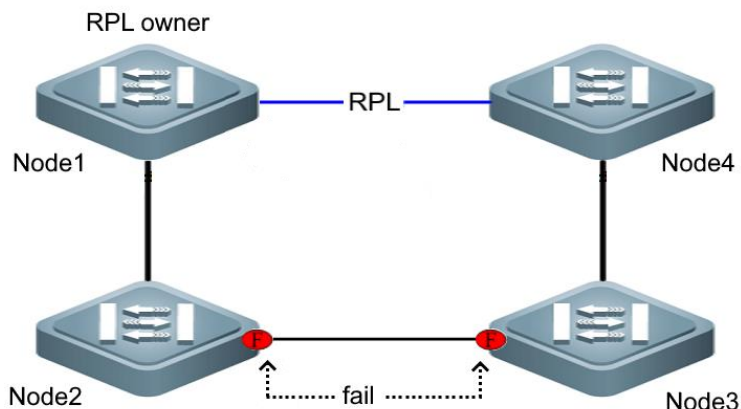


11.2.2 Link failure (Обрыв соединения)

- В случае обрыва соединения, нод блокирует оборванное соединение и использует сообщение gaps (SF) message для оповещения других нодов в кольце. Как показано на рисунке ниже, соединение между node2 и node3 оборвано, node2 и node3 блокируют разорванное соединение после истечения времени holdoff, и отправляют сообщения (SF) messages всем остальным нодам в кольце.

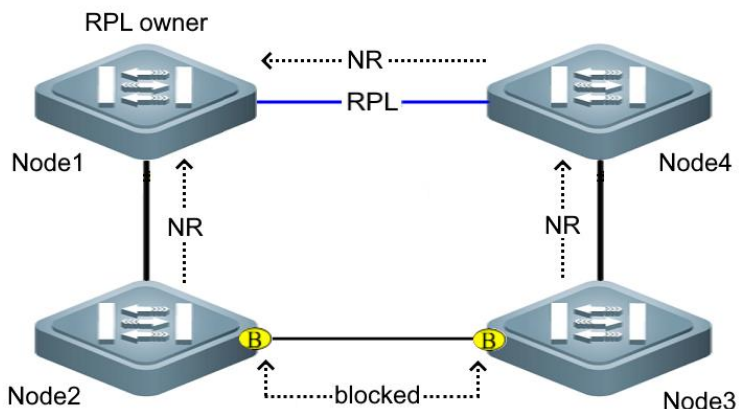


- После получения сообщения (SF) message нод-владелец RPL открывает порт RPL. Одновременно другие ноды обновляют свои таблицы MAC адресов и переходят в защищенный статус (protection state).

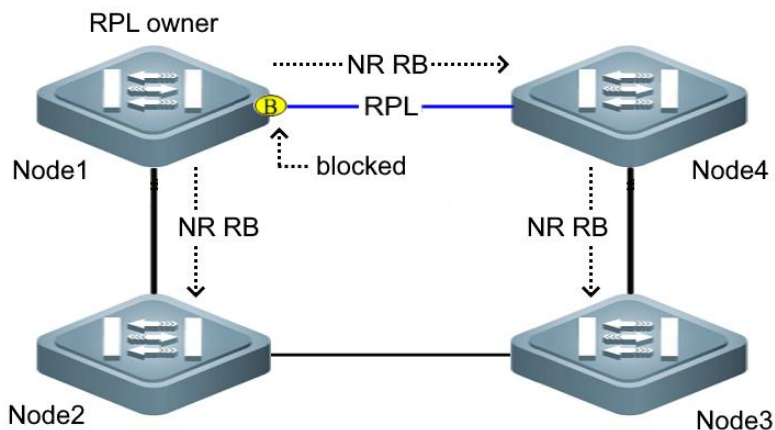


11.2.3 Link recovery (Восстановление соединения)

- При восстановлении соединения нод оставляет данное соединение заблокированным и отправляет сообщение (NR) message об отсутствии обрыва (см. рис. ниже).



- По истечении времени guard timer и получения первого сообщения (NR) message, нод-владелец RPL запускает отсчет времени WTR timer.
- По истечении времени WTR timer, нод-владелец RPL блокирует RPL и отправляет сообщения (NR, RB) messages.
- После получения этих сообщений, другие ноды обновляют свои таблицы MAC адресов. Нод прекращает отправлять сообщения (NR) и открывает заблокированный порт, кольцо возвращается в свое обычное состояние (см. рис. ниже).

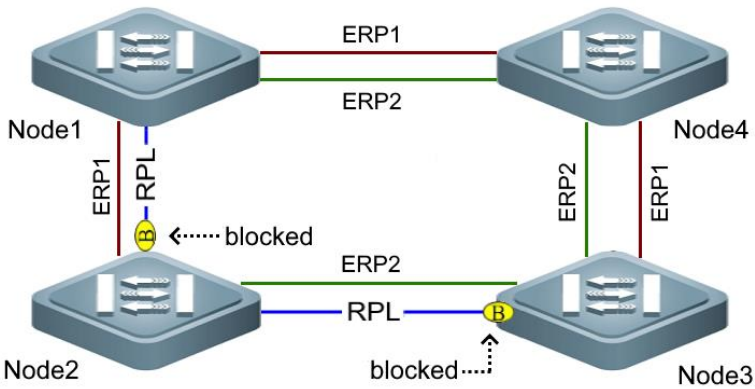


11.3 ERPS Technical features (Возможности ERPS)

11.3.1 ERPS load balancing (Распределение нагрузки)

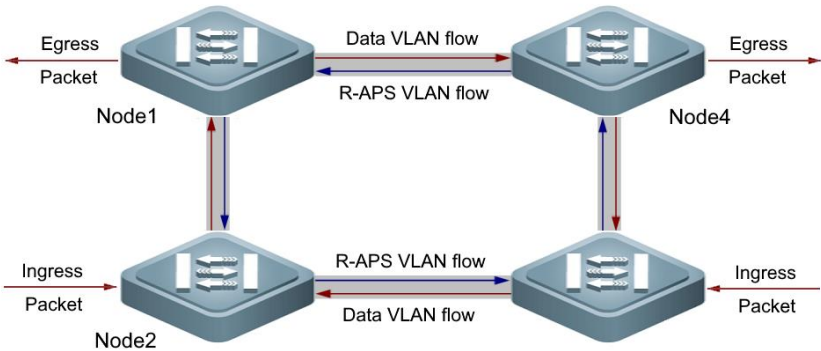
При конфигурации нескольких instances и нескольких ERPS колец в одной физической сети с кольцевой топологией, различные ERPS кольца направляют трафик различных VLANs (защищенных VLANs). При организации различной топологии трафика данных различных VLANs в сети, требуется распределить нагрузку. Как показано на рисунке ниже, физическая сеть с кольцевой топологией имеет поддерживает обмен данными с двумя instances и двумя ERPS кольцами. Разные VLANs защищены двумя ERPS кольцами. Нод2

является владельцем RPL нода erp1 и нод3 является владельцем RPL node ERP2. Согласно конфигурации, разные VLANs могут блокировать разные соединения, для обеспечения распределения нагрузки в одном кольце.



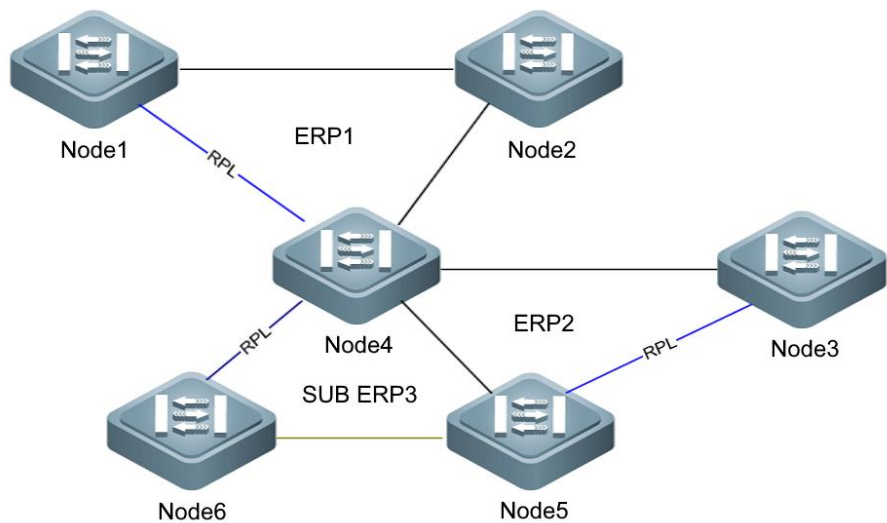
11.3.2 Good safety (Безопасность)

В ERPs существует два типа VLANs in ERPS: raps VLAN и data VLAN. Raps VLAN используется только для передачи протоколов (protocol message) ERPS. ERPS обрабатывает только протоколы полученные от raps VLAN и игнорирует протоколы полученные от data VLAN, таким образом повышается безопасность ERPS.



11.3.3 Multi ring intersection and tangent (Пересечение колец)

ERPs поддерживает возможность добавления множества пересекающихся и касающихся колец в пределах одного нода (node4), как показано на рисунке ниже. Такая возможность на много увеличивает устойчивость сети.



11.4 ERPS Protocol command (Команды ERPS)

Таблица команд ERPS протокола.

Команда	Описание	CLI режим
erps predefine configuration (ring-node rpl-owner-node)	Включение предопределенной конфигурации ERPs	Global configuration mode
no erps predefine configuration	Отключение предопределенной конфигурации ERPs	Global configuration mode
erps <1-8>	Создание ERPs instance	Global configuration mode

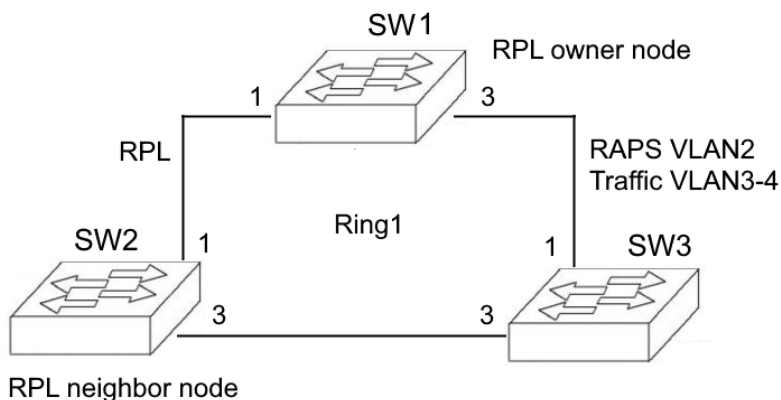
no erps <1-8>	Удаление ERPs instance	Global configuration mode
node-role (interconnection none-interconnection)	Конфигурация роли нода в ERPs кольце, взаимосвязанный или невязанный нод.	ERPS mode
ring <1-32>	Создать ERPs кольцо	ERPS mode
no ring <1-32>	Удалить ERPs кольцо	ERPS mode
ring <1-32> ring-mode (major-ring sub-ring)	Режим ERPs кольца (главное или подкольцо).	ERPS mode
ring <1-32> node-mode (rpl-owner-node rpl-neighbor-node ring-node)	Конфигурация режима нода ERPs, (владелец RPL, соседний RPL, обычный нод).	ERPS mode
ring <1-32> raps-vlan <2-4094>	Конфигурация VLAN протокола ERPs кольца.	ERPS mode
no ring <1-32> raps-vlan	Удаление VLAN протокола ERPs кольца.	ERPS mode
ring <1-32> traffic-vlan <1-4094>	Конфигурация данных VLAN ERPs кольца.	ERPS mode
no ring <1-32> traffic-vlan <1-4094>	Удаление данных VLAN ERPs кольца.	ERPS mode
ring <1-32> (rpl-port rl-port) IFNAME	Конфигурация порта ERPs кольца, (RPL порт или обычный порт).	ERPS mode
no ring <1-32> (rpl-port rl-port)	Удаление порта ERPs кольца.	ERPS mode
ring <1-32> revertive-behaviour (revertive non-revertive)	Конфигурация сценария восстановления ERPs кольца (восстанавливаемое или невосстанавливаемое).	ERPS mode
ring <1-32> hold-off-time <0-10000>	Настройка времени hold off time ERPs кольца.	ERPS mode
no ring <1-32> hold-off-time	Восстановление значения по умолчанию hold off time ERPs кольца.	ERPS mode
ring <1-32> guard-time <10-2000>	Настройка времени guard time ERPs кольца.	ERPS mode
no ring <1-32> guard-time	Восстановление значения по умолчанию guard time ERPs кольца.	ERPS mode

ring <1-32> wtr-time <1-12>	Настройка WTR time ERPs кольца.	ERPS mode
no ring <1-32> wtr-time	Восстановление значения по умолчанию WTR time ERPs кольца.	ERPS mode
ring <1-32> wtb-time <1-10>	Настройка WTB time ERPs кольца.	ERPS mode
no ring <1-32> wtb-time	Восстановление значения по умолчанию WTB time ERPs кольца.	ERPS mode
ring <1-32> raps-send-time <1-10>	Настройка времени отправления protocol message ERPs кольца.	ERPS mode
no ring <1-32> raps-send-time	Восстановление значения по умолчанию времени отправления protocol message ERPs кольца.	ERPS mode
ring <1-32> (enable disable)	Открытие или закрытие ERPs кольца.	ERPS mode
ring <1-32> forced-switch IFNAME	Принудительное включение порта ERPs кольца ring.	ERPS mode
ring <1-32> clear forced-switch	Отключение принудительного включения порта ERPs кольца.	ERPS mode
ring <1-32> manual-switch IFNAME	Ручное включение порта ERPs кольца.	ERPS mode
ring <1-32> clear manual-switch	Отключение ручного включения порта ERPs кольца.	ERPS mode
ring <1-32> clear recovery	Ручное восстановление при сценарии невозстанавливаемого ERPs кольца или до истечения времени WTR / WTB.	ERPS mode
show erps	Просмотр описания всех ERPs instances и колец коммутатора.	Privileged mode
show erps <1-8>	Просмотр одного из ERPs instance или колец коммутатора.	Privileged mode

11.5 ERPS Typical application (Применение ERPS)

11.5.1 Single ring example (Пример одиночного кольца)

Как показано на рисунке ниже, ноды коммутаторов SW1, SW2 и SW3 одиночного ERPs кольца ring1, порты 1 и 3 каждого нода являются портами ERPs кольца. VLAN 2 кольца, данные VLAN 3 и 4 нода коммутатора SW1 являются владельцами RPL owner node, нод коммутатора SW2 является RPL neighbor node, и соединение между коммутаторами SW1 и SW2 является соединением RPL link.



Конфигурация коммутатора sw1:

```
Switch>enable
Switch#configure terminal
Create ERPs protocol and datavlan
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-4
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола protocol и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 1 и одного ERPs кольца ring 1:

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode rpl-owner-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

Конфигурация коммутатора sw2:

```
Switch>enable
Switch#configure terminal
Creating ERPs protocol and data VLAN
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-4
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола protocol и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)#exit
```

```
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 1 и одного ERPs кольца ring 1:

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode rpl-neighbor-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

Конфигурация коммутатора sw3:

```
Switch>enable
Switch#configure terminal
Creating ERPs protocol and data VLAN
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-4
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола protocol и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)#exit
```

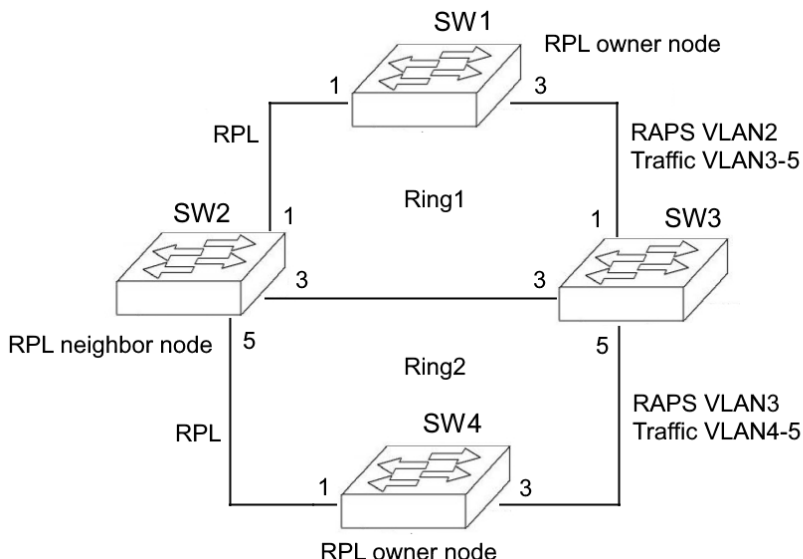
Конфигурация ERPs instance 1 и одного ERPs кольца ring 1:

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode ring-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

11.5.2 Multi ring example (Пример нескольких колец)

Как показано на рисунке ниже, ноды коммутаторов SW1, SW2 и SW3 главного ERPs кольца ring1, порты 1 и 3 нодов коммутаторов SW1, SW2 и SW3 являются портами главного кольца. VLAN 2 главного кольца ring1, данные VLAN 3, 4 и 5 нода коммутатора SW1 являются владельцами RPL owner node, нод коммутатора SW2 является RPL neighbor node, и соединение между коммутаторами SW1 и SW2 является главным соединением кольца ring1 RPL link.

Ноды коммутаторов SW2, SW3 и SW4 ERPs подкольца (sub ring) RING2. Порты 5 нодов коммутаторов SW2 и SW3 и порты 1 и 3 нода коммутатора SW4 являются портами sub ring RING2. VLAN3 подкольца (sub ring) RING2, данные VLAN 4 и 5, нод коммутатора SW4 является подкольцом RING2 RPL owner node, соединение между коммутаторами SW2 и SW4 является соединением RPL link подкольца RING2.



Конфигурация коммутатора sw1:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-5
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 1 и ERPs главного кольца ring 1

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode rpl-owner-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 traffic-vlan 5
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

Конфигурация коммутатора sw2:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-5
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
```

```

Switch(config-ge1/3)# switchport trunk allowed vlan add 2
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
Switch(config-ge1/3)#exit
Switch(config)# interface ge1/5
Switch(config-ge1/5)# switchport mode trunk
Switch(config-ge1/5)# switchport trunk allowed vlan add 3
Switch(config-ge1/5)# switchport trunk allowed vlan add 4
Switch(config-ge1/5)# switchport trunk allowed vlan add 5
Switch(config-ge1/5)#exit

```

Конфигурация ERPs instance 1 и ERPs главного кольца ring 1 и ERPs подкольца sub ring 2:

```

Switch(config)#erps 1
Switch(config-erps-1)# node-role interconnection
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode rpl-neighbor-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 traffic-vlan 5
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#ring 2
Switch(config-erps-1)# ring 2 ring-mode sub-ring
Switch(config-erps-1)# ring 2 node-mode ring-node
Switch(config-erps-1)# ring 2 raps-vlan 3
Switch(config-erps-1)# ring 2 traffic-vlan 4
Switch(config-erps-1)# ring 2 traffic-vlan 5
Switch(config-erps-1)# ring 2 rpl-port ge1/5
Switch(config-erps-1)# ring 2 enable
Switch(config-erps-1)#exit

```

Конфигурация коммутатора sw3:

```
Switch>enable
```

```
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
```

```
Switch(config-vlan)#vlan 2-5
```

```
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
```

```
Switch(config-ge1/1)# switchport mode trunk
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
```

```
Switch(config-ge1/1)#exit
```

```
Switch(config)# interface ge1/3
```

```
Switch(config-ge1/3)# switchport mode trunk
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
```

```
Switch(config-ge1/3)#exit
```

```
Switch(config)# interface ge1/5
```

```
Switch(config-ge1/5)# switchport mode trunk
```

```
Switch(config-ge1/5)# switchport trunk allowed vlan add 3
```

```
Switch(config-ge1/5)# switchport trunk allowed vlan add 4
```

```
Switch(config-ge1/5)# switchport trunk allowed vlan add 5
```

```
Switch(config-ge1/5)#exit
```

Конфигурация ERPs instance 1 и ERPs главного кольца ring 1 и ERPs подкольца sub ring 2:

```
Switch(config)#erps 1
```

```
Switch(config-erps-1)# node-role interconnection
```

```
Switch(config-erps-1)#ring 1
```

```
Switch(config-erps-1)# ring 1 ring-mode major-ring
```

```
Switch(config-erps-1)# ring 1 node-mode ring-node
```

```

Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 traffic-vlan 5
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#ring 2
Switch(config-erps-1)# ring 2 ring-mode sub-ring
Switch(config-erps-1)# ring 2 node-mode ring-node
Switch(config-erps-1)# ring 2 raps-vlan 3
Switch(config-erps-1)# ring 2 traffic-vlan 4
Switch(config-erps-1)# ring 2 traffic-vlan 5
Switch(config-erps-1)# ring 2 rpl-port ge1/5
Switch(config-erps-1)# ring 2 enable
Switch(config-erps-1)#exit

```

Конфигурация коммутатора sw4:

```

Switch>enable
Switch#configure terminal

```

Создание ERPs протокола и данных VLAN:

```

Switch(config)#vlan database
Switch(config-vlan)#vlan 3-5
Switch(config-vlan)#exit

```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```

Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4

```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
Switch(config-ge1/3)#exit
```

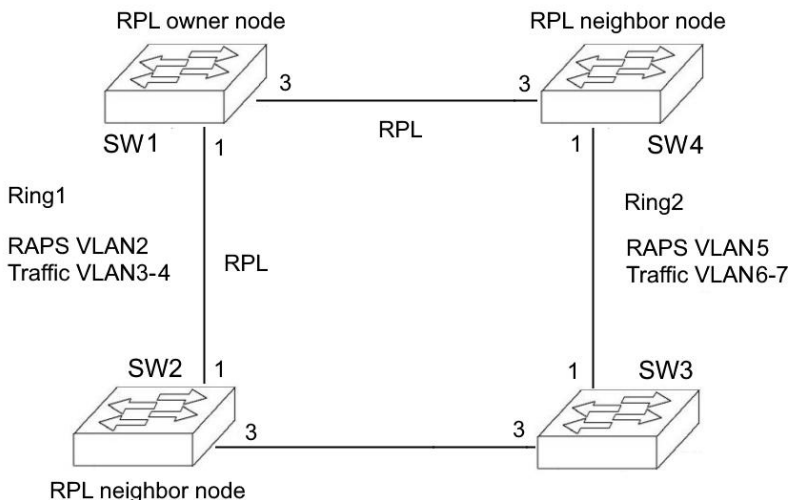
Конфигурация ERPs instance 1 и ERPs подкольца sub ring 2:

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 2
Switch(config-erps-1)# ring 2 ring-mode sub-ring
Switch(config-erps-1)# ring 2 node-mode rpl-owner-node
Switch(config-erps-1)# ring 2 raps-vlan 3
Switch(config-erps-1)# ring 2 traffic-vlan 1
Switch(config-erps-1)# ring 2 traffic-vlan 4
Switch(config-erps-1)# ring 2 traffic-vlan 5
Switch(config-erps-1)# ring 2 rpl-port ge1/1
Switch(config-erps-1)# ring 2 rl-port ge1/3
Switch(config-erps-1)# ring 2 enable
Switch(config-erps-1)#exit
```

11.5.3 Multi instance load balancing example (Пример распределения нагрузки)

Как показано на рисунке ниже, ноды коммутаторов SW1, SW2, SW3 и SW4 образуют одиночное кольцо ring1 и ERPs instance 1. Порты 1 и 3 каждого нода являются портами ERPs кольца. Протокол VLAN2 кольца, данные VLAN 3 и 4, нод коммутатора SW1 является владельцем RPL owner node, нод коммутатора SW2 является соседним RPL neighbor node, RPL link - соединение между коммутаторами SW1 и SW2.

Ноды коммутаторов SW1, SW2, SW3 и SW4 образуют одиночное кольцо RING2 и ERPs instance 2. Порты 1 и 3 каждого нода являются портами ERPs кольца. Протокол VLAN5 кольца, данные VLAN 6 и 7, нод коммутатора SW1 является владельцем RPL owner node, нод коммутатора SW4 является соседним RPL neighbor node, RPL link - соединение между коммутаторами SW1 и SW4.



Конфигурация instance 1:

Конфигурация коммутатора sw1:

```
Switch>enable
```

```
Switch#configure terminal
```

Конфигурация instance для создания ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
```

```
Switch(config-vlan)#vlan 2-4
```

```
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
```

```
Switch(config-ge1/1)# switchport mode trunk
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
```

```
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
```

```
Switch(config-ge1/1)#exit
```

```
Switch(config)# interface ge1/3
```

```
Switch(config-ge1/3)# switchport mode trunk
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
```

```
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
```

```
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 1 и ERPs одиночного кольца ring 1:

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode rpl-owner-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

Конфигурация коммутатора sw2:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-4
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 1 и ERPs одиночного кольца ring 1:

```
Switch(config)#erps 1
```

```
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode rpl-neighbor-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

Конфигурация коммутатора sw3:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-4
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 1 и ERPs одиночного кольца ring 1:

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
```

```
Switch(config-erps-1)# ring 1 node-mode ring-node
Switch(config-erps-1)# ring 1 raps-vlan 2
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

Конфигурация коммутатора sw4:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 2-4
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 2
Switch(config-ge1/1)# switchport trunk allowed vlan add 3
Switch(config-ge1/1)# switchport trunk allowed vlan add 4
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 2
Switch(config-ge1/3)# switchport trunk allowed vlan add 3
Switch(config-ge1/3)# switchport trunk allowed vlan add 4
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 1 и ERPs одиночного кольца ring 1:

```
Switch(config)#erps 1
Switch(config-erps-1)#ring 1
Switch(config-erps-1)# ring 1 ring-mode major-ring
Switch(config-erps-1)# ring 1 node-mode ring-node
Switch(config-erps-1)# ring 1 raps-vlan 2
```

```
Switch(config-erps-1)# ring 1 traffic-vlan 1
Switch(config-erps-1)# ring 1 traffic-vlan 3
Switch(config-erps-1)# ring 1 traffic-vlan 4
Switch(config-erps-1)# ring 1 rpl-port ge1/1
Switch(config-erps-1)# ring 1 rl-port ge1/3
Switch(config-erps-1)# ring 1 enable
Switch(config-erps-1)#exit
```

Конфигурация instance 2:

Конфигурация коммутатора sw1:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 5-7
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
Switch(config-ge1/1)# switchport trunk allowed vlan add 6
Switch(config-ge1/1)# switchport trunk allowed vlan add 7
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
Switch(config-ge1/3)# switchport trunk allowed vlan add 6
Switch(config-ge1/3)# switchport trunk allowed vlan add 7
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 2 и ERPs одиночного кольца ring 2:

```
Switch(config)#erps 2
Switch(config-erps-2)#ring 2
Switch(config-erps-2)# ring 2 ring-mode major-ring
Switch(config-erps-2)# ring 2 node-mode rpl-owner-node
Switch(config-erps-2)# ring 2 raps-vlan 5
Switch(config-erps-2)# ring 2 traffic-vlan 1
```

```
Switch(config-erps-2)# ring 2 traffic-vlan 6
Switch(config-erps-2)# ring 2 traffic-vlan 7
Switch(config-erps-2)# ring 2 rpl-port ge1/3
Switch(config-erps-2)# ring 2 rl-port ge1/1
Switch(config-erps-2)# ring 2 enable
Switch(config-erps-2)#exit
```

Конфигурация коммутатора sw2:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 5-7
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
Switch(config-ge1/1)# switchport trunk allowed vlan add 6
Switch(config-ge1/1)# switchport trunk allowed vlan add 7
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
Switch(config-ge1/3)# switchport trunk allowed vlan add 6
Switch(config-ge1/3)# switchport trunk allowed vlan add 7
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 2 и ERPs одиночного кольца ring 2:

```
Switch(config)#erps 2
Switch(config-erps-2)#ring 2
Switch(config-erps-2)# ring 2 ring-mode major-ring
Switch(config-erps-2)# ring 2 node-mode ring-node
Switch(config-erps-2)# ring 2 raps-vlan 5
Switch(config-erps-2)# ring 2 traffic-vlan 1
Switch(config-erps-2)# ring 2 traffic-vlan 6
Switch(config-erps-2)# ring 2 traffic-vlan 7
```

```
Switch(config-erps-2)# ring 2 rpl-port ge1/1
Switch(config-erps-2)# ring 2 rl-port ge1/3
Switch(config-erps-2)# ring 2 enable
Switch(config-erps-2)#exit
```

Конфигурация коммутатора sw3:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 5-7
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
Switch(config-ge1/1)# switchport trunk allowed vlan add 6
Switch(config-ge1/1)# switchport trunk allowed vlan add 7
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
Switch(config-ge1/3)# switchport trunk allowed vlan add 6
Switch(config-ge1/3)# switchport trunk allowed vlan add 7
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 2 и ERPs одиночного кольца ring 2:

```
Switch(config)#erps 2
Switch(config-erps-2)#ring 2
Switch(config-erps-2)# ring 2 ring-mode major-ring
Switch(config-erps-2)# ring 2 node-mode ring-node
Switch(config-erps-2)# ring 2 raps-vlan 5
Switch(config-erps-2)# ring 2 traffic-vlan 1
Switch(config-erps-2)# ring 2 traffic-vlan 6
Switch(config-erps-2)# ring 2 traffic-vlan 7
Switch(config-erps-2)# ring 2 rpl-port ge1/1
Switch(config-erps-2)# ring 2 rl-port ge1/3
```

```
Switch(config-erps-2)# ring 2 enable
Switch(config-erps-2)#exit
```

Конфигурация коммутатора sw4:

```
Switch>enable
Switch#configure terminal
```

Создание ERPs протокола и данных VLAN:

```
Switch(config)#vlan database
Switch(config-vlan)#vlan 5-7
Switch(config-vlan)#exit
```

Конфигурация режима VLAN порта кольца как trunk и добавление ERPs протокола и данных VLAN:

```
Switch(config)# interface ge1/1
Switch(config-ge1/1)# switchport mode trunk
Switch(config-ge1/1)# switchport trunk allowed vlan add 5
Switch(config-ge1/1)# switchport trunk allowed vlan add 6
Switch(config-ge1/1)# switchport trunk allowed vlan add 7
Switch(config-ge1/1)#exit
Switch(config)# interface ge1/3
Switch(config-ge1/3)# switchport mode trunk
Switch(config-ge1/3)# switchport trunk allowed vlan add 5
Switch(config-ge1/3)# switchport trunk allowed vlan add 6
Switch(config-ge1/3)# switchport trunk allowed vlan add 7
Switch(config-ge1/3)#exit
```

Конфигурация ERPs instance 2 и ERPs одиночного кольца ring 2:

```
Switch(config)#erps 2
Switch(config-erps-2)#ring 2
Switch(config-erps-2)# ring 2 ring-mode major-ring
Switch(config-erps-2)# ring 2 node-mode rpl-neighbor-node
Switch(config-erps-2)# ring 2 raps-vlan 5
Switch(config-erps-2)# ring 2 traffic-vlan 1
Switch(config-erps-2)# ring 2 traffic-vlan 6
Switch(config-erps-2)# ring 2 traffic-vlan 7
Switch(config-erps-2)# ring 2 rpl-port ge1/3
Switch(config-erps-2)# ring 2 rl-port ge1/1
Switch(config-erps-2)# ring 2 enable
Switch(config-erps-2)#exit
```

12. AAA Configure (Конфигурация AAA)

В данной главе описывается конфигурация протоколов 802.1x и RADIUS коммутатора для предотвращения неавторизованного подключения к сети. Глава содержит следующие разделы:

- 802.1x Introduction (введение)
- RADIUS Introduction (введение)
- Configure 802.1x (конфигурация)
- Configure RADIUS (конфигурация)

AAA обозначает аутентификацию (authentication), авторизацию (authorization) и учет (accounting). Протокол обеспечивает три функции безопасности: authentication, authorization и billing. Конфигурация AAA является видом управления сетевой безопасностью. В данном случае сетевая безопасность сводится к контролю доступа (пользователи имеющие доступ к сети, сервисы доступные пользователям, учет пользователей сети).

- Authentication: получение доступа пользователем.
- Authorization: сервисы доступные пользователям.
- Accounting: учет использования сетевых ресурсов пользователем.

Обычно используются все возможности протокола AAA, включая 802.1x клиентов, оборудование поддерживающее authentication и hyperboss. Клиенты 802.1x устанавливаются на ПК для выхода в интернет. Для подключения к сети пользователи должны использовать клиента 802.1x для аутентификации. Пользователь получает запрос от клиента и передает имя пользователя (user name) и пароль (password) системе system hyperboss для аутентификации. Коммутатор не проводит аутентификацию. Hyperboss получает запрос от коммутатора, предоставляет актуальные данные и пропускает пользователей, прошедших аутентификацию.

Протокол 802.1x используется для коммуникации между 802.1x клиентом и коммутатором, протокол radius используется для коммуникации между коммутатором и hyperboss.

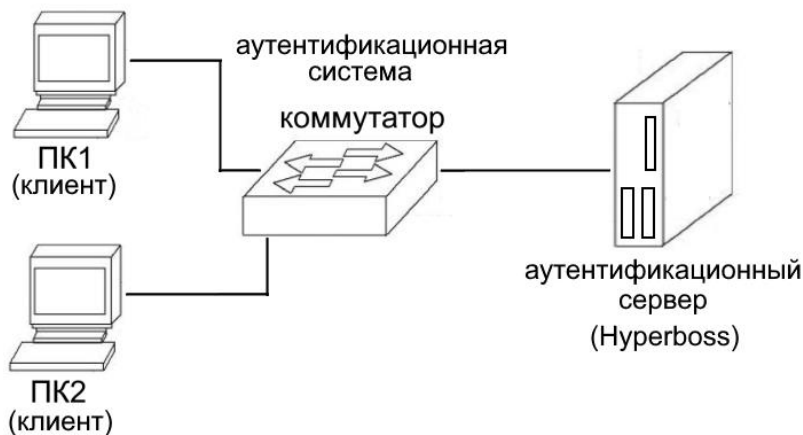
12.1 802.1x introduction (Введение)

Протокол аутентификации 802.1x основан на управления доступом порта. Порт может быть логическим, физическим с использованием MAC адреса или VLAN ID. Сетевые коммутаторы реализуют протокол 802.1x основанный на MAC адресе.

Протокол 802.1x является протоколом второго уровня layer-2. Коммутатор и ПК пользователя должны находиться в одной подсети, пакет протокола не может выходить за пределы сегмента сети. Аутентификация 802.1x воспринимает сервер клиента и должна быть сервером для всех пользователей. До аутентификации пользователя, только аутентификационные данные могут пройти через порт коммутатора. После успешной аутентификации поток данных сможет пройти через порт коммутатора, т.е. для доступа к сети пользователь должен пройти аутентификацию.

12.1.1 802.1x Equipment composition (Подключение оборудования)

802.1x оборудование делится на три группы: клиент (supplicant system), аутентификационная система (authenticator system) и аутентификационный сервер (authentication server). Схема подключения приведена ниже.



Клиент обращается к оборудованию с запросом на доступ к сети (обычно это ПК пользователя). 802.1x клиентское ПО (client software), которое отвечает за клиентскую часть протокола 802.1x должно быть установлено на ПК пользователя. Клиент инициирует запрос на 802.1x аутентификацию к серверу (authentication server) для подтверждения имени и пароля (user name, password). После успешной аутентификации пользователь получает доступ к сети.

Аутентификационная система относится к аутентификационному оборудованию (коммутатору). Аутентификационная система управляет доступом пользователя к сети путем изменения статуса логического порта пользователя (на основе MAC адреса). Если логический порт пользователя не авторизован, пользователь не может подключиться к сети, если логический порт пользователя авторизован, то пользователь получает доступ к сети.

Аутентификационная система соединяет клиента и аутентификационный сервер (authentication server). Аутентификационная система запрашивает идентификационную информацию пользователя, передает её на аутентификационный сервер и далее результат аутентификации от сервера к клиенту. Аутентификационная система требует реализации протокола 802.1x со стороны сервера и пользователя, а протокола radius со стороны клиента и аутентификационного сервера. Протокол radius клиента аутентификационной системы инкапсулирует информацию EAP отправленную клиентом 802.1x и отправляет её аутентификационному серверу. Информация EAP распаковывается из пакета протокола radius отправленного аутентификационным сервером и переданного клиенту 802.1x через 802.1x сервер.

Аутентификационный сервер это оборудование, которое проводит аутентификацию пользователя. Аутентификационный сервер получает идентификационную информацию пользователя от аутентификационной системы и верифицирует её. При успешной аутентификации, аутентификационный сервер авторизует аутентификационную систему и позволяет пользователю подключиться к сети. Если аутентификация не состоялась, аутентификационный сервер сообщает это аутентификационной системе и пользователь не получает доступа к сети. Аутентификационный сервер и аутентификационная система взаимодействуют по расширенному EAP RADIUS протоколу.

12.1.2 Protocol package (Пакеты протокола)

Поток аутентификационных данных передаваемый протоколом 802.1x по сети имеет кадровый формат eapol (EAP over LAN). Вся идентификационная информация пользователя (включая user name и password) инкапсулируется в EAP (extended authentication protocol), и EAP инкапсулируется в eapol кадр. Имя пользователя представлено в EAP в виде текста, тогда как пароль представлен в зашифрованном формате MD5.

Формат кадра Eapol представлен таблице ниже. Тип PAE Ethernet является типом Ethernet протокола с числом eapol и значением 0x888e. Версия протокола это версия eapol со значением 1. Тип пакета относится к формату кадра eapol. Длина пакета это длина содержания кадра eapol, содержание пакета относится к содержанию кадра eapol.

Формат кадра EAPOL:

	Octet Number
PAE Ethernet Type	1-2
Protocol Version	3
Packet Type	4
Packet Body Length	5-6
Packet Body	7-N

Коммутатор использует три кадра eapol протокола:

- Eapol start: значение типа пакета 1. Аутентификация инициализирует кадр. Когда пользователь запрашивает аутентификацию, формируется кадр и отправляется от клиента коммутатору.
- Eapol logoff: значение типа пакета 2. Кадр запроса на выход отправляется для уведомления коммутатора о том, что клиент не нуждается в подключении к сети.
- EAP packet: значение типа пакета 0. Информационный кадр аутентификации используется для передачи аутентификационной информации.

Формат пакета EAP представлен таблице ниже. Код относится к типу EAP пакета, включая запрос, ответ, успех и неудача. Идентификатор определяет соответствие ответа и запроса. Длина соответствует длине пакета EAP, включая заголовок. Дата соответствует данным EAP пакета.

EAP пакеты бывают четырех типов:

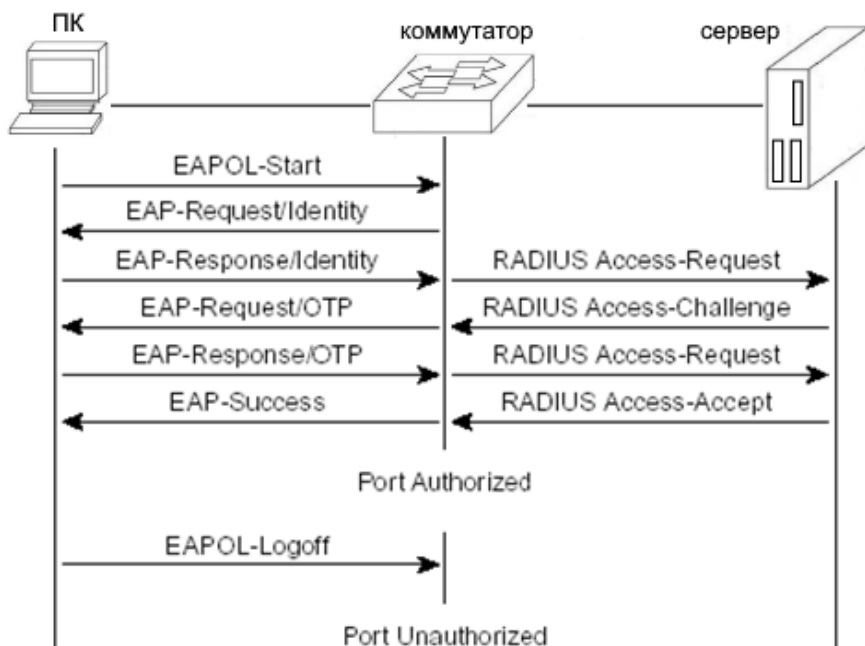
- EAP request: Значение кода 1, EAP пакет запроса отправляется коммутатором клиенту для запроса user name и password.
- EAP response: Значение кода 2, EAP пакет ответа отправляется клиентом коммутатору, содержит user name и password.
- EAP success: Значение кода 3, EAP пакет success отправляется коммутатором клиенту, информирует об успешной аутентификации пользователя.
- EAP failure: Значение кода 4. EAP пакет failure отправляется коммутатором клиенту, информирует о неудачной аутентификации пользователя.

Формат пакета EAP:

	Octet Number
Code	1
Identifier	2
Length	3-4
Data	5-N

12.1.3 Protocol flow interaction (Взаимодействие протоколов)

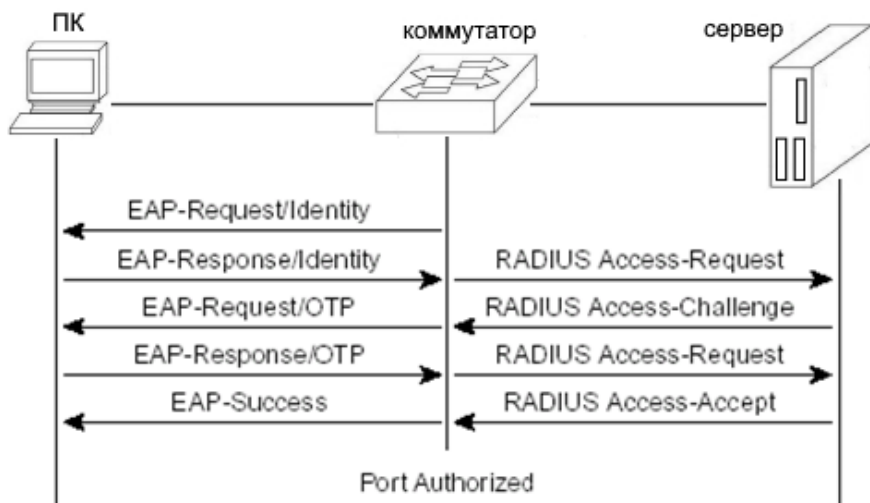
Когда на коммутаторе включен протокол 802.1x а порт находится в статусе auto, все пользователи, подключенные к порту должны пройти процедуру аутентификации перед получением доступа в сеть. Взаимодействие протоколов показано на рисунке ниже.



Когда пользователю необходимо подключиться к сети, клиент отправляет коммутатору eapol start для запроса аутентификации. После получения ответа, коммутатор отправляет EAP запрос пользователю о user name, клиент возвращает EAP ответ. Коммутатор извлекает информацию EAP и инкапсулирует её в пакет radius и отправляет его аутентификационному серверу, который запрашивает пароль пользователя. Коммутатор отправляет запрос пароля пользователя EAP клиенту, клиент возвращает EAP ответ. Коммутатор инкапсулирует информацию EAP в пакет radius и отправляет аутентификационному серверу, который аутентифицирует пользователя согласно имени и паролю. Если коммутатор получает успешный статус EAP аутентификации, то клиент будет уведомлен об этом. Получение клиентом EAP success, означает успешную аутентификацию и возможность доступа к сети.

По окончании сеанса подключения к сети, клиент отправляет уведомление eapol logoff коммутатору, коммутатор переводит логический порт в неавторизованный статус. Одновременно с этим пользователь отключается от сети.

Для предотвращения ненормального отключения клиента, коммутатор задействует процедуру повторной аутентификации. На коммутаторе может быть установлен временной интервал повторной аутентификации. По истечении этого времени, коммутатор инициализирует повторную аутентификацию. При успешной аутентификации, пользователь может продолжить пользоваться сетью. Взаимодействие протоколов показано на рисунке ниже.



12.1.4 802.1x port status (Статусы 802.1x порта)

Физический порт коммутатора может иметь четыре статуса: n/a state (неопределенный), auto state (авто), force authorized state (принудительно авторизованный) и force unauthorized (принудительно неавторизованный). Если на коммутаторе отключен протокол 802.1x, все порты находятся в неопределенном статусе (n/a). Перед тем как установить порт коммутатора в статус авто (auto state), авторизованный (force authorized) или неавторизованный (force unauthorized), следует включить 802.1x протокол.

Когда порт коммутатора находится в статусе n/a, все пользователи подключенные к порту могут получить доступ в сеть без аутентификации. Полученные от порта пакеты протокола 802.1x

коммутатор отбрасывает.

Когда порт коммутатора находится в статусе force authorized, все пользователи подключенные к порту могут получить доступ в сеть без аутентификации. Когда коммутатор получает от порта пакеты eapol start, коммутатор возвращает обратно EAP success пакеты. Полученные от порта пакеты протокола 802.1x коммутатор отбрасывает.

Когда порт коммутатора находится в статусе force unauthorized, все пользователи подключенные к порту не могут получить доступ в сеть, аутентификационные запросы не проходят. Полученные от порта пакеты протокола 802.1x коммутатор отбрасывает.

Когда порт коммутатора находится в статусе auto, все пользователи подключенные к порту должны пройти аутентификацию для получения доступа в сеть. Работа протокола 802.1x показана на рисунках выше. Для прохождения аутентификации пользователем порт коммутатора должен быть переведен в статус auto.

Когда порт коммутатора находится в статусе auto, одновременно включается функция anti ARP Spoofing. Функция anti ARP Spoofing может контролировать пакеты данных источника MAC, источника IP из группы IP пакетов с информацией предоставленной клиентом в процессе аутентификации, и пакетами данных отправителя IP и MAC из ARP пакетов с информацией предоставленной клиентом в процессе аутентификации и пересланных портом, в случае несовпадения пакеты будут сброшены. Для настройки этой функции, клиент должен иметь статический IP адрес. Если IP адрес назначается динамически протоколом DHCP, то для использования этой функции следует активировать DHCP snooping protocol. Для получения более подробной информации, следует обратиться к разделу IP MAC binding configuration.

12.2 RADIUS introduction (Введение)

Протокол Radius поддерживает EAP расширение для обеспечения взаимодействия между коммутатором и аутентификационным сервером. Протокол Radius соответствует модели клиент/сервер (client / server). Коммутатору необходимо реализовать radius клиент, а аутентификационный сервер должен реализовать radius сервер.

Для обеспечения безопасности взаимодействия коммутатора и аутентификационного сервера и предотвращения взаимодействия с

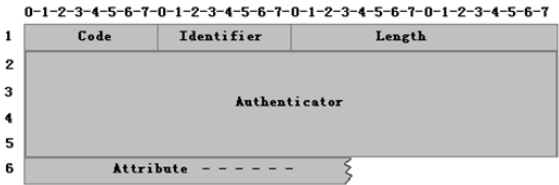
неавторизованными коммутаторами и серверами, коммутатор и аутентификационным сервер должны аутентифицировать друг друга. Коммутатор и аутентификационный сервер должны иметь один и тот же ключ. При отправлении пакетов протокола RADIUS, должны генерироваться краткие сообщения (message summaries) с использованием HMAC алгоритма согласно ключу. Когда коммутатор и аутентификационным сервер получают пакеты протокола RADIUS, краткое сообщение (message summaries) должно быть верифицировано ключем. При удачной верификации, пакет протокола RADIUS считается легальным, в противном случае (нелегальный) пакет протокола RADIUS сбрасывается.

12.2.1 Protocol package (Пакеты протокола)

Протокол Radius основан на UDP. Протокол Radius может инкапсулировать аутентификационную и биллинговую (billing) информацию. Ранний аутентификационный порт radius - 1645, текущий порт - 1812. Ранний radius billing порт - 1646, текущий порт - 1813.

Так как протокол radius хостируется на UDP, существует механизм задержки передачи (timeout retransmission). Для увеличения надежности коммуникации между аутентификационной системой и сервером radius, предусмотрено две схемы сервера radius, в т.ч. механизм ожидания сервера (standby server mechanism).

Формат сообщения (radius message) показан на рисунке ниже. Код относится к типу сообщения протокола radius. Индикатор идентификации используется для определения соответствия запроса и ответа. Длина - это длина всего сообщения (включает заголовок сообщения). Аутентификатор – строка длиной 16 байт, случайный номер пакетов запроса и дайджест сообщений сгенерированный MD5 для ответных пакетов. Атрибут относится к атрибуту пакета протокола radius.



В сети используются следующие пакеты протокола RADIUS:

- Access request (запрос доступа): значение кода - 1, пакет запроса аутентификации отправлен от системы аутентификации к аутентификационному серверу. Пакет инкапсулирует имя пользователя и пароль.

- Access accept (разрешение доступа): значение кода - 2. Пакет ответа отправленный от аутентификационного сервера к системе аутентификации указывает, что аутентификация прошла успешно.

- Access reject (отказ доступа): значение кода - 3. Пакет ответа отправленный от аутентификационного сервера к системе аутентификации указывает, что аутентификация не пройдена.

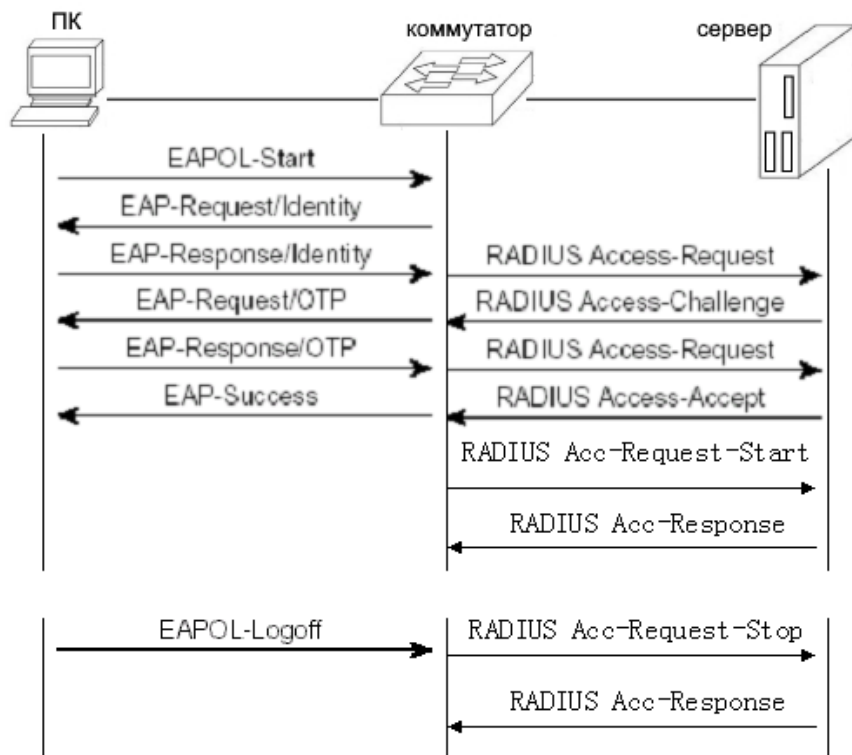
- Access challenge (вызов доступа): значение кода - 11, Пакет ответа отправленный от аутентификационного сервера к системе аутентификации указывает, что аутентификационному серверу требуется больше информации (имя пользователя, пароль и т.п.).

- Accounting request (счет запроса): значение кода - 4. Пакет запроса биллинга (billing) отправленный от системы аутентификации к аутентификационному серверу включает в себя стартовый и конечный пакеты биллинга (start billing package). Пакет инкапсулирует биллинговую информацию.

- Accounting response (счет ответа): значение кода - 5. Пакет ответа биллинга отправленный от аутентификационного сервера к системе аутентификации указывает, что информация о биллинге получена.

12.2.2 Protocol flow interaction (Взаимодействие протокола)

Когда пользователь инициирует процедуру аутентификации, система аутентификации и аутентификационный сервер взаимодействуют по протоколу radius. После успешной аутентификации или перехода пользователя в offline, системе аутентификации необходимо отправить пакет radius charging к аутентификационному серверу. Взаимодействие потоков протокола, когда система аутентификации не отправляет пакеты биллинга, показана на рисунке ниже.



При аутентификации пользователя, коммутатор инкапсулирует в сообщении запроса имя пользователя и отправляет его аутентификационному серверу. Сервер отвечает на запрос access challenge запросом пароля. Коммутатор запрашивает пароль пользователя (клиента). Клиент инкапсулирует пароль в EAP. После получения EAP, коммутатор инкапсулирует его в access request и отправляет аутентификационному серверу. Аутентификационный сервер аутентифицирует пользователя, и в случае успеха отправляет коммутатору принятие (access accept). После получения этого сообщения коммутатор уведомляет клиента об успешной аутентификации. Одновременно он отправляет запрос accounting request для уведомления аутентификационного сервера о начале обмена, а сервер отправляет обратно ответ accounting response.

Если пользователь не планирует подключаться к сети, коммутатор получает уведомление о переходе в offline, то коммутатор

отправляет запрос accounting request для уведомления аутентификационного сервера об окончании биллинга. Пакет инкапсулирует биллинговую информацию, и аутентификационный сервер отправляет обратно сообщение accounting response.

12.2.3 User authentication method (Методы аутентификации)

Протокол Radius предлагает три метода аутентификации:

- PAP (Password Authentication Protocol) Текстовый ввод имени пользователя и пароля. Коммутатор передает имя пользователя и пароль серверу radius с помощью пакетов протокола radius. Сервер radius проверяет базу данных на соответствие имени и пароля. При совпадении аутентификация считается пройденной успешно, в противном случае аутентификация считается не пройденной.

- CHAP (Challenge Handshake Authentication Protocol) Генерация 16и байтного случайного кода пользователя (используется для подключения к Internet). Пользователь шифрует случайный код, пароль и домены для генерации ответа и передают имя пользователя и ответ коммутатору. Коммутатор передает имя пользователя, ответ и оригинальный 16и байтный случайный код серверу radius. Radius проверяет базу данных коммутатора на наличие имени пользователя, получает тот же зашифрованный пароль пользователя. Далее шифрует его в соответствии с переданным 16и байтным случайным кодом, и проверяет результат на соответствие с переданным ответом. При совпадении аутентификация считается пройденной успешно, в противном случае аутентификация считается не пройденной.

- EAP (Extensible Authentication Protocol) Без участия коммутатора в процессе аутентификации, коммутатор используется только для передачи между пользователем и сервером radius. Когда пользователь запрашивает доступ в Internet, коммутатор запрашивает имя пользователя и передает его серверу radius. Сервер radius генерирует 16и байтный случайный код для пользователя и сохраняет его. Пользователь шифрует случайный код, пароль и домены для генерации ответа и передает имя пользователя и ответ коммутатору, который передает их серверу radius. Radius проверяет базу данных коммутатора на наличие имени пользователя, получает тот же зашифрованный

пароль пользователя. Далее шифрует его в соответствии с переданным 16и байтным случайным кодом, и проверяет результат на соответствие с переданным ответом. При совпадении аутентификация считается пройденной успешно, в противном случае аутентификация считается не пройденной.

12.3 Configure 802.1x (Конфигурация 802.1x)

Данный раздел содержит подробное описание конфигурации протокола 802.1x in и включает в себя следующие подразделы:

- 802.1x Default configuration (конфигурация по умолчанию)
- Turn 802.1x on and off (включение / отключение)
- Configure 802.1x port status (статусы портов)
- Configure re authentication mechanism (повторная аутентификация)
- Configure the maximum number of port access hosts (управление доступом хостов)
- Configure interval and number of retransmissions (интервалы передачи)
- Configure port as transport port (транспортные порты)
- Configure 802.1x client version number (номер версии клиента)
- Configure whether to check the client version number (проверка номера версии клиента)
- Configure authentication method (метод аутентификации)
- Configure whether to check the timing package of the client (проверка тайминга пакетов клиента)
- 802.1x display information (просмотр информации)

12.3.1 802.1xDefault configuration (Конфигурация по умолчанию)

Конфигурация протокола 802.1x по умолчанию представлена ниже:

- 802.1x is closed (протокол отключен)
- Status of all ports is N/A (статус всех портов не определен)
- Recertification mechanism is turned off, and the interval between recertification is 3600 seconds (ресертификация отключена, временной интервал 3600 сек)
- Maximum number of access hosts for all ports is 100 (максимальное

- число доступа к хостам – 100 для всех портов)
- Timeout interval for resending EAP request is 30 seconds (временной интервал отправки запросов EAP 30 сек)
 - Number of times to resend EAP request after timeout is 3 (количество отправок запросов EAP – 3)
 - Waiting time for user authentication failure is 60 seconds (Время ожидания повторной попытки аутентификации после неудачной попытки 60 сек)
 - Time interval of timeout retransmission at the server is 10 seconds (временной интервал повторной отправки серверу – 10 сек)

Команда возврата коммутатора к настройкам по умолчанию в режиме global config mode:

```
Switch(config)#dot1x default
```

12.3.2 Turn 802.1x on and off (Включение и отключение 802.1x)

Команда включения протокола 802.1x в режиме global config mode:

```
Switch(config)#dot1x
```

Команда отключения протокола 802.1x в режиме global config mode:

```
Switch(config)#no dot1x
```

При отключении 802.1x, все порты возвращаются в статус N/A (не определен).

12.3.3 Configure 802.1x port status (Выбор статуса портов)

Перед установкой статуса порта убедитесь что протокол 802.1x включен. Если все пользователи должны пройти аутентификацию для получения доступа к сети, то порт должен быть установлен в автоматический статус (auto state).

Команда установки автоматического статуса порта GE1/1 (auto state) в режиме interface configuration mode и включение функции anti ARP Spoofing:

```
Switch(config-ge1/1)dot1x control auto
```

Возможные причины по которым функция anti ARP Spoofing не активируется:

- Ресурсы системы CFP исчерпаны.
- Функция ACL фильтрации включена на текущем интерфейсе.
- Функция DHCP snooping включена на текущем интерфейсе.
- Текущий интерфейс 3го уровня (three-layer) или находится в режиме trunk.

Команда установки статуса порта GE1/1 в force authorized status в режиме interface configuration mode:

```
Switch(config-ge1/1)dot1x control force-authorized
```

Команда установки статуса порта GE1/1 в force unauthorized state в режиме interface configuration mode:

```
Switch(config-ge1/1)dot1x control force-unauthorized
```

Команда установки статуса порта GE1/1 в N/A status в режиме interface configuration mode:

```
Switch(config-ge1/1)no dot1x control
```

Внимание: если порт имеет привязку к MAC адресу, то он не может быть переведен в статус auto, force authorized или force unauthorized status.

12.3.4 Configure re authentication (Конфигурация реаутентификации)

Для предотвращения аномального перехода клиента в offline относительно коммутатора и аутентификационного сервера, в коммутаторе предусмотрена процедура реаутентификации (повторной аутентификации), которая инициируется коммутатором через определенный промежуток времени (re authentication interval).

Команда запуска реаутентификации в режиме global config mode:

```
Switch(config)#dot1x reauthenticate
```

Команда отключения реаутентификации в режиме global config mode:

```
Switch(config)#no dot1x reauthenticate
```

Команда установки интервала реаутентификации в режиме global config mode:

```
Switch(config)#dot1x timeout re-authperiod <interval>
```

Внимание: интервал реаутентификации не должен быть слишком коротким, в противном случае снижается пропускная способность сети и производительность CPU (процессора) коммутатора.

12.3.5 Configure number of port access hosts (Управление хостами)

Каждый порт коммутатора может управлять доступом к определенному числу хостов. Эта функция ограничивает доступ к сети пользователям, которые могут использовать для неавторизованного доступа к сети несколько хостов. Порт может контролировать до 100 хостов (по умолчанию), их количество может быть уменьшено. Если будет установлено значение 0, то порт будет запрещать доступ к сети всем пользователям.

Команда установки максимального количества хостов доступа порта GE1/1 в режиме interface configuration mode:

```
Switch(config-ge1/1)dot1x support-host <number>
```

12.3.6 Configure interval and number of retransmissions (Интервалы передачи)

Стандарты протокола 802.1x определяют временные интервалы и количество передач (ретрансляций) необходимых для организации взаимодействия. Коммутатор использует стандартные значения этих параметров. Пользователям не рекомендуется изменять временные интервалы и количество ретрансляций при использовании периода TX относящегося к передаче пакетов запросов EAP протокола. Параметр Max req указывает количество ретрансляций запросов EAP request. Временной интервал используется для указания времени ожидания реаутентификации пользователем. Интервал Server timeout определяет временной промежуток ретрансляции коммутатором пакетов radius аутентификационному серверу. Интервал Sup timeout определяет временной промежуток отправки коммутатором пакетов EAP запросов клиенту.

Команда установки интервалов и числа ретрансляций в режиме global config mode:

```
Switch(config)#dot1x timeout tx-period <interval>  
Switch(config)#dot1x max-req <number>  
Switch(config)#dot1x timeout quiet-period <interval>  
Switch(config)#dot1x timeout server-timeout <interval>  
Switch(config)#dot1x timeout supp-timeout <interval>
```

12.3.7 Configure port as transport port (Транспортный порт)

Когда коммутатор не переключается в режим 802.1x аутентификации, а другие коммутаторы в подсети находятся в режиме 802.1x аутентификации, порт связанный с клиентом и аутентификационным коммутатором может быть сконфигурирован как транспортный порт, аутентификационный eapol пакет может быть отправлен от клиента к аутентификационному коммутатору. Аналогичным образом может быть реализована 802.1x аутентификация для других коммутаторов и клиентов.

Команда перевода в транспортный режим передачи порта GE1/1 в interface configuration mode:

```
Switch(config-ge1/1)dot1x transmit-port
```

Команда перевода в обычный режим передачи порта GE1/1 в interface configuration mode:

```
Switch(config-ge1/1)no dot1x transmit-port
```

12.3.8 Configure 802.1x client version number (Номер верси клиента)

Только клиент с версией не ниже установленной может пройти аутентификацию. На коммутаторе установлена версия клиента 2.0 (по умолчанию).

Команда установки версии клиента в режиме global config mode:

```
Switch(config)# dot1x client-version <string>
```

12.3.9 Configure check the client version number (Проверка номера версии клиента)

Функция проверки номера версии клиента включена по умолчанию. Если данная функция включена, то во время проведения аутентификации, коммутатор сначала должен проверить номер версии клиента.

Команда включения функции проверки номера версии клиента в режиме global config mode:

```
Switch(config)# dot1x check-version open
```

12.3.10 Configure authentication method (Выбор метода аутентификации)

Режим аутентификации выбирается клиентом и разделяется на основную и расширенную аутентификацию. Коммутатор может быть настроен на применение одного из методов в первую очередь. Если клиент выбирает метод аутентификации несоответствующий установленному на коммутаторе, то после нескольких неудачных попыток аутентификации, клиенту будет предложено начать аутентификацию по другому методу.

Команда выбора расширенного метода аутентификации (extended mode) в режиме global config mode:

```
Switch(config)# dot1x extended
```

12.3.11 Configure to check the timing package of the client (Проверка тайминга пакета клиента)

Включение данной функции после настройки тайминга протокола 802.1x, активирует проверку всех пакетов на прохождение аутентификации.

Команда настройки проверки тайминга пакетов клиента в режиме global config mode:

```
Switch(config)# dot1x check-client
```

12.3.12 Display 802.1x information (Информация 802.1x)

Команда просмотра информации о настройках 802.1x в режимах normal mode / privileged mode. Команда dot1x выводит для просмотра информацию о всех настройках 802.1x, включая конфигурацию портов. Команда dot1x interface выводит для просмотра информацию о всех клиентах имеющих доступ к порту:

```
Switch#show dot1x
```

```
Switch#show dot1x interface
```

12.4 Configure RADIUS (Конфигурация RADIUS)

Данный раздел содержит подробное описание конфигурации протокола RADIUS in и включает в себя следующие подразделы:

- Radius default configuration (конфигурация по умолчанию)
- Configure the IP address of the authentication server (IP адрес сервера аутентификации)
- Configure shared key (конфигурация общего ключа)
- Start and close billing (включение/отключение биллинга)
- Configure radius port and attribute information (конфигурация порта и информация атрибута)
- Configure radius roaming function (настройка роуминга)
- Display radius information (информация radius)

12.4.1 RADIUS Default configuration (Конфигурация по умолчанию)

Конфигурация протокола RADIUS по умолчанию представлена ниже:

- IP адрес основного и резервного аутентификационного сервера не установлены (IP адрес 0.0.0.0)
- Ключ shared не сконфигурирован, (string shared ключа не задан)
- Биллинг включен (по умолчанию)
- UDP аутентификационный порт – 1812, UDP биллинговый порт - 1813
- Значение атрибута nasport - 0xc353, значение nasporttype - 0x0f, значение nasportserver - 0x02

12.4.2 IP address of the authentication server (IP адрес сервера аутентификации)

Для обеспечения взаимодействия протокола radius между коммутатором и сервером аутентификации, необходимо прописать в коммутаторе IP адрес сервера аутентификации. На практике может использоваться один или два сервера аутентификации (основной и резервный). Если в коммутаторе прописано два IP адреса аутентификационных серверов, то коммутатор может взаимодействовать с резервным сервером, если соединение с основным сервером по какой-либо причине прервано.

Команда ввода IP адреса основного сервера аутентификации в режиме global config mode:

```
Switch(config)#radius-server host <ip-address>
```

Команда ввода IP адреса резервного сервера аутентификации в режиме global config mode:

```
Switch(config)#radius-server option-host <ip-address>
```

12.4.3 Configure shared key (Конфигурация общего ключа)

Между коммутатором и аутентификационным сервером требуется взаимная аутентификация, для чего необходим общий ключ (shared key), который должен быть установлен на коммутаторе и сервере аутентификации.

Команда ввода конфигурации общего ключа (shared key) в режиме global config mode:

```
Switch(config)#radius-server key <string>
```

12.4.4 Start and close billing (включение/отключение биллинга)

Коммутатор прекращает отправку пакетов протокола radius серверу аутентификации по окончании успешной аутентификации или когда пользователь переходит в режим offline. Обычно в практических случаях функция биллинга (billing) включена.

Команда запуска биллинга в режиме global config mode:

```
Switch(config)#radius-server accounting
```

Команда остановки биллинга в режиме global config mode:

```
Switch(config)#no radius-server accounting
```

12.4.5 Configure radius port and attribute information (порт radius и информация атрибута)

Внимание: не рекомендуется вносить изменения в конфигурацию порта radius и информацию атрибута.

Команда изменения конфигурации radius UDP порта в режиме global config mode:

```
Switch(config)#radius-server udp-port <port-number>
```

Команды изменения информации атрибута (radius attribute) в режиме global config mode:

```
Switch(config)#radius-server attribute nas-portnum <number>
```

```
Switch(config)#radius-server attribute nas-porttype <number>
```

```
Switch(config)#radius-server attribute service-type <number>
```

12.4.6 Configure radius roaming function (Функция роуминга)

В случае когда клиент привязан к Mac, IP адресу или VLAN изменяет свою локацию, перестают выполняться условия аутентификации по протоколу 802.1x. Функция роуминга (radius roaming function) позволяет обойти требования привязки клиента к Mac, IP адресу или VLAN и делает возможной процедуру аутентификации согласно 802.1x протоколу.

Команда включения функции роуминга в режиме global config mode:

```
Switch(config)#radius-server roam
```

Команда отключения функции роуминга в режиме global config mode:

```
Switch(config)#no radius-server roam
```

12.4.7 Display Radius information (Информация Radius)

Команда просмотра информации о настройках radius в режимах normal mode / privileged mode:

```
Switch#show radius-server
```

12.4.8 Configuration example (Пример конфигурации)

Запустить протокол 802.1x, установить порт GE1/1 в режим auto, сконфигурировать основной сервер аутентификации (IP 198.168.80.111) и общий ключ (shared key) коммутатора как ABCDEF.

```
Switch#  
Switch# dot1x  
Switch#config t  
Switch(config)#radius-server host 198.168.80.111  
Switch(config)#radius-server key abcdef  
Switch(config)# interface ge1/1  
Switch(config-ge1/1)# dot1x control auto
```

12.4 TACACS+ Introduction (Протокол TACACS+)

TACACS+ протокол аутентификации и авторизации предоставляет возможности управления авторизацией пользователей, верификации легитимности пользователей и авторизацией команд. При активированной аутентификации TACACS +, пользователю необходимо верифицировать имя пользователя и пароль через сервер TACACS+ для доступа к коммутатору. Пройти аутентификацию для доступа к коммутатору возможно только с корректными данными имени пользователя и паролем.

TACACS+ разделяет права пользователей на два уровня: обычные пользователи и привилегированные пользователи. Обычные пользователи имеют право использовать обычный режим (normal mode) интерфейса ввода CLI команд, привилегированные пользователи имеют доступ ко всем режимам интерфейса ввода CLI команд. Выполнение команд основано на правах различного уровня. Когда пользователь вводит команду (кроме enable end и exit), на сервере

TACACS+ проводится верификация пользователя. Если верификация не пройдена, то выполнения команды не происходит.

Аутентификация и авторизация TACACS+ применима только для терминалов telnet и SSH, но не применима для управления через консоль (console terminals). При доступе к коммутатору через терминалы telnet и SSH, имя пользователя и пароль должны быть верифицированы. Доступ к командной строке CLI может быть получен только после верификации имени пользователя и пароля. Доступ к SSH возможен только для привилегированных пользователей. Аутентификация TACACS+ также применяется для доступа к web-интерфейс, но только для верификации пароля привилегий и разрешений без авторизации команд.

Функция TACACS+ на коммутаторе по умолчанию отключена. Telnet, SSH или web используют многопользовательский режим управления. После включения функции TACACS+ доступна настройка многопользовательского режима, но сам многопользовательский режим на практике не используется.

Таблица команд аутентификации TACACS + протокола.

Команда	Описание	CLI режим
tacacsplus enable	Включение TACACS+	Global configuration mode
tacacsplus disable	Отключение TACACS + function	Global configuration mode
tacacsplus host <server-ip>	IP адрес основного сервера с поддержкой IPv4 и IPv6. Рекомендуется использовать ACS (Cisco)	Global configuration mode
tacacsplus option-host <server-ip>	IP адрес standby сервера с поддержкой IPv4 and IPv6. Рекомендуется использовать ACS (Cisco)	Global configuration mode
tacacsplus key WORD	Конфигурация общего ключа для шифрования передаваемых данных (должен соответствовать конфигурации сервера).	Global configuration mode

tacacsplus auth-type (PAP CHAP)	Выбор метода аутентификации (включая PAP и chap). PAP (по умолчанию), поле инкапсулирует secret. Chap инкапсулирует проверку кода MD5 secret.	Global configuration mode
show tacacsplus	Просмотр конфигурации TACACS+	Global configuration mode
no tacacsplus host	Сброс IP адреса основного сервера	Global configuration mode
no tacacsplus key	Сброс общего ключа (shared key)	Global configuration mode

13. GMRP Configure (Конфигурация GMRP)

В данной главе описывается конфигурация протокола GMRP (GARP Multicast Registration Protocol) – протокол групповой регистрации GARP. Глава содержит следующие разделы:

- GMRP Introduction (введение)
- Configure GMRP (конфигурация GMRP)
- Display GMRP (информация GMRP)

13.1 GMRP Introduction (Протокол GMRP)

Протокол GMRP (GARP Multicast Registration Protocol) – протокол групповой регистрации основанный на GARP, используется для обслуживания групповой регистрации информации в коммутаторе. Все коммутаторы поддерживающие GMRP могут получать групповую информацию от других коммутаторов, обновляя локальную групповую информацию и передавать её другим коммутаторам. Такой механизм обмена информацией в локальной сети обеспечивается всем оборудованием поддерживающим GMRP.

Когда хост предполагает вступить в группу multicast, он направляет GMRP запрос-сообщение. Коммутатор добавляет порт

получивший запрос-сообщение GMRP в multicast группу и направляет запрос-сообщение GMRP в VLAN которому принадлежит принявший сообщение порт. Источник multicast в VLAN имеет информацию о членах группы multicast. Когда источник multicast отправляет сообщение в группу multicast, коммутатор направляет сообщение multicast порту соединенному с членами multicast группы для реализации функции multicast второго уровня (two-layer) в VLAN.

13.2 GMRP configurtion (Протокол GMRP)

Основная конфигурация GMRP включает в себя:

- Open GMRP (включение GMRP)
- View GMRP (просмотр информации GMRP)

Перед включением порта GMRP, должен быть запущен глобальный (global) GMRP протокол.

13.2.1 GMRP settings (Установки GMRP)

Таблица команд установок GMRP

Команда	Описание	CLI режим
set gmrp enable disable	Включение / Отключение global VLAN GMRP	Global configuration mode
set gmrp enable vlan <vlan-id>	Включение global specific VLAN GMRP	Global configuration mode
set gmrp registration{fixed forbidden normal} <if-name>	Конфигурация интерфейса регистраци режима multicast	Global configuration mode
set gmrp timer {join leave nleaveall} <time-value>	Установка временных интервалов таймеров	Global configuration mode
set port gmrp enable <if-name>	Включение функции GMRP для порта	Global configuration mode
set port gmrp disable <if-name>	Отключение функции GMRP для порта	Global configuration mode

13.2.2 View GMRP information (Просмотр GMRP информации)

Для просмотра и проверки установленной конфигурации GMRP в режиме privileged mode используйте команды приведенные в таблице ниже:

Таблица команд просмотра установок GMRP

Команда	Описание	CLI режим
show gmrp configuration	Просмотр информации о конфигурации GMRP	Privileged mode
show gmrp machine	Просмотр информации GMRP state machine	Privileged mode
show gmrp statistics vlanid	Просмотр статистики GMRP выбранных vlan id.	Privileged mode
show gmrp timer <ifname>	Просмотр информации таймера выбранного порта.	Privileged mode

13.2.3 GMRP Configuration example (Пример конфигурации)

Для реализации динамической регистрации и обновления информации multicast между коммутаторами необходимо включить GMRP протокол на коммутаторе. Ниже представлена схема подключения коммутаторов.



Конфигурация коммутатора sw1:

Запуск global GMRP

```
Switch(config)# set gmrp enable
```

Запуск GMRP на Gigabit Ethernet порту GE1/1

```
Switch(config)# set port gmrp enable ge1/1
```

```
Switch(config)#
```

Конфигурация коммутатора sw2:
Запуск global GMRP
Switch(config)# set gmrp enable
Запуск GMRP на Gigabit Ethernet порты GE1/2
Switch(config)# set port gmrp enable ge1/2
Switch(config)#

14. IGMP SNOOPING Configure (IGMP SNOOPING)

При использовании одноадресной передачи (unicast) пакетов нескольким получателем в сети, требуется их копирование для каждого конечного получателя. Количество передаваемых пакетов будет линейно возрастать вместе с количеством получателей, что вызовет повышение нагрузки на хосты, коммутаторы и сетевое оборудование и негативно повлияет на пропускную способность сети. Режим групповой передачи multicast становится более актуальным всяизи с распространением использования многоточечной коммуникации с большим объемом передаваемых данных (таких как multi-point видео конференции и т.п.). Режим групповой передачи multicast позволяет более эффективно использовать сетевые ресурсы.

Коммутаторы поддерживающие функцию IGMP snooping рассчитаны на применение и обслуживание групповой передачи multicast. Функция IGMP snooping осуществляет мониторинг IGMP пакетов в сети для реализации динамического отслеживания IP multicast и MAC адресов. Глава содержит следующие разделы:

- IGMP SNOOPING introduction (введение)
- IGMP SNOOPING configuration (конфигурация)
- IGMP SNOOPING configuration example (пример конфигурации)

14.1 IGMP SNOOPING Introduction (Протокол GMRP)

В обычных сетях multicast пакеты обрабатываются как вещательные пакеты в подсетях, что может вызвать рост трафика и возникновение задержек. Когда на коммутаторе включена функция

IGMP snooping, осуществляется мониторинг IGMP пакетов в сети для динамического отслеживания IP multicast и MAC адресов, обслуживание списка IP multicast MAC адресов выходного порта, что снижает сетевой трафик путем отправки данных multicast на выходной порт port.

14.1.1 IGMP SNOOPING Processing (IGMP SNOOPING процессы)

IGMP snooping является сетевым протоколом второго уровня (layer-2), который осуществляет отслеживание процессинга пакетов IGMP протокола в коммутаторе, обслуживает multicast группу на принимающем порту, VLAN ID и multicast адрес пакетов IGMP и отправку этих пакетов. Только порты относящиеся к multicast группе могут принимать данные multicast. Всё это снижает сетевой трафик и повышает пропускную способность сети.

Multicast группа включает адрес, порт-член, VLAN ID и время.

Формирование IGMP snooping multicast группы является обучаемым процессом. Когда порт коммутатора получает пакет отчета IGMP, IGMP snooping генерирует новую multicast группу, порт получающий пакет отчета IGMP добавляется в multicast группу. Когда коммутатор получает пакет запроса IGMP (если multicast группа уже существует), порт получающий запрос IGMP добавляется в multicast группу, в противном случае пакет запроса IGMP будет перенаправлен. IGMP snooping поддерживает механизм выхода из IGMP V2. Если IGMP snooping настроен на быстрый выход, то принимающий порт может выйти из multicast группы немедленно при получении leave пакета IGMP V2. Если настроено время задержки (leave timeout), то принимающий порт может выйти из группы multicast по истечении указанного временного интервала.

Существует два механизма обновления IGMP snooping. Один из них описан выше. В большинстве случаев IGMP snooping удаляет группы multicast по истечении времени age time. При вступлении multicast группы в IGMP snooping фиксируется время joining time, когда заданное время (age time) истекает, multicast группа удаляется из коммутатора.

Когда порт получает leave пакет, он удаляется из multicast группы немедленно, что может нарушить последовательность прохождения данных в сети, т.к. этому порту может быть подключено сетевое оборудование без поддержки функции IGMP snooping.

Отправление команды leave может оказать влияние на другое сетевое оборудование, которое выражается в невозможности принимать потоки данных multicast. Для предотвращения этого эффекта предусмотрен механизм fast leave timeout (настройка времени timeout). После получения leave пакета, порт ожидает окончания временного интервала для выхода из multicast группы, что гарантирует сохранение целостности потока данных в сети.

14.1.2 Dynamic Multicast Layer 2 (Динамический Multicast)

Multicast MAC адрес записан в таблице (forwarding table) layer 2 памяти устройства и может быть получен через IGMP snooping dynamic learning.

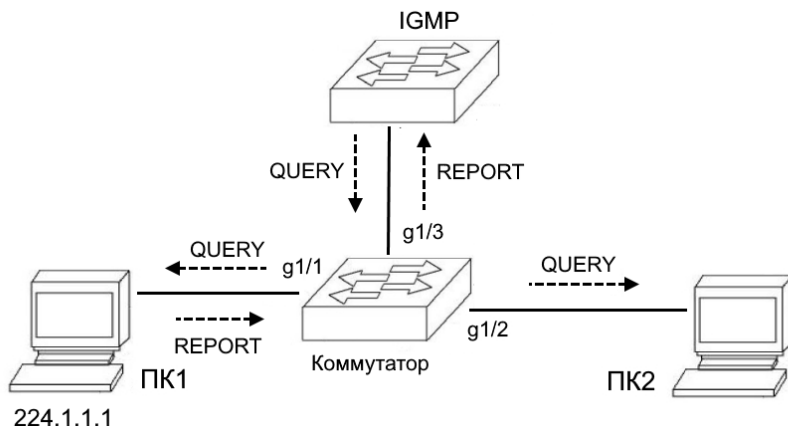
При отключении на коммутаторе IGMP snooping, аппаратная часть layer-2 оборудования таблица multicast forwarding находится в режиме unregistered forwarding mode и multicast MAC адрес не может быть динамически изменен. В таблице multicast forwarding table аппаратной части layer-2 оборудования отсутствуют записи, и все потоки данных multicast обрабатываются как вещательные.

Для повышения эффективности управления multicast трафиком в сети на коммутаторе может быть включен протокол IGMP snooping (если сеть имеет multicast окружение). В этом случае таблица multicast forwarding table layer-2 оборудования находится в режиме registered forwarding mode. Коммутатор может запоминать multicast MAC адрес путем получения пакетов IGMP протокола по сети. Только потоки multicast layer-2 соответствуют записям таблицы multicast forwarding table аппаратной части оборудования layer-2.

14.1.3 Join a group (Включение в группу)

Когда хост включается в multicast группу, он отправляет пакет IGMP report, который определяет multicast группу в которую будет включен хост. Когда коммутатор получает пакет IGMP query, он направляет пакет всем портам-членам VLAN. Когда хост получает пакет IGMP query, он возвращает пакет IGMP report. Когда коммутатор получает пакет IGMP report, он создает layer-2 multicast вход. Порт

получающий пакеты IGMP query и порт отправляющий IGMP report будут добавлены в layer-2 multicast вход и станут его выходными портами.



Как показано на рисунке выше, предполагается что все устройства в подсети относятся к VLAN 2. На роутере запущен протокол igmpv2 и отправляются пакеты IGMP query. Хост 1 должен вступить в multicast группу 224.1.1.1. После получения пакета IGMP query от порта 1/3, коммутатор запомнит этот порт и направит пакет портам 1/1 и 1/2. Хост 1 отправляет обратно пакет IGMP report после получения пакета IGMP. Хост 2 не отправляет пакет IGMP report т.к. не должен вступать в multicast группу. После получения пакета IGMP report от порта 1/1, коммутатор направит пакет query от порта 1/3 и создаст layer-2 multicast вход (если вход не существует).

Таблица элементов layer-2 multicast входа

Layer 2 multicast address	VLAN ID	Output port list
01:00:5e:01:01:01	2	1/1 , 1/3

Как показано на рисунке ниже, хост 1 вступил в multicast группы 224.1.1.1, а хост 2 должен вступить multicast группу 224.1.1.1. Когда хост 2 получает пакет IGMP query, он отправляет обратно пакет IGMP report. После получения пакета IGMP report от порта 1/2, коммутатор направит пакет query от порта 1/3, и добавляет пакет от порта 1/2 к layer-2 multicast входу.

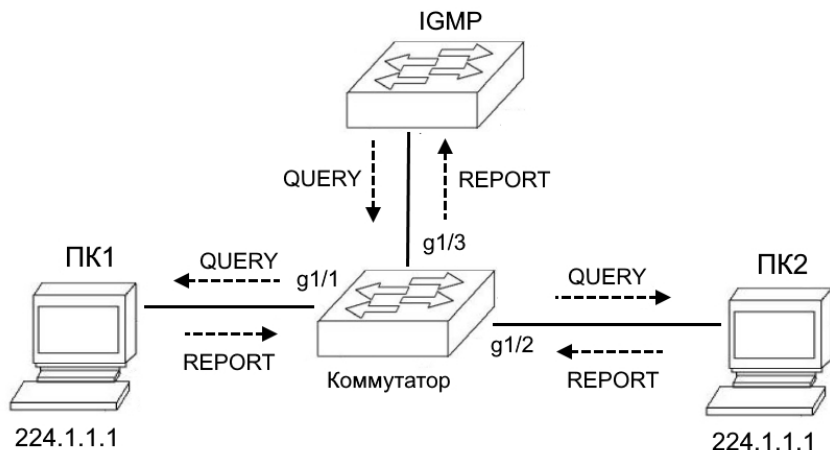


Таблица элементов layer-2 multicast входа

Layer 2 multicast address	VLAN ID	Output port list
01:00:5e:01:01:01	2	1/1 , 1/2 , 1/3

14.1.4 Left a group (Выход из группы)

Для того чтобы сформировать стабильную среду multicast, устройства с запущенным протоколом IGMP (такие как роутеры) регулярно отправляют пакеты IGMP query всем хостам. Хосты, вступившие или вступающие в multicast группу, отправят IGMP report в ответ на полученный запрос IGMP query.

Хост может покинуть multicast группу двумя путями – активным и пассивным. При активном варианте хост отправляет роутеру пакет IGMP leave. При пассивном варианте хост не отправляет роутеру пакет IGMP report после получения запроса IGMP query от роутера.

В зависимости от способа выхода хоста из multicast группы, существует два варианта выхода из layer-2 multicast entry входа порта коммутатора – по истечении времени timeout и путем получения пакета IGMP leave.

Когда коммутатор не получает пакет IGMP report от порта multicast группы более определенного времени, порт удаляется из

соответствующего layer 2 multicast entry входа. Если layer 2 multicast entry вход не имеет порта, то layer 2 multicast entry вход удаляется.

Когда на коммутаторе активирована функция быстрого выхода (fast leave), если порт получает пакет IGMP leave от multicast группы, порт удаляется из соответствующего layer 2 multicast entry входа. Если layer 2 multicast entry вход не имеет порта, то layer 2 multicast entry вход тоже удаляется.

Быстрый выход (Fast leave) обычно используется когда один порт соответствует одному хосту. Если одному порту соответствует более одного хоста, то может быть сконфигурирован timeout выход для обеспечения непрерывности потока передаваемых multicast данных в сети.

14.1.5 IGMP Interrogator (Следование IGMP)

Оборудование multicast с поддержкой 3-layer работает в сети как генератор запросов IGMP queries. Оборудование multicast с поддержкой layer 2 только принимает сообщения IGMP и может создавать и обслуживать таблицу forwarding table для реализации layer 2 multicast. В сетях без оборудования layer 3 multicast, оборудование layer 3 multicast не может создавать запросы IGMP queries. Для мониторинга сообщений IGMP оборудование layer-2 multicast должно быть сконфигурировано для выполнения запросов IGMP query. Оборудование layer-2 multicast не только выполняет запросы IGMP query, но и осуществляет мониторинг сообщений IGMP messages для формирования и обслуживания таблицы forwarding table для реализации layer 2 multicast.

Оборудование IGMP layer-2 играет роль роутера для следования запросов IGMP query, регулярно отправляя сообщения IGMP query, принимая и обслуживая сообщения IGMP report от пользователей, создает таблицу forwarding table layer-2 multicast. Релевантные параметры сообщений query отправляемые IGMP могут быть настроены пользователем.

- Включить функцию query для выбранного VLAN.
- Определить версию IGMP query;
- Определить версию IGMP используемую для отправленных сообщений query (V1, V2 или V3);
- Задать IP адрес источника query;

- Задать IP адрес источника сообщений query отправленных query оборудованием;
- Задать временной интервал отправки сообщений query;
- Задать временной интервал отправки сообщений global query.

14.1.6 IGMP snooping Multicast filtering (Фильтрация IGMP snooping)

Функция IGMP snooping может контролировать обслуживание multicast, нагрузку и эффективно предотвращать прохождение нелегальных потоков данных multicast. Применение правил для интерфейса основанных на настройках правил фильтрации multicast, позволяет вводить ограничения для участников определенных групп.

14.2 IGMP SNOOPING configuration (GMRP конфигурация)

14.2.1 IGMP snooping default configuration (IGMP начальная конфигурация)

- IGMP snooping – функция отключена по умолчанию;
- Таблица multicast forwarding table в режиме unregistered forwarding mode;
- Функция Fast leave – отключена по умолчанию;
- Значение временного интервала fast leave timeout - 300 сек;
- Значение age time группы multicast порта report - 400 сек;
- Значение age time группы multicast порта query - 300 сек.

14.2.2 Open and close IGMP snooping (Включение и отключение)

Протокол IGMP snooping может быть включен в режиме global или separate; IGMP snooping может быть включен или отключен для VLAN только в режиме on global.

Включение global IGMP SNOOPING

```
Switch#configure terminal
```

```
Switch(config)#ip igmp snooping
```

Включение IGMP snooping для VLAN

```
Switch#configure terminal
```

```
Switch(config)#ip igmp snooping vlan <vlan-id>
```

Отключение GMP SNOOPING

```
Switch#configure terminal
```

```
Switch(config)#no ip igmp snooping
```

Отключение VLAN GMP SNOOPING

```
Switch#configure terminal
```

```
Switch(config)#no ip igmp snooping vlan <vlan-id>
```

14.2.3 Configure lifetime (Установка lifetime)

Установка lifetime multicast группы (миллисекунды)

```
Switch#configure terminal
```

```
Switch(config)#ip igmp snooping group-membership-timeout <interval>  
vlan <vlan-id>
```

Установка lifetime query группы (миллисекунды)

```
Switch#configure terminal
```

```
Switch(config)#ip igmp snooping query-membership-timeout <interval>  
vlan <vlan-id>
```

14.2.4 Configure fast-leave (Настройка fast-leave)

Включение fast leave для VLAN

```
Switch#configure terminal
```

```
Switch(config)#ip igmp snooping fast-leave vlan <vlan-id>
```

Отключение fast-leave

```
Switch#configure terminal
```

```
Switch(config)#no ip igmp snooping fast-leave vlan <vlan-id>
```

Установка wait time для fast leave

```
Switch#configure terminal
```

```
Switch(config)# ip igmp snooping fast-leave-timeout <interval> vlan <vlan-  
id>
```

Восстановление wait time для fast leave (значение по умолчанию)
Switch#configure terminal
Switch(config)#no ip igmp snooping fast-leave-timeout vlan <vlan-id>

14.2.5 Configure MROUTER (Настройка MROUTER)

Конфигурация порта static query port
Switch#configure terminal
Switch#interface ge1/6
Switch(config-ge1/6)#ip igmp snooping mrouter vlan [vlan-id]

14.2.6 Configure IGMP snooping query port (Настройка порта query)

Конфигурация порта static query port
Switch#configure terminal
Switch(config)#interface ge1/6
Switch(config-ge1/6)#ip igmp snooping mrouter vlan [vlan-id]

14.2.7 Configure IGMP snooping query function (Настройка query)

Включение функции IGMP snooping query для vlan1
Switch#configure terminal
Switch(config)#ip igmp sno
Switch(config)#ip igmp snooping querier vlan 1

14.2.8 Configure IGMP snooping multicast filtering (Настройка фильтрации)

Настройка порта GE1/1 на фильтрацию адреса 235.0.0.1
Switch#configure terminal
Switch(config)#ip igmp snooping filter-rule 1 deny 235.0.0.1
Switch(config)#interface ge1/1
Switch(config-ge1/1)#ip igmp snooping filter-group 1

14.2.9 Display information (Просмотр информации)

Просмотр конфигурации IGMP snooping

```
Switch#show ip igmp snooping
```

Просмотр конфигурации VLAN

```
Switch#show ip igmp snooping vlan <vlan-id>
```

Просмотр информации aging report multicast группы

```
Switch#show ip igmp snooping age-table group-membership
```

Просмотр информации aging query

```
Switch#show ip igmp snooping age-table query-membership
```

Просмотр информации forwarding multicast группы

```
Switch#show ip igmp snooping forwarding-table
```

Просмотр информации mrouter

```
Switch#show ip igmp snooping mrouter
```

Просмотр текущей информации системы system (включая конфигурацию IGMP snooping)

```
Switch#show running-config
```

14.2.10 IGMP SNOOPING Configuration example (Пример конфигурации)

Включение функции IGMP snooping на коммутаторе, пользователи ПК1, ПК2 и ПК3 включены в multicast группу.



```
Switch#config t
Switch(config)#ip igmp snooping
Switch(config)#vlan database
Switch(config-vlan)#vlan 200
Switch(config-vlan)#exit
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode access
Switch(config-ge1/1)#switchport access vlan 200
Switch(config)#interface ge1/2
Switch(config-ge1/2)#switchport mode access
Switch(config-ge1/2)#switchport access vlan 200
Switch(config)#interface ge1/3
Switch(config-ge1/3)#switchport mode access
Switch(config-ge1/3)#switchport access vlan 200
Switch(config)#ip igmp snooping vlan 200
Switch(config)#ip igmp snooping group-membership-timeout 60000 vlan
200
```

15. MVR Configuration (Конфигурация MVR)

15.1 MVR introduction (Введение)

Multicast VLAN registration (MVR) применяется для обслуживания multicast потоков в сетях провайдером. MVR позволяет подписчикам подключаться к портам и пользоваться или отключать multicast потоки в multicast VLAN, делиться потоком данных одного multicast VLAN с другими VLANs. MVR имеет два назначения: эффективно и безопасно направлять потоки между VLANs; поддерживать динамическое включение и исключение из multicast групп.

Режимы работы MVR аналогичны IGMP snooping. Одновременно могут использоваться две функции. MVR поддерживает включение и исключение из multicast группы, а процессом включения и исключения других групп управляет IGMP snooping. Разница заключается в том что IGMP snooping может направлять потоки multicast только в один VLAN, а MVR нескольким разным VLAN.

15.2 MVR Configure (Конфигурация MVR)

Таблица команд установок MVR

Команда	Описание	CLI режим
mvr (enable disable)	Включение/отключение global MVR	Global configuration mode
no mvr	Сброс всех настроек MVR	Global configuration mode
mvr group A.B.C.D	Создание multicast IP адреса	Global configuration mode
no mvr group A.B.C.D	Удаление multicast IP адреса	Global configuration mode
mvr group A.B.C.D <1-256>	Конфигурация multicast IP адреса и группы адресов MVR	Global configuration mode
mvr vlan <1-4094>	Выбор VLAN для получения данных multicast	Global configuration mode
no mvr vlan	Восстановление vlan1 для получения данных multicast по умолчанию	Global configuration mode
mvr-interface (enable disable)	Включение MVR интерфейса	Interface configuration mode
show mvr	Просмотр конфигурации MVR	Privileged mode

15.3 MVR Configuration example (Пример конфигурации)

Топология подключения приведена на рисунке ниже. Пользователи ПК1 и ПК2 принадлежат VLAN10 и VLAN20 соответственно. Пользователи ПК1 и ПК2 смотрят один и тот же контент. Диапазон контента 225.1.1.1 ~ 225.1.1.64, MVR VLAN – 100.



Требуется сконфигурировать VLAN, запустить global IGMP snooping, сконфигурировать MVR VLAN, диапазон MVR programr группы и запустить MVR global:

```

Switch#configure terminal
Switch(config)#ip igmp snooping
Switch(config)# mvr enable
Switch(config)#mvr vlan 100
Switch(config)#mvr group 225.1.1.1 64
Switch#
  
```

Конфигурация портов пользователей GE1/1, GE1/2, порт uplink GE1/3:

```

Switch#configure terminal
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode hybrid
Switch(config-ge1/1)# switchport hybrid native vlan 10
Switch(config-ge1/1)#switchport hybrid allowed vlan add 100 egress-
tagged disable
Switch(config-ge1/1)#mvr enable
Switch(config-ge1/1)#

Switch#configure terminal
Switch(config)#interface ge1/2
Switch(config-ge1/2)#switchport mode hybrid
Switch(config-ge1/2)# switchport hybrid native vlan 20
Switch(config-ge1/2)#switchport hybrid allowed vlan add 100 egress-
tagged disable
Switch(config-ge1/2)#mvr enable
Switch(config-ge1/2)#
  
```

```
Switch#configure terminal
Switch(config)#interface ge1/3
Switch(config-ge1/3)# switchport access vlan 100
Switch(config-ge1/3)#
```

16. DHCP SNOOPING Configure (DHCP SNOOPING)

В сетях с динамическим доступом хост получает IP адрес и другие параметры сети от DHCP сервера. DHCP snooping это протокол защиты от атак путем просмотра сообщений DHCP, динамической привязки MAC адреса клиента и IP адреса назначенного DHCP сервером клиенту и фильтрации сообщений ARP (атаки) коммутатором.

Коммутатор поддерживающий функцию DHCP snooping, может эффективно противостоять ARP атакам. DHCP snooping просматривает DHCP сообщения в сети и привязывает порт к порту информацию ARP.

Четыре физических порта соединенных с DHCP сервером могут быть сконфигурированы для предотвращения взаимодействия с неизвестным сервером или внешней сетью. Глава содержит следующие разделы:

- DHCP SNOOPING introduction (введение)
- DHCP SNOOPING configuration (конфигурация)
- DHCP SNOOPING configuration example (пример конфигурации)

16.1 DHCP SNOOPING introduction (Введение)

ARP протокол может создавать уязвимости для безопасности сети по причине простой реализации «trust» механизма. Когда хост получает сообщения ARP атаки содержащие ложную MAC информацию, то он напрямую перезаписывает локальную таблицу ARP кэша без ограничений, что приводит к направлению потока данных к атакующей стороне. Таким образом, привязка ARP информации к порту реализуется на коммутаторах уровня layer-2 switch, которые могут эффективно отфильтровывать сообщения ARP атаки и предотвратить попадание таких сообщений хосту. Если неизвестный DHCP сервер

подключится к сети, то это приведет к нарушениям порядка присвоения IP адресов. Протокол DHCP snooping обеспечивает физический порт связью с сервером. Невыделенный физический порт не может перенаправлять пакеты DHCP протокола от DHCP сервера, что снижает вероятность подключения неизвестного сервера к сети.

16.1.1 DHCP SNOOPING Processing (Процесс DHCP SNOOPING)

Протокол DHCP snooping только просматривает сообщения dhcprequest, dhcpack и dhcprelease и не получает другие типы DHCP сообщений, связывает IP и MAC адреса согласно этим сообщениям.

Основной global DHCP snooping коммутатор отвечает за прием DHCP сообщений, IP сообщений от UDP портов 67 и 68.

16.1.2 DHCP SNOOPING Binding table (Таблица привязки)

Записи таблицы привязки DHCP snooping индексируются MAC адресами (включая тип записи), IP и MAC адресом, информацией об интерфейсе, времени задержки и использования. Существует два типа: req и ACK. Вход req указывает что было получено сообщение dhcprequest, а не сообщение dhcpack. В это время запускается таймер отсчета времени задержки, (временной интервал составляет 10 сек по умолчанию). Если в течение 10 сек сообщение dhcpack не получено, тип req таблицы привязки будет оменен. Тип ACK указывает что получено сообщение dhcpack, записанный IP адрес назначен сервером. В это время запускается таймер отсчета времени использования (lease timer). Временной интервал (lease value) определяется DHCP сервером и содержится в сообщении dhcpack. После обновления соединения таймер перезапускается. По истечении временного интервала lease, вход (entry) таблицы привязки отменяется. В информации об интерфейсах записывается интерфейс к которому подключен клиент, таким образом интерфейс сообщает о привязке IP и MAC адресов.

При получении сообщения dhcprequest, создается таблица привязки с типом req, записывается IP адрес, MAC адрес, информация интерфейса и запускается таймер задержки (по умолчанию 10 сек).

При получении сообщения dhcprequest, когда таблица привязки

с типом req уже существует, обновляется entry и перезапускается таймер задержки.

При получении сообщения dhcprequest, когда существует таблица привязки с типом ACK, записывается информация об интерфейсе.

При получении сообщения dhcpack, таблица привязки с типом req, записывается IP адрес назначенный сервером в сообщении dhcpack, останавливается таймер задержки и запускается таймер использования (lease timer).

Если при получении сообщения dhcpack таблица привязки имеет тип req, то сообщение сбрасывается.

Если при получении сообщения dhcpack, таблица привязки имеет тип ACK и интерфейс был изменен, отменяется и обновляется тип таблицы привязки интерфейса.

Если интерфейс не изменяется, а IP адрес назначенный сервером изменяется, отменяется и обновляется тип таблицы привязки.

Если интерфейс и IP адрес назначенный сервером не изменяется, то происходит обновление и перезапуск таймера использования (lease timer).

По истечении времени задержки (delay timer), тип req таблицы привязки отменяется.

По истечении времени использования (lease timer), тип ACK таблицы привязки отменяется.

16.1.3 Physical port of the linked server (Физический порт сервера)

Протокол DHCP snooping определяет физический порт сервера для связи с ним. Сообщения DHCP принимаются только от этого выделенного порта. Если в сети присутствует несколько DHCP серверов, то предложения назначения IP адресов от неопределенных портов серверов будут отфильтрованы. Определенный физический порт отвечает за упорядоченное распределение IP адресов в сети для того чтобы предотвратить распространение незапланированного пула адресов от неизвестных серверов, что может создать проблемы клиентам для подключения к сети.

16.2 DHCP SNOOPING Configuration (Конфигурация)

16.2.1 DHCP SNOOPING Default configuration (Настройки по умолчанию)

Протокол DHCP SNOOPING отключен по умолчанию.

Время задержки (delay timer, при типе req таблицы привязки) 10 сек по умолчанию.

16.2.2 Global on and off (Глобальное включение и отключение)

Протокол DHCP snooping для интерфейса может быть включен только после включения Global DHCP snooping. Протокол DHCP snooping для всех интерфейсов может быть отключен только после отключения Global DHCP snooping.

Команды включения global DHCP snooping:

```
Switch#configure terminal
Switch(config)#ip dhcp snooping [IF_LIST]
```

Параметром является список физических портов определенного DHCP сервера, может быть выбрано до четырех портов. Порты в списке разделяется символом «,»(GE1/1, GE1/4, GE1/5).

Команды отключения global DHCP snooping:

```
Switch#configure terminal
Switch(config)#no ip dhcp snooping
```

16.2.3 Interface on and off (Включение и отключение для интерфейса)

Включение на порту DHCP SNOOPING

```
Switch#configure terminal
Switch(config)#interface ge1/1
Switch(config-ge1/1)#dhcp snooping
```

Отключение на порту DHCP SNOOPING

```
Switch#configure terminal  
Switch(config)#interface ge1/1  
Switch(config-ge1/1)#no dhcp snooping
```

16.2.4 Interface on and off OPTION82 (Включение и отключение OPTION82)

Включение на порту DHCP SNOOPING OPTION82

```
Switch#configure terminal  
Switch(config)#interface ge1/1  
Switch(config-ge1/1)#dhcp snooping option82
```

Отключение на порту DHCP SNOOPING OPTION82

```
Switch#configure terminal  
Switch(config)#interface ge1/1  
Switch(config-ge1/1)#no dhcp snooping option82
```

Конфигурация на порту circuit ID DHCP snooping option82

```
Switch#configure terminal  
Switch(config)#interface ge1/1  
Switch(config-ge1/1)# dhcp snooping option82 circuit-id vlan111
```

Отключение на порту circuit ID DHCP snooping option82

```
Switch#configure terminal  
Switch(config)#interface ge1/1  
Switch(config-ge1/1)# no dhcp snooping option82 circuit-id
```

16.2.5 Display information (Просмотр информации)

Просмотр конфигурации DHCP snooping

```
Switch#show dhcp snooping
```

Просмотр таблицы привязки DHCP snooping

```
Switch#show dhcp snooping binding-table
```

Просмотр текущей конфигурации системы включая DHCP snooping

```
Switch#show running-config
```

16.3 DHCP SNOOPING Configuration example (Пример конфигурации)

Включить DHCP snooping на коммутаторе уровня layer-2. Пользователи ПК1, ПК2 и ПК3 получают динамические IP адреса и параметры сети от DHCP сервера. Запуск функции DHCP snooping option82 на интерфейсах пользователей ПК1, ПК2 и ПК3, circuit ID значение AAA, информация ARP динамически связана с интерфейсами.



```
Switch#configure terminal
Switch(config)#ip dhcp snooping ge1/5
Switch(config)#interface ge1/1
Switch(config-ge1/1)#dhcp snooping
Switch(config-ge1/1)#dhcp snooping option82
Switch(config-ge1/1)#dhcp snooping option82 circuit-id aaa
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#dhcp snooping
Switch(config-ge1/2)#dhcp snooping option82
Switch(config-ge1/2)#dhcp snooping option82 circuit-id aaa
Switch(config-ge1/2)#interface ge1/3
Switch(config-ge1/3)#dhcp snooping
Switch(config-ge1/3)#dhcp snooping option82
Switch(config-ge1/3)#dhcp snooping option82 circuit-id aaa
Switch(config-ge1/3)#end
```

Просмотр конфигурации DHCP

```
Switch#show dhcp snooping
```

DHCP Snooping enabled globally

DHCP Server interface: ge1/5

Enable interface: ge1/1 ge1/2 ge1/3

Option 82 interface: ge1/1(Circuit ID: aaa) ge1/2(Circuit ID: aaa)

ge1/3(Circuit ID: aaa)

Switch#

```
Switch#show dhcp snooping binding-table
```

IP	MAC	FLAG	PORT	LEASE
192.168.1.100	00:11:5b:34:42:ad	ACK	ge1/1	23:59:58
192.168.1.101	00:11:64:52:13:5d	ACK	ge1/2	23:50:01
192.168.1.102	00:11:80:4d:a2:46	ACK	ge1/3	20:34:45

Возможные проблемы при настройке DHCP Snooping могут быть вызваны следующими причинами:

- Превышение ресурсов CFP.
- Global DHCP snooping не будет работать, если на интерфейсах включена функция фильтрации ACL.
- Global DHCP snooping не будет работать, если на интерфейсах включена функция привязки IP и MAC адресов.
- Функция фильтрации ACL включена для текущего интерфейса.
- Функция 802.1x anti ARP Spoofing включена для текущего интерфейса.
- Интерфейс имеет уровень 3-layer или является trunk интерфейсом.

17. DHCP CLIENT Configuration (DHCP Клиент)

DHCP (Dynamic Host Configuration Protocol) основан на режиме работы клиент / сервер. DHCP клиент делает возможным получение адреса и шлюза интерфейсом коммутатора уровня layer-3 через DHCP сервер.

17.1 DHCP CLIENT Configuration (Конфигурация клиента)

Включение функции DHCP на интерфейсе vlan1

```
switch(config)#  
Switch(config)#interface vlan1  
Switch(config-vlan1)#dhcp client enable
```

Обновление IP адреса интерфейса vlan1

```
switch(config)#  
Switch(config)#interface vlan1  
Switch(config-vlan1)#dhcp client renew
```

Сброс IP адреса интерфейса vlan1

```
switch(config)#  
Switch(config)#interface vlan1  
Switch(config-vlan1)#dhcp client release
```

18. DHCP RELAY Configuration (DHCP RELAY)

18.1 DHCP RELAY introduction (Введение)

DHCP (Dynamic Host Configuration Protocol) является расширенной версией протокола BOOTP. Протокол динамически конфигурирует сетевое окружение хостов, которое подразделяется на сервер и клиента. Сервер управляет IP данными в сети, обрабатывает запросы клиентов, динамически конфигурирует TCP / IP окружение клиента. Для работы протокола DHCP в сети должен присутствовать один сервер. Сервер может принимать запросы DHCP от хостов находящихся в сети и взаимодействоватьс TCP / IP параметрами. Существует два режима работы: автоматический и динамический. В автоматическом режиме клиент один раз получает IP адрес, который далее использует постоянно. В динамическом режиме клиент получает IP адрес с ограниченным сроком использования. По истечении срока использования (lease) IP адрес должен быть сброшен, также возможно продление срока использования или получение нового IP.

Динамический режим эффективно решает проблему дефицита и актуализации IP-адресов.

Процесс работы протокола DHCP:

Если клиент первой раз регистрируется в сети и не имеет данных IP, он генерирует сообщение `discover message` с адресом источника 0.0.0 и адресом назначения 255.255.255.255. Если сервер не отвечает, то генерируется четыре сообщения `discover request` через определенные интервалы времени.

Когда сервер получает `discover` сообщение, он выбирает `idle IP` для отправки ответного сообщения `offer message` клиенту.

Если в сети несколько серверов, клиент получит несколько ответных сообщений `offer message`. Обычно клиент выбирает первое полученное ответное сообщение и отправляет уведомление остальным серверам о получении IP адреса.

Если клиент обнаруживает что IP адрес используется протоколом ARP, он отправляет сообщение об отказе (`decline message`) серверу и запускает процесс `discover` повторно.

После получения запроса (`request message`), сервер отправляет сообщение `ACK message` клиенту для подтверждения возможности использования (`lease`).

В течение времени использования (`lease`) клиенту обычно не требуется повторно проходить процесс `discover`. Перед тем как истечет период времени пользования `lease`, следует отправить запрос `request` серверу для обновления IP. Сервер сделает попытку разрешить клиенту использовать оригинальный IP. В случае отсутствия ограничений, сервер отправит сообщение `ACK` для подтверждения. Если IP адрес уже используется другим клиентом, сервер отправит сообщение `NACK` и отклонит запрос на обновление.

Клиент может отправить сообщение о сбросе (`release message`) для прерывания использования до истечения времени `lease`.

При запуске рабочая станция отправляет запрос, который будет повторен по истечении половины времени пользования `lease`. При отсутствии подтверждения, использование IP может быть продолжено, запрос будет повторен по истечении 3/4 времени пользования `lease`. При отсутствии подтверждения, IP не может более использоваться.

Сообщение `discover message` создается в режиме `broadcast mode` и распространяется только в пределах сегмента сети, роутер не передает это сообщение за пределы сегмента сети. Когда сервер и

клиент находятся в разных сегментах сети, и клиент не обладает информацией об IP установках сетевого окружения и места нахождения роутера, сообщение discover message не может быть передано серверу. Для решения этой проблемы существует функция DHCP relay, которая дает возможность роутеру транслировать сообщения протокола DHCP, что делает возможным работу протокола в разных сегментах сети.

18.2 DHCP RELAY Configuration (Конфигурация DHCP RELAY)

Функция DHCP relay связана с интерфейсом, который выполняет передачу сообщений протокола DHCP между сегментами сети и отвечает за релевантность конфигурации в режиме interface mode.

18.2.1 DHCP relay of interface (Включение DHCP relay)

Команда включения DHCP relay для интерфейса в режиме interface configuration mode:

```
DHCP relay < ip-address-1> [ip-address-2]
```

Команда отключения DHCP relay для интерфейса в режиме interface configuration mode:

```
no DHCP relay
```

По умолчанию протокол DHCP relay отключен.

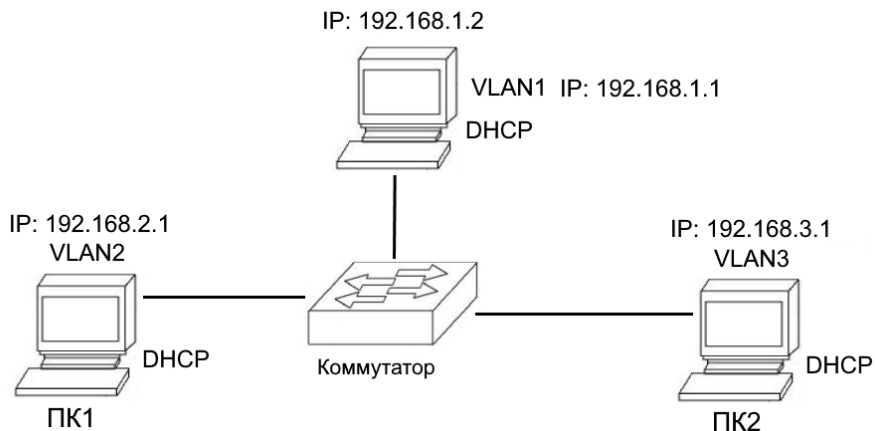
18.2.2 Display information (Просмотр информации)

Просмотр информации о конфигурации протокола DHCP relay

```
Switch#show dhcp relay
```

18.2.3 DHCP RELAY configuration example (Пример конфигурации)

Коммутатор 1 направляет запросы (DHCP relay request) и отвечает коммутатору 2. Пользователи ПК1 и ПК2 получают IP адреса доступа через DHCP серверы из разных сегментах сети.



Команды конфигурации:

```
Switch>en
Switch# configure terminal
Switch(config)# vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#vlan 3
Switch(config-vlan)#exit
Switch(config)#ip interface vlan 2
Switch(config)#ip interface vlan 3
Switch(config)#int ge1/2
Switch(config-ge1/2)#sw access vlan 2
Switch(config-ge1/2)#int ge1/3
Switch(config-ge1/3)#sw access vlan 3
Switch(config-ge1/3)# interface vlan2
Switch(config-vlan2)#ip address 192.168.2.1/24
Switch(config-vlan2)#dhcp relay 192.168.1.2
Switch(config-vlan2)#interface vlan3
Switch(config-vlan3)#ip address 192.168.3.1/24
Switch(config-vlan3)#dhcp relay 192.168.1.2
```

Просмотр текущей конфигурации:

```
show running-config
```

Просмотр конфигурации DHCP relay:

```
show dhcp relay
```

19. DHCP SERVER Configuration (DHCP Сервер)

19.1 DHCP SERVER introduction (Введение)

DHCP (Dynamic Host Configuration Protocol) является расширенной версией протокола BOOTP. Протокол динамически конфигурирует сетевое окружение хостов, которое подразделяется на сервер и клиента. Сервер управляет IP данными в сети, обрабатывает запросы клиентов, динамически конфигурирует TCP / IP окружение клиента. Для работы протокола DHCP в сети должен присутствовать один сервер. Сервер может принимать запросы DHCP от хостов находящихся в сети и взаимодействовать с TCP / IP параметрами. Существует два режима работы: автоматический и динамический. В автоматическом режиме клиент один раз получает IP адрес, который далее использует постоянно. В динамическом режиме клиент получает IP адрес с ограниченным сроком использования. По истечении срока использования (lease) IP адрес должен быть сброшен, также возможно продление срока использования или получение нового IP. Динамический режим эффективно решает проблему дефицита и актуализации IP адресов.

Подробно процесс работы протокола описан в п.18.1 настоящего Руководства.

19.2 DHCP SERVER configuration (Конфигурация DHCP SERVER)

Настройки функции DHCP server производятся в режимах global mode, interface mode, address pool mode, включая команду запуска, настройку пула адресов, настройки global setting и т.п.

Настройка DHCP server включает в себя следующие пункты:

- Включение функции global DHCP server
- Запуск получения сообщений DHCP server интерфейсом
- Настройка пула адресов (address pool)
- Настройка диапазона пула адресов (address pool range)
- Настройка пула адресов маски подсети (address pool subnet mask)
- Настройка времени пользования пулом адресов (address pool lease)

- Настройка шлюза пула адресов по умолчанию (address pool default gateway)
- Настройка пула адресов DNS сервера (address pool DNS server)
- Настройка пула адресов для исключения в ручную (address pool to exclude addresses manually)

19.2.1 Start global DHCP server (Включение global DHCP server)

Команда включения в режиме global configuration mode:

IP DHCP server

Команда отключения в режиме global configuration mode:

no IP DHCP server

По умолчанию функция DHCP server отключена.

19.2.2 Start interface to receive DHCP server message (Включение приема сообщений DHCP server интерфейсом)

Команда включения в режиме interface configuration mode:

DHCP server listen

Команда отключения в режиме interface configuration mode:

no DHCP server listen

По умолчанию функция приема сообщений DHCP server отключена.

19.2.3 Configure address pool (Настройка пула адресов)

Команда для создания пула в режиме global configuration mode:

DHCP server pool <pool name> (ввести имя пула)

Команда для удаления пула в режиме global configuration mode:

no DHCP server pool <pool name> (ввести имя пула)

Параметр <pool name> имя пула нужен для идентификации пула, может содержать максимум 16 знаков.

По умолчанию пула адресов не существует. Для создания пула в

режиме address pool configuration mode введите имя пула, *не используйте актуальные адреса*.

19.2.4 Configure address pool range (Настройка диапазона пула адресов)

Команда ввода адресов начала и конца диапазона в режиме address pool configuration mode:

range <low address> <high address>

Команда удаления диапазона в режиме address pool configuration mode:
no range

Параметр <low address> - начальный адрес диапазона, разделяется точками в десятичном формате.

Параметр <High address> - конечный адрес диапазона, разделяется точками в десятичном формате.

По умолчанию диапазон пула адресов не существует. При вводе диапазона пула адресов, каждый динамически назначаемый адрес будет принадлежать данному диапазону.

19.2.5 Configure address pool subnet mask (Настройка пула адресов маски подсети)

Команда ввода адресов в режиме address pool configuration mode:
subnet mask <address>

Параметр <address> - адрес маски подсети, разделяется точками в десятичном формате, длина маски может варьироваться.

По умолчанию адрес: 255.255.255.0. длина 24 бита.

19.2.6 Configure address pool lease (Настройка времени использования пула адресов)

Команда конфигурации lease в режиме address pool configuration mode:
leave [<days> <hours> <minutes> |infinite]

Параметры <days> количество дней (диапазон 0 – 999), <Hours> количество часов (диапазон 0 – 23), <Minutes> количество минут (диапазон 0 – 59). <Infinite> - время использования lease не ограничено. По умолчанию значение lease - 8 дней.

19.2.7 Configure address pool default gateway (Настройка шлюза пула адресов)

Команда конфигурации адреса default gateway по умолчанию в режиме address pool configuration mode:

```
default router <IP address>
```

Команда Default route отменяет адрес пула default gateway.

Параметр <IP address> IP адрес шлюза по умолчанию, разделяется точками в десятичном формате, должен находиться в пределах пула адресов того же сегмента подсети.

По умолчанию адрес шлюза default gateway не настроен.

19.2.8 Configure address pool DNS server (Настройка DNS сервера пула адресов)

Команда конфигурации адресов DNS сервера в режиме address pool configuration mode:

```
DNS server <IP address1> [IP address2]
```

Команда no DNS server отменяет пул адресов DNS сервера.

Параметры <IP address1> и <IP address2> IP адреса DNS серверов в десятичном формате. Может быть назначен один или два DNS сервера. При назначении двух серверов, назначение производится одной командой. При использовании двух команд сохраняется IP адрес DNS сервера введенный последней командой, адреса введенные предыдущей командой удаляются.

По умолчанию адреса DNS серверов не настроены.

19.2.9 Manual address exclude from address pool (Исключение адресов из пула)

Команда исключения адреса из пула адресов в режиме address pool configuration mode:

```
exclude address <IP address>
```

Команда исключения диапазона адресов из пула адресов:

```
exclude address <low address> <high address>
```

Команда восстановления исключенного адреса из пула адресов:

```
no exclude address < IP address >
```

Команда восстановления исключенного диапазона адресов из пула:

```
no exclude address <low address> <high address>
```

Команда восстановления всех исключенных адресов из пула адресов:

```
no exclude address
```

Параметр <IP address> исключенный IP адрес, разделяется точками в десятичном формате. Одна команда исключает один адрес из пула. <Low address> и <high address> начальный и конечный IP адреса исключаемого диапазона из пула в десятичном формате. Одна команда исключает несколько последовательных адресов из пула. При ручном исключении адресов, адреса удаляются из пула на неограниченное время. Исключение адресов из пула означает, что данные адреса не будут распределяться динамически.

По умолчанию исключенные адреса не назначены.

19.2.10 Configure option82 (Подключение функции option82)

Команда конфигурации circuit ID функции option82 в режиме address pool configuration mode:

```
option82 circuit ID <circuit ID>
```

Параметр <circuit ID> строка (string), максимальная длина 64.

19.2.11 Clear assigned address table entries (Очистка таблицы адресов)

Команда удаления из таблицы всех назначенных адресов в режиме Exec configuration mod:

```
clear DHCP server address {[IP address] | [all]}
```

19.2.12 Clear conflicting address table entries detected (Очистка таблицы конфликтующих адресов)

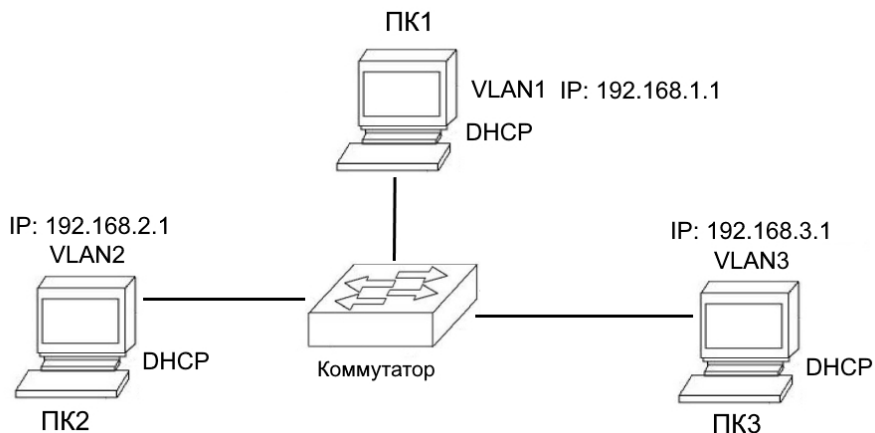
Конфликтующий адрес (адреса) автоматически определяются и заносятся в таблицу.

Команда удаления из таблицы всех конфликтующих адресов в режиме Exec configuration mod:

```
clear DHCP server address conflict {[IP address] | [all]}
```

19.2.13 DHCP SERVER configuration example (Пример конфигурации)

Создать конфигурацию пула адресов для клиентов подсетей vlan1, vlan2 и vlan3, таким образом чтобы коммутатор работал как DHCP сервер и мог назначать IP адреса соответствующих сегментов подсети клиентам.



Команды конфигурации:

```
Switch>en
Switch# configure terminal
Switch(config)#ip dhcp server
Switch(config)# vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#vlan 3
Switch(config-vlan)#exit
Switch(config)#ip interface vlan 2
Switch(config)#ip interface vlan 3
Switch(config)#int ge1/2
Switch(config-ge1/2)#sw access vlan 2
Switch(config-ge1/2)#int ge1/3
Switch(config-ge1/3)#sw access vlan 3
Switch(config-ge1/3)#interface vlan2
Switch(config-vlan2)#ip address 192.168.2.1/24
Switch(config-vlan2)#dhcp server listen
Switch(config-vlan2)#interface vlan3
Switch(config-vlan3)#ip address 192.168.3.1/24
Switch(config-vlan3)#dhcp server listen
Switch(config-vlan3)#interface vlan1
Switch(config-vlan1)#dhcp server listen
Switch(config)#dhcp server pool a
Switch(config-dhcp)#range 192.168.2.1 192.168.2.20
Switch(config-dhcp)#lease 2 0 0
Switch(config-dhcp)#default-router 192.168.2.1
Switch(config-dhcp)#dns-server 1.1.1.1 2.2.2.2
Switch(config-dhcp)#exclude-address 192.168.2.10
Switch(config-dhcp)#exit
Switch(config)# dhcp server pool b
Switch(config-dhcp)#range 192.168.3.2 192.168.3.20
Switch(config-dhcp)#default-router 192.168.3.1
Switch(config-dhcp)#exit
Switch(config)# dhcp server pool c
Switch(config-dhcp)#range 192.168.1.2 192.168.1.20
Switch(config-dhcp)#default-router 192.168.1.1
Switch(config-dhcp)#exit
```

Просмотр конфигурации DHCP сервера

Просмотр конфигурации:

`show running-config`

Просмотр глобальной конфигурации:

`show dhcp server`

Просмотр (одного) пула адресов:

`show dhcp server pool [pool-name]`

Просмотр соответствующей таблицы адресов:

`show dhcp server address`

20. Configure ACL (Конфигурация ACL)

Задача обеспечения безопасности при доступе к сети является одной из наиболее важных для администратора. Коммутатор поддерживает функцию фильтрации ACL filtering и может обеспечивать безопасность при подключении пользователей к сети. Путем конфигурации правил ACL rules, коммутатор осуществляет фильтрацию потока входящих данных согласно этим правилам и обеспечивает необходимую безопасность доступа. Глава содержит следующие разделы:

- ACL Introduction to resource library (Введение, библиотека)
- Introduction to ACL filtering (Фильтрация)
- ACL repository configuration (Конфигурация хранилища)
- ACL filtering configuration (Конфигурация фильтрации)
- ACL configuration example (Пример конфигурации)

20.1 ACL resource library Introduction (Введение библиотека)

Хранилище ACL (access list control) представляет собой набор из нескольких групп определенных правил доступа. Хранилище ACL не осуществляет управление и направление потоков данных, но разделяет конфликтующие правила. После обращения приложения к библиотеке хранилища ACL, приложение получает возможность управлять

потоками данных согласно установленным правилам ACL ресурсов. Функция ACL может быть применена на порту доступа, обслуживанию доступа, QoS и т.п.

Ресурсы библиотеки ACL включают в себя стандартные правила группы IP rule group (номер группы 1 ~ 991300 ~ 1999), расширенные правила IP rule group (номер группы 100 ~ 1992000 ~ 2699), группу IP MAC group (номер группы 700 ~ 799) и группу ARP group (номер группы 1100 ~ 1199). Приоритет при конфликте правил устанавливается автоматически для каждой группы правил. Когда пользователь устанавливает правило ACL, система установит приоритет согласно правилам приоритизации.

На практике, когда пакет данных проходит через порт, коммутатор сравнивает поля в каждом правиле с соответствующими полями в пакете данных. Когда одновременно применяется несколько правил, происходит действие соответствующее первому правилу. Соответствующее правило определяет передать или сбросить пакет. Таким образом, определяется соответствие поля в правиле и аналогичного поля в передаваемом пакете данных. Только тогда правило ACL определит дальнейшее действие - разрешить или отклонить операцию.

В коммутаторе правила в пределах одной группы сортируются автоматически. Автоматическая сортировка является комплексным действием. В процессе сортировки правила с широким диапазоном переносятся в конец, а правила с узким диапазоном переносятся в начало. Размер диапазона определяется правилами ограничения. Правила ограничений в основном отражены в адресе wildcard и количестве неадресных полей. Wildcard является bit строкой. Адрес IP имеет 4 байта, MAC адрес имеет 6 байт. Bits'1' означает отсутствие соответствия, bits'0' означает соответствие. Неадресные поля содержатся в wildcard и относятся к типу протокола IP и протоколу порта. Их длина унифицирована и соответствует длине поля в байтах, требуется только произвести подсчет количества полей. Большому числу битов wildcard соответствует значение '0'.

Преимущество необходимости автоматической сортировки правил показывает фильтрация доступа к порту. Например если требуется отклонить адрес 10.10.10.0/16 с которого обращаются к сегменту сети и дать доступ адресу 192.168.1.0/24, необходимо сконфигурировать следующие правила:

Access list 1 permit 192.168.1.0 0.0.255 - rule 1

Access list 1 deny 10.10.10.0 0.0.255.255 - Rule 2

Далее описываются правила rule 1 and rule 2.

Данные правила конфликтуют т.к. адрес правила rule 1 включен в адрес правила rule 2, т.е. одно правило разрешает, а второе отклоняет. Исходя из принципа фильтрации ACL, следует согласовать последовательность применения двух правил: сначала Rule 1 а затем rule 2. В этом случае коммутатор использует функцию автоматической сортировки в независимости от конфигурации заданной пользователем. Когда пересылается пакет данных от адреса 192.168.1.1, сначала применяется первое правило, далее второе правило. При совпадении с обоими правилами, производится действие в соответствии с первым правилом (forwarding). Если forwarding адрес 10.1, то только первый адрес сбрасывается.

Если автоматическая сортировка отключена, то пользователь может сконфигурировать другую последовательность применения правил сначала rule 2, затем rule 1:

Access list 1 deny 10.10.10.0 0.0.255.255 - Rule 2

Access list 1 permit 192.168.1.0 0.0.255 - rule 1

Т.к. первичное правило rule 2 содержит следующее правило rule 1, то может возникнуть ситуация когда пакеты соответствующие rule 1 и rule 2 будут постоянно обрабатываться по правилу rule 2.

В коммутаторе, биты '0.0.255.255' wildcard означают: bits'1' отсутствие соответствия, bits'0' - соответствие. Видно что биты wildcard правила rule 2 -'0.0.255.255', должны соответствовать двум байтам (16 bits), биты wildcard правила rule 1 - '0.0.255', и три байта (24 bits) должны соответствовать. Таким образом правило rule 1 начинает работать после правила rule 2. При расширенном IP, сортировка должна учитывать больше полей, таких как тип протокола IP, коммуникацию портов и т.п. При этом принцип сортировки остается неизменным, чем больше существует ограничений, тем меньше вес правила. Сортировка правил происходит в фоновом режиме, команды пользователя могут быть показаны в последовательности сконфигурированной пользователем.

ACL поддерживает фильтрацию следующих полей: IP источника, IP назначения, тип протокола IP (TCP, UDP, OSPF), порт источника (такой как 161) и порт назначения. Пользователь может задать разные правила для доступа исходя из различных конкретных

требований.

В коммутатор допускает установку множества правил одновременно: для фильтрации порта доступа, доступа для обслуживания или фильтрацию на двух портах.

20.2 ACL Filter Introduction (Введение ACL фильтр)

Функция ACL filtering применяется на входном порту коммутатора. Поток данных на входном порту коммутатора должен соответствовать определенным правилам для реализации функции фильтрации. Процесс ACL filtering выполняется со скоростью передачи данных в линии, что не оказывает негативного влияния на скорость передачи данных через коммутатор.

Когда функция ACL filtering не активирована на порту коммутатора, все потоки данных через порт вне зависимости от соответствия правилам пропускаются. Когда функция ACL filtering включена на порту коммутатора, то пропускаются только потоки данных, которые соответствуют заданным правилам. Не соответствующие правилам (отклоненные) потоки данных сбрасываются.

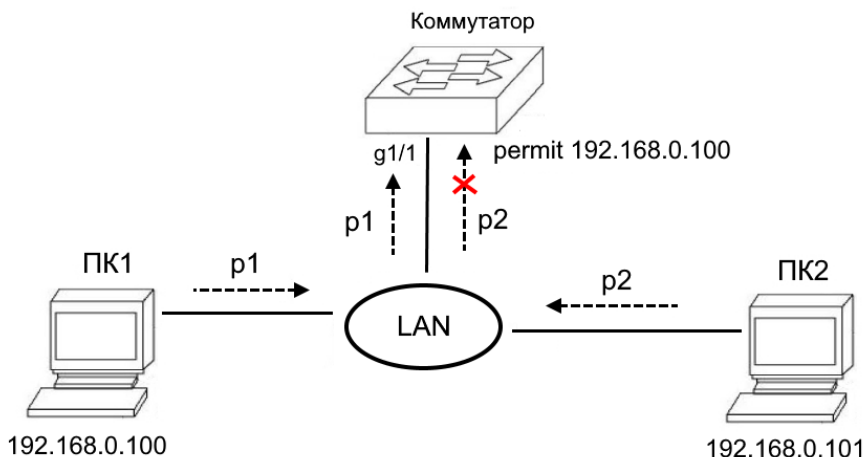
При конфигурации функции ACL filtering на порту, можно выбрать несколько групп правил ACL (rule groups). После выбора, группа правил импортируется в CFP порта. Если отсутствует правило разрешения или запрета для всех пакетов IP протокола в группе правил, то будет добавлено правило запрета для всех IP протоколов при записи в CFP. При изменении правил в хранилище ACL, правила записанные в CFP будут изменены автоматически.

Например, если группа правил содержит только одно правило: `access list 1 permit 192.168.1.0 0.0.255`. По умолчанию, правило отклоняющее все пакеты IP протокола будет скрыто. На практике, два правила будут импортированы в CFP порта. Во время процесса фильтрации потока, только потоки данных от источников с диапазоном адресов от 192.168.1.0 до 192.168.1.255 будут пропущены через порт, все остальные потоки данных не будут пропущены.

Например, существует два правила: `access list 1 deny 192.168.1.0 0.0.255` и `access list 1 permit any`. В этом случае, согласно правилу разрешено пропустить все пакеты данных IP протокола, скрытые правила отсутствуют. На практике, два правила будут

импортированы в CFP порта. В процессе фильтрации потока данных, только данные от источника с диапазоном адресов от 192.168.1.0 до 192.168.1.255 будут отфильтрованы, остальные потоки данных будут пропущены.

На рисунке ниже приведен пример ACL фильтрации. Выбрать группу 1 правил ACL для порта коммутатора g1/1. В группе существует только одно правило для доступа list 1 permission 10.10.10.100. К порту коммутатора g1/1 хотят получить доступ два пользователя из локальной сети. IP адреса пользователя 1 - 10.10.10.100 и пользователя 2 - 10.10.10.101. Только пользователь 1 может получить доступ из локальной сети к порту коммутатора g1/1, одновременно пользователь 2 не может получить доступ порту коммутатора g1/1. Поток данных P1 отправленный пользователем 1 может быть получен протом коммутатора g1/1, в то время как поток данных P2 отправленный пользователем 2 будет сброшен портом коммутатора g1/1.



В том случае если одна или несколько групп правил относятся к одному порту, они будут отсортированы автоматически если сортировка между двумя группами пересекается.

Если пользователь относится к группе правил и группа изменяется, то порт относящийся к этой группе автоматически отвечает за конфигурацию пользователя. Необходимость обновлять конфигурацию порта отсутствует.

20.3 ACL Repository configuration (ACL хранилище)

По умолчанию предустановленные правила в коммутаторе отсутствуют. Библиотека коммутатора поддерживает четыре типа правил ACL: стандартные (standard IP rules), расширенные (extended IP rules), IP MAC группа и ARP группа.

Standard IP rules: управляют потоками данных на основе IP адреса источника данных.

Команда: `access list < groupid > {deny | permit} < source >`

Groupid – список группы управления доступом, стандартные IP ACL поддерживают группы 1 - 99 или 1300 - 1999.

Deny / permit – запрет / разрешение на передачу пакетов данных.

Source - IP источника, имеет три способа ввода:

- 1) A.b.c.d wildcard – управление IP адресом из сегмента сети;
- 2) Any – эквивалент a.b.c.d 255.255.255.255;
- 3) Host a.b.c.d - эквивалент a.b.c.d 0.0.0

Wildcard – определяет связанные биты, для связанных значение '0', для несвязанных значение '1'.

Extended IP rule: standard IP rule с расширенными возможностями. Управление направлением пакетов данных может осуществляться на основе IP адреса источника, IP адреса получателя, типа IP протокола и сервисного порта.

Команда: `access list < groupid > {deny | permit} < protocol > < source > [EQ < srcport >] < destination > [destport] < TCP flag >`

Groupid: список группы управления доступом, поддерживает группы 100 - 199 или 2000 - 2699.

Deny / permit - запрет / разрешение на передачу пакетов данных.

Protocol – тип протокола IP layer (TCP, UDP и т.д.), допускается введение 6 (TCP). При отсутствии необходимости контролировать протоколы вводится IP или 0.

Source: IP источника, имеет три способа ввода:

- 1) A.b.c.d wildcard - управление IP адресом из сегмента сети;
- 2) Any - эквивалент a.b.c.d 255.255.255.255;
- 3) Host a.b.c.d - эквивалент a.b.c.d 0.0.0

Srcport – при использовании протоколов TCP или UDP, возможно управление портом источником пакетов данных на основе сервисного имени порта WWW или цифрового значения (например 80).

Destination: IP адрес получателя, имеет три способа ввода:

- 1) A.b.c.d wildcard - управление IP адресом из сегмента сети;
- 2) Any - эквивалент a.b.c.d 255.255.255.255;
- 3) Host a.b.c.d - эквивалент a.b.c.d 0.0.0

Destport: при использовании протоколов TCP or UDP, возможно управление портом получателем пакетов данных на основе сервисного имени порта WWW или цифрового значения (так же как при srcport).

TCP flag: при использовании протокола TCP возможно управление пакетами данных соответствующим полем TCP field. Опциональные параметры ACK, fin, PSH, RST, syn и urg.

IP MAC rules: Группа IP MAC group может управлять источником (source) получателем (destination) MAC адреса и source destination IP адресом IP пакетов.

Команда: access list < groupid > {deny | permit} < SRC MAC > IP < SRC IP > < DST IP >

Groupid: номер группы управления доступом, Extended IP ACL поддерживает группы 700 - 799.

Deny / permit - запрет / разрешение на передачу пакетов данных.

SRC MAC - MAC адрес источника, имеет три способа ввода:

- 1) НННН. НННН. НННН wildcard управляет MAC адресом из одного сегмента;
- 2) Any – эквивалент НН НННН. НННН FFFF. FFFF. FFFF.
- 3) Host a.b.c.d – эквивалент НН НННН. НННН 0000.0000.0000

SRC IP - IP адрес источника.

DST IP - IP адрес получателя, имеет три способа ввода:

- 1) A.b.c.d wildcard - управление IP адресом из сегмента сети;
- 2) Any - эквивалент a.b.c.d 255.255.255.255;
- 3) Host a.b.c.d - эквивалент a.b.c.d 0.0.0

ARP rules - Группа ARP group может управлять типом ARP пакетов, IP и Mac адресом отправителя.

Команда: access list < groupid > {deny | permit} ARP {ARP type} < sender MAC > < sender IP >

Groupid - номер группы управления доступом, extended IP ACL поддерживает группы 1100 - 1199.

Deny / permit - запрет / разрешение на передачу пакетов данных.

ARP type – относится к any | reply | request. Any не управляет ARP пакетами, reply отвечает за управление ARP пакетами, request пакет

запроса управляющий ARP пакетами.

Sender MAC – MAC адрес отправителя ARP пакетов, имеет три способа ввода.

- 1) НННН. НННН. Нhhh wildcard управляет MAC адресом из одного сегмента;
- 2) Any – эквивалент НН НННН. НННН FFFF. FFFF. FFFF.
- 3) Host a.b.c.d – эквивалент НН НННН. НННН 0000.0000.0000

Sender IP – IP адрес отправителя ARP пакетов, имеет три способа ввода.

- 1) A.b.c.d wildcard - управление IP адресом из сегмента сети;
- 2) Any - эквивалент a.b.c.d 255.255.255.255;
- 3) Host a.b.c.d - эквивалент a.b.c.d 0.0.0

Прочие команды:

show access-list [groupid]

Просмотр списка правил текущей конфигурации ACL для определенной группы groupid, при отсутствии данного параметра выводится весь список правил.

no access-list <groupid >

Удаление определенного списка правил, при вводе параметра Groupid удаляются все правила группы.

20.4 Time period based ACL (Период времени ACL)

Иногда пользователю требуется чтобы некоторые правила ACL действовали в течение определенного промежутка времени, т.н. фильтрация по временному периоду. Пользователь может задать один или несколько временных периодов, далее обращаться к этим периодам через их имя или соответствующее им правило. При ACL фильтрации основанной на временном периоде правило применяется в течение ограниченного промежутка времени.

Если для правила не задан временной период, то система позволит это сделать незамедлительно, но при этом правило не начнет действовать пока пользователь не закончит конфигурирование и не наступит совпадение выбранного временного периода с системным временем.

Существует два способа ввода временного периода:

- (1) Относительный (relative time): от определенного времени до определенного времени в определенный день;
- (2) Точный (absolute time): ввод в формате от определенного часа и минуты в определенный день, месяц и год до определенного часа и минуты в определенный день, месяц и год.

Таблица команд конфигурации:

Команда	Описание	Режим CLI
time-range WORD cycle-time from <0-23> <0-59> to <0-23> <0-59>	Ввод времени relative time в относительном формате.	Global configuration mode
time-range WORD cycle-time days from <0-6> to <0-6>	Ввод недели в относительном формате.	Global configuration mode
time-range WORD cycle-time from <0-23> <0-59> to <0-23> <0-59> days from <0-6> to <0-6>	Ввод времени и недели в относительном формате.	Global configuration mode
time-range WORD utter-time from <2000-2100> <1-12> <1-31> <0-23> <0-59> to <2000-2100> <1-12> <1-31> <0-23> <0-59>	Ввод точной даты и времени в абсолютном формате.	Global configuration mode
no time-range WORD cycle-time	Отмена всех относительных временных периодов.	Global configuration mode
no time-range WORD utter-time	Отмена всех временных периодов в абсолютном формате.	Global configuration mode
no time-range WORD	Отмена всех временных периодов в абсолютном и относительном форматах.	Global configuration mode
no time-range	Отмена всех временных периодов.	Global configuration mode
show time-range WORD cycle-time	Просмотр всех периодов в относительном формате.	Privileged mode
show time-range WORD utter-time	Просмотр всех периодов в абсолютном формате.	Privileged mode
show time-range WORD	Просмотр всех временных периодов.	Privileged mode

show time-range	Просмотр всех временных периодов.	Privileged mode
time-acl (<1-99> <100-199> <1300-1999> <2000-2699> <700-799> <1100-1199>) time-range WORD	Ввод правила ACL rule действующего во временной период.	Global configuration mode
no time-acl (<1-99> <100-199> <1300-1999> <2000-2699> <700-799> <1100-1199>) time-range (WORD)	Отмена правила ACL rule действующего во временной период.	Global configuration mode
show time-acl (<1-99> <100-199> <1300-1999> <2000-2699> <700-799> <1100-1199>) time-range	Просмотр правила ACL rule действующего во временной период.	Privileged mode
show time-acl all time-range	Просмотр всех временного периода действия правила ACL rule.	Privileged mode

Примечание:

- При конфигурации одновременно нескольких временных периодов (в относительном формате) и их взаимодействии действует принцип «или». В любой момент времени системное время и отсчет временных периодов активны;
- При конфигурации одновременно нескольких временных периодов (в абсолютном формате) и их взаимодействии действует принцип «или». В любой момент времени системное время и отсчет временных периодов активны;
- При конфигурации одновременно временных периодов в относительном и абсолютном формате, они связаны и активизируются при совпадении с системным временем;
- Может быть сконфигурировано до 256 временных периодов. Один временной период может быть сконфигурирован с 256 периодами в относительном формате и абсолютном формате. Одно правило ACL rule может относиться к 256 временных периодов. Временной период активируется, если связанное с ним правило ACL rule применено для интерфейса.

20.5 ACL Filter configuration (Конфигурация ACL фильтра)

По умолчанию функция ACL фильтрации отключена на портах.

Команда в режиме layer 2 interface configuration mode:

```
access-group <groupId>
```

GroupId - Номер ACL группы связанной с портом.

Причинами сбоя при конфигурации ACL фильтрации на порту могут быть:

- Слишком большое количество правил для ACL группы;
- Превышение аппаратных возможностей коммутатора.

Просмотр конфигурации ACL фильтрации на порту:

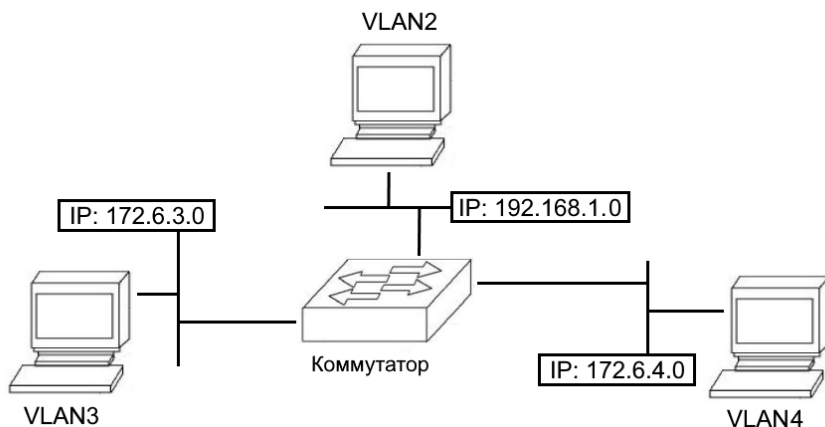
```
show access-group
```

Команда отключения функции ACL фильтрации на порту:

```
no acl-group <groupId>
```

20.6 ACL Filter configuration example (Пример конфигурации)

Коммутатор подключен к трем подсетям, назначен ACL, сетевой адрес заблокированного источника 192.168.1.0. Разрешено прохождение трафика других сетей. К порту коммутатора 1/1 подключен сегмент сети 192.168.1.0 (см. рисунок ниже).



Команды конфигурации:

```
Switch#config t
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#vlan 3
Switch(config)#interface ge1/1
Switch(config-ge1/1)# switchport mode access
Switch(config-ge1/1)#switchport access vlan 2
Switch(config)#interface vlan2
Switch(config-vlan2)#ip add 192.168.1.1/24
Switch(config)#interface ge1/2
Switch(config-ge1/2)#switchport mode access
Switch(config-ge1/2)#switchport access vlan 3
Switch(config)#interface vlan3
Switch(config-vlan2)#ip add 172.16.3.1/24
Switch(config)#access-list 10 deny 192.168.1.0 0.0.0.255
Switch(config)#access-list 10 permit any
Switch(config)#interface ge1/1
Switch(config-ge1/1)#access-group 10
Switch(config-ge1/1)#exit
Switch(config)#interface ge1/2
Switch(config-ge1/2)#access-group 10
```

21. TCP/IP Basic configuration (Конфигурация TCP/IP)

Для обеспечения коммуникации сетевых устройств коммутаторы уровня layer-2 с функцией сетевых настроек поддерживают управление основными настройками TCP / IP протокола. Глава содержит следующие разделы:

- Configure VLAN interface (Конфигурация VLAN интерфейса)
- Configure ARP (Конфигурация ARP)
- Configure static routing (Конфигурация статической маршрутизации)
- IP routing configuration example (Пример конфигурации)

21.1 Configure VLAN interface (VLAN интерфейс настройка)

В коммутаторе каждый интерфейс уровня layer-3 связан с VLAN, т.о. интерфейс уровня layer-3 также называется VLAN interface. Создание и удаление VLAN интерфейса производится вручную. Коммутатор допускает создание до 4094 VLANs, но только 32 подсети. Интерфейс подсети может быть создан или удален пользователем или удален в ручную или удален одновременно с VLAN подсети.

Каждый VLAN интерфейс имеет своё имя. Имя интерфейса VLAN содержится в строке "VLAN" следующей за номером VLAN ID. Например, имя интерфейса уровня layer 3 VLAN 1 - "vlan1", имя интерфейса уровня layer 3 VLAN 4094 – «vlan4094».

Также как и порты, VLAN интерфейсы имеют статусы управления (management) и соединения (link). Вначале статус управления (management) для VLAN интерфейса недоступен. Как только VLAN интерфейс создан, он автоматически получает статус management. Статус link VLAN интерфейса относится к портам принадлежащим VLAN корреспонденту интерфейса. До тех пор пока статус порта link активен, статус link VLAN интерфейса тоже активен. Если все порты VLAN не активны, статус link VLAN интерфейса тоже не активен.

В настройках VLAN интерфейса можно прописать IP адрес, префикс сегмента подсети (network prefix) соединенного с интерфейсом (который может быть преобразован в маску подсети). Вначале коммутатор имеет только один IP адрес относящийся к VLAN интерфейсу. Перед началом конфигурации IP адресов следует создать VLANs и добавить связанные с ними порты. По умолчанию коммутатор имеет интерфейс vlan1 с назначенным IP адресом 10.10.10.1/24, пользователь может изменить IP адрес интерфейса vlan1. Другие интерфейсы VLANs (отличные от vlan1) не имеют установленных по умолчанию IP адресов.

Таблица команд конфигурации IP адресов VLAN интерфейса:

Команда	Описание	Режим CLI
Ip interface vlan <2-4094>	Создать VLAN интерфейс	Global configuration mode

No Ip interface vlan <2-4094>	Удалить VLAN интерфейс	Global configuration mode
ip address <ip-prefix>	Установить IP адрес VLAN интерфейса. Параметры включают IP адрес интерфейса, префикс сегмента подсети. Если VLAN интерфейс уже имеет IP адрес, следует сначала удалить оригинальный адрес, а затем установить требуемый IP адрес в формате a.b.c.d/m.	Interface configuration mode
no ip address [ip-prefix]	Удалить IP адрес VLAN интерфейса. Если параметр определен, то он должен быть таким же как параметр заданный при установке, в противном случае команда будет некорректной. Формат параметра: a.b.c.d/m.	Interface configuration mode
show interface [if-name]	Просмотр информации VLAN интерфейса, в т.ч. IP адрес, MAC адрес, статус (management, link) и т.п. Параметром является имя VLAN интерфейса. Если параметр не определен, следует посмотреть информацию о всех портах и VLAN интерфейсах.	Normal mode, privileged mode
show running-config	Просмотр текущей конфигурации системы, в т.ч. конфигурации VLAN интерфейса.	privileged mode

Пример:

Сконфигурировать подсеть 193.1.1.0 на интерфейсе vlan3, префикс подсети - 24 (маска подсети 255.255.255.0), IP адрес интерфейса 193.1.1.1, посмотреть информацию интерфейса vlan3.

Команды:

```
switch(config)#interface vlan3
switch(config-vlan3)#ip address 193.1.1.1/24
switch(config-vlan3)#end
switch#show interface vlan3
```

21.2 Configure ARP (ARP конфигурация)

Протокол ARP (address resolution protocol) обеспечивает указание пути от IP адреса к соответствующему MAC адресу. Когда источник отправляет кадры данных Ethernet получателю из того же VLAN, получатель определяется согласно 48и битному Ethernet MAC адресу, получатель определяет необходимо ли принять пакет данных на основании MAC адреса.

Предполагается что хосты А и В двух соседних сегментов сети осуществляют коммуникацию через коммутатор. Перед отправкой данных хосту В, хост отправляет сообщение запрос ARP request интерфейсу коммутатора, который напрямую связан с хостом А. После получения ответа ARP response, хост отправляет пакеты данных интерфейсу. После получения пакета данных, коммутатор отсылает запрос ARP request message хосту В, после получения ответа ARP response message от хоста В, коммутатор отправляет пакет данных хосту В.

ARP кэш коммутатора, называется таблицей ARP table, в которой хранятся записи путей от IP адреса к MAC адресу в связанной сети. Каждый пункт таблицы ARP table имеет срок действия (по умолчанию 20 мин.). Если коммутатор не получает запросы ARP request или ответные сообщения (response message) от IP адреса в период действия записи, соответствующая запись удаляется из таблицы ARP table.

21.2.1 Configure static ARP (Конфигурация статического ARP)

Существует два типа ввода данных в таблицу ARP – статический (static ARP) и динамический (dynamic ARP). В случае Static ARP данные в таблицу ARP table вводятся или удаляются пользователем с помощью соответствующих команд, система не обновляет и не удаляет данные из таблицы автоматически. В случае Dynamic ARP данные в таблицу ARP table вносятся системой согласно полученным запросам ARP request или ответным пакетам. Система автоматически создает, удаляет и обновляет записи в реальном времени без участия пользователя, но пользователь имеет возможность удалять записи из таблицы вручную.

По умолчанию на коммутаторе отключен режим ввода static ARP. Когда удаляется VLAN интерфейс или изменяется IP сегмента подсети, и при статическом и при динамическом режиме ввода в таблицу ARP table исходные записи сегмента подсети будут удалены.

Таблица команд конфигурации static ARP:

Команда	Описание	Режим CLI
arp <ip-address> <mac-address>	Конфигурация ввода данных в таблицу static ARP. Первый параметр – IP адрес принадлежащий сегменту подсети. Второй параметр - MAC адрес, должен быть unicast MAC адресом. Формат MAC адреса hhhh НННН. Hhhh, например 0010.5cb1 7825.	Global configuration mode
no arp <ip-address>	Отмена ввода данных в таблицу ARP, включая удаление IP адреса.	Global configuration mode

21.2.2 View ARP information (Просмотр информации ARP)

Таблица команд просмотра конфигурации ARP:

Команда	Описание	Режим CLI
show arp	Просмотр информации о таблице ARP, включая содержание по каждому пункту таблицы.	Normal mode, privileged mode
show running-config	Просмотр текущей конфигурации системы, в т.ч. конфигурации ARP.	privileged mode

21.3 Configure static routing (Конфигурация статической маршрутизации)

Статическая маршрутизация (Static routing) позволяет направлять пакеты данных от источника к получателю по маршруту заданному пользователем. Пользователь может сконфигурировать

статический маршрут как маршрут по умолчанию для отправки пакетов данных к шлюзу.

Статический маршрут задается вручную администратором, и совместим с сетями, имеющими простую структуру. Для нормальной работы коммутатора администратору достаточно только правильно задать статический маршрут. Статическая маршрутизация не отнимает много сетевых ресурсов т.к. автоматически обновление маршрута не происходит.

Маршрут является статическим по умолчанию. Маршрут по умолчанию используется тогда, когда не найден соответствующий подходящий маршрут. В маршрутной таблице маршрут по умолчанию записан для сети как 0.0.0.0/0 (маска 0.0.0.0). Если в таблице маршрутов отсутствует назначение (получатель сообщения) и отсутствует маршрут по умолчанию, то сообщение ICMP источнику показывает недоступность сетевого адреса назначения, далее сообщение сбрасывается. В современных сетях с сотнями коммутаторов, использование протокола динамической маршрутизации может потребовать расходования значительных ресурсов пропускной способности сети. Применение функции default routing уменьшает время затраченное на маршрутизацию пересылаемых данных и как следствие делает более быстрым взаимодействие большого количества подключенных к сети пользователей.

Коммутатор позволяет создавать множество статических маршрутов для одного получателя, но только один из них может быть использован для передачи данных одновременно. По умолчанию коммутатор не имеет заранее сконфигурированных статических.

Таблица команд конфигурации статической маршрутизации:

Команда	Описание	Режим CLI
ip route <ip-prefix> <nexthop-address>	Создание статического маршрута. Первый параметр определяет IP сегмента и длину префикса сети, Второй параметр определяет следующий hop IP адреса.	Global configuration mode

ip route <ip-address> <mask-address> <nexthop-address>	Создание статического маршрута. Первый параметр определяет IP адрес сегмента сети, второй параметр определяет маску подсети сегмента, третий параметр определяет следующий hop IP адреса.	Global configuration mode
no ip route <ip-prefix> [nexthop-address]	Удаление статического маршрута. Первый параметр определяет IP адрес сегмента сети, второй параметр определяет следующий hop IP адреса. При отсутствии первого параметра отменяется второй параметр. При наличии второго параметра, соответствующий путь к сегменту сети и следующий hop будут удалены.	Global configuration mode
no ip route <ip-address> <mask-address> [nexthop-address]	Удаление статического маршрута. адрес сегмента сети, второй параметр определяет маску подсети сегмента, третий параметр определяет следующий hop IP адреса. При отсутствии третьего параметра, все маршруты соответствующие сегменту сети будут удалены. При наличии третьего параметра, путь к сегменту сети и следующий hop будут удалены.	Global configuration mode

Команда	Описание	Режим CLI
show ip route [<ip-address> <ip-prefix>]	Просмотр конфигурации активных маршрутов. Можно выбрать все маршруты, маршруты соответствующие сегменту, статические маршруты.	Normal mode, privileged mode

show ip route database	Просмотр конфигурации всех маршрутов (активных и неактивных).	Normal mode, privileged mode
show running-config	Просмотр текущей конфигурации системы, в т.ч., конфигурации статических маршрутов.	privileged mode

Пример:

Установить сеть назначения 200.1.1.0, маска подсети 255.255.255.0, следующий hop to 10.1.1.2.

Команды конфигурации:

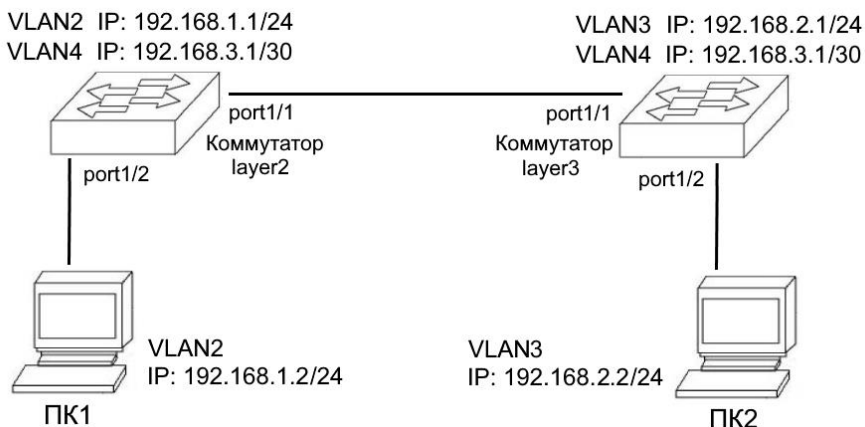
Switch(config)#ip route 200.1.1.0 255.255.255.0 10.1.1.2
или Switch(config)#ip route 200.1.1.0/24 10.1.1.2

Удалить статический маршрут IP адрес назначения 200.1.1.0, маска подсети 255.255.255.0, следующий hop of 10.1.1.2.

Команды конфигурации:

Switch(config)#no ip route 200.1.1.0/24
или Switch(config)#no ip route 2001.1.0/24 10.1.1.2
или Switch(config)#no ip route 200.1.1.0 255.255.255.0
или Switch(config)#no ip route 200.1.1.0 255.255.255.0 10.1.1.2

21.4 TCP/IP Basic configuration example (Пример конфигурации)



21.4.1 L3 layer interface (Интерфейс уровня L3)

Задать конфигурацию интерфейса уровня 3-layer на коммутаторе 1, соответствующую одновременно vlan2 и IP адресу 192.168.1.1/24.

Команды конфигурации:

```
Switch#config t
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#exit
Switch(config)#interface ge1/2
Switch(config-ge1/2)#switch access vlan 2
Switch(config)#ip interface vlan 2
Switch(config)#interface vlan2
Switch(config-vlan2)#ip add 192.168.1.1/24
```

Для верификации пользователь 1 может проверить ping IP адреса интерфейса уровня layer 3 соответствующего vlan2 коммутатора 1.

21.4.2 Static routing (Статическая маршрутизация)

Для получения доступа к коммутатору 1 пользователь 2 должен пройти маршрут от коммутатора 2 до коммутатора 1.

Команды конфигурации:

```
Switch#config t
Switch(config)#ip route 192.168.2.0 255.255.255.0 192.168.3.2
The configuration on switch 2 is as follows:
Switch#config t
Switch(config)#ip route 192.168.1.0/24 192.168.3.1
```

Для верификации пользователь 2 может проверить ping коммутатора 1.

21.4.3 ARP (Протокол ARP)

Сконфигурировать static ARP пользователя 1 только для доступа из vlan2. MAC адрес пользователя 1 is 00:00:00:00:01.

Команды конфигурации коммутатора 1:

```
Switch#config t
Switch(config)#arp 192.168.1.2 0000.0000.0001
```

Для верификации пользователь 1 может проверить ping IP адреса интерфейса уровня layer 3 соответствующего vlan2 коммутатора 1.

22. Configure SNMP (Конфигурация SNMP)

Протокол SNMP предназначен для дистанционного управления коммутатором. В данной главе описывается конфигурация протокола SNMP и содержит следующие разделы:

- SNMP Introduction (Введение)
- SNMP configuration (Конфигурация)
- SNMP configuration example (Пример конфигурации)

22.1 SNMP introduction (Введение)

SNMP является наиболее распространенным сетевым протоколом управления в настоящее время. Протокол поддерживает управление ошибками (fault management), биллингом (billing management), конфигурацией (configuration management), представлением (performance management) и безопасностью (security management). Протокол обеспечивает формат согласования между программами управления сетью и агентами управления сетью (network management agent).

SNMP включает в себя следующие элементы: рабочую станцию управления (management workstation), агента управления (management agent), информационную базу управления (management information base) и сетевой протокол управления (network management protocol). По отношению к коммутатору, агент управления это подключенный сервер рабочей станции управления. Информация рабочей станции доступна сетевому агенту представляется в форме MIB от информационной базы управления.

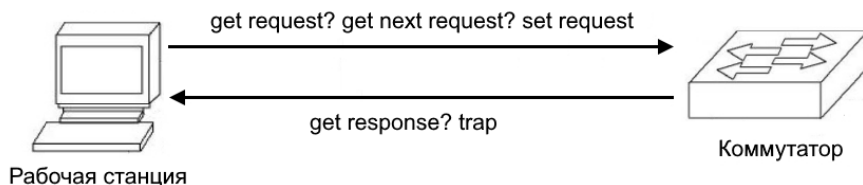
SNMP выполняет три операции: получение (get), установка (set) и захват (trap). Операция получения (get operation) позволяет станции управления получать значение агента. Операция установки (set operation) позволяет рабочей станции устанавливать значение агента. Операция захвата позволяет агенту сообщать о событиях рабочей станции управления.

Сообщения Trap messages посылаются рабочей станции когда на коммутаторе происходят события. Эти события включают холодный старт (cold start), горячий старт (hot start), установку соединения (link up) и разрыв соединения (link down) на портах, ошибку общего имени аутентификации, статус STP и т.д.

В настоящее время существует три версии протокола: SNMPv1, SNMPv2 и SNMPv3. Последующие версии являются доработанными предыдущими версиями с улучшенными функциями и повышенной безопасностью. Коммутатор поддерживает все три версии протокола SNMP и может анализировать все три версии пакетов протокола SNMP. При отправке сообщения trap message возможно использовать любую версию протокола SNMPv1, SNMPv2 или SNMPv3.

Коммутатор поддерживает RFC, bridge и private MIB объекты. Протокол SNMP позволяет полностью управлять коммутатором. Коммутатор поддерживает следующие MIB: RFC 1213, RFC 1493, RFC 1724, RFC 1850, RFC 1907, RFC 2233, RFC 2571, RFC 2572, RFC 2573, RFC 2574, RFC 2575, RFC 2674 и другие общие MIB.

На рисунке ниже представлено взаимодействие между рабочей станцией и агентом управления по протоколу SNMP. Рабочая станция получает доступ к агенту управления коммутатором путем отправки сообщений запроса (SNMP get request, GetNext request, getbulk request и set request), получения или установки значения объекта MIB коммутатора. Агент управления отправляет ответное сообщение SNMP рабочей станции управления. При возникновении каких-либо событий на коммутаторе, агент управления коммутатором направляет сообщение SNMP trap message рабочей станции.



22.2 SNMP configuration (Конфигурация SNMP)

Конфигурация SNMP включает общую настройку коммутатора, рабочей станции и информационной системы SNMP. По умолчанию коммутатор имеет имя common body, которое является публичным параметром. Для конфигурации доступно до 8 имен common bodies. По умолчанию коммутатор не имеет рабочей станции trap.

Таблица команд конфигурации SNMP:

Команда	Описание	Режим CLI
snmp community <community-name> {ro rw}	Конфигурация общего имени доступа к управлению сетью. Команда является интерактивной. Пользователь может вводить требуемое имя и давать разрешение на чтение/запись.	Global configuration mode
no snmp community <community-name>	Удалить выбранное имя SNMP community name.	Global configuration mode
snmp trap <notify-name> host <ipaddress> version {1 2c 3}	Добавление или изменение получателя SNMP trap. Команда является интерактивной, имя д.б. уникальным. При внесении изменений в имя возможно изменить trap target item. Хост является trap, возможные версии - SNMPv1, SNMPv2c, SNMPv3. По умолчанию для порта назначения команда 162.	Global configuration mode
no snmp trap <notify-name>	Удаление выбранного SNMP trap.	Global configuration mode
snmp system information <contact location name> <information-string>	Конфигурация системы информации (включает contact, location и имя).	Global configuration mode
no snmp system information <contact location name >	Удаление конфигурации системы информации.	Global configuration mode
show snmp community	Просмотр всех текущих публичных body names и статуса разрешений чтения/записи информации.	Normal mode / privileged mode
show snmp trap	Просмотр всех текущих trap names и target IP адресов, версии информации соответствующего trap.	Normal mode / privileged mode
show snmp system information	Просмотр установок SNMP системы информации.	Normal mode / privileged mode

22.3 SNMP configuration example (Пример конфигурации)

Сконфигурировать common name named private с разрешением на чтение/запись. Проверить отправку SNMP trap named, IP адрес назначения 10.10.10.10, версия протокола SNMPv1.

В конфигурацию дополнительно включить контактную информацию E-mail: networks@acb.com, местонахождение Shenzhen.

Команды конфигурации system name и дополнительной информации switch:

```
Switch#config t
Switch(config)#snmp community private rw
Switch(config)#snmp system information contact E-
mail:networks@abc.com
Switch(config)#snmp system information location Shenzhen
Switch(config)#snmp system information name switch
```

23. Configure RMON (Конфигурация RMON)

Протокол RMON (remote monitoring) предназначен для мониторинга потока данных в сети. В данной главе описывается конфигурация протокола RMON и содержит следующие разделы:

- RMON Introduction (Введение)
- RMON configuration (Конфигурация)
- RMON configuration example (Пример конфигурации)

23.1 RMON introduction (Введение)

Протокол RMON (remote monitoring) в основном используется для дистанционного мониторинга потоков данных циркулирующих в пределах сегмента сети или в некоторых случаях во всей сети. В настоящее время данный сетевой стандарт управления является наиболее распространенным. RMON является расширенной разновидностью протокола SNMP MIB, и наиболее усовершенствованной версией стандарта MIB II. RMON использует SNMP для наиболее эффективного мониторинга удаленного сетевого оборудования.

Система мониторинга RMON состоит из двух частей: детектора (агент или монитор) и станции управления. Агент RMON хранит сетевую информацию в RMON MIB, который непосредственно встроен в сетевое оборудование (роутеры, коммутаторы и т.п.). Станция управления использует SNMP для получения информации о данных RMON.

Оборудование поддерживает использование четырех общедоступных групп RMON:

- (1) Statistics (статистическая группа) – обеспечивает статистическими данными каждого интерфейса. Большинство объектов это счетчики, которые записывают информацию полученную мониторами отинтерфейса.
- (2) History (группа история) – сохраняет данные определенного интерфейса в фиксированные интервалы времени.
- (3) Alarm group (тревожная группа): сохраняет определенные данные всех интерфейсов в фиксированные интервалы времени и сравнивает их с установленными пороговыми значениями. При совпадении условий соответствующие события триггируются.
- (4) Event group (группа события) – фиксирует события, возможно сохранение логов или фиксация.

23.2 RMON Configuration (Конфигурация RMON)

Команды RMON включают конфигурацию четырех групп, команды просмотра конфигурации и данных.

Таблица команд конфигурации RMON:

Команда	Описание	Режим CLI
rmon statistics <1-100> (owner WORD)	Включение статистической группы определенного номера (интерактивная команда). Ввод серийного номера (номер статистической группы от 1 до 100) и владельца (опциональный параметр).	Port configuration mode
no rmon statistics <1-100>	Сброс конфигурации определенной статистической группы.	Port configuration mode

rmon history <1-100> buckets <1-100> interval <1-3600> (owner WORD)	Конфигурация параметров группы истории определенного номера (интерактивная команда). Ввод серийного номера (номера группы от 1 до 100), количества запросов (максимальное количество сохраненных данных от 1 до 100), длительность временного интервала (1 – 3600 сек), и владельца.	Port configuration mode
no rmon history <1-100>	Сброс конфигурации определенной группы историй.	Port configuration mode
rmon alarm <1-60> WORD <1-3600> (absolute delta) rising-threshold <1-2147483647> <1-60> falling-threshold <1-2147483647> <1-60> (owner WORD)	Конфигурация параметров тревожной группы определенного номера (интерактивная команда). Ввод номера группы (от 1 до 60), объекта мониторинга (id MIB нода), длительности временного интервала (1 – 3600 сек), метода сравнения (по абсолютному или относительному значению), верхнего порога, предельного номера события, нижнего порога, нижнего номера события и владельца. Значения параметров верхнего и нижнего порогов от 1 до 2147483647. Номера событий задаются в пределах от 1 до 60.	Global configuration mode
no rmon alarm <1-60>	Сброс конфигурации определенной тревожной группы.	Global configuration mode
rmon event <1-60> (log log-trap WORD none trap WORD) (description WORD) (owner WORD)	Конфигурация параметров определенной группы событий (интерактивная команда). Ввод номера группы (от 1 до 60), типа события (log, log trap, none и trap), имени common body, описания и владельца.	Global configuration mode

	При выборе типа события log trap или trap, должно быть определено имя common body (в противном случае конфигурация имени common body будет проигнорирована).	
no rmon event <1-60>	Сброс конфигурации определенной группы событий.	Global configuration mode
show rmon (statistics history-control alarm event) config	Просмотр конфигурации RMON (интерактивная команда), пользователь может просматривать информацию об объекте.	Global configuration mode
show rmon statistics-data interface IFNAME	Просмотр данных статистической группы RMON, требуется ввести имя интерфейса.	Global configuration mode
show rmon history-data interface IFNAME	Просмотр данных группы истории RMON, требуется ввести имя интерфейса.	Global configuration mode

23.3 RMON Sconfiguration example (Пример конфигурации)

Сконфигурировать статистическую группу с номером 10 для порта GE1/1, и владельца «tereco».

Сконфигурировать группу истории с номером 2 для порта GE1/8. Максимальное значение данных 80, временной интервал 1 мин., владелец отсутствует.

Сконфигурировать группу событий с номером 1, запись log, владелец отсутствует.

Сконфигурировать группу событий с номером 3, отправление trap, имя common body общедоступное, владелец отсутствует.

Сконфигурировать тревожную группу с номером 5 для мониторинга количества байт принятых каждым портом. При количестве байт превышающем 1000 за 30 сек., сформировать trap alarm, если количество байт меньше 10, записать log. Владелец отсутствует.

Команды конфигурации коммутатора:

```
Switch#configure terminal
Switch(config)#interface ge1/1
Switch(config-ge1/1)#rmon statistics 10 owner tereco
Switch(config-ge1/1)#exit
Switch(config)#interface ge1/8
Switch(config-ge1/8)#rmon history 2 buckets 80 interval 60
Switch(config-ge1/8)#exit
Switch(config)#rmon event 1 log
Switch(config)#rmon event 3 trap public
Switch(config)#rmon alarm 5 1.3.6.1.2.1.2.2.1.10 30 delta rising-threshold
1000 3 falling-threshold 10 1
```

24. Cluster configuration (Конфигурация кластера)

Коммутатор поддерживает функцию управления кластером, которая позволяет реализовать объединение сетевых устройств в группу для управления одним устройством. Глава описывает конфигурацию данной функции и содержит следующие разделы:

- Introduction to cluster management (Введение)
- Configuration management device (Управляющее устройство)
- Configure member devices (Группа управляемых устройств)
- Cluster management display (Просмотр конфигурации)
- Typical configuration examples (Пример конфигурации)

24.1 Cluster management introduction (Введение)

24.1.1 Cluster definition (Определение кластера)

Кластер это группа сетевого оборудования, которая может управляться одним устройством. Функция Cluster management помогает решить проблему централизованного управления большим количеством распределенных сетевых устройств.

Преимущества кластера является возможность пользоваться публичными IP адресами, и упрощение конфигурации задач управления.

Для управления и обслуживания коммутаторов входящих в кластер, администратору сети достаточно сконфигурировать публичные IP адреса на одном из коммутаторов кластера.

Командный коммутатор (command switch) отвечает за конфигурацию публичных IP адресов и управление остальными коммутаторами членами (member switch) кластера. Для конфигурации и управления коммутаторами внутри кластера используются следующие протоколы:

- NDP (Neighbor Discovery Protocol)
- NTDP (Neighbor Topology Discovery Protocol)
- Cluster (Cluster Management Protocol)

К рабочим процессам кластера относится определение топологии, создание кластера и его обслуживание. Определение топологии и обслуживание являются независимыми процессами. Процесс определения топологии должен начинаться до формирования кластера. Принципы работы следующие:

- Все устройства получают информацию о соседних устройствах по протоколу NDP, включая версию прошивки, имя хоста, MAC адрес, имя порта и пр.
- Управляющее устройство получает информацию о других устройствах по протоколу NTDP, включая диапазон hop number range (определенного пользователем), информацию о соединениях, и далее определяет кандидата в кластер исходя из полученных данных топологии.
- Управляющее устройство завершает операции по добавлению кандидатов в кластер на основании информации собранной протоколом NTDP и покидает кластер.
- Все сообщения кластера относятся к сообщениям уровня layer-2 Ethernet messages. Формат и взаимодействие процессов определяется стандартами и техническими требованиями YDT 1692-2007 для коммутаторов Ethernet с поддержкой функции cluster management.

24.1.2 Cluster role (Роль кластера)

Различные роли формируются в соответствии со статусом и функциями каждого сетевого устройства кластера. Пользователь имеет возможность конфигурировать роли:

1) Командный коммутатор (Command switch), единственный коммутатор, который может управлять и конфигурировать кластером, имеет публичный IP адрес в кластере:

- Создает кластеры;
- Обнаруживает и определяет кандидатов в кластер на основании информации протоколов NDP (neighbor discovery protocol) и NTDP (neighbor Topology Discovery Protocol);
- Управляет обслуживанием кластера. Пользователь может добавлять или удалять кандидатов в кластер;
- Обеспечивает канал управления кластером после окончания формирования кластера.

2) Коммутатор-член (Member switch), один из управляемых коммутаторов в кластере:

- Имеет статус кандидата до вступления в кластер;
- Не имеет публичного IP адреса;
- Управление осуществляется через агента командного коммутатора.

3) Коммутатор-кандидат (Candidate switch) потенциально имеет возможность стать членом кластера, для вхождения в кластер коммутатор обязательно должен иметь статус кандидата.

4) Независимый коммутатор (Independent switch), функция cluster function на коммутаторе отключена. Может работать в соответствии со следующими ролями:

- В процессе создания кластера пользователь назначает независимый коммутатор управляющим кластером. Каждый кластер должен иметь только один командный (управляющий кластером) коммутатор. После назначения командного коммутатора, он на основе полученной информации определяет кандидатов в кластер. Пользователь имеет возможность добавлять дополнительных кандидатов в кластер;
- После добавления в кластер кандидаты становятся членами кластера;

- После удаления членов кластера, они снова становятся кандидатами;
- Командный коммутатор снова становится кандидатом только при удалении кластера.

24.1.3 NDP introduction (Протокол NDP)

Протокол NDP используется для получения информации о непосредственно связанных соседних устройствах, включая соединение порта, имя устройства, версию прошивки и т.п. Принцип работы протокола следующий:

- Оборудование с запущенным протоколом NDP периодически отправляет NDP сообщения соседним устройствам, которые содержат информацию (соединение порта, имя устройства, версию прошивки и т.п.) и время обновления информации NDP на принимающем устройстве. Одновременно принимаются, но не отправляются сообщения NDP от соседних устройств;
- Оборудование с запущенным протоколом NDP хранит и обслуживает таблицу информации о соседних устройствах (NDP neighbor information table), дает доступ к таблице соседним устройствам. При обнаружении нового соседнего устройства NDP сообщение от него принимается впервые, а table entry будет добавлена в таблицу NDP neighbor information table. Если информация NDP полученная от соседнего устройства отличается от предыдущей информации, соответствующие данные в таблице NDP будут обновлены. Если информация не изменилась, то только время обновления (aging time) будет обновлено. Если информация NDP отправленная соседним устройством не будет получена вовремя, то соответствующие пункты таблицы neighbor table будут удалены автоматически.

24.1.4 NTDP introduction (Протокол NTDP)

Протокол NTDP используется для сбора информации о каждом устройстве и их сетевых соединениях в заданном диапазоне. Протокол NTDP обеспечивает управляемые устройства информацией для

включения в кластер и осуществляет сбор информации о топологии подключения устройств с учетом выбранных (specified) hops.

NDP обеспечивает протокол NTDP обновляемой информацией из таблицы. NTDP отправляет и пересылает запросы о топологии в соответствии с её изменениями для сбора информации NDP каждого устройства в сети и согласовании её со всеми соседними устройствами. После сбора информации, управляющее устройство может использовать информацию для выполнения своих функций. Когда протокол NDP на устройствах членах обнаруживает изменения в соседних устройствах, он уведомляет управляющее оборудование об изменениях сообщением handshake message. Управляющее устройство может запустить протокол NTDP для уточнения топологии и вовремя отразить изменения, произошедшие в сети.

Управляющее устройство может проводить сбор информации о сетевой топологии на регулярной основе, а пользователь может инициировать процедуру вручную путем введения соответствующих команд. Процесс сбора управляющим устройством информации о топологии сети происходит следующим образом:

- Управляющее устройство с соответствующего порта регулярно отправляет запросы NTDP о топологии.
- Получившее запрос устройство немедленно отправляет ответное сообщение с данными топологии управляющему оборудованию, копия запроса NTDP направляется соседним устройствам. Ответное сообщение с данными топологии содержит основную информацию об устройстве и информацию NDP всех соседних устройств.
- После получения запроса, соседнее оборудование будет проделывать ту же операцию пока запрос о топологии не получат все соседние устройства в определенном диапазоне hop number.

Когда запрос топологии распространится по сети и все сетевые устройства получают запросы, должны быть вовремя отправлены ответные сообщения. Для предотвращения перегрузки сети и управляющего устройства должны быть предприняты следующие меры:

- После получения запроса о топологии, устройства не отправляют ответное сообщение немедленно, а задерживают отправку на некоторое время.
- Каждый порт устройства с включенной функцией NTDP (кроме первого порта) отправляет ответное сообщение только после отправки сообщения предыдущим портом.

24.1.5 Cluster management (Управление кластером)

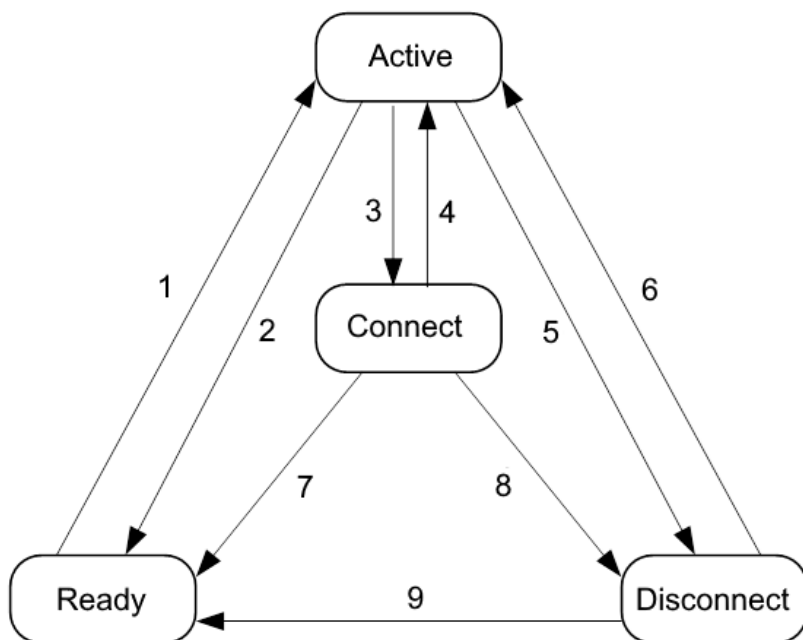
1) Кандидат на включение в кластер:

- Перед созданием кластера пользователь назначает управляющее устройство. Управляющее устройство находит и определяет кандидатов на основании данных протоколов NDP и NTDP и автоматически добавляет их в кластер. Пользователь имеет возможность добавить кандидатов вручную.
- После успешного включения кандидатов в кластер, устройства становятся членами кластера и получают последовательные номера и информацию от управляющего оборудования (IP адрес и т.п.).

2) Внутренняя коммуникация кластера:

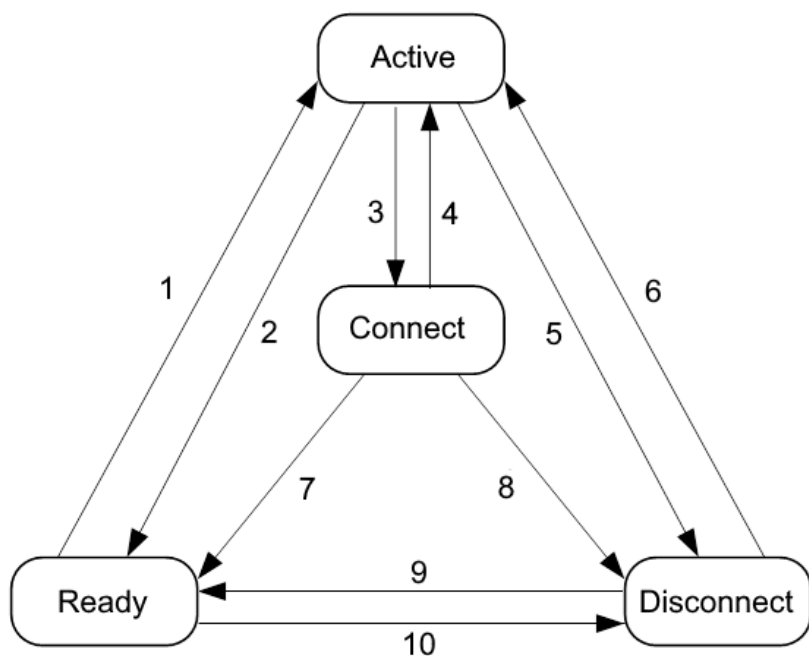
- Управляющее устройство и члены кластера осуществляют коммуникацию в масштабе реального времени с использованием handshake сообщений.

Статус соединений между управляющим устройством и членами кластера показан на рисунке ниже.



Преобразование статуса командного коммутатора:

- 1 member joining – вступление члена
- 2 member deletion – удаление члена
- 3 the handshake signal is not received for three consecutive times – сигнал handshake не получен за три попытки
- 4 receive handshake signal – получение сигнала handshake
- 5 recovery request received – получен запрос восстановления
- 6 interruption recovery, re registration passed – прервано восстановление
- 7 member deletion – удаление члена
- 8 the state remains for more than the specified time – статус не изменился за определенное время
- 9 delete member – удаление члена



Миграция статуса коммутатора члена:

- 1 join the cluster – вступление в кластер
- 2 exit the cluster – выход из кластера
- 3 the handshake signal is not received for three consecutive times - сигнал handshake не получен за три попытки

- 4 receive handshake signal - получение сигнала handshake
- 5 After receiving the join request – получен запрос на включение в кластер
- 6 interrupt recovery and re register - прервано восстановление
- 7 exit the cluster - выход из кластера
- 8 the state remains for more than the specified time or the join request message is received – статус сохраняется дольше положенного времени или получен запрос на вступление
- 9 exit the cluster - выход из кластера
10. Configuration recovery – восстановление конфигурации

Для получения статуса инициации коммутатору дается команда собрать основную информацию об устройствах, идентифицировать устройство как коммутатор-кандидат. Удаление члена в любом статусе приведет к возвращению статуса коммутатора к статусу кандидата.

- При успешном создании кластера и переходе кандидата в члены, управляющее устройство идентифицирует статус члена и сохраняет статус как активный. Устройство член тоже идентифицирует и сохраняет свой статус как активный.
- Управляющее устройство и устройство член регулярно обмениваются сообщениями handshake messages друг с другом. После получения сообщения handshake message от члена, управляющее устройство не реагирует и сохраняет статус члена как active state.
- Если управляющее устройство не получает сообщение handshake message отправленное членом в течение трех раз, то статус устройства члена изменится с active на connect. Если устройство член не получит сообщение handshake message отправленное управляющим устройством в течение трех раз, то статус устройства члена изменится с active на connect.
- Если управляющее устройство получит сообщение handshake message отправленное членом со статусом connect в течение определенного времени, то статус члена изменится обратно на active, в противном случае статус устройства изменится на disconnect, и управляющее устройство будет считать его отключенным. Если устройство член получит сообщение handshake message отправленное управляющим устройством в течение определенного промежутка времени, то статус устройства члена изменится с connect на active, в противном случае статус устройства изменится на disconnect.

При восстановлении прерванного соединения между управляющим устройством и устройством членом, устройство член изменит статус с `disconnected` на вступление в кластер. При успешном вступлении управляющее устройство и устройство член изменят статус на `active`.

При изменении топологии устройство член передаст информацию об изменении управляющему устройству через сообщение `handshake message`.

24.1.6 Manage VLAN (Управление VLAN)

Управление VLAN ограничивает функции управления кластером (cluster management), для конфигурации доступны следующие функции:

- Сообщения управления в кластере (NDP, NTDP, handshake messages) ограничиваются функциями VLAN и отделяются от других сообщений в целях повышения безопасности.
- Устройства управления и члены реализуют внутреннюю коммуникацию через VLAN управление. Требуется чтобы кандидаты соединялись друг с другом через каскадный управляющий порт, таким образом допускается соединение через VLAN.
- Если порт не позволяет проходить управлению VLAN, то устройства соединенные с портом не могут вступить в кластер. Таким образом, кластеризация различает порты соединенные с кандидатами, управляющим устройством, включая каскадные порты, что разрешает прохождение VLAN управления.
- Только когда по умолчанию VLAN ID порта соединяет управляющее устройство с устройством членом (кандидатом) и каскадный порт управляется VLAN, сообщение содержащее конфигурацию management VLAN пропускается без tag, в противном случае сообщение management VLAN должно проходить только с tag (см. раздел "configuring VLAN").

24.2 Cluster configuration introduction (Введение)

Перед началом конфигурации кластера, пользователь должен уточнить роли и функции различного оборудования входящего в кластер, определить релевантные функции оборудования и создать план коммуникации устройств кластера.

Конфигурация кластера	
Конфигурация управляющего оборудования	Включение функции NDP в системе и на порту
	Конфигурация параметров NDP
	Включение функции NTDP в системе и на порту
	Конфигурация параметров NTDP
	Ручная конфигурация информации NTDP
	Включение функции cluster function
	Создание кластера
	Конфигурация взаимодействия членов внутри кластера
	Конфигурация управления членами кластера
Конфигурация устройств-членов	Включение функции NDP в системе и на порту
	Включение функции NTDP в системе и на порту
	Ручная конфигурация информации NTDP
	Включение функции cluster function
Конфигурация взаимодействия членов кластера	

Внимание:

После создания кластера, кластер не будет расформирован при отключении функции NDP или NTDP на управляющем устройстве или устройстве члене, но будет нарушено нормальное функционирование кластера.

24.3 Configuration management device (Конфигурация управляющего устройства)

24.3.1 Enable NDP function (Включение функции NDP)

Таблица команд включения функции NDP на устройстве и порте:

Команда	Описание	Режим CLI
ndp global enable	Включение global NDP на коммутаторе. Функция отключена по умолчанию.	Configuration mode
ndp enable	Включение функции NDP на порту. Функция отключена на всех портах по умолчанию.	Interface configuration mode

Внимание:

- Функция NDP (global и port) должна быть одновременно включена на устройстве и порту для корректной работы.
- Функция NDP не поддерживает работу с агрегированными портами (aggregation ports).
- Для предотвращения получения информации управляющим устройством о топологии устройств не входящих в кластер, рекомендуется отключить функцию NDP на портах соединенных с такими устройствами.

24.3.2 Configure NDP parameters (Конфигурация параметров NDP)

Таблица команд конфигурации параметров NDP:

Команда	Описание	Режим CLI
ndp aging-timer <aging-time>	Ввод параметра aging time отправляемого сообщения NDP. Значение 180 сек по умолчанию.	Configuration mode
ndp hello-timer <hello-time>	Ввод временного интервала для отправки сообщений NDP. Значение 60 сек по умолчанию.	Configuration mode

Внимание:

Обычно параметр aging time NDP сообщения на принимающем устройстве не должен быть меньше чем временной интервал sending time NDP, в противном случае может возникнуть нестабильность в работе таблицы соседних портов NDP.

24.3.3 Enable NTDP function (Включение функции NTDP)

Таблица команд включения функции NTDP на устройстве и порте:

Команда	Описание	Режим CLI
ntdp global enable	Включение global NTDP на коммутаторе. Функция отключена по умолчанию.	Configuration mode
ntdp enable	Включение функции NTDP на порту. Функция отключена на всех портах по умолчанию.	Interface configuration mode

Внимание:

- Функция NTDP (global и port) должна быть одновременно включена на устройстве и порту для корректной работы.
- Функция NTDP не поддерживает работу с агрегированными портами (aggregation ports).
- Для предотвращения получения информации управляющим устройством о топологии устройств не входящих в кластер, рекомендуется отключить функцию NTDP на портах соединенных с такими устройствами.

24.3.4 Configure NTDP parameters (Конфигурация параметров NTDP)

Таблица команд конфигурации параметров NTDP:

Команда	Описание	Режим CLI
---------	----------	-----------

Команда	Описание	Режим CLI
ntdp hop <hop-value>	Конфигурация получения топологии. По умолчанию макс число hops между ближайшим и самым удаленным устройством - 3.	Configuration mode
ntdp timer <interval-time>	Установка временного интервала получения топологии. Значение 1 мин по умолчанию.	Configuration mode
ntdp timer hop-delay <time>	Установка времени ожидания определенных устройств до отправки ответного сообщения о топологии на первый порт. Значение 200 мс по умолчанию.	Configuration mode
ntdp timer port-delay <time>	Установка времени задержки на порту устройства для отправки запросов о топологии. Значение 20 мс по умолчанию.	Configuration mode

24.3.5 Configure NTPD information (Конфигурация информации NTPD)

После создания кластера, управляющее устройство начинает периодически собирать информацию о топологии. Кроме того, пользователь в любой момент может вручную собирать информацию NTPD (вне зависимости от того существует ли кластер или нет). Для повышения эффективности управления оборудованием пользователь может в реальном времени инициировать процесс сбора информации NTPD.

Таблица команд запуска сбора параметров NTPD:

Команда	Описание	Режим CLI
ntdp explore	Ручной запуск процесса сбора информации о топологии.	Normal mode, privileged mode

24.3.6 Enable cluster function (Функция создание кластера)

Таблица команд включения функции создания кластера:

Команда	Описание	Режим CLI
cluster enable	Включение функции создания кластера. Отключена по умолчанию.	Configuration mode

24.3.7 Establish cluster (Создание кластера)

Управление VLAN ограничивает функции управления кластером (cluster management), при конфигурации VLAN могут быть реализованы следующие функции:

- Сообщения управления в кластере (NDP, NTDP, handshake messages) ограничиваются функциями VLAN и отделяются от других сообщений в целях повышения безопасности.
- Устройства управления и члены реализуют внутреннюю коммуникацию через VLAN управление.

Таблица команд создания VLAN кластера:

Команда	Описание	Режим CLI
cluster management-vlan <vlan-id>	Назначение управляющего VLAN. По умолчанию назначен VLAN1.	Configuration mode

Внимание:

Если устройство принадлежит кластеру, то изменение управляющего VLAN не допускается.

Если устройство не принадлежит кластеру, то:

- Проверьте наличие VLAN и переходите к следующему шагу.
- Проверьте все интерфейсы. Если VLAN которому принадлежит интерфейс и управляющий VLAN не совпадают, то включите а затем отключите протоколы global NDP и NTDP на коммутаторе, далее включите их снова.

- Проверьте наличие VLAN интерфейса уровня layer-3. В случае его отсутствия необходимо создать соответствующий VLAN интерфейс уровня layer-3.
- Установите MAC интерфейса layer-3 как dev_ID. Если VLAN успешно установлен, а новый layer-3 интерфейс невозможно создать, то MAC vlan1 использует dev_ID.
- Если был сконфигурирован управляющий VLAN, а пользователь напрямую удаляет VLAN из базы данных VLAN, то управляющий VLAN будет автоматически установлен как vlan1, протоколы global NDP, NTDP и кластер будут отключены.

Перед созданием кластера, пользователь должен создать частный (private) диапазон IP адресов для устройств-членов кластера. При вступлении кандидата в кластер, управляющее устройство динамически определяет private IP адрес из диапазона для использования и назначает его кандидату для коммуникации с кластером, управления и обслуживания управляемого оборудования (устройств-членов).

Таблица команд создания диапазона (пула) IP адресов:

Команда	Описание	Режим CLI
cluster ip-pool <IP/MASK>	Ввод диапазона private IP адресов на управляющем устройстве для устройств-членов кластера.	Configuration mode

Внимание:

- IP адреса и пул адресов кластера VLAN интерфейса управляющего устройства и устройств-членов не могут быть сконфигурированы в одном сегменте сети, в противном случае кластер не будет работать корректно.
- Конфигурация возможна только если устройство не входит в кластер.
- Используйте управляющий VLAN для определения соответствующего layer-3 интерфейса. Отсутствие layer-3 интерфейса приведет к возникновению сбоя (устройство не сможет быть использовано как командный коммутатор кластера). Если имеется layer-3 интерфейс, то можно конфигурировать основные адреса пула. Если происходит сбой конфигурации, то нарушается и конфигурация пула адресов.

При создании кластера управляющее устройство по умолчанию отсутствует.

Таблица команд создания кластера:

Команда	Описание	Режим CLI
cluster build <name>	Создание кластера вручную, назначение управляющего устройства, ввод имени кластера.	Configuration mode
cluster auto-build <name>	Автоматическое создание кластера. Автоматическое добавление в кластер найденных кандидатов в соответствии с диапазоном hop number.	Configuration mode
cluster delete <name>	Удаление кластера.	Configuration mode
cluster stop auto-add member	Остановка автоматического добавления <i>новых</i> устройств-членов в кластер. Ранее добавленные устройства остаются в кластере.	Configuration mode

Внимание:

- Перед созданием кластера пользователь может только назначить управляющий VLAN. После включения устройств в кластер, пользователь не может изменять управляющий VLAN. Если необходимо изменить управляющий VLAN после создания кластера, необходимо удалить кластер на управляющем устройстве, далее изменить управляющий VLAN и заново создать кластер.
- По соображениям безопасности не рекомендуется конфигурировать управляющий VLAN как VLAN ID порта связывающего управляющее устройство, устройства-члены и каскадный порт.
- Если по умолчанию VLAN ID порта связывающего управляющее устройство с устройством-членом и каскадными портами является управляющим VLANs, допускается прохождение сообщений управляющих VLANs без меток (labels). В противном случае, управляющее устройство, порт соединяющий управляющее устройство с устройствами-членами и каскадными портами должны быть

сконфигурированы для разрешения пропускать сообщения управляющих VLANs без меток (labels). Для уточнения конфигурации см. главу "VLAN".

- Диапазон private IP адресов устройств-членов кластера может быть сконфигурирован на управляющем устройстве только до создания кластера. Если был создан кластер, система не может изменять диапазон IP адресов.

24.3.8 Configure member interaction (Взаимодействие членов кластера)

Управляющее устройство в кластере взаимодействует с устройствами-членами в масштабе реального времени с использованием сообщений handshake messages для установления статуса соединения между ними. Временной интервал отправки сообщений handshake messages и время удержания оборудования могут быть установлены на управляющем устройстве. Данная конфигурация будет применена ко всем устройствам-членам кластера одновременно.

Таблица команд настройки параметров взаимодействия устройств в кластере:

Команда	Описание	Режим CLI
cluster timer <interval-time>	Установка временного интервала отправки сообщений handshake messages. Значение по умолчанию 10с.	Configuration mode
cluster holdtime <hold-time>	Установка времени эффективного удержания оборудования. Значение по умолчанию 60с.	Configuration mode

24.3.9 Configure cluster member management (Настройка управления членами кластера)

Пользователь может вручную выбрать или удалить кандидата на включение в кластер на управляющем устройстве.

На управляющем устройстве может быть выполнена операция определенного кандидата в кластер, в противном случае произойдет ошибка.

Таблица команд включения / удаления устройств-кандидатов в кластер:

Команда	Описание	Режим CLI
cluster add member mac-address <mac-address>	Включение устройства-кандидата в кластер.	Configuration mode
cluster delete member mac-address <mac-address>	Удаление устройства-кандидата в кластер.	Configuration mode

24.4 Configure member devices (Конфигурация устройств-членов)

24.4.1 Enable NDP function (Включение функции NDP)

См. п. 24.3.1

24.4.2 Enable NTDP function (Включение функции NTDP)

См. п. 24.3.3

24.4.3 Configure manual collection NTDP information (Включение функции NDP)

См. п. 24.3.6

24.5 Configure access cluster members (Конфигурация доступа устройств-членов)

После корректной конфигурации функций NDP, NTDP и кластера для управления и мониторинга через управляющее устройство могут быть настроены устройства-члены кластера. Возможно подключиться к интерфейсу определенного устройства-члена через управляющее устройство для настройки и управления.

Таблица команд подключения к интерфейсу устройства-члена кластера:

Команда	Описание	Режим CLI
cluster switch-to member <member-number>	Включение на управляющем устройстве интерфейса управления устройства-члена кластера.	Normal mode, privileged mode

Внимание:

Подключение Telnet используется для взаимосвязи между управляющим устройством и устройствами-членами, следует учитывать что:

- на подключаемом оборудовании должен быть запущен telnet сервер (команда «telnet server enable»), в противном случае подключение не произойдет.
- Если введенный номер устройства-члена не существует, будет получено сообщение об ошибке.
- Если при подключении на подключаемом оборудовании уже используется telnet, то произойдет сбой передачи команд.

24.6 Display cluster configuration (Просмотр конфигурации)

Таблица команд просмотра конфигурации кластера:

Команда	Описание	Режим CLI
show ndp[interface <ifname>]	Просмотр информации о конфигурации NDP.	Normal mode, privileged mode

Команда	Описание	Режим CLI
reset ndp statistics [interface <ifname>]	Очистить статистику NDP.	Configuration view
show ntdp	Просмотр информации о конфигурации ntdp.	Normal mode, privileged mode
show ntdp device-list	Просмотр информации собранной протоколом ntdp.	Normal mode, privileged mode
show ntdp single-device mac-address <mac-address>	Просмотр информации ntdp о конкретном устройстве.	Normal mode, privileged mode
show cluster	Просмотр статуса и статистики кластера к которому принадлежит оборудование.	Normal mode, privileged mode
show cluster topology	Просмотр информации о топологии кластера.	Normal mode, privileged mode
show cluster candidates [mac-address <mac-address>]	Просмотр информации устройства-кандидата.	Normal mode, privileged mode
show cluster members [<member-number>]	Просмотр информации об устройствах-членах кластера.	Normal mode, privileged mode

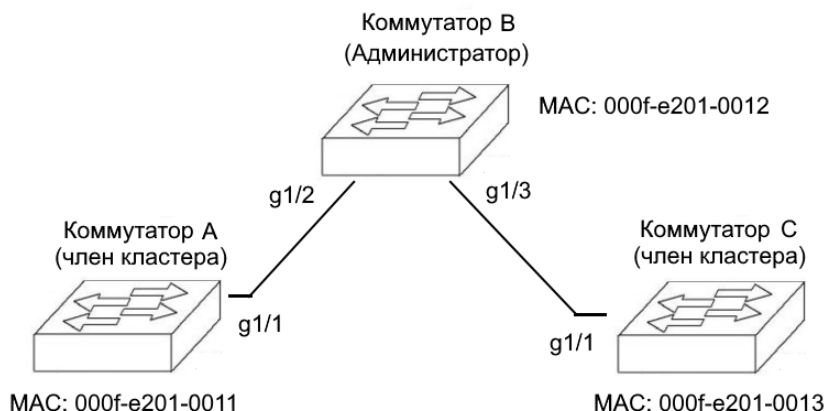
24.7 Cluster management example (Пример управления кластером)

Сетевые подключения:

Кластер ABC состоит из трех коммутаторов, VLAN10 - управляющий VLAN. Коммутатор В является управляющим устройством

(Administrator), коммутаторы А и С являются устройствами-членами. Основная группа ip адресов кластера – пул адресов 10.0.0.1, поддерживает 8 устройств.

Схема подключения:



Конфигурация оборудования:

Команды конфигурации коммутатора А (устройства-члена):

Конфигурация управляющего VLAN

```
[SwitchA] cluster management-vlan 10
```

```
[SwitchA] interface ge1/1
```

```
[SwitchA-ge1/1] switch access vlan 10
```

Включение функций global NDP и NDP на порту GE1/1

```
[SwitchA] ndp global enable
```

```
[SwitchA] interface ge1/1
```

```
[SwitchA-ge1/1] ndp enable
```

Включение функций global ntdp и ntdp на порту GE1/1

```
[SwitchA] ntdp global enable
```

```
[SwitchA] interface ge1/1
```

```
[SwitchA-ge1/1] ntdp enable
```

Создание кластера

```
[SwitchA] cluster enable
```

Конфигурация коммутатора С (устройства-члена)

Конфигурация всех устройств-членов проводится аналогичным образом, т.е. конфигурации коммутаторов А и С одинаковая.

Конфигурация устройства управления проводится аналогичным образом, но имеет незначительные отличия.

Команды конфигурации коммутатора В (управляющее устройство):

Конфигурация управляющего VLAN

```
[SwitchB] cluster management-vlan 10
```

```
[SwitchB] interface ge1/2
```

```
[SwitchB-ge1/2] switch access vlan 10
```

```
[SwitchB] interface ge1/3
```

```
[SwitchB-ge1/3] switch access vlan 10
```

Включение функций global NDP, ntdp и NDP, ntdp на портах GE1/2 и GE1/3 соответственно

```
[SwitchB] ndp global enable
```

```
[SwitchB] ntdp global enable
```

```
[SwitchB] interface ge1/1
```

```
[SwitchB-ge1/2] ndp enable
```

```
[SwitchB-ge1/2] ntdp enable
```

```
[SwitchB] interface ge1/3
```

```
[SwitchB-ge1/3] ndp enable
```

```
[SwitchB-ge1/3] ntdp enable
```

Конфигурация aging time сообщений NDP на принимающем оборудовании 200 сек

```
[SwitchB] ndp timer aging 200
```

Конфигурация временного интервала отправляемого сообщения NDP 70 сек

```
[SwitchB] ndp timer hello 70
```

Выбор максимального значения hops для сбора топологии (значение 2)

```
[SwitchB] ntdp hop 2
```

Конфигурация времени задержки пересылки сообщения запроса сбора топологии на первом порту оборудования 150 мс

```
[SwitchB] ntdp timer hop-delay 150
```

Конфигурация времени задержки пересылки сообщения запроса сбора топологии на других портах оборудования 15 мс

```
[SwitchB] ntdp timer port-delay 15
```

Конфигурация временного интервала сбора топологии 3 мин

```
[SwitchB] ntdp timer 3
```

Создание кластера

```
[SwitchB] cluster enable
```

Конфигурация частного private диапазона IP адресов для устройств-членов 10.0.0.1 - 10.0.0.9

```
[SwitchB] cluster ip-pool 10.0.0.1 8
```

Конфигурация данного коммутатора как управляющего оборудования, создание кластера с именем ABC. Кандидаты автоматически становятся членами кластера

```
[SwitchB] cluster autobuild abc
```

После добавления всех коммутаторов в кластер можно отключить функцию автодобавления auto join cluster function

```
[SwitchB] cluster stop auto-add member
```

25. SNTP Configuration (Конфигурация SNTP)

Глава описывает конфигурацию протокола SNTP (Simple Network Time Protocol), который является упрощенной версией протокола NTP (Network Time Protocol) и содержит следующие разделы:

- Introduction (Введение)
- Configuration SNTP (Конфигурация протокола)
- Display SNTP (Просмотр конфигурации)

25.1 SNTP introduction (Введение)

В настоящее время широко используются протоколы для проведения процедуры синхронизации сетевого времени, в т.ч. протокол NTP (Network Time Protocol). Упрощенной версией этого протокола является протокол SNTP (Simple Network Time Protocol).

Протокол NTP может быть развернут на различных платформах и операционных системах, использовать различные алгоритмы, но он критичен к таким параметрам сети как network delay и jitter. Протокол может обеспечить точность 1-50 мс, обеспечить поддержку аутентификации высокого уровня защиты. Протокол NTP является

комплексным и требовательным к системным параметрам.

В протоколе SNTP (Simple Network Time Protocol) используется простой алгоритм расчета времени, при котором высокая точность достигается через 1 сек, что достаточно в большинстве практических случаев.

С момента, когда сообщения протоколов SNTP message и NTP message начинают соответствовать, клиент SNTP сформированный коммутатором становится полностью совместимым с NTP сервером.

25.2 Configure SNTP (Конфигурация SNTP)

25.2.1 Default SNTP settings (Начальные установки)

Параметры протокола SNTP установленные по умолчанию:

Параметр	Значение по умолчанию
SNTP status (статус)	SNTP service отключен
NTP Server (сервер)	No (отсутствует)
Synchronization interval of SN (интервал синхронизации)	1800 sec (сек)
Local time zone (часовой пояс)	+8, East 8

Включение / отключение SNTP в режиме global configuration mode:

Switch# configure terminal

Включить SNTP

Switch(config)# sntp enable

Отключить SNTP

Switch(config)# sntp disable

25.2.2 Configure SNTP server (Конфигурация SNTP сервера)

С момента, когда сообщения протоколов SNTP message и NTP message начинают полностью соответствовать, клиент SNTP становится совместимым с NTP сервером. Т.к. в сети существует множество

серверов NTP, возможно выбрать один с минимальной сетевой задержкой (NTP сервер на коммутаторе).

Уникальный адрес NTP сервера может быть определен как <http://www.time.edu.cn/> или <http://www.ntp.org/> например: 192.43.244.18 (time. NIST. Gov).

В коммутатор может быть заведено до трех адресов серверов. Коммутатор использует первый адрес сервера для синхронизации времени. Если синхронизации не происходит, коммутатор использует второй адрес сервера и т.д.

Команды добавления адресов серверов в режиме global configuration mode:

```
Switch# configure terminal
Switch(config)# sntp server 210.72.145.44
```

Если в коммутаторе уже прописаны три адреса серверов, то при попытке добавления очередного адреса произойдет сбой. Для добавления нового адреса необходимо удалить один из существующих адресов.

Команды удаления адресов серверов:

Удалить все адреса серверов

```
Switch(config)# no sntp server
```

Удалить адрес сервера

```
Switch(config)# no sntp server 210.72.145.44
```

25.2.3 Configure the interval of SNTP synchronization (Конфигурация времени синхронизации)

Время SNTP клиента должно быть синхронизировано с часами NTP сервера для возможности корректировки.

Команды конфигурации:

```
Switch# configure terminal
Switch(config)# sntp interval 60
```

Установка интервала синхронизации времени в секундах, в диапазоне 60 - 65535 сек. По умолчанию установлено значение 1800 сек, для восстановления этого значения предусмотрена следующая команда:

```
Switch(config)# no sntp interval
```

25.2.4 Configure local time zone (Конфигурация часового пояса)

После установления коммуникации по SNTP протоколу устанавливается время по Гринвичу (Greenwich mean time - GMT). Для установки местного времени необходимо уточнить часовой пояс, по умолчанию установлен Dongba District (КНР).

Команды конфигурации:

Switch# configure terminal

Установка часового пояса West Zone 8

Switch(config)# sntp time-zone -8

Команда восстановления часового пояса по умолчанию (8 east zone)

Switch(config)# no sntp time-zone

25.2.5 SNTP information display (Просмотр конфигурации)

Команды просмотра конфигурации:

Switch# show sntp

Switch# show running-config

26. IGMP Configuration (Конфигурация IGMP)

Глава описывает конфигурацию протокола IGMP (Internet Group Management Protocol), который является частью технологии multicast и осуществляет связь между хостом и multicast роутером. Глава содержит следующие разделы:

- Introduction (Введение)
- Configuration IGMP (Конфигурация протокола)
- Configuration example (Пример конфигурации)

26.1 IGMP introduction (Введение)

Протокол IGMP (Internet Group Management Protocol) является частью технологии multicast и осуществляет связь между хостом и

multicast роутером. Хост сообщает multicast группе о подключении или отключении к локальному роутеру. Multicast роутер может периодически запрашивать локальный хост о его статусе получения для определенной группы. Протокол IGMP имеет три версии: версия 1 (rfc1112), версия 2 (rfc2236) и версия 3 (rfc3376). Протокол IGMP основывается на режиме ответов на запросы.

Роутер multicast использует цикл сообщений запросов для определения хоста в подсети и имеет определенную группу приема. Хост использует ответные сообщения (report message) для уведомления локального роутера multicast группы о вступлении или оставлении группы (membership report). Когда хост планирует вступить в определенную группу, нет необходимости ожидать периодического сообщения запроса (query message), и хост может немедленно отправить сообщение group membership report роутеру. Когда роутер запрашивает членство в хосте подсети, хост использует механизм поддержки для отправки сообщения report message со случайной задержкой от 1 до 10 сек. При получении той же группой сообщения report message от других хостов, хост прекращает отправку сообщений member report своей собственной группы для снижения трафика IGMP в подсети.

Если в подсети есть несколько роутеров, необходимо проводить сравнение IP адресов для выбора одного из них, чтобы снизить количество сообщений запросов IGMP query messages распространяемых в подсети.

Версия 1 поддерживает только запрос и ответ. Когда член покидает группу, multicast роутер останавливает пересылку потоков multicast данных только по истечении времени задержки групповой информации. Запросы зависят от выбранных протоколов.

Версия 2 поддерживает немедленную отправку информации о группе членов путем добавления сообщения departure message. Когда последний член покидает группу, роутеру отправляется сообщение departure message, и роутер отправляет специальный групповой group query message для определения других принимающих хостов в подсети. Пока соответствующее сообщение group membership message не получено (после отправки конфигурации пересылки), группа становится поврежденной, останавливается пересылка потоков данных multicast. Версия 2 эффективно снижает задержку отправления с помощью специального механизма, также Версия 2 поддерживает механизм

избирательного поиска. Когда главный запрос query message получен другими роутерами подсети, выбирается один с наименьшим IP адресом как поисковый (путем сравнения их собственных IP адресов и запускается таймер). При каждом получении сообщения запроса query message, таймер перезапускается. При обнулении таймера, начинается новый процесс выбора. Для установки времени ответа поддерживающего механизма хоста, поле максимального времени ответа добавляется в сообщение протокола Версии 2. Производящий запрос может контролировать подавление времени ответа хоста путем добавления его значения в сообщение запроса query message. Когда хост получает сообщение запроса query message, он использует максимальное время ответа response time для генерации задержки запуска таймера. Когда получено сообщение группы membership message других хостов в течение времени задержки, его собственное сообщение report message задерживается.

Версия 3 поддерживает multicast определенного источника. Сообщения запросов Query messages разделяются на главные (general query), специальные (specific group query) и специальные от источника (group specific source query). Ответные сообщения (report message) увеличивают поле оригинальной группы адресов в до группы записей. Каждая группа записей содержит адрес группы и список адресов соответствующих источников. Whether to join a specific source or not is distinguished by the Тип параметра group record type определяет включение источника в список выделенных источников. Версия 3 использует pim-ssm для реализации source specific multicast.

26.2 IGMP configuration (Конфигурация IGMP)

Функции IGMP в основном относятся к интерфейсам: конфигурация соответствующих параметров, значения таймеров каждого интерфейса и поддержание необходимой конфигурации интерфейса включают в себя:

- Включение функции IGMP для интерфейса;
- Конфигурация группового фильтра для контроля доступа интерфейса;
- Конфигурация группового фильтра интерфейса при выходе из группы;

- Конфигурация количества запросов определенной группы для интерфейса;
- Конфигурация интервала запросов определенной группы для интерфейса;
- Конфигурация таймера интерфейса при отсутствии запросов;
- Конфигурация таймера запросов интерфейса;
- Конфигурация максимального времени ответа интерфейса;
- Конфигурация параметров интерфейса;
- Конфигурация версии протокола интерфейса.

26.2.1 Start IGMP function of interface (Включение функции IGMP)

Включение / отключение IGMP в режиме interface configuration mode:

Включить SNTP

IP IGMP

Отключить SNTP

по IP IGMP

По умолчанию протокол IGMP отключен.

Команда позволяет интерфейсу отправлять и принимать сообщения протокола IGMP. Если команда IP multicast маршрутизации запустила поддержку функции multicast в режиме global configuration mode, будет автоматически запущен протокол IGMP на соответствующем интерфейсе после ввода команды запуска IP PIM sparse mode на интерфейсе в режиме interface configuration mode. Аналогично, команда отключения по IP PIM sparse mode в режиме interface configuration mode автоматически выключит протокол IGMP на соответствующем интерфейсе.

26.2.2 Configure the group filter access control list for the interface (Конфигурация фильтра группового списка доступа)

Команды конфигурации в режиме interface configuration mode:

Конфигурация фильтра списка группового доступа на интерфейсе:

IP IGMP access group {< ACL ID > | < ACL name >}

Отмена списка фильтра группового доступа на интерфейсе:
no IP IGMP access group

Параметр < ACL ID > показывает определенный номер стандартного списка доступа, может изменяться в пределах от 1 до 99. Параметр < ACL name > - имя списка стандартной группы контроля доступа. Функция ACL выполняет фильтрацию таблицы multicast адресов интерфейса. Команды конфигурации приведены в разделе ACL configuration. По умолчанию ACL группа фильтрации не сконфигурирована.

26.2.3 Configure the access control list filtered by the interface leaving the group (Выход из группы фильтра группового списка доступа)

Команды конфигурации в режиме interface configuration mode:

Конфигурация фильтра списка группового доступа на интерфейсе при выходе из группы:

```
IP IGMP immediate leave group list {< acl-id1 > | < acl-id2 > | < ACL name >}
```

Отмена выхода списка фильтра группового доступа на интерфейсе:
no IP IGMP immediate leave

Параметр < ACL Id1 > показывает определенный номер стандартного списка доступа, может изменяться в пределах от 1 до 99. Параметр < Acl-id2 > номер расширенного диапазона стандартного списка доступа, который может находиться в пределах 1300 – 1999. Параметр < ACL name > - имя списка стандартной группы контроля доступа. Функция ACL выполняет фильтрацию таблицы multicast адресов интерфейса, который должен быть отфильтрован, когда интерфейс получает сообщение об отправлении departure messages для версий 2 и 3. Команды конфигурации приведены в разделе ACL configuration. По умолчанию покидающая ACL группа фильтрации не сконфигурирована.

26.2.4 Number of specific group queries for the configuration interface (Конфигурация количества запросов группы интерфейса)

Команды конфигурации в режиме interface configuration mode:

Конфигурация количества запросов группы интерфейса:

IP IGMP last member query count < count >

Конфигурация количества запросов группы к значению по умолчанию:
no IP IGMP last member query count

Параметр < count > указывает количество запросов группы или источника, которые должны быть отправлены при получении сообщения отправления departure message, изменяется в пределах 2-7. Считается, что группа не имеет принимающего хоста пока не получено сообщение report message в течении этого периода. Значение параметра по умолчанию – 2.

26.2.5 Configure the specific group query interval of the interface (Конфигурация интервалов запросов группы интерфейса)

Команды конфигурации в режиме interface configuration mode:

Конфигурация интервала запросов группы интерфейса:

IP IGMP last member query interval < interval >

Конфигурация интервала запросов группы к значению по умолчанию:
no IP IGMP last member query interval

Параметр < interval > - временной интервал передачи запросов группы или определенного источника в диапазоне 1000 – 25500мс. Когда получено сообщение departure message, необходимо несколько раз отправить запрос определенной группы или источника. Команда устанавливает временной интервал между отправками сообщений запросов query messages. Пока не получено ответное сообщение report message на несколько запросов, считается что группа или источник не имеют принимающего хоста. Значение параметра по умолчанию - 1 сек.

26.2.6 Configure the non querier timer time of the interface (Конфигурация таймера отсутствия запросов)

Команды конфигурации в режиме interface configuration mode:

Установка времени таймера:

IP IGMP query timeout < time >

Установка времени таймера к значению по умолчанию:

no IP IGMP query timeout

Параметр < time > определяет время таймера в течении которого должно быть получено сообщение запроса, изменяется в пределах 60 - 300. Если в подсети присутствует несколько коммутаторов, то один из них должен быть определен для отправки запросов query message. При инициализации сети, все коммутаторы отправляют запросы query messages по умолчанию. Определение отправляющего коммутатора производится на основании запросов query messages и сравнении IP адресов, т.е. выбирается устройство с наименьшим IP адресом. Если сообщения запросов отсутствуют, то запускается таймер. Запуск таймера происходит каждый раз при получении сообщения запоса query message. Если время таймера истекает, то происходит переназначение отправляющего запросы устройства. Значение параметра по умолчанию - 255 сек.

26.2.7 Configure the query timer interval of the interface (Конфигурация интервалов отправки запросов)

Команды конфигурации в режиме interface configuration mode:

Конфигурация интервалов отправки запросов:

IP IGMP query interval < interval >

Конфигурация интервалов отправки запросов к значению по умолчанию:

no IP IGMP query interval

Параметр < interval > определяет временной интервал между отправкой сообщений запросов, изменяется в пределах 1-18000. При запуске протокола IGMP на интерфейсе, запускается таймер периодичности отправки сообщений запросов query messages в

подсети. Если в подсети присутствует несколько коммутаторов, то один из них должен быть определен как поисковый. Поисковый коммутатор использует таймер для периодической отправки сообщений запроса query messages. Остальные коммутаторы запускают таймеры задержки timeout timer. При истечении времени таймеров задержки timeout, заново выбирается поисковый коммутатор. Значение параметра по умолчанию - 125 сек.

26.2.8 Configure the maximum response time of the interface (Конфигурация максимального времени ответа интерфейса)

Команды конфигурации в режиме interface configuration mode:

Конфигурация максимального времени ожидания ответа на сообщение запроса query message:

IP IGMP query Max response time < time >

Конфигурация максимального времени ожидания ответа на сообщение запроса query message к значению по умолчанию:

no IP IGMP query Max response time

Параметр < time > определяет максимальное время ответа ожидания интерфейса при отправлении сообщения для версии 2 и 3, изменяется в пределах 1 – 240 сек, при этом точность измерения этого параметра для сообщений запроса query message составляет 0.1 сек. Протокол IGMP использует режим ответа на запрос (query response mode) для получения информации о соответствующей группе принимающего хоста подсети. Когда коммутатор отправляет сообщение запроса query message, сетевые хосты ответят группе сообщением report message. Для ограничения количества сообщений протокола IGMP, в сообщение запроса вводится параметр максимального времени ожидания (maximum response time), а получающий сообщение хост генерирует случайное значение согласно параметру maximum response time. Случайное значение задержки используется для отправки ответного сообщения. Когда сообщения group report message отправленные другими хостами принимаются в течении времени задержки, сообщения group report message задерживаются, что эффективно снижает вероятность возникновения конфликтов между сообщениями IGMP group reports. Значение параметра maximum response time для версии 2

может достигать величины 128, для версии 3 значение параметра может достигать большей величины. Значение параметра по умолчанию - 10 сек.

26.2.9 Configure interface parameters (Конфигурация параметров интерфейса)

Команды конфигурации в режиме interface configuration mode:

Конфигурация параметров интерфейса:

IP IGMP robustness variable < value >

Конфигурация параметров интерфейса к значению по умолчанию:

no IP IGMP robustness variable

Параметр < value > определяет жизнеспособность интерфейса, значение изменяется в пределах 2 - 7. Этот параметр характеризует устойчивость протокола. Способность адаптации к потере пакетов протокола и значение параметра влияет на количество ретрансляций и временные интервалы соответствующих таймеров. Сообщение запроса query message для версии 3 содержит значение QRV, которое используется для синхронизации параметра с устройствами не производящими запросов. Значение параметра по умолчанию - 2.

26.2.10 Configure protocol version of the interface (Конфигурация версии протокола интерфейса)

Команды конфигурации в режиме interface configuration mode:

Конфигурация версии протокола интерфейса:

IP IGMP version < version >

Конфигурация версии протокола к значению по умолчанию:

no IP IGMP version

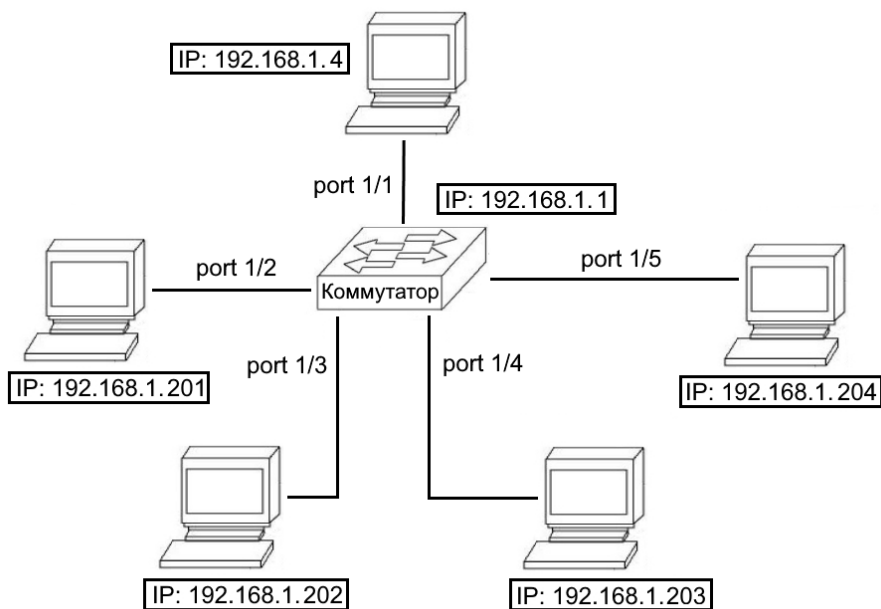
Параметр < version > определяет версию используемого протокола IGMP, значение изменяется в пределах 1 - 3. Значение параметра по умолчанию – 3.

26.3 IGMP configuration example (Пример конфигурации)

Обычно протокол IGMP работает совместно с протоколом маршрутизации multicast routing protocol (PIM-SM). См. пример конфигурации PIM-SM. Запуск протоколов multicast (IP multicast routing) и PIM-SM protocol (IP PIM sparse mode) на интерфейсе автоматически запускает соответствующий IGMP протокол без отдельной команды конфигурации. При приеме multicast, возможно просматривать соответствующую информацию о группе IGMP и статус интерфейса.

Когда хост источник и принимающий хосты находятся в одной подсети, поток данных multicast не перенаправляется через все сегменты сети, и multicast маршрутизация уровня 3 не является необходимой, протокол IGMP и функция IGMP snooping может быть включена отдельно для реализации multicast в подсети. При этом протокол PIM-SM должен быть запущен для записи таблицы multicast на устройстве уровня layer 3.

Схема конфигурации устройств приведена ниже:



Команды конфигурации:

```
Switch# configure terminal
Switch(config)#interface vlan2
Switch(config-vlan2)#ip address 192.168.1.1/24
Switch(config-vlan2)#ip igmp
Switch(config-vlan2)#ip pim sparse-mode passive
Switch(config-vlan2)#interface ge1/1
Switch(config-ge1/1)#switch access vlan 2
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#switch access vlan 2
Switch(config-ge1/2)#interface ge1/3
Switch(config-ge1/3)#switch access vlan 2
Switch(config-ge1/3)#interface ge1/4
Switch(config-ge1/4)#switch access vlan 2
Switch(config-ge1/4)#interface ge1/5
Switch(config-ge1/5)#switch access vlan 2
```

Команды просмотра конфигурации:

```
show ip igmp group
show ip igmp interface
```

27. PIM-SM Configuration (Конфигурация PIM-SM)

Глава описывает конфигурацию протокола PIM-SM (Sparse Mode Protocol Independent Multicast), который является частью технологии multicast и используется для определения топологии сети. Протокол осуществляет связь между роутерами. Глава содержит следующие разделы:

- Introduction (Ведение)
- Configuration PIM-SM (Конфигурация протокола)
- Configuration example (Пример конфигурации)

27.1 PIM-SM introduction (Введение)

Протокол PIM-SM (Sparse Mode Protocol Independent Multicast - Sparse Mode) является частью технологии multicast и используется для определения топологии сети. Протокол осуществляет связь между роутерами. В сравнении с передачей unicast точка-точка, передача multicast точка-многоточка используется для разварачивания дерева распределения, таким образом что поток данных источника может передаваться по дереву распределения к группе членов хостов.

Существует две формы дерева распределения. Одна из них основана на источнике. Каждый источник имеет дерево распределения для каждой группы. Кратчайший путь от источника к группе членов (активное дерево распределения) называется SPT. Другая форма основана на информации о точке распределения. Все источники группы делятся информацией о дереве распределения.

Информация источника должна быть зарегистрирована в точке конвергирования (RP: rendezvous point) перед тем как быть переданной по дереву распределения до принимающего хоста. Это дерево называется RPT. Технология дерева распределения протокола PIM-SM использует одностороннюю форму обмена информацией о дереве, т.е. поток данных может передаваться только от источника к получателю, но не наоборот.

Конечная цель протокола PIM-SM это применение к Wan, т.е. к отдельным получателям представленным в сети, режим которых отличается от pruning mode режима dense mode. Для создания дерева распределения протокол PIM-SM использует pruning mode, а принимающий хост использует режим explicit join mode (IP multicast mode).

Shared tree – это дерево распределения, которое распространяется всеми источниками группы. Когда различным источникам требуется отправить потоки данных, источник должен быть зарегистрирован в точке агрегации, и точка агрегации отправляет данные получателю. Кратчайший путь от источника к точке агрегации начинается с данных. Когда несколько источников зарегистрированы в точке конвергенции для отправки потоков данных, точка конвергенции становится «бутылочным горлышком» сети. В этот момент протокол PIM-SM позволяет каждому источнику и группе использовать свой кратчайший путь без прохождения точки конвергенции.

Для небольших сетей достаточно одной точки конвергенции. В случае большой сети, для распределения задач необходимо конфигурировать несколько RPS, но статический метод конфигурации не позволяет учитывать динамические изменения сети. Протокол PIM-SM обеспечивает функционирование инструментария RP bootstrap. Кандидаты BSRs и RPS конфигурируются в сети, при этом кандидаты BSR выбирают BSR. Кандидаты RP регулярно отправляют BSR данные маршрутизации собственной группы. BSR осуществляет сбор RP информации и объединяют её с информацией bootstrap, которая периодически распространяется по сети. Каждый роутер обладает информацией bootstrap. При вступлении в группу осуществляется поиск bootstrap информации, соответствующей группе RP и отправляется сообщение о вступлении в группу. Если отсутствует информация bootstrap соответствующая группе RP, используется hash алгоритм для маршрутизации группы RP и отправляется сообщение о вступлении в группу. Механизмы подтверждения и распространения информации по сети pim-dr работают одновременно. Роутер используется в процессе вхождения в дерево распределения distribution tree, и в процессе пересылки потоков данных. Протокол PIM-SM стал основным для multicast маршрутизации внутри домена. Расширенный протокол PIM-SSM дает возможность пользователям получать сервисы непосредственно от multicast источников, что превышает возможности традиционного протокола PIM-SM.

Коммутаторы уровня layer 3 дают возможность применять концепцию безопасности при использовании протоколов multicast. Технология Multicast может быть разделена на три части: обнаружение источника, обнаружение multicast получателя и получение топологии. Одна часть multicast протокола это протокол взаимодействия между хостом и роутером (IGMP); другая часть это протокол взаимодействия между роутерами (PIM-SM). Протоколы Multicast взаимодействия между хостом и роутером и между роутером и роутером используют интерфейс уровня 3-layer для описания направлений входа и выхода multicast маршрутизации. Для коммутаторов уровня layer-3, в случае выходного интерфейса уровня layer-3 multicast маршрута и портов VLAN уровня layer-2 соответствующие интерфейсу порты станут выходными. Если существует только один VOD хост исходящего потока данных VLAN, все остальные выходные порты VLAN не являются необходимыми. Таким образом, протокол IGMP snooping protocol

выполняет мониторинг сообщений IGMP messages порта и записывает их. Когда поток данных отправляется из VLAN, то он только направляется запрашивающему порту. Соответственно информация портов уровня layer 2 должна быть записана коммутатором уровня layer 3.

При взаимодействии роутеров, для обслуживания дерева распределения multicast, необходимо отправить сообщение о вступлении pruning message. Если отправляющие и получающие порты сообщений PIM протокола находятся под мониторингом, таблица отправляющего интерфейса multicast также записывает соответствующие layer-2 порты. Когда потоки данных проходят через multicast дерево, данные будут отправлены только портам VLAN получившим сообщение join message. PIM snooping это протокол уровня layer-2, который проводит мониторинг сообщений PIM протокола. Коммутатор уровня layer 3 может одновременно обеспечить работу протоколов multicast маршрутизации IGMP snooping и PIM snooping.

Таблица multicast маршрутизации содержит следующую информацию: адрес Multicast источника, адрес входного интерфейса multicast группы, список выходных портов, список выходных интерфейсов.

При наличии потока данных, ветка дерева распределения multicast от источника до принимающего хоста через порт уровня layer 2, где однократно появилось сообщение протокола (быстрее чем во VLAN) соответствует интерфейсу уровня layer 3.

Протоколы IGMP snooping и PIM snooping являются протоколами динамического мониторинга, которые записывают информацию портов уровня layer-2 согласно сообщению протокола. Коммутатор версии v3 Version 0 обеспечивает только статическую конфигурацию функции очистки multicast маршрутизации портов уровня layer 2. Данная функция называется multicast security. Через команды статической конфигурации функция определяет что выходной поток данных multicast отправлен только определенным портам VLAN уровня layer 2, соответствующим выходному интерфейсу. По умолчанию поток данных будет исходящим для всех портов VLAN уровня layer-2, соответствующим выходному интерфейсу.

27.2 PIM-SM configuration (Конфигурация PIM-SM)

Протокол PIM-SM должен быть запущен последовательно на каждом интерфейсе после запуска функции multicast уровня layer 3. Соответствующие команды выполняются в режимах global configuration mode и interface configuration mode.

Конфигурация протокола PIM-SM включает:

- Запуск функции multicast маршрутизации
- Конфигурация размера таблицы multicast маршрутизации
- Конфигурация TTL объема multicast интерфейса
- Запуск функции PIM SM на интерфейсе
- Конфигурация пассивного режима интерфейса
- Конфигурация приоритета интерфейса
- Конфигурация сообщений Hello message интерфейса (не содержит информации gen id)
- Конфигурация интервала таймера Hello timer
- Конфигурация времени удержания соседей интерфейса
- Конфигурация списка фильтрации соседей интерфейса
- Конфигурация адреса источника сообщения unicast registration
- Конфигурация предельного количества зарегистрированных сообщений registered messages
- Конфигурация проверки регистрации RP
- Конфигурация таймера времени ингибирования
- Конфигурация таймера Kat timer
- Конфигурация фильтрации адреса источника
- Конфигурация контрольной суммы зарегистрированного сообщения в режиме Cisco
- Конфигурация статических RP адресов
- Конфигурация кандидата RP
- Отклонение установленного приоритета RP
- Конфигурация сообщения c-gr-adv message в режиме Cisco
- Конфигурация кандидатов BSR
- Конфигурация времени таймера JP timer
- Конфигурация переключения SPT (shortest path tree)
- Конфигурация SSM (source specific multicast)

27.2.1 Start multicast routing (Запуск функции multicast routing)

Команды запуска/отключения функции multicast routing в режиме global configuration mode:

Команда запуска функции:

IP multicast routing

Команда отключения функции:

no IP multicast routing

Функция multicast routing function отключена по умолчанию, при включении запускаются протоколы PIM-SM и IGMP.

27.2.2 Configure multicast routing table capacity (Конфигурация размера таблицы маршрутизации)

Команды конфигурации таблицы в режиме global configuration mode:

Команда установки размера и порогового значения таблицы:

IP multicast route limit < limit > [threshold]

Команда восстановления значений по умолчанию:

no IP multicast route limit

Параметр < limit > устанавливает размер таблицы multicast routing table, изменяется в пределах 1-2147483647. Параметр [threshold] определяет пороговое значение тревоги alarm для таблицы, изменяется в пределах 1-2147483647, данный параметр является опциональным. При превышении порогового значения alarm threshold, коммутатор формирует специальное сообщение prompt message. Значение по умолчанию: 2147483647.

27.2.3 Configure TTL value of multicast interface (Конфигурация значения TTL multicast интерфейса)

Команды конфигурации в режиме interface configuration mode:

Команда установки значения TTL multicast интерфейса:

IP multicast TTL threshold < threshold >

Команда отключения TTL multicast интерфейса:
no IP multicast TTL threshold

Команда восстановления значения TTL multicast интерфейса не предусмотрена.

Параметр < threshold > определяет значение TTL multicast интерфейса, изменяется в пределах 0-255 (значение 256 является некорректным). Команда конфигурации значения TTL интерфейса используется для обслуживания списка выходного интерфейса таблицы multicast маршрутизации. Если параметр TTL находится в пределах 0-255, это значение будет передано актуальному multicast интерфейсу. Если указано значение 256, то multicast интерфейс использует значение TTL по умолчанию - 1. Значение по умолчанию: 256.

27.2.4 Start interface PIM SM function (Запуск функции PIM SM на интерфейсе)

Команды запуска/отключения функции в режиме interface configuration mode:

Команда запуска функции PIM SM Protocol:
IP PIM spark mode

Команда отключения функции PIM SM Protocol:
no IP PIM spark mode

По умолчанию протокол PIM SM на интерфейсе отключен. При включении протокола PIM SM на интерфейсе возможно отправление и получение сообщений протокола и участие в процессе выбора RP (при конфигурации bootstrapping). При запуске протокола PIM SM на интерфейсе автоматически запускается протокол IGMP protocol (нет необходимости вводить команду запуска IP IGMP).

27.2.5 Configure passive mode of interface (Конфигурация пассивного режима интерфейса)

Команды включения/отключения пассивного режима интерфейса в режиме interface configuration mode:

Команда включения пассивного режима:

IP PIM spark mode passive

Команда отключения пассивного режима:

no IP PIM spark mode passive

По умолчанию протокол PIM SM отключен. При включении протокола PIM SM на интерфейсе с помощью команды IP PIM spark mode возможно отправление и получение сообщений протокола. Если интерфейс находится в пассивном режиме passive mode, то интерфейс не отправляет сообщения протокола, не участвует в процессе регистрации, не принимает участие в процессе выбора RP, а только осуществляет мониторинг сообщений протокола protocol messages.

27.2.6 Configure interface priority (Конфигурация приоритета интерфейса)

Команды конфигурации в режиме interface configuration mode:

Команда установки приоритета при участии интерфейса в процессе выбора Dr election:

IP PIM Dr priority < priority >

Команда возврата приоритета к значению по умолчанию:

no IP PIM Dr priority [priority]

Параметр < priority > определяет приоритет интерфейса при участии в процессе выбора Dr election, изменяется в пределах 0-4294967294. Если в локальной сети находятся несколько коммутаторов, один из них выбирается как Dr с помощью протокола Hello protocol, который отвечает за отправление сообщения join pruning message в вышестоящий поток данных upstream. В процессе выбора DR, в первую очередь учитывается приоритетность (выбирается высшая). В случае одинакового значения приоритета, как Dr выбирается коммутатор с большим IP адресом. Значение по умолчанию 1.

27.2.7 Configure interface Hello message (Конфигурация сообщения Hello message)

Команды конфигурации сообщения в режиме interface configuration mode:

Команда конфигурации исключения параметра genid domain из сообщения Hello message:

IP PIM exclude genid

Команда конфигурации включения параметра genid domain в сообщение Hello message:

no IP PIM exclude genid

По умолчанию сообщение интерфейса имеет конфигурацию включающую параметр hello genid domain. Тип 20 в списке опций сообщения Hello message – значение сгенерированного ID, который не имеет значения подписи 32-bit, сгенерированного случайным образом при запуске интерфейса. Оно используется для быстрой идентификации при перезапуске соседнего устройства и уменьшения задержки установки информации RP для перезапуска роутера. Коммутатор удерживает значение genid value в сообщении Hello message и обновляет значение после получения нового сообщения Hello message. Если genid изменился, считается что соседнее устройство перезапустилось и нуждается в своевременном получении релевантной информации, вместо ожидания окончания отсчета таймера до начала самостоятельного сбора информации, что уменьшает необходимое время для восстановления при изменении сети.

27.2.8 Configure interface Hello timer interval (Конфигурация интервала Hello таймера)

Команды конфигурации таймера в режиме interface configuration mode:

Команда ввода временного интервала таймера интерфейса:

IP PIM Hello interval < interval >

Команда возврата временного интервала к значению по умолчанию:

no IP PIM Hello interval

Параметр `< interval >` определяет временной интервал отправки сообщений Hello message интерфейсом, изменяется в пределах 1-65535. Значение интервала Hello timer влияет на величину holdtime value, которая должна превышать интервал Hello interval в 3.5 раза. Если значение интервала holdtime не сконфигурировано или его значение меньше чем hello interval, то следует установить его с учетом значения Hello interval. Значение по умолчанию 30 сек.

27.2.9 Configure the holding time of neighbors on the interface (Конфигурация holding time интерфейса)

Команды конфигурации времени удержания в режиме interface configuration mode:

Команда установки времени отправления сообщений Hello messages:

IP PIM Hello holdtime `< value >`

Команда возврата временного интервала к значению по умолчанию:

no IP PIM Hello holdtime

Параметр `< value >` определяет время существования соседних устройств, изменяется в пределах 1-65535. Если сообщение Hello message не получено в течении периода hold time, считается что соседние устройства более не существуют. Значение параметра должно превышать временной интервал Hello interval в 3.5 раза. Если значение параметра меньше чем 3.5 Hello interval, то Hello interval будет использован как значение времени удержания hold time. Значение по умолчанию: 105 сек (3.5 Hello interval).

27.2.10 Configure neighbor list filtering of the interface (Конфигурация списка фильтрации интерфейса)

Команды конфигурации списка фильтрации соседних устройств на интерфейсе в режиме interface configuration mode:

Команда конфигурации списка фильтрации соседних устройств ACL интерфейса:

IP PIM neighbor filter {`< ACL ID >` | `< ACL name >`}

Команда сброса списка фильтрации соседних устройств ACL интерфейса:

```
no IP PIM neighbor filter {< ACL ID > | < ACL name >}
```

Параметр < ACL ID > указывает номер списка стандартного контроля доступа, изменяется в пределах 1 – 99. Параметр < ACL name > указывает имя списка стандартного контроля доступа. Функция ACL используется для фильтрации по списку соседних устройств интерфейса. Когда соседнее устройство отфильтровывается функцией ACL, оно удаляется из списка фильтрации интерфейса. Более подробная информация о настройке функции приведена в разделе ACL configuration. По умолчанию функция фильтрации соседних устройств neighbor filtering ACL отключена.

27.2.11 Configure the source address of unicast registration message (Конфигурация адреса источника unicast registration message)

Команды конфигурации адреса источника сообщения в режиме global configuration mode:

Команда конфигурации адреса источника или интерфейса сообщения регистрации unicast registration:

```
IP PIM register source {< IP > | < if name >}
```

Команда сброса конфигурации адреса источника сообщения:

```
no PIM register source
```

Параметр < IP > устанавливает адрес источника сообщения регистрации unicast registration message, вводится в десятичном формате с разделением точками.

Параметр < If name > указывает имя интерфейса-источника сообщений unicast registration message, который является интерфейсом уровня 3 (eg: vlan2). Адрес используется для регистрации источника (отправителя) сообщений.

По умолчанию адрес источника сообщений не установлен, используется адрес интерфейса принявшего пакет данных.

27.2.12 Configure the number limit of registered message (Конфигурация предельного количества зарегистрированных сообщений)

Команды конфигурации предельного количества зарегистрированных сообщений в режиме global configuration mode:

Команда конфигурации порогового значения количества зарегистрированных сообщений полученных в течении 1 сек:

IP PIM register rate limit < limit >

Команда сброса порогового значения количества зарегистрированных сообщений:

no IP PIM register rate limit

Параметр < limit > устанавливает максимальное число зарегистрированных сообщений которые могут быть получены в течении 1 сек, изменяется в пределах 1-65535. Сообщение registration message обычно инкапсулирует пакет данных. При регистрации большой поток данных достигает первого роутера источника. В это время поток данных протокола PIM SM Protocol может быть ограничен для предотвращения перегрузки роутера. Используется специальный механизм ограничения, т.к. протоколу PIM SM достаточно только наличие потока данных для запуска создания SPT от источника до RP. При использовании ограничения потока данных, производится подсчет количества пакетов в 1 сек. Если временной интервал между двумя пакетами данных превышает 1 сек, то ограничение не применяется. По умолчанию ограничение количества зарегистрированных пакетов не установлено (значение порога 0).

27.2.13 Check RP when configuring registration (Проверка RP при конфигурации регистрации сообщений)

Команда проверки доступности RP в режиме global configuration mode:

IP PIM register RP reachability

Команда отмены проверки доступности RP:

no IP PIM register RP reachability

По умолчанию функция проверки доступности RP отключена.

27.2.14 Configure registration inhibit timer (Конфигурация таймера подавления регистрации)

Команды конфигурации таймера подавления регистрации в режиме global configuration mode:

Команда установки значения времени таймера подавления:

IP PIM register suppression < value >

Команда отключения таймера подавления:

no IP PIM register suppression

Параметр < value > определяет значение времени таймера, изменяется в пределах 1-65535. По умолчанию установлено значение 60 сек.

27.2.15 Configure and register Kat timer (Конфигурация Kat таймера регистрации)

Команды конфигурации таймера в режиме global configuration mode:

Команда установки значения времени Kat таймера регистрации для RP:

IP PIM RP register Kat < value >

Команда отключения таймера Kat:

no IP PIM RP register Kat

По умолчанию значения восстановления Kat таймера (s, g). Параметр < value > определяет значение времени таймера Kat timer (s, g) созданное источником RP, изменяется в пределах 1-65535. По умолчанию установлено значение 185 сек (3 * register-suppression + register-probe).

27.2.16 Registration address configuration (Конфигурация регистрации адреса)

Команды конфигурации регистрации адреса в режиме global configuration mode:

Команда конфигурации списка фильтрации источников регистрационных сообщений:

IP PIM accept register list {< acl-id1 > | < acl-id2 > | < ACL name >}

Команда отключения списка фильтрации:

no IP PIM accept register

Параметр < Acl-id1 > определяет размер расширенного списка контроля доступа, изменяется в пределах 100-199.

Параметр < Acl-id2 > определяет размер расширенной части списка контроля доступа, изменяется в пределах 2000-2699.

Параметр < ACL name > указывает имя списка контроля доступа. Список контроля доступа используется для фильтрации адреса источника зарегистрированного в RP. Если адрес источника RP отфильтрован ACL, соответствующие (s, g) входы будут удалены из таблицы маршрутизации multicast routing table. Более подробная информация о настройке функции приведена в разделе ACL configuration. По умолчанию функция регистрации источника ACL отключена.

27.2.17 Configure the checksum of the registered message in Cisco mode (Конфигурация проверки контрольной суммы в режиме Cisco)

Команда конфигурации контрольной суммы зарегистрированного сообщения в режиме global configuration mode:

IP PIM Cisco register checksum [group list {< ACL Id1 > | < ACL Id2 > | < ACL name >}]

Команда конфигурации контрольной суммы незарегистрированного сообщения:

no IP PIM Cisco register checksum [group list {< ACL Id1 > | < ACL Id2 > | < ACL name >}]

Параметр < Acl-id1 > определяет размер стандартного списка контроля доступа, изменяется в пределах 1-99.

Параметр < Acl-id2 > определяет размер расширенной части списка контроля доступа, изменяется в пределах 1300-1999.

Параметр < ACL name > указывает имя стандартного списка контроля доступа. Контрольная сумма в режиме Cisco mode вычисляет контрольную сумму всего регистрационного сообщения (включая пакет данных). Контрольная сумма основного регистрационного сообщения включает только заглавную часть сообщения протокола PIM SM protocol

и заглавную часть регистрационного сообщения. Если список контроля доступа сконфигурирован, протоколу ACL требуется отфильтровать соответствующие группы адресов при отправлении регистрационного сообщения. Контрольная сумма Cisco может использоваться протоколом ACL, но основная контрольная сумма не может быть использована протоколом ACL. По умолчанию контрольная сумма Cisco не сконфигурирована.

27.2.18 Configure static RP address (Конфигурация статического RP адреса)

Команда конфигурации статического RP в режиме global configuration mode:

```
IP PIM RP address < IP > [< acl-id1 > | < acl-id2 > | < ACL name >]
```

Команда сброса статического RP:

```
no IP PIM RP address < IP > [< acl-id1 > | < acl-id2 > | < ACL name >]
```

Параметр < IP > устанавливает адрес статического RP, вводится в десятичном формате с разделением точками.

Параметр < acl-id1 > определяет размер стандартного списка контроля доступа, изменяется в пределах 1-99.

Параметр < Acl-id2 > определяет размер расширенной части списка контроля доступа, изменяется в пределах 1300-1999.

Параметр < ACL name > указывает имя стандартного списка контроля доступа. По умолчанию статический RP не сконфигурирован.

27.2.19 Configure candidate RP (Конфигурация кандидата RP)

Команды конфигурации кандидатов RP в режиме global configuration mode:

Команда конфигурации интерфейса c-rp:

```
IP PIM RP candidate < if name > [group list {< ACL ID > | < ACL name >}  
| interval < value1 > | priority < Value2 >]
```

Команда сброса атрибута C-RP:

```
no IP PIM RP candidate [if name]
```

Параметр `< if name >` указывает имя интерфейса сконфигурированного как C-RP, который является интерфейсом уровня 3 (eg: `vlan2`).

Параметр `< ACL ID >` указывает номер стандартного списка контроля доступа, изменяется в пределах 1-99.

Параметр `< ACL name >` указывает имя стандартного списка контроля доступа.

Параметр `< Value1 >` определяет величину временного интервала периода от C-RP `unicast c-rp-adv` сообщения к BSR, изменяется в пределах 1-16383.

Параметр `< Value2 >` устанавливает приоритет C-RP, изменяется в пределах 0-255. Он включен в информационные сообщения `bootstrap`, распространяется по домену для выбора RP, соответствующего группе. По умолчанию значение параметра `value1` составляет 60сек, значение параметра `value2` 192.

27.2.20 Configure ignore RP set priority (Конфигурация приоритета отклонения RP)

Команда установки приоритета отклонения (при поиске RP используется алгоритм `hash algorithm mapping`) в режиме `global configuration mode`:

```
IP PIM ignore RP set priority
```

Команда отмены установки приоритета отклонения RP:

```
no IP PIM ignore RP set priority
```

По умолчанию в установках RP включен порядок поиска приоритета отклонения RP.

27.2.21 Configure c-rp-adv message in Cisco mode (Конфигурация c-rp-adv сообщения в режиме Cisco)

Команда конфигурации формата сообщения `c-rp-adv` для Cisco BSR в режиме `global configuration mode`:

```
IP PIM CRP Cisco prefix
```

Команда конфигурации формата сообщения c-rp-adv для стандарта RFC:

```
no IP PIM CRP Cisco prefix
```

По умолчанию сообщение c-rp-adv message имеет формат для стандарта RFC. В этом формате Cisco BSR не принимает сообщения c-rp-adv messages. Префикс CNT домена - 0. Если по умолчанию отсутствует групповой адрес, то префикс CNT домена - 1. Групповой адрес - 224.0.0.0. По умолчанию стандарт RFC имеет префикс 0 если поле префикса CNT не содержит групповой адрес. Команда конфигурации совместима с BSR. Роутерам Cisco необходимо отправлять сообщения c-rp-adv, которые могут быть распознаны Cisco BSR.

27.2.22 Configure candidate BSR (Конфигурация кандидата BSR)

Команда конфигурации интерфейса, приоритета и вычисления длины маски хэша кандидата BSR в режиме global configuration mode:

```
IP PIM BSR candidate < if name > [< hash mask len > | < priority >]
```

Команда сброса идентификации кандидата BSR интерфейсом:

```
no IP PIM BSR candidate [if name]
```

Параметр < if name > указывает имя интерфейса кандидата BSR, который является интерфейсом уровня 3 (eg: vlan2).

Параметр < Hash mash len > определяет длину маски в алгоритме hash algorithm, изменяется в пределах 0-32.

Параметр < Priority > устанавливает приоритет кандидата BSR в процессе выбора BSR, изменяется в пределах 0-255. Когда кандидат BSR выбирается BSR, длина его маски и приоритет передаются в сообщении bootstrap message. По умолчанию значение приоритета 0, значение параметра hash mask len 10.

27.2.23 Configure JP timer interval (Конфигурация JP таймера)

Команда установки времени join pruning таймера в режиме global configuration mode:

IP PIM JP timer < time >

Команда установки времени join pruning таймера на значение по умолчанию:

no IP PIM JP timer [time]

Параметр < time > определяет временной интервал периодичности отправки сообщений pruning message коммутатором, изменяется в пределах 1-65535. Коммутатор отслеживает статус развертывания multicast дерева распределения с помощью JP таймера. Временной интервал JP таймера должен в 3.5 превышать время holding time, информация о времени удержания holding time содержится в сообщении join pruning message. По умолчанию установлено время таймера 60 сек.

27.2.24 Configure SPT switching (Конфигурация переключения SPT)

Команда конфигурации переключения SPT в режиме global configuration mode:

IP PIM SPT threshold [group list <ACL Id1> | <ACL Id2> | < ACL name >]

Команда сброса SPT переключения:

no IP PIM SPT threshold [group list < ACL Id1 > | < ACL Id2 > | < ACL name >]

Параметр < ACL Id1 > определяет размер стандартного списка контроля доступа, изменяется в пределах 1-99.

Параметр < Acl-id2 > определяет размер расширенной части списка контроля доступа, изменяется в пределах 1300-1999.

Параметр < ACL name > указывает имя стандартного списка контроля доступа. Когда функция переключения SPT switching включена, список контроля доступа используется для фильтрации (*, g) таблицы маршрутизации multicast routing table. Функция переключения SPT может быть сконфигурирована через ACL. По умолчанию функция переключения SPT switching отключена.

27.2.25 Configure SSM (Конфигурация SSM)

Команда конфигурации SSM (source specific multicast) в режиме global configuration mode:

```
IP PIM SSM {default | range {< ACL ID > | < ACL name >}}
```

Команда отключения SSM:

```
no IP PIM SSM
```

Параметр < ACL ID > указывает номер стандартного списка контроля доступа, изменяется в пределах 1-99.

Параметр < ACL name > указывает имя стандартного списка контроля доступа. По умолчанию функция SSM отключена.

27.2.26 Configure multicast security (Конфигурация multicast безопасности)

Команда конфигурации в режиме global configuration mode:

```
IP mroute < SRC addr > < GRP addr > < IIF name > < OIF name >  
remove < if name >
```

Для multicast маршрута (s, g), определенный порт уровня layer-2 в соответствующем VLAN выходного интерфейса не перенаправляет поток данных.

Команда добавления выходного порта в маршрут multicast route (s, g) выходного интерфейса соответствующего VLAN:

```
IP mroute < SRC addr > < GRP addr > < IIF name > < OIF name > Add <  
if name >
```

Параметр < SRC addr > указывает адрес источника multicast source в маршруте multicast route (s, g).

Параметр < GRP addr > представляет адрес группы multicast group в маршруте multicast routing (s, g).

Параметр < IIF name > указывает входной интерфейс в маршруте multicast routing (s, g).

Параметр < OIF name > указывает выходной интерфейс в маршруте multicast routing (s, g).

Параметр < If name > указывает что выходной интерфейс в маршруте multicast routing (s, g) соответствует порту уровня layer-2 VLAN.

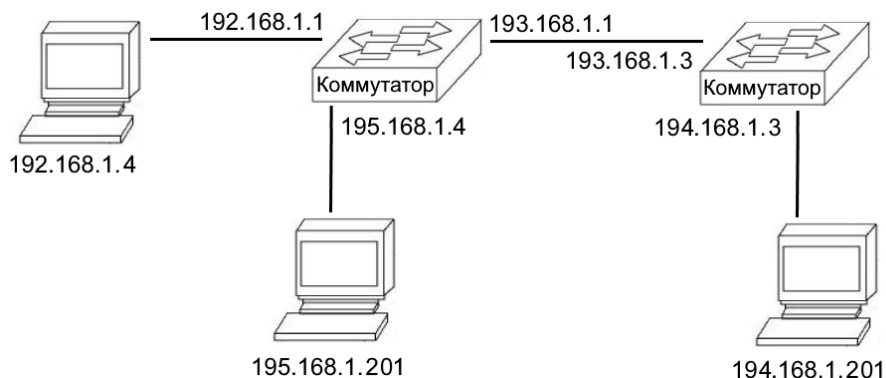
Команда `remove command` экранирует порт уровня layer-2 от VLAN в котором находится выходной интерфейс, т.о. поток данных выходного интерфейса не будет перенаправлен портом уровня layer-2.

Команда `add command` добавляет порт уровня layer-2 к VLAN в котором находится выходной интерфейс, т.о. поток данных от выходного интерфейса будет перенаправлен портом уровня layer-2.

По умолчанию все порты VLAN уровня layer-2 соответствующие выходному интерфейсу будут направлять потоки данных multicast.

27.3 PIM-SM configuration example (Пример конфигурации)

Конфигурация соединения оборудования представлена на рисунке ниже:



Команды запуска протокола PIM-SM для обеспечения возможности принимающему хосту получать поток данных multicast от хоста источника:

Команды конфигурации коммутатора A:

```
Switch#configure terminal
Switch(config)#ip multicast-routing
Switch(config)#ip pim rp-address 193.168.1.3
Switch(config)#interface vlan2
Switch(config-vlan2)#ip address 192.168.1.1/24
Switch(config-vlan2)#ip pim sparse-mode
```

```
Switch(config-vlan2)#interface vlan3
Switch(config-vlan3)#ip address 193.168.1.1/24
Switch(config-vlan3)#ip pim sparse-mode
Switch(config-vlan3)#interface vlan5
Switch(config-vlan5)#ip address 195.168.1.1/24
Switch(config-vlan5)#ip pim sparse-mode
```

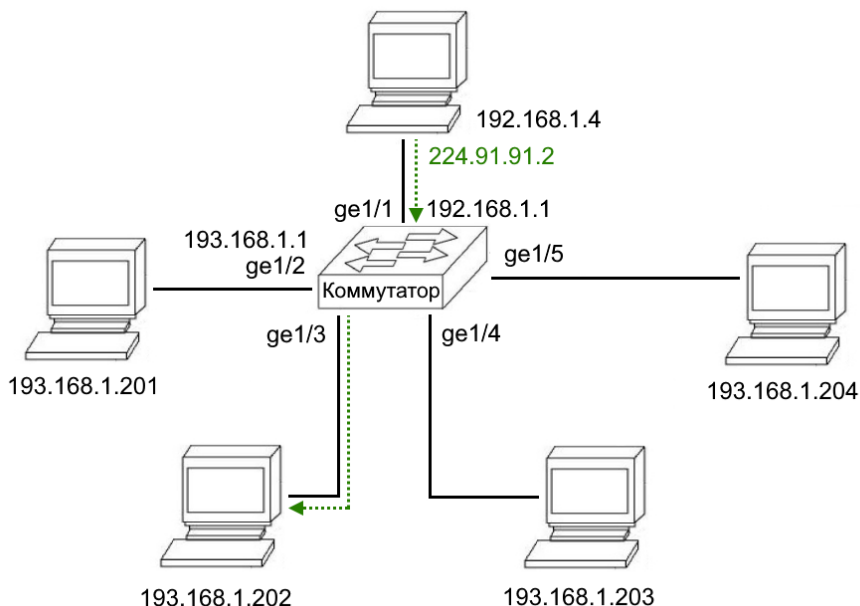
Команды конфигурации коммутатора B:

```
Switch#configure terminal
Switch(config)#ip multicast-routing
Switch(config)#ip pim rp-address 193.168.1.3
Switch(config)#interface vlan3
Switch(config-vlan3)#ip address 193.168.1.3/24
Switch(config-vlan3)#ip pim sparse-mode
Switch(config-vlan3)#interface vlan4
Switch(config-vlan4)#ip address 194.168.1.3/24
Switch(config-vlan4)#ip pim sparse-mode
```

Команды просмотра конфигурации PIM-SM:

```
show ip pim sparse-mode interface
show ip pim sparse-mode neighbor
show ip pim sparse-mode local-members
show ip pim sparse-mode mroute
show ip pim sparse-mode rp mapping
```

Функция безопасности multicast security function module обеспечивается протоколом маршрутизации multicast routing protocol. После запуска протокола маршрутизации multicast routing protocol, проводится конфигурация статуса forwarding state порта уровня layer-2 выходного интерфейса согласно особенностям пути маршрутизации. Как показано на рисунке ниже, vlan2 соединен с источником multicast source 192.168.1.4. Vlan3 включает 4 хоста, которые соединены с портами GE1/2, GE1/3, GE1/4 и GE1/5 соответственно. Как видно на рисунке, только хост 193.168.202 отвечает на запросы multicast группы 224.91.91.2, а три другие хоста не участвуют в запросах multicast on-demand, т.о. поток данных направленный от трех портов уровня layer-2 не является необходимым.



Команды конфигурации коммутатора для обеспечения защиты портов портов GE1/2, GE1/4 и GE1/5:

```
Switch#configure terminal
Switch(config)#ip multicast-routing
Switch(config)#ip pim rp-address 193.168.1.1
Switch(config)#interface vlan2
Switch(config-vlan2)#ip address 192.168.1.1/24
Switch(config-vlan2)#ip pim sparse-mode
Switch(config-vlan2)#interface ge1/1
Switch(config-ge1/1)#switchport access vlan 2
Switch(config-ge1/1)#interface vlan3
Switch(config-vlan3)#ip address 193.168.1.1/24
Switch(config-vlan3)#ip pim sparse-mode
Switch(config-vlan3)#interface ge1/2
Switch(config-ge1/2)#switchport access vlan 3
Switch(config-ge1/2)#interface ge1/3
Switch(config-ge1/3)#switchport access vlan 3
Switch(config-ge1/3)#interface ge1/4
Switch(config-ge1/4)#switchport access vlan 3
Switch(config-ge1/4)#interface ge1/5
```

```
Switch(config-ge1/5)#switchport access vlan 3
Switch(config-ge1/5)#exit
Switch(config)#ip mroute 192.168.1.4 224.91.91.2 vlan2 vlan3 remove
ge1/2
Switch(config)#ip mroute 192.168.1.4 224.91.91.2 vlan2 vlan3 remove
ge1/4
Switch(config)#ip mroute 192.168.1.4 224.91.91.2 vlan2 vlan3 remove
ge1/5
```

28. RIP Configuration (Конфигурация RIP)

Глава описывает конфигурацию протокола RIP (routing information protocol) – протокола динамической маршрутизации. Протокол использует алгоритм distance vector algorithm и в основном применяется в небольших сетях. Глава содержит следующие разделы:

- Introduction (Введение)
- Configuration RIP (Конфигурация протокола)
- Configuration example (Пример конфигурации)

28.1 RIP introduction (Введение)

RIP (routing information protocol) – протокола динамической маршрутизации. Протокол использует алгоритм distance vector algorithm и в основном применяется в небольших сетях.

Сообщения протокола RIP инкапсулированы в сообщениях UDP и используют UDP port 520. Главный принцип протокола RIP – использование hop для измерения расстояния до хоста и увеличение количества hop на единицу для каждого роутера, и таким образом производить вычисление длины и выбора пути. Протокол RIP поддерживает до 15и hops (максимальное количество), а значение 16 hops отмечается как недостижимое. Rip использует метод передачи всей таблицы маршрутов для синхронизации маршрутной информации роутеров в сети, сообщения содержащие маршрутную информацию обновляются каждые 30 сек. Если сообщение с обновленной маршрутной информацией от соседнего устройства не получено в

течении 180 сек, то оно отмечается как недоступное. Если сообщение с обновленной маршрутной информацией от соседнего устройства не получено в течении 120 сек, то путь будет удален.

Работа протокола Rip основана на простых принципах и поэтому его не сложно использовать, но его применение может вызвать определенные проблемы из-за возможного возникновения петлевых соединений. Для предотвращения возникновения петлевых соединений протокол RIP использует механизм горизонтальной сегментации (horizontal segmentation mechanism) чтобы обеспечить корректное взаимодействие роутеров. Горизонтальная сегментация предполагает, что принимающий интерфейс не обнарудует обновления маршрутов. Т.е. горизонтальная сегментация с инверсией публикует обновления маршрутов со стороны принимающего интерфейса с маркировкой их веса как недоступный (unreachable), т.о. соседний роутер може быстро обнаружить петлевое соединение не дожидаясь информации о весе маршрута.

Таблица маршрутизации должна включать адрес назначения (хоста или сети), адрес следующего hop, отправляющего интерфейса, вес маршрута, таймер (таймер сбрасывается при получении обновления маршрутизации) и отметку маршрутизации.

Одновременно с запуском протокола RIP, отправляется сообщение с запросом полной таблицы маршрутизации (full table request message) в форме broadcast (RIP-1) или multicast (RIP-2). Связанные роутеры отправляют свои таблицы маршрутизации в ответ на сообщение запроса request message. После получения ответного сообщения response message, роутер последовательно обрабатывает маршруты и преобразует свою собственную маршрутную таблицу, при появлении нового маршрута, срабатывает триггер сообщения об обновлении. После завершения ряда процессов обновления, протокол RIP поводит финальную обработку, и каждый роутер в сети получает самую последнюю информацию о маршрутизации. После стабилизации сети, протокол RIP начинает осуществлять рассылку таблицы маршрутизации соседним устройствам каждые 30 сек. Каждый роутер обрабатывает свою собственную маршрутную информацию в соответствии с полученными сообщениями об обновлении и выбирает самый оптимальный маршрут. Протокол RIP использует механизм временной задержки для обработки маршрутов которые не обновлялись длительное время для обеспечения корректности

маршрутов в реальном времени.

Протокол RIP наиболее часто используется в локальных сетях организаций и региональных сетях с простой структурой, протокол RIP не подходит для использования в крупных сетях со сложной структурой.

28.2 RIP configuration (Конфигурация RIP)

После запуска протокола RIP становится возможным конфигурировать функции и атрибуты протокола. Соответствующие команды выполняются в режимах RIP configuration mode и interface configuration mode.

Конфигурация протокола RIP включает:

- Запуск протокола RIP и режима configuration mode
- Запуск RIP интерфейса
- Конфигурация передачи unicast message
- Конфигурация рабочего статуса интерфейса
- Конфигурация весов маршрутов по умолчанию
- Конфигурация дистанции управления
- Конфигурация таймера
- Конфигурация версии
- Введение внешней маршрутизации
- Конфигурация фильтрации маршрутизации
- Конфигурация дополнительных весов маршрутов
- Конфигурация версии RIP интерфейса
- Конфигурация статуса приема-передачи интерфейса
- Конфигурация горизонтальной сегментации
- Аутентификация сообщений
- Конфигурация веса интерфейса

28.2.1 Start RIP and enter RIP configuration mode (Запуск и режим конфигурации RIP)

Команды запуска/отключения протокола RIP в режиме global configuration mode:

Команды запуска протокола и режима конфигурации:

```
router rip
```

Команда отключения протокола:

```
no router rip
```

Протокол RIP отключен по умолчанию.

28.2.2 Enable RIP interface (Запуск RIP интерфейса)

Протокол RIP может определять некоторые интерфейсы, конфигурировать сеть как RIP, отправлять и получать сообщения протокола интерфейса RIP.

Команда запуска интерфейса в режиме configuration mode:

```
network < network address >
```

Команда отключения интерфейса:

```
no network < network address >
```

Параметр имеет два формата: a.b.c.d/m и a.b.c.d., первый формат определяет сетевой IP и длину маски, второй формат – сетевой IP и маску. По умолчанию, при запуске протокола параметр отсутствует.

После запуска протокола rip для его работы должен быть определен сегмент сети, т.е. протокол работает с интерфейсом определенного сегмента сети. Интерфейсы не относящиеся к данному сегменту сети не получают и не отправляют информацию протокола о путях и маршрутах. Интерфейсы, не относящиеся к данному сегменту сети, для протокола RIP не существуют. Параметр содержащий сетевой адрес может быть включен или отключен и назначен как IP адрес интерфейса. Соответствующие команды устанавливают интерфейс для сегмента сети с определенным адресом. Например, если IP адрес интерфейса 192.160.1.1, с помощью команд следует установить 192.160.1.1/24. Команда для просмотра полученной конфигурации сети 192.160.1.0/24: show running config.

28.2.3 Configure unicast message transmission (Конфигурация передачи unicast сообщений)

Версия 1 RIP протокола использует формат broadcast для обмена сообщениями, версия 2 использует формат multicast (224.0.0.9) для обмена сообщениями. Если соединение не поддерживает broadcast при запущенном протоколе Rip, то для обмена сообщениями необходимо определить специфический unicast адрес.

Команды конфигурации peer unicast IP адреса в режиме RIP configuration mode:

```
neighbor < IP address >
```

Команда сброса peer unicast IP адреса:

```
no neighbor < IP address >
```

Параметр IP address определяет unicast IP адрес. По умолчанию протокол RIP не отправляет сообщения на какой-либо unicast адрес.

28.2.4 Configure the working state of the interface (Конфигурация рабочего статуса интерфейса)

При работе RIP протокола в сети дотаточно маршрутизации RIP интерфейса, но не обязательно направлять RIP маршрут на этот интерфейс. Соответствующие сетевые команды определяют отправляемые и получаемые интерфейсом сообщения RIP протокола. Команды passive interface нужны для определения маршрута интерфейса и блокировки вещания интерфейса.

Команда установки статуса passive интрфейса в режиме Rip configuration mode:

```
passive interface < if name >
```

Команда сброса статуса passive интерфейса:

```
no passive interface < if name >
```

Параметр if name указывает имя интерфейса (например vlan1, vlan2...). По умолчанию все включенные RIP интерфейсы не находятся в статусе passive.

28.2.5 Configure default routing weights (Конфигурация веса маршрута)

При указании внешнего маршрута, необходимо определить его вес (весовой коэффициент). Если вес маршрута не определен, то используется значение веса по умолчанию.

Команда установки веса пути по умолчанию в режиме Rip configuration mode:

```
default metric < metric >
```

Команда отмены установки веса пути по умолчанию:

```
no default metric [Metric]
```

При восстановлении внешнего пути, по умолчанию используется значение веса 1.

Параметр metric value изменяется в пределах 1-16 (более чем 1 и менее чем 16). По умолчанию значение параметра metric value 1. Для установки значения по умолчанию предусмотрена команда default metric command.

28.2.6 Configure management distance (Конфигурация дистанции управления)

Каждый протокол имеет приоритет соглашения agreed priority. Дистанция управления это приоритет маршрута при реализации стратегии маршрутизации. При наличии двух одинаковых маршрутов (по двум разным протоколам) до одной точки назначения, предпочтение отдается более короткому пути management distance.

Команда установки значения параметра distance в режиме Rip configuration mode:

```
distance < distance >
```

Команда установки значения параметра по умолчанию:

```
no distance [distance]
```

Параметр value distance изменяется в пределах 1-255

По умолчанию значение параметра distance value 120. Для установки значения по умолчанию предусмотрена команда no distance command.

28.2.7 Configure timer (Конфигурация таймера)

Протокол Rip имеет три таймера. Первый (update) таймер отвечает за рассылку полной таблицы маршрутизации по всем RIP интерфейсам каждые 30 сек. Второй (timeout) таймер отслеживает периодичность получения обновленных таблиц маршрутизации для каждого маршрута каждые 180 сек (отмечаются значением metric 16). Третий (garbage) таймер отвечает за удаление из таблицы маршрутизации путей со значением metric 16 не получивших обновление в течении 120 сек.

Команда установки значений времени таймеров в режиме Rip configuration mode:

```
timers basic < update > < timeout > < garbage >
```

Команда установки значений времени таймеров по умолчанию:

```
no timers basic
```

Параметр < update > определяет периодичность обновления полной таблицы маршрутизации. Параметр < timeout > участвует в определении путей маршрутизации не получивших обновления. Параметр < garbage > участвует в определении путей маршрутизации отмеченных как недоступные и подлежащих удалению. Время таймеров изменяется в пределах 231-1.

По умолчанию значение параметра update 30 сек, значение параметра timeout 180 сек (по истечении которого путь отмечается как недоступный invalid), значение параметра garbage 120 сек (по истечении которого недоступный путь удаляется).

28.2.8 Configuration version (Версия конфигурации)

В настоящее время доступны две версии Rip протокола: version 1 (rfc1058) и version 2 (rfc2453). Текущая версия прописана в специальном поле сообщения протокола.

Команда установки версии протокола в режиме Rip configuration mode:

```
version < version >
```

Команда установки версии протокола по умолчанию:

```
no version [version]
```

Параметр version может иметь значения 1 или 2.
По умолчанию установлена version 2.

28.2.9 Introducing external routing (Введение внешней маршрутизации)

Протокол Rip позволяет пользователям использовать маршрутную информацию других протоколов в таблице маршрутизации. Протокол Rip может использовать данные протоколов connected, static, OSPF, IS-IS и BGP.

Команды конфигурации в режиме Rip configuration mode:

```
redistribute {kernel | connected | static | OSPF | Isis | BGP} [Metric < metric > | route map < route map name >]
```

Команды сброса конфигурации:

```
no redistribute {kernel | connected | static | OSPF | Isis | BGP} [Metric < metric > | route map < route map name >]
```

Первый параметр указывает имя импортируемого протокола (включая прямое соединение connection, статическое static, OSPF, is is и BGP). Второй параметр определяет вес при импорте, изменяется в пределах 1-16. Третий параметр указывает имя связанной маршрутной карты route map. Маршрутная карта конфигурируется в режиме global configuration mode (см. раздел command manual).

По умолчанию функция использования внешней маршрутизации протоколом Rip отключена.

28.2.10 Configure routing filtering (Конфигурация фильтра маршрутизации)

Протокол Rip поддерживает функцию фильтрации маршрутизации. Функция обеспечивает фильтрацию получаемых и отправляемых маршрутов с помощью специальной таблицы контроля доступа access control list, таблицы префикса адресов address prefix list, и настройки политики правил policy rules.

Команды конфигурации функции фильтрации в режиме Rip configuration mode:

Команда включения таблицы доступа для контроля входа и выхода интерфейса:

```
distribute list < ACL name > {in | out} [if name]
```

Команда отключения таблицы контроля доступа:

```
no distribute list < ACL name > {in | out} [if name]
```

Параметр < ACL name > указывает имя соответствующей таблицы доступа. Параметр [if name] указывает Rip интерфейс, к которому применяется фильтрация. Значения {in | out} определяют принимаемые и отправляемые маршруты (receiving или publishing routes).

Команда включения таблицы префиксов prefix list:

```
distribute list prefix < pre name > {in | out} [if name]
```

Команда отключения таблицы префиксов:

```
no distribute list prefix < pre name > {in | out} [if name]
```

Параметр < pre name > указывает имя соответствующей таблицы префиксов. Параметр [if name] указывает Rip интерфейс, к которому применяется фильтрация. Значения {in | out} определяют принимаемые и отправляемые маршруты (receiving или publishing routes).

По умолчанию Rip протокол не фильтрует какие-либо принятые или отправленные маршруты.

Таблицы доступа Access list и префиксов prefix list конфигурируются в режиме global configuration mode (см. раздел command manual).

28.2.11 Configure additional routing weight (Конфигурация дополнительного веса маршрута)

Дополнительный вес (весовой коэффициент) маршрута это предустановленное добавочное значение к весу пути Rip протокола на входе и выходе. Дополнительный вес не изменяет непосредственно вес маршрута в маршрутной таблице, но добавляет предустановленную величину к весу маршрута при его получении или отправке интерфейсом.

Команда ввода дополнительного веса в интерфейс при отправке-получении маршрута в режиме Rip configuration mode:

```
offset list < ACL name > {in | out} < offset > [if name]
```

Команда сброса дополнительного веса в интерфейсе при отправке-получении маршрута:

```
no offset list < ACL name > {in | out} < offset > [if name]
```

Параметр < ACL name > указывает имя таблицы доступа интерфейса. Значения {in | out} определяют добавление дополнительного веса к принимаемым и отправляемым маршрутам. Параметр < offset > определяет значение дополнительного веса, изменяется в пределах 0-16. Параметр [if name] указывает Rip интерфейс, на котором применяется добавление дополнительного веса.

По умолчанию при получении сообщения, значение дополнительного веса для каждого пути 1, при отправке сообщения значение дополнительного веса для каждого пути 0.

28.2.12 Configure RIP version of interface (Конфигурация версии RIP интерфейса)

Протокол RIP имеет две версии сообщений: RIP-1 и RIP-2. Используемая протоколом версия может быть сконфигурирована на интерфейсе. На принимающей стороне могут быть использованы три конфигурации: только сообщения версии RIP-1, только сообщения версии RIP-2 или обе версии сообщений RIP-1 и RIP-2 одновременно. На отправляющей стороне могут быть использованы следующие конфигурации: отправка сообщений RIP-1, отправка сообщений RIP-2 (в режиме broadcast mode), отправка сообщений RIP-2 (в режиме multicast mode), отправка сообщений RIP-1 RIP-2 одновременно. Сообщения RIP-2 имеют два типа отправки: broadcast и multicast. Использование multicast не может оградить хост (в той же сети) не использующий RIP протокол от получения сообщений RIP протокола, но защищает хост использующий RIP-1 от неправильной обработки маршрутов RIP-2 с маской подсети.

Команда конфигурации интерфейса для приема сообщений только одной версии в режиме interface configuration mode:

```
IP rip receive {version 124eip}
```

Параметр: Version 1 или version 2.

Команда конфигурации интерфейса для приема сообщений двух версий (version 1 и version 2):

IP rip receive version {1 2 | 2 1}

Параметр: 1 2 или 2 1 (допустимы оба варианта).

Команда конфигурации интерфейса по умолчанию для принимаемых сообщений:

no IP rip receive version [1 | 2 | 1 2 | 2 1]

По умолчанию установлено: version 2 multicast mode.

Команда конфигурации интерфейса для отправки сообщений только одной версии:

IP rip send version {1 | 2 | 1-compatible}

Параметр: Version 1 или version 2; 1-compatible означает, что интерфейс version 2 отправляет сообщения совместимые с version 1 (сообщения broadcast чаще чем multicast).

Команда конфигурации интерфейса для отправки сообщений двух версий (version 1 и version 2):

IP rip send version {1 2 | 2 1}

Параметр: 1 2 или 2 1 (допустимы оба варианта).

Команда конфигурации интерфейса по умолчанию для отправляемых сообщений:

no IP rip send version [1 | 2 | 1-compatible | 1 2 | 2 1]

По умолчанию установлено: version 2 multicast mode.

28.2.13 Configure the transceiver status of the interface (Конфигурация статуса интерфейса)

После запуска RIP интерфейса с помощью команд в режиме rip mode, возможно установить статус для отправки и получения сообщений протокола (принимать сообщения протокола или отправлять сообщения протокола в режиме interface mode).

Команда конфигурации интерфейса для приема сообщений протокола в режиме interface configuration mode:

IP rip receive packet

Команда конфигурации интерфейса не отправлять сообщения протокола:

no IP rip receive packet

Команда конфигурации интерфейса для отправки сообщений протокола:

IP rip send packet

Команда конфигурации интерфейса не отправлять сообщения протокола:

no IP rip send packet

По умолчанию установлен статус интерфейса для получения и отправки сообщений протокола.

Внимание: Команды запускают работу RIP протокола в сети. Интерфейс в сети отправляет и принимает сообщения протокола, путь интерфейса включен в маршрутную таблицу. Пассивный статус интерфейса предполагает запрет на отправку и получение сообщений протокола интерфейсом, но при этом путь интерфейса сохраняется в маршрутной таблице. IP rip принимает пакеты и IP rip отправляет пакеты команд одновременно учитывая отправляет или принимает сообщения протокола интерфейс.

28.2.14 Configure horizontal segmentation (Конфигурация горизонтальной сегментации)

Горизонтальная сегментация означает, что путь полученный от интерфейса не был отправлен этим интерфейсом. Горизонтальная сегментация с инверсией (toxic inversion) означает что путь полученный от интерфейса отправляется этим интерфейсом, но его значение metric value - 16. Горизонтальная сегментация может до определенной степени предотвращать образование петлевых соединений. Горизонтальная сегментация с инверсией обладает большей эффективностью по сравнению с обычной горизонтальной сегментацией, при этом прямая маркировка не доступна. Обычно в сетях pVMA, необходимо запрещать горизонтальную сегментацию для выстраивания корректного маршрута.

Команда включения функции horizontal division и режима toxic inversion на интерфейсе в режиме interface configuration mode:

IP rip split horizon [poisoned]

Команда отключения функции horizontal division function на интерфейсе:
no IP rip split horizon

Параметр [no poisoned] включает функцию горизонтальная сегментация. Параметр [poisoned] включает функцию горизонтальная сегментация с инверсией (toxic inversion).

По умолчанию выбрана функция горизонтальная сегментация с инверсией (toxic inversion).

28.2.15 Message authentication (Аутентификация сообщений)

Версия протокола RIP-1 не поддерживает функцию аутентификации сообщений (message authentication), в то время как версия RIP-2 поддерживает аутентификацию сообщений. Существует два метода аутентификации сообщений: plaintext аутентификация и MD5 аутентификация. Одновременная интерпретация незашифрованных аутентификационных данных системой аутентификации plaintext не может гарантировать обеспечение безопасности и не может быть использована в сети с высокими требованиями безопасности. Установка пароля может быть разделена на обычную цепочку ключ - ключ (key и key chain). Обычный ключ сохраняет независимую строку (independent string). Цепочка ключа управляет ID, содержимым, полученным значением lifetime и отправкой lifetime ключа (см. раздел Command Reference Manual для управления цепочкой ключа).

Команда установки метода аутентификации plaintext или MD5 в режиме interface configuration mode:

IP rip authentication mode {text | MD5}

Команда сброса установленного метода аутентификации:

no IP rip authentication mode [text | MD5]

Параметр text определяет метод аутентификации plaintext. Параметр MD5 определяет метод аутентификации MD5.

По умолчанию функция аутентификации отключена.

Команда установки строки password string для аутентификации:

IP rip authentication string < password >

Команда отключения строки password string для аутентификации:

no IP rip authentication string [password]

Параметр: 16 byte authentication password.

Команда установки цепочки ключа key chain для аутентификации:

IP rip authentication key chain < key chain name >

Команда отключения цепочки ключа key chain для аутентификации:

no IP rip authentication key chain [key chain name]

Параметр < key chain name > указывает имя соответствующей цепочки ключа. Цепочка ключа конфигурируется в режиме global configuration mode (см. раздел command manual).

28.2.16 Configure interface weight (Конфигурация веса интерфейса)

Команда конфигурации веса интерфейса в режиме interface configuration mode:

IP rip metric < metric >

Команда установки веса интерфейса по умолчанию:

no IP rip metric

Параметр metric определяет значение добавленного веса к пути интерфейса, значение изменяется в пределах 1-16. По умолчанию установлено значение 1.

28.2.17 Display information (Просмотр конфигурации)

Команды просмотра конфигурации протокола в режиме normal mode или privileged mode:

Команда просмотра информации о всех включенных протоколах:

show ip protocols

Команда просмотра информации о RIP протоколе:

show ip protocols rip

Команда просмотра путей RIP:

```
show ip rip
```

Команда просмотра базы данных RIP:

```
show ip rip database
```

Команда просмотра количества entries базы данных RIP:

```
show ip rip database count
```

Команда просмотра информации RIP интерфейса:

```
show ip rip interface [if name]
```

Параметр if name указывает имя интерфейса

Команда просмотра текущей конфигурации коммутатора, включая конфигурацию протокола RIP в режиме privileged mode:

```
show running config
```

Команда просмотра текущей конфигурации протокола RIP:

```
Show rip command
```

28.3 RIP configuration example (Пример конфигурации)

Три коммутатора соединены в пары с двумя сегментами сети соответственно. Протокол Rip обеспечивает взаимодействие между тремя ПК. Конфигурация соединения оборудования представлена на рисунке ниже:



Команды конфигурации коммутатора 1:

```
Switch# configure terminal
Switch(config)#router rip
Switch(config-rip)#network 192.168.1.0/24
Switch(config-rip)#network 10.1.1.0/24
Switch(config-rip)#network 10.1.2.0/24
```

Команды конфигурации коммутатора 2:

```
Switch# configure terminal
Switch(config)#router rip
Switch(config-rip)#network 192.168.2.0/24
Switch(config-rip)#network 10.1.2.0/24
Switch(config-rip)#network 10.1.3.0/24
```

Команды конфигурации коммутатора 3:

```
Switch# configure terminal
Switch(config)#router rip
Switch(config-rip)#network 192.168.3.0/24
Switch(config-rip)#network 10.1.1.0/24
Switch(config-rip)#network 10.1.3.0/24
```

Команды просмотра конфигурации RIP:

```
show ip protocols rip
show ip rip database
show ip rip interface
```

29. RIPng Configuration (Конфигурация RIPng)

Глава описывает конфигурацию протокола RIPng (относительно простого внутреннего шлюзового протокола), который применяется в IPv6 сетях. Протокол RIPng в основном применяется в сетях небольших предприятий и в региональных сетях с простой структурой. Глава содержит следующие разделы:

- Introduction (Введение)
- Configuration RIPng (Конфигурация протокола)
- Configuration example (Пример конфигурации)

29.1 RIPng introduction (Введение)

Протокол RIPng применяется в IPv6 сетях. Т.к. RIPng протокол является относительно простым (проще чем OSPFv3 и IS-IS) в конфигурации, обслуживании и управлении, то он находит широкое применение на практике.

При создании IPv6 сетей, протокол динамической маршрутизации должен обеспечивать точность и эффективность маршрутизации при передаче информации и сообщений IPv6. Т.о. используются все преимущества RIP и его разновидностей IETF (RIP next generation) для модификации IPv6 сетей. RIPng является необходимым протоколом технологии для обеспечения маршрутизации IPv6 сетей. Протокол RIPng имеет ряд отличий от RIP для применения в сетях IPv6, которые приведены ниже:

- RIPng использует port 521 UDP (RIP использует 520) для отправки и получения маршрутной информации.
- Адрес назначения RIPng использует длину префикса 128 bit (mask length).
- RIPng использует адрес 128 bit IPv6 как адрес следующего hop.
- RIPng использует link local address fe80:: / 10 как адрес источника для отправки сообщения обновления маршрутной информации RIPng.
- RIPng отправляет маршрутную информацию периодически методом multicast, и использует ff02:: 9 как адрес multicast роутера в пределах локальных соединений.
- RIPng сообщение состоит из заглавной и множественной маршрутной таблицы. В одном RIPng сообщении, наибольшее число RTE определяется согласно значению MTU интерфейса.

29.2 RIPng configuration (Конфигурация RIPng)

После запуска протокола RIPng становится возможным конфигурировать функции и атрибуты протокола. Соответствующие команды выполняются в режимах RIPng configuration mode и interface configuration mode.

Конфигурация протокола RIPng включает:

- Запуск протокола RIPng и режима configuration mode

- Запуск RIPv6 интерфейса
- Конфигурация соседних устройств для отправки обновлений
- Конфигурация рабочего статуса интерфейса
- Конфигурация весов маршрутов по умолчанию
- Конфигурация таймера
- Введение внешней маршрутизации
- Конфигурация фильтрации маршрутизации
- Конфигурация дополнительных весов маршрутов
- Конфигурация горизонтальной сегментации
- Конфигурация веса интерфейса

29.2.1 Start RIPv6 and enter RIPv6 configuration mode (Запуск и режим конфигурации RIPv6)

Команды запуска/отключения протокола RIPv6 в режиме global configuration mode:

Команды запуска протокола и режима конфигурации:

```
router IPv6 rip
```

Команда отключения протокола:

```
no router IPv6 rip
```

Протокол RIPv6 отключен по умолчанию.

29.2.2 Enable RIPv6 interface (Запуск RIPv6 интерфейса)

Протокол RIPv6 может определять некоторые интерфейсы, конфигурировать сеть как RIPv6, отправлять и получать сообщения протокола интерфейса RIPv6.

Команда запуска интерфейса в режиме configuration mode:

```
IPv6 Router rip
```

Команда отключения интерфейса:

```
no IPv6 Router rip
```

По умолчанию протокол RIPv6 после запуска отключен на всех интерфейсах.

После запуска протокола RIPv6, следует определить интерфейс на котором будет работать протокол.

29.2.3 Configure specific neighbors to send updates (Конфигурация соседних устройств для отправки обновлений)

Команды выбора соседних устройств для отправки им обновлений в режиме RIPv6 configuration mode:

Команда конфигурации IPv6 адреса для выбранных соседних устройств:
neighbor < neighbor address > < ifname >

Команда отключения функции для соседнего устройства:
no neighbor < neighbor address > < ifname >

Параметр < neighbor address > это IPv6 адрес выбранного соседнего устройства.

29.2.4 Configure the working state of the interface (Конфигурация рабочего статуса интерфейса)

При работе RIPv6 протокола в сети достаточно маршрутизации RIPv6 интерфейса, но не обязательно направлять RIPv6 маршрут на этот интерфейс. Команды passive interface нужны для определения маршрута интерфейса и блокировки вещания интерфейса.

Команда установки статуса passive интерфейса в режиме RIPv6 configuration mode:

passive interface < if name >

Команда сброса статуса passive интерфейса:

no passive interface < if name >

Параметр if name указывает имя интерфейса (например vlan1, vlan2...). По умолчанию все включенные RIPv6 интерфейсы не находятся в статусе passive.

29.2.5 Configure default routing weights (Конфигурация веса маршрута)

При указании внешнего маршрута, необходимо определить его вес (весовой коэффициент). Если вес маршрута не определен, то используется значение веса по умолчанию.

Команда установки веса пути по умолчанию в режиме Ripng configuration mode:

```
default metric < metric >
```

Команда отмены установки веса пути по умолчанию:

```
no default metric [Metric]
```

При восстановлении внешнего пути, по умолчанию используется значение веса 1.

Параметр metric value изменяется в пределах 1-16 (более чем 1 и менее чем 16). По умолчанию значение параметра metric value 1. Для установки значения по умолчанию предусмотрена команда no default metric command.

29.2.6 Configure timer (Конфигурация таймера)

Протокол Ripng имеет три таймера. Первый (update) таймер отвечает за рассылку полной таблицы маршрутизации по всем RIPng интерфейсам каждые 30 сек. Второй (timeout) таймер отслеживает периодичность получения обновленных таблиц маршрутизации для каждого маршрута каждые 180 сек (отмечаются значением metric 16). Третий (garbage) таймер отвечает за удаление из RIPng таблицы маршрутизации путей со значением metric 16 не получивших обновление в течении 120 сек.

Команда установки значений времени таймеров в режиме Ripng configuration mode:

```
timers basic < update > < timeout > < garbage >
```

Команда установки значений времени таймеров по умолчанию:

```
no timers basic
```

Параметр < update > определяет периодичность обновления полной таблицы маршрутизации. Параметр < timeout > участвует в определении

путей маршрутизации не получивших обновления. Параметр `< garbage >` участвует в определении путей маршрутизации отмеченных как недоступные и подлежащих удалению. Время таймеров изменяется в пределах 231-1.

По умолчанию значение параметра `update` 30 сек, значение параметра `timeout` 180 сек (по истечении которого путь отмечается как недоступный `invalid`), значение параметра `garbage` 120 сек (по истечении которого недоступный путь удаляется).

29.2.7 Introducing external routing (Введение внешней маршрутизации)

Протокол `Ripng` позволяет пользователям использовать маршрутную информацию других протоколов в таблице маршрутизации. Протокол `Ripng` может использовать данные протоколов `connected`, `static`, `OSPF`, `IS-IS` и `BGP`.

Команды конфигурации в режиме `Ripng configuration mode`:

```
redistribute {kernel | connected | static | ospf6 | Isis | BGP} [Metric < metric > | route map < route map name >]
```

Команды сброса конфигурации:

```
no redistribute {kernel | connected | static | ospf6 | Isis | BGP} [Metric < metric > | route map < route map name >] Cancel incoming route
```

Первый параметр указывает имя импортируемого протокола (включая прямое соединение `connection`, статическое `static`, `OSPF`, `is is` и `BGP`). Второй параметр определяет вес при импорте, изменяется в пределах 1-16. Третий параметр указывает имя связанной маршрутной карты `route map`. Маршрутная карта конфигурируется в режиме `global configuration mode` (см. раздел `command manual`).

По умолчанию функция использования внешней маршрутизации протоколом `Ripng` отключена.

29.2.8 Configure routing filtering (Конфигурация фильтра маршрутизации)

Протокол `Ripng` поддерживает функцию фильтрации маршрутизации.

Функция обеспечивает фильтрацию получаемых и отправляемых маршрутов с помощью специальной таблицы контроля доступа `access control list`, таблицы префикса адресов `address prefix list`, и настройки политики правил `policy rules`.

Команды конфигурации функции фильтрации в режиме `Ripng configuration mode`:

Команда включения таблицы доступа для контроля входа и выхода интерфейса:

```
distributed list < ACL name > {in | out} [if name]
```

Команда отключения таблицы контроля доступа:

```
no distributed list < ACL name > {in | out} [if name]
```

Параметр `< ACL name >` указывает имя соответствующей таблицы доступа. Параметр `[if name]` указывает `Ripng` интерфейс, к которому применяется фильтрация. Значения `{in | out}` определяют принимаемые и отправляемые маршруты (`receiving` или `publishing routes`).

Команда включения таблицы префиксов `prefix list`:

```
distributed list prefix < pre name > {in | out} [if name]
```

Команда отключения таблицы префиксов:

```
no distributed list prefix < pre name > {in | out} [if name]
```

Параметр `< pre name >` указывает имя соответствующей таблицы префиксов. Параметр `[if name]` указывает `Ripng` интерфейс, к которому применяется фильтрация. Значения `{in | out}` определяют принимаемые и отправляемые маршруты (`receiving` или `publishing routes`).

По умолчанию `Ripng` протокол не фильтрует какие-либо принятые или отправленные маршруты.

Таблицы доступа `Access list` и префиксов `prefix list` конфигурируются в режиме `global configuration mode` (см. раздел `command manual`).

29.2.9 Configure additional routing weight (Конфигурация дополнительного веса маршрута)

Дополнительный вес (весовой коэффициент) маршрута это предустановленное добавочное значение к весу пути `Ripng` протокола на входе и выходе. Дополнительный вес не изменяет непосредственно

вес маршрута в маршрутной таблице, но добавляет предустановленную величину к весу маршрута при его получении или отправке интерфейсом.

Команда ввода дополнительного веса в интерфейс при отправке-получении маршрута в режиме Ripng configuration mode:

```
offset list < ACL name > {in | out} < offset > [if name]
```

Команда сброса дополнительного веса в интерфейсе при отправке-получении маршрута:

```
no offset list < ACL name > {in | out} < offset > [if name]
```

Параметр < ACL name > указывает имя таблицы доступа интерфейса. Значения {in | out} определяют добавление дополнительного веса к принимаемым и отправляемым маршрутам. Параметр < offset > определяет значение дополнительного веса, изменяется в пределах 0-16. Параметр [if name] указывает Ripng интерфейс, на котором применяется добавление дополнительного веса.

По умолчанию при получении сообщения, значение дополнительного веса для каждого пути 1, при отправке сообщения значение дополнительного веса для каждого пути 0.

29.2.10 Configure horizontal segmentation (Конфигурация горизонтальной сегментации)

Горизонтальная сегментация означает, что путь полученный от интерфейса не был отправлен этим интерфейсом. Горизонтальная сегментация с инверсией (toxic inversion) означает что путь полученный от интерфейса отправляется этим интерфейсом, но его значение metric value - 16. Горизонтальная сегментация может до определенной степени предотвращать образование петлевых соединений. Горизонтальная сегментация с инверсией обладает большей эффективностью по сравнению с обычной горизонтальной сегментацией, при этом прямая маркировка не доступна. Обычно в сетях pBMA, необходимо запрещать горизонтальную сегментацию для выстраивания корректного маршрута.

Команда включения функции horizontal division и режима toxic inversion на интерфейсе в режиме interface configuration mode:

IPv6 rip split horizon [poisoned]

Команда отключения функции horizontal division function на интерфейсе:
no IPv6 rip split horizon

Параметр [no poisoned] включает функцию горизонтальная сегментация. Параметр [poisoned] включает функцию горизонтальная сегментация с инверсией (toxic inversion).

По умолчанию выбрана функция горизонтальная сегментация с инверсией (toxic inversion).

29.2.11 Configure interface weight (Конфигурация веса интерфейса)

Команда конфигурации веса интерфейса в режиме interface configuration mode:

IPv6 rip metric offset < metric >

Команда установки веса интерфейса по умолчанию:

no IPv6 rip metric offset < metric >

Параметр metric определяет значение добавленного веса к пути интерфейса, значение изменяется в пределах 1-16. По умолчанию установлено значение 1.

29.2.12 Display information (Просмотр конфигурации)

Команды просмотра конфигурации протокола в режиме normal mode или privileged mode:

Команда просмотра информации о RIPv6 протоколе:

show IPv6 rip

Команда просмотра информации о путях RIPv6:

show IPv6 route rip

29.3 RIPng configuration example (Пример конфигурации)



Команды конфигурации коммутатора S1:

```
Switch(config)#  
Switch(config)#int vlan1  
Switch(config-vlan1)#ipv6 address 3ffe:506::1/64  
Switch(config-vlan1)#exit  
Switch(config)#router ipv6 rip  
Switch(config-router)#exit  
Switch(config)#int vlan1  
Switch(config-vlan1)#  
Switch(config-vlan1)#ipv6 router rip  
Switch(config-vlan1)#
```

Команды конфигурации коммутатора S2:

```
Switch(config)#  
Switch(config)#int vlan1  
Switch(config-vlan1)#ipv6 address 3ffe:506::2/64  
Switch(config-vlan1)#exit  
Switch(config)#router ipv6 rip  
Switch(config-router)#exit  
Switch(config)#int vlan1  
Switch(config-vlan1)#  
Switch(config-vlan1)#ipv6 router rip  
Switch(config-vlan1)#
```

30. OSPF Configuration (Конфигурация OSPF)

Глава описывает конфигурацию протокола OSPF (open shortest path first), работа протокола основана на алгоритме статуса соединения link state algorithm. Протокол имеет высокую скорость конвергенции и способен работать в крупномасштабных сетях. Глава содержит следующие разделы:

- Introduction (Введение)
- OSPF Configuration (Конфигурация протокола)
- Configuration example (Пример конфигурации)

30.1 OSPF introduction (Введение)

Протокол OSPF (open shortest path first) основан на алгоритме статуса соединения link state algorithm. Протокол имеет высокую скорость конвергенции и способен работать в крупномасштабных сетях.

Роутеры с поддерживающие протокол OSPF обслуживают базу данных статуса соединений link state database (LSDB), которая описывает топологию всей автономной системы как карту. Когда базы данных всех роутеров синхронизированы, каждый роутер вычисляет кратчайший путь к ноду назначения в автономной системе со своей точки зрения и обслуживает свою собственную маршрутную таблицу. Когда изменяется топология сети, роутеру достаточно инкапсулировать изменившийся статус соединения link state в сообщении обновления link state update (LSU) и передать его. Все роутеры заново синхронизируют локальную базу данных и пересчитывают маршрут. Каждый роутер публикует статус соединения link state (LSA), отслеживает и собирает описание топологии LSDB всей сети. После преобразования данных в весовую диаграмму, может быть использован алгоритм SPF для вычисления таблицы маршрутов.

В вещательной сети, каждый роутер должен транслировать информацию о своем собственном статусе другим роутерам, настраиваются множественные парные соединения, которые вносят большое количество передаваемых сообщений, которые не являются необходимыми. Протокол OSPF согласован с designated router (DR) и backup designated router (BDR). Роутер отправляет информацию о

соединении к DR, который собирает, ранжирует, а затем отправляет её другим роутерам. Т.о. достигается эффективное уменьшение количества смежных соединений между роутерами в вещательной сети.

Протокол OSPF поддерживает пять типов сообщений:

- Hello message сообщение периодически отправляется соседним устройствам для обнаружения, обслуживания и поддержки выбора Dr. Сообщение содержит несколько атрибутов интерфейса. Для создания соединений соседних устройств необходимо наличие некоторых параметров сообщения Hello message.
- DD message (database description) сообщение используется для описания собственного LSDB в процессе синхронизации, включая заголовок каждого LSA. Заголовок LSA head делает возможным уникальное определение одного LSA, и peer роутер может определить этот LSA, в противном случае запросить LSA повторно.
- LSR message (link state request) сообщение необходимо чтобы два роутера после обмена сообщениями DD messages, могли определить который из LSAS пропущен смежным роутером. В это время необходимо отправить сообщение LSR message для запроса полного LSA. Для запроса необходим только заголовок LSA.
- LSU message (link state update) библиотека множества LSAS.
- Lsack message (link state acknowledgement) сообщение подтверждает получение сообщения LSU message для гарантированной передачи информации о соединениях link. Подтверждается заголовком LSA head.

Router ID concept обеспечивает уникальную идентификацию роутера в автономной системе.

Area (область, регион). Если протокол OSPF работает в крупной сети, большое количество роутеров приводит к увеличению размера LSDB, росту времени синхронизации, росту времени вычисления маршрутов и повышенной нагрузки процессора CPU. Чем крупнее сеть, тем чаще происходит изменение топологии, и тем больше требуется времени на её изменение. Роутер тратит много времени на передачу сообщений и вычисление маршрутов, что приводит к неоправданной перегрузки сети. Протокол OSPF реализует концепцию разделения роутеров по разным участкам сети (регионам). Т.о. LSDB обеспечивает синхронизацию и вычисляет маршруты только для региона (участка сети). Взаимодействие между регионами сети выполняется роутером

boundary router (ABR). В этом случае, количество роутеров в регионе будет ограничено, и LSDB будет иметь меньший размер и потреблять меньше сетевых ресурсов. При изменении топологии будет значительно уменьшено время вычисления маршрутов и увеличится скорость конвергенции. Концепция разделения большой сети на регионы эффективно группирует крупномасштабные сети и обеспечивает маршрутизацию на коротких расстояниях в пределах каждого региона. Роутеры между регионами взаимодействуют на уровне backbone region (регионы с ID 0). Таким образом, все не backbone areas области должны быть соединены с областью backbone area, а значит, ABR имеет только один интерфейс для соединения с областью backbone area. Если при планировании сети, не backbone area не может быть соединена с областью backbone area, должно быть создано виртуальное соединение (логический путь). Т.е. ABR в области backbone area и ABR в не backbone area создают соединение точка-точка через область передачи (transmission area). Информация о внутренней маршрутизации домена области backbone через виртуальное соединение будет распространена в области не backbone area.

30.2 OSPF configuration (Конфигурация OSPF)

После запуска протокола OSPF, в режиме конфигурации OSPF становится возможным настраивать и управлять соответствующими функциями протокола. Команды конфигурации протокола OSPF доступны в режиме OSPF configuration mode и interface configuration mode.

Конфигурация протокола OSPF включает:

- Запуск протокола OSPF и режима конфигурации OSPF mode
- Запуск интерфейса
- Определение хоста
- Конфигурация ID роутера router ID
- Конфигурация смежных точек adjacency points
- Отключение отправки сообщений интерфейсом
- Конфигурация таймера SPF timer
- Конфигурация дистанции управления
- Введение внешней маршрутизации
- Конфигурация типа сетевого интерфейса

- Конфигурация интервалов отправки сообщений Hello message
- Конфигурация времени существования соседнего роутера
- Конфигурация интервала ретрансляции
- Конфигурация интервала задержки интерфейса
- Конфигурация приоритета интерфейса при выборе Dr
- Конфигурация важности сообщений при отправке интерфейсом (cost of sending messages)
- Заполнение значения MTU value для сообщений DD отправленных интерфейсом
- Конфигурация аутентификации сообщений интерфейса
- Конфигурация регионального виртуального соединения (regional virtual link)
- Конфигурация агрегации региональной маршрутизации
- Конфигурация аутентификации региональных сообщений
- Конфигурация области Stub area
- Конфигурация области NSSA area
- Конфигурация агрегации внешней маршрутизации
- Конфигурация весов внешних маршрутов по умолчанию

30.2.1 Start OSPF and enter OSPF mode (Запуск протокола OSPF и режима конфигурации)

Одновременно может быть запущено несколько процессов OSPF, каждый из процессов идентифицируется по ID. При запуске протокола OSPF необходимо определить какой процесс начат. Если параметры отсутствуют, то процесс имеет номер 0.

Команда запуска процесса OSPF (с номером ID) в режиме global configuration mode:

```
router OSPF [process ID]
```

Команда остановки процесса OSPF с номером ID:

```
no router OSPF [process ID]
```

Параметр process ID определяет номер запущенного процесса, изменяется в пределах 216-1. Если параметр process ID не определен, то запускается процесс OSPF с номером 0.

По умолчанию протокол OSPF отключен.

30.2.2 Enable interface (Включение интерфейса)

Назначение протокола OSPF – разделение по уровням целой автономной системы на разные регионы для создания концептуальной иерархической модели сети. В автономной системе роутеры принудительно формируются в группы, а зона является логической. Когда разные интерфейсы роутера принадлежат разным регионам, пересекающимся регионам, то роутеры называют связанными boundary router ABR. Каждый сегмент сети в котором запущен протокол OSPF может принадлежать только определенной зоне, соответственно каждый интерфейс роутера с запущенным протоколом OSPF должен принадлежать определенной зоне. Каждая зона имеет свой ID номер, зона с номером 0 является зоной backbone area. Маршрутная информация между разными регионами передается через граничный роутер border router. В отличие от протокола Rip, при работе протокола OSPF на интерфейсе должен быть определен его регион действия.

Команда выбора интерфейса для работы протокола OSPF в определенной зоне в режиме OSPF configuration mode:

```
network < network address > area < area ID >
```

Команда отключения протокола на определенном интерфейсе определенной зоны:

```
no network < network_ address> area <area-id>
```

Параметр network address имеет два формата: a.b.c.d/m и a.b.c.d. Первый формат определяет IP сети и длину маски mask length, второй формат определяет IP сети и маску. Параметр Area ID тоже имеет два формата: a.b.c.d и интегральный integer. Первый использует десятичный формат с разделением точками, второй использует цифровой формат, изменяется в пределах 0 или 232-1.

По умолчанию после запуска протокола OSPF, интерфейс автоматически не создается.

30.2.3 Specify host (Определение хоста)

Команда определения пути хоста в режиме OSPF configuration mode:

```
host < IP address > area < area ID > [cost < cost >]
```

Команда сброса пути хоста:

```
no host < IP address > area < area ID > [cost < cost >]
```

Параметр IP address использует формат a.b.c.d для указания хоста назначения в определенной зоне, является типом корня в представлении соединения роутера. Параметр Area ID имеет два формата: a.b.c.d и интегральный integer. Первый использует десятичный формат с разделением точками, второй использует цифровой формат, изменяется в пределах 0 или 232-1. Параметр Cost определяет вес (значение) соединения, является опциональным параметром.

По умолчанию параметр cost не определен, имеет значение 0.

30.2.4 Configure router ID (Конфигурация ID роутера)

Для однозначной идентификации роутера в автономной системе используется номер ID в 32-bit формате. ID номер роутера может быть сконфигурирован вручную. В процессе конфигурации необходимо обеспечить уникальность ID номеров всех роутеров в автономной системе. Если номер роутера не определен, то используется IP адрес интерфейса обратной связи loopback interface. Если loopback не имеет IP адреса, то в качестве ID выбирается больший из IP адресов текущего интерфейса. Для обеспечения стабильности функционирования протокола OSPF, ID номер роутера должен быть определен уже на этапе планирования сети.

Команда конфигурации ID номера роутера в режиме OSPF configuration mode:

```
OSPF router ID < router ID >
```

Команда сброса ID номера роутера:

```
no OSPF router ID
```

Команда сброса ID номера роутера:

```
no router ID [router ID]
```

Параметр router ID имеет формат a.b.c.d.

По умолчанию после запуска протокола OSPF, будет автоматически сгенерирован параметр router ID. При этом учитываются следующие правила: в первую очередь учитывается router ID заданный с помощью

команд; при его отсутствии выбирается IP адрес интерфейса обратной связи loorback; при его отсутствии больший из IP адресов текущего интерфейса; при его отсутствии выбирается 0.0.0.0.
Две группы команд имеют одинаковое назначение.

30.2.5 Configure adjacency points (Конфигурация смежных точек)

Сообщения протокола OSPF используют для взаимодействия адреса multicast 224.0.0.5 или 224.0.0.6. Когда соединение не поддерживает вещание pBMA, для использования unicast и взаимодействия с сообщениями протокола должны быть сделаны некоторые настройки конфигурации протокола OSPF. Доступно ручное указание IP адреса и соответствующего значения атрибута для противоположной стороны соединения.

Команда определения смежных точек и установки свойств в режиме OSPF configuration mode:

```
neighbor < IP address > [priority < prio > | poll interval < deadtime > | cost < cost >]
```

Команда сброса смежных точек и втрибутов противоположной стороны соединения:

```
no neighbor < IP address > [priority < prio > | poll interval < deadtime > | cost < cost >]
```

Параметр IP address определяет адрес противоположной стороны соединения, имеет формат a.b.c.d. Параметр Prio определяет приоритет пира peer priority, изменяется в пределах 0-255. Параметр Deadtime определяет время разрыва соединения end-to-end down. Если отсчет времени остановлен, то сообщение Hello message не будет отправлено соединению end-to-end, изменяется в пределах 1 или 216-1. Параметр Cost определяет вес соединения с противоположной стороны, изменяется в пределах 1 или 216-1.

По умолчанию установлено значение приоритета priority 1 (при значении 0 не происходит участия в процессе выбора Dr election), интервал Poll - 120 сек, значение Cost 10.

30.2.6 Prohibit the interface from sending messages (Запрет на отправку сообщений)

В простых сетях, интерйс протокола OSPF представляет только сегмент сети между двумя устройствами для передачи данных, переводит интерфейс в пассивный статус и блокирует передачу сообщений Hello message для этого соединения, которое не влияет на информацию о маршруте интерфейса.

Команда перевода интерфейса в пассивный статус в режиме OSPF configuration mode:

```
passive interface < if name >
```

Команда перевода интерфейса в обычный режим:

```
no passive interface < if name >
```

Параметр if name – имя интерфейса уровня layer-3 (vlan1, vlan2...).

По умолчанию после запуска протокола OSPF, включенные интерфейсы не находятся в пассивном статусе. После перевода интерфейса с запущенным протоколом OSPF в пассивный статус, может быть опубликован прямой путь интерфейса, но сообщения протокола OSPF на интерфейсе будут заблокированы, и интерфейс не сможет создать соединение с соседним устройством. В некоторых случаях это помогает эффективно экономить сетевые ресурсы.

30.2.7 Configure SPF calculation time (Вычисление времени SPF)

Когда происходят изменения в базе данных статуса соединений LSDB (link state database) протокола OSPF, требуется произвести перерасчет кратчайшего пути. Если кратчайший путь вычисляется незамедлительно для каждого изменения, то на это требуется много ресурсов, что снижает производительность роутера. Путем конфигурации значений времени задержки delay и удержания hold SPF вычисления, может быть снижена частота SPF вычислений (вызванная слишком частыми изменениями сети), и как следствие удастся снизить расходование ресурсов системы в единицу времени и повысить производительность роутера.

В процессе SPF вычисления используется таймер, который запускает каждое следующее вычисление в соответствии со временем

ингибирования inhibition time. Процесс SPF вычисления запускается по истечении времени таймера, пересчитывается время ингибирования с момента последнего SPF вычисления. Если заданное время ингибирования превышено, то для запуска таймера используется заданное время задержки delay time. Если заданное время ингибирования не превышено, то оно используется для вычисления необходимого времени задержки delay time. Если время задержки имеет значение меньше заданного, для вычисления используется заданное время, в противном случае для запуска SPF вычисления используется вычисленное время задержки.

Команда установки значений задержки и удержания (delay и hold) интервала SPF вычисления в режиме OSPF configuration mode:

```
timers SPF < delay > < hold >
```

Команда установки значений задержки и удержания по умолчанию:

```
no timers SPF
```

Параметр delay устанавливает время задержки необходимое для SPF вычисления. Параметр Hold устанавливает временной интервал между двумя SPF вычислениями.

По умолчанию установлены значения: задержки delay 5 сек, удержания Hold 10 сек.

30.2.8 Configure management distance (Конфигурация дистанции управления)

На роутере может быть одновременно запущено несколько протоколов маршрутизации. Дистанция управления нужна для выбора необходимой маршрутной информации при работе нескольких протоколов маршрутизации. В том случае если разные протоколы находят один и тот же путь, то наименьшая дистанция управления является предпочтительной.

Команда конфигурации дистанции управления в режиме OSPF configuration mode:

```
distance < distance >
```

Команда возврата к конфигурации по умолчанию:

```
no distance < distance >
```

Команда конфигурации различных типов дистанции управления:

```
distance OSPF {intra area < distance > | inter area < distance > | external  
< distance >}
```

Команда возврата трех типов дистанции управления к значениям по умолчанию:

```
no distance OSPF
```

Параметр distance изменяется в пределах 1-255. Параметр Intra area относится к дистанции управления маршрутов в домене. Параметр Inter area относится к дистанции управления маршрутов между доменами. Параметр External указывает дистанцию управления внешних маршрутов.

По умолчанию значение параметра management distance протокола OSPF 110. Значение параметра маршрутов intra domain route, inter domain route и external route 0.

30.2.9 Introducing external routing (Введение внешней маршрутизации)

На роутере может быть одновременно запущено несколько протоколов динамической маршрутизации, различные протоколы могут делиться маршрутной информацией. Протокол OSPF использует маршрутные данные других протоколов для внешних маршрутов автономной системы, которые вводятся автономной системой связанной с роутером boundary router ASBR. При введении внешнего маршрута, можно определять такие атрибуты как вес weight и тип веса weight type.

Протокол OSPF имеет разделение на четыре типа маршрутизации: внутridoменная маршрутизация (intra domain routing); междудоменная маршрутизация (inter domain routing); внешняя маршрутизация 1-го типа (external route type-1); внешняя маршрутизация 2-го типа (external route type-2). Оба типа внешней маршрутизации описывают пути к точке назначения вне автономной системы. Внешняя маршрутизация типа 1 до другого IGPS. Считается что протокол OSPF имеет высокую надежность и имеет вес сравнимый с весом маршрута внутри автономной системы. Таким образом значение cost внешнего маршрута оценивается суммой значений cost

пути от самого роутера к ASBR и cost от ASBR до точки назначения. Внешняя маршрутизация типа 2 до другого eGPS. Считается что протокол OSPF не имеет высокую надежность и его значение cost на много превышает вес маршрута внутри автономной системы и их значения cost сильно различаются. Таким образом, значение cost внешнего маршрута оценивается по значению cost пути от ASBR до точки назначения, а значение cost пути от роутера до ASBR не учитывается.

Команда конфигурации параметров в режиме OSPF configuration mode:

```
redistribute {kernel | connected | static | rip | Isis | BGP} [Metric < metric > | metric type < type > | route map < route map name > | tag < tag >]
```

Команда возврата параметров к значениям по умолчанию:

```
no redistribute {kernel | connected | static | rip | Isis | BGP} [Metric < metric > | metric type < type > | route map < route map name > | tag < tag >]
```

Параметр 1 указывает тип вводимого внешнего маршрута, включая прямое соединение, static, rip, IS-IS and BGP. Параметр 2 определяет значение веса weight value (устанавливается при введении внешнего маршрута), изменяется в пределах 0 или 224-1. Параметр 3 определяет тип вводимого внешнего маршрута (type-1 или type-2), Type-1 соответствует IGP route, type-2 соответствует EGP route. Параметр 3 указывает имя соответствующей маршрутной карты. Маршрутная карта конфигурируется в режиме global configuration mode (см. соответствующий раздел). Параметр 4 - tag, изменяется в пределах 0 или 232-1, является атрибутом внешней маршрутизации.

По умолчанию протоколы внешней маршрутизации отключены.

30.2.10 Configure the network type of the interface (Конфигурация типа сетевого интерфейса)

Протокол OSPF сам выбирает роутер, каждый роутер описывает топологию соответствующей сети и отправляет её другим роутерам. Протокол OSPF разделяет сети (в зависимости от соединений интерфейса) на четыре типа согласно протоколу уровней соединений:

- broadcast type (протоколы уровня Ethernet, FDDI, и т.д.);
- nBMA non broadcast множественного доступа multiple access type (протоколы уровня fr, ATM, HDLC, X.25, и т.д.);
- точка-многоточка point to multipoint type (по умолчанию отсутствует, должен быть принудительно сконфигурирован с помощью других типов сетей). В большинстве случаев для соединения точка-многоточка изменяется протокол nBMA;
- точка-точка point-to-point type (протоколы уровня PPP, LAPB и POS).

В сетях broadcast networks без возможности множественного доступа, тип интерфейса может быть сконфигурирован как nBMA. Если не все роутеры имеют прямой доступ к сети nBMA, интерфейс может иметь конфигурацию типа точка-многоточка.

Сеть nBMA имеет полное соединение через протокол OSPF, также доступны non broadcast и multi-point. Для сети типа точка-многоточка отсутствует необходимость полного соединения. Для NBMA необходимо выбрать Dr, в сети точка-многоточка DR отсутствует. В сети NBMA сообщения multicast message передаются через назначенное соседнее устройство. В сети точка-многоточка используются сообщения unicast message и multicast message.

Команда конфигурации типа сети для соединения интерфейса в режиме interface configuration mode:

```
ip ospf network < type >
```

Команда возврата к конфигурации по умолчанию:

```
no ip ospf network
```

Параметр Type определяет тип сети: broadcast, non broadcast, точка-точка, точка-многоточка [non broadcast]; Первый тип – сеть broadcast, второй тип – сеть non broadcast (nBMA), третий тип – сеть point-to-point, четвертый тип – сеть point-to-multipoint. Сети точка-многоточка подразделяются на broadcast и non broadcast сети. Non broadcast соседние устройства не могут быть обнаружены автоматически и требуют отдельного определения.

По умолчанию установлен тип сети broadcast network.

30.2.11 Configure Hello message sending interval (Конфигурация интервалов отправки Hello сообщений)

Сообщения Hello message периодически отправляются соседнему роутеру для обслуживания взаимодействия между роутерами, выбора DR и BDR. Интервал отправки сообщений Hello message может быть установлен вручную, с учетом определенной частоты обмена сообщениями между устройствами сети. Значение времени таймера Hello timer должно быть обратно пропорционально скорости конвергенции роутера и нагрузки сети.

Команда установки времени таймера Hello timer в режиме interface configuration mode:

```
ip ospf Hello interval < seconds >
```

Команда возврата к конфигурации значения по умолчанию:

```
no ip ospf Hello interval
```

Параметр seconds определяет время между отправками двух сообщений Hello message, изменяется в пределах 1 или 216-1.

По умолчанию для сетей broadcast и точка-точка установлено значение 10 сек; для pBMA установлено значение более 30 сек.

30.2.12 Configure neighbor router expiration time (Конфигурация времени существования соседнего роутера)

Команда конфигурации времени существования в режиме interface configuration mod:

```
ip ospf dead interval < seconds >
```

Команда возврата к значению по умолчанию:

```
no ip ospf dead interval
```

Параметр seconds изменяется в пределах 1 или 216-1, соседнее устройство считается недоступным, если сообщение Hello message не получено в течение заданного периода времени. При каждом получении сообщения hello message происходит перезапуск таймера.

По умолчанию для сети broadcast и точка-точка установлено значение 40 сек. Для сети pBMA и точка-многоточка установлено значение 120

сек. При изменении типа сети значения Hello interval и dead interval возвращаются к значениям по умолчанию.

30.2.13 Configure retransmission time (Конфигурация интервала ретрансляции)

OSPF это протокол надежного статуса соединения (reliable link state protocol), который отражает факт взаимодействия сообщений запроса LSU messages и ответа со стороны подключенного к другому концу линии оборудования. При получении подтверждающего сообщения, считается что обновление статуса соединения link state получено. Если подтверждающее сообщение не получено в течении интервала ретрансляции retransmission interval, сообщение LSA будет передано соседнему оборудованию. Интервал ретрансляции может быть сконфигурирован вручную. Длительность интервала должна быть больше чем время необходимое для передачи сообщения между двумя роутерами в направлении туда и обратно. Если интервал выбрать слишком коротким, то это вызовет ненужную ретрансляцию.

Команда конфигурации интервала ретрансляции интерфейса в режиме interface configuration mode:

```
ip ospf retransmit interval < seconds >
```

Команда установки значения интервала по умолчанию:

```
no ip ospf retransmit interval
```

Параметр seconds устанавливает длительность интервала ретрансляции в секундах, изменяется в пределах 1 и 216-1, используется в случае если подключенное оборудование не получило сообщение LSA.

По умолчанию установлено значение 5 сек.

30.2.14 Configure interface delay (Конфигурация интервала задержки интерфейса)

В сообщении обновления статуса соединения LSU message, каждый статус соединения содержит время домена time domain. Перед передачей необходимо увеличить задержку трансляции передающего

интерфейса. Этот параметр в основном учитывает время необходимое интерфейсу для отправки сообщений и особенно критичен для низкоскоростных сетей.

Команда установки времени задержки передачи интерфейсом в режиме interface configuration mode:

```
ip ospf transmit delay < seconds >
```

Команда установки значения времени задержки по умолчанию:

```
no ip ospf transmit delay
```

Параметр seconds устанавливает значение времени задержки в секундах, изменяется в пределах 1 и 216-1, указывает на то, что значение должно быть увеличено для age domain LSA отправленного интерфейсом.

По умолчанию установлено значение 1 сек.

30.2.15 Configure the priority of interface in Dr election (Конфигурация приоритета интерфейса при выборе Dr)

Для предотвращения повторной передачи информации о соединении в режиме точка-точка в сетях broadcast, необходимо выбрать назначенный роутер DR и BDR, который будет отвечать за отправку информации о соединении в пределах сегмента сети. Приоритет интерфейса указывает на его значение при выборе Dr. В случае возникновения конфликта, интерфейс с большим приоритетом будет считаться определяющим. Если значение приоритета интерфейса 0, то он не будет принимать участия в процессе выбора Dr. В противном случае интерфейсы становятся кандидатами. Каждый роутер обладает своей собственной информацией о приоритете и его собственный DR отправляет сообщения Hello message в сеть, и в конечном итоге выбирается один Dr с наибольшим приоритетом. В случае равного приоритета, роутер с наивысшим ID получает преимущество.

При сбое Dr, роутерам в сети необходимо пройти повторную процедуру выбора Dr, что занимает некоторое время, в течении которого, могут произойти ошибки вычисления маршрутов. Задача BDR обеспечить плавный переход к новому Dr, т.е. BDR является резервным Dr. Его выбор происходит одновременно с выбором Dr. Этот процесс обеспечивает смежное взаимодействие с другими роутерами в сети.

В отличие от BDR, Dr обеспечивает сбор информации и публикацию нодов в сети. BDR обеспечивает только синхронизацию смежных устройств. При отключении Dr, BDR немедленно становится DR и обеспечивает сбор информации в пределах сегмента сети. В это же время запускается новый процесс выбора BDR, но этот процесс не оказывает влияния на вычисление маршрута.

Команда установки приоритета интерфейса при выборе Dr в режиме interface configuration mode:

```
ip ospf priority < prio >
```

Команда установки значения приоритета по умолчанию:

```
no ip ospf priority
```

Параметр prio устанавливает значение приоритета при выборе Dr, изменяется в пределах 0-255. При значении приоритета 0 интерфейс не участвует в процессе выбора.

По умолчанию установлено значение 1.

30.2.16 Configure the cost of sending messages on the interface (Конфигурация важности сообщений при отправке интерфейсом)

В сетях управление трафиком осуществляется на основе настройки соединений и разделения их по важности. Значение cost интерфейса соответствует значению cost отправляемых им сообщений. Если эти значения не сконфигурированы вручную, протокол OSPF автоматически вычислит значение интерфейса interface cost согласно baud rate интерфейса.

Команда установки значения параметра в режиме interface configuration mode:

```
ip ospf cost < cost >
```

Команда установки значения параметра по умолчанию:

```
no ip ospf cost
```

Параметр cost определяет значение сообщения отправленного интерфейсом, изменяется в пределах 1 и 216-1.

По умолчанию установлено значение 10.

30.2.17 MTU value for DD message sent by interface (Значение MTU value для сообщений DD отправленных интерфейсом)

Команда отключения проверки значения MTU в сообщении DD message в режиме interface configuration mode:

```
ip ospf MTU ignore
```

Команда включения проверки значения MTU в сообщении DD message:

```
no ip ospf MTU ignore
```

По умолчанию проверка значения MTU включена.

30.2.18 Configure interface message authentication (Аутентификация сообщений интерфейса)

Протокол OSPF поддерживает режимы plaintext mode и MD5 mode для аутентификации сообщений интерфейса.

Команда конфигурации режима аутентификации в режиме interface configuration mode:

```
ip ospf authentication < mode >
```

Команда отключения режима аутентификации:

```
no ip ospf authentication
```

Параметр no parameter устанавливает режим аутентификации plaintext.

Параметр Message digest устанавливает режим аутентификации MD5.

Параметр Null отключает аутентификацию.

Команда устанавливает режим plaintext для аутентификации строки password string:

```
ip ospf authentication key < password >
```

Команда отключения режима plaintext для аутентификации строки password string:

```
no ip ospf authentication key
```

Параметр password определяет строку password string при plaintext аутентификации.

Команда устанавливает режим MD5 для аутентификации password:

```
ip ospf message digest key < key ID > MD5 < password >
```

Команда отключения режима MD5 для аутентификации password:
no ip ospf message digest key < key ID >

Параметр key ID используется для сортировки key chain, изменяется в пределах 1-255. Параметр Password указывает строку password string. По умолчанию аутентификация отключена.

30.2.19 Configure regional virtual link (Конфигурация регионального виртуального соединения)

Протокол OSPF использует принцип разделения роутеров по уровням на разные группы в автономной системе. Эти группы называются регионами. Все регионы не являются равными и параллельными и имеют иерархическое взаимодействие. Среди них выделяется регион 0.0.0.0, который называется осевым или позвоночным (backbone region). Другие не backbone regions должны обмениваться своими доменными маршрутами через регион backbone region. Т.о. все не backbone области должны быть соединены с backbone областью, соответственно, один интерфейс ABR находится в области 0. Если какие-либо области не могут гарантировать физический путь к области backbone по причине ограничения сетевой топологии, то должны быть сконфигурированы виртуальные соединения для создания логического пути. Оба конца виртуального маршрута принадлежат ABR, средняя часть пути не проходит через backbone area и называется областью передачи (transmission area). При конфигурации виртуального соединения необходимо определить ID области передачи (transmission area) и ID оппозитного ABR, которые конфигурируются на ABR с обоих концов в первую очередь.

Вычисление пути через область передачи transmission area и создание виртуального соединения является логическим эквивалентом соединения точка-точка между конечными точками маршрута. Т.о. параметры интерфейса могут быть сконфигурированы на физическом интерфейсе, и может быть запущена функция аутентификации.

Сообщение unicast message передается между ABRs. Роутер передающий сообщение unicast message через область transmission area передает его как обычное IP сообщение. Т.о. подтверждается создание логического соединения в области передачи transmission area, и сообщения протокола могут передаваться между двумя ABRs.

Команда конфигурации transmission area и peer ID виртуального соединения в режиме OSPF configuration mode:

```
Area < area ID > virtual link < router ID >
```

Команда конфигурации режима аутентификации виртуального соединения:

```
authentication <mode>
```

Команда конфигурации пароля аутентификации виртуального соединения как plaintext:

```
authentication-key <password> |
```

Команда конфигурации пароля аутентификации виртуального соединения как MD5:

```
message-digest-key <key-id> md5 <password> |
```

Команда конфигурации Hello interval виртуального соединения:

```
hello-interval <seconds> |
```

Команда конфигурации failure time виртуального соединения:

```
dead-interval <seconds> |
```

Команда конфигурации retransmission interval виртуального соединения:

```
retransmit-interval <seconds> |
```

Команда конфигурации interface delay виртуального соединения:

```
transmit-delay <seconds> |
```

Команда удаления виртуального соединения:

```
no area < area ID > virtual link < router ID > [authentication <mode> |  
authentication-key <password> | message-digest-key <key-id> md5  
<password> | hello-interval <seconds> | dead-interval <seconds> |  
retransmit-interval <seconds> | transmit-delay <seconds>]
```

Параметр area ID определяет ID области передачи, вводится в десятичном формате с разделением точками a.b.c.d или целым числом, изменяется в пределах 0 или 232-1. Параметр Router ID определяет ID peer роутера виртуального соединения, вводится в десятичном формате с разделением точками a.b.c.d. Конфигурация параметров аутентификации и отправляющего интерфейса являются опциональными возможностями (см. соответствующие разделы).

По умолчанию виртуальное соединение не сконфигурировано.

30.2.20 Configure regional routing aggregation (Конфигурация агрегации региональной маршрутизации)

Команда конфигурации масштаба агрегирования в режиме OSPF configuration mode:

```
Area < area ID > range < IP prefix > [advertise | not advertise]
```

Команды отмены масштаба агрегирования:

```
no area < area ID > range < IP prefix > [advertise | not advertise]
```

Параметр area ID соответствует ID области, определяет агрегацию маршрутов в области, вводится в десятичном формате с разделением точками a.b.c.d или целым числом, изменяется в пределах 0 или 232-1. Параметр IP prefix определяет диапазон агрегации, вводится в формате a.b.c.d/m. Параметры advertise и not advertise указывают на необходимость или отсутствие необходимости отправлять диапазон агрегации (IP prefix). По умолчанию отправляется информация об изначальных сетевых маршрутах.

30.2.21 Configure regional message authentication (Конфигурация аутентификации региональных сообщений)

Тип аутентификации всех роутеров в одной области должен быть последовательным. Аутентификационные password strings всех роутеров в пределах одного сегмента сети также должны быть последовательными. Конфигурация аутентификации области только запускает функцию аутентификации (plaintext или MD5), и password использует соответствующее значение конфигурации интерфейса (см. раздел Аутентификация сообщений интерфейса).

Команда конфигурации региональной аутентификации в режиме OSPF configuration mode:

```
Area < area ID > authentication [message digest]
```

Команда отключения региональной аутентификации:

```
no area < area ID > authentication
```

Параметр area ID указывает ID области подлежащей аутентификации, вводится в десятичном формате с разделением точками a.b.c.d или целым числом, изменяется в пределах 0 или 232-1. Параметр none

устанавливает plaintext аутентификацию, параметр message digest устанавливает MD5 аутентификацию, оба параметра являются опциональными.

По умолчанию область аутентификации не сконфигурирована.

30.2.22 Stub area configuration (Конфигурация Stub области)

Команда конфигурации роутера в stub area в режиме OSPF configuration mode:

```
Area < area ID > stub [no-summary]
```

Команда отключения режима stub area:

```
no area < area ID > stub [no summary]
```

Команда установки значения cost по умолчанию для передачи ABR маршрутизации в stub area:

```
Area < area ID > default cost < cost >
```

Команда установки значения cost по умолчанию:

```
no area < area ID > default cost
```

Параметр area ID указывает ID области с атрибутом stub, вводится в десятичном формате с разделением точками a.b.c.d или целым числом, изменяется в пределах 0 или 232-1. Параметр No summary указывает на то, что маршруты внутри домена не введены в область stub area.

Первая группа команд предназначена для конфигурации роутера находящегося в области stub area. Вторая группа команд предназначена для конфигурации ABR с интерфейсом соединенным с областью stub area.

По умолчанию область stub area не сконфигурирована.

30.2.23 Configure NSSA area (Конфигурация NSSA области)

Команда конфигурации возможностей NSSA в режиме OSPF configuration mode:

```
Area < area ID > NSSA [options]
```

Команда удаления области NSSA zone:

```
no area < area ID > NSSA [options]
```

Параметр area ID указывает ID области. Параметр options определяет возможности NSSA (см. соответствующий список). По умолчанию область NSSA не сконфигурирована.

30.2.24 Configure external route aggregation (Конфигурация агрегации внешних маршрутов)

Роутеры, вводимые другими протоколами, ведут передачу один за другим в формате LSU type-5. Для определения диапазона префикса prefix range используется команда агрегации. Маршруты в пределах этого диапазона подавляются за исключением агрегированных. При большом числе внешних маршрутов, возможно эффективное уменьшение масштаба LSDB.

Команда конфигурации возможностей и масштаба агрегации в режиме OSPF configuration mode:

```
summary address < IP prefix > [not advertise | tag < tag >]
```

Команда отмены агрегации внешних маршрутов:

```
no summary address < IP prefix > [not advertise | tag < tag >]
```

Параметр IP prefix определяет диапазон агрегации маршрутов, вводится в формате a.b.c.d/m. Параметр Not advertised означает что агрегированные маршруты не будут передаваться. Параметр Tag устанавливает значение tag value, изменяется в пределах 0 или 232-1. По умолчанию установлено значение 0, введенные внешние маршруты не агрегированы.

30.2.25 Configure default weights for external routes (Конфигурация весов внешних маршрутов)

Если команда перераспределения не определяет значения при введении внешних маршрутов, используются значения весовых коэффициентов по умолчанию.

Команда конфигурации значений весов внешних маршрутов по умолчанию в режиме OSPF configuration mode:

```
default metric < metric >
```

Команда отмены значений весов внешних маршрутов по умолчанию:
no default metric [Metric]

Параметр metric value изменяется в пределах 0 или 224-1.
По умолчанию установлено значение 1.

30.2.26 Display information (Просмотр конфигурации)

В режиме normal mode или privileged mode:

Команды просмотра информации о протоколе OSPF:

show ip protocols

show ip protocols OSPF

Команда просмотра информации о процессах протокола OSPF:

show ip ospf [process ID]

Параметр process ID указывает номер процесса, изменяется в пределах 0 или 216-1.

Команда просмотра информации ABR:

show ip ospf border routes

Команда просмотра информации LSDB:

Command: show ip ospf database < type >

Параметр type относится ко всем типам LSA и краткой информации.

Команда просмотра информации OSPF интерфейса:

show ip ospf interface [if name]

Параметр if name указывает имя интерфейса.

Команда просмотра таблицы маршрутизации OSPF:

show ip ospf route [count]

Параметр count указывает общее количество entries таблицы маршрутизации.

Команда просмотра информации о виртуальных соединениях OSPF:

show ip ospf virtual links

Команда просмотра информации о соседних устройствах OSPF:

show ip ospf neighbor [options]

Параметр options см. описание команд.

В режиме privileged mode:

Команда просмотра текущей конфигурации коммутатора, включая информацию о протоколе OSPF:

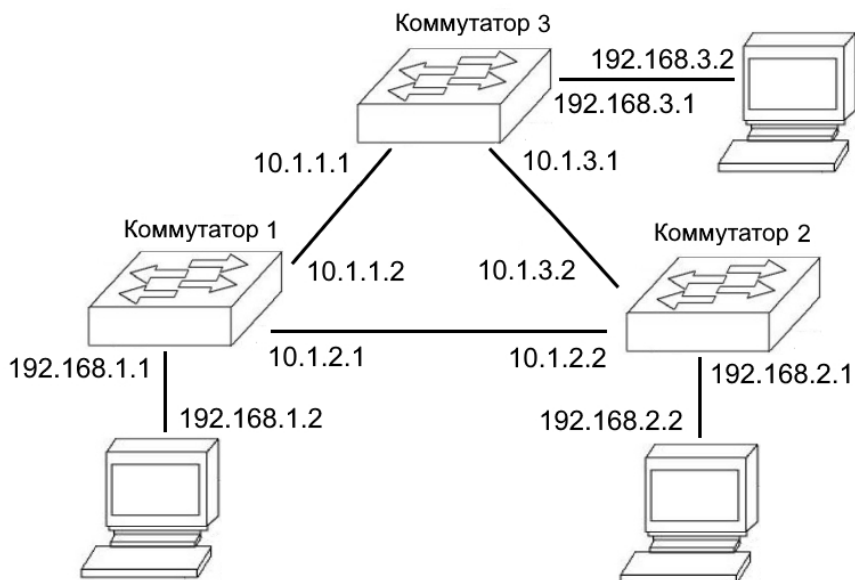
```
show running config
```

Команда просмотра текущей конфигурации протокола OSPF:

```
show running config OSPF
```

30.3 OSPF configuration example (Пример конфигурации)

Три коммутатора соединены в пары, в шесть сетевых сегментов соответственно. Протокол OSPF обеспечивает сетевое взаимодействие между тремя ПК. Требуется, чтобы интерфейс находился в области area 0



Команды конфигурации коммутатора 1:

```
Switch#configure terminal
Switch(config)#router ospf 100
Switch(config-ospf-100)#network 10.1.1.0/24 area 0
Switch(config-ospf-100)#network 10.1.2.0/24 area 0
Switch(config-ospf-100)#network 192.168.1.0/24 area 0
```

Команды конфигурации коммутатора 2:

```
Switch#configure terminal
Switch(config)#router ospf 100
Switch(config-ospf-100)#network 10.1.2.0/24 area 0
Switch(config-ospf-100)#network 10.1.3.0/24 area 0
Switch(config-ospf-100)#network 192.168.2.0/24 area 0
```

Команды конфигурации коммутатора 3:

```
Switch#configure terminal
Switch(config)#router ospf 100
Switch(config-ospf-100)#network 10.1.1.0/24 area 0
Switch(config-ospf-100)#network 10.1.3.0/24 area 0
Switch(config-ospf-100)#network 192.168.3.0/24 area 0
```

Команды просмотра конфигурации OSPF:

```
show ip ospf database
show ip ospf interface
show ip ospf neighbor
show ip route ospf
show ip ospf route
```

31. OSPFv3 Configuration (Конфигурация OSPFv3)

Глава описывает конфигурацию протокола OSPFv3 (open shortest path first) версии 3, работа протокола маршрутизации основана на IPv6 (rfc5340, rfc2740). Глава содержит следующие разделы:

- Introduction (Введение)
- OSPFv3 Configuration (Конфигурация протокола)
- Configuration example (Пример конфигурации)

31.1 OSPFv3 introduction (Введение)

Протокол OSPF (open shortest path first) является внутренним шлюзовым протоколом, основан на алгоритме статуса соединения link state algorithm. В настоящее время протокол OSPFv2 (версия 2) используется для работы с IPv4, а протокол OSPFv3 (версия 3) используется для работы с IPv6 (rfc5340, rfc2740).

Протокол маршрутизации OSPFv3 является независимой модифицированной версией протокола OSPFv2. Главная задача протокола OSPFv3 обеспечить маршрутизацию независимо от уровня сети. Для решения этой задачи была переработана внутренняя маршрутная информация протокола OSPFv3.

Отличия версий протокола OSPFv3 и OSPFv2:

- OSPFv3 не вставляет информацию основанную на IP в заголовок сообщения в начале пакета и данные о статусе соединения link state announcement (LSA).
- OSPFv3 использует информацию независимую от сетевых протоколов для представления ключевых задач IP сообщений, таких как идентификация маршрутных данных LSA.

Основные принципы используемые протоколом OSPFv3:

- Последовательность сообщений Hello message, статуса state machine, LSDB, механизма миграции flooding mechanism, вычисление маршрута и т.д.
- Разделение автономной системы на один или несколько регионов посредством соединений LSA (link), публикация маршрутов в форме объявления.
- Использование для взаимодействия между устройствами сообщений OSPFv3 для унификации маршрутной информации внутри области.
- Сообщения OSPFv3 инкапсулируются в сообщениях IPv6 и могут быть отправлены в unicast и multicast форматах.

Таблица типов сообщений протокола OSPFv3

Тип сообщения	Назначение сообщения
Hello Message	Обнаружение и обслуживание соседнего оборудования OSPFv3, отправляются периодически.

DD Message (Database Description packet)	Содержит краткую информацию о LSDB, используется для синхронизации баз данных между двумя устройствами.
LSR Message (Link State Request Packet)	Запросы LSA устройства с противоположной стороны. Устройство отправляет сообщение LSR соседнему устройству после успешного обмена сообщениями DD message.
LSU Message (Link State Update Packet)	Отправляет требуемое LSA устройству с противоположной стороны.
LSAck Message (Link State Acknowledgment packet)	Подтверждение получения LSA.

Таблица типов LSA

Тип LSA	Назначение функции LSA
Router-LSA (Type1)	Устройство генерирует LSA для зоны OSPFv3 интерфейса, которое описывает статус соединения, распространяется в зоне, которой принадлежит.
Network-LSA (Type2)	Генерируется Dr, описывает статус соединения, распространяется в зоне, которой принадлежит.
Inter-Area-Prefix-LSA (Type3)	Генерируется ABR, описывает маршрут сегмента сети в зоне, производит уведомление релевантных зон.
Inter-Area-Router-LSA (Type4)	Генерируется ABR, описывает маршрут к ASBR, производит уведомление релевантных зон за исключением зоны в которой находится ASBR.
AS-external-LSA (Type5)	Генерируется ASBR, описывает внешний маршрут, производит уведомление всех зон за исключением stub area и NSSA area.

NSSA LSA (Type7)	Генерируется ASBR, описывает внешний маршрут, распространяется только в зоне NSSA area.
Link-LSA (Type8)	Каждое устройство генерирует LSA для каждого соединения, описывает соединение местного адреса и префикс адреса соединения IPv6, обеспечивает опции соединения, которые будут отправлены в сеть LSA. Распространяется только в пределах данного соединения.
Intra-Area-Prefix-LSA (Type9)	Каждое устройство и Dr генерирует одно или более сообщения LSAS, распространяется в пределах данной зоны. - LSA генерируется оборудованием, описывает префикс адреса IPv6 связанного с маршрутом LSA. - LSA генерируется Dr, описывает префикс адреса IPv6 связанного с сетью LSA.

31.2 OSPFv3 configuration (Конфигурация OSPFv3)

31.2.1 Start OSPFv3 (Запуск протокола OSPFv3)

Одновременно может быть запущено несколько процессов OSPFv3, каждый из процессов идентифицируется по ID. При запуске протокола OSPFv3 необходимо определить какой процесс начат. Если параметры отсутствуют, то процесс имеет номер 0.

Команда запуска процесса OSPFv3 (с номером ID) в режиме global configuration mode:

```
router IPv6 OSPF [process ID]
```

Команда остановки процесса OSPF с номером ID:

```
no IPv6 Router OSPFv3 [process ID]
```

Параметр process ID определяет номер запущенного процесса, изменяется в пределах 1 или 216-1. Если параметр process ID не определен, то запускается процесс OSPFv3 с номером 0. По умолчанию протокол OSPFv3 отключен.

31.2.2 Configure router ID (Конфигурация ID роутера)

Для однозначной идентификации роутера в автономной системе используется номер ID в 32-bit формате. ID номер роутера может быть сконфигурирован вручную. В процессе конфигурации необходимо обеспечить уникальность ID номеров всех роутеров в автономной системе. Если номер роутера не определен, то используется IP адрес интерфейса обратной связи loopback interface. Если loopback не имеет IP адреса, то в качестве ID выбирается больший из IP адресов текущего интерфейса. Для обеспечения стабильности функционирования протокола OSPFv3, ID номер роутера должен быть определен уже на этапе планирования сети.

Команда конфигурации ID номера роутера в режиме OSPFv3 configuration mode:

```
router ID < router ID >
```

Команда сброса ID номера роутера:

```
no router ID
```

Команда сброса ID номера роутера:

```
no router ID [router ID]
```

Параметр router ID имеет формат a.b.c.d.

По умолчанию после запуска протокола OSPFv3, будет автоматически сгенерирован параметр router ID. При этом учитываются следующие правила: в первую очередь учитывается router ID заданный с помощью команд; при его отсутствии выбирается IP адрес интерфейса обратной связи loopback; при его отсутствии больший из IP адресов текущего интерфейса; при его отсутствии выбирается 0.0.0.0.

Две группы команд имеют одинаковое назначение.

31.2.3 Prohibit the interface from sending messages (Запрет на отправку сообщений)

В простых сетях, интерфейс протокола OSPFv3 представляет только сегмент сети между двумя устройствами для передачи данных, переводит интерфейс в пассивный статус и блокирует передачу сообщений Hello message для этого соединения, которое не влияет на информацию о маршруте интерфейса.

Команда перевода интерфейса в пассивный статус в режиме OSPFv3 configuration mode:

```
passive interface < if name >
```

Команда перевода интерфейса в обычный режим:

```
no passive interface < if name >
```

Параметр if name – имя интерфейса уровня layer-3 (vlan1, vlan2...).

По умолчанию после запуска протокола OSPFv3, включенные интерфейсы не находятся в пассивном статусе. После перевода интерфейса с запущенным протоколом OSPFv3 в пассивный статус, может быть опубликован прямой путь интерфейса, но сообщения протокола OSPFv3 на интерфейсе будут заблокированы, и интерфейс не сможет создать соединение с соседним устройством. В некоторых случаях это помогает эффективно экономить сетевые ресурсы.

31.2.4 Configure SPF calculation time (Вычисление времени SPF)

Когда происходят изменения в базе данных статуса соединений LSDB (link state database) протокола OSPFv3, требуется произвести перерасчет кратчайшего пути. Если кратчайший путь вычисляется незамедлительно для каждого изменения, то на это требуется много ресурсов, что снижает производительность роутера. Путем конфигурации значений времени задержки delay и удержания hold SPF вычисления, может быть снижена частота SPF вычислений (вызванная слишком частыми изменениями сети), и как следствие удастся снизить расходование ресурсов системы в единицу времени и повысить производительность роутера.

В процессе SPF вычисления используется таймер, который запускает каждое следующее вычисление в соответствии со временем

ингибирования inhibition time. Процесс SPF вычисления запускается по истечении времени таймера, пересчитывается время ингибирования с момента последнего SPF вычисления. Если заданное время ингибирования превышено, то для запуска таймера используется заданное время задержки delay time. Если заданное время ингибирования не превышено, то оно используется для вычисления необходимого времени задержки delay time. Если время задержки имеет значение меньше заданного, для вычисления используется заданное время, в противном случае для запуска SPF вычисления используется вычисленное время задержки.

Команда установки значений задержки и удержания (delay и hold) интервала SPF вычисления в режиме OSPFv3 configuration mode:

```
timers SPF < delay > < hold >
```

Команда установки значений задержки и удержания по умолчанию:

```
no timers SPF
```

Параметр delay устанавливает время задержки необходимое для SPF вычисления. Параметр Hold устанавливает временной интервал между двумя SPF вычислениями.

По умолчанию установлены значения: задержки delay 5 сек, удержания Hold 10 сек.

31.2.5 Configure management distance (Конфигурация дистанции управления)

На роутере может быть одновременно запущено несколько протоколов маршрутизации. Дистанция управления нужна для выбора необходимой маршрутной информации при работе нескольких протоколов маршрутизации. В том случае если разные протоколы находят один и тот же путь, то наименьшая дистанция управления является предпочтительной.

Команда конфигурации дистанции управления в режиме OSPFv3 configuration mode:

```
distance < distance >
```

Команда возврата к конфигурации по умолчанию:

```
no distance < distance >
```

Команда конфигурации различных типов дистанции управления:

```
distance OSPFv3 {intra area < distance > | inter area < distance > |  
external < distance >}
```

Команда возврата трех типов дистанции управления к значениям по умолчанию:

```
no distance OSPFv3
```

Параметр distance изменяется в пределах 1-255. Параметр Intra area относится к дистанции управления маршрутов в домене. Параметр Inter area относится к дистанции управления маршрутов между доменами. Параметр External указывает дистанцию управления внешних маршрутов.

По умолчанию значение параметра management distance протокола OSPFv3 110. Значение параметра маршрутов intra domain route, inter domain route и external route 0.

31.2.6 Introducing external routing (Введение внешней маршрутизации)

На роутере может быть одновременно запущено несколько протоколов динамической маршрутизации, различные протоколы могут делиться маршрутной информацией. Протокол OSPFv3 использует маршрутные данные других протоколов для внешних маршрутов автономной системы, которые вводятся автономной системой связанной с роутером boundary router ASBR. При введении внешнего маршрута, можно определять такие атрибуты как вес weight и тип веса weight type.

Протокол OSPFv3 имеет разделение на четыре типа маршрутизации: внутридоменная маршрутизация (intra domain routing); междудоменная маршрутизация (inter domain routing); внешняя маршрутизация 1-го типа (external route type-1); внешняя маршрутизация 2-го типа (external route type-2). Оба типа внешней маршрутизации описывают пути к точке назначения вне автономной системы. Внешняя маршрутизация типа 1 до другого IGPS. Считается что протокол OSPFv3 имеет высокую надежность и имеет вес сравнимый с весом маршрута внутри автономной системы. Таким образом, значение cost внешнего маршрута оценивается суммой

значений cost пути от самого роутера к ASBR и cost от ASBR до точки назначения. Внешняя маршрутизация типа 2 до другого eGPS. Считается что протокол OSPFv3 не имеет высокую надежность и его значение cost на много превышает вес маршрута внутри автономной системы и их значения cost сильно различаются. Таким образом, значение cost внешнего маршрута оценивается по значению cost пути от ASBR до точки назначения, а значение cost пути от роутера до ASBR не учитывается.

Команда конфигурации параметров в режиме OSPFv3 configuration mode:

```
redistribute {kernel | connected | static | rip | Isis | BGP} [Metric < metric > | metric type < type > | route map < route map name > | tag < tag >]
```

Команда возврата параметров к значениям по умолчанию:

```
no redistribute {kernel | connected | static | rip | Isis | BGP} [Metric < metric > | metric type < type > | route map < route map name > | tag < tag >]
```

Параметр 1 указывает тип вводимого внешнего маршрута, включая прямое соединение, static, rip, IS-IS and BGP. Параметр 2 определяет значение веса weight value (устанавливается при введении внешнего маршрута), изменяется в пределах 0 или 224-1. Параметр 3 определяет тип вводимого внешнего маршрута (type-1 или type-2), Type-1 соответствует IGP route, type-2 соответствует EGP route. Параметр 3 указывает имя соответствующей маршрутной карты. Маршрутная карта конфигурируется в режиме global configuration mode (см. соответствующий раздел). Параметр 4 - tag, изменяется в пределах 0 или 232-1, является атрибутом внешней маршрутизации.

По умолчанию протоколы внешней маршрутизации отключены.

31.2.7 Configure the network type of the interface (Конфигурация типа сетевого интерфейса)

Протокол OSPFv3 сам выбирает роутер, каждый роутер описывает топологию соответствующей сети и отправляет её другим роутерам. Протокол OSPFv3 разделяет сети (в зависимости от соединений интерфейса) на четыре типа согласно протоколу уровней соединений:

- broadcast type (протоколы уровня Ethernet, FDDI, и т.д.);
- NBMA non broadcast множественного доступа multiple access type (протоколы уровня fr, ATM, HDLC, X.25, и т.д.);
- точка-многоточка point to multipoint type (по умолчанию отсутствует, должен быть принудительно сконфигурирован с помощью других типов сетей). В большинстве случаев для соединения точка-многоточка изменяется протокол NBMA;
- точка-точка point-to-point type (протоколы уровня PPP, LAPB и POS).

В сетях broadcast networks без возможности множественного доступа, тип интерфейса может быть сконфигурирован как NBMA. Если не все роутеры имеют прямой доступ к сети NBMA, интерфейс может иметь конфигурацию типа точка-многоточка.

Сеть NBMA имеет полное соединение через протокол OSPFv3, также доступны non broadcast и multi-point. Для сети типа точка-многоточка отсутствует необходимость полного соединения. Для NBMA необходимо выбрать Dr, в сети точка-многоточка DR отсутствует. В сети NBMA сообщения multicast message передаются через назначенное соседнее устройство. В сети точка-многоточка используются сообщения unicast message и multicast message.

Команда конфигурации типа сети для соединения интерфейса в режиме interface configuration mode:

```
IPv6 ospf network < type >
```

Команда возврата к конфигурации по умолчанию:

```
no IPv6 ospf network
```

Параметр Type определяет тип сети: broadcast, non broadcast, точка-точка, точка-многоточка [non broadcast]; Первый тип – сеть broadcast, второй тип – сеть non broadcast (NBMA), третий тип – сеть point-to-point, четвертый тип – сеть point-to-multipoint. Сети точка-многоточка подразделяются на broadcast и non broadcast сети. Non broadcast соседние устройства не могут быть обнаружены автоматически и требуют отдельного определения.

По умолчанию установлен тип сети broadcast network.

31.2.8 Configure Hello message sending interval (Конфигурация интервалов отправки Hello сообщений)

Сообщения Hello message периодически отправляются соседнему роутеру для обслуживания взаимодействия между роутерами, выбора DR и BDR. Интервал отправки сообщений Hello message может быть установлен вручную, с учетом определенной частоты обмена сообщениями между устройствами сети. Значение времени таймера Hello timer должно быть обратно пропорционально скорости конвергенции роутера и нагрузки сети.

Команда установки времени таймера Hello timer в режиме interface configuration mode:

```
IPv6 ospf Hello interval < seconds >
```

Команда возврата к конфигурации значения по умолчанию:

```
no IPv6 ospf Hello interval
```

Параметр seconds определяет время между отправками двух сообщений Hello message, изменяется в пределах 1 или 216-1.

По умолчанию для сетей broadcast и точка-точка установлено значение 10 сек; для pBMA установлено значение более 30 сек.

31.2.9 Configure neighbor router expiration time (Конфигурация времени существования соседнего роутера)

Команда конфигурации времени существования в режиме interface configuration mod:

```
IPv6 ospf dead interval < seconds >
```

Команда возврата к значению по умолчанию:

```
no IPv6 ospf dead interval
```

Параметр seconds изменяется в пределах 1 или 216-1, соседнее устройство считается недоступным, если сообщение Hello message не получено в течение заданного периода времени. При каждом получении сообщения hello message происходит перезапуск таймера.

По умолчанию для сети broadcast и точка-точка установлено значение 40 сек. Для сети pBMA и точка-многоточка установлено значение 120

сек. При изменении типа сети значения Hello interval и dead interval возвращаются к значениям по умолчанию.

31.2.10 Configure retransmission time (Конфигурация интервала ретрансляции)

OSPFv3 это протокол надежного статуса соединения (reliable link state protocol), который отражает факт взаимодействия сообщений запроса LSU messages и ответа со стороны подключенного к другому концу линии оборудования. При получении подтверждающего сообщения, считается что обновление статуса соединения link state получено. Если подтверждающее сообщение не получено в течение интервала ретрансляции retransmission interval, сообщение LSA будет передано соседнему оборудованию. Интервал ретрансляции может быть сконфигурирован вручную. Длительность интервала должна быть больше чем время необходимое для передачи сообщения между двумя роутерами в направлении туда и обратно. Если интервал выбрать слишком коротким, то это вызовет ненужную ретрансляцию.

Команда конфигурации интервала ретрансляции интерфейса в режиме interface configuration mode:

```
IPv6 ospf retransmit interval < seconds >
```

Команда установки значения интервала по умолчанию:

```
no IPv6 ospf retransmit interval
```

Параметр seconds устанавливает длительность интервала ретрансляции в секундах, изменяется в пределах 1 и 216-1, используется в случае если подключенное оборудование не получило сообщение LSA.

По умолчанию установлено значение 5 сек.

31.2.11 Configure interface delay (Конфигурация интервала задержки интерфейса)

В сообщении обновления статуса соединения LSU message, каждый статус соединения содержит время домена time domain. Перед передачей необходимо увеличить задержку трансляции передающего

интерфейса. Этот параметр в основном учитывает время необходимое интерфейсу для отправки сообщений и особенно критичен для низкоскоростных сетей.

Команда установки времени задержки передачи интерфейсом в режиме interface configuration mode:

```
IPv6 ospf transmit delay < seconds >
```

Команда установки значения времени задержки по умолчанию:

```
no IPv6 ospf transmit delay
```

Параметр seconds устанавливает значение времени задержки в секундах, изменяется в пределах 1 и 216-1, указывает на то, что значение должно быть увеличено для age domain LSA отправленного интерфейсом.

По умолчанию установлено значение 1 сек.

31.2.12 Configure the priority of interface in Dr election (Конфигурация приоритета интерфейса при выборе Dr)

Для предотвращения повторной передачи информации о соединении в режиме точка-точка в сетях broadcast, необходимо выбрать назначенный роутер DR и BDR, который будет отвечать за отправку информации о соединении в пределах сегмента сети. Приоритет интерфейса указывает на его значение при выборе Dr. В случае возникновения конфликта, интерфейс с большим приоритетом будет считаться определяющим. Если значение приоритета интерфейса 0, то он не будет принимать участия в процессе выбора Dr. В противном случае интерфейсы становятся кандидатами. Каждый роутер обладает своей собственной информацией о приоритете и его собственный DR отправляет сообщения Hello message в сеть, и в конечном итоге выбирается один Dr с наибольшим приоритетом. В случае равного приоритета, роутер с наивысшим ID получает преимущество.

При сбое Dr, роутерам в сети необходимо пройти повторную процедуру выбора Dr, что занимает некоторое время, в течении которого, могут произойти ошибки вычисления маршрутов. Задача BDR обеспечить плавный переход к новому Dr, т.е. BDR является резервным Dr. Его выбор происходит одновременно с выбором Dr. Этот процесс обеспечивает смежное взаимодействие с другими роутерами в сети.

В отличие от BDR, Dr обеспечивает сбор информации и публикацию нодов в сети. BDR обеспечивает только синхронизацию смежных устройств. При отключении Dr, BDR немедленно становится DR и обеспечивает сбор информации в пределах сегмента сети. В это же время запускается новый процесс выбора BDR, но этот процесс не оказывает влияния на вычисление маршрута.

Команда установки приоритета интерфейса при выборе Dr в режиме interface configuration mode:

```
IPv6 ospf priority < prio >
```

Команда установки значения приоритета по умолчанию:

```
no IPv6 ospf priority
```

Параметр prio устанавливает значение приоритета при выборе Dr, изменяется в пределах 0-255. При значении приоритета 0 интерфейс не участвует в процессе выбора.

По умолчанию установлено значение 1.

31.2.13 Configure the cost of sending messages on the interface (Конфигурация важности сообщений при отправке интерфейсом)

В сетях управление трафиком осуществляется на основе настройки соединений и разделения их по важности. Значение cost интерфейса соответствует значению cost отправляемых им сообщений. Если эти значения не сконфигурированы вручную, протокол OSPFv3 автоматически вычислит значение интерфейса interface cost согласно baud rate интерфейса.

Команда установки значения параметра в режиме interface configuration mode:

```
IPv6 ospf cost < cost >
```

Команда установки значения параметра по умолчанию:

```
no IPv6 ospf cost
```

Параметр cost определяет значение сообщения отправленного интерфейсом, изменяется в пределах 1 и 216-1.

По умолчанию установлено значение 10.

31.2.14 MTU value for DD message sent by interface (Значение MTU value для сообщений DD отправленных интерфейсом)

Команда отключения проверки значения MTU в сообщении DD message в режиме interface configuration mode:

```
IPv6 ospf MTU ignore
```

Команда включения проверки значения MTU в сообщении DD message:

```
no IPv6 ospf MTU ignore
```

По умолчанию проверка значения MTU включена.

31.2.15 Configure regional virtual link (Конфигурация регионального виртуального соединения)

Протокол OSPFv3 использует принцип разделения роутеров по уровням на разные группы в автономной системе. Эти группы называются регионами. Все регионы не являются равными и параллельными и имеют иерархическое взаимодействие. Среди них выделяется регион 0.0.0.0, который называется осевым или позвоночным (backbone region). Другие не backbone regions должны обмениваться своими доменными маршрутами через регион backbone region. Т.о. все не backbone области должны быть соединены с backbone областью, соответственно, один интерфейс ABR находится в области 0. Если какие-либо области не могут гарантировать физический путь к области backbone по причине ограничения сетевой топологии, то должны быть сконфигурированы виртуальные соединения для создания логического пути. Оба конца виртуального маршрута принадлежат ABR, средняя часть пути не проходит через backbone area и называется областью передачи (transmission area). При конфигурации виртуального соединения необходимо определить ID области передачи (transmission area) и ID оппозитного ABR, которые конфигурируются на ABR с обоих концов в первую очередь.

Вычисление пути через область передачи transmission area и создание виртуального соединения является логическим эквивалентом соединения точка-точка между конечными точками маршрута. Т.о. параметры интерфейса могут быть сконфигурированы на физическом интерфейсе, и может быть запущена функция аутентификации.

Сообщение unicast message передается между ABRs. Роутер передающий сообщение unicast message через область transmission area передает его как обычное IP сообщение. Т.о. подтверждается создание логического соединения в области передачи transmission area, и сообщения протокола могут передаваться между двумя ABRs.

Команда конфигурации transmission area и peer ID виртуального соединения в режиме OSPFv3 configuration mode:

```
Area < area ID > virtual link < router ID >
```

Команда конфигурации Hello interval виртуального соединения:

```
hello-interval <seconds> |
```

Команда конфигурации failure time виртуального соединения:

```
dead-interval <seconds> |
```

Команда конфигурации retransmission interval виртуального соединения:

```
retransmit-interval <seconds> |
```

Команда конфигурации interface delay виртуального соединения:

```
transmit-delay <seconds> |
```

Команда удаления виртуального соединения:

```
no area < area ID > virtual link < router ID > | hello-interval <seconds> |  
dead-interval <seconds> | retransmit-interval <seconds> | transmit-delay  
<seconds>]
```

Параметр area ID определяет ID области передачи, вводится в десятичном формате с разделением точками a.b.c.d или целым числом, изменяется в пределах 0 или 232-1. Параметр Router ID определяет ID peer роутера виртуального соединения, вводится в десятичном формате с разделением точками a.b.c.d. Конфигурация параметров аутентификации и отправляющего интерфейса являются опциональными возможностями (см. соответствующие разделы).

По умолчанию виртуальное соединение не сконфигурировано.

31.2.16 Configure regional routing aggregation (Конфигурация агрегации региональной маршрутизации)

Команда конфигурации масштаба агрегирования в режиме OSPFv3 configuration mode:

Area < area ID > range < IPv6 prefix > [advertise | not advertise]

Команды отмены масштаба агрегирования:

no area < area ID > range < IPv6 prefix > [advertise | not advertise]

Параметр area ID соответствует ID области, определяет агрегацию маршрутов в области, вводится в десятичном формате с разделением точками a.b.c.d или целым числом, изменяется в пределах 0 или 232-1. Параметр IPv6 prefix определяет диапазон агрегации, вводится в формате X: X: X: X / m. Параметры advertise и not advertise указывают на необходимость или отсутствие необходимости отправлять диапазон агрегации (IPv6 prefix). По умолчанию отправляется информация об изначальных сетевых маршрутах.

31.2.17 Stub area configuration (Конфигурация Stub области)

Команда конфигурации роутера в stub area в режиме OSPFv3 configuration mode:

Area < area ID > stub [no-summary]

Команда отключения режима stub area:

no area < area ID > stub [no summary]

Команда установки значения cost по умолчанию для передачи ABR маршрутизации в stub area:

Area < area ID > default cost < cost >

Команда установки значения cost по умолчанию:

no area < area ID > default cost

Параметр area ID указывает ID области с атрибутом stub, вводится в десятичном формате с разделением точками a.b.c.d или целым числом, изменяется в пределах 0 или 232-1. Параметр No summary указывает на то, что маршруты внутри домена не введены в область stub area.

Первая группа команд предназначена для конфигурации роутера находящегося в области stub area. Вторая группа команд предназначена для конфигурации ABR с интерфейсом соединенным с областью stub area.

По умолчанию область stub area не сконфигурирована.

31.2.18 Configure default weights for external routes (Конфигурация весов внешних маршрутов)

При импортировании внешних маршрутов, не используется команда для установки значения red metric по умолчанию.

Команда конфигурации значений весов внешних маршрутов по умолчанию в режиме OSPFv3 configuration mode:

```
default metric < metric >
```

Команда отмены значений весов внешних маршрутов по умолчанию:

```
no default metric [Metric]
```

Параметр metric value изменяется в пределах 0 или 224-1.

По умолчанию установлено значение 1.

31.2.19 Display information (Просмотр конфигурации)

В режиме privileged mode:

Команда просмотра информации о процессах протокола OSPF:

```
show IPv6 ospf [process ID]
```

Параметр process ID указывает номер процесса, изменяется в пределах 0 или 216-1.

Команда просмотра информации LSDB:

```
Command: show IPv6 ospf database < type >
```

Параметр type относится ко всем типам LSA и краткой информации.

Команда просмотра информации OSPFv3 интерфейса:

```
show IPv6 ospf interface [if name]
```

Параметр if name указывает имя интерфейса.

Команда просмотра таблицы маршрутизации OSPF:

```
show ip ospf route
```

Команда просмотра информации о виртуальных соединениях OSPF:

```
show IPv6 ospf virtual links
```

Команда просмотра информации о соседних устройствах OSPF:

```
show IPv6 ospf neighbor [options]
```

Параметр options см. описание команд.

31.3 OSPFv3 configuration example (Пример конфигурации)

Два коммутатора соединены, как показано на рисунке ниже:



Команды конфигурации коммутатора S1:

```
Switch(config)#  
Switch(config)#int vlan1  
Switch(config-vlan1)#ipv6 address 3ffe:506::1/64  
Switch(config-vlan1)#exit  
Switch(config)#router ipv6 ospf  
Switch(config-router)#router-id 1.1.1.1  
Switch(config-router)#exit  
Switch(config)#int vlan1  
Switch(config-vlan1)#  
Switch(config-vlan1)#ipv6 router ospf area  
Switch(config-vlan1)#
```

Команды конфигурации коммутатора S2:

```
Switch(config)#  
Switch(config)#int vlan1  
Switch(config-vlan1)#ipv6 address 3ffe:506::2/64  
Switch(config-vlan1)#exit  
Switch(config)#router ipv6 ospf  
Switch(config-router)# router-id 5.5.5.5  
Switch(config-router)#exit  
Switch(config)#int vlan1  
Switch(config-vlan1)#  
Switch(config-vlan1)#ipv6 router ospf area 0  
Switch(config-vlan1)#
```

Команда просмотра конфигурации OSPFv3:

```
Switch#show ipv6 ospf neighbor
```

OSPFv3 Process (*null*)

Neighbor ID	Pri	State	Dead Time	Interface	Instance ID
5.5.5.5	1	Full/DR	00:00:39	vlan1	0

```
Switch#
```

32. BGP Configuration (Конфигурация BGP)

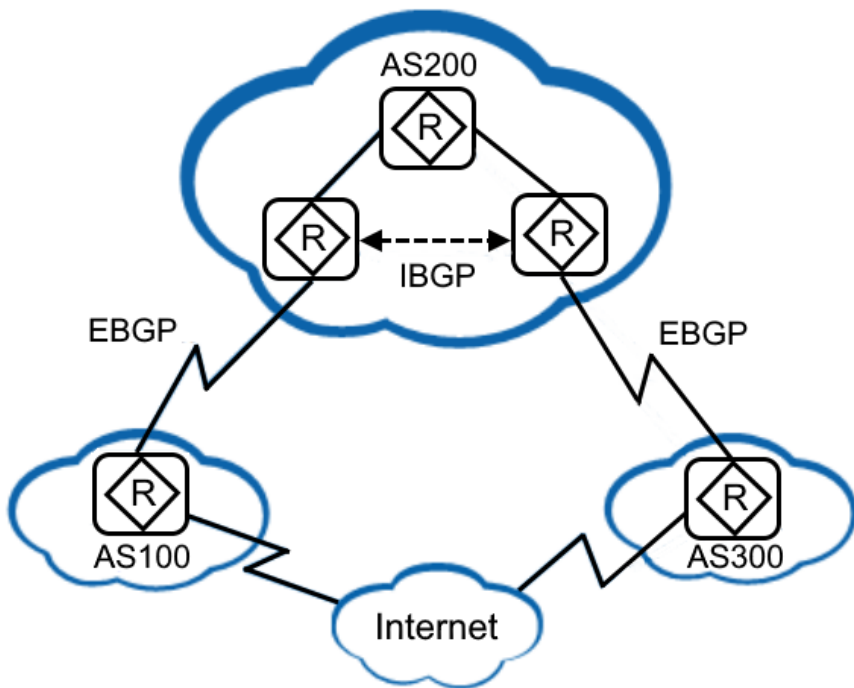
Глава описывает конфигурацию протокола BGP (border gateway protocol). BGP – протокол дистанционной векторной маршрутизации, который обеспечивает доступность маршрутов между автономными системами и при этом выбирает наилучший путь. Глава содержит следующие разделы:

- Introduction (Введение)
- BGP Configuration (Конфигурация протокола)
- Configuration example (Пример конфигурации)

32.1 BGP introduction (Введение)

Border Gateway Protocol (BGP) - протокол дистанционной векторной маршрутизации, который обеспечивает доступность маршрутов между автономными системами и при этом выбирает наилучший путь. Существуют четыре ранние версии протокола: BGP-1 (rfc1105), BGP-2 (rfc1163), BGP-3 (rfc1267) и BGP-4 (rfc1771). После 2006 года в сетях unicast IPv4 используется версия BGP-4 (rfc4271), в прочих сетях (таких как IPv6) используется версия MP-BGP (rfc4760).

Протокол BGP имеет разделение на ebgp (внешний / external BGP) и ibgp (внутренний / internal BGP) режимы работы (см. рис. ниже).



- EBGP – протокол BGP запущенный между разными автономными системами. Для предотвращения возникновения петлевых соединений между автономными системами, при получении маршрута от ebgp peer оборудование BGP сбрасывает маршрут с номером локальной автономной системы.

- IBGP - протокол BGP запущенный в пределах одной автономной системы. Для предотвращения возникновения петлевых соединений в автономной системе, оборудование BGP не анонсирует маршрут полученный от iBGP пиров для других iBGP пиров, и создает полное соединение со всеми iBGP пирами. Для решения проблемы возникновения слишком большого количества соединений iBGP пиров, протокол BGP создает routing reflector и BGP alliance.

Существует пять видов сообщений, шесть видов state machines и пять принципов интерактивных процессов создания обновления и удаления BGP пиров.

BGP пиры взаимодействуют друг с другом используя пять видов сообщений, из которых только сообщения `keepalive messages` отправляются периодически:

- Open message: используются для создания соединения между BGP пирами.
- Update message: используются для обмена маршрутной информацией между пирами.
- Notification message: используются для прерывания BGP соединения.
- Keepalive message: используются для обслуживания BGP соединения.
- Route refresh message: используются для запроса повторной отправки маршрутной информации пиром после изменения маршрутной политики. Только оборудование BGP, которое поддерживает обновление маршрутов, может отправить ответное сообщение на этот запрос.

Существует шесть типов статусов `state machines` для обеспечения процесса взаимодействия BGP пиров: `idle`, `connect`, `active`, `open message sent (opensent)`, `open message confirmed (openconfirm)` и `connection established (established)`. В процессе создания BGP пиров, используется три видимых всеми статуса: `idle`, `active` и `established`.

Оборудование BGP добавляет оптимальный маршрут в таблицу маршрутов BGP routing table для формирования маршрута BGP route. После того как оборудование BGP создаст соединение с соседним пиром, вступают в силу следующие принципы взаимодействия:

- BGP маршрут полученный от iBGP пира, оборудование BGP предоставляет его только ebgp пиру.
- BGP маршруты полученные от ebgp пиров, оборудование BGP предоставляет их всем ebgp и iBGP пирам.
- Если существует несколько маршрутов к одному адресу назначения, оборудование BGP предоставляет только оптимальный маршрут к peer.
- При обновлении маршрута, оборудование BGP отправляет только обновленный BGP маршрут.
- BGP оборудование получает все маршруты отправленные пирами.

Протоколы BGP и IGP используют разные маршрутные таблицы на одном устройстве. Для реализации взаимодействия между разными автономными системами, протокол BGP должен взаимодействовать с протоколом IGP, таким образом, маршрутные таблицы BGP и IGP взаимно вводятся друг в друга.

Протокол BGP не ищет маршруты самостоятельно, т.е. прочие маршруты должны быть введены в маршрутную таблицу BGP для обеспечения взаимодействия между автономными системами. Когда автономная система должна передать маршрутную информацию другой автономной системе, edge роутер автономной системы вводит маршруты IGP в маршрутную таблицу BGP. Когда протокол BGP вводит маршруты IGP, для улучшения планирования сети используется стратегия маршрутизации для фильтрации маршрута. Для определения момента получения трафика автономной системой устанавливается маршрутный атрибут, или значение Med для ebgp пира.

Протокол BGP поддерживает импорт и сетевые соединения при вводе маршрутов:

- При обычном импорте маршрутов в маршрутную таблицу BGP используется методы протоколов rip, OSPF, Isis и др согласно типу протокола. Для обеспечения эффективности ввода IGP маршрута, могут использоваться static route и direct route.
- При сетевом методе используется последовательный ввод уже существующих маршрутов в маршрутные таблицы IP и BGP один за другим, что обеспечивает большую точность, чем при обычном импорте.

Ввод BGP маршрутизации протоколом IGP.

Когда автономная система должна получить маршруты от другой автономной системы, edge роутер автономной системы вводит маршруты BGP в маршрутную таблицу IGP. Для того чтобы избежать получения устройством автономной системы большого количества маршрутов BGP одновременно, при получении протоколом IGP маршрутов BGP, может быть использована маршрутная политика для фильтрации маршрутов и установки маршрутных атрибутов.

32.2 BGP configuration (Конфигурация BGP)

32.2.1 Start BGP and enter BGP configuration mode (Запуск протокола BGP)

Для идентификации протокола BGP используется нумерация, при запуске протокола необходимо определить автономную систему на которой запускается протокол BGP.

Команда запуска протокола BGP в режиме global configuration mode:

```
router BGP < as number >
```

Команда отключения протокола BGP и удаление его конфигурации:

```
no router BGP < as number >
```

Параметр as number определяет номер автономной системы, изменяется в пределах 1-4294967295.

По умолчанию протокол BGP отключен.

32.2.2 Inject routing information into BGP protocol (Ввод маршрутной информации в протокол BGP)

Сразу после запуска протокола маршрутная информация BGP отсутствует. Существует два способа ввода маршрутной информации в BGP. Первый способ предоставляет возможность ручного ввода маршрутной информации в BGP с помощью команд. Второй способ ввода использует взаимодействие через протокол IGP (маршрутная информация вводится протоколом IGP). Протокол BGP сообщает о введенных маршрутах соседним устройствам. Ниже описан ручной метод ввода маршрутной информации.

Команда установки конфигурации сетевой информации для публикации BGP спикером в режиме BGP configuration mode:

```
network < network number > mask < mask > [route map < map tag >]  
[backdoor]
```

Команда отмены установки сетевой информации:

```
no network < network number > mask < mask > [route map] [backdoor]
```

32.2.3 Configure BGP timer (Конфигурация таймера BGP)

Протокол BGP использует таймер времени существования keepalive для обслуживания и обеспечения эффективного соединения с пирами и таймер holdtime для определения состояния пиров. Когда установлено BGP соединение между BGP спикерами с обеих сторон, выбирается меньшее значение таймера holdtime, в то время как значение таймера keepalive должно быть выбрано менее 1/3 от значений holdtime и установленного keepalive.

Команда установки сетевых таймеров BGP в режиме BGP configuration mode:

```
timers BGP < keepalive > < holdtime >
```

Команда установки значений таймеров по умолчанию:

```
no timers BGP < keepalive > < holdtime >
```

По умолчанию установлены значения keepalive таймера 60 сек и holdtime таймера 180 сек.

32.2.4 Configure BGP and IGP synchronization (Синхронизация протоколов BGP и IGP)

Публикация маршрутной информации которая пересекает одну автономную систему на пути к другой автономной системе, разрешается в том случае когда гарантируется, что все роутеры автономной системы получили маршрутную информацию. В противном случае (если роутер с запущенным протоколом IGP не получил маршрутную информацию) при прохождении данных через автономную систему, роутер может сбросить пакеты данных, что повлечет за собой потерю данных (феномен образования маршрутной черной дыры).

Необходимо гарантировать, что бы все роутеры автономной системы получали внешнюю маршрутную информацию. Этот процесс называется синхронизацией протоколов BGP и IGP. Синхронизацию выполняет BGP speaker путем перераспределения всех роутеров между протоколами BGP и IGP, таким образом роутеры автономной системы получат необходимую маршрутную информацию.

Существуют два случая, когда механизм синхронизации протокола BGP может быть отключен:

- Маршрутная информация не проходит через автономную систему (автономная система является терминальной (terminal as));
- На всех роутерах автономной системы запущен протокол BGP и все спикеры BGP speakers создают взаимосвязи с полным соединением (спикеры BGP speakers создают смежные взаимосвязи).

Команда запуска механизма синхронизации маршрутной информации протоколов BGP и IGP в режиме BGP configuration mode:
synchronization

Команда отмены синхронизации маршрутной информации протоколов BGP и IGP:
no synchronization

32.2.5 Configure interaction between BGP and IGP (Конфигурация взаимодействия между протоколами BGP и IGP)

Ввод маршрутной информации сгенерированной протоколом IGP в протокол BGP. Для представления маршрута по умолчанию используются команды приведенные ниже.

Команда представления маршрутной информации другого протокола протоколу BGP в режиме BGP configuration mode:
redistribute < protocol type > [route map < map tag >]

Команда отключения функции и отмены параметров конфигурации:
no redistribute < protocol type > [route map < map tag >]

По умолчанию функция представления маршрута отключена.

32.2.6 Selection of BGP optimal path (Выбор оптимального пути BGP)

Выбор оптимального пути является важной частью протокола BGP. Ниже описывается процесс выбора оптимального маршрута протоколом BGP. Выбор осуществляется в следующей последовательности.

В случае если entry маршрутной таблицы является некорректным, то участие в выборе оптимального маршрута не произойдет.

- Выбор локального роутера с высоким значением high pref атрибута;
- Выбор пути сгенерированного спикером BGP speaker;

Роутеры выбранные BGP спикером используют сгенерированные сетью команды представления и агрегации (redistribute command и aggregate command).

- Выбор маршрута с кратчайшей длиной автономной системы;
- Выбор маршрута с наименьшим значением атрибута;
- Выбор маршрута с наименьшим значением Med value;
- Приоритет пути ebgp path выше чем у iBGP path, или приоритеты одинаковые;
- Выбор пути с наименьшим значением IGP metric до следующего hop;
- Если маршрут предоставлен ebgp, выбирается более ранний маршрут;
- Выбирается маршрут BGP спикера с наименьшим ID роутера;
- Выбирается маршрут с наибольшей длиной выбирающего кластера;
- Выбирается маршрут с наибольшим адресом соседнего устройства.

Команда включения функции выбора оптимального BGP маршрута в режиме BGP configuration mode:

```
BGP bestpath < as path ignore | compare confirmed aspath | compare routerid | don compare originator ID | Med [confirmed | missing as worst | remove rcv Med | remove send Med] | tie break on age >
```

Команда отключения функции выбора оптимального BGP маршрута:

```
no BGP bestpath < as path ignore | compare confirmed aspath | compare routerid | don compare originator ID | Med [confirmed | missing as worst | remove rcv Med | remove send Med] | tie break on age >
```

32.2.7 Configure BGP routing aggregation (Агрегация маршрутизации BGP)

Протокол BGP-4 поддерживает CIDR, таким образом допускается создание таблицы агрегации для уменьшения маршрутной таблицы BGP. Таблица агрегации BGP добавляется в маршрутную

таблицу BGP только когда существует действующий маршрут в пределах диапазона агрегации.

Команда создания таблицы агрегации BGP IPv4/IPv6 в режиме configuration mode:

```
aggregate address < IP address > / < mask length > [as set | summary only]
```

Команда удаления таблицы агрегации:

```
no aggregate address < IP address > / < mask length > [as set | summary only]
```

32.2.8 Configure BGP routing reflector (Конфигурация BGP routing reflector)

Для увеличения скорости конвергенции маршрутной информации, все BGP спикеры автономной системы обычно устанавливают полное соединение для взаимодействия (два смежных BGP спикера устанавливают взаимосвязь). Большое количество BGP спикеров в автономной системе приводит к повышению нагрузки, увеличению количества выполняемых задач по сетевой конфигурации, что негативно сказывается на быстродействии сети.

Функция routing reflector использует два метода для уменьшения числа iBGP пиров в автономной системе. Функция Routing reflector уменьшает количество соединений iBGP пиров в автономной системе. BGP спикер переводится в режим routing reflector, который разделяет iBGP пиры автономной системы на две категории: клиент и не клиент.

Функция routing reflector реализованная в автономной системе действует по следующим правилам:

Сначала конфигурируется routing reflector и определяется её клиент из группы, устанавливается соединение и взаимосвязь. Клиент внутреннего маршрута группы routing reflector не должен создавать соединения взаимосвязанные с другими BGP спикерами вне группы.

В автономной системе полное взаимосвязанное соединение создается между iBGP пирами не клиентов. В данном случае, iBGP пиры не клиентов включают следующие варианты: Между несколькими routing reflectors в группе; Routing reflectors и группы в одной группе BGP

спикеры не участвуют в работе функции routing reflector (обычно эти BGP спикеры не поддерживают функцию routing reflector); Между routing reflectors в одной группе и routing reflectors в другой группе.

Правила процессинга маршрута полученного route reflector следующие:

- Обновление маршрута полученное от ebgp спикера отправляется всем клиентам и не клиентам;
- Обновление маршрута полученное от клиента отправляется другим клиентам и не клиентам;
- Обновление маршрута полученное от iBGP не клиента отправляется всем их клиентам.

Команда конфигурации устройства как routing reflector и его клиентов в режиме BGP configuration mode:

```
neighbor [peer address | peer tag] route reflector client
```

Команда удаления установленного клиента:

```
no neighbor [peer address | peer tag] route reflector client
```

32.2.9 Configure BGP as Federation (Конфигурация BGP Federation)

Функция BGP Federation предоставляет ещё один способ уменьшить количество соединений iBGP пиров в автономной системе.

Автономная система разделяется на несколько автономных подсистем sub autonomous systems, объединенных в альянс путем установления унифицированного alliance ID (присвоенный номер). Для внешних объединений, альянс является автономной системой со своим определенным номером. Внутри альянса, полное соединение iBGP пира устанавливается между BGP спикерами в пределах автономной подсистемы, ebgp соединение устанавливается между BGP спикерами автономных подсистем. Несмотря на это, при обмене информацией о next_Hop, Med и local_Path атрибутах (таких как префикс) данные параметры остаются неизменными.

Команда конфигурации ID Federation автономной системы в режиме BGP configuration mode:

```
BGP confirmation identifier < as number >
```

Команда удаления ID Federation:

```
no BGP confirmation identifier
```

32.2.10 Configure BGP management distance (Конфигурация дистанции управления BGP)

Параметр management distance указывает надежность источника маршрутной информации, изменяется в пределах 1 - 255. Чем выше значение параметра, тем ниже надежность.

BGP устанавливает разные значения management distances для различных источников маршрутной информации, которые подразделяются на external distance (внешние), internal distance (внутренние) и local distance (локальные):

- External distance: получают от ebgp пиров.
- Internal distance: получают от iBGP пиров peers.
- Local distance: получают от пиров, но при этом считается что лучше получить маршрутную информацию от IGP.
- Manage distances: обычно эти маршруты предоставляются служебными сетевыми командами (network backdoor command).

Команда установки management distance для различных типов BGP маршрутов в режиме BGP configuration mode:

```
distance BGP <external distance> < internal distance > < local distance >
```

Команда установки значения management distance по умолчанию:

```
no distance BGP < external distance > < internal distance > < local distance >
```

32.2.11 Configure BGP routing update mechanism (Конфигурация механизма обновления BGP маршрутизации)

Механизм обновления BGP маршрутизации включает в себя два этапа: периодическое обновление (regular scan update) и обновление при произошедшем изменении (event triggered update). При Scan update по расписанию BGP использует таймер для периодического запуска сканирования для обновления маршрутной таблицы. При event triggered update, если пользователь вносит изменение в конфигурацию BGP (изменение next hop BGP маршрута), автоматически запускается сканирование для обновления маршрутной таблицы.

Команда установки таймера периодического сканирования протокола BGP в режиме BGP configuration mode:

BGP scan time < time >

Команда отключения таймера периодического сканирования:

Command: no BGP scan time [time]

По умолчанию установлено значение таймера периодического сканирования 60 сек.

32.2.12 Configure BGP local autonomous system (Конфигурация BGP локальной автономной системы)

Локальная автономная система BGP используется для конфигурации отличной от реальной автономной системы (основанной на роутере) для определенного пира. Локальная автономная система имеет виртуализацию и пир-оборудование, таким образом при изменениях BGP конфигурации, пир-оборудование не изменяется для создания соединений BGP.

Миграция и объединение автономной системы для больших сетей может обеспечить отсутствие влияния на конфигурацию прочего связанного оборудования.

При работе BGP протокола, когда локальное устройство создает BGP соединение с пиром, оно сообщает пир-оборудованию номер локальной автономной системы через открытое сообщение. Пир-оборудование проверяет соответствует ли BGP анонсированная автономная система локальной автономной системе сконфигурированной пиром. При несоответствии автономных систем, BGP соединение будет разорвано. По умолчанию локальная автономная система BGP соединения является реальной автономной системой BGP. Путем конфигурации локальной автономной системы для пира, локальное оборудование будет использовать её для замены реальной автономной системы и создания BGP соединения с пиром.

Команда установки номера локальной автономной системы для BGP пира в режиме BGP configuration mode:

neighbor [peer address | peer tag] local as < as number >

Одновременно пир-оборудование может использовать локальную автономную систему как удаленную локальную систему для создания BGP соединения с оборудованием.

Команда удаления конфигурации локальной автономной системы:
no neighbor [peer address | peer tag] local as < as number >

32.3 BGP configuration example (Пример конфигурации)

Два коммутатора соединены, как показано на рисунке ниже:



Команды конфигурации коммутатора S1:

```
Switch(config)#  
Switch(config)#int vlan1  
Switch(config-vlan1)#ip address 192.168.0.1/64  
Switch(config-vlan1)#exit  
Switch(config)#router bgp 1  
Switch(config-router)#bgp router-id 1.1.1.1  
Switch(config-router)#neighbor 192.168.0.2 remote-as 1  
Switch(config-router)#neighbor 192.168.0.2 interface vlan1  
Switch(config-router)# neighbor 192.168.0.2 next-hop-self  
Switch(config-router)#exit
```

Команды конфигурации коммутатора S2:

```
Switch(config)#  
Switch(config)#int vlan1  
Switch(config-vlan1)#ip address 192.168.0.2/64  
Switch(config-vlan1)#exit  
Switch(config)#router bgp 1  
Switch(config-router)#bgp router-id 2.2.2.2  
Switch(config-router)#neighbor 192.168.0.1 remote-as 1  
Switch(config-router)#neighbor 192.168.0.1 interface vlan1  
Switch(config-router)# neighbor 192.168.0.1 next-hop-self  
Switch(config-router)#exit
```

Команда просмотра конфигурации BGP:

```
Switch#show bgp neighbors
```

BGP neighbor is 192.168.0.2, remote AS 1, local AS 1, internal link

BGP version 4, remote router ID 2.2.2.2

BGP state = Established, up for 00:02:09

Last read 00:02:09, hold time is 90, keepalive interval is 30 seconds

Neighbor capabilities:

Route refresh: advertised and received (old and new)

Address family IPv4 Unicast: advertised and received

Received 7 messages, 0 notifications, 0 in queue

Sent 6 messages, 0 notifications, 0 in queue

Route refresh request: received 0, sent 0

Minimum time between advertisement runs is 5 seconds

For address family: IPv4 Unicast

BGP table version 2, neighbor version 2

Index 1, Offset 0, Mask 0x2

NEXT_HOP is always this router

Community attribute sent to this neighbor (both)

1 accepted prefixes

0 announced prefixes

Connections established 1; dropped 0

Local host: 192.168.0.1, Local port: 53557

Foreign host: 192.168.0.2, Foreign port: 179

Nexthop: 192.168.0.1

Nexthop global: fe80::82:44ff:fe13:4803

Nexthop local: ::

BGP connection: non shared network

```
Switch#
```

33. VRRP Configuration (Конфигурация VRRP)

Глава описывает конфигурацию протокола VRRP (Virtual Router Redundancy Protocol). VRRP - протокол резервирования виртуального маршрутизатора. Глава содержит следующие разделы:

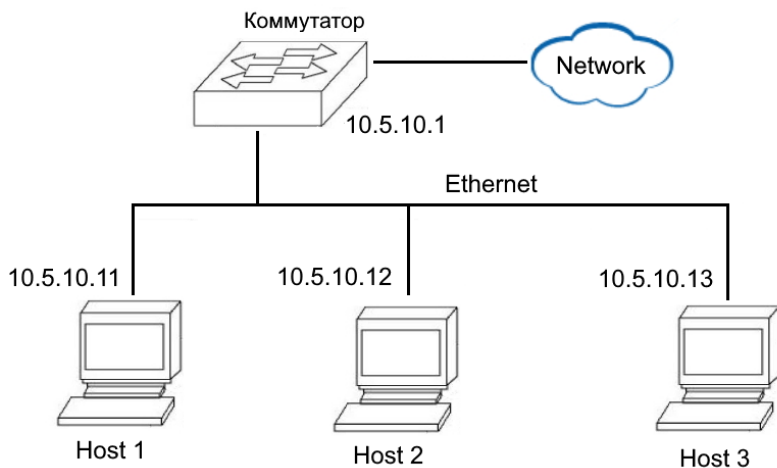
- Introduction (Введение)
- VRRP Configuration (Конфигурация протокола)
- Configuration example (Пример конфигурации)

33.1 VRRP introduction (Введение)

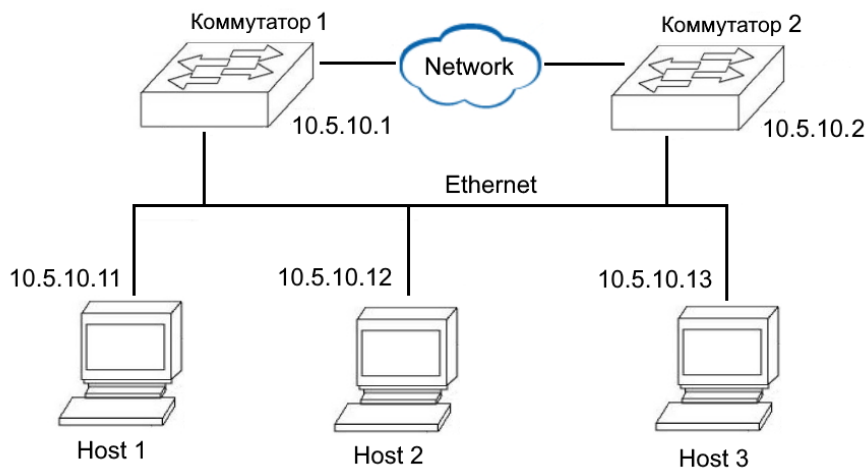
VRRP (протокол резервирования виртуального маршрутизатора) - сетевой протокол применяется для повышения доступности маршрутизатора, используемого в роли шлюза по умолчанию. Работа протокола базируется на объединении нескольких маршрутизаторов в виртуальное устройство на логическом уровне, что позволяет, в случае выхода из строя одного из маршрутизаторов, не производить переконфигурацию клиентских устройств.

33.1.1 VRRP summary (Принцип работы протокола VRRP)

На рисунке ниже показана типовая схема сети intranet. Один интерфейс коммутатора соединен с внешней сетью, а другой интерфейс соединен с внутренней сетью. IP адрес интерфейса, соединенного с внутренней сетью 10.5.10.1. Хосты 1, 2 и 3 имеют IP адреса в пределах одного сегмента сети 10.5.10.0/24. Шлюз по умолчанию сконфигурирован на хостах 1, 2 и 3. Точки next hop до коммутатора и IP адрес next hop 10.5.10.1. Следовательно, когда хост отправляет сообщение на IP адрес за пределами сети, то он будет соответствовать маршруту по умолчанию и будет отправлен коммутатору. Далее коммутатор перенаправляет сообщение из внешней сети соответствующему хосту. Таким образом хост реализует взаимодействие с внешней сетью.



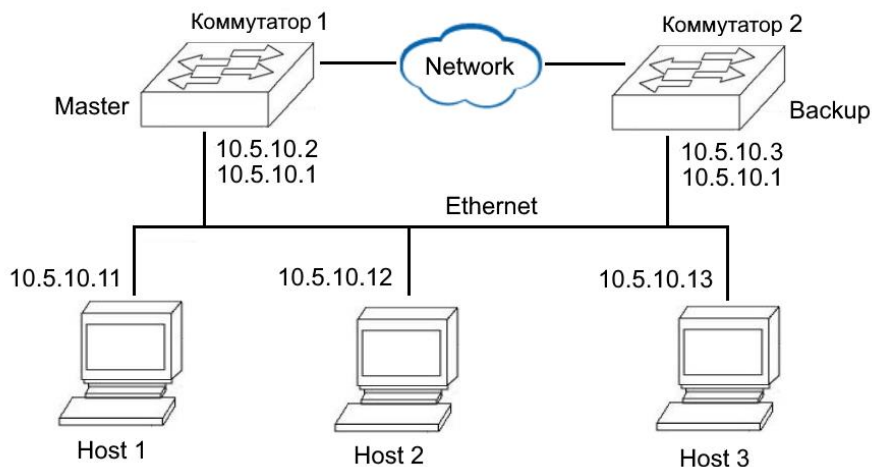
В приведенной выше схеме, коммуникация между хостами и внешней сетью возможна только через единственный коммутатор. При отказе коммутатора соединение всех хостов с внешней сетью разрывается. Для решения этой проблемы необходимо увеличить число коммутаторов до двух и более и использовать протокол динамической маршрутизации OSPF или rip между хостами и коммутатором, как показано на рисунке ниже.



Когда на хосте запущен протокол динамической маршрутизации, хост может получать информацию обо всех маршрутах к внешней сети. При коммуникации с внешней сетью, хост находит маршрут согласно IP адресу назначения сообщения и получает next hop чтобы определить куда отправить сообщение (коммутатору 1 или коммутатору 2). При сбое одного из коммутаторов, маршрут хоста может быть изменен в течение очень короткого времени, и next hop маршрута будет указан коммутатору без сбоя, так что соединение между хостом и внешней сетью не будет прервано.

В некоторых случаях не представляется возможным применить протокол динамической маршрутизации на хосте по причине слишком высокой нагрузки создаваемой протоколом. Работа протокола динамической маршрутизации на хосте также создает дополнительный трафик данных в сети, кроме того, некоторые хосты вообще не поддерживают работу протокола.

Использование специально разработанного протокола VRRP позволяет решить проблему в случае сбоя работы единственного коммутатора. Как показано на рисунке ниже, коммутатор 1 и коммутатор 2 представляют собой виртуальный роутер. Коммутаторы имеют реальные IP адреса интерфейсов и общий виртуальный IP адрес 10.5.10.1.



По умолчанию адрес шлюза хоста установлен как виртуальный IP адрес 10.5.10.1. Когда коммутатор 1 является виртуальным мастер-устройством, коммуникация между хостом и внешней сетью проходит через него. В случае если на коммутаторе 1 происходит сбой, коммутатор 2 начинает выполнять функции виртуального мастер-устройства, при этом он обеспечивает коммуникацию хоста с внешней сетью. При использовании протокола VRRP, на хосте необходимо установить только шлюз по умолчанию, без запуска прочих протоколов. В этом случае нагрузка на хост минимальна, к нагрузке добавляется только небольшая часть потока данных протокола VRRP.

33.1.2 VRRP term (VRRP пояснения и терминология)

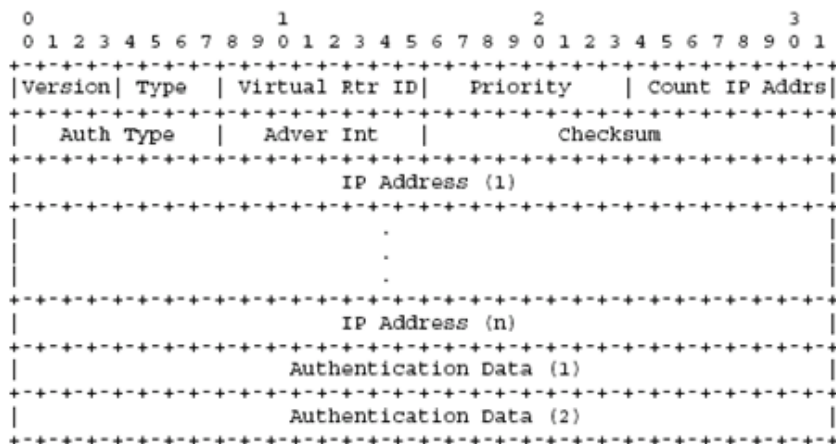
Ниже приведена расшифровка встречающихся в данной главе понятий и определений:

- VRRP (Virtual Router Redundancy Protocol) - протокол резервирования виртуального маршрутизатора, применяется для повышения доступности маршрутизатора, используемого в роли шлюза по умолчанию, повышает надежность и устойчивость работы сети.
- VR (Virtual router) – виртуальный объект основанный на интерфейсе подсети, включает в себя VRID (virtual router identifier) идентификационный номер виртуального роутера и IPадрес. Данный IP адрес называется virtual IP address, используется как шлюз хоста по умолчанию.
- VRRP Router – VRRP роутер с запущенным протоколом VRRP, может входить в виртуальный роутер.
- IP Address Owner – VRRP роутер, чей виртуальный IP адрес тот же что и реальный IP адрес интерфейса.
- Virtual Router Master – виртуальный мастер-роутер, ответственный за пересылку проходящих пакетов данных уровня layer 3 и отвечающий на запросы ARP IP адреса виртуального роутера. Виртуальный мастер-роутер это VRRP роутер, который является IP address owner.

- Virtual Router Backup – виртуальный резервный роутер, не пересылает пакеты данных уровня layer 3 и не отвечает на запросы ARP IP адреса виртуального роутера. В случае сбоя работы виртуального мастер-роутера его функции переходят к виртуальному резервному роутеру.
- Коммутатор может иметь несколько интерфейсов, VRRP протокол может быть запущен на нескольких интерфейсах подсети.
- Виртуальный роутер может быть создан на интерфейсе подсети.
- Для идентификации виртуального роутера используется VRID (virtual router identifier).

33.1.3 VRRP Protocol interaction (Взаимодействие протокола VRRP)

Пакет VRRP протокола инкапсулирован в пакете IP, заголовок VRRP сообщения показан на рисунке ниже.



1. Поле заголовка MAC VRRP пакета.

Источник MAC адреса: виртуальный MAC адрес виртуального роутера 00-00-5e-00-01 - {VRID}, где VRID – идентификационный номер виртуального роутера. Например, если VRID виртуального роутера 1, то виртуальный MAC адрес 00-00-5e-00-01-01.

MAC адрес назначения: VRRP multicast MAC адрес 01-00-5e-00-00-12.

2. Поле заголовка VRRP пакета.

IP адрес источника: первичный IP адрес интерфейса который отправляет VRRP пакет.

IP адрес назначения: multicast IP адрес 224.0.0.18. Пересылка пакетов Layer 3 не разрешена.

TTL: 255 для защиты от дистанционных атак VRRP пакетами.

Protocol: 112

3. Поле заголовка VRRP.

Version: 2

Type: тип VRRP пакета. Поддерживается только тип 1: 1 --- advertisement, VRRP пакет уведомления.

VRID: идентификатор виртуального роутера.

Priority: приоритет отправки VRRP виртуального роутера.

Count IP addr: номер виртуального IP адреса. Виртуальный роутер может иметь несколько виртуальных IP адресов.

Auth type: метод аутентификации между VRRP роутерами входящими в виртуальный роутер.

Advertisement interval: временной интервал между оповещениями. По умолчанию значение параметра 1 сек.

Checksum: контрольная сумма, вычисляется на основании версии заголовка VRRP.

IP address (ES): один или более виртуальных IP адресов.

Authentication data: данные аутентификации.

4. Приоритет VRRP.

Каждый VRRP роутер принадлежащий виртуальному роутеру должен иметь свой приоритет. Значение параметра приоритета изменяется в пределах 0 - 255, значения 0 и 255 имеют особое назначение. Значения приоритета выбираются из диапазона 1 - 254, по умолчанию установлено значение 100. Чем выше значение параметра, тем выше приоритет, устройство с наивысшим приоритетом становится мастер-роутером.

VRRP роутер IP address owner, входящий в виртуальный роутер, имеет значение параметра приоритета 255.

Когда виртуальный мастер-роутер уведомляет другие резервные роутеры о том, что он больше не является мастер-устройством, путем отправки VRRP пакета с приоритетом 0, запускается процесс

быстрого назначения одного из резервных роутеров новым мастер-роутером.

5. VRRP аутентификация.

VRRP протокол поддерживает три метода аутентификации, из которых выбирается метод наиболее подходящий требованиям безопасности сети.

0 --- No Authentication (отсутствие аутентификации)

No certification (отсутствие сертификации)

1 --- Simple Text Password (простой текстовый пароль)

Simple password authentication (простая аутентификация)

2 --- IP Authentication Header (IP аутентификация)

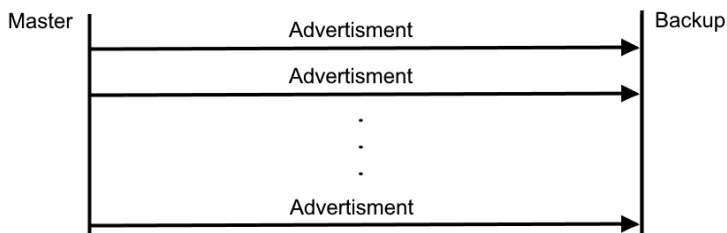
При IP аутентификации заголовок вычисляется методом message digest через HMAC-MD5.

Отсутствие аутентификации или простая текстовая аутентификация могут быть использованы в сети с низкой безопасностью, метод HMAC аутентификации может быть использован в сети с высокой безопасностью.

Для методов аутентификации 0 и 2, поля аутентификационных данных содержат 0, для метода аутентификации 1 поля аутентификационных данных содержат пароль password. Для метода аутентификации 2 краткое сообщение содержит IP идентификационное поле, которое добавлено в IP заголовок.

6. Взаимодействие пакетов VRRP.

VRRP протокол поддерживает только один тип пакетов - advertisement. В виртуальном роутере, виртуальный мастер-роутер отправляет пакеты-объявления announcement packet с интервалом advertisement interval (по умолчанию 1 сек). Виртуальный резервный роутер определяет статус миграции согласно полученным VRRP пакетам уведомления. Взаимодействие по протоколу между мастер-роутером и резервным роутером показано на рисунке ниже.



33.1.4 Election of virtual master router (Выбор виртуального мастер-роутера)

В виртуальном роутере при выборе виртуального мастер-роутера определяющими являются следующие факторы:

- IP address owner.

Если VRRP роутер является IP address owner (IP адрес интерфейса тот же что и виртуальный IP адрес, при нормальной работе роутера), то он является виртуальным мастер-роутером.

- VRRP priority.

При нормальной работе VRRP роутер с наивысшим приоритетом становится виртуальным мастер-роутером. Значение параметра приоритета изменяется в пределах 1 - 254. Значение приоритета IP address owner роутера 255. Когда виртуальный мастер-роутер уведомляет виртуальный резервный роутер, что он больше не является мастер-роутером, VRRP пакту присваивается приоритет 0.

- Актуальный IP адрес интерфейса.

При одинаковом значении параметра приоритета VRRP роутер с большим актуальным IP адресом интерфейса становится виртуальным мастер-роутером.

Переключение между режимами Active / standby виртуального роутера происходит при условиях описанных ниже.

1) При сбое виртуального мастер-роутера переключение произойдет в двух случаях:

- Если виртуальный мастер-роутер в статусе active отправит VRRP пакет с приоритетом 0. После получения этого пакета резервный роутер отправит skew_. Если VRRP пакет виртуального мастер-роутера не получен в течение определенного времени, произойдет переключение на новый виртуальный мастер-роутер. Переключение происходит в течение 1 сек.

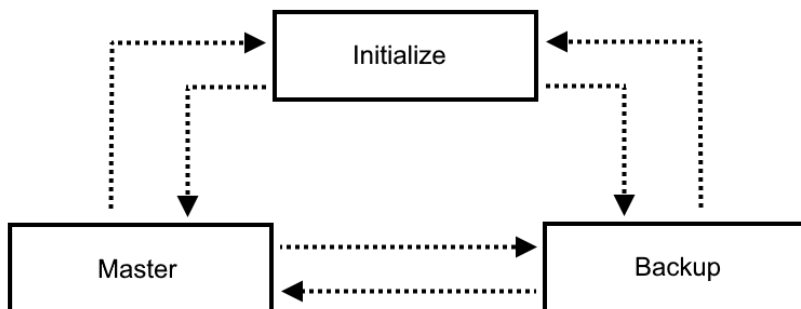
- Если виртуальный мастер-роутер не находится в статусе active, виртуальный резервный роутер перейдет в статус виртуального мастер-роутера если не получит VRRP пакет в течение периода времени master down interval. $\text{Master_Down_Interval} = (3 * \text{Advertisement_Interval}) + \text{Skew_Time}$. $\text{Skew_Time} = ((256 - \text{Priority}) / 256)$.

2) Если виртуальный мастер-роутер не является IP address owner, а в сети появляется IP address owner роутер, то новый роутер станет виртуальным мастер-роутером и произойдет active / standby переключение.

3) Если в сети появляется VRRP роутер и его приоритет выше чем приоритет виртуального мастер-роутера и он находится в режиме preemption mode (configuration variable preempt_mode), то новый роутер станет виртуальным мастер-роутером и произойдет active / standby переключение.

33.1.5 Status of virtual router (Статус виртуального роутера)

В виртуальном роутере каждый VRRP роутер проходит процедуру изменения статуса как state machine. Миграция (изменение статуса) state machine показана на рисунке ниже.



1) Initialize status (статус инициализации).

В статусе инициализации виртуальный роутер ожидает начальное событие. При получении сообщения о начальном событии запускаются следующие процессы:

- Если VRRP роутер имеет приоритет 255 (т.е. является IP address owner), то роутер становится виртуальным мастер-роутером и получает статус master.
- В противном случае роутер становится виртуальным резервным роутером и получает статус backup.

Действия роутера после получения статуса master:

- Отправка VRRP пакета уведомления (notification package).
- Отправка ARP запроса включая виртуальный IP адрес и соответствующий виртуальный MAC адрес.
- Установка интервала таймера advertisement_Interval

Действия роутера после получения статуса backup :

- Установка интервала таймера master_Down_Interval

2) Backup status (статус резервный).

Задача роутера в статусе backup отслеживать статус и работоспособность виртуального мастер-роутера, чтобы в случае сбоя начать выполнять его функции. При получении сообщения об отключении мастер-роутера происходит сброс таймера master_Down_Timer и возврат резервного роутера в статус инициализации initialize state.

При окончании времени таймера master_Down_ резервный роутер становится виртуальным мастер-роутером и получает статус master state.

При получении VRRP пакета уведомления notification package возможны следующие варианты:

- Если поле приоритета в VRRP пакете 0, устанавливается интервал master_Down_Timer и skew_Time.
- В противном случае если приоритет VRRP пакета превышает или равен приоритету VRRP роутера router, сбрасывается master_Down_Timer.
- В противном случае происходит сброс VRRP пакета.

3) Master status (статус мастер).

Задача VRRP роутера в статусе master state осуществлять пересылку пакетов данных layer 3 проходящих через виртуальный роутер.

В случае возникновения сбоя сбрасывается таймер adver_Timer, отправляется VRRP пакет уведомления notification package с приоритетом 0, роутер переходит в статус инициализации initialize state.

При окончании времени таймера adver_ отправляется VRRP пакет уведомления notification packet и сбрасывается таймер adver_Timer.

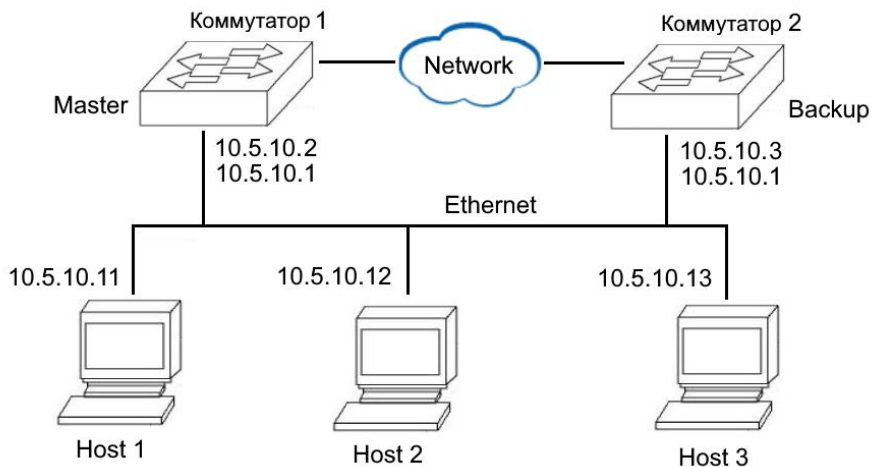
При получении VRRP пакета уведомления notification package возможны следующие варианты:

- Если значение приоритета в пакете 0, отправляется VRRP пакет уведомления notification package и сбрасывается таймер adver_Timer.
- В противном случае, если приоритет в пакете выше чем приоритет VRRP роутера или приоритеты имеют одинаковое значение, но IP адрес пакета выше чем первичный IP адрес интерфейса VRRP роутера, таймер adver_Timer будет сброшен, установлен таймер master_Down_Timer, роутер переходит в статус backup state.
- В противном случае происходит сброс VRRP пакета уведомления notification packet.

33.1.6 VRRP track (VRRP трекинг)

Протокол VRRP может определять внутренние сбои виртуального роутера (такие как обрыв соединения интерфейса в области нахождения виртуального роутера или сбой работы VRRP роутера), но не может определить внешние сбои виртуального роутера. Когда происходит сбой вне виртуального роутера, виртуальный роутер не может выбрать виртуальный главный роутер, что приводит к прерыванию передачи данных в сети. Функция трекинга VRRP tracking способна решить эту проблему. VRRP роутер отслеживает некоторые внешние события. Когда происходит внешний сбой, VRRP роутер изменяет порядок действий и выбирает новый виртуальный главный роутер для обеспечения непрерывной передачи данных в сети.

Как показано на рисунке ниже, когда внешний интерфейс виртуального мастер-роутера коммутатора 1 теряет соединение, при отключенной функции VRRP трекинга (коммутатор 1 не может определить внешний сбой и продолжает оставаться виртуальным мастер-роутером) и хост не может получить доступ к внешней сети. При включенной функции VRRP трекинга, коммутатор 1 может обнаружить внешний сбой и изменить последовательность действий и выбрать новый виртуальный мастер-роутер. Коммутатор 1 становится виртуальным резервным роутером, а коммутатор 2 становится виртуальным мастер-роутером, таким образом хост продолжает получать доступ к внешней сети.



Функция трекинга VRRP обеспечивает отслеживание интерфейса, маршрута и прохождения Ping. VRRP роутер отслеживает интерфейс находящийся за пределами виртуального роутера. При обрыве соединения индицируется наличие внешнего сбоя. VRRP роутер отслеживает маршрут содержащийся в маршрутной таблице. Если маршрут не существует или не активен, индицируется наличие внешнего сбоя. VRRP роутер отслеживает прохождение Ping к связанному с ним устройству. Если устройство не отвечает на Ping в течение определенного времени, то это указывает на то что произошел внешний сбой.

Виртуальный роутер может отслеживать интерфейс, маршрут и Ping одновременно. Каждый тип трекинга может отследить несколько событий. Как только происходит отслеживаемое событие, виртуальный роутер фиксирует внешний сбой. При отсутствии отслеживаемых событий индицируется нормальная работа виртуального роутера (отсутствие внешних сбоев).

33.2 VRRP configuration (Конфигурация VRRP)

Раздел конфигурации VRRP содержит следующие пункты:

- Create and delete virtual routers (создание и удаление виртуального роутера).
- Configure the virtual IP address of the virtual route (конфигурация IP адреса виртуального роутера).
- Configure parameters of virtual router (конфигурация параметров).
- Configure VRRP tracking (конфигурация функции VRRP трекинга).
- Start and shut down the virtual router (запуск и остановка виртуального роутера).
- View VRRP information (просмотр конфигурации).

33.2.1 Create and delete virtual routers (Создание и удаление виртуальных роутеров)

Виртуальный роутер создается на основе интерфейса подсети и должен иметь соответствующий идентификационный номер VRID. По умолчанию виртуальный роутер в системе не создан.

Когда необходимость в использовании виртуального роутера отпадает, то он может быть удален. Если виртуальный роутер запущен, то перед удалением он должен быть остановлен.

Таблица команд создания и удаления виртуального роутера:

Команда	Описание	Режим CLI
router vrrp <vrid>	Создание виртуального роутера и вход в режим VRRP configuration mode. Если BP уже существует, то сразу осуществляется вход в режим VRRP configuration mode. Значение параметра VRID изменяется в пределах 1 – 255.	Global configuration mode
no router vrrp [vrid]	Удаление виртуального роутера. VRID – значение параметра.	Global configuration mode

33.2.2 Configure the virtual IP address of virtual router (Конфигурация виртуального IP адреса виртуального роутера)

Виртуальный роутер должен иметь виртуальный IP адрес. Теоретически виртуальный роутер может иметь более одного виртуального IP адреса, но на практике это не применяется. По умолчанию коммутатор не имеет виртуального IP адреса.

Таблица команд конфигурации виртуального IP адреса:

Команда	Описание	Режим CLI
virtual-ip <virtual-ip> <backup master>	Устанавливает виртуальный IP адрес виртуального роутера.	VRRP Configuration mode
no virtual-ip	Удаляет виртуальный IP адрес виртуального роутера.	VRRP Configuration mode

Внимание:

- Конфигурация виртуального IP адреса производится когда виртуальный роутер находится в выключенном состоянии. При включенном виртуальном роутере сконфигурировать виртуальный IP адрес невозможно.
- Конфигурируемый виртуальный IP адрес должен принадлежать тому же сегменту сети что и первичный IP адрес интерфейса, в противном случае процедура не будет успешной.

33.2.3 Configure parameters of virtual router (Конфигурация параметров виртуального роутера)

Параметры виртуального роутера включают в себя priority (приоритет), preemption mode (режим упреждения), notification interval (интервал уведомления), authentication method (метод аутентификации) и authentication data (аутентификационные данные). Значения параметров по умолчанию приведены в таблице ниже.

Таблица значений параметров конфигурации виртуального роутера по умолчанию:

Параметр	Значение
priority	100
Preempt Mode	TRUE
Notification interval	1 sec
Authentication method	No certification
Authentication data	None

При конфигурации виртуального роутера интервал уведомления (notification interval), метод аутентификации (authentication method) и данные аутентификации (authentication data) должны быть сконфигурированы одинаково, в то время как приоритет и упреждающий режим (preemption mode) могут быть сконфигурированы одинаково.

Приоритет имеет два параметра конфигурации: configuration priority и operation priority. В большинстве случаев, operation priority использует значение configuration priority, но если VRRP роутер является IP address owner, то значение параметра operation priority установлено 255, а параметр configuration priority не используется.

Аутентификация имеет два метода реализации: отсутствие аутентификации (no authentication) и простую аутентификацию с помощью пароля (password authentication), метод IP аутентификации не используется.

Таблица команд конфигурации виртуального роутера:

Команда	Описание	Режим CLI
priority < priority-value >	Выбор значения параметра приоритета виртуального роутера. Изменяется в пределах 1 – 254.	VRRP onfiguration mode
preempt-mode {false true}	Установка режима preemption mode виртуального роутера. True означает preemption и false означает no preemption.	VRRP onfiguration mode
advertisement-interval <interval>	Установка интервала notification interval виртуального роутера. Изменяется в пределах 1 – 255 сек.	VRRP onfiguration mode

authentication none	Установка метода аутентификации виртуального роутера. Без аутентификации.	VRRP onfiguration mode
authentication simple-password <key>	Установка метода аутентификации виртуального роутера . Аутентификация по простому паролю, необходимо ввести данные аутентификации т.е.password.	VRRP onfiguration mode

Внимание:

- Конфигурация параметров виртуального роутера производится когда виртуальный роутер находится в выключенном состоянии. При включенном виртуальном роутере сконфигурировать параметры виртуального роутера невозможно.

33.2.4 Configure VRRP tracking (Конфигурация VRRP трекинга)

В настоящее время коммутатор поддерживает только функцию трекинга интерфейса VRRP (interface tracking function). VRRP роутер может осуществлять трекинг интерфейса уровня layer-2 или агрегированного интерфейса. По умолчанию коммутатор не имеет интерфейса для трекинга.

Если VRRP роутер является IP address owner, администратор может настроить VRRP трекинг, но на практике функция VRRP tracking не будет работать. В данном случае, если у виртуального роутера произойдет внешний сбой, то не произойдет назначение нового виртуального мастер-роутера. Для использования функции VRRP tracking, не следует конфигурировать виртуальный роутер как IP address owner.

Для корректной работы функции VRRP tracking администратор должен определить интерфейс и включить виртуальный роутер. В этом случае при разрыве соединения интерфейса, VRRP роутер будет считать, что произошел внешний сбой. Значение параметра operation priority виртуального роутера устанавливается как значение source priority минус значение priority value. В процессе взаимодействия пакета протокола VRRP, может быть назначен новый главный виртуальный роутер. После восстановления сбоя, когда отслеживаемые интерфейсы

имеют исправные соединения link up, то operation priority виртуального роутера будет установлен как configuration priority.

Таблица команд конфигурации VRRP tracking:

Команда	Описание	Режим CLI
circuit-failover <if-name> <priority-value>	Установка интерфейса для отслеживания виртуальным роутером.	VRRP onfiguration mode
no circuit-failover	Удаление интерфейса отслеживаемого виртуальным роутером.	VRRP onfiguration mode

Внимание:

- Конфигурация параметров функции VRRP tracking производится когда виртуальный роутер находится в выключенном состоянии. При включенном виртуальном роутере сконфигурировать параметры функции VRRP tracking невозможно.

33.2.5 Start and shut down the virtual router (Включение и отключение виртуального роутера)

После создания виртуального роутера, виртуального IP адреса и установки соответствующих параметров, виртуальный роутер находится в статусе инициализации и для дальнейшей работы его следует включить. Включение виртуального роутера запускает протокол, происходит отправка сообщения о событии startup event, коммутатор state machine изменяет свой статус на master state (главный) или or backup state (резервный). Выключение виртуального роутера останавливает работу протокола, отправляет сообщение о событии shutdown event, статус коммутатора изменяется обратно на initialize state (инициализация).

Перед включением виртуального роутера следует убедиться, что создан виртуальный IP адрес. Когда виртуальный роутер находится во включенном состоянии и при этом есть необходимость изменить виртуальный IP адрес или иные параметры, следует сначала

выключить виртуальный роутер и после этого вносить изменения в конфигурацию. После проведения (изменения) конфигурации следует включить виртуальный роутер.

Таблица команд включения / выключения виртуального роутера:

Команда	Описание	Режим CLI
enable	Включение (запуск) виртуального роутера.	VRRP onfiguration mode
disable	Выключение виртуального роутера.	VRRP onfiguration mode

33.2.6 View VRRP information (Просмотр конфигурации)

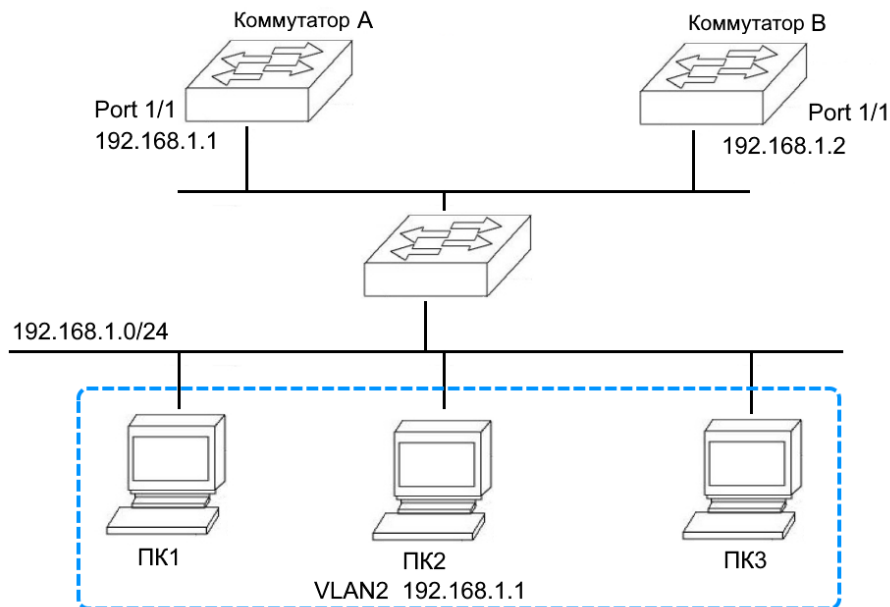
С помощью соответствующих команд возможно просматривать информацию о текущем статусе и конфигурации протокола VRRP.

Таблица команд просмотра текущей конфигурации параметров и протокола VRRP:

Команда	Описание	Режим CLI
show vrrp [vrid]	Просмотр информации о виртуальном роутере. При отсутствии параметра vrid будет выведена информация обо всех виртуальных роутерах.	Normal mode / privileged mode
show running-config	Просмотр текущей конфигурации системы, в т.ч. информации о конфигурации протокола VRRP.	privileged mode

33.3 VRRP Configuration example (Пример конфигурации)

На двух коммутаторах запущен протокол VRRP для обеспечения 3-layer маршрутизации, резервирования LAN для пользователей и предотвращения сбоев в сети. Коммутатор 1 имеет статус master, коммутатор 2 имеет статус backup (резервный). Схема подключения приведена на рисунке ниже:



Команды конфигурации коммутатора A:

```
Switch#configure terminal
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#exit
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport access vlan 2
Switch(config-ge1/1)#exit
Switch(config)#ip interface vlan 2
Switch(config)#interface vlan2
```

```
Switch(config-vlan2)#ip address 192.168.1.1/24
Switch(config-vlan2)#exit
Switch(config)#router vrrp 1
Switch(config-vrrp)# virtual-ip 192.168.1.1 master
Switch(config-vrrp)#enable
```

Команды конфигурации коммутатора B:

```
Switch#configure terminal
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#exit
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport access vlan 2
Switch(config-ge1/1)#exit
Switch(config)#ip interface vlan 2
Switch(config)#interface vlan2
Switch(config-vlan2)#ip address 192.168.1.2/24
Switch(config-vlan2)#exit
Switch(config)#router vrrp 1
Switch(config-vrrp)# virtual-ip 192.168.1.1 backup
Switch(config-vrrp)#enable vrrp
```

Команды просмотра конфигурации протокола VRRP:

```
show running-config
show vrrp
show vrrp 1
```

34. VLLP Configuration (Конфигурация VLLP)

Глава описывает конфигурацию протокола VLLP (VRRP layer-2 loop protection protocol). VLLP – протокол защиты layer-2 от образования петлевых соединений. Глава содержит следующие разделы:

- Introduction (Введение)
- VLLP Configuration (Конфигурация протокола)
- Configuration example (Пример конфигурации)

34.1 VLLP introduction (Введение)

VLLP (VRRP layer-2 loop protection protocol) – частный протокол, разработанный для защиты от заикливания (образования петель) layer-2 при использовании VRRP протокола. Когда коммутаторы уровня layer-3 используют VLLP для блокировки заиклиненных сообщений layer-2 and layer-3, производится взаимное уведомление двух коммутаторов об этом. Когда истекает время таймера блокировки порта block state times, то порт переходит в статус forward state.

VLLP протокол осуществляет мониторинг и обслуживание порта в режиме loop state путем немедленного уведомления об изменении статуса порта и истечении времени таймера, тем самым обеспечивая разрыв петлевого соединения вовремя. Протокол VLLP использует метод запроа-ответа для сбора информации о статусе порта. Когда на двух коммутаторах уровня layer-3 запущен VLLP протокол, один из них автоматически становится отправителем, а другой получателем. Отправитель отвечает за регулярную отправку сообщений запросов, а получатель отвечает за их получение и отправку сообщений ответов в обратном направлении. Когда получатель изменяет статус порта, необходимо отправить сообщение link state change message для уведомления отправителя о соответствующем изменении. Протокол VLLP работает совместно с протоколом VRRP для отслеживания статуса релевантных портов VLAN. Использование протокола VLLP необходимо для коммутаторов VRRP, в то время как коммутаторы уровня layer-2 не нуждаются в использовании этого протокола (для защиты от петлевых соединений).

Основные принципы VLLP протокола:

- Оборудование VLLP: VLAN с работающим протоколом VLLP называется VLLP оборудованием.
- VLLP порт: порт участвующий во взаимодействии протокола VLLP в соответствующем VLAN и VLLP оборудовании.
- Статус VLLP оборудования: оборудование VLLP имеет два статуса (sender отправитель и and receiver получатель).
- Отправитель Sender: оборудование VLLP в статусе sender осуществляет периодическую отправку сообщений запросов query message.
- Получатель Receiver: оборудование VLLP в статусе receiver формирует ответные сообщения в случаях изменения статуса соединения.
- Статус порта VLLP: VLLP порт имеет три статуса STP (disable отключен, block заблокирован и forward пересылает).
- Главный порт (Main port): VLLP порт назначается главным когда VLLP оборудование находится в статусе отправителя (sender state); VLLP отправитель имеет только один главный порт. VLLP порты с высоким приоритетом и участвующие в маршрутизации (mapping relationship) имеют преимущество при выборе главного порта.
- Preempt VLLP port mapping relationship: VLLP порты соединенные с другими коммутаторами и взаимодействующие с использованием сообщений протокола VLLP.

Принципы выбора получателя из VLLP оборудования:

- По критерию высокого приоритета;
- По критерию наибольшего MAC адреса (при одинаковом приоритете).

Принципы выбора главного порта (Main port):

1. Порт должен иметь соединение (link up);
2. VLLP порт должен иметь взаимосвязь (mapping relationship).

Если ни один VLLP порт не удовлетворяет условиям, то главный порт отсутствует;

Если существует несколько портов удовлетворяющих условиям, то главным выбирается один порт с высшим приоритетом; При одинаковом приоритете выбирается один первичный порт для создания mapping relationship.

Принципы определения статуса VLLP порта:

- Если VLLP оборудование является отправителем, то главный порт имеет статус forward;
- Если VLLP оборудование является отправителем и имеет статус соединения link up, то статус VLLP порта без взаимосвязи mapping relationship – forward;
- Если VLLP оборудование является отправителем и имеет статус соединения link up, то статус VLLP порта с взаимосвязью mapping relationship – block;
- Если VLLP оборудование является отправителем и имеет статус соединения link up, то статус VLLP порта - forward;
- VLLP порт при отсутствии соединения link down имеет статус disabled;

Сообщения протокола VLLP инкапсулируются в кадре MAC frame.

Destination MAC Address (6 bytes)						
Source MAC Address (6 bytes)						
0x8100		Prio	Vlan ID (12 bits)	Ethernet Type (2 bytes)		
Version	Type					
Priority	Query Inter				Port1(2 bytes)	
					Main port	
Reserved (4 bytes)						
Port2		Link state	STP state			

Существует три типа сообщений протокола VLLP:

- Link status query message LQ
- Link state response message La
- Link state change message LC

Диапазон значений каждого поля в формате сообщения:

- Des MAC: fixed at 00:09: CA: FF: FF: FF
- SRC MAC: MAC of VLAN sending VLLP protocol message
- Ethernet type: fixed to 0x268
- Version: currently 1
- Type: LQ is 1; La is 2; LC is 3;
- Port1: index value of the port sending VLLP protocol message
- Priority: VLLP equipment priority, value 1 ~ 255

Query interval: interval of VLLP sender query timer, 5 seconds by default
Main port: index value of VLLP sender's main port, only in LQ message;
Reserved: the reserved domain is zero
Port2: the index value of the port whose state changes in LC message.
The LQ is consistent with port2 and port1 in La message
Link state: the link state of port2. Link up is 1 and link down is 2;
STP state: the STP state of port2. Disable is 1, block is 2, forward is 3

Принципы работы протокола VLLP:

Для начала работы следует сконфигурировать VLLP оборудование и запустить протокол VLLP на VLAN и то же самое проделать на оборудовании и VLAN с противоположной стороны линии. Таким образом протокол VLLP формирует объединение пары отправителя и получателя (sender – receiver). После запуска протокола обе стороны отправляют друг другу сообщения LQ messages. Когда VLLP оборудование получает сообщение LQ message, происходит выбор получателя на основании приоритета содержащегося в сообщении и MAC адреса противоположной стороны. По итогам одна из сторон становится получателем (receiver) и перестает отправлять сообщения LQ message, но продолжает отвечать на сообщения LQ message отправителя (sender). Если не получено сообщение LQ message (по истечении времени таймера получателя), получатель возвращается в статус отправителя и начинает отправку сообщений LQ message.

Порт, участвующий во взаимодействии сообщений VLLP протокола, принадлежащий VLAN, должен быть сконфигурирован как VLLP порт. VLLP порт должен поддерживать layer-2 (включая транк группу), но при этом не допускается конфигурировать членов trunk группы как VLLP порты. Сообщения протокола VLLP отправляются и принимаются через VLLP порт. Сконфигурированный VLLP порт VLAN осуществляющий обмен с со связанным оборудованием (с противоположного конца линии) с запущенным VLLP протоколом составляют пару mapping relationship. Они определяют mapping relationship с противоположным портом путем обмена сообщениями запроса query message и сообщениями ответа response message, вычисляют возможное образование петлевых соединений в сети согласно mapping relationship и статусу соединения link state. Состояние

STP state VLLP порта обслуживается согласно принципу определения VLLP порта для предотвращения возникновения петель в топологии.

VLAN с запущенным протоколом VLLP может содержать несколько VLLP портов. Они могут иметь или не иметь физическое соединение с противоположным оборудованием. Если порт принадлежит нескольким VLAN, один и тот же VLLP порт будет появляться в нескольких VLLP устройствах. Протокол VLLP динамически собирает информацию об изменениях статуса соединения link state VLLP порта и статуса STP state противоположного VLLP порта для своевременного вычисления петлевых соединений в сети и эффективного противодействия этому явлению.

Если существует несколько VLAN с идентичной конфигурацией портов, и на них требуется запустить VLLP протокол, и каждый layer-2 порт должен отправлять и принимать сообщения VLLP протокола циркулирующих в разных VLAN, то это может привести к перегрузке коммутатора. Для решения этой проблемы предлагается использовать концепцию вспомогательных VLAN. Таким образом, выделяется первичный VLAN с запущенным VLLP протоколом, а остальные VLAN с аналогичной конфигурацией добавляются к нему как дополнительные VLAN. В результате петлевые соединения вычисляются протоколом VLLP первичного VLAN, и записывается статус порта. При конфигурации дополнительных VLAN убедитесь что VLLP порты главного VLAN также являются портами дополнительных VLAN, и все порты дополнительных VLAN являются портами основного VLAN. Когда layer-2 порт добавляется к дополнительному VLAN, и он одновременно принадлежит главному VLAN, то статусом порта управляет главный VLAN. Если порт не принадлежит главному VLAN, то его статусом невозможно управлять, при этом формируется тревожное сообщение.

34.2 VLLP configuration (Конфигурация VLLP)

Раздел конфигурации VLLP содержит следующие пункты:

- Create VLLP device on layer 3 interface (создание VLLP оборудования на layer 3 интерфейсе).
- Enable VLLP device (включение VLLP оборудования).
- Create VLLP port on layer 2 interface (создание VLLP порта на layer 2 интерфейсе).

- Configure VLLP device priority (конфигурация приоритета VLLP оборудования).
- Configure VLLP device query timer interval (конфигурация интервала таймера запросов).
- Configure attached VLANs (конфигурация дополнительных VLAN).
- Configure VLLP port priority (конфигурация приоритета VLLP порта).

34.2.1 Create VLLP device on layer 3 interface (создание VLLP оборудования на layer 3 интерфейсе)

Команда создания VLLP оборудования в режиме global configuration mode:

```
router vllp < if name >
```

Команда удаления VLLP оборудования:

```
no router vllp < if name >
```

Параметр if name должен совпадать с именем интерфейса layer-3 (например: vlan1, vlan2...).

По умолчанию протокол VLLP отключен.

34.2.2 Enable VLLP device (Включение VLLP оборудования)

Команда включения VLLP оборудования в режиме vllp configuration mode:

```
vllp enable
```

Команда отключения VLLP оборудования:

```
vllp disable
```

По умолчанию после создания VLLP оборудование отключено.

34.2.3 Create VLLP port on layer 2 interface (Создание VLLP порта на layer 2 интерфейсе)

Команда создания VLLP порта в режиме vllp configuration mode:

```
vllp port < if name > Create vllp port
```

Команда удаления VLLP порта:

```
no vllp port < if name >
```

Параметр if name должен совпадать с именем интерфейса layer-2 (например: GE1 / 1, trunk1...).

По умолчанию протокол VLLP не применяется к порту layer-2 port. Протокол VLLP не может быть применен если интерфейс layer-2 является членом trunk member.

34.2.4 Configure VLLP device priority (Конфигурация приоритета VLLP оборудования)

Команда конфигурации приоритета в режиме vllp configuration mode:

```
vllp priority < priority >
```

Команда отмены конфигурации приоритета к значению по умолчанию:

```
no vllp priority [priority]
```

Параметр priority value изменяется в пределах 1 - 255. Параметр используется для определения (выбора) получателя.

По умолчанию установлено значение параметра 100.

34.2.5 Configure VLLP device query timer interval (Конфигурация интервала таймера запросов)

Команда конфигурации времени таймера в режиме vllp configuration mode:

```
vllp query interval < interval >
```

Команда отмены конфигурации таймера к значению по умолчанию:

```
no vllp query interval [interval]
```

Параметр interval value изменяется в пределах 1 - 255. Конфигурация принимает установленное значение если VLLP оборудование находится в статусе отправителя (sender) или возвращается в этот статус.

По умолчанию установлено значение 5 сек.

34.2.6 Configure attached VLANs (Конфигурация дополнительных VLAN)

Команда конфигурации дополнительных VLAN в режиме vllp configuration mode:

```
vllp dependency < if name >
```

Команда удаления дополнительных VLAN:

```
no vllp dependency < if name >
```

Параметр if name должен соответствовать имени интерфейса layer-3 (например: vlan1, vlan2...).

По умолчанию дополнительные интерфейсы не сконфигурированы.

34.2.7 Configure VLLP port priority (Конфигурация приоритета VLLP портов)

Команда конфигурации приоритета vllp портов в режиме configuration mode:

```
vllp port < if name > priority < priority >
```

Команда отмены приоритета к значению по умолчанию:

```
no vllp port < if name > priority [priority]
```

Параметр if name должен соответствовать имени интерфейса layer-2 (например: GE1 / 1, trunk1...), Значение параметра priority изменяется в пределах 1 - 255. Параметр Priority используется отправителем vllp для выбора первичного порта primary port.

По умолчанию установлено значение параметра 100.

34.2.8 View VLLP information (Просмотр конфигурации)

Команда просмотра списка VLLP оборудования в режиме normal mode или privileged mode:

```
show vllp
```

Команда детального просмотра VLLP оборудования:

```
show vllp < if name >
```

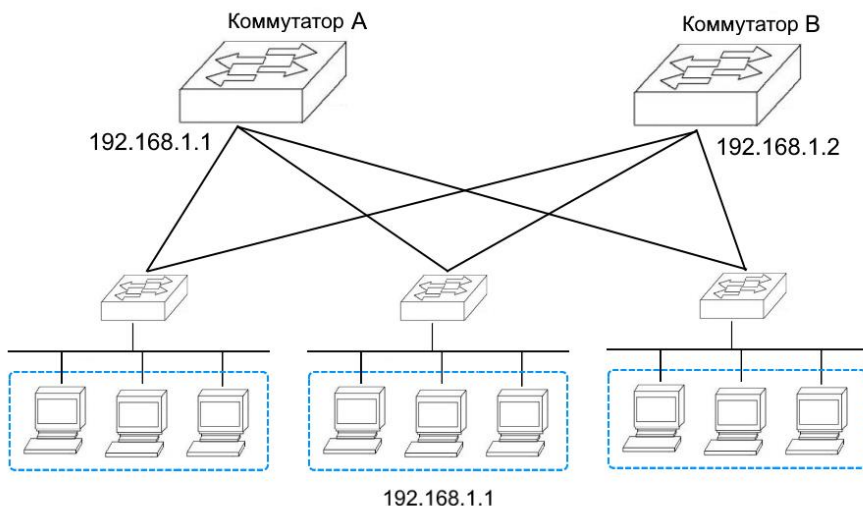
Параметр if name должен соответствовать имени интерфейса layer-3 (например: vlan1, vlan2...).

Команда просмотра mapping relationship каждого VLLP порта port относящегося к работе VLLP протокола:

```
show vllp map
```

34.3 VLLP Configuration example (Пример конфигурации)

Схема подключения приведена на рисунке ниже:



Команды конфигурации коммутатора A:

```
Switch#configure terminal
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config-vlan)#exit
Switch(config)#ip interface vlan 2
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport access vlan 2
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#switchport access vlan 2
Switch(config-ge1/2)#interface ge1/3
Switch(config-ge1/3)#switchport access vlan 2
```

```
Switch(config-ge1/3)#interface vlan2
Switch(config-vlan2)#ip address 192.168.1.1/24
Switch(config-vlan2)#exit
Switch(config)#router vrrp 1
Switch(config-vrrp)# virtual-ip 192.168.1.1 master
Switch(config-vrrp)#enable
Switch(config-vrrp)#exit
Switch(config)#router vllp vlan2
Switch(config-vllp)#vllp port ge1/1
Switch(config-vllp)#vllp port ge1/2
Switch(config-vllp)#vllp port ge1/3
Switch(config-vllp)#vllp enable
```

Команды конфигурации коммутатора B:

```
Switch#config t
Switch(config)#vlan database
Switch(config-vlan)#vlan 2
Switch(config)#ip interface vlan 2
Switch(config-vlan)#exit
Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport access vlan 2
Switch(config-ge1/1)#interface ge1/2
Switch(config-ge1/2)#switchport access vlan 2
Switch(config-ge1/2)#interface ge1/3
Switch(config-ge1/3)#switchport access vlan 2
Switch(config-ge1/3)#interface vlan2
Switch(config-vlan2)#ip address 192.168.1.2/24
Switch(config-vlan2)#exit
Switch(config)#router vrrp 1
Switch(config-vrrp)# virtual-ip 192.168.1.1 backup
Switch(config-vrrp)#virtual-interface vlan1
Switch(config-vrrp)#enable
Switch(config-vrrp)#exit
Switch(config)#router vllp vlan2
Switch(config-vllp)#vllp port ge1/1
Switch(config-vllp)#vllp port ge1/2
Switch(config-vllp)#vllp port ge1/3
Switch(config-vllp)# enable
```

Команды просмотра конфигурации протокола VLLP:

```
show vllp
show vllp <if-name>
show vllp map
```

35. Configure policy routing (Политика маршрутизации)

Глава описывает механизм и конфигурацию политики маршрутизации на основе пользовательских установок. Глава содержит следующие разделы:

- Introduction (Введение)
- Policy routing Configuration (Конфигурация политики)
- Configuration example (Пример конфигурации)

35.1 Policy routing introduction (Введение)

Политика маршрутизации использует механизм основанный на настройках, которые устанавливает прользователь. В отличие от простого направления сообщений на основе информации маршрутной таблицы (IP адрес назначения сообщения), политика маршрутизации использует ряд информационных атрибутов сообщения, такие как адрес назначения, адрес источника сообщения и прочую дополнительную информацию.

35.2 Policy routing configuration (Конфигурация политики маршрутизации)

Раздел конфигурации политики маршрутизации содержит следующие пункты:

- Create a new policy route (создание новой маршрутной политики).
- Insert a policy rout (введение маршрутной политики).
- Delete a policy rout (удаление маршрутной политики).
- Move a policy route (перемещение маршрутной политики).

- View policy routing information (просмотр информации маршрутной политики).

35.2.1 Create a new policy route (Создание новой политики маршрутизации)

Команда создания новой политики в режиме global configuration mode:
policy route <ID> <SIP> <DIP> <next-hop>

Параметр ID указывает номер создаваемой политики, изменяется в диапазоне 1 - 100.

Параметры Sip, dip определяют один из трех методов ввода IP источника и назначения:

- 1) A.b.c.d wildcard может управлять IP адресом из сегмента сети;
- 2) Any – эквивалентно a.b.c.d 255.255.255.25
- 3) Host a.b.c.d – эквивалентно a.b.c.d 0.0.

Wildcard определяет какие биты должны соответствовать или не соответствовать '0', '1'.

Параметр Next hop указывает адрес хоста соответствующий next hop, вводится в формате a.b.c.d

35.2.2 Insert a policy route (Введение политики маршрутизации)

Команда ввода новой политики в режиме global configuration mode:
policy route insert <ID> <SIP> <DIP> <next-hop> before <EXIST_ID>

Параметр ID указывает номер вводимой политики, изменяется в диапазоне 1 - 100.

Параметры Sip, dip определяют один из трех методов ввода IP источника и назначения:

- 1) A.b.c.d wildcard может управлять IP адресом из сегмента сети;
- 2) Any – эквивалентно a.b.c.d 255.255.255.25
- 3) Host a.b.c.d – эквивалентно a.b.c.d 0.0.

Wildcard определяет какие биты должны соответствовать или не соответствовать '0', '1'.

Параметр Next hop указывает адрес хоста соответствующий next hop, вводится в формате a.b.c.d

Параметр EXIST_ ID указывает какая политика была введена ранее, изменяется в диапазоне 1 - 100.

35.2.3 Delete a policy route (Удаление политики маршрутизации)

Команда удаления маршрутной политики в режиме global configuration mode:

```
no policy route <ID>
```

Параметр ID указывает номер удаляемой политики, изменяется в диапазоне 1 – 100.

35.2.4 Move a policy route (Перемещение политики маршрутизации)

Команда перемещения политики в режиме global configuration mode:

```
policy route move <ID> (before|after) <TO_ID>
```

Параметр ID определяет перемещаемую политику.

Параметр (before|after) определяет очередность перемещения (до или после).

35.2.5 View policy routing information (Просмотр информации)

Команда просмотра информации о политике маршрутизации в режиме normal mode или privileged mode:

```
show policy route
```

35.3 Policy routing configuration example (Пример конфигурации)

Команды конфигурации IP адреса источника как 192.168.3.100, с использованием gateway 192.168.0.20:

```
Switch#configure terminal
```

```
Switch#(config)#policy route 1 host 192.168.3.100 any 192.168.0.20
```

Команды конфигурации IP адреса получателя как 192.168.10.100, с использованием gateway of 192.168.2.1.

```
Switch#configure terminal
```

```
Switch#(config)#policy route 2 any host 192.168.10.100 192.168.2.1
```

Команды конфигурации IP адреса источника как 192.168.3.100, IP адреса получателя как 10.10.10.100. Шлюз gateway of 192.168.5.1.

```
Switch#configure terminal
```

```
Switch#(config)#policy route 3 host 192.168.3.100 host 10.10.10.100  
192.168.5.1
```

36. Configure system log (Логи системы)

Глава описывает конфигурацию логов системы и содержит следующие разделы:

- Introduction (Введение)
- System log Configuration (Конфигурация логов)

36.1 System log introduction (Введение)

Модуль логов системы является важной частью коммутатора. Он используется для записи информации об операциях всей системы, отклонениях от нормального режима работы, некорректных действиях оператора, помощи администратору при мониторинге системы и своевременного определения состояния системы. Модуль системных логов отвечает за управление всей информацией о системе: запущенных модулях, сбору, классификации, хранению, показу и выводу логов.

Модуль системных логов поддерживает важную функцию отладки системы (debugging function). Данная функция помогает администратору осуществлять мониторинг работы сети, производить отладку, диагностировать сбои в сети и т.п. Администратор может легко определить и связаться с нужным разделом, локализовать дефектное

сетевое оборудование путем просмотра логов отфильтрованных с помощью функции debugging.

Раздел конфигурации логов системы содержит следующие пункты:

- Format of log information (формат логов).
- Storage of logs (хранение логов).
- Display of logs (просмотр логов).
- Debugging tool (инструментарий отладки).

36.1.1 Format of log information (Формат логов)

Формат логов выглядит следующим образом:

Timestamp priority: module name: log content

время события приоритет: модуль: содержание

timestamp пробел priority, priority / пробел module name, module name / пробел log content.

Например: 2006/05/20 13:56:34 Warning: MSTP: Port up notification received for port ge1/2

В этом log сообщении timestamp 2006 / 05 / 20 13:56:34; приоритет warning; модуль MSTP; содержание лога port up notification received for port GE1 / 2.

1) Time stamp

Формат timestamp: year / month / day hour: minute: second.

Система поддерживает 24-часовой формат (от 0 до 23ч).

Timestamp фиксирует время создания лога, используется системное время коммутатора. Системное время устанавливается на предприятии изготовителе и может быть изменено администратором. Системное время не сбивается при отключении устройства.

2) Priority

Приоритет определяет важность события зафиксированного логом. Согласно важности события разделяются на четыре уровня. Порядок очередности приоритета выстроен от высокого к низкому: critical, warning, informational и debugging. Описание уровней приоритета приведено в таблице ниже:

Таблица уровней приоритета логов:

Приоритет	Описание
Critical	Критический (серьезная ошибка)
Warning	Предупреждение (обычная ошибка, уведомление, очень важная рекомендация).
Informational	Важная рекомендация (обычная рекомендация, диагностическая информация)
Debugging	Информация для отладки

3) Module name

Модуль указывает на название модуля к которому относится событие зафиксированное в логе. В таблице ниже приведен список основных модулей к которым могут относиться события зафиксированные в логах:

Таблица модулей указанных в логах:

Модуль	Описание
CLI	Command line interface module Интерфейс командной строки
MSTP	Multi instance spanning tree protocol module Протокол MSTP
VLAN	VLAN functional module Функционал VLAN
ARP	ARP functional module Функция ARP
IP	IP functional module Функционал IP адресов
ICMP	ICMP functional module Функция ICMP
UDP	UDP functional module Функционал UDP
TCP	TCP functional module Функционал TCP

4) Log content

Содержание лога – фраза или предложение отражающее основной смысл события зафиксированное в логе.

36.1.2 Storage of logs (Хранение логов)

Существует три варианта хранения логов:

- Хранение в памяти;
- Хранение в NVM (Node Version Manager);
- Хранение на сервере.

Для хранения логов в памяти существует четыре отдельные таблицы согласно приоритету важности логов. В каждой таблице хранятся логи одного приоритета (разделены по четырем категориям). Каждая таблица имеет 1K entries для хранения информации. Когда таблица полностью заполняется, последующие логи записываются вместо самых старых логов. Недостатком этого метода хранения является то, что при перезапуске системы сохраненная информация стирается и администратор не может определить проблему если имело место отключение системы.

Для сохранения наиболее важной информации (которая имеет критический уровень и уровень предупреждения), существует возможность хранить логи в NVM системы. В этом случае после перезапуска системы, логи хранившиеся в NVM могут быть восстановлены, таким образом администратор получает возможность диагностировать проблему по которой произошел сбой. Недостатком данного способа хранения является ограничение, вызванное небольшой емкостью NVM.

Оптимальным способом хранения информации является хранение логов на сервере, который реализуется с помощью протокола syslog protocol. Логи могут быть отправлены на сервер в масштабе реального времени, сервер сохраняет информацию и предоставляет доступ к ней через интерфейс. Основными достоинствами данного метода хранения является доступность информации пользователям и большая емкость хранилища размещенного на сервере.

На данный момент обеспечивается только возможность хранения логов в памяти устройства, хранение информации в NVM и на сервере недоступно.

36.1.3 Display of logs (Просмотр информации)

Существует два способа просмотра логов: ручной просмотр (manual display) и просмотр в реальном времени (real-time display). В ручном режиме для просмотра логов пользователь вводит соответствующие команды. В режиме просмотра в реальном времени поступающая информация выводится непосредственно на терминал, и пользователь может увидеть её вовремя.

В ручном режиме пользователь может просматривать весь список логов сразу. В этом режиме в начало списка помещаются последние записанные логи, таким образом пользователь сначала видит записи о событиях, которые произошли недавно.

В режиме реального времени для просмотра текущей информации, пользователь должен запустить дисплей терминала (real-time display switch of the terminal). Если коммутатор включен, то генерируемые логи не только записываются в таблицу, но и выводятся на терминал. Если коммутатор выключен, то логи не будут отображаться в реальном времени. На данный момент система может выводить логи в реальном времени только на консольный терминал, вывод информации на telnet терминал недоступен.

36.1.4 Debugging tool (Инструментарий отладки)

Для отладки оборудования и сетевых настроек система предлагает пользователю набор диагностического инструментария. Соответствующие инструменты могут отслеживать пакеты данных отправляемые и получаемые модулями системы, изменения статусов модулей системы и т.п., что дает возможность администратору следить за процессами, происходящими в системе. В случае возникновения нештатной ситуации работе в сети или оборудования, инструменты отладки отследят изменения и помогут решить проблему.

Инструменты отладки обеспечивают отслеживание состояния кромматоров и предоставляют необходимую информацию для своевременного реагирования на ситуацию администратору. При возникновении сбоя в работе сети или оборудования администратор задействует определенный процесс для поиска и локализации проблемы в конкретном модуле или системе.

Когда отлаживаемый коммутатор включен, система генерирует соответствующую информацию (логи), которая будет записана в соответствующую таблицу. Приоритет логов отладки имеет информационный характер. При включенном терминале информация о логах будет выведена в реальном времени. Если отлаживаемый коммутатор отключен, система не генерирует соответствующие логи.

36.2 System log configuration (Конфигурация логов системы)

Раздел конфигурации логов системы содержит следующие пункты:

- Configure terminal real-time display switch (конфигурация просмотра в режиме реального времени).
- View log information (просмотр логов).
- Configure debugging switch (конфигурация инструментов отладки).
- View debugging information (просмотр информации отладки).

36.2.1 Configure terminal real-time display switch (Конфигурация просмотра в режиме реального времени)

Режим просмотра информации в реальном времени по умолчанию отключен, генерируемые логи записываются в таблицу, но не выводятся на терминал. Некоторая информация (логи) не ограничивается системой коммутатора, эта информация всегда выводится на консольный терминал в реальном времени.

На данный момент система может выводить логи в реальном времени только на консольный терминал, вывод информации на telnet терминал недоступен.

Когда пользователь применяет команды для записи текущей конфигурации системы в файл конфигурации, конфигурация просмотра информации в реальном времени не сохранится в файле текущей конфигурации системы. После перезапуска системы эти настройки не сохранятся и их надо будет восстанавливать.

Таблица команд конфигурации режима просмотра информации в реальном времени:

Команда	Описание	Режим CLI
log stdout	Включение терминала для просмотра логов в реальном времени.	Global configuration mode
no log stdout	Отключение терминала для просмотра логов в реальном времени.	Global configuration mode

36.2.2 Set log level (Установка уровня логов)

Таблица команд установки уровня важности логов:

Команда	Описание	Режим CLI
log trap <[alerts critical debugging emergencies errors informational notifications warnings]>	Установка уровня логов	Global configuration mode

36.2.3 View log information (Просмотр логов)

Таблица команд для просмотра логов:

Команда	Описание	Режим CLI
show log	Просмотр всех логов	Global configuration mode

36.2.4 Configure debugging switch (Конфигурация инструментов отладки)

Система предоставляет пользователю инструменты отладки, которые обеспечивают отслеживание состояния оборудования и различных модулей системы. Ниже представлены упрощенные

команды для каждого модуля из списка. Для просмотра полного формата команд следует обратиться к соответствующим разделам руководства пользователя.

Когда пользователь применяет команды для записи текущей конфигурации системы в файл конфигурации, конфигурация просмотра информации в реальном времени не сохранится в файле текущей конфигурации системы. После перезапуска системы эти настройки не сохранятся и их надо будет восстанавливать.

Таблица команд конфигурации инструментов отладки:

Команда	Описание	Режим CLI
debug ip ...	Включение отладки отправления и получения IP пакетов.	Privileged mode
no debug ip ...	Отключение отладки отправления и получения IP пакетов.	Privileged mode
debug ip icmp ...	Включение отладки отправления и получения ICMP пакетов.	Privileged mode
no debug ip icmp ...	Отключение отладки отправления и получения ICMP пакетов.	Privileged mode
debug ip arp ...	Включение отладки отправления и получения ARP пакетов.	Privileged mode
no debug ip arp ...	Отключение отладки отправления и получения ARP пакетов.	Privileged mode
debug ip udp ...	Включение отладки отправления и получения UDP пакетов.	Privileged mode
no debug ip udp ...	Отключение отладки отправления и получения UDP пакетов.	Privileged mode
debug ip tcp ...	Включение отладки отправления и получения TCP пакетов.	Privileged mode

no debug ip tcp ...	Отключение отладки отправления и получения TCP пакетов.	Privileged mode
debug mstp ...	Включение отладки и диагностики протокола MSTP.	Privileged mode
no debug mstp ...	Отключение отладки и диагностики протокола MSTP.	Privileged mode
debug igmp snooping ...	Включение отладки и диагностики функции IGMP snooping.	Privileged mode
no debug igmp snooping ...	Отключение отладки и диагностики функции IGMP snooping.	Privileged mode
debug dhcp snooping ...	Включение отладки и диагностики протокола DHCP snooping.	Privileged mode
no debug dhcp snooping ...	Отключение отладки и диагностики протокола DHCP snooping.	Privileged mode
no debug all	Отключение отладки и диагностики на всей системе.	Privileged mode

36.2.5 View debugging information (Просмотр информации отладки)

Таблица команд просмотра информации отладки:

Команда	Описание	Режим CLI
show debugging [dhcp snooping erps igmp snooping ip mstp rip]	Просмотр конфигурации инструментов отладки. Если параметры не введены, то показана конфигурация всех модулей. При вводе одного параметра выводится информация о соответствующем модуле. При вводе IP, выводится информация о IP, ICMP, ARP, UDP и TCP модулях.	Normal mode / privileged mode

36.3 SYSLOG configuration (Конфигурация функции SYSLOG)

Протокол SYSLOG является стандартным протоколом для управления логами оборудования, протокол имеет простую архитектуру и широкое распространение.

Раздел конфигурации протокола SYSLOG содержит следующие пункты:

- SYSLOG introduction (введение).
- SYSLOG configuration (конфигурация протокола).
- SYSLOG configuration example (пример конфигурации).

36.3.1 SYSLOG introduction (Введение)

SYSLOG является стандартным протоколом для управления логами оборудования, протокол разделяется на три части.

Первая часть отвечает за определение подмодулей для разделения информации генерируемых логов (согласно соответствующим модулям), определение уровней важности логов для мониторинга состояния оборудования. Сбор различной информации об оборудовании осуществляется на основании определенных условий.

Вторая часть представляет собой файл конфигурации, который определяет как распорядиться собранной информацией. Логи могут храниться на локальном устройстве, могут быть отправлены на определенный сетевой сервер или направлены определенному пользователю и т.п. Файл конфигурации определяет как сохранять информацию сгенерированную устройством.

Третья часть отвечает за отправку сообщений протокола SYSLOG в соответствии с форматом сообщений определяемым RFC. Система коммутатора и рабочее соглашение протокола SYSLOG вместе составляют систему log module (лог модуль). Первая часть протокола SYSLOG состоит из каждого функционального подмодуля (sub module) коммутатора и отправляет логи каждого уровня системе log module, обслуживает таблицы логов в соответствии с одним из четырех уровней. Вторая часть протокола SYSLOG распределяет информацию в соответствии с log module системы. В первую очередь информация выводится в реальном времени на serial port terminal или вручную через terminal display коммутатор; Во вторую очередь информация

сохраняется в памяти в виде четырех таблиц разных уровней; В третью очередь записывается информация высокого уровня важности на NVM для её сохранения в случае отключения питания; В четвертую очередь информация в виде сообщения SYSLOG отправляется на удаленный сервер для хранения и последующей сортировки. Подмодуль системы SYSLOG реализует только третью часть протокола SYSLOG - передачу логов системы на сервер.

36.3.2 SYSLOG configuration (Конфигурация SYSLOG)

Таблица команд конфигурации протокола SYSLOG:

Команда	Описание	Режим CLI
syslog open <server-ip>	Открыть адрес сервера SYSLOG; Параметр server IP – IP адрес сервера.	Global configuration mode
syslog close	Закрыть адрес сервера SYSLOG.	Global configuration mode
log syslog	Запустить протокол SYSLOG.	Global configuration mode

36.3.3 SYSLOG configuration example (Пример конфигурации SYSLOG)

Сконфигурировать IP адрес сервера SYSLOG как 192.168.0.201.

Команды конфигурации коммутатора:

```
Switch#configure terminal
Switch(config)#syslog open 192.168.0.201
Switch(config)# log syslog
```

Команды просмотра конфигурации протокола SYSLOG:

```
Switch#show syslog
Syslog is opened!
server ip address: 192.168.0.201
udp destination port: 514
severity level: debugging
local device name: switch
```

37. IPv6 Basic configuration (Базовая конфигурация IPv6)

Глава описывает основные функции протокола IPv6, включая IPv6 layer 2 Forwarding и IPv6 ND. Глава содержит следующие разделы:

- Introduction (Введение)
- Configure IPv6 basic functions (Конфигурация основных функций протокола IPv6)
- Configure IPv6 Neighbor Discovery Protocol (Конфигурация функции ND - поиск соседних устройств)
- Display and maintenance of IPv6 (Обслуживание и просмотр информации IPv6)

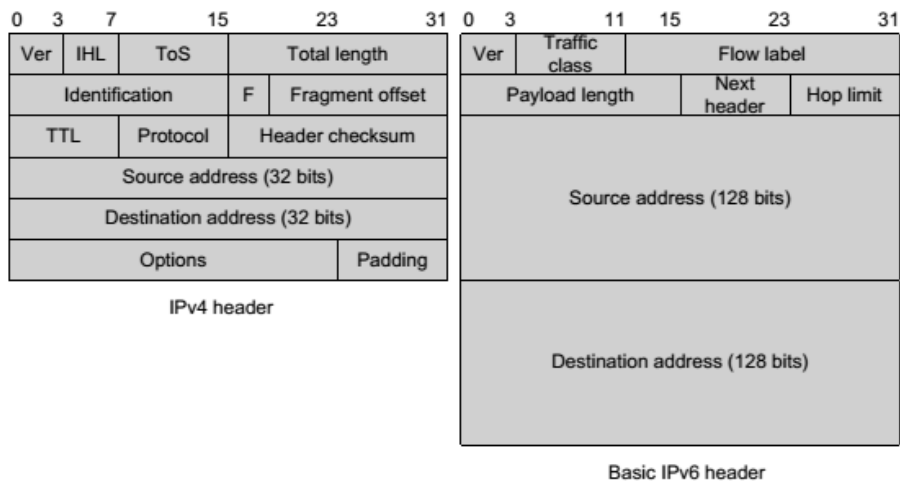
37.1 IPv6 introduction (Введение)

Протокол IPv6 (Internet Protocol version 6) является вторым поколением стандартного сетевого протокола IPng (IP next generation). Протокол соответствует спецификациям разработанным IETF (Internet Engineering Task Force) и улучшенной версии протокола IPv4. Наиболее важным отличием протокола IPv6 от IPv4 является увеличение длины IP адреса с 32 бит до 128 бит.

37.1.1 IPv6 Protocol features (Возможности протокола IPv6)

1) Simplified header format (упрощенный формат заголовка)

Длина заголовка основного сообщения протокола IPv6 уменьшается путем уменьшения или перемещения полей заголовка сообщения IPv4. Протокол IPv6 использует фиксированную длину заголовка основного сообщения, что упрощает процессинг сообщений IPv6 оборудованием и повышает эффективность перенаправления сообщений. Несмотря на то, что длина IP адреса протокола IPv6 в четыре раза превышает длину IP адреса протокола IPv4, длина заголовка основного сообщения IPv6 составляет только 40 бит, что в два раза меньше длины заголовка сообщения протокола IPv4 (исключая опционное поле).



2) Sufficient address space (адресное пространство)

Длина IP адреса источника и назначения в IPv6 составляет 128 бит (16 байт) и обеспечивает более чем 3.4×10^{38} возможных вариантов пространств адресов, что может потребовать разделение адресов на множество уровней для публичной и частной сети.

3) Hierarchical address structure (иерархия адресов)

Адресное пространство протокола IPv6 поддерживает иерархическую структуру, которая помогает быстро производить поиск маршрута. Одновременно, с помощью агрегации маршрутов, может быть эффективно уменьшено потребление сетевых ресурсов занимаемых маршрутной таблицей протокола IPv6.

4) Address auto configuration (автоматическая конфигурация адресов)

Для упрощения конфигурации хостов, протокол IPv6 поддерживает статусную и бесстатусную конфигурацию адресов:

- Статусная конфигурация предполагает получение IPv6 адресов и соответствующую информацию от сервера (например DHCP сервера).
- Бесстатусная конфигурация предполагает что хост автоматически конфигурирует IPv6 адреса и дополнительную информацию в соответствии со своим собственным link layer адресом и префиксом публикуемым роутером.

Одновременно хост может формировать локальный адрес связи согласно собственному link layer адресу и префиксу по умолчанию (fe80:: / 10) для обеспечения коммуникации с другими хостами.

5) Built in security (встроенная безопасность)

Протокол IPv6 использует IPSec как стандартное расширение заголовка, которое может обеспечить безопасность соединения точка-точка. Данная функция является стандартом для решения проблем сетевой безопасности и повышения возможностей взаимодействия при применении протоколов IPv6 на разных устройствах.

6) Support QoS (поддержка QoS)

Поле flow label заголовка сообщения протокола IPv6 используется для идентификации трафика и позволяет оборудованию идентифицировать сообщение первого класса first-class message и обеспечивать специальный процессинг.

7) Enhanced neighbor discovery mechanism (механизм поиска соседних устройств)

Механизм neighbor discovery протокола IPv6 реализуется через сообщения группы ICMPv6 (Internet control message protocol for IPv6), которая управляет взаимодействием информации между соседними узлами (узлы на одном соединении). Он заменяет протокол ARP (address resolution protocol), icmpv4 роутер перенаправляет icmpv4 сообщения и обеспечивает ещё ряд других функций.

8) Flexible extended message header (гибкий заголовок сообщения)

Протокол IPv6 сбрасывает опциональное поле заголовка сообщения IPv4 и вводит различные расширенные заголовки сообщений, которые не только повышают эффективность процессинга, но и значительно увеличивают гибкость и расширяют возможности протокола IPv6. Опциональное поле заголовка сообщения протокола IPv4 имеет размер 40 байт, в то время как размер гибкого заголовка сообщения протокола IPv6 ограничено только размером IPv6 сообщения.

37.1.2 IPv6 Address introduction (IPv6 адреса)

1) Представление IPv6 адреса

IPv6 адреса представлены в виде последовательностей чисел в шестнадцатиричном виде (16 бит), разделенных двоеточием (:). Каждый IPv6 адрес разделен на восемь групп. 16 бит каждой группы представлены четырьмя шестнадцатиричными номерами, группы разделены двоеточиями, как показано ниже: 2001:0000:130f:0000:0000:09c0:876a:130b. Для упрощения "0" в IPv6 адресе может быть записан в двух видах:

- Первый "0" в каждой группе может быть пропущен и адрес будет иметь следующий вид: 2001:0:130f: 0:0:9c0:876a:130b.
- Если адрес содержит две и более группы с последовательностью 0, то 0 может быть заменен на "::" и адрес будет иметь следующий вид: 2001:0:130f:: 9c0:876a:130b

Внимание:

Символ '::' может быть использован в IPv6 адресе только один раз. В противном случае при конвертации символа ':' в 0 для восстановления 128 бит адреса, количество нулей представленных символом '::' не может быть определено.

IPv6 адрес состоит из двух частей: префикса и идентификатора интерфейса. Префикс соответствует части поля номера сети IPv4 адреса, ID интерфейса соответствует номеру хоста в IPv4 адресе.

Префикс выражается как IPv6 адрес / длина префикса. IPv6 адрес имеет одну из форм приведенных выше, а длина префикса – десятичное число, указывающее сколько цифр находится слева от IPv6 адреса.

2) Классификация IPv6 адреса

В большинстве случаев IPv6 адрес может иметь три типа: unicast address, multicast address и anycast address.

- Unicast address: используется для однозначной идентификации интерфейса, аналогично IPv4 unicast address. Сообщение с данными отправленное на unicast адрес будет передано интерфейсу, соответствующему этому адресу.
- Multicast address: используется для идентификации группы интерфейсов (обычно принадлежат разным нодам), аналогично IPv4 multicast address. Сообщение с данными отправленное на multicast address будет передано интерфейсам, соответствующим этому адресу.

- Anycast address: используется для идентификации группы интерфейсов (обычно принадлежат разным нодам). Сообщение с данными отправленное на anycast address будет передано ближайшему от источника интерфейсу нода (определяется протоколом маршрутизации), соответствующему этому адресу.

Адреса типа broadcast address отсутствуют в протоколе IPv6. Функция broadcast адреса реализуется через multicast адрес.

Тип адреса IPv6 определяется по формату префикса (несколькими первыми цифрами адреса). Соответствие между главным типом адреса (main address type) и форматом префикса (format prefix) представлено в таблице ниже.

地址类型		格式前缀（二进制）	IPv6 前缀标识
单播地址	未指定地址	00...0 (128 bits)	::/128
	环回地址	00...1 (128 bits)	::1/128
	链路本地地址	1111111010	FE80::/10
	站点本地地址	1111111011	FEC0::/10
	全球单播地址	其他形式	-
组播地址		11111111	FF00::/8
任播地址		从单播地址空间中进行分配，使用单播地址的格式	

3) Типы unicast адресов

Существует много типов IPv6 unicast addresses, включая global unicast address, link local address и site local address.

- Глобальный global unicast address является эквивалентом публичного IPv4 public network address и предоставляется провайдером сети. Этот тип адреса допускает агрегацию маршрутных префиксов, которые ограничивают количество entries в в таблице глобальных маршрутов.
- Link local address используется для обеспечения коммуникации между нодами по протоколу neighbor discovery и бесстатусной автоматической конфигурации. Сообщения с данными использующими link local address как адрес источника или адрес назначения не будут направлены другим соединениям.
- Site local address аналогичен частному адресу IPv4. Сообщения с данными использующими local address как адрес источника или адрес назначения не будут направлены другим сайтам (эквивалент частной сети).

- Loopback address представляет собой unicast address 0:0:0:0:0:1 (упрощен как:: 1) и не может быть ассоциирован с физическим интерфейсом. Его функция аналогична функции loopback address в протоколе IPv4, т.е. нод используется для отправки сообщений IPv6 самому себе.

- Unspecified address представляет собой адрес ':: ' который не может быть ассоциирован с каким либо нодом. До того как нод получит действующий IPv6 адрес, должно быть заполнено поле источника отправленного IPv6 сообщения, но он не может быть использован как адрес получателя сообщения IPv6.

4) Multicast адрес

Адреса multicast приведенные ниже зарезервированы для специальных целей.

地址	应用
FF01::1	节点本地范围所有节点组播地址
FF02::1	链路本地范围所有节点组播地址
FF01::2	节点本地范围所有路由组播地址
FF02::2	链路本地范围所有路由组播地址
FF05::2	站点本地范围所有路由组播地址

Существуют и другие виды multicast адресов, такие как адреса запрашиваемые нодом (solicited node). В основном адрес используется для получения link layer адреса соседнего нода для реализации определения дублирования адресов. Каждый unicast или anycast IPv6 адрес соответствует запрошенному адресу нода. Формат адреса имеет вид FF02:0:0:0:0:1:FFXX:XXX, где Ff02:0:0:0:0:1: FF - фиксированный формат 104 бит; 20: XXXX – последние 24 бита unicast или anycast IPv6 адреса.

5) IEEE EUI-64 формат идентификатора интерфейса

Идентификатор интерфейса IPv6 unicast адреса используется для идентификации соединенного интерфейса. Для IPv6 unicast адреса в основном требуется идентификатор интерфейса 64 бит. Формат идентификатора интерфейса IEEE eui-64 рассчитывается из MAC адреса (link layer address) интерфейса. Размер идентификатора интерфейса IPv6 адреса составляет 64 бита, в то время как MAC адрес

имеет размер 48 бит. Таким образом необходимо ввести шестнадцатиричное число fffe (1110) в середину MAC адреса (после наибольшего 24го бита). Для обеспечения получения уникального идентификатора интерфейса из MAC адреса, устанавливается universal / local (U / L) bit (7ой бит от наибольшего) - "1". В итоге получается идентификатор интерфейса соответствующий формату eui-64.

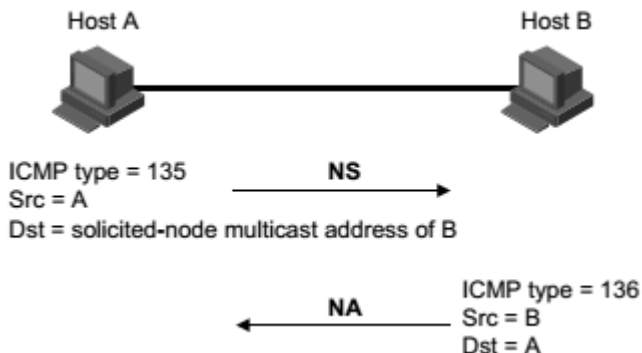


37.1.3 IPv6 neighbor discovery introduction (Введение IPv6 поиск соседних устройств)

Протокол IPv6 neighbor discovery использует пять типов сообщений ICMPv6 для выполнения следующих задач: address resolution (разрешение адреса), верификация доступности соседних устройств, определение дубликации адресов (duplicate address detection), поиск маршрута / префикса (router discovery / prefix discovery), автоматическое конфигурирование адреса и перенаправление (automatic address configuration and redirection). Типы сообщений ICMPv6 используемые протоколом neighbor discovery представлены ниже.

1) Address resolution

Получение адреса link layer соседнего нода по соединению (аналогично ARP протокола IPv4) через сообщение запроса соседнего устройства ns и сообщение уведомления Na. Как показано на рисунке ниже нод А должен получить link layer адрес нода В.



- Нод A отправляет multicast ns сообщения, адрес источника NS сообщения это IPv6 адрес интерфейса нода A, адрес назначения это multicast адрес запрашиваемого нода (нода B). Содержание сообщения включает в себя link layer адрес нода A.
- После получения NS сообщения нод B определяет адрес назначения сообщения как multicast адрес запрашивающего нода соответствующего IPv6 адресу. Далее нод B может получить link layer адрес нода A и вернуть Na сообщение в режиме unicast, включая свой собственный link layer адрес.
- Нод A может получить link layer адрес нода B через полученное Na сообщение.

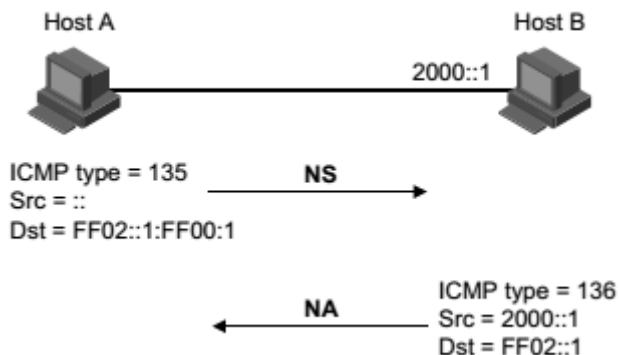
2) Verify neighbor reachability

После получения link layer адреса соседнего нода, возможно определить доступность соседнего нода через сообщение запроса ns и сообщение уведомления Na.

- Нод отправляет ns сообщение, где адрес назначения это IPv6 адрес соседнего нода.
- Если получено сообщение подтверждения от соседнего нода, считается что соседнее устройство доступно. В противном случае соседнее устройство считается недоступным.

3) Duplicate address detection

Когда нод получает IPv6 адрес, необходимо использовать механизм определения дубликации адресов для проверки адреса на использование другими нодами (аналогично ARP протокола IPv4). Механизм определения основан на обмене NS и Na сообщениями, как показано на рисунке ниже.



- Нод А отправляет ns сообщение. Адрес источника NS сообщения не выделен ::, адрес назначения – multicast адрес запрашиваемого нода, соответствует обнаруженному IPv6 адресу. Сообщение содержит определяемый IPv6 адрес.
- Если нод В уже использует IPv6 адрес, На сообщение будет возвращено, при этом сообщение содержит свой собственный IPv6 адрес.
- Когда нод А получает На сообщение от нода В, известно что IPv6 адрес уже используется. В противном случае это означает что адрес не используется, и нод А может использовать этот IPv6 адрес.

4) Router discovery / prefix discovery автоматическая конфигурация адреса

Router discovery / prefix discovery обнаружение маршрута / префикса относится к префиксу соседнего роутера нода и сети полученного RA сообщения и другим параметрам конфигурации.

Автоматическая конфигурация адреса без статуса относится к автоматической конфигурации IPv6 адреса нода согласно информации полученной при обнаружении router discovery / prefix discovery. Механизм Router discovery / prefix discovery реализован через сообщения роутера (RS запросы и RA уведомления).

- При запуске нода отправляется запрос роутеру путем формирования RS сообщения запроса префикса и прочей конфигурационной информации нода.
- Роутер возвращает RA сообщение, включая информацию о префиксе (RA сообщения отправляются периодически).

- Нод использует префикс адреса и другие конфигурационные параметры из возвращенного роутером RA сообщения для автоматической конфигурации IPv6 адреса и прочей информации интерфейса.

Опциональная информация о префиксе включает в себя не только префикс адреса, но и предпочтительное время существования и актуальность префикса. После получения периодического RA сообщения, нод обновляет предпочтительное время существования префикса согласно полученному сообщению.

Автоматически сгенерированный адрес может использоваться в пределах времени существования, по истечении которого автоматически сгенерированный адрес будет удален.

5) Redirection function

При запуске хоста, в таблице маршрутов по умолчанию может существовать только один путь к шлюзу. При таких условиях шлюз по умолчанию отправляет сообщение ICMPv6 хосту источнику для уведомления о перенаправлении, чтобы хост выбрал лучший next hop для отправки последующих сообщений (аналогично функции ICMP перенаправления сообщений протокола IPv4).

Оборудование отправит сообщение хосту ICMPv6 о перенаправлении при возникновении следующих условий:

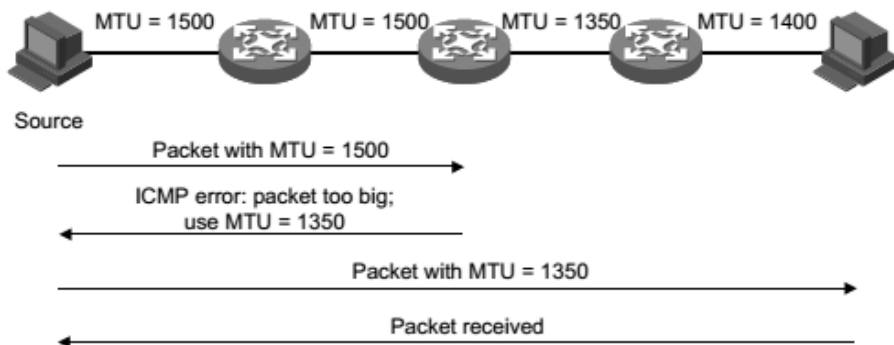
- Один и тот же интерфейс получает и отправляет сообщения с данными;
- Выбранный путь не был создан или изменен сообщением ICMPv6 redirection message;
- Выбранный путь не является маршрутом по умолчанию;
- Перенаправленное IPv6 сообщение с данными не содержит увеличенного маршрутного заголовка (routing extension header).

37.1.4 IPv6 PMTU discovery (Поиск IPv6 PMTU)

На протяжении маршрута передачи сообщений от источника к получателю соединение может иметь разные MTU (максимальная единица передачи). В протоколе IPv6, когда длина сообщения превышает MTU соединения, для облегчения процессинга на промежуточном оборудовании и рационального использования

ресурсов сети производится фрагментация сообщения со стороны источника.

Задача механизма PMTU (path MTU) - поиск наименьшего MTU на маршруте от источника к получателю. Процесс поиска PMTU показан на рисунке ниже.



- 1) Хост источник использует собственный механизм MTU для фрагментации сообщений, и далее отправляет сообщение хосту получателю.
- 2) Когда промежуточное оборудование получает сообщение для перенаправления и обнаруживает что значение MTU поддерживаемое интерфейсом меньше чем длина сообщения, то сообщение сбрасывается и источнику отправляется сообщение ICMPv6 об ошибке, включая MTU интерфейса, на котором произошел сбой.
- 3) После получения сообщения об ошибке, хост источник использует MTU переданный в сообщении для фрагментации и повторной отправки сообщения.
- 4) Для предотвращения многократного повторения операции по фрагментации сообщения необходимо произвести определение минимального MTU на пути от источника к получателю.

37.1.5 Protocol specification (Спецификации протокола)

RFC 1881: IPv6 Address Allocation Management

RFC 1887: An Architecture for IPv6 Unicast Address Allocation

RFC 1981: Path MTU Discovery for IP version 6
RFC 2375: IPv6 Multicast Address Assignments
RFC 2460: Internet Protocol, Version 6 (IPv6) Specification
RFC 2461: Neighbor Discovery for IP Version 6 (IPv6)
RFC 2462: IPv6 Stateless Address Autoconfiguration
RFC 2463: Internet Control Message Protocol (ICMPv6)for the Internet Protocol Version 6

IPv6 Specification
RFC 2464: Transmission of IPv6 Packets over Ethernet Networks
RFC 2526: Reserved IPv6 Subnet Anycast Addresses
RFC 3307: Allocation Guidelines for IPv6 Multicast Addresses
RFC 3513: Internet Protocol Version 6 (IPv6) Addressing Architecture
RFC 3596: DNS Extensions to Support IP Version 6

37.2 IPv6 basic configuration tasks (Конфигурация IPv6)

37.2.1 Configure IPv6 unicast address (Конфигурация unicast адреса)

Глобальный IPv6 global unicast адрес определяется ручными установками. Существует два способа установки IPv6 link local адреса:

- Автоматическая генерация: при включенном порте VLAN устройство автоматически генерирует link local адрес для интерфейса согласно префиксу link local адреса prefix (fe80:: / 10) и link layer адресу интерфейса;
- Link: ручная конфигурация IPv6 пользователем.

Команда	Описание	Режим CLI
ipv6 address <ipv6-address>/<prefix-length>	Ручная установка IPv6 адреса. По умолчанию локальный адрес генерируется автоматически согласно MAC адресу VLAN интерфейса layer 3.	Global configuration mode

37.3 Configure IPv6 Neighbor Discovery Protocol (Конфигурация NDP)

37.3.1 Configure relevant parameters of RA message (Конфигурация параметров RA сообщения)

Пользователь может конфигурировать параметры RA сообщений отправляемых интерфейсом, в т.ч. интервал отправления сообщений в зависимости от ситуации. Одновременно пользователь может конфигурировать релевантные параметры RA сообщений для уведомления хоста. Когда хост получает RA сообщение, он может использовать эти параметры для выполнения соответствующих операций. Параметры конфигурации сообщения и их значения представлены в таблице ниже.

Таблица параметров и значений RA сообщения:

Параметр	Описание
Cur Hop Limit	При отправлении сообщений IPv6, хост использует значения этих параметров для заполнения поля hop limit в заголовке IPv6 сообщения. Одновременно значение параметра используется как поле hop limit в ответном сообщении оборудования.
Prefix Information	После получения хостом информации префикса (опубликованной оборудованием), может быть выполнена stateless автоматическая конфигурация и другие операции.
Managed address configuration flag bit (M flag)	Используется для определения когда хост применяет stateful динамическую фильтрацию для получения IPv6 адреса. Если установлен flag bit 1, хост получит IPv6 адрес путем stateful автоконфигурации (как DHCP server). В противном случае IPv6 адрес будет получен путем stateless автоконфигурации, т.о. IPv6 адрес будет сгенерирован согласно собственной информации о префиксе link layer адреса опубликованной роутером.

Other configuration flag bits (O flag)	Используется для определения когда хост применяет автоконфигурацию для получения другой информации (кроме информации о IPv6 адресе). При других конфигурациях установлен flag bit 1, хост получает информацию кроме IPv6 адреса путем stateful автоконфигурации (как DHCP server). В противном случае другая информация будет получена путем stateless автоконфигурации.
Router lifetime (Router Lifetime)	Используется для установки времени, когда роутер, отправляющий RA сообщения является роутером хоста по умолчанию. Согласно значению параметра lifetime в принятом RA сообщении, хост может определить считать ли роутер публикующий RA сообщения роутером по умолчанию.
Neighbor request message retransmission interval (Retrans Timer)	Если устройство не получит ответ на отправленное NS сообщение в течение заданного времени, то устройство повторно отправит NS сообщение.
Time to keep neighbors reachable (Reachable Time)	Если обнаружено, что соседнее устройство недоступно, считается что это связано с установкой параметра reachability time. Когда установленное время истекло, необходимо отправить сообщение соседнему устройству для подтверждения доступности устройства.
Link maximum transmission unit (Link MTU)	Опция MTU используется в RA сообщении для подтверждения того, что все ноды соединения используют одинаковое значение параметра MTU. В основном используется когда нод может не иметь информации о соединении MTU. Другие сообщения neighbor discovery messages должны отсутствовать.

Команда конфигурации hop limit в режиме VLAN interface mode:

```
ipv6 nd cur-hop-limit value
```

По умолчанию число hops публикуемых роутером ограничено до 64 hops.

Команда конфигурации подавления публикации RA сообщений в режиме VLAN interface mode:

```
ipv6 nd send-ra
```

По умолчанию функция подавления публикации RA сообщений включена.

Команды конфигурации максимального и минимального времени интервалов отправки RA сообщений в режиме VLAN interface mode:

```
ipv6 nd max-ra-interval value
```

```
ipv6 nd min-ra-interval value
```

По умолчанию установлены максимальный интервал отправки RA сообщений 600 сек, минимальный интервал отправки RA сообщений 198 сек.

Внимание:

- При периодической отправке RA сообщений, временной интервал отправки сообщений выбирается в пределах между максимальным (maximum time interval) и минимальным (minimum time interval) значениями параметра публикации RA сообщений.
- Значение минимального интервала должно быть не более чем 0.75 значения параметра максимального интервала отправки сообщений.

Команда конфигурации префикса RA сообщений в режиме VLAN interface mode:

```
ipv6 nd prefix X::X::X::X/M (valid-lifetime preferred-lifetime (off-link | no-autoconfig))
```

По умолчанию информация префикса RA сообщения не сконфигурирована. IPv6 адрес интерфейса отправляющего RA сообщение будет использован как префикс RA сообщения.

Команда установки flag bit адреса в режиме VLAN interface mode:

```
ipv6 nd managed-config-flag
```

По умолчанию установлено значение flag bit адреса 0, т.о. хост получает IPv6 адрес через stateless автоконфигурацию.

Команда установки иного значения flag bit адреса в режиме VLAN interface mode:

```
ipv6 nd other-config-flag
```

По умолчанию иная конфигурация flag bits установлена на 0, т.о. хост получает информацию через stateless автоконфигурацию.

Команда конфигурации параметра lifetime роутера в RA сообщении в режиме VLAN interface mode:

```
ipv6 nd ra-lifetime value
```

По умолчанию параметр lifetime роутера в RA сообщении имеет значение 1800 сек.

Команда конфигурации интервала отправки повторного сообщения запроса (retransmission interval) соседнему устройству в режиме configuration mode:

```
ipv6 nd base retrans-timer value
```

По умолчанию установлено значение отправки ns сообщения интерфейсом 1000 мс.

Команда конфигурации интервала ретрансляции роутера в RA сообщении в режиме VLAN interface mode:

```
ipv6 nd retrans-timer value
```

По умолчанию значение параметра времени ретрансляции в RA сообщении установлено 0.

Команда конфигурации времени доступности соседних устройств (keep neighbors reachable) в режиме configuration mode.

```
ipv6 nd base reachable-time value
```

По умолчанию значение параметра keep neighbors reachable для интерфейса установлено 30000 мс.

Команда конфигурации поля времени доступности соседних устройств (keep neighbors reachable) RA сообщения режиме VLAN interface mode.

```
ipv6 nd reachable-time value
```

По умолчанию значение параметра поля keep neighbors reachable RA сообщения отправленного интерфейсом 0.

Команда конфигурации значения параметра link MTU size в режиме VLAN interface configuration mode:

```
ipv6 nd link-mtu value
```

По умолчанию значение параметра link MTU в RA сообщении отправленном интерфейсом установлено 0.

Когда хост источника отправляет сообщение интерфейса, производится сравнение значений MTU интерфейса и соединения link MTU. Если длина сообщения превышена более чем в два раза, то для сегментации сообщения будет использовано минимальное значение MTU.

37.3.2 Configure the number of neighbor request messages (Конфигурация количества запросов)

После получения интерфейсом IPv6 адреса, он начинает отправлять сообщения запроса соседним устройствам (neighbor request message) для проверки возможного дублирования адресов. Если ответ не получен в течение определенного времени (конфигурируется командами IPv6 nd retrans timer command), то интерфейс продолжит отправлять сообщения neighbor request message. Когда количество запросов достигнет установленного значения, а ответ так и не будет получен, считается что полученный адрес является допустимым.

Таблица команд конфигурации количества запросов.

Команда	Описание	Режим CLI
ipv6 nd dad attempts <value>	По умолчанию значение параметра количества запросов соседним устройствам установлено 1. При значении параметра 0, функция проверки дублирования адреса отключена.	Configuration mode

37.3.3 IPv6 Static routing configuration (IPv6 статическая конфигурация)

Таблица команд статической конфигурации IPv6

Команда	Описание	Режим CLI
ipv6 route <X:X::X:X/M> (<X:X::X:X> <ifName>) <distance>	Configure IPv6 static routing.	Configuration mode

37.4 IPv6 Display and maintenance (Просмотр конфигурации)

After completing the above configuration, execute the show command in the privileged view to display the operation of IPv6 after configuration, and verify the effect of configuration by viewing the display information.

Таблица команд просмотра конфигурации.

Команда	Описание	Режим CLI
show ipv6 ndp nc	Просмотр информации о соседних устройствах.	Privileged mode
show ipv6 interface (<ifName>) brief	Просмотр информации IPv6 интерфейса, которая может быть сконфигурирована совместно с IPv6 адресом.	Privileged mode
show ipv6 route (database)	Просмотр маршрутов IPv6.	Privileged mode

38. MLD Snooping configuration (Конфигурация MLD Snooping)

Когда технология Internet использует unicast для отправки одних и тех же пакетов данных большому количеству получателей в сети, возникает необходимость копировать эти пакеты для каждого конечного получателя, что приводит к увеличению количества отправляемых пакетов данных пропорционально количеству получателей. В результате происходит увеличение нагрузки на хосты, коммутаторы, роутеры и прочее сетевое оборудование, что расходует ресурсы сети и негативно сказывается на эффективности работы сети. При увеличении требований к качеству передаваемого видеоконтента (например multi-point видеоконференции), для групповой коммуникации multi-point все чаще применяется режим передачи multicast для повышения эффективности использования сетевых ресурсов.

Для обслуживания режима multicast в коммутаторе реализована функция MLD snooping function. Функция MLD snooping осуществляет

мониторинг MLD пакетов в сети и реализует динамическое отслеживание IPv6 Multicast MAC-адресов.

Глава описывает концепцию реализации и конфигурацию функции MLD snooping и содержит следующие разделы:

- Introduction (Введение)
- MLD snooping configuration (Конфигурация)
- Configuration example (Пример конфигурации)

38.1 MLD SNOOPING introduction (Введение)

В наиболее распространенных сетях пакеты multicast обрабатываются как broadcast пакеты в подсети, которые приводят к увеличению сетевого трафика и скапливанию данных. Если на коммутаторе запущена функция MLD snooping, то она динамически отслеживает IPv6 Multicast MAC адреса, обслуживает список выходных портов IPv6 Multicast MAC адресов, и направляет потоки multicast данных на выходной порт, что снижает сетевой трафик.

38.1.1 MLD snooping process (Процесс MLD snooping)

MLD snooping является сетевым протоколом уровня layer-2, который отслеживает прохождение пакетов MLD протокола через коммутатор, обслуживает группу multicast в соответствии с принимающим портом, VLAN ID и multicast адресом пакетов MLD протокола, персылает пакеты MLD протокола. Только порты входящие в multicast группу могут принимать потоки данных multicast, что снижает сетевой трафик и позволяет эффективно расходовать сетевые ресурсы. Multicast группа включает multicast group-адрес, порты члены, VLAN ID и время существования (age time).

Формирование MLD snooping multicast группы производится следующим образом: когда порт коммутатора получает пакет MLD report, функция MLD snooping создает новую multicast группу, и порт получающий пакет MLD report включается в multicast группу. Когда коммутатор получает пакет ответа MLD query (если группа multicast уже существует), порт получающий пакет MLD query также будет включен в multicast группу, в противном случае пакет MLD query будет

перенаправлен. Функция MLD snooping также поддерживает механизм исполнения MLD V2. Если функция MLD snooping сконфигурирована с возможностью быстрого выхода, то её принимающий порт может выйти из multicast группы сразу при получении пакета done packet MLD V2. Если сконфигурирован параметр fast leave timeout определяющий время задержки, то принимающий порт может выйти из multicast группы только по истечении установленного времени.

Функция MLD snooping поддерживает два механизма обновления, один из которых описан ниже. В большинстве случаев функция MLD snooping удаляет multicast группы с истекшим временем существования (age time). Функция MLD snooping фиксирует время создания multicast группы. По истечении установленного временем существования (age time) механизм exchange opportunity удаляет multicast группу из коммутатора.

Когда порт получает done protocol пакет, то он немедленно удаляется из multicast группы, что может привести к нарушению потока данных в сети т.к. к порту может быть подключено оборудование не поддерживающее функцию MLD snooping, которое связано с другими устройствами получающими потоки данных multicast. Отправка одним из устройств пакета done может оказать влияние на другие устройства, которые не смогут получить поток multicast данных. Использование времени задержки (fast leave timeout mechanism) для выхода из multicast группы позволяет предотвратить данную ситуацию. Конфигурация времени ожидания отправки (departure waiting time) производится через fast leave timeout. После получения пакета leave packet, порт некоторое время ожидает fast leave timeout, после чего удаляется из multicast группы, что обеспечивает непрерывность потока multicast данных в сети.

38.1.2 Layer 2 Dynamic Multicast (Динамический Multicast)

Таблица multicast forwarding table получает данные о multicast MAC адресах оборудования уровня layer 2 от MLD snooping dynamic learning динамическим образом. IPv6 Multicast MAC адреса также динамически поступают от протокола MLD snooping.

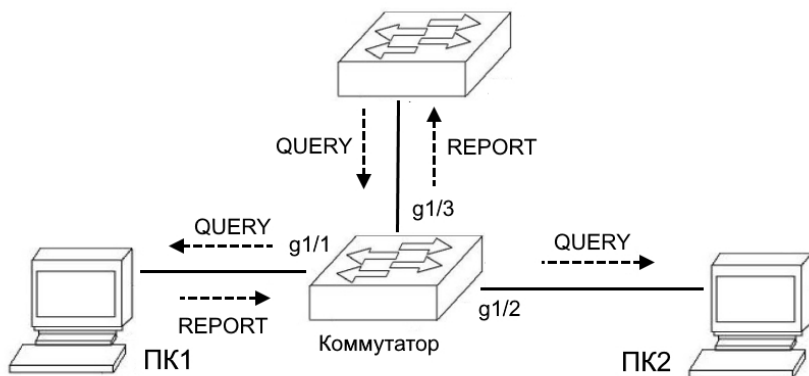
Если на коммутаторе отключена функция MLD snooping, таблица multicast forwarding table оборудования уровня layer-2

находится в режиме `unregistered forwarding mode`, и multicast MAC адрес не может быть получен динамически. В таблице `multicast forwarding table` оборудования уровня `layer-2` все потоки `layer-2 multicast` данных обрабатываются как `broadcast` данные.

На коммутаторе может быть включена функция `MLD snooping` для эффективного управления multicast трафиком в случае, если сеть имеет внешнее окружение поддерживающее multicast. Одновременно таблица `multicast forwarding table` оборудования уровня `layer-2` находится в режиме `registered forwarding mode`. Коммутатор может получать данные о multicast MAC адресах путем получения пакетов `MLD` протокола по сети. Передаются только `layer-2 multicast` потоки с учетом данных таблицы `multicast forwarding table` оборудования уровня `layer-2`.

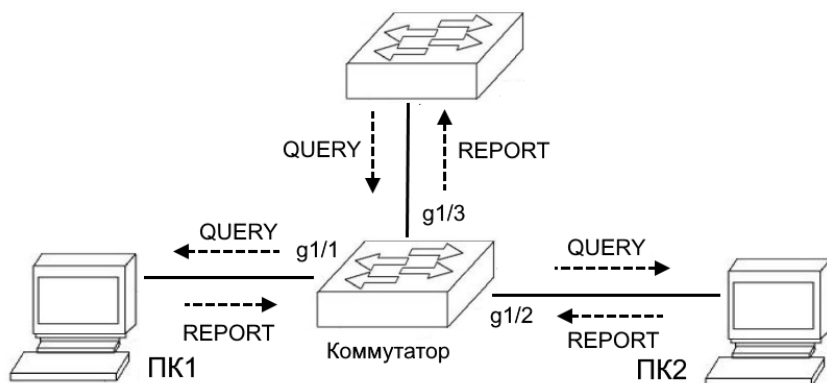
38.1.3 Join a group (Включение в группу)

Для того чтобы вступить в multicast группу хост отправляет пакет `MLD report packet`, который определяет multicast группу в которую нужно вступить. Когда коммутатор получает `MLD query` пакет, он перенаправляет пакет всем другим портам `VLAN`. Когда хост порта вступает в multicast группу, группа получает `MLD query` пакет запроса и возвращает `MLD report` пакет. Когда коммутатор получает `MLD report` пакет, он создает `layer-2 multicast entry`. Порт получающий `MLD query` пакет и отправляющий `MLD report` пакет будет добавлен в `layer-2 multicast entry` и станет его выходным портом.



Как показано на рисунке выше предполагается, что подсеть имеет VLAN 2. На роутере запущен MLDv2 протокол который периодически отправляет MLD query пакеты. Host 1 должен вступить в multicast группу ff15:: 1. После получения MLD query пакета от порта 1/3, коммутатор запишет порт и перенаправит пакет портам 1/1 и 1/2. После получения MLD query пакета Host 1 отправляет обратно MLD report пакет. Host 2 не отправляет MLD report пакет т.к. он не вступает в multicast группу. После получения MLD report пакета от порта 1/1 коммутатор перенаправит пакет от query порта 1/3 и создаст layer-2 multicast entry (считается что entry не существует). Layer-2 multicast entry имеет следующий вид:

Layer 2 multicast адрес	VLAN ID	Output port list
33:33:00:00:00:01	2	1/1 , 1/3



Как показано на рисунке выше, исходные условия те же что и в предыдущем случае. Host 1 вступил в multicast группу ff15:: 1, но и host 2 тоже должен вступить в multicast группу ff15:: 1. Когда host 2 получает MLD query пакет и отправляет обратно MLD report пакет, коммутатор перенаправит пакет от query порта 1/3 после получения MLD report пакета от порта 1/2, и порт 1/2 будет добавлен в layer 2 multicast entry. Layer 2 multicast entry станет иметь следующий вид:

Layer 2 multicast адрес	VLAN ID	Output port list
33:33:5e:00:00:01	2	1/1 , 1/2 , 1/3

38.1.4 Leave a group (Выход из группы)

Для формирования устойчивой multicast среды, устройства (роутеры) с запущенным протоколом MLD с определенным интервалом отправляют пакеты MLD query всем хостам. Хосты, которые должны вступить в multicast группу, после получения сообщений MLD query отправляют обратно пакеты MLD report.

Существует два способа чтобы хост смог покинуть multicast группу: активный и пассивный. Активный выход означает то что хост отправляет MLD leave пакет роутеру. Пассивный выход означает то что хост после получения сообщений MLD query от роутера не отправляет обратно сообщение MLD report.

В соответствии со способом выхода из multicast группы существует два способа покинуть layer-2 multicast entry на порту коммутатора: выход по истечении времени timeout и выход при получении MLD done пакета.

Когда коммутатор не получает MLD report пакет multicast группы от порта в течении определенного времени, то порт удаляется из соответствующего layer-2 multicast entry. Если layer-2 multicast entry не имеет порта, то layer-2 multicast entry будет удален.

В случае когда быстрый выход (fast leave) на коммутаторе отключен, а порт получает MLD leave пакеты multicast группы, порт будет удален из соответствующего layer-2 multicast entry. Если layer-2 multicast entry не имеет порта, то layer-2 multicast entry будет удален.

Быстрый выход в основном используется в случае когда один порт соединен с хостом. Когда с портом соединено более одного хоста может быть сконфигурирован параметр fast leave timeout (время ожидания) для обеспечения непрерывности передачи потоков multicast данных в сети.

38.2 MLD SNOOPING Configuration (Конфигурация MLD SNOOPING)

38.2.1 MLD SNOOPING default configuration (Конфигурация MLD SNOOPING по умолчанию)

Протокол MLD snooping отключен по умолчанию, таблица multicast forwarding table оборудования layer-2 находится в режиме unregistered forwarding mode.

Начальные установки по умолчанию:

Режим Fast leave отключен.

Время fast leave timeout - 300 сек.

Время существования age time multicast группы для report порта - 400 с.

Время существования age time multicast группы для query порта - 300 с.

38.2.2 On and off MLD SNOOPING (Включение и отключение MLD SNOOPING)

Включение протокола MLD snooping protocol может быть выполнено в глобальным или частным образом. Протокол MLD snooping может быть включен или отключен на VLAN только когда MLD snooping включен глобально.

Команды глобального включения протокола global MLD snooping:

```
Switch#configure terminal
```

```
Switch(config)#ipv6 mld snooping
```

Команды включения протокола global MLD snooping на VLAN:

```
Switch#configure terminal
```

```
Switch(config)#ipv6 mld snooping vlan <vlan-id>
```

Команды глобального отключения протокола global MLD snooping:

```
Switch#configure terminal
```

```
Switch(config)#no ipv6 mld snooping
```

Команды включения протокола global MLD snooping на VLAN:

```
Switch#configure terminal
```

```
Switch(config)#no ipv6 mld snooping vlan <vlan-id>
```

38.2.3 Configure lifetime (Конфигурация lifetime)

Команды конфигурации параметра lifetime multicast группы:

```
Switch#configure terminal
Switch(config)#ipv6 mld snooping group-membership-timeout <interval>
vlan <vlan-id>
```

Значение параметра устанавливается в мс.

Команды конфигурации параметра query group lifetime:

```
Switch#configure terminal
Switch(config)#ipv6 mld snooping query-membership-timeout <interval>
vlan <vlan-id>
```

Значение параметра устанавливается в мс.

38.2.4 Configure fast-leave (Конфигурация fast-leave)

Команды включения функции fast-leave на VLAN:

```
Switch#configure terminal
Switch(config)#ipv6 mld snooping fast-leave vlan <vlan-id>
```

Команды отключения функции fast-leave на VLAN:

```
Switch#configure terminal
Switch(config)#no ipv6 mld snooping fast-leave vlan <vlan-id>
```

Команды конфигурации параметра fast leave wait time:

```
Switch#configure terminal
Switch(config)# ipv6 mld snooping fast-leave-timeout <interval> vlan
<vlan-id>
```

Команды установки значения по умолчанию параметра fast leave wait time:

```
Switch#configure terminal
Switch(config)#no ipv6 mld snooping fast-leave-timeout vlan <vlan-id>
```

38.2.5 Configure MROUTER (Конфигурация MROUTER)

Команды конфигурации static query port:

```
Switch#configure terminal
```

```
Switch#interface ge1/6
Switch(config-ge1/6)#ipv6 mld snooping mrouter vlan [vlan-id]
```

38.2.6 Configure VLAN querier (Конфигурация VLAN querier)

Команды конфигурации и запуска query function на VLAN 1:

```
Switch#configure terminal
Switch(config)#ipv6 mld snooping querier vlan 1
```

38.2.7 Display information (Просмотр конфигурации)

Команда просмотра информации о конфигурации протокола MLD snooping:

```
Switch#show ipv6 mld snooping
```

Команда просмотра информации конфигурации VLAN:

```
Switch#show ipv6 mld snooping vlan <vlan-id>
```

Команда просмотра информации aging report multicast группы:

```
Switch#show ipv6 mld snooping age-table group-membership
```

Команда просмотра информации aging query:

```
Switch#show ipv6 mld snooping age-table query-membership
```

Команда просмотра forwarding information multicast группы:

```
Switch#show ipv6 mld snooping forwarding-table
```

Команда просмотра информации mrouter:

```
Switch#show ipv6 mld snooping mrouter
```

Команда просмотра текущей конфигурации системы, включая конфигурацию протокола MLD snooping:

```
Switch#show running-config
```

38.3 MLD Snooping Configuration example (Пример конфигурации)

Включить функцию MLD snooping на коммутаторе, пользователи ПК1, ПК2 и ПК3 должны войти в multicast группу. Схема подключения приведена на рисунке ниже:



Команды конфигурации коммутатора:

```
Switch#config t
Switch(config)#vlan database
Switch(config-vlan)#vlan 200
Switch(config)#ipv6 mld snooping

Switch(config)#interface ge1/1
Switch(config-ge1/1)#switchport mode access
Switch(config-ge1/1)#switchport access vlan 200

Switch(config)#interface ge1/2
Switch(config-ge1/2)#switchport mode access
Switch(config-ge1/2)#switchport access vlan 200

Switch(config)#interface ge1/3
Switch(config-ge1/3)#switchport mode access
Switch(config-ge1/3)#switchport access vlan 200

Switch(config)#ipv6 mld snooping group-membership-timeout 60000 vlan
200
```

39. POE configuration (Конфигурация POE)

Коммутатор поддерживает технологию POE (power over Ethernet) – возможность передачи питания оконечным устройствам по кабелю витой пары. Глава описывает дополнительные возможности этой функции и содержит следующие разделы:

- Introduction (Введение)
- POE Configuration (Конфигурация PoE)

39.1 POE introduction (Введение)

Технология POE (power over Ethernet) относится к устройствам осуществляющим передачу и получающим питание по кабелю витой пары оконечным устройствам (PD - powered device) таким как IP телефоны, беспроводные точки доступа, web камеры и т.п.

1) Преимущества технологии PoE:

- Надежность: централизованное питание, удобство управления.
- Простота подключения: отсутствие необходимости подачи внешнего питания, использование одного кабеля.
- Стандарты: соответствует стандартам IEEE 802.3af и 802.3at, применение глобально унифицированного интерфейса.
- Широкие перспективы применения: использование для IP телефонии, питания беспроводных точек доступа (wireless AP), webcam и т.п.

2) Совместимость PoE устройств:

- Система PoE включает источники PoE (PSE) и потребителей PoE (PD).
- Pse power supply (источники PoE) обеспечивают питанием всю систему PoE, разделяются на внешние (external power supply) и встроенные (built-in power supply).
- PSE (power sourcing equipment) - единое устройство. Каждое PSE управляет интерфейсом POE независимо. PSE находит и определяет потребителей PD подключенных к интерфейсу, классифицирует PD и подает им питание. При обнаружении отключения потребителя PD, PSE останавливает подачу питания.

Ethernet интерфейс с возможностью подачи POE называется PoE интерфейсом (включая FE и GE).

- PD – устройство получающее питание от PSE. Такие устройства подразделяются на стандартные и нестандартные PD. К стандартным PD относятся устройства с поддержкой IEEE 802.3af и 802.3at стандартов. PD оборудование допускает подключение других источников питания для обеспечения резервирования питания.

39.2 POE configuration (Конфигурация PoE)

39.2.1 Manual PoE configuration (Ручная конфигурация PoE)

Ниже представлены команды включения/отключения и просмотра конфигурации PoE на интерфейсе:

Команда	Описание	Режим CLI
poe enable	Включение PoE на интерфейсе. По умолчанию PoE на интерфейсе включено.	Interface configuration mode
no poe enable	Отключение PoE на интерфейсе. По умолчанию PoE на интерфейсе включено.	Interface configuration mode
show poe	Просмотр конфигурации PoE для всех интерфейсов.	User mode or privileged user mode

39.2.2 PoE Policy configuration (Конфигурация политики PoE)

Ниже представлены команды конфигурации политики PoE:

- Включение/отключение политики PoE на интерфейсе;
- Конфигурация политики PoE на интерфейсе;
- Просмотр конфигурации политики PoE.

Команда	Описание	Режим CLI
poe policy enable	Включение политики PoE на интерфейсе. По умолчанию политика PoE на интерфейсе отключена.	Interface configuration mode
no poe policy enable	Отключение политики PoE на интерфейсе. По умолчанию политика PoE на интерфейсе отключена.	Interface configuration mode
poe policy shutdown clock <clock-value> week-day <day-value>	Установка параметров политики PoE на интерфейсе. Команда может быть применена несколько раз. Значения параметров по умолчанию: политика PoE на интерфейсе отключена, часы указывают конкретное время или временной диапазон в 24-м формате (например значение 1 соответствует промежутку между 1 и 2, значение 20-23 соответствует диапазону от 20 до 0). Значение параметра дня недели указывает день недели (например значению 3 соответствует среда, диапазон 1-7 соответствует периоду понедельник-суббота).	Interface configuration mode
no poe policy shutdown clock <clock-value> week-day <day-value>	Сброс параметров политики PoE на интерфейсе.	Interface configuration mode
show poe policy <if-name>	Просмотр конфигурации политики PoE на интерфейсе.	User mode or privileged user mode

39.2.3 PDQuery configuration (Конфигурация запросов PD)

Ниже представлены команды конфигурации запросов подключенного оборудования PD:

- Установка количества запросов оборудования PD;
- Установка периодичности запросов оборудования PD;
- Просмотр конфигурации PDQuery.

Команда	Описание	Режим CLI
poe pd-ip-address <ip-address> no poe pd-ip-address	Установка или удаление IP адреса PD оборудования подключенного к интерфейсу. По умолчанию IP адрес не установлен. После установки IP адреса PD оборудования, система периодически начинает отправлять запросы. Если PD не отвечает определенное число раз, то коммутатор перезапускает PD путем отключения PoE.	Interface configuration mode
poe pd-query-interval <interval> no poe pd-query-interval	Установка временного интервала отправки сообщений запроса PD. По умолчанию установлено значение 5с.	Interface configuration mode
poe pd-timeout-number <number> no poe pd-timeout-number	Установка количества запросов, на которые может быть не получен ответ PD. По умолчанию установлено значение 3.	Interface configuration mode
poe pd-boot-time <time> no poe pd-boot-time	Установка времени перезагрузки PD. По умолчанию установлено значение 120с.	Interface configuration mode
show poe pd-information	Просмотр информации обо всех подключенных PD устройствах	User mode, privileged user mode

4
230731