



کتاب

**CCNP (Routing and Switching)**

به زبان فارسی



# CCNP

## Routing and Switching

نویسنده:

مهندس ابوالفضل هاشمی

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

# CCNP Route & Switch Fast Track

نویسند:

مهندس ابوالفضل هاشمی

[A.Hashemi70@gmail.com](mailto:A.Hashemi70@gmail.com)

09134232203

## فهرست

|    |       |                                                      |
|----|-------|------------------------------------------------------|
| ۱۱ | ..... | مقدمه                                                |
| ۱۲ | ..... | تعاریف                                               |
| ۱۳ | ..... | EIGRP (Enhance Interior Gateway Routing Protocol) -۱ |
| ۱۳ | ..... | EIGRP Introduction -۱-۱                              |
| ۱۵ | ..... | EIGRP Authentication -۲-۱                            |
| ۱۶ | ..... | EIGRP Stop Neighbors Connectivity -۳-۱               |
| ۱۶ | ..... | Administrative Distance -۴-۱                         |
| ۱۷ | ..... | EIGRP Router ID -۵-۱                                 |
| ۱۷ | ..... | EIGRP Relationship -۶-۱                              |
| ۱۸ | ..... | Feasible & Reported Distance -۷-۱                    |
| ۱۹ | ..... | Change Metric -۸-۱                                   |
| ۲۰ | ..... | Convergence EIGRP -۹-۱                               |
| ۲۱ | ..... | EIGRP Variance -۱۰-۱                                 |
| ۲۱ | ..... | EIGRP Access Control List -۱۱-۱                      |
| ۲۱ | ..... | Access Control List -۱-۱۱-۱                          |
| ۲۲ | ..... | Prefix List -۲-۱۱-۱                                  |
| ۲۳ | ..... | Route-map -۳-۱۱-۱                                    |
| ۲۳ | ..... | EIGRP Route Summarization -۱۲-۱                      |
| ۲۴ | ..... | EIGRP Default Route -۱۳-۱                            |
| ۲۵ | ..... | EIGRP Manipulation -۱۴-۱                             |
| ۲۶ | ..... | EIGRP Named Mode -۱۵-۱                               |

|          |                                     |       |
|----------|-------------------------------------|-------|
| ۲۶ ..... | OSPF (Open Shortest Path First)     | -۲    |
| ۲۸ ..... | OSPF Timers                         | -۱-۲  |
| ۲۹ ..... | OSPF Router-ID                      | -۲-۲  |
| ۲۹ ..... | OSPF MTU                            | -۳-۲  |
| ۳۰ ..... | OSPF Authentication                 | -۴-۲  |
| ۳۱ ..... | Frame Relay Introduction            | -۵-۲  |
| ۳۱ ..... | Type of Networks in OSPF            | -۶-۲  |
| ۳۳ ..... | OSPF LSA                            | -۷-۲  |
| ۳۴ ..... | OSPF LSA Type-1 (Router-LSA)        | -۸-۲  |
| ۳۵ ..... | OSPF LSA Type-2 (Network-LSA)       | -۹-۲  |
| ۳۶ ..... | :OSPF LSA Type-3 (Summary-LSA)      | -۱۰-۲ |
| ۳۷ ..... | OSPF LSA Type-4 (ASBR -Summary-LSA) | -۱۱-۲ |
| ۳۸ ..... | OSPF LSA Type-5 (ASBR-External-LSA) | -۱۲-۲ |
| ۳۸ ..... | OSPF Neighbor Relationship          | -۱۳-۲ |
| ۴۰ ..... | OSPF Metric Calculation             | -۱۴-۲ |
| ۴۱ ..... | OSPF Filtering                      | -۱۵-۲ |
| ۴۲ ..... | OSPF Route Summarization            | -۱۶-۲ |
| ۴۲ ..... | OSPF Default Route                  | -۱۷-۲ |
| ۴۳ ..... | OSPF Stub Area                      | -۱۸-۲ |
| ۴۵ ..... | OSPF Virtual Link                   | -۱۹-۲ |
| ۴۶ ..... | OSPF Path Manipulation              | -۲۰-۲ |
| ۴۶ ..... | Redistribution                      | -۲۱-۲ |

|          |                                       |       |
|----------|---------------------------------------|-------|
| ٤٨ ..... | Prevention Loops to Redistribute OSPF | -٢٢-٢ |
| ٥٠ ..... | Redistribute OSPF Path Manipulation   | -٢٣-٢ |
| ٥١ ..... | Routing Border Gateway Protocol       | -٣    |
| ٥٣ ..... | iBGP & eBGP                           | -١-٣  |
| ٥٤ ..... | BGP Neighbor Relationship             | -٢-٣  |
| ٥٤ ..... | eBGP                                  | -٣-٣  |
| ٥٥ ..... | iBGP                                  | -٤-٣  |
| ٥٥ ..... | BGP Neighbor State                    | -٥-٣  |
| ٥٦ ..... | BGP Packet Type                       | -٦-٣  |
| ٥٧ ..... | BGP Table                             | -٧-٣  |
| ٥٨ ..... | BGP AS_PATH Attribute                 | -٨-٣  |
| ٥٩ ..... | Route into BGP                        | -٩-٣  |
| ٦٠ ..... | Outgoing BGP Scenario                 | -١٠-٣ |
| ٦٠ ..... | BGP Route Advertisement               | -١١-٣ |
| ٦٢ ..... | Prevent Black Hole in BGP             | -١٢-٣ |
| ٦٣ ..... | Peer Group                            | -١٣-٣ |
| ٦٣ ..... | BGP Filtering                         | -١٤-٣ |
| ٦٥ ..... | BGP Path Attribute                    | -١٥-٣ |
| ٦٧ ..... | BGP Maximum-paths                     | -١٦-٣ |
| ٦٨ ..... | BGP Multi-Protocol                    | -١٧-٣ |
| ٦٩ ..... | IP Service Level Agreement (IP-SLA)   | -٤    |
| ٧٠ ..... | Policy-Based Routing                  | -٥    |

|    |                                            |        |
|----|--------------------------------------------|--------|
| ୮୧ | .....Network Time Protocol (NTP)           | -୨     |
| ୮୩ | ..... IPv6 Introduction                    | -୮     |
| ୮୪ | ..... RIPng, RIPv2                         | -୧-୮   |
| ୮୫ | ..... EIGRP IPv6                           | -୨-୮   |
| ୮୫ | ..... OSPF v3 IPv6                         | -୩-୮   |
| ୮୮ | ..... LSA Type                             | -୧-୩-୮ |
| ୮୮ | .....IPv6 Static Route & Access-List       | -୪-୮   |
| ୮୮ | .....Redistribute IPv6 & IPV4              | -୫-୮   |
| ୮୮ | ..... IPv4 & IPv6 translation              | -୨-୮   |
| ୮୯ | .....Tunneling                             | -୮-୮   |
| ୯୩ | .....NAT64                                 | -୮-୮   |
| ୯୪ | ..... DNS64                                | -୯-୮   |
| ୯୬ | ..... VPN                                  | -୮     |
| ୯୬ | ..... IPsec Introduction                   | -୧-୮   |
| ୯୮ | .....Dynamic Multipoint VPN (DMVPN)        | -୨-୮   |
| ୯୮ | .....GRE Introduction                      | -୩-୮   |
| ୯୮ | ..... NHRP Introduction                    | -୪-୮   |
| ୯୯ | .....DMVPN Configuration                   | -୫-୮   |
| ୧୦ | ..... VRF (Virtual Routing and Forwarding) | -୯     |
| ୧୧ | ..... Switch Introduction                  | -୧୦    |
| ୧୨ | .....Etherchannel                          | -୧୧    |
| ୧୩ | ..... PAgP vs LACP                         | -୧-୧୧  |

|           |                                            |         |
|-----------|--------------------------------------------|---------|
| ٩٤ .....  | Etherchannel Layer                         | -٢-١١   |
| ٩٤ .....  | Etherchannel Misconfiguration Guard        | -٣-١١   |
| ٩٤ .....  | Virtual Trunking Protocol (VTP)            | ١٢-     |
| ٩٨ .....  | VTP Pruning                                | -١-١٢   |
| ٩٩ .....  | Multi-Layer Switch                         | -١٣     |
| ١٠٠ ..... | Cisco Express Forwarding (CEF)             | -١٤     |
| ١٠٢ ..... | CEF vs TCAM                                | -١-١٤   |
| ١٠٣ ..... | Switch Software Introduction               | -١٥     |
| ١٠٣ ..... | Switch Redundancy                          | -١٤     |
| ١٠٥ ..... | Virtual Switching System (VSS)             | -١-١٤   |
| ١٠٥ ..... | DHCP Server                                | -١٧     |
| ١٠٤ ..... | Basic Manage Switch Port                   | -١٨     |
| ١٠٧ ..... | Port Mirroring                             | -١٩     |
| ١٠٨ ..... | Common Spanning Tree 802.1d                | -٢٠     |
| ١١١ ..... | 802.1d Optimize and Security               | -١-٢٠   |
| ١١٣ ..... | 802.1w Rapid Spanning Tree Protocol (RSTP) | -٢-٢٠   |
| ١١٥ ..... | Per Vlan Spanning Tree (PVST)              | -٣-٢٠   |
| ١١٤ ..... | Multiple Spanning Tree (MST) 802.1s        | -٤-٢٠   |
| ١٢٠ ..... | High Availability                          | -٢١     |
| ١٢١ ..... | Hot Standby Router Protocol (HSRP)         | -١-٢١   |
| ١٢٢ ..... | HSRP State Machine                         | -١-١-٢١ |
| ١٢٢ ..... | HSRP Modify Preemption                     | -٢-١-٢١ |
| ١٢٣ ..... | HSRP Authentication                        | -٣-١-٢١ |

|           |                                                    |         |
|-----------|----------------------------------------------------|---------|
| ۱۲۳ ..... | HSRP Object Tracking                               | -۴-۱-۲۱ |
| ۱۲۳ ..... | Virtual Router Redundancy Protocol (VRRP)          | -۲-۲۱   |
| ۱۲۵ ..... | Gateway Load Balancing Protocol (GLBP)             | -۳-۲۱   |
| ۱۲۵ ..... | Active Virtual Gateway (AVG)                       | -۱-۳-۲۱ |
| ۱۲۸ ..... | Supervisor & Route Processor Redundancy            | -۲۲     |
| ۱۲۹ ..... | Route Processor Redundancy (RPR)                   | -۱-۲۲   |
| ۱۳۰ ..... | RPR+                                               | -۲-۲۲   |
| ۱۳۰ ..... | SSO (State SwitchOver)                             | -۳-۲۲   |
| ۱۳۱ ..... | SSO with Non-Stop Forwarding (NSF)                 | -۴-۲۲   |
| ۱۳۲ ..... | POE (Power over Ethernet)                          | -۲۳     |
| ۱۳۳ ..... | Cisco Inline Power (ILP)                           | -۱-۲۳   |
| ۱۳۳ ..... | IEEE 802.3af                                       | -۲-۲۳   |
| ۱۳۴ ..... | IEEE 802.3at                                       | -۳-۲۳   |
| ۱۳۴ ..... | Cisco Universal Power over Ethernet (UPOE)         | -۴-۲۳   |
| ۱۳۵ ..... | Voice Vlan                                         | ۲۴-     |
| ۱۳۶ ..... | Voice QoS                                          | -۱-۲۴   |
| ۱۳۸ ..... | QoS ابزارهای                                       | -۲-۲۴   |
| ۱۳۹ ..... | QoS Models                                         | -۳-۲۴   |
| ۱۴۲ ..... | Configure QoS                                      | -۴-۲۴   |
| ۱۴۳ ..... | Auto QoS                                           | -۵-۲۴   |
| ۱۴۳ ..... | Securing Switch Access                             | -۲۵     |
| ۱۴۳ ..... | Authentication, Authorization and Accounting (AAA) | -۱-۲۵   |

|           |                                         |         |
|-----------|-----------------------------------------|---------|
| ۱۴۴ ..... | Relationship between AAA Server and NAS | -۱-۱-۲۵ |
| ۱۴۵ ..... | AAA IOS Configuration                   | -۲-۱-۲۵ |
| ۱۴۶ ..... | Securing Switch Access (Storm Control)  | -۲-۲۵   |
| ۱۴۸ ..... | DHCP Snooping                           | -۳-۲۵   |
| ۱۴۹ ..... | Dynamic Arp Inspection                  | -۴-۲۵   |
| ۱۵۲ ..... | Security VLAN                           | -۲۶     |
| ۱۵۲ ..... | VLAN Access List                        | -۱-۲۶   |
| ۱۵۳ ..... | Private VLAN                            | -۲-۲۶   |

## مقدمه

دوره CCNP (Cisco Certified Network Profesional) یکی از دوره‌های شرکت سیسکو است که زیر مجموعه‌هایی مانند Security، Voice، Switch، Route و... را داراست. در این مجموعه به شرح مباحث اصولی Route و Switch پرداخته می‌شود. منبع این مباحث از ترجمه ویدیوهای مربوط به شرکت INE که توسط آقای Keith Bogart ارائه شده است، سایت شرکت سیسکو ([www.Cisco.com](http://www.Cisco.com)) و همچنین اضافه شدن مطالب جامع‌تری از استانداردها کامل گردیده است. سلسله مراتب به صورت توضیح پروتکل‌های مسیریابی OSPF، EIGRP و BGP است که پس از آن موضوعات پراکنده IPv6، NTP و دیگر موارد شرح داده شده است. در آخر مطالب مربوط به Switching و امنیت آن اضافه گردیده است.

## تعاریف

**Hello packet:** بسته‌های سلام است که معمولاً در پروتکل‌های مسیریابی بین دو همسایه برای برقراری ارتباط ارسال می‌شود که با یکدیگر آشنا شوند.

**VLSM:** یا همان Variable Length Subnet Masking در شبکه‌ها معمولاً برای جلوگیری از هدر رفتن آدرس‌های IP و یا ایجاد یک شبکه کوچک استفاده می‌شود. به‌طور مثال 192.168.1.1 با 255.255.255.0 به‌صورت 255.255.255.192 استفاده می‌شود.

**Classful:** آدرس‌هایی هستند که subnet mask آنها استاندارد است و در یکی از کلاس‌های A, B, C, D و یا E قرار دارند.

**Classless:** آدرس‌هایی هستند که از هر subnet mask دلخواهی می‌توان استفاده نمود. در واقع از VLSM استفاده می‌شود.

**Wildcard:** به‌صورت برعکس subnet mask است. Subnet mask = 255.255.255.0 و wild card = 0.0.0.255.

**Auto Summary:** این عمل در مسیریاب‌ها باعث خلاصه نمودن بسته‌های مسیریابی ارسالی می‌شود به‌نحوی که بسته‌های Classless را ارسال نمی‌کند.

**Frame Relay:** این دستگاه بین دو اینترفیس یک Virtual circuit ایجاد می‌کند. بدین معنا که با تنظیم آن می‌توان نشان داد که هر بسته از پورتهای که وارد می‌شود از چه پورتهای خارج شود.

**Multi-protocol label switching (MPLS):** روشی است برای یک ارتباط مطمئن و معتبر به‌صورت real time برای برنامه‌های کاربردی که پرمزینه است. این تکنیک مسیریابی اطلاعات را از یک گره به سمت دیگر بر اساس برجسب‌های مسیر کوتاه‌تر به جای آدرس‌های طولانی شبکه هدایت می‌کند، بنابراین از جست‌وجوی پیچیده در یک جدول مسیریابی اجتناب می‌کند و سرعت را افزایش می‌دهد. فرض کنید بین دو مسیریاب بین ISP و مشتری ارتباط برقرار است که هر دو از آدرس‌های مشابه استفاده می‌کنند. برای اینکه از چند دستگاه استفاده نشود و آدرس‌های مشابه با هم تداخل نداشته باشند هر مسیریاب چندین جدول مسیریابی مجزا ایجاد می‌کند و آدرس‌های مربوط به هر سازمان را در خود جای می‌دهد. این جدول VRF نام دارد. در داخل MPLS VPN از Virtual Routing & Forwarding (VRF) استفاده می‌شود.

VPLS: یا Virtual Private LAN Service یک راه بر پایه اترنت است که برای ارائه چند نقطه به چند نقطه از طریق شبکه IP یا MPLS تأمین می‌شود. همچنین اجازه می‌دهد تا سایت‌های پراکنده جغرافیایی دامنه‌های همه بخشی اترنت خود را به اشتراک بگذارند. در واقع VPLS یک تکنولوژی VPN است.

Split-horizon: روشی است برای جلوگیری حلقه در مسیریابی‌های Distance-Vector به این شکل که به مسیریاب اجازه نمی‌دهد بسته‌های دریافتی از مسیریاب‌های قبلی به همان اینترفیس فرستاده شود.

Commit information rate (CIR): میزان پهنای باندی است که گارانتی شده است و از حذف شدن بسته جلوگیری می‌شود به عبارتی Packet Drop وجود ندارد.

## ۱- EIGRP (Enhance Interior Gateway Routing Protocol)

قبل از اینکه به شرح پروتکل EIGRP پرداخته شود، انواع پروتکل مسیریابی به اختصار توضیح داده می‌شود. پروتکل‌های مسیریابی می‌تواند کاربرد آن داخلی و یا بین شبکه‌ای باشد. به پروتکل‌هایی که داخل شبکه کاربرد دارند، Interior Gateway Protocol (IGP) گویند و به پروتکل‌هایی که بین دو شبکه مسیریاب را جابه‌جا می‌کنند، Exterior Gateway Protocol (EGP) گویند. انواع پروتکل‌های داخل شبکه‌ای به شرح زیر است:

- Link State (OSPF): هر مسیریاب کل توپولوژی را می‌داند.
- Distance Vector (RIP): هیچ چیزی از توپولوژی شبکه نمی‌داند.
- Advanced Distance Vector (EIGRP): این پروتکل Classless و مخصوص شرکت cisco است.

### ۱-۱ EIGRP Introduction

از روش DUAL استفاده می‌کند که ابتدا بهترین مسیر loop-free و سریعترین مسیر و بعد از آن backup را انتخاب می‌کند. این پروتکل از سه مرحله تشکیل می‌شود:

- Neighbor relationship: ارتباط با همسایه‌ها جهت برقراری ارتباط.
- Topology exchange information: به اشتراک گذاری نقشه شبکه جهت بدست آوردن بهترین مسیر که فقط اطلاعات همسایه‌ها ارسال می‌شود.
- DUAL to populate the routing table : Diffusing Update Algorithm این الگوریتم با ارسال بسته‌های به روزرسانی جدول مسیریابی را به روز می‌نماید.

Hello packet (بسته سلام جهت برقراری ارتباط با دستگاه دیگر) را در دو موقع ارسال می‌کند new neighbor (پیدا کردن همسایه جدید) و maintain (اعلام زنده بودن مسیریاب).

برای این پروتکل می توان درگاه خروجی را برای مسیریابی مشخص نمود که فقط از آن درگاه مسیریابی انجام دهد. در تنظیمات پروتکل EIGRP یک شماره سیستمی خودمختار به نام Autonomous System Number یا ASN وارد می شود که طراحی شده است برای دریافت بسته های مسیریابی که فقط در این شبکه autonomous یا AS قرار دارند. این الگوریتم بسته هایی که در یک AS قرار دارند را برای مسیریاب های دیگر به اشتراک می گذارد و اولویت بالاتری برای آنها لحظ می کند. اگر AS متفاوت بود بسته های مخصوص به این شبکه ها به اشتراک گذاشته نمی شود.

ویژگی Auto-summary برای خلاصه کردن مسیرهایی است که به صورت subnet یا Variable Length Subnet Masking (VLSM) است که به صورت پیش فرض فعال است.

Subnet mask = 255.255.255.0 و wild card = 0.0.0.255.

پروتکل EIGRP ۵ نوع بسته ارسالی دارد که شامل موارد زیر است:

- بسته Hello: این بسته برای پیدا نمودن همسایه ها ارسال می شود. از آدرس چند پخشی 244.0.0.10 به صورت پیش فرض ارسال می شود. در این پروتکل هر ۵ ثانیه یکبار و در شبکه های WAN با سرعت ۱۵۴۴ مگابیت هر ۶۰ ثانیه یکبار ارسال می شود. این بسته شامل هیچ اطلاعات مسیریابی نیست.
  - بسته Acknowledgment: این بسته در جواب رسیدن بسته های به روزرسانی ارسال می شود. این بسته به صورت unicast برای فرستنده ارسال می شود.
  - بسته Update: این بسته شامل اطلاعات مسیریابی است که برای مقصد ارسال می شود. این بسته ها به صورت unicast برای همسایه مجاور و جدید ارسال می شود. اگر Metric یک لینک تغییر پیدا نمود، بسته به روزرسانی به صورت چند پخشی 224.0.0.10 ارسال می شود. در جواب برای اطمینان از دریافت بسته ACK ارسال می شود.
  - بسته Query: این بسته برای پیدا نمودن feasible successors استفاده می شود. این بسته همیشه به صورت چند پخشی است. در واقع successors بهترین مسیر است و feasible successors مسیر بعدی و یا جایگزین successors است.
  - بسته Reply: این بسته پاسخی برای بسته های Query است و به صورت unicast ارسال می شود.
- برای همسایگانی که به صورت دستی یا static neighbors هستند نیز از IP Unicast استفاده می شود.

مواردی که باید به آن توجه داشت:

- بسته‌ها با primary IP ارسال می‌شود و هر دو مسیر یاب باید در یک دامنه باشند.
- Interface نباید در حالت passive باشد.
- K-value ها در مسیر یاب‌ها باید یکی باشند. در کل پارامترها باید یکی باشند. پیش فرض  $k_1=k_3=1, k_2=k_4=k_5=0$ . این پارامترها به معنای:

K1 = Bandwidth modifier

K2 = Load modifier

K3 = Delay modifier

K4 = Reliability modifier

K5 = Additional Reliability modifier

K6 = for jitter and energy

- با استفاده از K-Value ها می‌توان Metric را محاسبه نمود. اگر  $k_5=0$  کسر  $k_5/$  (reliability+k4) برابر ۱ در نظر گرفته می‌شود..

$$\text{Metric} = 256 * [k_1 * \text{bandwidth} + (K_2 * \text{bandwidth}) / (256 - \text{load}) + k_3 * \text{delay}] * [k_5 / (\text{reliability} + k_4)]$$

با دستور k k4 k3 k2 k1 to metric weights مقادارها تغییر می‌کند.

## ۲-۱ EIGRP Authentication

پیش از تنظیم باید به موارد زیر توجه داشت:

- ۱- تعریف پسورد برای مسیر یاب‌ها باید یکسان باشد.
- ۲- Pre-shared باید یکسان باشد.
- ۳- فقط MD5 پشتیبانی می‌شود.
- ۴- این عمل از DoS جلوگیری می‌کند ولی در مقابل Man in the Middle ضعیف است.
- ۵- Config آن در دو سطح Global و Interface انجام می‌شود.

Key chain name of key->key name of key->key-string pass

ip authentication mode EIGRP ASnumber MD5->ip authentication key-chain EIGRP ASnumber name of key

۶- Key Chain بر مبنای Time هست و برای زمان‌های مختلف قابل تنظیم است.

۷- تنظیم زمان‌ها برای پسوندها باید یکی باشد.

۸- برای تنظیم زمان‌ها بهتر است از NTP استفاده شود ولی با clock set می‌توان زمان را تنظیم نمود.

۹- Key Chain Name نیازی نیست یکی باشد و محلی است ولی ترتیب پسوندها و خود پسوندها باید یکی باشد.

۱۰- دستور نمایش تنظیمات show key chain name of chain و show clock.

### ۳-۱ EIGRP Stop Neighbors Connectivity

برای متوقف نمودن ارسال بسته‌های EIGRP به همسایه‌ها از دستور redistributed connected استفاده می‌شود. همچنین می‌توان Interface را با دستور passive-interface در داخل تنظیمات EIGRP برای ارسال بسته‌های آن خاموش نمود. اگر redistribute را وارد شود نیازی به دستور passive نیست.

### ۴-۱ Administrative Distance

Administrative Distance مؤلفه‌ای است برای انتخاب بهترین مسیر با پروتکل‌های مختلف برای مقصد یکسان که به عنوان مقداری برای نشان دادن ارزش هر مسیر تعریف می‌شود. این مقادیر در جدول زیر آورده شده است.

| Administrative distance | Routing Protocol                      |
|-------------------------|---------------------------------------|
| 0                       | Directly connected interface          |
| 1                       | Static route out an interface         |
| 1                       | Static route to next-hop address      |
| 3                       | DMNR - Dynamic Mobile Network Routing |
| 5                       | EIGRP summary route                   |
| 20                      | External BGP                          |
| 90                      | Internal EIGRP                        |
| 100                     | IGRP                                  |
| 110                     | OSPF                                  |

|     |                                       |
|-----|---------------------------------------|
| 115 | IS-IS                                 |
| 120 | Routing Information Protocol (RIP)    |
| 140 | Exterior Gateway Protocol (EGP)       |
| 160 | On Demand Routing (ODR)               |
| 170 | External EIGRP                        |
| 200 | Internal BGP                          |
| 250 | Next Hop Resolution Protocol (NHRP)   |
| 254 | Default static route learned via DHCP |
| 255 | Unknown and unused                    |

### ۵-۱ EIGRP Router ID

Router ID در بسته‌های ارسالی دیده می‌شود بدین شکل که مسیریاب با قرار دادن Router ID خود در این فیلد به همسایه نشان می‌دهد که خود این بسته را ایجاد کرده است. Router ID از external routing loops جلوگیری می‌کند. بدین شکل که اگر فیلد Router ID بسته دریافتی با Router ID خود مسیریاب یکی باشد بدین معنا است که خود این بسته را ایجاد کرده و دوباره بدست خود رسیده است. این مقدار با استفاده از loop back تعریف می‌شود. اگر loop back تعریف نگردد خود مسیریاب بزرگترین IP را به عنوان router ID استفاده می‌کند. اگر دو مسیریاب نام‌های یکسان داشته باشند در همسایه بودن آنها تاثیری نخواهد داشت ولی در external routing جواب نمی‌دهد و دچار مشکل می‌شود.

### ۶-۱ EIGRP Relationship

در این قسمت بسیاری از مفاهیم و همچنین روابط بین مسیریاب‌ها آورده شده است.

NBMA: non broadcast multi access نمونه‌ای است که بسته‌های همسایگی پروتکل EIGRP دچار مشکل می‌شوند. در واقع این وسایل از ارسال چند پخش‌ی جلوگیری می‌شود. مثالی از آن frame relay است. در این وسایل EIGRP نمی‌تواند بسته‌های update را ارسال کند چون لایه دو از چندپخش‌ی حمایت نمی‌کند. به همین خاطر باید به صورت تک تک ارسال شود که پهنای باند را بیش از حد مصرف می‌کند به همین دلیل باید پهنای باند محدودی تعریف شود. به‌طور مثال ۵۰٪. دستور به صورت زیر است:

Interface serial 0/0/0 multipoint

Ip add 10.1.1.1 255.255.255.0

Frame-relay interface-dlci 102

Frame-relay interface-dlci 103

Bandwidth 100

Ip bandwidth-percent EIGRP 1 50

Split-horizon: روشی است برای جلوگیری حلقه در مسیریابی‌های Distance-Vector به این شکل که به مسیریاب اجازه نمی‌دهد بسته‌های دریافتی از مسیریاب‌های قبلی به همان اینترفیس فرستاده شود. به صورت پیش فرض در EIGRP، split-horizon فعال است. برای تنظیم آن از دستور `ip split-horizon EIGRP <asn>` در داخل اینترفیس استفاده می‌شود.

بعد از روابط همسایه‌ای EIGRP، Topology Table تغییر می‌کند. ولی قبل از آن باید شبکه برای هر مسیریاب تعریف شود که به دو شکل `network command` و `redistributed connect` انجام می‌شود. EIGRP با استفاده از بسته‌های `update` و `ACK` جدول خود را تغییر می‌دهد.

بسته Update شامل موارد زیر است:

- prefix
- prefix length
- metric: bandwidth, delay, load, reliable
- Non metric: hop count, MTU.

باید توجه داشت قالب و Administrative Distance بسته‌های `internal route` و `external route` تفاوت دارند.

همچنین باید توجه داشت در پروتکل EIGRP بسته‌های `multi cast` منتظر `ACK` به صورت `unicast` می‌ماند.

## ۷-۱ Feasible & Reported Distance

پس از ساخت `topology table` مسیریاب باید `routing table` را بسازد که برای ساختن آن باید دو تصمیم را بررسی نماید:

- ۱- تصمیم بگیرد این ردیف که از دریافت بسته‌ها رسیده است در جدول باعث ایجاد حلقه می‌شود یا خیر.
- ۲- تصمیم بگیرد که اگر این ردیف صحیح است باید در کجای جدول مسیریابی (`routing table`) قرار بگیرد.

Feasible Distance: مقداری است که از متغیرهای bandwidth و... می آید و درون خود مسیر یاب ذخیره می شود.

Reported Distance: مقدار FD است که مسیر یاب قبلی برای این مسیر یاب ارسال نموده است.

## ۸-۱ Change Metric

برای اینکه پروتکل را مجبور نماییم از یک مسیر دیگری route کند می توان Metric را تغییر داد. دو راه برای این کار وجود دارد:

۱- استفاده از bandwidth یا delay: که بهتر است delay را تغییر داد چون با تغییر bandwidth

پارامترهای دیگر مانند QoS تغییر می کند. این تغییر مقدار FD را تغییر می دهد نه RD.

۲- استفاده از Offset: این روش نیاز به ACL دارد و هم مقدار FD را تغییر می دهد هم RD. دستور آن به این شکل است:

```
(Config)#access-list <name> <permit | deny> <IP> <wildcard>
```

```
(Config-router)#offset-list <name access list> <in | out> <value>
```

باید توجه داشت که Administrative Distance (AD) قبل از Metric در اولویت است.

Successor: به مسیری گفته می شود که به عنوان بهترین مسیر برای ارسال بسته به مقصد مشخص انتخاب می شود. این مسیر با توجه به الگوریتم EIGRP بدست می آید. بدین صورت محاسبه می شود که مسیر یاب یک نقشه از همسایه ها و اینکه کدام همسایه چه شبکه ای را پشتیبانی می کند با استفاده از بسته های مسیریابی ایجاد می کند. پس از ایجاد این شبکه هزینه مسیرها محاسبه می شود و جدول مسیریابی کامل می شود.

Feasible Successor: به عنوان مسیر جایگزین (Backup) انتخاب می شود. برای انتخاب آن اول AD مورد بررسی قرار می گیرد. سپس RD در نظر گرفته می شود. اگر مقدار FD و RD با هم برابر باشد هر دو مسیر به عنوان اصلی انتخاب می شود و اگر  $RD < FD$  بود حلقه ایجاد نمی شود و به عنوان مسیر جایگزین انتخاب می شود. اگر  $FD > RD$  باشد ممکن است حلقه باشد یا نباشد و باید مورد بررسی قرار گیرد و به هر حال به عنوان مسیر جایگزین نگهداری می شود.

## ۹-۱ - Convergence EIGRP

همگرایی در این پروتکل به معنای ایجاد جدول مسیریابی است. همگرایی در EIGRP Topology دو حالت Passive و Active دارد که passive به معنای این است که جست‌وجو تمام شده است و جدول نهایی به روز شده است. حالت active بدین معنا است که در حال تکمیل جدول است و یا ارتباط قطع شده است و در حالت جست‌وجو برای مسیر جدید است.

اگر ارتباط با successor قطع شود و feasible successor هم نداشته باشد آن‌گاه برای یافتن مسیر جدید Query به همسایگان ارسال می‌شود. اگر همسایگان در مدت زمان معینی که پیش‌فرض ۳ دقیقه است پاسخ دهند بهترین مسیر انتخاب و جایگزین می‌شود. اگر بعد از این مدت زمان جوابی دریافت نشود، این مسیر از جدول مسیریابی حذف و ارتباط reset می‌شود. همسایگان نیز مانند مسیریاب اصلی به همسایگان پایین دستی خود پیغام query ارسال می‌کنند و صبر می‌کنند همه جواب‌ها یا replyها دریافت شود و سپس به مسیریاب اصلی reply می‌دهند. در این مدت زمان ۳ دقیقه مدت زمان ۹۰ ثانیه که فرا رسد SIAQuery ارسال می‌شود که به همسایگان این مطلب را مجدداً یادآوری کند که به این حالت Stuck-in-Active گویند. اگر همسایگان جواب دهند که هنوز منتظرند ۳ دقیقه از نو شروع می‌شود در غیر این صورت بعد از ۳ دقیقه مسیر حذف می‌شود. این روند برای همسایگان پایین دستی نیز هست.

دستورات برای تغییر زمان‌های پیش‌فرض بدین شکل است:

```
(Congi-route)#timers active time <time>
```

```
(Config-if)#ip hello-interval EIGRP <as> <time>
```

```
(Config-if)#ip hold-time EIGRP <as> <time>
```

راه‌هایی که برای محدود کردن ارسال بسته‌های EIGRP است در زیر تشریح داده شده است:

۱- Stub router: این تنظیم برای زمانی که مسیریاب خارج از شبکه است مثلاً برای شعبه که سرعت پایین است استفاده می‌شود. یا برای edge استفاده می‌شود. در این حالت این مسیریاب مسیرهای EIGRP را یاد نمی‌گیرد و به دیگر مسیریاب‌ها نمی‌فرستد و همسایگان زمانی که بفهمند این مسیریاب Stub است دیگر Query به این مسیریاب ارسال نمی‌کنند. به طور پیش‌فرض Stub فقط connected و summary را ارسال می‌کند و حالت‌های دیگر مانند static، redistributed، receive-only باید تنظیم شود.

۲- Summarization: بسته‌هایی که classless هستند را ارسال نمی‌نماید.

## ۱-۱۰-۱ EIGRP Variance

در جدول مسیریابی ریدیف‌هایی که Cost/Metric یکسانی دارند، به عبارت دیگر برای رسیدن به یک مقصد مشخص چندین مسیر با اولویت برابر وجود دارد، می‌توان تعداد آنها را با  $\text{maximum-path} \langle \# \rangle$  معین نمود. پروتکل‌های EIGRP، OSPF و RIP با استفاده از Cisco Express Forwarding (CEF) می‌تواند بین آنها Load-Balancing انجام دهد. این Load-Balancing بر مبنای Flow هست نه بر مبنای بسته و این بدان معناست که در هر ارتباط از یک مسیر استفاده می‌شود و بسته‌های مربوط به این ارتباط فقط از این مسیر عبور می‌کنند. در صورت قطع ارتباط و ایجاد ارتباط جدید ممکن است مسیر دیگر استفاده شود. قابل توجه است که تصمیم‌گیری برای ارسال یک جریان صورت می‌گیرد و نه برای هر بسته از یک جریان. در پروتکل‌های RIP و OSPF اگر دو ریدیف با Cost/Metric برابر نباشند از Load-Balancing استفاده نمی‌کنند، ولی پروتکل EIGRP با استفاده از متغیری به نام Variance از ریدیف‌های متفاوت برای Load-Balancing استفاده می‌کند. در واقع Feasible Successor به عنوان Load-Balancer استفاده می‌شود. مقدارهای variance به صورت زیر است:

۱-  $\text{Variance} = 1$ : مقدار پیش‌فرض است و بدان معناست که فقط از ریدیف‌های مساوی در Load-Balancing استفاده شود.

۲-  $128 < \text{Variance} < 1$ : اگر  $(\text{Variance} * \text{FD}) > \text{RD}$  باشد و مقدار  $\text{FD} > \text{RD}$  باشد، به عبارتی حلقه هم نباشد آن خط را در Load-Balancing شرکت می‌دهد.

## ۱-۱۱-۱ EIGRP Access Control List

### Access Control List ۱-۱۱-۱

ابتدا مفهوم Access Control List (ACL) و سپس کاربرد آن در مسیریابی شرح داده می‌شود. ACL یا لیست کنترل دسترسی در واقع فیلتری است که تعریف می‌شود کدام بسته‌های به روزرسانی مسیریابی و یا هر بسته دیگر اجازه عبور یا عدم عبور به داخل و یا خارج شبکه دارد. باید توجه داشت که ACL به تنهایی کاربردی ندارد و با ترکیب دستور یا دستورات دیگر کامل می‌شود. همچنین ACL با ترکیب wildcard mask کار می‌کند و نام آن می‌تواند عدد یا عبارت متنی باشد. ACL شامل دو نوع Standard و Extended است. ACL می‌تواند شامل خط‌هایی از دستورات باشد که به هر خط Access Control Entries (ACE) گفته می‌شود. به صورت پیش‌فرض آخر هر ACL دستور deny any نهفته است ولی قابل مشاهده نیست. به این نکته نیز باید توجه داشت که اولین تطابق در لیست ACL باعث می‌شود که ریدیف‌های بعد از آن بررسی نشود.

Wildcard mask -> subnet

0.0.3.255 -> 255.255.252.0

در واقع با استفاده از ACL می‌توان مشخص نمود که کدام بسته به روزرسانی وارد نقشه مسیریابی شود و یا اینکه کدام بسته به روزرسانی به همسایه‌ها ارسال شود. با استفاده از دستور زیر می‌توان این کار را انجام داد:

```
(Config)#access-list <name or number> <permit | deny> <address> <wildcard> standard
```

```
(Config)#access-list <name or number> <permit | deny> <address> <wildcard> extended
```

برای استفاده نمودن در مسیریابی از دستور زیر استفاده می‌شود:

```
(Config-route)#distributed-list <ACL number> <in | out>
```

### Prefix List -۲-۱۱-۱

روش دیگر برای فیلتر کردن مسیریابی استفاده از Prefix-list است. این روش قابلیت‌های بیشتری نسبت به ACL دارد مانند اینکه نیازی نیست IP حتماً تطابق داشته باشد و می‌توان بازه کمتر و یا بیشتر در آن قرار داد، ولی همانند ACL مقادیر permit یا deny دارد و همانند آن برای نام‌گذاری از عبارت متنی یا عدد استفاده می‌شود. برای تنظیم آن از دستور زیر می‌توان استفاده نمود:

```
(Config)#ip prefix-list <list-name> <sequence sequence-number> <deny | permit> <prefix-IP/prefix-length> <ge ge-value> <le le-value> <eq eq-value>
```

در این تنظیم <sequence-number> 1 است. همچنین در این مورد خاص هم IP و هم length باید با ردیف مربوطه یکسان باشد. ge مخفف greater than or equal و le مخفف less than or equal و eq مخفف equal برای مقدار prefix length است.

نکته‌ای که باید توجه داشت این است که برای تعریف default route یا مسیر پیش فرض از عبارت 0.0.0.0/0 و برای همه مسیرها از عبارت 0.0.0.0 le 32 استفاده می‌شود.

همانطور که گفته شد prefix list نیز مانند ACL به تنهایی کاربرد ندارد. برای استفاده کردن در router EIGRP از دستور زیر استفاده می‌شود:

```
(Config-route)#distributed-list prefix <name of prefix-list> <in | out> <interface>
```

این روش بیشتر برای BGP استفاده می‌شود. کلاً مزیت آن نسبت به ACL استفاده از prefix length و شماره خط یا sequence number است.

### Route-map -۳-۱۱-۱

روش دیگر برای فیلتر کردن استفاده از route-map است. این روش پیچیده‌تر از روش‌های قبل است و مانند برنامه نویسی if/then/else عمل می‌کند. هر route-map می‌تواند یک یا چند خط یا همان sequence number داشته باشد. دستور آن به شکل زیر شرح داده شده است:

```
(Config)#route-map <name> <permit | deny> <sequence number>
```

به صورت پیش فرض برای اولین بار اگر خطی تعریف شود و sequence number خالی باشد عدد ۱۰ قرار می‌گیرد.

```
(Config-route-map)#match <match method>
```

اگر از match استفاده نشود عبارت به معنای any برداشت می‌شود.

اگر خط مورد نظر با خط sequence number پایین تر تطابق نداشته باشد، sequence number بعدی بررسی می‌شود.

Route-map مانند prefix و ACL نیز از عبارت‌های permit یا deny استفاده می‌کند. تفاوت مهمی که در داخل route-map قابل شهود است این است که می‌توان از ACL یا prefix داخل آن استفاده نمود که در ادامه شرح داده شده است:

```
ACL | prefix = permit => route-map= run [permit | deny]
```

```
ACL | prefix = deny => route-map= doesn't run and go to next route-map sequence
```

```
(Config-route)#distributed-list route-map <name of route-map> <in | out> <interface name>
```

همانطور که قابل فهم است، عبارت permit در route-map به معنای اجرای دستور اجازه عبور و عدم عبور است و معنای deny به مفهوم عدم اجرای دستور اجازه عبور و عدم عبور است.

### EIGRP Route Summarization -۱۲-۱

Route summarization یا خلاصه کردن مسیریابی معمولاً برای شبکه‌های بزرگ استفاده می‌شود که باعث کاهش اندازه شبکه، کاهش جدول مسیریابی و کاهش محدوده فرستادن query می‌شود و برای هر مسیریاب به طور جداگانه می‌توان استفاده نمود.

فرض کنید شبکه‌ای با آدرس شبکه‌های زیر وجود دارد. در این صورت می‌توان به جای این آدرس‌ها تنها از یک آدرس استفاده نمود.

192.168.0.0/24    192.168.1.0/24    الی    192.168.31.0/24

در آدرس‌های بالا به دلیل اینکه ۳ بیت ابتدایی قسمت سوم از سمت چپ ثابت است و بقیه متغیر می‌توان آن را به صورت زیر خلاصه نمود. در واقع \* ثابت و # متغیر است.

192.168.\*\*\*#####.0/19

مسیرها را به دو صورت manual و auto می‌توان تنظیم نمود که پس از آن همسایه‌ها down و سپس up می‌شوند.

دستور manual:

```
(Config-if)#ip summary-address EIGRP asn <ip prefix> <mask>
```

دستور auto:

```
(Config-route)#auto-summary
```

موقعی که summary فعال باشد در جدول مسیریابی ردیفی با همان ip summary اضافه می‌شود که با مقصد null0 برای جلوگیری از حلقه نمایش داده می‌شود همچنین AD=5 قرار می‌دهد. مشکلی که در اینجا رخ خواهد داد این است که اگر مسیری مانند پروتکل eBGP با AD=20 وارد مسیریاب شود مشکل رخ خواهد داد. زیرا AD برای EIGRP Internal برابر 90 است و این مقدار بیشتر از 20 است، در صورتیکه برای خلاصه مسیر AD=5 در نظر گرفته می‌شود و مسیر EIGRP Summarization را جایگزین می‌نماید. نکته‌های قابل توجه دیگر این است که مسیریاب فقط IP‌های local خود را خلاصه می‌کند و بقیه را یاد می‌گیرد ولی خلاصه نمی‌کند. این امر فقط برای دستور network خلاصه‌سازی می‌شود و برای redistributed این دستور کارایی ندارد. رفتار خلاصه‌سازی برای class full routing است و برای شبکه‌های discontinuous کارایی ندارد. تفاوت شبکه‌های contiguous و discontinuous در این است که شبکه contiguous باید آدرس‌های IP در همان کلاس IP و به صورت پشت سر هم باشد ولی در discontinuous نیازی نیست.

### ۱-۱۳- EIGRP Default Route

Default route یا مسیر پیش فرض زمانی استفاده می‌شود که مسیری با هیچ یک از ردیف‌های جدول مسیریابی تطابق نداشته باشد. همچنین اگر در مسیریابی مسیر پیش فرض تعریف شده باشد و بخواهد آن را به دیگر مسیریاب‌ها با پروتکل EIGRP ارسال نماید می‌توان از آن استفاده نمود. برای تعریف مسیر پیش فرض به سه روش قابل انجام است:

۱- static default route with EIGRP: استفاده از یک مسیر پیش فرض دستی که دستور آن به شکل زیر است:

```
(Config)#ip route 0.0.0.0 0.0.0.0 <next-hop-ip> <out-going-int>  
(Config-route)#redistribute static metric [bandwidth] [delay] [reliability] [load]  
[MTU]
```

در دستور بالا به دلیل استفاده از metric بهتر است ولی می توان از دستور زیر نیز استفاده نمود.

```
(Config-route)#network 0.0.0.0
```

۲- Default network: اغلب از این روش و برای شبکه classful استفاده می شود. به بقیه مسیریاب ها این مسیر را انتقال می دهد. این روش ابتدا تشخیص می دهد که شبکه classful است و سپس مسیر را ارسال می کند. دستور آن به شرح زیر است:

```
(Config)#ip default-network <net-id>  
(Config-route)#network <net-id>
```

۳- Summary-address: این روش سه پیش فرض دارد ۱- همسایگان ریست می شوند ۲- فقط مسیر پیش فرض پخش می شود ۳- به صورت پیش فرض آدرس خلاصه شده به صورت محلی نصب می شود. دستور آن به شرح زیر است:

```
(Config-if)#ip summary-address EIGRP <asn> 0.0.0.0 0.0.0.0
```

با دستور show ip route می توان مسیر پیش فرض را مشاهده نمود.

## ۱-۱۴- EIGRP Manipulation

با توجه به سیاست های مدیریتی ممکن است مدیر شبکه بخواهد مسیرهای پروتکل EIGRP را به صورت دستی تغییر دهد. سه روش برای دستکاری مسیرها در EIGRP وجود دارد. ۱- change Metric: مانند تغییر delay و Bandwidth ۲- route filtering: مانند فیلتر کردن قسمتی از مسیرها که مسیریاب مسیر دیگری را انتخاب کند. مشکلی که وجود دارد این است که مسیر Backup از بین می رود. ۳- route summarization: خلاصه کردن یک مسیر که مسیریاب مسیری را انتخاب می کند که بیشترین تطابق را داشته باشد که در این روش backup وجود دارد.

## ۱-۱۵ EIGRP Named Mode

حالتی از تنظیم مسیریاب است که همه دستورات در این حالت زده می‌شود. به عبارتی نیازی نیست که به قسمت تنظیمات مربوط به interface level وارد شد. البته فقط برای IPV4 کاربرد دارد. در کل تنظیمات EIGRP دو حالت Classic و Named دارد. در حالت Named یک نام Process انتخاب می‌شود و ASN با دستور address-family وارد می‌شود. به طور مثال همه دستورات را می‌توان به صورت زیر در حالت Named وارد نمود:

```
(Config)#router EIGRP Cisco
```

```
(Config-router)#address-family ipv4 autonomic-system 123
```

```
(Config-router-af)#network 10.10.10.0 0.0.0.255
```

```
(Config-router-af)#af-interface <interface-id>
```

```
(Config-router-af-interface)#summary-address 10.10.0.0/16
```

```
(Config-router-af)#topology base
```

```
(Config-router-af-topology)#no auto-summary
```

حالت topology حالت جدیدی است که می‌توان دستورات مربوط به EIGRP topology table را وارد نمود.

## ۲- OSPF (Open Shortest Path First)

قبل از اینکه به شرح این پروتکل پرداخته شود تفاوت OSPF با EIGRP بررسی می‌شود.

۱- پروتکل OSPF از نوع Link State (LS) و EIGRP از نوع Distance Vector (DV) است. در DV از

الگوریتم Bellman Ford و لی در LS از الگوریتم Dijkstra استفاده می‌شود. باید توجه داشت که EIGRP در واقع نوع پیشرفته‌ای از DV است و ممکن است بعضی از عیب‌های آن را نداشته باشد.

۲- در DV اطلاعات مسیریاب قبلی کمتر ارسال می‌شود در RIP فقط مسیریاب بعدی را می‌بیند. EIGRP مقادیر Kها را ارسال می‌نمود. ولی OSPF مقادیری مانند P2P بودن، broadcast بودن و ... را ارسال می‌نماید.

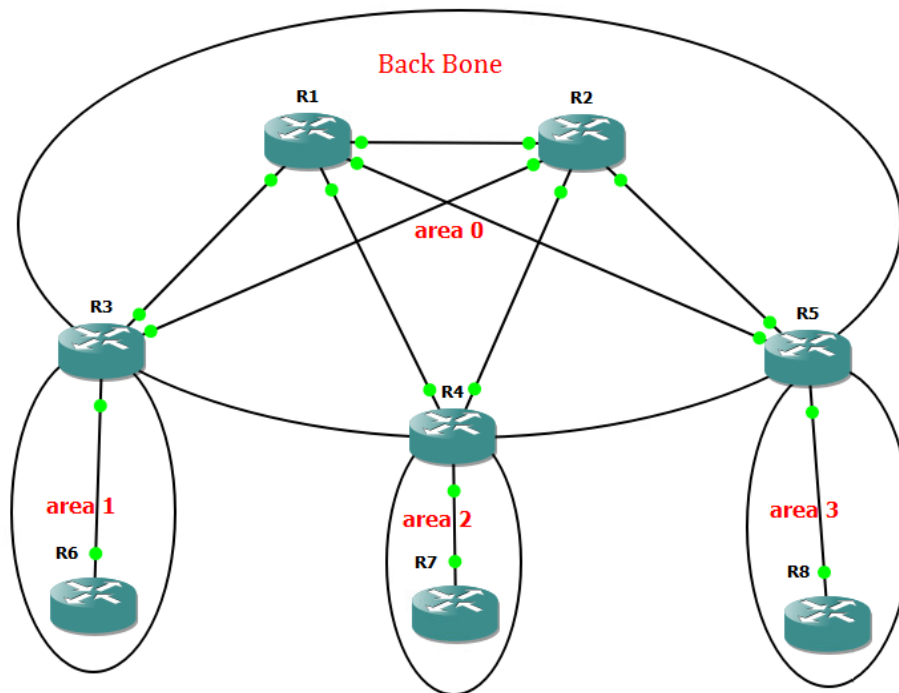
۳- در DV اطلاعات مسیریاب‌های قبلی ارسال نمی‌شوند. در EIGRP فقط شبکه همسایه‌ها را بدست می‌آورد ولی در OSPF اطلاعات همسایگان همسایه‌ها نیز ارسال می‌شود و هر مسیریاب یک تصویری از کل شبکه دارد.

۴- در DV اطمینان مسیرها لحاظ نمی‌گردد ولی OSPF هر نیم ساعت مسیریاب اصلی بسته‌ای برای زنده ماندن مسیر ارسال می‌کند. زمان پاسخ‌گویی حداکثر یک ساعت است و بعد از آن اگر پاسخی دریافت نشود مسیر مورد نظر از کل مسیریاب‌ها حذف می‌گردد.

مرحله تشکیل جدول مسیریابی OSPF را می توان به سه قسمت تقسیم کرد:

- ۱- پوشش همسایه ها و ارتباط بین آنها (Neighbor Discovery)
- ۲- ساختار تغییر پایگاه داده (Topology Database Exchange)
- ۳- محاسبه بهترین مسیر (Route Computation)

OSPF مانند EIGRP از hello packet استفاده می کند. این بسته سلام از طریق آدرس IP Multicast 224.0.0.5 ارسال می شود و از استاندارد IP Protocol 89 استفاده می نماید. نقشه OSPF یک ساختار سلسله مراتبی دارد که در بالای شبکه Back Bone ، در وسط مسیریاب های ABR یا Area Border Router و در نهایت مسیریاب های ناحیه بندی شده قرار دارند. پیشنهاد شرکت سیسکو این است که حداکثر ۴۰ مسیریاب در یک



ناحیه قرار گرفته شود.

ناحیه یا Area نیز شبیه ASN در EIGRP برای جدا کردن شبکه است با این تفاوت که در area شبکه سلسله مراتبی است ولی ASN شبکه کاملاً تخت است. همچنین area مربوط به لینک یا اینترفیس است. فقط مسیریاب هایی که در یک area باشند همدیگر را می بینند. البته مسیریاب های ABR از این قانون مستثنی هستند. در واقع وظیفه این مسیریاب ها ایجاد ارتباط بین ناحیه پایین و شبکه Back Bone است. اگر ساختار OSPF به صورت تخت در نظر گرفته شود که همه در یک ناحیه قرار گرفته باشند به دلیل اینکه هر مسیریاب نقشه کل شبکه را دارد CPU مسیریاب ها

بالا می‌رود و دچار مشکل می‌شود. نکته‌ای که باید به آن توجه کرد این است که در OSPF فیلتر کردن و یا summarization وجود ندارد (البته در ادامه کمی تشریح داده می‌شود).

دستورات برای تنظیم OSPF:

به صورت محلی است (Config)#router OSPF <process-id>

(Config-router)#network <net-id> <wildcardmask> area <area number>

(R1)#sho ip OSPF database

(R1)#sho ip OSPF neighbor

OSPF برای فعال شدن نیاز به router id دارد که انتخاب آن به سه روش زیر امکان پذیر است:

۱- وارد کردن دستی

۲- انتخاب loopback interface فعال با بالاترین شماره IP

۳- انتخاب loopback interface غیر فعال با بالاترین شماره IP

در صورت تغییر router-id بعد از تنظیم OSPF باید دستور clear ip OSPF process وارد شود تا دوباره شروع به بررسی گردد.

## OSPF Timers ۱-۲

OSPF بیشتر از EIGRP پیچیده است. در EIGRP ۳ نوع بسته اصلی ۱- Update ۲- Hello ۳- Query وجود دارد و ۲ بسته دیگر پاسخ به آنها است. ولی در OSPF ۶ نوع بسته وجود دارد. در EIGRP پارامترهایی همچون K-value، ASN و آدرس شبکه در مسیریاب‌ها باید یکی باشد ولی در OSPF به شرح زیر است:

۱- Hello Interval

۲- Dead Interval

۳- Area ID

۴- Net ID

۵- Stub Area Flag

۶- Authentication

برای اینکه OSPF بتواند بسته چندپخش ارسال کند نیاز به فعال شدن OSPF (network یا ip OSPF area) دارد. همچنین لینک مربوطه نباید در حالت غیر فعال باشد (not passive-interface).

• در بسته‌های Hello پارامترهای زیر وجود دارد:

- OSPF router ID
- List of neighbors reachable on the interface
- router priority
- DR ip address(designated route)
- BDP ip address

زمان پیش‌فرض برای Hello Interval=10 و Dead Interval = 4\*Hello ثانیه است که به صورت خودکار در ۴ ضرب می‌شود. Hello Interval می‌تواند کمتر از ثانیه باشد همچنین از مکانیزم Bi-Directional Forwarding Detection نیز می‌توان استفاده نمود که به صورت خودکار با توجه به پارامترهای شبکه بهترین زمان ممکن را انتخاب و بسته را ارسال می‌نماید. دستورات برای تغییر زمان به شرح زیر است:

```
(Config-if)#ip OSPF hello-interval <value>
(Config-if)#ip OSPF dead-interval <value>
(Config-if)#ip OSPF dead-interval minimal hello-multiplier <value>
(R1)#Show ip OSPF interface <interface>
```

## ۲-۲-۲ OSPF Router-ID

Router-id در پروتکل OSPF خیلی بیشتر از پروتکل EIGRP حائز اهمیت است. در نبود router-id در پروتکل EIGRP مشکلی در مسیرهای داخلی یا internal route ایجاد نمی‌شود و همسایه‌ها به درستی با یکدیگر ارتباط برقرار می‌کنند، ولی در external route مشکل ساز می‌شود. حال نقش Router-ID در پروتکل OSPF به اختصار توضیح داده می‌شود:

- ۱- اگر در یک ناحیه router-id مسیریاب‌ها با هم برابر باشد در ارتباط همسایگی دچار مشکل می‌شوند و پیام duplicate router-id داده می‌شود.
- ۲- اگر ناحیه‌ها متفاوت باشند برابری router-id تشخیص داده شده و مسیریاب‌ها LSAها را متمایز می‌کنند و پیام OSPF Flood War می‌دهند.
- ۳- در پروتکل OSPF نیز router-id مانند EIGRP تعریف می‌شود.

## ۲-۳-۲ OSPF MTU

MTU یا Maximum Transmission Unit مقدار حجم بسته‌ای است که مسیریاب‌ها به آن اجازه عبور می‌دهند. در شبکه‌ها معمولاً این مقدار به صورت پیش‌فرض 1500 Bytes است و اگر بیشتر باشد مسیریاب‌ها بنا به هدر Bit

Don't Fragment اگر 0 باشد بسته را تکه تکه یا به اصطلاح fragment و اگر ۱ باشد drop می کنند. در پروتکل OSPF اگر MTU ساینز تعریف شده در تنظیمات با هم برابر نباشند مسیریاب همسایه به حالت EXSTART می رود و سپس Down می شود و پیام Too Many Retransmission داده می شود.

## ۲-۴- OSPF Authentication

همانند پروتکل EIGRP این پروتکل نیز از قابلیت احراز اصالت استفاده می کند که با کمی تفاوت شرح داده می شود. به دو دلیل از OSPF authentication استفاده می شود: ۱- برای فیلتر کردن مسیر ۲- برای امنیت بالا. برای تنظیم و فعال کردن احراز اصالت سه مرحله وجود دارد: ۱- authentication enable فعال نمودن احراز اصالت. ۲- type of authentication تعیین نوع احراز اصالت. ۳- enable per interface فعال نمودن آن برای هر اینترفیس. پروتکل OSPF مانند پروتکل EIGRP انعطاف پذیر نیست زیرا: ۱- نمی توان برای پسوندها زمان انتخاب کرد و ۲- نمی توان از chain key استفاده نمود و فقط برای MD5 می توان چند پسوندد انتخاب نمود. در واقع سه مدل احراز هویت انجام می گیرد: ۱- بدون احراز هویت ۲- به صورت متن ساده یا plain text ۳- به صورت هشینگ از MD5. برای تنظیم آنها می توان هم در حالت interface و هم در حالت router این کار را انجام داد. دستورات آن به شرح زیر است:

```
(Config-if)#ip OSPF authentication -> plain text
```

```
(Config-if)#ip OSPF authentication message-digest -> MD5
```

```
(Config-if)#ip OSPF authentication-key <key-value> -> plain text
```

```
(Config-if)#ip OSPF message-digest-key <key-number> md5 <key-value> -> MD5
```

```
(Config-router)#area <area-number> authentication -> plain text
```

```
(Config-router)#area <area-number> authentication message-digest -> MD5
```

```
(R1)#show ip interface <interface-type> <interface-number>
```

```
(R1)#debug ip OSPF adj
```

```
(R1)#debug ip OSPF hello
```

## ۵-۲ Frame Relay Introduction

مکانیزمی است که در شبکه‌های WAN استفاده می‌شود و بجای اینکه بسته‌ها در لایه ۲ به صورت 802.3 بسته‌بندی یا encapsulate شود به صورت Frame Relay با بیت کمتر و به صورت DLCI یا Data Link Connection Identifier بسته‌بندی می‌شود. هنگام روشن شدن مسیر یاب‌ها، مبدأ یا SRC و مقصد یا DES را به یکدیگر ارسال می‌نمایند تا متوجه شوند که بسته از کدام DLCI وارد شده و آن را در جدول invers arp ذخیره کنند. اگر چند مسیر یاب در یک DLCI با آدرس‌های متفاوت وجود داشته باشد باید از SUB interface استفاده شود. به همین ترتیب بسته‌های broadcast, multicast وجود ندارد. در این صورت بسته‌ها باید تک تک به صورت unicast ارسال شود. همچنین این امر باید در مسیر یاب‌ها بر روی اینترفیس‌ها تنظیم شود. شبکه‌های Frame Relay به ۳ مدل تقسیم می‌شوند: ۱- full mesh - ۲ Hub & Spoke - ۳ Dual Hub & Spoke. باید توجه داشت که DLCI به صورت ۱۰ بیتی است و از ۱۶ تا ۱۰۰۷ قابل استفاده است. به عبارت دیگر Frame Relay از یک شناسه حلقه مجازی استفاده می‌کند تا فریم‌های شبکه را به یک مدار مجازی دائمی (PVC) permanent virtual circuit یا یک مدار مجازی موقت بر اساس تقاضا (SVC) switched virtual circuit اختصاص دهد.

## ۶-۲ Type of Networks in OSPF

در پروتکل OSPF اطلاعات لینک کمک می‌کند تا تشخیص داده شود بسته‌ها به صورت Unicast یا Multicast ارسال شوند. انواع شبکه‌ها به صورت زیر است.

| Network Type                      | DR/BDR? | Hello Unicast/Multicast? | Hello/Dead Intervals? |
|-----------------------------------|---------|--------------------------|-----------------------|
| broadcast                         | Yes     | Multicast                | 10/40                 |
| non-broadcast                     | Yes     | Unicast                  | 30/120                |
| point-to-point                    | No      | Multicast                | 10/40                 |
| point-to-multipoint               | No      | Multicast                | 30/120                |
| point-to-multipoint non-broadcast | No      | Unicast                  | 30/120                |
| Loopback                          | No      | N/A                      | N/A                   |

۱- Broadcast: به صورت پیش فرض بر روی شبکه LAN است. ارسال بسته‌ها به صورت همگانی قابلیت پیدا کردن همسایگان به صورت خودکار را می‌دهد. از DR (Designated Router) و BDR (Backup DR) پشتیبانی می‌کند. در واقع DR مسیریابی است که ترافیک شبکه را کاهش می‌دهد و تنها مسیریابی است که جدول پایگاه داده کامل OSPF را به دیگر مسیر یاب‌ها ارسال می‌کند و BDR جایگزین یا Backup آن

است. زمان Hello و Dead، 10 و 40 است. ابتدا یک بسته ارسال می‌شود که همه مسیر یاب‌ها می‌توانند آن را مشاهده نمایند. اگر شبکه Non-Broadcast تشخیص داده شود مدیر شبکه می‌تواند با دستور IP OSPF Network Broadcast به صورت دستی آن را تغییر می‌دهد.

۲- Non-Broadcast: در شبکه Frame Relay یا Multicast به صورت پیش فرض در این حالت است. پروتکل OSPF درباره داخل Frame Relay اطلاع ندارد ولی طرز عملکرد آن را می‌داند. اگر در Frame Relay، sub interface تعریف شود باید نوع بسته‌بندی یا encapsulation مشخص شود. در این حالت همسایگان به صورت خودکار پیدا نمی‌شوند و باید به صورت دستی در یک طرف آن وارد شود. زمان Hello و Dead، 30 و 120 است. دو روش برای تنظیم این مدل شبکه وجود دارد.

۱. روش اول:

```
(Conf)#interface serial 1/0
(Conf-if)#encapsulate frame-relay
(Conf-if)#ip address <ip-address>
(Conf-if)#ip OSPF <OSPF-proses-id> area <area-number>
```

۲. روش دوم:

```
(Conf)#interface serial 1/0.101 multipoint
(Conf-subif)#ip address <ip-address>
(Conf-subif)#ip OSPF <OSPF-proses-id> area <area-number>
(Conf-subif)#frame-relay map ip <ip-address> <DLCI>
(Config-router)#neighbor <neighbor-ip>
```

دستور زیر برای دیدن مشخصات Frame Relay است.

```
(R1)#sho frame-relay map
(R1)#sho frame-relay pvc
```

۳- Point-To-Point: در این حالت DR و BDR وجود ندارد. همسایه‌ها به صورت خودکار پیدا می‌شوند. زمان Hello و Dead، 10 و 40 است. مثال‌هایی از این حالت PPP و HDLC است.

```
(Config)#interface serial 1/0.102 point-to-point
(Config-subif)#ip address <ip-address>
```

```
(Config-subif)#ip OSPF <opsf-proses-id> area <area-id>
(Config-subif)#frame-relay interface-dlci <dlci-number>
(Config-fr-dlci)#exit
```

همچنین با دستور ip OSPF network point-to-point می توان این کار را انجام داد.

۴- Point-To-Multipoint: این حالت پیش فرض هیچ نوع شبکه ای نیست. معمولاً برای مکانهایی استفاده می شود که یک مسیریاب با چند مسیریاب دیگر به صورت P2P ارتباط دارد. در این حالت DR و BDR وجود ندارد. همسایه ها به صورت خود کار همدیگر را پیدا می کند. زمان Hello و Dead، 30 و 120 است.

```
(Config)#interface serial 1/0.103 multipoint
(Config-subif)#ip OSPF <opsf-proses-id> area <area-id>
(Config-subif)#ip OSPF network point-to-multipoint
(Config-subif)#frame-relay interface-dlci <dlci-number-router X>
(Config-subif)#frame-relay interface-dlci <dlci-number-router Y>
```

گاهی اوقات پیش می آید که شبکه از حالت Broadcast به حالت Point-To-Multipoint برده می شود به طور مثال زمانی که بعضی از مسیریاب ها شلوغ هستند می توان با تنظیم Cost اولویت آنها را مشخص نمود. در این حالت دیگر به وضعیت لینک نگاه نمی شود و تنها Cost در نظر گرفته می شود.

```
(Config-router)#neighbor <ip-address> cost <value>
```

۵- Point-To-Multipoint Non-Broadcast: همسایه ها به صورت خود کار همدیگر را پیدا نمی کند. باید به صورت دستی این شبکه تنظیم شود.

```
(Config-if)#ip OSPF network point-to-multipoint non-broadcast
```

## ۷-۲ OSPF LSA

اگر شبکه یک سازمان مجموعی از ناحیه ها در نظر گرفته شود آن گاه دو نوع پروتکل شبکه وجود دارد که به اختصار توضیح داده می شود:

۱- Interior Gateway Protocol (IGP): این نوع پروتکل فقط در یک autonomous system یا در یک ناحیه از شبکه کارایی دارد.

۲- Exterior Gateway Protocol (EGP): این نوع پروتکل ارتباط بین چند autonomous system یا چند ناحیه را برعهده دارد.

همه پروتکل‌های شبکه قبل از این که اطلاعات مسیرها را در جدول مسیریابی خود ذخیره کنند در جدولی به نام Data Base یا Topology ذخیره می‌کنند و پس از محاسبات آن را در جدول مسیریابی ذخیره می‌نمایند. که به طور خاص در پروتکل OSPF به آن Link State Data Base می‌گویند. زمانی که بر روی یک اینترفیس پروتکل OSPF فعال می‌شود این جدول به صورت خودکار ایجاد می‌گردد. باید توجه کرد که هر جدول مخصوص به یک AS است.

در این پروتکل هر مسیریاب ریشه درختی است که بر روی آن الگوریتم SPF یا کوتاه‌ترین مسیر اول پیمایش شود، اجرا می‌شود. LSA یا Link State Advertisement بسته‌هایی هستند که این جدول و پروتکل را کامل می‌کنند. انواع آن به شرح زیر است:

| Types of LSA in OSPF |                         |
|----------------------|-------------------------|
| Type 1               | Router LSA              |
| Type 2               | Network LSA             |
| Type 3               | Network Summary LSA     |
| Type 4               | ASBR Summary LSA        |
| Type 5               | AS External LSA         |
| Type 6               | Group Membership LSA    |
| Type 7               | NSSA External LSA       |
| Type 8               | External Attributes LSA |
| Type 9-11            | Opaque LSA              |

| LSA in CCNP |                     |
|-------------|---------------------|
| Type 1      | Router LSA          |
| Type 2      | Network LSA         |
| Type 3      | Network Summary LSA |
| Type 4      | ASBR Summary LSA    |
| Type 5      | AS External LSA     |
| Type 7      | NSSA External LSA   |

## ۲-۸- SPF LSA Type-1 (Router-LSA)

پروتکل OSPF با استفاده از بسته به روزرسانی یا Update اطلاعات همسایه‌های خود را جابه‌جا می‌کند. به این بسته‌ها Hello نیز گفته می‌شود. هر مسیریاب DB خود را به همسایه‌ها و آدرس‌های چند پخش 224.0.0.5 و (DR) 224.0.0.6 در ناحیه خود ارسال می‌کند. همچنین همسایه‌ها این بسته‌ها را به دیگر مسیریاب‌های موجود در آن ناحیه ارسال می‌کنند. به این بسته‌ها Router LSA گفته می‌شود. با استفاده از دستور زیر این بسته‌ها را می‌توان مشاهده نمود:

(R1)#show ip OSPF database [router (advertise router)]

در قسمت router link states تمامی بسته‌های نوع ۱ را می‌توان مشاهده کرد.

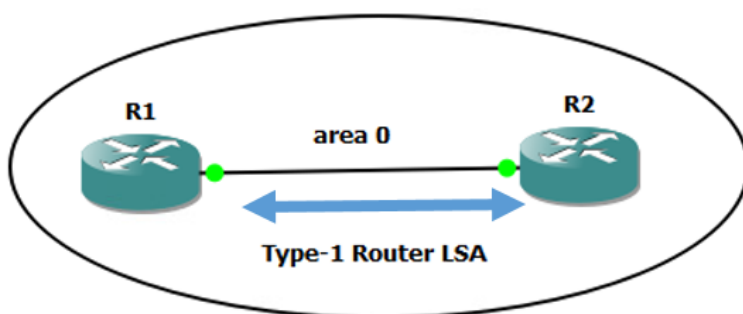
Link state id یک شناسه برای لینک و برابر router-id OSPF است.

Link state data عبارتی برای توضیح لینک است به طور مثال نوع ارتباط و Cost لینک را نشان می‌دهد.

OSPF Network Stub: اگر ایترفیسی از مسیریاب به جایی متصل باشد که هیچ مسیریابی بسته‌های به روزرسانی

و یا OSPF Hello را ارسال نکند، به آن ناحیه OSPF Network Stub گفته می‌شود.

لینک point-to-point یک stub network و link connected to another router است.



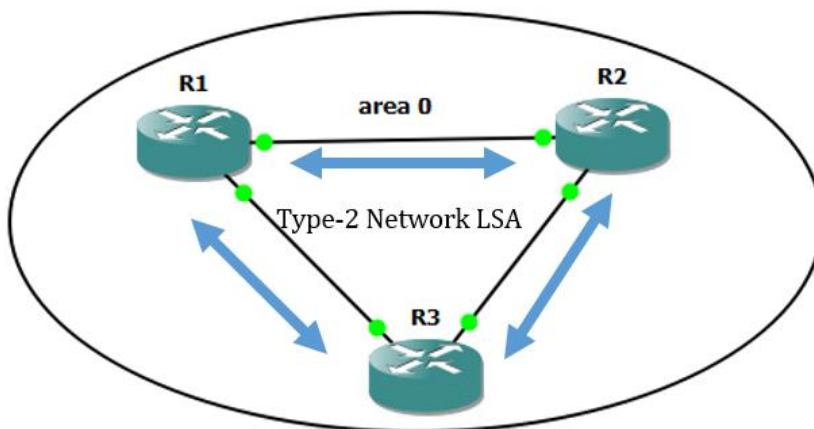
## ۹-۲ OSPF LSA Type-2 (Network-LSA)

این بسته‌ها فقط برای شبکه‌های Broadcast و Non-Broadcast استفاده می‌شود که DR/BDR دارند. تنها DR این بسته‌ها را ارسال می‌کند که حاوی NET-ID و SubnetMask است تا بقیه مسیریاب‌ها بتوانند SPF را ایجاد کنند.

نحوه انتخاب DR و BDR به یکی از روش‌های زیر است:

- اولین مسیریابی که روشن می‌شود خود را DR انتخاب کرده و به بقیه اطلاع می‌دهد.
- مسیریاب با بیشترین اولویت بین ۱ و ۲۵۵ که با دستور `ip OSPF priority <Value>` می‌توان مشخص کرد DR می‌شود.
- مسیریاب با بزرگترین router-id و priority که پیش‌فرض ۱ است به عنوان DR انتخاب می‌شود. اگر DR خاموش شد BDR جایگزین می‌شود و مسیریاب جایگزین برای BDR که الان DR شده است، انتخاب

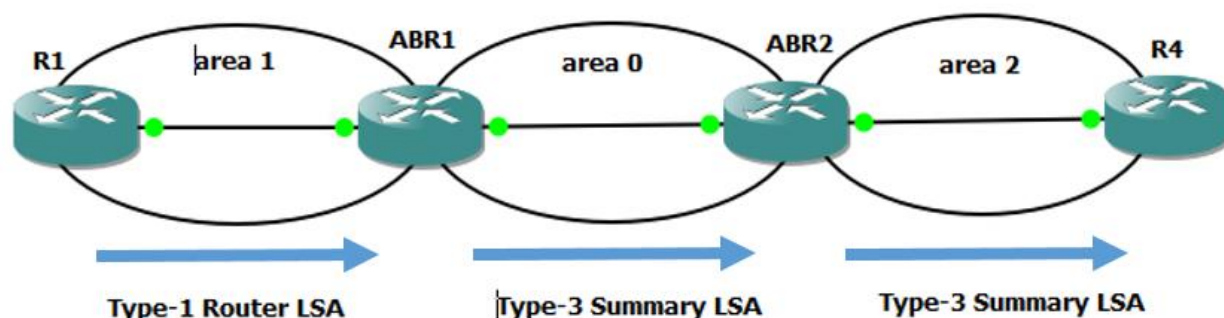
می شود. اگر Parority تغییر نمایند مسیر یاب DR تغییر نمی کند چون باید SPF تغییر کند که موجب سربار ترافیکی می شود. فقط در صورت تغییر اسم و یا Router-ID این عملیات انجام می شود.



## ۲-۱۰- OSPF LSA Type-3 (Summary-LSA)

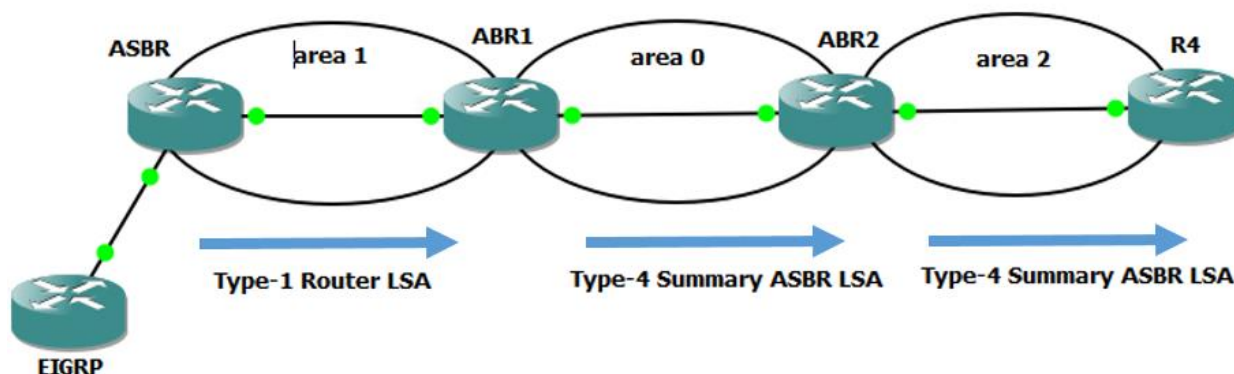
این نوع بسته ها توسط مسیر یاب های ABR فرستاده می شود تا اطلاعات مربوط (Net-ID، Mask و Cost from ABR) به دیگر ناحیه ها AS را جابه جا کنند. به دلیل سربار Memory و CPU و همچنین کاهش Link State Data Base (LSDB) ناحیه بندی صورت می گیرد. Type-1 و Type-2 فقط در یک ناحیه فرستاده می شود و بین ناحیه ها ارسال نمی گردد. هنگام ارسال بسته های نوع یک مشخص می شود که کدام مسیر یاب Area Border Router (ABR) است. مسیر یاب ها فقط می دانند چه Net-ID و Mask و همچنین کدام مسیر یاب ABR است. همچنین ABR مشخصه LSID را از شماره Net-ID بدست آورده و اضافه می کند. ABR شماره RID را اضافه می کند تا مشخص شود کدام مسیر یاب این را پخش کرده است چون ممکن است مسیریابی از طریق مسیر یاب های دیگر مانند این ارسال کرده باشند. با دستور `Show ip OSPF database <summary>` می توان این پیام ها را مشاهده نمود.

تعداد LSA ها را نیز می توان محدود کرد که با دستور `max-lsa <number>` قابل تنظیم است، ولی ابتدا پیغام warning/log می دهد، سپس منتظر می ماند و بعد از مدتی ارتباط با همسایگان را قطع می کند.



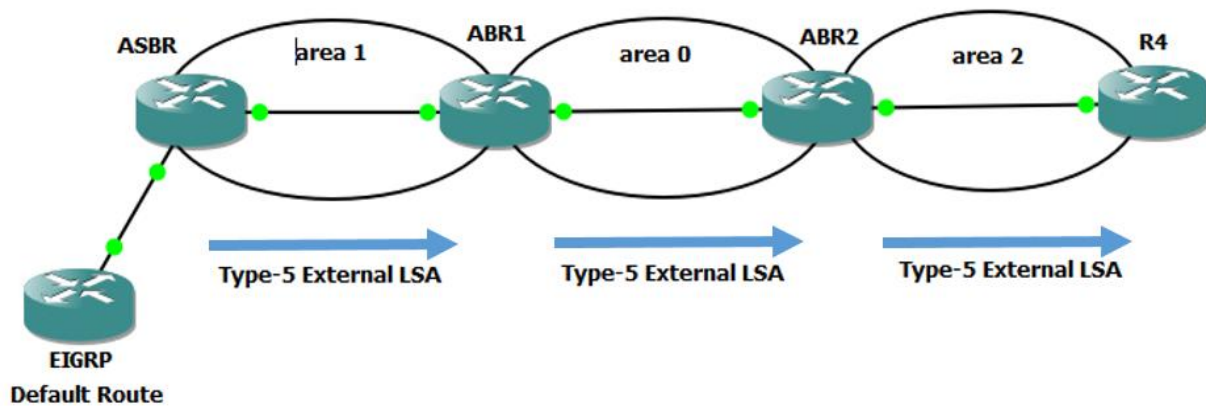
## ۲-۱۱- OSPF LSA Type-4 (ASBR -Summary-LSA)

زمانی که بسته های Type-1 از مسیریاب ASBR به دیگر ناحیه ها فرستاده می شود، مسیریاب هایی که این بسته را دریافت می کنند، آن را به Type-4 یا خلاصه بسته مسیریاب ASBR تبدیل می کنند. در واقع مسیریاب های ABR از این بسته استفاده می کنند تا ASBR را به دیگر مسیریاب ها اطلاع دهند. باید توجه داشت که مسیریاب ASBR خود از این بسته استفاده نمی کند. هنگامی که مسیریاب ASBR عمل redistribute را انجام می دهد. مسیریاب هایی که در ناحیه ASBR نیستند، نمی توانند ASBR را پیدا کنند، زیرا در درخت SPF وجود ندارد. پس مسیریاب های ABR بسته Type-4 LSA را ایجاد می کنند که به مسیریاب های داخلی اعلام کنند برای دسترسی به ASBR باید از طریق من بسته ارسال کنند.



## ۱۲-۲ OSPF LSA Type-5 (ASBR-External-LSA)

این بسته‌ها توسط مسیریاب ASBR ایجاد می‌شود تا مسیرهای خارجی وارد شده به پروتکل OSPF را رواج دهند. این بسته‌ها دو زیر مجموعه دارد که برای redistribution استفاده می‌شود. type-1 تمامی مسیریاب‌ها cost را در ترجمه می‌توانند تغییر دهند. Type-2 که به صورت پیش فرض است و فقط ASBR می‌تواند cost را تغییر دهد. در Cisco IOS مسیر پیش فرض یا Default Route از مسیر خارجی به OSPF مجاز نمی‌باشد.



## ۱۳-۲ OSPF Neighbor Relationship

هر مسیریابی باید LSDB یکسانی در یک ناحیه داشته باشد. مسیریاب‌های داخلی یک ناحیه فقط بسته‌های Type-1, 2 تشخیص می‌دهند و مسیریاب ABR نیز به غیر از Type-1 و Type-2 بسته Type-3 را نیز تشخیص می‌دهد.

بین مسیریاب‌ها بسته‌های زیر جابه‌جا می‌شود:

۱- Hello: این بسته از نوع یک و شامل موارد زیر است:

- Router ID
- Hello/dead interval
- Neighbors
- Area ID
- Authentication data
- Router priority
- DR IP address
- BDR IP address

۲- Database Description (DD or DBD): مانند نامتناسب بودن MTU که در حالت fully

adjacent یا همگرایی کامل قابل مشاهده نیست، ولی در حالت initiate می‌توان مشاهده نمود.

مسیریاب‌ها توسط این بسته از LSDB یکدیگر با خبر می‌شوند.

۳- Link-State Request: بسته‌های Link-State Request، Link-State Update و Link-State Acknowledgment به دلیل تغییر در LSDB انجام می‌گیرد.

۴- Link-State Update

۵- Link-State Acknowledgment

حالت‌هایی که یک همسایه در طی زمان Fully adjacent وجود دارد که بدون DR است به شرح زیر توضیح داده شده است:

- ۱- Down: در همسایه‌ها دیده نمی‌شود ولی در syslog پیغام می‌دهد.
- ۲- Attempt: در حالت Static Neighbor است. در این حالت بسته hello ارسال می‌شود، ولی جوابی دریافت نمی‌شود.
- ۳- Init: در این حالت بسته hello دریافت می‌شود، ولی مسیریاب خودش را در بین فیلد همسایگان نمی‌بیند. در واقع هنگامی که OSPF فعال می‌شود این اتفاق می‌افتد.
- ۴- 2way: در این حالت بسته hello دریافت می‌شود و مسیریاب خودش را در بین فیلد همسایگان می‌بیند. در این حالت این بسته برای انتخاب DR یا BDR و برای تشخیص یکدیگر استفاده می‌شود.
- ۵- ExStart: قبل از اینکه DBD ارسال شود مسیریاب‌ها به این حالت می‌روند و روند انتخاب Master/Slave پیش می‌آید که فقط DBD Header یا به عبارت بسته‌های توضیحی خالی پایگاه‌داده ارسال می‌شود که بدنه آن خالی است. OSPF از پروتکل TCP استفاده نمی‌کند چون در لایه پایین IP قرار دارد، ولی مکانیزمی شبیه TCP برای Sequence number و Acknowledgment دارد که Master مسئول ایجاد Seq num است. Master مسیریاب با Router-ID بزرگتر است. هر مسیریاب دوست دارد Master باشد و flag مربوط به master را یک قرار می‌دهد.
- ۶- Exchange: بعد از حالت ExStart مسیریاب‌ها با استفاده از IP Multicast شروع به DBD Exchange می‌کنند. در این حالت بسته‌های مربوط به توضیح پایگاه‌داده دارای بدنه و اطلاعات است.
- ۷- Loading: بعد از Exchange سه مرحله req->upd->ack اتفاق می‌افتد.
- ۸- Fully: این مرحله به معنای ثابت شدن همسایه‌ها است.

حال اگر DR/BDR در شبکه وجود داشته باشد، مشابه روند قبل مراحل طی می‌شود، با این تفاوت که مسیریاب‌ها LSDB خود را با DR یا BDR بررسی می‌کنند. فقط DR و BDR هستند که IP Multicast 224.0.0.6 دریافت

و بررسی می کنند. پروتکل OSPF هر ۳۰ دقیقه یکبار کل LSA را پخش می نماید و زمان LSA را صفر قرار می دهد. اگر نیازی به Flush باشد آن را ۳۶۰۰ قرار می دهد.

## ۲-۱۴- OSPF Metric Calculation

در این قسمت Metric و یا همان Cost در پروتکل OSPF بررسی می شود. این معیار براساس پهنای باند محاسبه می شود و می توان آن را به صورت دستی نیز تنظیم نمود. انتخاب بهترین مسیر به این شکل انجام می شود که ابتدا تمام LSA ها جمع آوری شده، سپس Shortest Path First (SPF) یا اول مسیر کوتاه تر، با کمترین Cost ممکن محاسبه می گردد. محاسبه Cost برای شبکه های Inter-area یا بین ناحیه و Intra-area یا داخل ناحیه متفاوت است. در محدوده Intra-area لیستی از cost ها وجود دارد که براساس آن درخت SPF ایجاد می شود. در محدوده Inter-area که ABR وجود دارد قضیه متفاوت است ولی به صورت کلی مانند Intra-area می توان محاسبه نمود. حتی اگر تمامی اطلاعات موجود باشد این امر از دیدگاه مسیریاب ها ساده نیست. مسیریاب های ABR با رسیدن Type-1,2 LSA درباره یک محدوده اطلاعات بدست می آورند و با استفاده از بسته Type-2 LSA مقدار Cost را به دیگر محدوده ها ارسال می کنند. همچنین اطلاع می دهند که برای رسیدن به مقصد موجود در این ناحیه باید از از من عبور کنید. اگر چندین مسیریاب ABR این آدرس مقصد را اطلاع دهند تمامی مقدارها جمع آوری شده و کمترین مقدار که همان بهترین مسیر است، محاسبه می شود و همه مسیریاب ها از آن استفاده می کنند.

مسیریاب ها ابتدا مسیرهای Intra-area سپس Inter-area و بعد از آن External را در اولویت برای ارسال بسته ها قرار می دهند. باید توجه داشت در حالت External و redistribution بسته های مربوط به شبکه نوع ۵ و ۷ ارسال و دریافت می گردد. Cost یا هزینه مسیر در پروتکل OSPF از فرمول reference-bandwidth/interface-bandwidth محاسبه می شود. در صورت نیاز سه راه برای تغییر Cost یا Metric وجود دارد:

۱- Change Reference Bandwidth: تغییر پهنای باند مبدأ که به صورت پیش فرض برابر 100Mbps است. برای تغییر آن از دستور `auto-cost reference-bandwidth <BW>` استفاده می شود و اگر بالای 100Mbps باشد cost برابر ۱ در نظر گرفته می شود. این خصوصیت باید برای هر اینترفیس به صورت جداگانه و دستی تنظیم شود.

۲- Setting Bandwidth: تنظیم دستی پهنای باند اینترفیس یکی دیگر از روش های تغییر Cost است. این روش توصیه نمی شود چون تنظیمات دیگر که براساس پهنای باند بدست می آید را تحت تاثیر قرار می دهد.

۳- Configure Cost Directly: تنظیم دستی cost به صورت مستقیم بهترین روش است، ولی از لحاظ مدیریتی امری پیچیده است. این تنظیم با دستور `ip OSPF cost <value>` تنظیم می شود.

## ۲-۱۵ OSPF Filtering

برای فیلتر نمودن یک مسیر مشخص صورت می گیرد. این امر فقط در بین ناحیه ها انجام پذیر است. قبل از آن که فیلتر مسیر تنظیم شود، باید تمام مسیر یاب ها LSDB یکسانی داشته باشند تا SPF به درستی صورت گیرد. عمل فیلترینگ توسط مسیر یاب ABR انجام می شود. مسیر یاب ABR مسیر را در لیست خود ذخیره می کند و اگر این مسیر فیلتر شود آن را ارسال نمی کند، بدین معنی که با استفاده از بسته های LSA Type-1,2 مسیرها را دریافت و ذخیره می کند، ولی در هنگام ایجاد بسته های LSA Type-3 آنها را که فیلتر شده اند ارسال نمی کند. هدف از فیلترینگ در پروتکل OSPF این است که: بسته ها از مسیر یاب خاص ارسال شوند و یا اینکه مسیر به هیچ کدام از مسیر یاب ها ارسال نشود. دستور تنظیم فیلتر نمودن به شرح زیر است:

```
(Config-router)#area <area-number> filter-list prefix <prefix-name> <in | out>
```

در اینجا in فرستادن به ناحیه و out دریافت از ناحیه است. باید توجه داشت ابتدا prefix-list باید تعریف شده باشد.

روش دیگر استفاده از area-range است که بیشتر برای summarization استفاده می شود.

```
(Config-router)#area <area-number> range <IP> <Mask> not-advertise
```

اگر area وارد شده مجاورت ABR نباشد هیچ تاثیری ندارد و فیلتر انجام نمی شود.

حال اگر در یک ناحیه خواسته شود فیلتر صورت گیرد، به این معنی است که بسته LSA دریافت می شود ولی در جدول routing قرار نمی گیرد. در واقع فیلتر بین OSPF DB و Routing Table قرار می گیرد. در OSPF DB تمام مسیر یاب ها، همه LSA ها از همه مسیر یاب ها را دریافت می کنند، ولی برای قرار دادن بسته در مسیر مشخص با فیلتر کردن دچار مشکل می شوند. این امر مانند EIGRP با دستور distribute-list انجام می گیرد. به عبارت ساده تر داخل یک ناحیه، به دلیل اینکه باید LSDB تمام مسیر یاب ها یکسان باشد، مسیر باید دریافت و در پایگاه داده ذخیره شود ولی در جدول مسیر یابی قرار نگیرد. به همین دلیل از دستور زیر استفاده می شود. باید توجه داشت که تنظیم این دستور در مسیر یاب باعث می شود که مسیر به هیچ عنوان از این مسیر یاب عبور نمی کند حتی اگر در جدول مسیر یاب های بعد از آن در جدول مسیر یابی وجود داشته باشد. پس باید این دستور در مسیر یاب صحیح صورت گیرد.

(Config-router)#distribute-list prefix <prefix-name> <in | out> <interface>

در فیلتر کردن مسیریاب بسته را دریافت و اطلاعات آن را یاد می‌گیرد ولی استفاده نمی‌کند. اگر بخواهیم به اطلاعات بسته هم توجه نکنند باید از دستور IN استفاده نمود، زیرا Distribute-list بین LSADB و Routing-table قرار دارد و فیلتر می‌تواند advertise یا انتقال اطلاعات مسیر را برگرداند.

## ۱۶-۲ OSPF Route Summarization

در ابتدا باید توجه داشت که هیچ‌کدام از مسیریاب‌هایی که در یک ناحیه هستند نمی‌توانند خلاصه‌سازی انجام دهند، زیرا در یک ناحیه تمامی LSDBها در تمامی مسیریاب‌ها یکسان است. دلایل خلاصه‌سازی در پروتکل OSPF می‌تواند برای کاهش LSDB و کاهش بار Memory و CPU یا دستکاری مسیر باشد. خلاصه‌سازی از دو طریق انجام می‌شود:

۱- Area Border Router (ABR): مسیریاب‌های ABR بسته‌های 4, 3-Type ارسال می‌کنند. مانند پروتکل EIGRP خلاصه‌سازی صورت می‌گیرد با این تفاوت که فقط در مسیریاب‌های ABR انجام پذیر است. در دستور زیر ناحیه خلاصه مورد نظر باید وارد شود:

(Config-router)#area <area-number> range <IP> <Mask> cost <cost>

۲- Autonomous System Boundary Router (ASBR): این مسیریاب مانند مسیریاب ABR است با این تفاوت که عمل redistribute نیز انجام می‌دهد. عمل redistribute در واقع ترجمه یک پروتکل مسیریابی به یک پروتکل مسیریابی دیگر است. این مسیریاب بسته‌های 5, 4, 3-Type ارسال می‌کند. Type-5 برای مسیریابی External و redistribute Net-ID انجام می‌شود. بدین گونه است که این مسیریاب بسته LSA با فیلدهای LSID که داخل آن Net-id و Mask قرار می‌گیرد، ایجاد می‌نماید. همچنین دو مؤلفه ASBR RID و Cost نیز در آن قرار می‌گیرد. به صورت پیش فرض Cost=20 است. با دستور زیر می‌توان در این مسیریاب خلاصه‌سازی نمود:

(Config-router)#summary-address <Net-ID> <Mask>

## ۱۷-۲ OSPF Default Route

همانند پروتکل RIGRP مسیر پیش فرض موقعی استفاده می‌شود که هیچ مسیر تطابقی در جدول مسیریابی نباشد. مسیر پیش فرض باعث کم شدن load منابع و همگرایی سریع می‌شود. ولی این امر باعث حذف بسته‌های ناشناخته و یا

فرستادن بسته‌های ناشناس به مسیر پیش فرض است. برای تعریف مسیر پیش فرض در پروتکل OSPF دو راه وجود دارد.

۱- Default-Information Originate: اگر به هر روشی مسیر پیش فرض تعیین شده باشد OSPF اجازه دارد تا بسته LSA با مسیر پیش فرض ایجاد و پخش کند. اگر مسیر پیش فرض وجود نداشت با استفاده از عبارت Always می‌توان این اجازه را داد.

```
(Config-router)#default-information originate [always] [metric <metric>] [metric-type <type>] [route-map <route-map>]
```

با وارد نمودن این دستور LSA Type-5 از نوع sub type-2 که cost=20 است، توسط ASBR ایجاد می‌شود.

۲- Stub Area: فرستادن تمام مسیرها به تمام مسیرهای امکان پذیر است اما راه ساده‌تری نیز وجود دارد. برای مثال ناحیه ۱ به ABR می‌گوید تو تنها راه ورود و خروج هستی پس مسیرها را در خود ذخیره کن و به مسیرهای داخلی ناحیه ۱ و ارسال نکن در این صورت به آن مسیرهای Stub گویند. مسیرهای Stub تمام مسیرهای ناشناس را به مسیرهای ABR می‌فرستند. مواردی که لحاظ می‌شود: ۱- مسیرهای ABR مسیر پیش فرض را با استفاده از Type-3 به Stub area ارسال می‌کند. ۲- بسته LSA ABR Type-5 را ارسال نمی‌کند. ۳- مسیرهای ABR ممکن است دیگر Type-3 LSA ها را ارسال نکنند. ۴- این امر در مسیرهای ABR و تمامی مسیرهای Stub باید تنظیم شود. در غیر اینصورت در ارتباط همسایگی تداخل پیش می‌آید. ۵- مسیرهای Stub نمی‌توانند redistribute external را انجام دهند چون نیاز به Type-5 دارند.

## OSPF Stub Area – ۱۸-۲

مسیرهای ABR بسته LSA Type-5 را فیلتر می‌کند و همچنین در Totally Stubby و Totally NSSA نیز LSA Type-3 را فیلتر می‌کند. مسیرهای ASBR وظیفه دارد تا پروتکل خارجی را به OSPF ترجمه کند و بسته LSA Type-5 را بسازد. اما اگر خواسته شود مسیری در ناحیه خاص فیلتر شود و یا از advertise کردن آن جلوگیری شود، نمی‌توان از area range یا distribute list و یا area filter list استفاده نمود، زیرا area range برای نوع ۳ یا Type-3 است. Prefix list بین LSA DB و route table است که باعث advertise می‌شود و filter list هم برای external جواب نمی‌دهد. راه حل این است که فقط روی مسیرهای ASBR دستور distribute-list prefix-list OUT را تنظیم نمود.

۱- Stub: در این حالت neighbor mismatch رخ می‌دهد و تغییر آن در بیت External Routing Capability ظاهر می‌شود. در واقع تمامی مسیرهای در این ناحیه از اطلاعات مسیریابی اطلاع ندارند و فقط می‌دانند که برای ارسال بسته باید آن را به مسیر یاب ABR ارسال نمود. در واقع همین برای آنها کافیست. برای تنظیم آن از دستور زیر استفاده می‌شود و بر روی همه مسیرهای باید تنظیم شود:

```
(Config-router)#area <area-id> stub
```

۲- Totally Stubby: فقط بر روی مسیر یاب ABR با دستور no-summary زده می‌شود و بقیه نیازی به no-summary ندارند. اما اگر در مسیرهای دیگر وارد شود مشکلی پیش نمی‌آید. در واقع مانند Stub است با این تفاوت که نیازی به دانستن اطلاعات مسیرهای خلاصه‌سازی شده هم نیست. دستور آن در زیر آمده است:

```
(Config-router)#area <area-id> stub no-summary
```

۳- Not-So-Stubby: اگر خواسته شود که مسیریابی داخل ناحیه Stub عمل ASBR انجام دهد که بتوان redistribution انجام داد از این تنظیم استفاده می‌شود که بسته LSA Type-7 را ایجاد می‌کند. این دستور در تمامی مسیرهای وارد می‌شود.

```
(Config-router)#area <area-id> nssa
```

۴- Totally Not-So-Stubby: همانند nssa است با این تفاوت که نیازی به دانستن اطلاعات مسیرهای خلاصه‌سازی شده هم نیست. این دستور در مسیر یاب ASBR وارد می‌شود. دستور آن به شرح زیر است:

```
(Config-router)#area <area-id> nssa no-summary
```

|              |                     |               |                                                    |
|--------------|---------------------|---------------|----------------------------------------------------|
| Stub         | Don't Filter Type-3 | Filter Type-5 | Cannot send type-7                                 |
| Totally Stub | Filter Type-3       | Filter Type-5 | Cannot send type-7                                 |
| NSSA         | Don't Filter Type-3 | Filter Type-5 | Can send Type-7=redistribute external to stub area |
| Totally NSSA | Filter Type-3       | Filter Type-5 | Can send Type-7                                    |

| Area              | Restriction                                                                                                                 |
|-------------------|-----------------------------------------------------------------------------------------------------------------------------|
| Normal            | None                                                                                                                        |
| Stub              | No Type 5 AS-external LSA allowed                                                                                           |
| Totally Stub      | No Type 3, 4 or 5 LSAs allowed except the default summary route                                                             |
| NSSA              | No Type 5 AS-external LSAs allowed, but Type 7 LSAs that convert to Type 5 at the NSSA ABR can traverse                     |
| NSSA Totally Stub | No Type 3, 4 or 5 LSAs except the default summary route, but Type 7 LSAs that convert to Type 5 at the NSSA ABR are allowed |

## ۲-۱۹-OSPF Virtual Link

مسیریابی که در یک ناحیه غیر صفر است اگر نیازی به دانستن اطلاعات ناحیه دیگر داشته باشد باید حتماً از طریق مسیریاب ABR که با ناحیه صفر در ارتباط است ارتباط برقرار نماید. بعضی از مواقع بنا به دلایلی همچون طراحی شبکه، قطعی ارتباط و یا نتوانستن اتصال به Backbone، نمی‌شود مستقیماً به مسیریاب ABR دسترسی پیدا نمود. Virtual Link راه حل مجازی است که این مشکل را حل کرده تا به صورت غیر مستقیم بتوان با مسیریاب ABR ارتباط برقرار نمود. در واقع این لینک مجازی، تونلی بین مسیریاب‌های بین راه تا مسیریاب ABR است. این تونل مجازی می‌تواند با مسیریاب ABR ارتباط برقرار نماید که مسیریاب مورد نظر نیز مسیریاب ABR می‌شود، زیرا DB یکسانی دارند. باید توجه داشت که مسیریاب ABR با مسیریاب مورد نظر به یکدیگر LSA به صورت unicast ارسال می‌کنند. در virtual link ارسال چندپخشی انجام نمی‌گیرد. پس از تنظیم virtual link مسیریاب‌ها Router-ID و Interface IP‌های یکدیگر را می‌دانند و به این صورت unicast انجام می‌گیرد. در virtual link در بسته‌های LSA تگ مربوط به Do Not Age (DNA) تنظیم می‌شود چون لینک مجازی بین چندین مسیریاب است. همچنین ناحیه‌ای که بر روی آن virtual link تنظیم شده است نباید ناحیه stub باشد. دستورات برای تنظیم به شرح زیر است:

```
(Config-router)#area <area-number> virtual-link <remote-RID>
```

ناحیه در اینجا حمل کننده Virtual-Link است.

```
(R1)#show ip OSPF border-router
```

Hello/Dead را نمی توان در قسمت interface-level تغییر داد. برای تغییر آن باید دستور area virtual-link را وارد نمود. بعد از up شدن لینک نمی توان hello را تغییر داد و تغییر دادن آن فقط برای شروع مفید است. در virtual link احراز هویت مانند لینک عادی به ۳ صورت امکان پذیر است: ۱- بدون احراز ۲- plain text ۳- MD5.

## ۲-۲۰- OSPF Path Manipulation

دستکاری مسیر در پروتکل OSPF همانند پروتکل EIGRP به چند روش صورت می گیرد، که شامل changing metric، summary route، filtering و changing metric-type است. دستکاری مسیر در داخل ناحیه و بین ناحیه تفاوت دارد.

۱- داخل ناحیه: فقط metric را می توان تغییر داد. از دستوراتی مانند auto-cost که از bandwidth محاسبه می گردد، یا تغییر bandwidth و یا تغییر OSPF cost در سطح interface می توان استفاده نمود.

۲- بین ناحیه: برای redundancy باید از دو مسیریاب استفاده نمود که با استفاده از summarization برای ایجاد high availability و یا filtering برای ایجاد load balancing می توان استفاده نمود.

اولویت در OSPF به صورت  $N1 > N2 > E1 > E2 > OIA > O$  است.

OSPF Intra area > OSPF Inter area > External Type 1 > External Type 2 > NSSA Type 1 > NSSA Type 2

## ۲-۲۱- Redistribution

برای ترجمه پروتکل های مسیریابی باید حداقل یک مسیریاب دارای خاصیت های زیر باشد:

- ۱- حداقل یک لینک در ارتباط با هر پروتکل مسیریابی باشد.
- ۲- برای هر پروتکل مسیریابی تنظیمات مربوط به آن بر روی مسیریاب باید تنظیم شود.
- ۳- بر روی هر پروسه پروتکل مسیریاب دستور redistribute وارد شود.

عمل Redistribute درباره جزئیات هر پروتکل چیزی نمی داند و از topology و یا Database هر پروتکل مسیریابی اطلاعاتی نمی فهمد. ترجمه فقط بر اساس جدول مسیریابی اطلاعات صورت می گیرد. به طور خاص ابتدا EIGRP و سپس OSPF توضیح داده می شود. برای ترجمه پروتکل EIGRP به OSPF باید داخل تنظیمات مربوط

به پروتکل OSPF دستور وارد شود. برای فهم آن باید مترجم پروتکل OSPF باشد. زمانی که یک پروتکل به پروتکل دیگر redistribute می‌شود COST یا Metric و دیگر موارد باید به گونه‌ای ترجمه شود که برای پروتکل دیگر قابل فهم باشد. در پروتکل‌های OSPF و EIGRP این مؤلفه بر اساس پهنای باند است ولی در پروتکل RIP بر اساس hop count است. ولی به هر حال از اطلاعات interface خروجی می‌گیرد. ۳ راه برای تنظیم مؤلفه Cost وجود دارد: ۱- EIGRP: از طریق default-metric در پروتکل EIGRP این مقدار را تعیین نمود. ۲- در redistribute command وارد نمود. ۳- با استفاده از route-map در قسمت redistribute می‌توان وارد نمود. به ترتیب اولویت route-map بعد redistribute بعد EIGRP پیشنهاد می‌شود. دستور تنظیم آن به صورت زیر است:

```
(Config-router)#redistribute <protocol> [process-ID | ASN] [metric BW Delay  
Reliable Load MTU] [match internal | NSSA-external | external-1 | external-2] [tag tag-  
value] [route-map name]
```

Tag برای انتخاب مسیر و سپس جدا کردن آن است که مقدار آن integer است.

```
(R1)#show ip EIGRP topology
```

```
(Config-router)#redistribute connected subnets
```

Connected منظور آن است که هیچ چیزی تنظیم نشده است و Subnets برای class-less بودن آن است. در ترجمه پروتکل OSPF به پروتکل EIGRP، مشابه عمل redistribute در قبل انجام می‌شود با این تفاوت که بسته‌های Type-5 یا Type-7 برای NSSA، ممکن است به بسته‌های type-4 تبدیل شود. در تنظیمات دستوری آن تفاوتی وجود دارد. اول اینکه metric-type نشان دهنده 1,2 subtype است و دوم اینکه subnets برای class full یا classless استفاده می‌شود. مقادیر پیش فرض برای ترجمه بدین گونه است که: ۱- metric: اگر منبع BGP بود برابر ۱ اگر OSPF با Process متفاوت بود برابر مقدار منبع و دیگر پروتکل‌ها برابر ۲۰ است. ۲- Type: برای NSSA از نوع ۷ و بغیر از آن نوع ۵ استفاده می‌شود ۳- subnets: به صورت پیش فرض class full است. با استفاده از Route-map و یا Distribute-list می‌توان مسیرها را فیلتر نمود. Route-map قابلیت‌های زیادی دارد که شامل موارد ذیل است: ۱- کدام مسیر انتخاب شود و یا prefix/length آن تغییر کند. ۲- metric و یا metric-type کدام مسیر را تغییر دهد. ۳- برای مسیرها tag مشخص می‌شود. با استفاده از دستور match در route-map قابلیت‌هایی زیر وجود دارد:

```
(Config-route-map)#match interface | metric | route-type | tag
```

(Config-route-map)#match ip address | next-hop | route-source

با استفاده از دستور set در route-map قابلیت‌هایی زیر وجود دارد:

(Config-route-map)#set metric | metric-type | tag

برچسب زدن یا همان Tag فواید زیادی دارد که یک نمونه آن جلوگیری از حلقه است بدین شکل که اگر مسیری دوباره به مکان قبلی برگردد با استفاده از tag شناخته می‌شود.

## ۲-۲-۲ Prevention Loops to Redistribute OSPF

حلقه مسیریابی بدین معنی است که مسیری به اشتباه در یک حلقه بین چندین مسیریاب ارسال شود. موقعی که redistribute انجام می‌شود، زمانی حلقه ایجاد می‌شود که دو مسیریاب ترجمه را بر عهده داشته باشند در این حالت ممکن است مسیری به اشتباه از طریق مسیریاب دیگر به عنوان مسیر بهتر انتخاب شود. زمانی این اتفاق رخ می‌دهد که administrative distance مسیر این پروتکل بهتر از پروتکل دیگر باشد در EIGRP هیچ موقع این اتفاق نمی‌افتد. البته به شرطی که بین دو پروتکل نباشد، زیرا internal route آن کمتر از همه پروتکل‌های IGP است و external route آن بیشتر از همه IGP‌ها است.

۱- Change Metric: یکی از راه‌های جلوگیری از حلقه‌های مسیریابی تغییر metric است. فرض کنید دو مسیریاب با دو پروتکل OSPF و RIP وجود داشته باشند تا عملیات Redistribute را انجام دهند. وقتی update از RIP به OSPF وارد شود metric بر اساس hop count است. اگر در مسیریاب دیگر این مسیر hop count کمتری وجود داشته باشد بسته‌ها داخل OSPF می‌چرخد و سپس به مسیریاب دیگر می‌رسد که مسافت بیشتری طول می‌کشد. و برای حل این مشکل باید بیشترین hop count را به عنوان metric قرار داد. ولی update از OSPF به RIP چون Administrative Distance کمتری دارد مشکلی ایجاد نمی‌شود.

۲- Change Administrative Distance: AD به صورت محلی برای مسیریاب‌ها می‌باشد و آن را آگاهی یا اعلام نمی‌کنند و بقیه مسیریاب‌ها نیز از آن اطلاعی ندارند. برای جلوگیری از حلقه مسیریابی بین RIP و OSPF باید OSPF AD بیشتر از RIP AD باشد. باید توجه داشت که برای redundancy از دو مسیریاب برای redistribute استفاده می‌شود و به همین دلیل امکان وجود حلقه وجود دارد و با AD می‌توان بین مسیریاب‌های ASBR این فیلتر را انجام داد. برای تنظیم آن بر روی پروتکل‌های مختلف از دستورات زیر استفاده می‌شود:

RIP-> distance <value>

EIGRP-> distance <internal AD> <external AD>

EIGRP-> distance <value> <ip src add> <wildcard mask> این فقط برای داخلی می باشد نه خارجی

OSPF->distance <external AD> <intra AD> <inter AD>

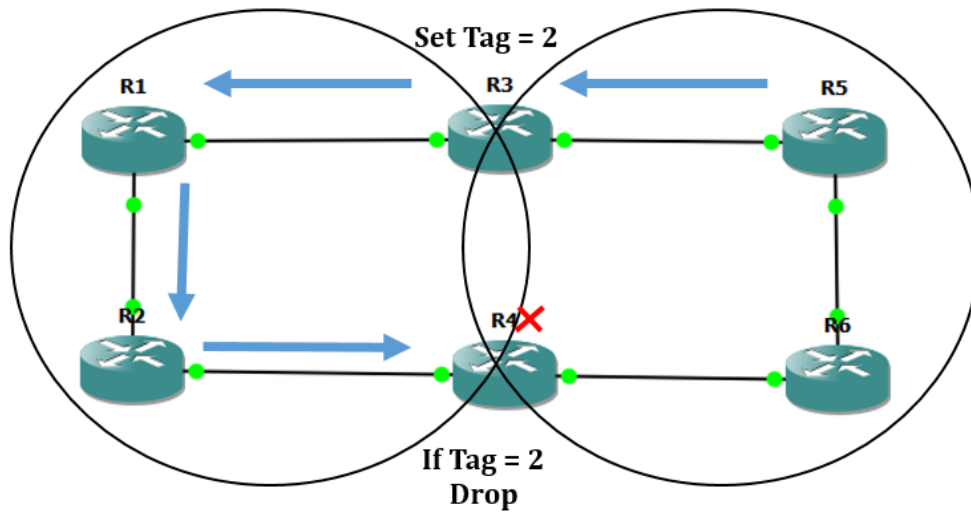
OSPF->distance <value> <ip src add> <wildcard mask>

دستور OSPF برای همه هست فقط به شرط اینکه سلسله مراتب  $0 > 0 \text{ IA} > 0 \text{ E1/E2}$  رعایت شود.

در OSPF اگر مسیر یکی باشد با دستور بالا همان AD را در نظر می گیرد و تغییر نمی کند. اگر HOST IP یعنی 32/ باشد تغییر می دهد. در دستور بالا از ACL نیز می توان استفاده نمود.

در کل زمانی که بین IGP1 و IGP2 ترجمه می شود و از IGP2 به IGP3 نیز ترجمه می شود مشکل زمانی به وجود می آید که از IGP2 به IGP3 دوباره بر می گردد که  $\text{external AD IGP3} < \text{external AD IGP2}$  باشد، که باعث حلقه می شود.

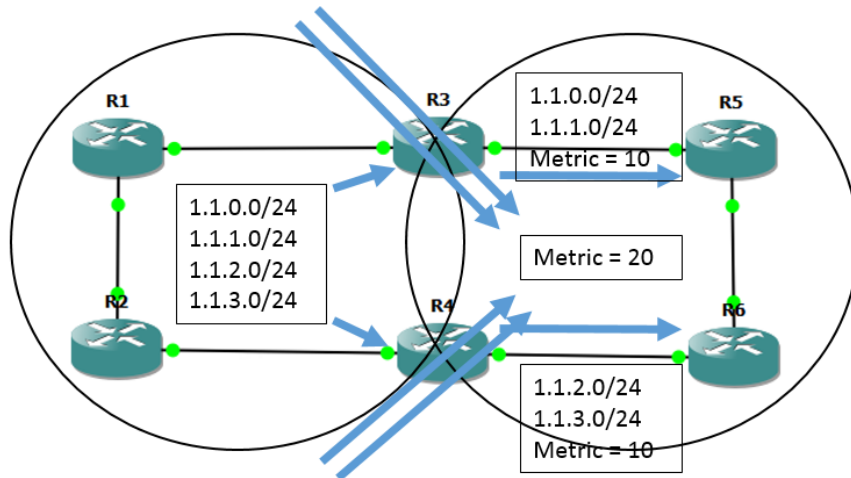
۳- Filtering (Set Tag or Subnet): استفاده از فیلتر کردن بر دو اساس می توان از حلقه جلوگیری نمود. یکی بر اساس Net-ID است که بسیار سخت است و اشتباهات انسانی در آن دخیل می شود. دیگری زدن برچسب که راحت تر است. زدن برچسب برای شناسایی، فیلتر کردن و جلوگیری از حلقه است. برچسب یا همان Tag را می توان در route-map، distribute-list و redistribute استفاده نمود. در اینجا برچسب زدن در هنگام پروسه redistribute انجام می گیرد. برای تنظیم زدن برچسب در route-map در redistribute می توان تمام مسیرها را برچسب زد تا مسیریاب دیگر که این برچسب را می بیند متوجه شود که قبلاً این مسیر redistribute شده است و نباید دوباره ارسال نماید. باید توجه داشت که برچسب بین EIGRP و OSPF باقی می ماند ولی در RIP باید دوباره زده شود.



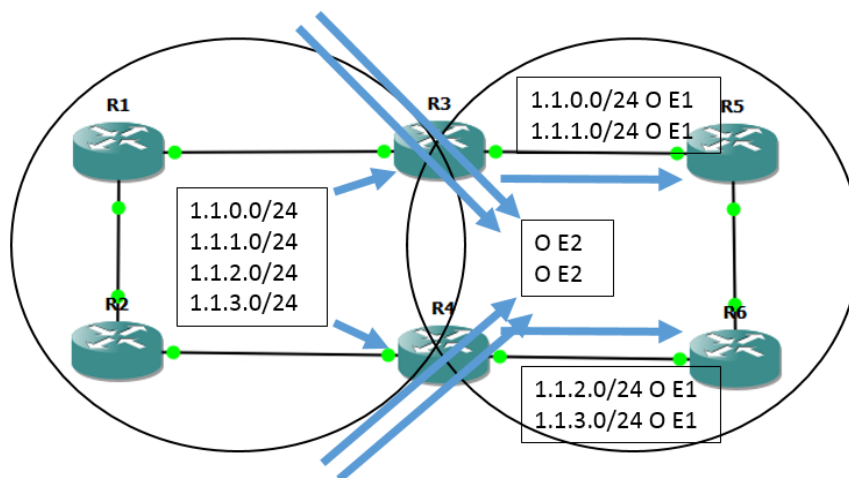
## ۲-۲۳- Redistribute OSPF Path Manipulation

زمانی می‌توان مسیر را به صورت دستی موقع redistribute تغییر داد که حداقل دو مسیر یاب وجود داشته باشد.

۱- با استفاده از تغییر Metric: این روش را می‌توان با استفاده از route-map، redistribute و default-metric انجام داد.

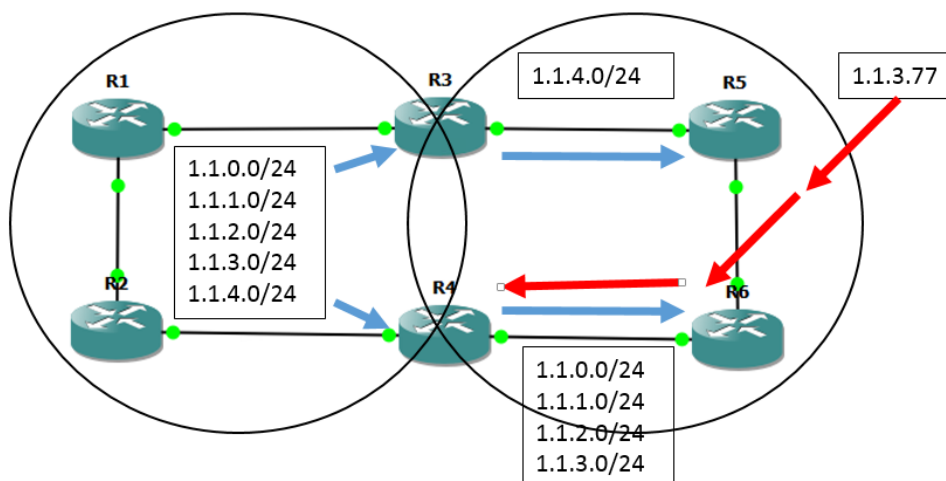


۲- با استفاده از تغییر Metric Type: این روش بدین گونه است که E1 اولویت بالاتری نسبت به E2 دارد.



۳- با استفاده از فیلتر کردن: این روش مناسب نیست چون redundancy را از بین می برد و با قطع مسیر یاب مسیر از دست می رود.

۴- با استفاده از Summarization: شبکه ای که بیشتر تطابق داشته باشد اولویت بالاتری دارد. مسیر یاب دیگر به عنوان Backup قرار می گیرد.



### ۳- Routing Border Gateway Protocol

این پروتکل برای Exterior gateway Protocol(EGP) طراحی شده است تا بین دو یا چند سرویس دهنده پیاده سازی شود و مسیرها را جابه جا نماید. پروتکل های IGP برای محاسبه کردن مسیرهای زیاد طراحی نشده اند، زیرا پروسه های زیادی دارند مانند SPF در پروتکل OSPF و همچنین پیچیدگی زیادی دارند. پروتکل EGP را می توان

برای شبکه‌های بزرگ Enterprise یا این که در قسمت Core routers که نیاز به دانستن جزئیات زیادی ندارند، و در سطح‌های پایین از پروتکل IGP استفاده نمود. همچنین اگر سازمانی با چند ISP و یا اینکه با یک ISP چندین ارتباط داشت، می‌توان از پروتکل BGP استفاده نمود. تفاوت‌ها و شباهت‌های پروتکل‌های IGP و BGP به شرح زیر است:

- ۱- پروتکل BGP مانند پروتکل IGP نیاز به ارتباط با همسایگان دارد ولی به صورت پیش فرض در پروتکل IGP به صورت خودکار این عمل انجام می‌گیرد، ولی در پروتکل BGP به صورت دستی باید صورت گیرد. در دستگاه‌های غیر سیسکو می‌توان خودکار تعریف نمود ولی در سیسکو این امکان وجود ندارد.
  - ۲- پروتکل BGP مانند پروتکل IGP، NET-ID، subnet و Next-Hop برای آن شبکه را ارسال می‌کند.
  - ۳- در پروتکل BGP همسایگان نیازی نیست از لحاظ فیزیکی و NET-ID همسایه باشند. تقریباً مانند virtual link است.
  - ۴- پروتکل‌های OSPF و EIGRP پروتکل‌های مخصوص به خود دارند، ولی پروتکل RIP از UDP و پروتکل BGP نیز از TCP 179 استفاده می‌کند.
  - ۵- پروتکل BGP با ارسال NET-ID و Subnet اطلاعاتی ارسال می‌کند که به آن Network Layer Reachability Information (NLRI) گویند.
  - ۶- پروتکل BGP مشخصه‌هایی برای ارسال IPV6، Multicast، انتخاب بهترین مسیر و Metric دارد.
  - ۷- پروتکل IGP بسیار سریع همگرا می‌شود، ولی پروتکل BGP نیاز به سریع همگرا بودن ندارد، ولی باید Scalable یا مقیاس پذیر باشد.
  - ۸- پروتکل BGP همچنین از الگوریتم Path Vector Logic استفاده می‌کند که شبیه Distance Vector است.
- در پروتکل IGP اگر چندین مسیر با مقصد یکسان وارد مسیریاب شود ابتدا AD، سپس Metric و در آخر Load Balance انجام می‌شود. در پروتکل BGP به صورت سلسله مراتبی، درختی و همچنین با عمق زیاد بررسی می‌شود. با بررسی این امر احتمال ضعیفی وجود دارد هزینه مسیرها یکسان شود و به صورت پیش فرض دو مسیر یکسان در جدول مسیریابی قرار نمی‌گیرد. لینک Backup در پروتکل EIGRP فقط مخصوص به خود مسیریاب است و advertise یا اعلام نمی‌شود. در پروتکل BGP بهترین مسیر advertise یا اعلام می‌شود و مسیری که دستی تغییر کرده است، اعلام نمی‌شود.

### ۳-۱- iBGP & eBGP

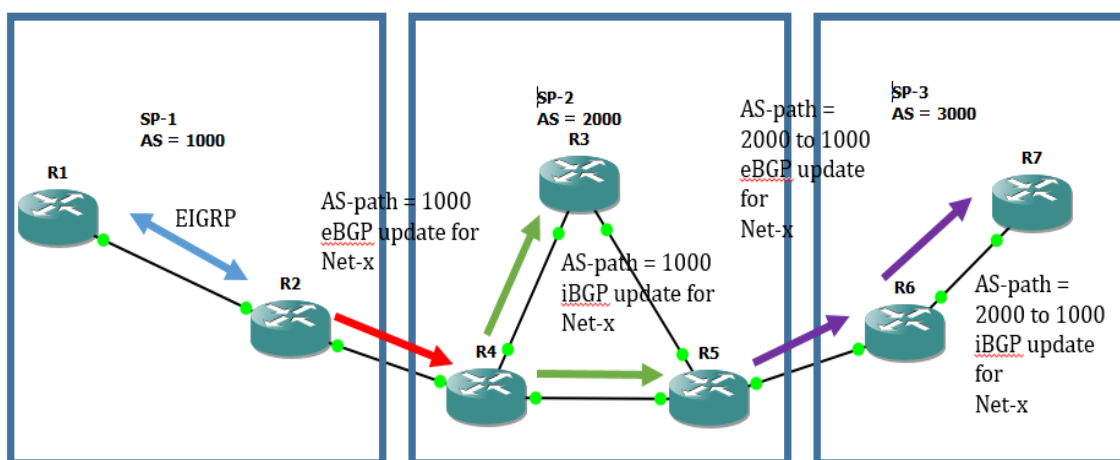
پروتکل EIGRP برای ارتباط نودی از شرکت که خارج از شرکت قرار دارد، طراحی شده است. در پروتکل BGP ارتباط با نودی است که خارج از شرکت و بر عهده سرویس دهنده دیگر است، طراحی شده است. حال اگر پروتکل BGP در ارتباط داخلی استفاده شود شماره Autonomous System مهم نیست، ولی در ارتباط خارجی حتماً باید Unique یا یکسان باشد. در پروتکل EIGRP بعد از وارد نمودن AS، همسایگان همدیگر را پیدا می‌کنند، ولی در پروتکل BGP باید آنها را به صورت دستی با AS مربوط به همسایه وارد نمود. اگر همسایه در همان AS باشد همسایه iBGP است. اگر در یک AS نباشد همسایه eBGP است. تفاوت‌هایی که وجود دارد در نحوه neighbor ship، نحوه ارسال Update و همچنین انتخاب بهترین مسیر است. پروتکل BGP بسته‌ها را بر خلاف پروتکل‌های EIGRP و OSPF به صورت unicast ارسال می‌کند. در بسته Update نام AS نیز ارسال می‌شود و فیلدی وجود دارد به نام AS-PATH که دنباله‌ای از update ها در آن قرار دارد. این فیلد از چپ به راست از آخرین به اولین ارسال کننده بسته Update قرار می‌گیرد. در پروتکل BGP دو نوع آدرس AS وجود دارد. Public که می‌تواند در بستر اینترنت قرار گیرد و Private که داخلی است. فیلد AS دارای 16 bit است.

Public -> 1-64495

Private -> 65512-65534

و شماره‌های 0, 64496-65511, 65535 رزرو است.

اگر بسته‌ای از private به سمت Public رفت با AS نامعتبر مشکل جایی ایجاد می‌شود که public با public دیگر در ارتباط است.



## ۲-۳ BGP Neighbor Relationship

پروتکل BGP نیز مانند پروتکل‌های دیگر از سه مرحله تشکیل شده است ۱- neighbor relationship ۲- exchange topology information ۳- run a best-path algorithm. در ارتباط با همسایه از TCP 179 و به صورت static استفاده می‌شود. پیچیدگی ارتباط بین همسایه‌ها از پروتکل EIGRP بیشتر و از پروتکل OSPF کمتر است. قبل از اینکه ارتباط بین همسایه‌ها ایجاد شود، مراحل باید انجام شود: ۱- ASN که در دستور وارد می‌شود باید در Local و remote-as همسایه دیگر برابر باشد. ۲- همسایه باید از طریق پروتکل مسیریابی IGP با هم در ارتباط باشند، اگر در ارتباط نباشند همسایه نمی‌شوند. ۳- router-id مسیریاب‌ها نباید برابر باشند. ۴- اگر authentication تنظیم شود پسورها باید یکسان باشند. ۵- TCP handshake باید به صورت کامل انجام شود. به طور مثال firewall بین آنها اجازه عبور دهد. ۶- دستور همسایه‌ها باید در دو مسیریاب وارد شود. مانند بقیه پروتکل‌ها router-id یا به صورت static و یا dynamic از اینترفیس‌ها انتخاب می‌شود. برای احراز اصالت از دستور neighbor <ip-add> password <key> استفاده می‌شود. باید توجه داشت که authentication در هدر TCP است نه پروتکل BGP. اگر حالت همسایه‌ها establish باشد به معنی برقرار بودن ارتباط است ولی هنوز update ارسال نمی‌کند.

## ۳-۳ eBGP

در ارتباط با همسایه به صورت eBGP فرض بر آن است که به صورت مستقیم با هم در ارتباطند. ابتدا در جدول مسیریابی به دنبال همسایه می‌گردد در صورت پیدا کردن از مسیر مشخص ارتباط برقرار می‌کند در غیر این صورت هیچ عملی انجام نمی‌شود. دستور آن به شکل زیر است:

```
(Config)#router bgp <ASN>
```

```
(Config-router)#neighbor <ip-add> remote-as <remote-ASN>
```

قبل از این که ارتباط همسایه‌ها ایجاد شود ابتدا جدول مسیریابی را مشاهده سپس اینترفیسی که بسته از آن خارج می‌شود را انتخاب و IP آن را به عنوان آدرس مبدأ در TCP قرار می‌دهد البته به صورت پیش فرض است. حال مشکلاتی وجود دارد. اگر چند ارتباط با همسایه‌های دیگر وجود داشت باید برای هر اینترفیس جداگانه دستور neighbor را وارد نمود، ولی برای هر ارتباط کل مسیرها ارسال می‌شود و bandwidth و Memory افزایش می‌یابد. اگر هم یک لینک تنظیم شود با قطع لینک ارتباط قطع می‌شود و fail over از بین می‌رود. بهترین راه انتخاب loopback است. در دستور neighbor آدرس loopback وارد می‌شود. در eBGP مؤلفه TTL=1 است، زیرا فرض بر آن است که مستقیم در ارتباط هستند. در نتیجه برای آن که چند مسیریاب بین آنها باشد باید از دستور ebgp multihop استفاده شود. در این حالت TTL=255 می‌شود. اگر loopback انتخاب شود و مسیری وجود نداشته

باشد، باید از static route استفاده شود. در استفاده از loopback مشکل دیگر این است که برای آدرس مقصد هنوز physical interface وارد می شود و به مسیریاب همسایه که وارد می شود به دلیل اینکه loopback به عنوان همسایه است و در حال حاضر آدرس physical interface وارد شده آن را drop می کند. برای حل این مشکل از دستور update-source استفاده می شود. باید توجه داشت که همچنین بر روی هر مسیریاب فقط یک ASN می توان تنظیم نمود.

```
(Config-router)#neighbor <ip-add> update-source loopback0
```

```
(Config-router)#neighbor <ip-add> ebgp-multihop
```

### ۳-۴- iBGP

از iBGP بیشتر موقعی استفاده می شود که ISP یا سرویس دهنده بین دو ISP دیگر باشد. به عبارتی به عنوان پل ارتباطی باشد که به آن عموماً transit AS گویند. تنظیمات آن مانند eBGP و نیازمندی های همسایه ها نیز مانند eBGP است تنها تفاوت آن TTL=255 است. در iBGP اگر مسیریابی بین دو مسیریاب دیگر باشد و بر روی آن BGP اجرا نشده باشد، مسیر می تواند مسیریابی شود، البته در همه مسیریاب ها، اگر مسیر در جدول route باشد اینکه که کدام پروتکل استفاده شده، مشکلی پیش نمی آید و مسیریابی صورت می گیرد.

### ۳-۵- BGP Neighbor State

حالت همسایه ها در پروتکل BGP به شرح زیر است:

۱- Idle: ارتباط با همسایه برقرار نشده است. دلیل آن می تواند عدم دسترسی و یا عدم وجود در جدول مسیریابی باشد.

۲- Connected: بعد از idle پیش می آید. منتظر برقرار شدن ارتباط TCP است. معمولاً نمی توان این حالت را مشاهده نمود، زیرا بسیار سریع رخ می دهد.

۳- Active: حالت خوبی نیست، زیرا که TCP Connection Fail اتفاق می افتد. در این حالت باید بسته ها را مشاهده و firewall یا ACL را بررسی نمود. پارامتر connect-retry فیلدی است که در هر حالت تغییر می کند. این زمان به صورت پیش فرض ۱۲۰ ثانیه است و نمی توان آن را تغییر داد. دلیل وجودی آن این است که بعد از این زمان ارتباط جدیدی با مقصد برقرار می کند و اگر ارتباط کامل شد این عدد پاک می شود در غیر این صورت بعد از این مدت دوباره درخواست ارتباط می دهد.

۴- OpenSent (OpenMessage): در این حالت ارتباط TCP کامل شده و در حال گفت و گو مسیریاب ها در رابطه با سازگاری و تنظیم پارامترهای با یکدیگر هستند.

۵- OpenConfirm: تمام شدن و قبول کردن تنظیمات در حالت OpenSent است.

۶- Established: مرحله آخر است و به معنی برقرار شدن کامل همسایگی است. در این حالت بسته‌های Update ارسال می‌شود.

### ۶-۳ BGP Packet Type

تمامی بسته‌های BGP بر روی TCP با پورت مقصد ۱۷۹ سوار می‌شوند. در قسمت Marker احراز اصالت متفاوت است. بسته‌های BGP به ۴ دسته تقسیم می‌شود:

|                            |                  |                |
|----------------------------|------------------|----------------|
| IP Header                  |                  |                |
| TCP Header                 |                  |                |
| Marker (All "Fs") 16-Bytes | Length (2-Bytes) | Type (1 Bytes) |
| BGP Data                   |                  |                |

۱- Open: نوع بسته ۱ است. عملیت 3 way TCP Hand Shaking است تا موقعی که Established شود.

|                            |                  |                  |           |
|----------------------------|------------------|------------------|-----------|
| IP Header                  |                  |                  |           |
| TCP Header                 |                  |                  |           |
| Marker (All "Fs") 16-Bytes | Length (2-Bytes) | <b>Type = 1</b>  |           |
| Version = 4                | My AS#           | Hold Time        | Router-ID |
| Optional Parameters Lenght |                  | BGP Capabilities |           |

۲- Update: پس از Open است. نوع بسته ۲ است. برای جابه‌جایی مسیرها بین مسیریاب‌ها است. در پروتکل‌های OSPF و EIGRP با تغییر فیلدهای age یا delay به مقدار ماکزیمم (poison route) می‌توان اختلال در پروتکل انجام داد، ولی در پروتکل BGP با استفاده از فیلد Unfeasible از این کار جلوگیری می‌شود. TLV= Type Link Value

|            |  |
|------------|--|
| IP Header  |  |
| TCP Header |  |

|                              |                           |                 |
|------------------------------|---------------------------|-----------------|
| Marker (All "Fs") 16-Bytes   | Length (2-Bytes)          | <b>Type = 2</b> |
| Unfeasible Routes Length     | Withdrawn Routes (if any) |                 |
| Total Path Attributes Length | Path Attributes (TLV)     |                 |
| NLRI Prefix Length           | NLRI Prefix               |                 |

۳- Notification: نوع بسته ۳ است. برای BGP ERROE است. معمولاً نتیجه آن، برقراری دوباره همسایگی است.

|                            |                  |                 |
|----------------------------|------------------|-----------------|
| IP Header                  |                  |                 |
| TCP Header                 |                  |                 |
| Marker (All "Fs") 16-Bytes | Length (2-Bytes) | <b>Type = 3</b> |
| Error Code                 | Error Subcode    |                 |
| Data                       |                  |                 |

۴- Keapalive: نوع بسته ۴ است و هیچ داده‌ای در آن نیست. مانند hello در پروتکل EIGRP است.

|                            |                  |                 |
|----------------------------|------------------|-----------------|
| IP Header                  |                  |                 |
| TCP Header                 |                  |                 |
| Marker (All "Fs") 16-Bytes | Length (2-Bytes) | <b>Type = 4</b> |

### ۳-۷ BGP Table

جدول پروتکل BGP شبیه به EIGRP Topology یا OSPF Database است. با دستور show ip bgp می‌توان این جدول را مشاهده نمود. پروتکل BGP به این که چه تعداد مسیر یاب در بین راه است اهمیت نمی‌دهد، ولی به این که چند AS بین راه است اهمیت می‌دهد. در جدول پروتکل BGP نمادهایی وجود دارد که در ادامه توضیح داده شده است.

- ۱- نماد >: نشان دهنده بهترین مسیر است.
- ۲- نماد\*: نشان دهنده مسیرهای صحیح و می توان آن را اعلام نمود.
- ۳- نماد R: مخفف RIB(Routing Information Base) است. نشان دهنده مسیر بد است به معنای اینکه مسیر بهتری نیز وجود دارد. همچنین به معنای یاد گرفتن مسیر است که در جدول مسیریابی قرار نمی گیرد. می تواند مشکل یا عدم دسترسی به مسیر باشد یا AD آن مناسب نباشد. در کل ممکن است به ۳ دلیل ایجاد شود: ۱- AD در BGP بیشتر از AD در جدول مسیریابی باشد. ۲- ممکن است حافظه مسیریاب پر شده باشد و memory failure اتفاق افتاده باشد، یا به دلیل کمبود منابع TCP connection نتواند انجام گیرد. ۳- جدول VRF (Virtual Routing and Forwarding) پر شده باشد، یا از حدی که مدیر سیستم تعریف کرده است تجاوز کرده باشد. این جدول مانند جدول مسیریابی است با این تفاوت که مدیر سیستم برای مشتری خاص تعریف می کند و به این منظور است که اگر دو IP مشابه از مشتریان وارد شبکه شد فقط در جدول مسیریابی خود قابل قبول باشد و تداخلی بر دیگر نداشته باشند. تعداد مسیرها در این جدول توسط مدیر سیستم قابل تنظیم است و اگر از حد آن گذشت، می تواند مسیرهای جدید را ذخیره نکند.
- ۴- نماد i: این نماد یعنی مسیر از iBGP یاد گرفته شده است.
- ۵- Next hop: به معنی ارسال بسته مسیر به این مسیریاب است. 0.0.0.0 این بدان معنا است که مسیر به صورت locally originate ایجاد شده است، به عبارت دیگر خود مسیریاب آن را ایجاد نموده است، یا مانند loop back باشد.
- ۶- Path: به معنای AS مسیری است که مسیر از آن جا می آید. اگر خالی باشد به معنای این است که خود مسیریاب ایجاد کننده و مالک آن است. همچنین ممکن است در همان AS باشد که خود در آن است.

### ۳-۸ BGP AS\_PATH Attribute

BGP برای ارسال بسته های update به همسایه ها مسیر را همراه با ویژگی های آن از جمله AS\_PATH ارسال می کند تا مسیر را توضیح دهد. به صورت پیش فرض مسیریاب های BGP از یک مسیر خاص، مکن است دو یا بیشتر بسته های Update را دریافت کنند. این بسته ها شامل AS\_PATH است که مسیریاب با توجه به آن بهترین مسیر را انتخاب می کند و هیچ اهمیتی به تعداد مسیریاب یا لینک نمی دهد، بلکه کوتاهترین AS\_PATH بهترین مسیر در جایگاه مقایسه AS\_PATH است. AS\_PATH دو مؤلفه دارد. AS\_SEQ یا صنفی که در آن AS ها به ترتیب قرار می گیرند. ترتیب آنها از آخرین AS به اولین AS است که از چپ به راست نوشته می شود. مؤلفه بعدی AS\_SET است. اگر عمل summarization رخ دهد، بسته فیلد AS\_SEQ را برابر مسیریاب خلاصه کننده و بقیه AS ها را داخل { } قرار می دهد. این AS ها ترتیب خاصی ندارند. به این صف از AS ها AS\_SET گویند. AS\_PATH برای جلوگیری

از حلقه مسیریابی نیز استفاده می‌شود به این صورت که اگر مسیریاب خود را بین ASها ببیند یعنی حلقه ایجاد شده است. اگر مسیریابی مسیر خلاصه را نداشته باشد ولی AS\_PATH خود را ببیند بسته را حذف می‌کند.

### ۳-۹- Route into BGP

پروتکل BGP برخلاف پروتکل IGP که با وارد نمودن شبکه‌های اطراف آن همسایه‌ها به صورت خودکار شناخته و شبکه‌ها را به هم اعلام می‌کنند، عمل نمی‌کند. در پروتکل BGP همسایه‌ها به صورت دستی وارد می‌شوند و شبکه‌ها به معنی اعلام به یکدیگر نیستند، بلکه به معنی ترجمه (redistribute) شبکه‌های IGP به BGP است. AD مربوط به  $iBGP=200$ ,  $eBGP=20$  است. روش‌هایی که می‌توان redistribute را انجام داد به شرح زیر است:

۱- استفاده از Network Command: این دستور در پروتکل‌های IGP و BGP متفاوت است در پروتکل IGP به معنای اجازه دادن برای شنیدن و وارد شدن اطلاعات مسیرهای IGP یا مقایسه با اینترفیس برای فعال کردن بر روی آن اینترفیس است. برای فعال کردن پروتکل IGP در network command وارد می‌شود، که فعال کردن به معنای فرستادن Hello و خودکار اضافه نمودن همسایه‌ها و اعلام کردن اطلاعات به هر همسایه‌ای که به صورت مستقیم با آن در ارتباط است. ولی در پروتکل BGP برخلاف همه موارد بالا به اینترفیس گوش نمی‌دهد و بر روی آن پروتکل فعال نمی‌شود. Hello ارسال نمی‌کند و همسایه به صورت دستی وارد می‌شود. در پروتکل BGP بدان معنا است که به جدول مسیریابی نگاه می‌کند، مسیرهای به غیر از BGP را مشخص می‌کند و تبدیل به پروتکل BGP می‌کند. سپس در داخل جدول BGP قرار می‌دهد و به همسایه‌ها اعلام می‌نماید.

`(Config-router)#network <IP> mask <subnet mask>`

Mask عبارت دلخواه است. بدون داشتن mask به منظور class full است و فقط class full را اعلام می‌کند.

در BGP دستور auto-summary نیز معنی متفاوتی دارد. اگر auto-summary فعال باشد و mask نیز در network command باشد تاثیری ندارد. اگر auto-summary فعال باشد و mask نیز در network command نباشد فقط شبکه‌هایی با مسیرهای class full اعلام می‌شود.

۲- استفاده از redistribute command: اگر مسیرها در جدول مسیریابی زیاد باشد، برای وارد کردن آنها زمان زیادی صرف می‌شود. در واقع تبدیل کردن به پروتکل BGP سخت باشد، با استفاده از ترجمه (redistribute) از پروتکل IGP به پروتکل BGP کار راحت‌تری است. از این دستور می‌توان استفاده نمود. تفاوت آن با network command این است که network اولویت بیشتری نسبت به redistribute دارد. البته در پروتکل BGP نگران metric نباید شد.

۳- استفاده از summarization یا aggregation: برای تنظیم آن از دستور زیر استفاده می‌شود:

(Config-router)#aggregate-address <prefix> <prefix-length> [summary-only]

Summary-only فقط برای اعلام نکردن به همسایه‌ها است و فقط برای استفاده شخصی مسیریاب است. این دستور به تنهایی کافی نیست و باید همراه network یا redistribute استفاده شود. زمانی از BGP استفاده می‌شود که در واقع redistribute انجام شود. وقتی summary انجام شود مشابه IGP در redistribute، پروتکل اجازه ندارد summary انجام دهد، مگر دستور summarization وارد شود.

### ۳-۱۰- Outgoing BGP Scenario

برای ارسال مسیره‌ها به خارج از شبکه داخلی ۴ سناریو وجود دارد:

۱- Single homed: ارتباط با یک ISP و با یک لینک باشد. در این حالت بهتر است از static route

استفاده شود. اگر از پروتکل BGP استفاده شود در واقع از حداقل فواید BGP استفاده می‌شود.

۲- Dual homed: ارتباط با یک ISP و با حداقل دو لینک باشد. در این صورت با یک مسیریاب دخالت

دارد. به عبارت دیگر لینک‌ها به صورت load balance و یا backup همدیگر هستند. در این حالت نیز

بهتر است از static route استفاده شود. اگر چندین مسیریاب باشد BGP زمانی بهتر است که مسیری بهتر

از مسیر دیگر تعریف شود و اگر تفاوتی بین مسیر نباشد BGP یا static route تفاوتی ایجاد نمی‌کند.

۳- Single multi homed: ارتباط با حداقل دو ISP و با هر ISP حداقل یک لینک باشد.

۴- Dual multi homed: ارتباط با حداقل دو ISP و با هر ISP حداقل دو لینک باشد.

### ۳-۱۱- BGP Route Advertisement

۱- پروتکل‌های مسیریابی تنها مسیرهایی را اعلام می‌کنند که بهترین مسیر است. پروتکل BGP تنها یک مسیر

بهتر را ارسال می‌کند و قطعاً یک مسیر بهتر بیشتر وجود ندارد. زیرا با مقایسه اثبات می‌شود که تنها یک مسیر

بهتر است.

۲- در پروتکل BGP مسیریاب‌ها، مسیرهایی که در AS خود یاد می‌گیرند اعلام نمی‌کنند. در واقع مسیریاب‌ها

iBGP learned را به iBGP Peers ارسال نمی‌کنند. در پروتکل EIGRP بهترین مسیره‌ها ارسال شده و

اگر چند مسیر وجود داشت یا load balance انجام می‌دهد یا backup قرار می‌دهد. در BGP فقط یک

مسیر بهتر وجود دارد.

۳- در پروتکل EIGRP در بسته به روزرسانی، next hop، خود مسیریاب ارسال کننده بسته قرار می‌گیرد. ولی

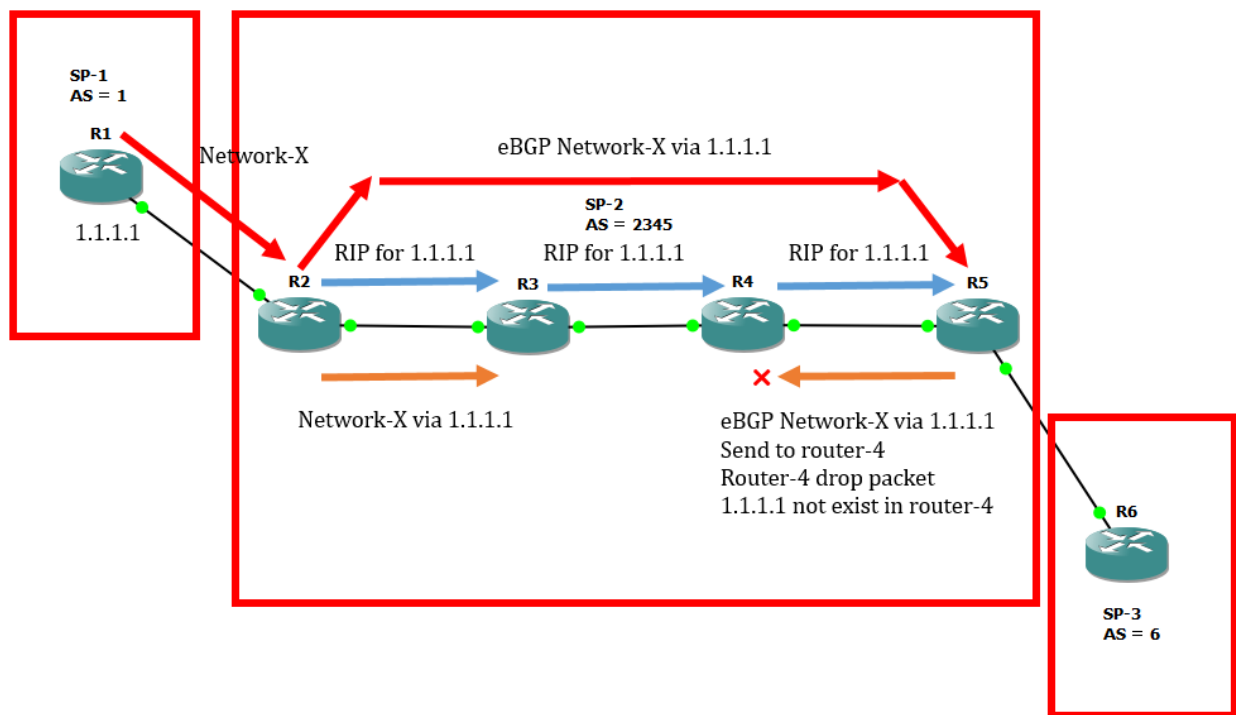
در پروتکل BGP لزوماً این طور نیست و ممکن است مسیریابی قرار گیرد که اصلاً مسیریاب گیرنده به آن

دسترسی نداشته باشد یا مسیریاب پشت سری باشد. این امر مشکل ساز می شود که در ادامه توضیح داده می شود.

۴- جلوگیری از حلقه: برای جلوگیری از حلقه اگر eBGP باشد که در AS\_PATH در صورت مشاهده کردن AS خود، مسیریاب تشخیص حلقه می دهد. اگر iBGP باشد که AS\_PATH تغییر نمی کند و بسته به روز رسانی نیز ارسال نمی شود. در نتیجه فقط بهترین مسیر ارسال می شود. در پروتکل RIP اگر تعداد hop ها بیشتر از 15 باشد حلقه تشخیص داده می شود. در پروتکل EIGRP به feasible condition و در پروتکل OSPF به LSDB نگاه می شود.

۵- Next hop: در eBGP مسیریاب next hop را تغییر و ip خود را قرار می دهد. در iBGP هیچ تغییری انجام نمی گیرد. همانطور که گفته شد این امر مشکل ساز می شود چون ممکن است next hop ناشناخته باشد. اگر از طریق IGP دسترسی وجود داشته باشد که هیچ در غیر این صورت ۳ حالت رخ می دهد. iBGP یاد می گیرد ولی در جدول مسیریابی قرار نمی دهد. iBGP یاد می گیرد ولی به همسایه ها ارسال نمی کند. در `show ip bgp <ip> <length>` به صورت inaccessible است. برای حل این مشکل می توان از دستور `neighbor <ip> next-hop-self` استفاده نمود و یا این که لینک eBGP ها را به iBGP ها ارسال نمود. می توان از static route نیز استفاده نمود ولی توصیه نمی شود.

۶- بسته های به روزرسانی که از طرف eBGP می رسد، iBGP نمی تواند آن را ارسال کند. به طور مثال در transit AS اگر بسته از eBGP وارد شود در transit AS بسته ها iBGP هستند و بعد وارد eBGP دیگر می شود. در این صورت بسته به روزرسانی در این بین از بین می رود، زیرا برای جلوگیری از حلقه همسایه های iBGP اطلاعات iBGP را به یکدیگر ارسال نمی کنند، که این مشکل ساز است. برای حل این مشکل ۱- اگر اولین مسیریاب در transit AS و آخرین مسیریاب با هم همسایه باشند بسته eBGP به مسیریاب آخر می رسد و AS بعدی نیز دریافت می کند. ولی در مسیریابی چون مسیریاب های ما بین در transit AS هیچ اطلاعاتی ندارند (به دلیل عدم ارسال برای حذف حلقه) بسته حذف می شود و به اصطلاح Black Hole ایجاد می شود.



### ۳-۱۲- Prevent Black Hole in BGP

در جدول مسیریابی زیر :

- O router-1 via router-2
- B Net-X via router-1

برای دسترسی به شبکه X باید به مسیریاب شماره ۱ بسته ارسال شود و برای دسترسی به مسیریاب ۱ باید به مسیریاب شماره ۲ بسته ارسال شود و چون بسته آدرس مقصد X دارد، مسیریاب شماره ۲ چون از چنین شبکه‌ای اطلاع ندارد آن را حذف می‌کند. در واقع Black Hole ایجاد می‌شود. برای جلوگیری از آن راه‌های زیر ارائه شده است:

۱- Synchronization: در BGP اگر مسیری از iBGP دریافت شد آن را اعلام نمی‌کند و در جدول مسیریابی قرار نمی‌دهد مگر این که دقیقاً در جدول مسیریابی از طریق IGP وجود داشته باشد. زمانی مفید است که مسیریاب‌هایی در این بین وجود داشته باشد که مشکل‌ساز شود. ولی زمانی که به طور مثال دومسیریاب وجود داشته باشد و مسیریاب‌های پایین دستی IGP داشته باشند و خواسته شود این مسیریاب‌ها از BGP هیچ اطلاعی نداشته باشند این قانون نباید تنظیم شود. دستور آن synchronization است که در BGP وارد می‌شود. اگر این قانون اعمال شود و BGP به IGP ترجمه شد توصیه نمی‌شود و در واقع بار سنگینی وارد شبکه می‌شود.

۲- Full Mesh: اگر همه با هم همسایه باشند و از شبکه‌ای که وارد می‌شود همه اطلاع داشته باشند، مشکل حل می‌شود و Black Hole ایجاد نمی‌شود.

۳- BGP Route Reflector: در CCNP گنجانده نشده است.

۴- BGP Confederation: در CCNP گنجانده نشده است.

### ۳-۱۳- Peer Group

فرض کنید تعداد بسیاری مسیریاب، همسایه‌های بسیار و هر کدام تنظیمات بسیار ولی یکسانی داشته باشند. در این صورت برای تنظیم مسیریاب‌ها زمان بسیاری و CPU بسیاری هنگام دریافت و ارسال بسته‌های update مصرف می‌شود. برای ساده‌تر نمودن آن می‌توان از peer group یا گروه همسایه‌ای استفاده نمود که دارای خاصیت‌های زیر هستند:

۱- می‌توان برای همه همسایه‌ها سیاست یکسان outbound تعریف کرد ولی برای inbound حتماً جداگانه باید تعریف شود.

۲- برای هر گروه، تمام تنظیمات یکسان و برای همه همسایه‌ها یکبار تعریف می‌گردد.

۳- بسته‌های به روزرسانی فقط یکبار ایجاد می‌شوند.

دستور تنظیم آن به صورت زیر است:

تعریف گروه:

```
(Config-router)#neighbor <name of group> peer-group
```

تنظیم دستورات برای همسایه‌ها

```
(Config-router)#neighbor <name of group> <set attribute>
```

عضو کردن همسایه در گروه

```
(Config-router)#neighbor <ip-add> peer-group <name of group>
```

باید توجه داشت که در روش معمول برای هر همسایه، بسته ایجاد، سیاست اعمال و برای هر کدام جداگانه ارسال می‌شود. ولی در این روش یک بسته با توجه به سیاست‌ها ایجاد و برای هر همسایه جداگانه ارسال می‌شود.

### ۳-۱۴- BGP Filtering

فیلتر کردن مسیرها در پروتکل BGP در CCNP ساده است و بر روی هر مسیریاب می‌توان انجام داد. فیلترینگ را می‌توان از جنبه‌های متفاوتی بررسی نمود. در Inbound مسیر قبل از وارد شدن به BGP Table فیلتر شود. در Outbound مسیر قبل از این که به همسایه‌ها اعلام شود فیلتر شود.

بعد از فعال کردن فیلتر باید ارتباط همسایه‌ها reset شود. البته می‌توان از آن جلوگیری نمود که در ادامه توضیح داده می‌شود. در EIGRP ارتباط همسایه‌ها قطع و دوباره به صورت خود کار وصل می‌شود در OSPF بعد از یک ساعت که مسیر فرستاده نشود حذف می‌شود. در IGP دو راه برای فیلتر کردن وجود داشت: ۱- interface ۲- inbound , outbound.

در BGP فقط برای همسایه مشخص می‌توان این عمل را انجام داد که به ۴ طریق صورت می‌گیرد: ۱- distribute-list ۲- prefix-list ۳- filter-list ۴- route-map.

۱. Prefix-list در OSPF به ساخت prefix-list اشاره می‌شود. در BGP به as-path access-list اشاره دارد.

۲. Filter-List ساده است ولی مانند access-list به تنهایی کارایی ندارد و به as-path access-list بر می‌گردد. ورودی آن یک regular expression است که به as-path بر می‌گردد که با علامت ^ شروع و با علامت \$ خاتمه پیدا می‌کند و علامت \_ به معنای space است. به طور مثال:  
(Config)#ip as-path access-list <number> [permit | deny] ^100\_  
اگر \_ نباشد به معنای شروع با ۱۰۰ است.

AS-path access-list فقط برای فیلتر استفاده نمی‌شود بلکه برای بعضی از ویژگی‌های مسیر که از AS-PATH استفاده می‌شود نیز می‌توان استفاده نمود. برای فیلتر کردن تمام مسیرها بجز AS خود مسیریاب با هر AS از \$^ استفاده می‌شود. همانطور که گفته شد برای اعمال فیلتر باید همسایه reset شود که این کار موجب قطع ارتباطات و سربار شبکه است. دو راه حل وجود دارد. Hard: به معنای پاک کردن کل پروسه همسایه است که پیشنهاد نمی‌شود. Soft: به معنای این است که همسایه‌ها ارتباط دارند و فقط مسیرها مجدداً ارسال می‌شوند. می‌توان جدولی ایجاد نمود که قبل از وارد شدن مسیرها به BGP Table ابتدا در این جدول که به shadow table معروف است قرار بگیرند و پس از آن به BGP Table انتقال یابند این امر باعث می‌شود که بعد از اعمال فیلتر معین شود کدام مسیر به BGP Table انتقال یابد. همچنان می‌توان با دستور IN به BGP Table و یا با دستور OUT به Route Table انتقال یا عدم انتقال داد.

```
(Config-router)#neighbor <ip-address> soft-reconfiguration [in | out]
(R1)#clear ip bgp <ip-address> soft [in | out]
(R1)#show ip bgp neighbors <ip-address> [received-routes | advertised-routes]
```

که با دستور soft out تمام مسیرها مجدداً ارسال می‌شود و با دستور soft in درخواست (route refresh message) ارسال داده می‌شود. اگر دستور soft-reconfiguration وارد نشود جدول ایجاد نمی‌شود.

| EIGRP                                                                              | OSPF                                                                                                                           | BGP                                                                                  |
|------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------|
| In= prevent incoming updates from entering EIGRP topology table                    | In = Allow incoming LSA into LSD and flood to peers, however prevent LSA from becoming a route in my own local routing table   | In = prevent updates from entering BGP table (applied against a nbr)                 |
| Out = prevent EIGRP routes in IP routing table from being advertised to EIGRP nbrs | Out = Only used on ASBRs to prevent redistribution of certain protocols/routes into external LSAs (do not create external LSA) | Out = prevent "best BGP route" for a given prefix from being advertised to neighbors |
| Optional = can specify interface                                                   | No interface Allowed                                                                                                           | No interface allowed (applied to BGP neighbors)                                      |
| Can use standard or extended ACL                                                   | Can use standard or extended ACL                                                                                               | standard and/or extended ACLs Allowed                                                |
| Standard = matching on prefix                                                      | Standard = matching on prefix                                                                                                  | Standard = matching on prefix                                                        |
| Ext = match on prefix and nbr sending us the route                                 | Ext = match on prefix and adv-rtr-id of LSA                                                                                    | Ext = match on prefix and subnet mask                                                |

### ۳-۱۵- BGP Path Attribute

پروتکل BGP ویژگی‌هایی دارد که از آنها برای جلوگیری از حلقه و انتخاب بهترین مسیر و ... استفاده می‌کند. هر Path Attribute (PA) یک ویژگی را توضیح می‌دهد. الگوریتم بهترین مسیر در سیسکو: ۱- WEIGHT: مسیری با بیشترین وزن به عنوان بهترین مسیر است که مخصوص شرکت سیسکو است و بر روی هر مسیریابی به صورت local است. ۲- LOCAL\_PREF: مسیری با بیشترین local preference انتخاب می‌شود به صورت پیش فرض ۱۰۰ است. ۳- network: بهتر از redistribute و بهتر از aggregate-address است. ۴- AS\_PATH: مسیر با کمترین طول بهتر است. با دستورات bgp bestpath as-path ignore می‌توان از آن صرف نظر نمود. اگر AS\_SET=1 بود یعنی از طول آن صرف نظر کن. ۵- Origin Type: IGP بهتر از EGP و بهتر از INCOMPLETE است. ۶- Multi-Exit Discriminator (MED): پایین‌ترین MED بهترین مسیر است. ۷- eBGP بهتر از iBGP است. ۸- Next-Hop: مسیریاب بعدی باید در دسترس باشد. کمترین metric در IGP به گره بعدی BGP، بهترین مسیر است.

در ادامه مؤلفه‌های بعدی در صورت انتخاب نشدن بهترین مسیر اجرا می‌شوند:

۹- اگر بهترین مسیر مشخص نشده باشد مسیری که چند مسیر برای رسیدن به آن است در BGP انتخاب می‌شود. ۱۰- اگر دو مسیر external باشند قدیمی‌ترین مسیر بهترین مسیر است. ۱۱- router-id: کمترین router-id بهترین مسیر است. ۱۲- اگر router-id یکسان باشد کمترین طول لیست کلاستر بهترین مسیر است ۱۳- کمترین آدرس همسایه بهترین مسیر است.

بهترین مسیر به ترتیب زیر انتخاب می‌شود:

با تغییر weight, local preference, as-path روی outbound انجام می‌شود و با تغییر as-path, MED تغییر بر روی inbound انجام می‌شود. مانند فیلترینگ باید ارتباط با همسایه‌ها ریست شود. ۱- Weight: وزن به عنوان قسمتی از بسته update نیست و به صورت محلی و فقط برای تصمیم‌گیری خود مسیریاب محلی است. به عنوان خصوصیتی برای مسیریاب‌های سیسکو است و به مسیریاب‌های دیگر اعلام نمی‌شود. به عبارت دیگر شبیه به AD است. حداکثر مقدار آن  $2^{16}-1$  است. به صورت پیش‌فرض برابر صفر است. در واقع مسیرهای IGP یادگیری شده را صفر در نظر می‌گیرد و مسیرهای locally originate برابر 32768 است. بر روی مسیرهای ورودی set می‌شود و تاثیر آن بر روی مسیرهای خروجی است. مقدار default را نمی‌توان تغییر داد و در route-map می‌توان از set weight استفاده نمود.

(Config-router)#neighbor route-map <per prefix>

(Config-router)#neighbor weight (all routes from this neighbor)

۲- Local Preference: بر عکس weight مقدار local preference در بسته update قرار دارد و به صورت محلی نمی‌باشد و به مسیریاب‌های دیگر اعلام می‌شود. این قسمت از بسته به صورت اختیاری در بسته update قرار می‌گیرد. در واقع بسته BGP Update در قسمت attribute به دو صورت انجام می‌گیرد. اجباری که as-path, next-hop و origin code قرار دارد. اختیاری که به دو صورت transitive که در همه AS‌ها باقی می‌ماند و non-transitive که فقط در همان AS باقی می‌ماند. Local preference به صورت پیش‌فرض 100 و جزء Non-Transitive است. حداکثر مقدار آن  $2^{32}-1$  است.

(Config-router)#bgp default local-preference <number>

برای این که بتوان در بسته بین AS‌ها باقی بماند باید از route-map -> set local-preference استفاده نمود.

۳- AS\_PATH length: با افزایش طول AS-PATH می‌توان انتخاب بهترین مسیر را تغییر داد و تاثیر بر روی Inbound و Outbound گذاشت. به این عمل as-path prepending یا همان افزایش طول گویند. این امر فقط برای افزایش طول ابتدای آن در AS نه جای دیگر و بهتر است که همان AS شروع قرار گیرد تا

مشکلی (تشخیص حلقه) پیش نیاید. با استفاده از دستور route-map در قسمت set as-path می توان آن را تغییر داد. در قسمت neighbor route-map [in | out] می توان آن را اعمال نمود. این عمل برای آن است که دسترسی به AS دیگری وجود ندارد و بتوان در AS خود تغییراتی ایجاد نمود. به اجبار AS های دیگر تحت تاثیر قرار می گیرند تا مسیریاب خود شما به عنوان primary , secondary قرار گیرد.

۴- Origin: مانند پروتکل IGP که برای تعریف مسیر از network و redistribute استفاده می شود در BGP نیز این اتفاق می افتد. Network command: از مسیرهای IGP یاد می گیرد که در جدول مسیریابی قرار دارند که به صورت igp نمایش داده می شود. Redistribute command: که به صورت ؟ یا unknown یا incomplete نمایش داده می شود. برای تغییر آن در route-map در قسمت set origin می توان اعمال نمود. اولویت آن به ترتیب ?>egp>igp است.

۵- MED: یا multi exit discriminator بدین معنی است که اگر تعداد لینک های متصل به یک AS از طریق یک و یا چند مسیریاب بیشتر از AS دیگر باشد اولویت مسیر بالاتری وجود دارد. این ویژگی در بسته های update به صورت اختیاری است و non-transitive است. تاثیر آن بر روی inbound است. فقط در یک AS باقی می ماند و هر چه کمتر باشد اولویت بالاتری دارد. مقدار پیش فرض صفر و آخرین مقدار  $2^{32-1}$  است. در کل مقدارهای weight و local-pref هر چه بالاتر باشند اولویت بالاتر است. در metric , MED هر چه پایین تر باشند اولویت بیشتر است. برای تغییر آن در route-map در قسمت set metric می توان اعمال نمود و برای اعمال neighbor route-map out تنظیم می شود. در آخر باید با دستور bgp always-compare-med برای اعمال آن تنظیم گردد. و با دستور show ip route یا show ip bgp می توان آن را مشاهده نمود.

### ۳-۱۶- BGP Maximum-paths

به حداکثر تعداد مسیر که از یک شبکه در جدول مسیریابی می تواند قرار گیرد گویند. از آن برای load balancing استفاده می شود. همچنین این خصوصیت را با maximum-paths می توان تعیین نمود. به طور مثال net-x فقط ۴ مسیر را در جدول مسیریابی داشته باشد. در BGP اگر تمامی ویژگی های به اختصار N WLLA OMNI یکسان باشد، می تواند چندین مسیر را در جدول مسیریابی داشته باشد، ولی تنها یک مسیر را اعلام می کند. به صورت پیش فرض این عدد ۳ است و load balancing به صورت محلی است و اعلام نمی شود.

### ۱۷-۳ BGP Multi-Protocol

نام تخصصی BGP معروف است به MBGP یا multi-protocol bgp بدین معنا که در لایه ۳ شبکه به غیر از پروتکل IPv4 و IPv6 دیگر پروتکل‌هایی که در لایه ۳ وجود دارد و یا در آینده می‌آید از آن حمایت می‌کند یعنی به گونه‌ای طراحی شده است که حال و آینده را را در بر می‌گیرد. در واقع کار NLRI این عمل است. از پروتکل‌هایی که در حال حاضر حمایت می‌کند: IPv4 unicast and multicast و IPv6 unicast و MPLS VPN است. روش کار بدین شکل است که هنگام ارسال open message در قسمت capability به همسایه پروتکل‌هایی که حمایت می‌کند را نشان می‌دهد و برای این که چندین پروتکل را حمایت کند از address families استفاده می‌شود.

در حقیقت به صورت پیش فرض از IPv4 استفاده می‌شود و تنظیم کردن آن قابل مشاهده نیست. با تعریف address families این قابلیت ایجاد می‌شود که با پروتکل‌های مختلف با همسایه‌ها صحبت شود. راه‌های مختلفی برای ارتباط وجود دارد: ۱- دو BGP Session برای هر پروتکل (IPv4, IPv6) به صورت جداگانه برای ارتباط و به روز رسانی ایجاد می‌شود. این امر برای redundancy خوب است. ۲- یک BGP Session برای ارتباط وجود داشته باشد ولی برای به روز رسانی از هر دو پروتکل حمایت می‌کند. در واقع یا session IPv4 و بسته‌ها IPv6 را بفهمد و یا session IPv6 و بسته‌ها IPv4 را بفهمد. که این دو روش بستگی به طراحی شبکه دارد و مزیت بر یکدیگر نیستند.

| When you type this                 | BGP really sees it like this       |
|------------------------------------|------------------------------------|
| !                                  | Router bgp 12                      |
| Router bgp 12                      | Neighbor 1.2.1.1 remote-as 200     |
| Neighbor 1.2.1.1 remote-as 200     | !                                  |
| Network 20.20.0.0 mask 255.255.0.0 | Address family ipv4 unicast        |
| !                                  | Neighbor 1.2.1.1 activate          |
|                                    | Network 20.20.0.0 mask 255.255.0.0 |

برای فهماندن همسایه از طریق IPv6 باید از Route-map استفاده شود. البته اگر به صورت connected با همسایه باشد نیازی نیست و برای همسایه با فاصله استفاده می‌شود. فیلد AFI در بسته به روز رسانی address family indicator است که نوع capability را نشان می‌دهد. SAFI نیز sub address family indicator است. تمامی این تنظیمات زمانی درست است که ارتباط IPv4, IPv6 بین همسایه‌ها باشد. در IPv6 دو نوع آدرس مشاهده می‌شود که محلی و عمومی می‌باشد.

## ۴- IP Service Level Agreement (IP-SLA)

ابزاری است که بر روی مسیریاب یا سوئیچ‌های لایه ۳ تنظیم می‌شود تا با استفاده از آن سرویس دهی را مشاهده و بررسی نمود. مسیریاب ترافیکی را بررسی، هدایت و آن را به مقصد یا ISP ارسال می‌کند. با استفاده از نتیجه آن یا reply مواردی مانند delay را اندازه‌گیری می‌کند و آن را به سیستم مانیتورینگ، مدیریت شبکه یا سیستم مدیریتی SNMP ارسال می‌کند. به این امر IP-SLA گفته می‌شود.

IP-SLA می‌تواند با سرویس‌های دیگر ارتباط برقرار کند به طور مثال با HSRP (Hot Standby Routing Protocol) به‌صورتی که اگر ارتباط یک مسیریاب با سرویس دهند قطع گردید، IP-SLA این را موضوع را تشخیص داده و به HSRP اطلاع می‌دهد و سپس مسیریاب دیگر وارد مدار می‌شود. یا با Policy Based Routing (PBR) به‌صورتی که اگر ارتباط برقرار بود آن‌گاه policy اعمال شود در غیر این صورت اعمال نشود. IP-SLA جزئیات زیادی دارد و CPU زیادی مصرف می‌کند.

Operation IP-SLA: به پروتکل یا عملیاتی که ارتباط را بررسی می‌کند گویند. مانند Ping یا HTTP و هر پروتکل دیگری که نوع بسته را مشخص می‌کند، استفاده می‌شود. باید توجه داشت که در یک زمان می‌توان از Operation‌های متفاوتی استفاده نمود.

Setting IP-SLA: تنظیماتی که بر روی Operation انجام می‌شود. مانند اینکه بسته چه زمانی، با چه فاصله زمانی و یا چه مقدار تأخیر داشته باشد.

IP-SLA یک فرستنده و گیرنده دارد. به گیرنده SLA Responder گویند. گیرنده گیرنده بسته به نوع عملکرد و پروتکل‌های پوشش دهنده می‌تواند Host و یا Router باشد.

تنظیمات مربوط به نوع خاص ICMP-Based IP-SLA:

```
(Config)#ip sla <sla-ops-number>
```

عبارت sla-ops-number شماره پروسه است.

```
(Config-ip-sla)#icmp-echo
```

در قسمت icmp-echo جزئیات را می‌توان وارد نمود.

```
(Config-ip-sla)#frequency <second>
```

عبارت frequency <second> اطلاعات مربوط به زمان ارسال یا تکرار آن است.

موارد بالا هنوز اجرا نمی‌شوند و باید زمان‌بندی شود.

```
(Config)#ip sla schedule <sla-ops-number>
```

```
(R1)#show ip sla configuration
```

```
(R1)#show ip sla statistics
```

PBR و HSRP به صورت مستقیم با ip-sla کار می کند. حال اگر خواسته شود نتیجه ارسال بررسی و از آن استفاده شود، باید از Track Object استفاده نمود. به طور مثال بعد از تشخیص خطا اطلاعات به سیستم مانیتورینگ ارسال شود. باید از عبارت delay برای مشخص نمودن مدت زمان خطا استفاده نمود.

```
(Config)#track <object-number> ip sla <sla-ops-number> <state | reachability>
```

```
(Config-track)#delay <down sec | up sec>
```

```
(Config)#route-map
```

```
(Config-route-map)#set ip next-hop verify-availability <ip-add> <sequence-number>  
track <track-id>
```

## ۵- Policy-Based Routing

سیستم PBR در وضعیت معمول نگاه آن به مقصد است و با مبدأ ارسالی ندارد. در PBR بسته ها کنترل می شود و قبل از مسیریابی در وضعیت normal هستند. PBR از route-map استفاده می کند. جزئیات route-map به شرح زیر ارائه است:

Route-map: این ابزار یا با PBR استفاده می شود یا مقصدی که بسته ارسال می شود را تعیین می کند. مقصد بسته یا IP (برای لایه ۲ یا وجود بیشتر از دو مسیریاب بین راه) و یا Interface (ارتباط p2p باشد) است.

Match command: این دستور به دو صورت است: ۱- استفاده از ACL ۲- تعیین حجم بسته بر اساس بایت.

Access-list: این دستور نیز به دو صورت است یا permit که اجازه می دهد PBR اجرا شود و یا deny که بسته drop نمی شود فقط اجازه اجرا شدن دارد.

Set command: این دستور ممکن است به روش های زیر انجام گیرد:

۱. Set ip next-hop: ممکن است همه پخش یا چند پخش یا برای Fast Ethernet یا Gig استفاده شود.

بعد از آن می توان بیشتر از یک آدرس وارد نمود.

۲. Set ip default next-hop: اگر مسیر در جدول مسیریاب نبود، آن گاه از این دستور استفاده می شود.

اگر مسیر در جدول مسیریابی نباشد آنگاه بر اساس PBR عمل می کند.

۳. Set interface: اگر برای اینترفیس اتفاقی بیافتد، این دستور بسته‌ها را برای بعدی ارسال می‌کند. می‌توان بیشتر از یک اینترفیس وارد نمود.

۴. Set default interface: اگر مسیر در جدول مسیریاب نباشد، آن گاه از این دستور استفاده می‌شود و به اینترفیس مشخص شده ارسال می‌کند.

برای اجرای PBR باید وارد اینترفیس شد و با دستور ip policy route-map آن را اجرا نمود.

باید توجه داشت که PBR برای بسته‌های ingress است و بسته‌هایی که خود مسیریاب آن را ایجاد می‌کند و egress است، کاربردی ندارد. برای بسته‌های locally generate by router itself باید از دستور ip local policy route-map استفاده نمود. با استفاده از route-map نیز می‌توان بین DSCP که برای کیفیت سرویس است را تغییر داد که با دستور set ip precedent و یا set ip tos می‌توان انجام داد.

## ۶- Network Time Protocol (NTP)

زمان برای بعضی از پروتکل‌ها بسیار حائز اهمیت است. اهداف استفاده از زمان شامل موارد زیر است: ۱- logging output در خروجی گرفتن از لاگ‌های ورود به سیستم اهمیت دارد. ۲- debug output در بررسی و تجزیه و تحلیل خروجی‌های سیستم اهمیت دارد. ۳- user show command در دستورهای نمایشی و راهنمایی کاربر حائز اهمیت است. ۴- network management / report tools در ابزارهای مدیریتی و گزارش‌گیری حائز اهمیت است. اغلب زمان با استفاده از باتری محافظت می‌شود. زمان را به روش‌های متفاوتی می‌توان تنظیم نمود:

۱- Manual: به صورت دستی قابل تنظیم است.

۲- NTP: این پروتکل زمان یا تاریخ را از سرور دریافت، تنظیم و به بقیه ارسال می‌کند.

۳- Simple Network Time Protocol (SNTP): فقط زمان و تاریخ را دریافت و تنظیم می‌کند.

۴- VINES(Virtual Integrated Network Service)

تفاوت NTP version 3,4: در ورژن ۳ فقط IPv4 را حمایت می‌کند ولی در ۴ هم IPv6 , IPv4 و

همچنین authentication وجود دارد.

NTP از پورت مبدأ و مقصد UDP 123 استفاده می‌کند. نحوه گرفتن زمان در NTP یا منبع زمان به صورت Atomic Clock، GPS، Radio Clock و Other Network Device است.

Stratum Level: به تعداد فاصله‌ها از منبع زمانی گویند. به طور مثال اگر سرور مستقیماً به منبع زمانی متصل باشد stratum-1 و اگر با فاصله یک hop از stratum-1 باشد stratum-2 گویند. در واقع سطح رسیدن به منبع زمانی است. به صورت پیش‌فرض اگر stratum level مشخص نشود stratum-8 تعیین می‌گردد.

حالت‌های node در NTP:

- ۱- NTP server یا NTP Master: سرور ارسال زمان و تاریخ است.
- ۲- NTP Client: نودی است که زمان و تاریخ را به صورت unicast از سرور دریافت می‌کند.
- ۳- NTP Peers یا Symmetric Mode: سنکرون کردن زمان بین سرورها است.
- ۴- NTP Broadcast / Multicast: ارسال بسته از طرف سرور به صورت Broadcast / Multicast است.

دستورات مربوط به تنظیم NTP server به شرح زیر است:

```
(R1)#clock set <time> <date>  
(Config)#ntp master [stratum]
```

موارد زیر اختیاری است

```
(Config)#ntp peer <ip-add>  
(Config)#ntp broadcast
```

دستورات مربوط به تنظیم NTP client به شرح زیر است:

```
(Config)#ntp server <ip-add>  
(Config)#ntp broadcast client
```

```
(R1)#show ntp status  
(R1)#show ntp associations
```

امنیت NTP نیز عامل مهمی در امنیت شبکه است. به دو صورت ACL و Authentication قابل استفاده است. در احراز هویت با استفاده از MD4 Hash و چندین کلید نیز می‌توان استفاده نمود. دستورات تنظیم آن به صورت زیر است:

```
(Config)#ntp authentication-key <key-number> md5 <key-string>  
(Config)#ntp trusted-key <key-number>  
(Config)#ntp authenticate
```

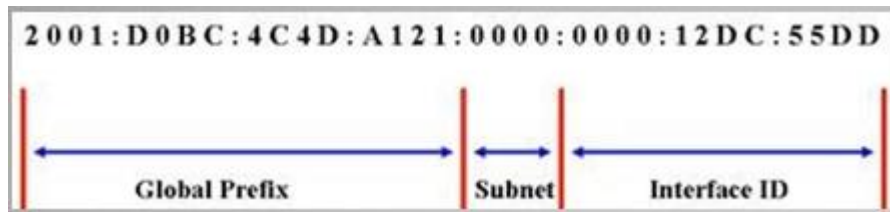
موارد بالا در سرور تنظیم می‌شود و در client نیز به همین ترتیب اما با تفاوت زیر:

```
(Config)#ntp server <ip> key <key-number>
```

## ۷- IPv6 Introduction

در این قسمت توضیحی در مورد IPv6 داده می‌شود. توجه داشته باشید که این قسمت پیش‌نیاز قسمت‌های بعد است. در ادامه به ویژگی‌های IPv6 پرداخته می‌شود.

۱. تعداد بیت برابر ۱۲۸ است.
۲. تعداد فیلد هدر آن نسبت به ipv4 کاهش یافته است.
۳. امنیت آن از لحاظ encryption و authentication افزایش یافته است.
۴. در ذات آن IPsec گنجانده شده است و به صورت مازول نیست.
۵. پویایی یا mobility وجود دارد. بدان معنی است که بدون هیچ مشکلی از شبکه خارج و داخل شبکه دیگر می‌شود.
۶. افزایش عملکرد Performance و کیفیت سرویس QoS گنجانده شده است.
۷. خاصیت auto configuration دارد که بدون نیاز به DHCP به صورت خودکار IP تنظیم می‌شود.



در IPv4 از ۴ قسمت ۸ بیتی که بین ۰ تا ۲۵۵ مقدار دهی می‌شود و به صورت octet است، استفاده می‌شود. در IPv6 از ۸ قسمت ۱۶ بیتی که بین ۰ تا  $2^{16}$  مقدار دهی می‌شود و به صورت hexadecimal است. در شکل بالا Interface ID یا همان Host ID ۶۴ بیت و بقیه نیز ۶۴ بیت است.

۱. FP (Former Prefix): به ۳ بیت اول آدرس FP می‌گویند. نوع آدرس را مشخص می‌کند. ممکن است آدرس global یا private باشد. به عنوان مثال FP در آدرس‌های Global Unicast 001 است.
۲. TLA\_ID (Top Level Aggregator Identifier): به ۱۳ بیت بعدی TLA می‌گویند. محل جغرافیایی یک IP Address را مشخص می‌کند. برای تشخیص اینکه مبدأ یک مسیر کجا بوده و از کجا می‌آید.
۳. NAP (Network Access Point): در سطح دنیای اینترنت ۱۲ تا NAP وجود دارد. که ۸ بیت در نظر گرفته می‌شود.

۴. NLA\_ID (Next Level Aggregator Identifier): به ۲۴ بیت بعدی NLA می گویند.
۵. SLA\_ID (Site Level Aggregator Identifier): به ۱۶ بیت بعدی SLA می گویند. که این سه مورد قبل محدوده های اینترنت هستند.

در IPv6 سه نوع ارسال بسته وجود دارد:

۱. Unicast: که به چند صورت link-local، آدرس های Private و محلی است که به صورت خودکار ساخته می شود در نبود DHCP که با FE80 شروع می شود. Global، مانند IPv4 در اینترنت است که به صورت IP Valid شناخته می شود. Unique Local Address: همان آدرس ها invalid برای استفاده به صورت محلی است که با FC و FD شروع می شوند. و دیگر انواع آن IPX, NASP, special هستند.
  ۲. Multicast: که با شروع FF در ابتدای IP است.
  ۳. Any cast: که به unicast آدرس بستگی دارد و به معنی ارسال یک به یک از آدرس های زیاد است.
- می توان در بیشتر تنظیمات سیسکو بجای IP از IPv6 استفاده نمود.

## ۷-۱- RIPng, IPv2

در واقع RIPng یا next generation با IPv6 کار می کند که شباهت ها و تفاوت های آن با RIP v2 به صورت زیر است:

- ۱- شباهت: استفاده از UDP، Distance Vector، AD=120، VLSM، Split Horizon، Poison Revers، 30 Sec Full Update، Hop Count Metric، Metric 16=∞، پشتیبانی از route tag و multicast Update Destination.
- ۲- تفاوت RIP v2 و RIPng: UDP=521, UDP=520، no auto summary IPv6، به خاطر نداشتن subnetting در IPv6، destination multicast address=FF02::9 یا همان 224.0.0.9، Link Local Next Hop، IPv6 AH/ESP authentication.

برای تنظیم RIPng ابتدا باید IPv6 Routing و سپس بر روی interface آن را فعال نمود. در IPv4 فقط یک RIP Process می توان بر روی همان مسیریاب ایجاد نمود. در IPv6 می توان از process name های متفاوت استفاده نمود. در IPv4 به صورت خودکار بسته ها بر روی interface ارسال و دریافت صورت می گیرد ولی در IPv6 بین interface ها ارتباط برقرار نمی شود، مگر اینکه دستور ipv6 unicast-routing که در سطح global است، تنظیم شود.

(Config)#ipv6 unicast-routing

```
(Config)#ipv6 router rip <name>  
(Config-if)#ipv6 address <ipv6>  
(Config-if)#ipv6 rip <name> enable
```

## ۲-۷ EIGRP IPv6

۱. همانطور که مشخص است پروتکل EIGRP ipv6 را اعلام می کند.
  ۲. از آدرس link-local neighbor به عنوان next-hop استفاده می شود. تفاوتی که با IPv4 دارد این است که در IPv4 همسایه ها یا مبدأ و مقصد باید در یک شبکه IP باشند، اما در IPv6 آدرس link-local باید در یک شبکه، ولی globally می تواند در شبکه متفاوت باشد. باید توجه داشت اگر آدرس تغییر نمود ارتباط همسایه ها بدون هیچ مشکلی ادامه می یابد.
  ۳. آدرس multicast = FF02::A استفاده می شود.
  ۴. در IPv6 از IPsec به عنوان authentication استفاده می شود.
  ۵. Auto summary در IPv6 معنایی ندارد چون subnetting وجود ندارد.
  ۶. از پروتکل لایه ۳ با شماره 88 استفاده می کند.
  ۷. از successor و feasible successor استفاده می شود.
  ۸. از DUAL برای به روزرسانی استفاده می شود.
  ۹. از همان metricهای مشابه IPv4 استفاده می شود.
  ۱۰. مؤلفه  $\text{infinity metric} = 2^{32-1}$  است.
- برای تنظیم آن ابتدا ipv6 unicast-routing فعال می شود. سپس دستور <ASN> ipv6 router EIGRP بر روی اینترفیس مورد نظر، فعال می شود. انتخاب router-id به همان شکل قبل و ۳۲ بیتی است و در آخر بر روی ipv6 EIGRP process دستور no shutdown وارد می شود.

## ۳-۷ OSPF v3 IPv6

- شباهت های آن با OSPF v2 به شرح زیر است:
۱. از protocol type 89 استفاده می شود.
  ۲. Link state است.
  ۳. از VLSM پشتیبانی می کند.
  ۴. LSA flooding دارد و تاریخ انقضاء مسیرها هر یک ساعت است.
  ۵. ساختار آن بر مبنای ناحیه است.
  ۶. مانند قبل packet types با کمی تفاوت وجود دارد.

۷. LSID ۳۲ بیتی دارد.

۸. از interface به عنوان metric استفاده می شود.

۹. مؤلفه infinity metric age = 3600 sec است.

۱۰. دوره ارسال تمامی مسیرها ۳۰ دقیقه است.

۱۱. آدرس FF02::5 , FF02::6 multicast = است.

۱۲. از IPv6 AH/ESP به عنوان authentication استفاده می کند.

تفاوت های با OSPF v2 به شرح زیر است:

۱. همسایه ها می توانند در یک subnet از شبکه نباشند.

۲. در IPv6 OSPF مفهومی به نام instance وجود دارد و در v3 از آن پشتیبانی شده است ولی در

دستگاه های cisco از آن پشتیبانی نمی شود. بدین معنا است که شبکه های OSPF را از هم جدا کرده

و هر کدام را یک instance نام گذاری می کند. حال هر مسیریاب می تواند به یک یا چند instance

دسترسی داشته باشد. این کار برای جدا کردن دسترسی ها بین مشتریان است. به صورت پیش فرض

instance-id = 0 است. اگر متفاوت باشد اعلامی در دستگاه نمی شود و فقط از طریق show run

قابل مشاهده است.

| LSA Name              | LS Type Code | Flooding Scope   | LSA Function Code |
|-----------------------|--------------|------------------|-------------------|
| Router LSA            | 0×2001       | Area Scope       | 1                 |
| Network LSA           | 0×2002       | Area Scope       | 2                 |
| Inter-Area-Prefix-LSA | 0×2003       | Area Scope       | 3                 |
| Inter-Area-Router-LSA | 0×2004       | Area Scope       | 4                 |
| AS-External-LSA       | 0×2005       | Area Scope       | 5                 |
| Group-Membership-LSA  | 0×2006       | Area Scope       | 6                 |
| Type-7-LSA            | 0×2007       | Area Scope       | 7                 |
| Link-LSA              | 0×2008       | Link-Local Scope | 8                 |
| Intra-Area-Prefix-LSA | 0×2009       | Area Scope       | 9                 |

## LSA Type-۱-۳-۷

۱. Router LSA: مسیریاب اعلام می کند که خود یک stub link است.
۲. Network LSA: این بسته در IPv4 توسط designated router ایجاد می شود. اگر آدرس IP اینترفیس تغییر کند درخت SPF نیز تغییر می کند و down می شود. این موضوع در IPv6 تغییر نمی کند و بسته Type-9 ایجاد می شود.
۳. Inter-Area-Prefix-LSA: همان summary در پروتکل OSPF است.
۴. Inter-Area-Router-LSA: همان مسیریاب ASBR است.
۵. AS-External-LSA: همان مسیریاب ABR است.
۶. Group-Membership-LSA: برای multicast استفاده می شود.
۷. Type-7-LSA: همان Type-7 در پروتکل OSPF است.
۸. Link-LSA: آدرس های Link-Local و Global به مسیریاب های DR و BDR ارسال می شود.
۹. Intra-Area-Prefix-LSA: بعد از ارسال Type-8، مسیریاب DR به بقیه مسیریاب ها نوع Type-8 را اعلام می کند.

تنظیم OSPF V3 نیز مانند پروتکل های قبل است. برای تنظیم آن ابتدا ipv6 unicast-routing فعال می شود. سپس `ipv6 router ospf <Process-id>` بر روی اینترفیس مورد نظر فعال می شود. انتخاب router-id نیز به همان شکل قبل و ۳۲ بیتی است. Summarization نیز بر روی مسیریاب های ABR و ASBR انجام می شود. در آخر نیز دستور `area <ASN> range <ipv6-prefix/length>` وارد می شود.

## IPv6 Static Route & Access-List -۴-۷

دستور تنظیم Static Route به شرح زیر است:

```
(Config)#ipv6 route <prefix/length> <interface | next-hop> <AD> tag <tag0value>
```

باید توجه داشت IPv6 access-list فقط عبارت نام می پذیرد. همچنین فقط extend دارد و بر روی یک access-list می توان صف تعریف نمود.

```
(Config)#ipv6 access-list <name>
```

```
(Config-ipv6-acl>#<permit | deny> <protocol> .....
```

```
(Config-if)#ipv6 traffic-filter <name> <in | out>
```

## ۵-۷ - Redistribute IPv6 & IPV4

در redistribute بین دو پروتکل هیچ کدام از جدول‌های topology و database یکدیگر اطلاعی ندارند، بلکه از جدول routing اطلاعات را استخراج می‌کنند. جدول مسیریابی IPv4 و IPv6 کاملاً از هم جدا هستند. در IPv6 نیز route-map با کمی متفاوت مانند match می‌توان استفاده نمود. AD در IPv6 همانند IPv4 است. مکانیزم‌های جلوگیری از حلقه شامل AD، Tag و filter-list است، همانند IPv4 است. نحوه استفاده از route-map در IPv6 به شرح زیر است:

۱. دستور route-type برای EIGRP Internal , OSPFv3 , RIPng وجود ندارد. ولی برای EIGRP External وجود دارد.
۲. دستور match فقط ACL و prefix-list قبول می‌نماید و برای IPv6 حتماً src=IPv6 و des=any باشد.
۳. نیاز به دستور subnet نیست چون IPv6 به صورت class full است.
۴. برای فعال کردن redistribute connected بر روی تمامی لینک‌ها باید از include-connected استفاده نمود.
۵. Loop back نیز به عنوان هاست شناخته می‌شود و باید از دستور network point-to-point برای تغییر آن استفاده نمود.
۶. از tag پشتیبانی می‌شود.

## ۶-۷ - IPv4 & IPv6 translation

در کل دو راه برای ترجمه یا تبدیل نسخه‌ها وجود دارد. Native به معنی استفاده ابتدا تا انتها از یک پروتکل یا tunneling که IPv6 داخل بسته IPv4 قرار گیرد و ارسال شود. برای صحبت کردن و تبدیل به همدیگر سه راه وجود دارد. در ادامه به شرح این ۳ راه پرداخته می‌شود.

۱. Dual IPv4, IPv6 Stacks: دو جدول جداگانه برای هر دو وجود داشته باشد. هر زمان از هر کدام نیاز شد استفاده شود. در لایه application می‌توان آن را تعریف نمود که هر برنامه از کدام استفاده نماید. به صورت پیش فرض در سیستم عامل ویندوز ترجیح داده می‌شود از IPv6 در صورت وجود استفاده نماید.
۲. Tunneling: در ادامه به شرح توضیح داده می‌شود.
۳. NAT Protocol Translation (NAT-PT): به معنای حذف کامل هدر IPv4 و ایجاد هدر IPv6 است.

## ۷-۷ Tunneling

به معنای بسته‌بندی یا encapsulate کردن IPv6 داخل IPv4 است. در صورت استفاده از آن سربار تنظیماتی IPv6 پایین می‌آید. به غیر از استفاده در مسیریاب حتی در هاست روش‌هایی وجود دارد. به طور مثال: 6in4, Tredo, ISATAP (Intra-Site Automatic Tunnel Addressing Protocol)

به دو صورت قابل انجام است:

۱. Point-to-Point: در داخل مسیریاب tunneling به معنای ایجاد یک اینترفیس منطقی مانند loopback

است که به آن Tunnel Interface گفته می‌شود.

۲. Point-to-Multipoint: مقصد تنظیم نمی‌شود و آدرس مقصد از IPv4 یا IPv6 استخراج می‌گردد.

نمونه از انواع Tunneling به شرح زیر است. همه از شماره پروتکل 41 که IPv6 encapsulation است، استفاده می‌نمایند.

۱. Manually Configured: به صورت Point-to-Point است و به نام MCT (Manually

Configure Tunnel) استفاده می‌شود. ابتدا یک اینترفیس Tunnel و یک لینک Virtual Point-

to-Point بین مبدأ و مقصد ایجاد می‌گردد. می‌توان بر روی آن نیز IPv6 IGP route اجرا نمود.

Generic Routing Encapsulation (GRE) نیز شبیه به همین است با این تفاوت که د هدر GRE

اضافه می‌گردد ولی در MCT هیچ هدر خاصی اضافه نمی‌گردد.

روش طراحی و تنظیم:

I. یک مسیریاب به عنوان مترجم انتخاب و مبدأ و مقصد مشخص شود. بهتر است از آدرس

Loopback به عنوان مبدأ استفاده شود.

II. یک Tunnel Interface که مقدار آن Integer است، ایجاد شود.

(Config)# interface tunnel <number>

III. مبدأ آن با دستور زیر تعیین شود.

(Config-int)# tunnel source <interface-type> <if-number | ipv4-add>

IV. مقصد آن با دستور زیر تعیین شود.

(Config-int)# tunnel destination <interface-type> <if-number | ipv4-add>

V. Mode tunnel که به صورت پیش فرض GRE تعیین شود. در غیر این صورت برای MCT

دستور زیر وارد شود.

(Config-int)# tunnel mode IPv6ip

VI. آدرس IPv6 به Tunnel interface داده شود.

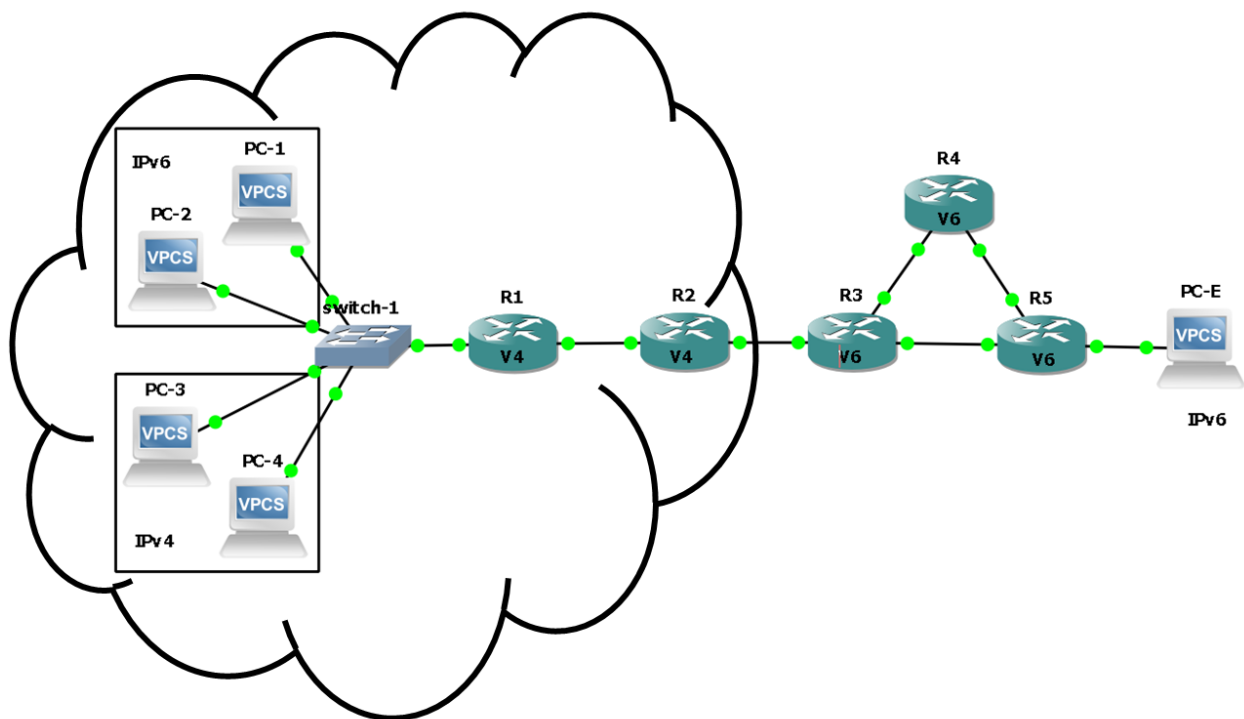
VII. مسیریابی IGP فعال شود. ابتدا باید مسیر پیش فرض <::/0> <tunnel> ipv6 route static و سپس در داخل tunnel پروتکل IGP را فعال نمود.

برای مشاهده جزئیات می توان از دستور زیر استفاده نمود.

(R1)#show interface tunnel

۲. Generic Routing Encapsulation (GRE): از دستور tunnel mode gre ip استفاده می شود.

این دستور باید در دو طرف یکسان باشد. در ادامه به شرح توضیح داده می شود.



۳. Intra-Site Automatic Tunnel Addressing Protocol (ISATAP): اگر بر روی هاست اجرا

شود دو جدول یا dual stack ایجاد می گردد. مشکل جایی پیش می آید که هاست از IPv6 استفاده نماید ولی default gateway router از IPv4 پشتیبانی کند. این حالت برای داخل شبکه محلی طراحی شده است ولی خارج از آن نیز می توان استفاده نمود. شباهت های که بین 6to4 و ISATAP است شامل: ۱- هر دو دارای dual stack می باشند. ۲- هر دو از embedded ipv4 استفاده می کنند با این تفاوت که در 6to4 از آدرس 2002:x:y:: استفاده می شود ولی در ISATAP از آدرس FE80::0000:5EFE:x:y استفاده می شود. ISATAP به صورت خودکار بر روی هاست یا

مسیریاب ایجاد می‌شود که از IPv4 موجود استفاده می‌کند. از پروتکل شماره 41 در بسته‌های IPv4 استفاده می‌کند تا IPv6-ISATAP را منتقل کند. در واقع شماره پروتکل 41 در تمامی پروتکل‌های tunneling استفاده می‌شود. مراحل انجام ISATAP به شرح زیر است:

۱. ابتدا باید به نام "ISATAP" رکوردی در DNSv4 ایجاد شود.
۲. هاست درخواستی برای دریافت ip isatap و پس از دریافت آدرس بسته‌ای به آن مسیریاب ارسال می‌کند.
۳. پس از دریافت بسته توسط مسیریاب بسته SLAAC (State Less Address Auto Configuration) به هاست ارسال می‌شود که در آن ipv6 برای استفاده در محیط global و بعضی اطلاعات تماسی وجود دارد. که این بسته router advertisement است. این عمل شبیه DHCP است ولی با این تفاوت که اطلاعاتی مانند leased time و دیگر موارد در آن نیست.

دلیل تخصیص IP به tunnel interface به شرح زیر است:

۱. تمامی اینترفیس‌ها چه فیزیکی و چه منطقی بنا به استفاده از IPv4 و IPv6 باید IP داشته باشند. در واقع باید پروتکل ip بر روی آنها فعال باشد.
۲. در ISATAP برای ارسال ipv6 به هاست نیاز به ipv6 دارد تا در همان range آدرس تخصیص یابد.

دستورات برای تنظیم به شرح زیر است:

```
(Config)#interface tunnel 0
(Config-if)#ipv6 address <ipv6-prefix> eui-64
```

برای تنظیم دو انتخاب وجود دارد یا به صورت دستی از serial interface استفاده شود یا از interface ID که به صورت خودکار قسمت ۵ و ۶ را برابر 0000:5EFE و قسمت ۷ و ۸ را برابر آدرس ipv4 قرار می‌دهد که از tunnel source می‌گیرد. باید توجه داشت که IPv6 باید global باشد تا بتوان با خارج از شبکه ارتباط داشت و باید آدرس شبکه 64/ داده شود.

```
(Config-if)#no ipv6 nd suppress-ra
```

این دستور برای ارسال router advertise است که به صورت پیش فرض غیر فعال است.

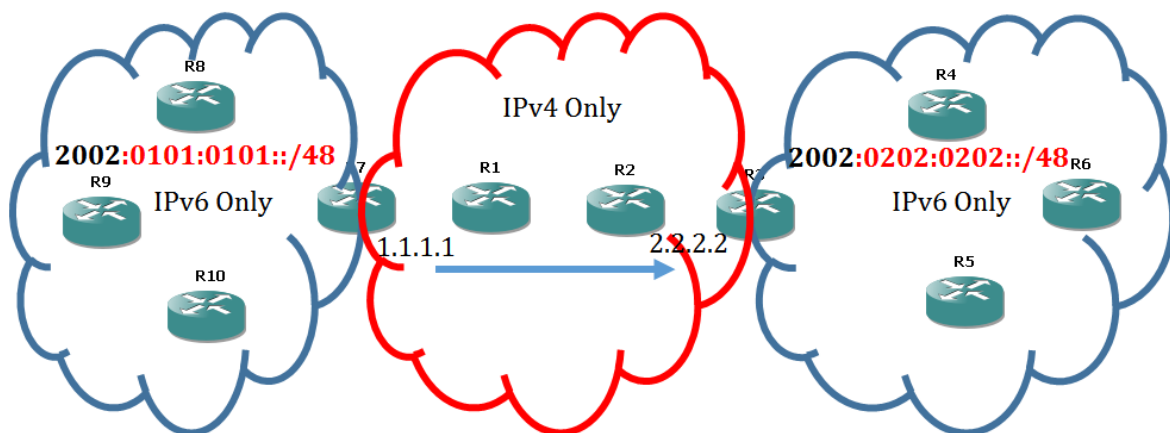
```
(Config-if)#tunnel source <ipv4-add>
(Config-if)#tunnel mode ipv6ip isatap
```

برای ارتباط بین دو مسیریاب هم مشابه 6to4 است با این تفاوت که آدرس ipv6 باید به صورت 0000:5EFE:x:y باشد.

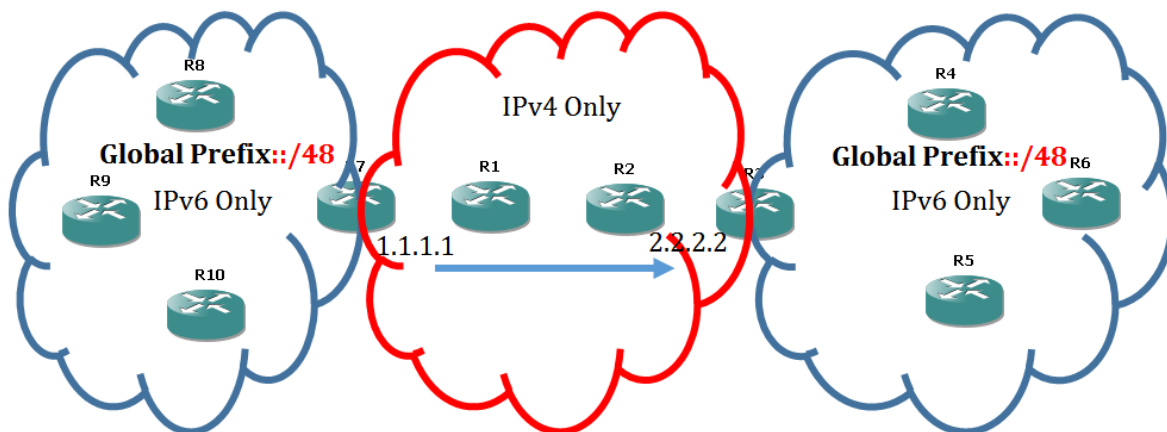
۴. 6to4: برای point-to-multipoint که به دو صورت 6to4 و ISATAP استفاده می‌شود. مسیریابی IGP IPv6 بر روی آن فعال نمی‌شود. باید از static route یا BGP استفاده نمود. برای کاهش ترافیک به دلیل نداشتن آدرس مقصد خوب است. مانند قبل باید default route تعریف شود و در Tunnel تنظیم گردد. در 6to4 آدرسی برای این کار رزرو شده است. این آدرس با 2002 شروع می‌شود و به صورت private است. در واقع این آدرس به گونه‌ای طراحی شده است که ۱۶ بیت اول 2002 و ۳۲ بیت بعد برابر IPv4 آدرس مقصد و ۱۶ بیت بعد آدرس subnet شبکه که از 0001 تا FFFF شماره گذاری می‌شود، قابل استفاده باشد. این روش به سادگی تنظیم می‌شود و برای شبکه محلی استفاده می‌گردد. روش دیگر، استفاده از آدرس‌های global است که پیچیده‌تر و فقط بر روی مسیریاب‌های لبه یا Edge آدرس بالا تنظیم می‌شود.

۱. ابتدا باید مبدأ آدرس که Loopback یا Physical Interface مشخص شود.
۲. از آدرس 2002:X:Y:Z::/64 برای شبکه محلی یا آدرس Global که X و Y برابر آدرس مبدأ IPv4 است و Z برابر subnet است، استفاده شود. اگر از آدرس global استفاده شود فقط مسیریاب‌های لبه از آدرس 2002 استفاده می‌کنند.
۳. Tunnel Interface ساخته شود. سپس آدرس IPv6 وارد شود. پس از Tunnel mode ipv6ip 6to4 انتخاب گردد.
۴. آدرس پیش‌فرض <next-hop> <tunnel> <ipv6> static ipv6 route تعیین گردد. اگر مسیریاب‌های بسیاری در قسمت private cloud باشد بهتر است از دستور ipv6 route <next-hop> <tunnel> <2002::/16> static استفاده شود. ولی اگر از آدرس global استفاده شود باید 48/ وارد نمود. در اینجا next-hop برای ارسال چندین مقصد استفاده می‌شود و آدرس مقصد در tunnel وارد نمی‌شود.
۵. مانند روش‌های قبل ipv6 unicast-routing و پروتکل مسیریابی بر روی تونل فعال شود.

Option-1 (Less Configuration Complexity)



Option-2 (More Configuration Complexity)

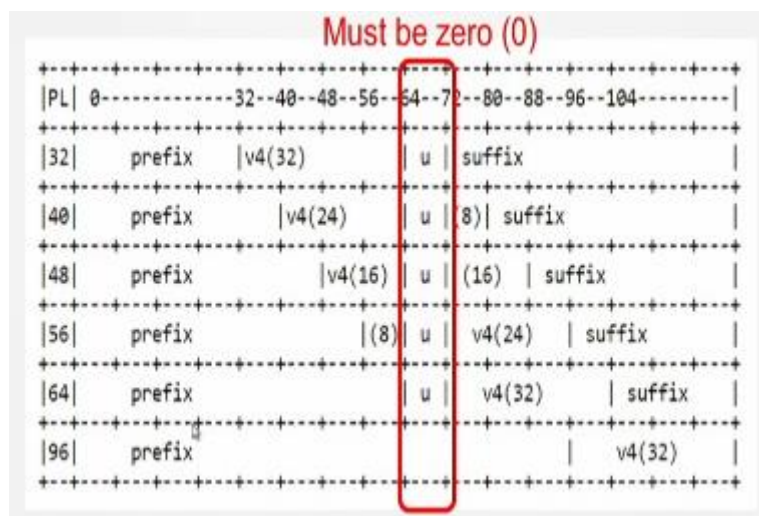


## NAT64 -۸-۷

به معنای حذف کامل Header IPv4 و اضافه نمودن Header IPv6 به بسته و بالعکس است. روش دیگری برای تبدیل آدرس است که به طور کامل با tunneling متفاوت است. زیرا یکی از headerها به طور کامل حذف می شود. طراحی اصلی آن تبدیل IPv4 به IPv6 است و از IPv4 به IPv6 را revers گویند. مسیر یاب ها و هاست ها دیگر به صورت dual stack نیستند.

## ۹-۷ - DNS64

در Ipv4 درخواست DNS به صورت A record ولی در Ipv6 به صورت AAAA یا quad A است. زمانی که درخواستی از Ipv4 به DNSv6 دریافت شود، پیغام یافت نشد داده می شود. در صورتیکه به DNS64 ارسال شود اگر در v6 یافت نشد در v4 جست و جو کرده و در صورت پیدا نمودن به همان صورت A یا AAAA پاسخ داده می شود. اگر درخواست v6 باشد به صورت <IPv6-prefix:IPv4-add> ارسال می شود. در NAT64 همانند tunneling آدرسی برای IPv4 embedded تعیین می گردد. این prefix به دوشکل است. Well-Known Prefix (WNP): غیر قابل مسیریابی است، زیرا global نیست و به صورت 64:FF9B::/96 است. Network Specific Prefix (NSP): که به صورت global و قابل مسیریابی است و به صورت 32/ یا 64/ است. در کل ترجمه IPv6 -> IPv4 وجود دارد. IPv4 + (/96 - /32) :: (WKP | NSP).



در بسته IPv6 باید قسمت u Octet صفر باشد. هر نوع prefix که استفاده شود این قسمت باید صفر شود. اگر WKP استفاده شود باید 64:FF9B::/96 استفاده شود. تفاوت بین NAT64 , NAT44 به این صورت است که ۱- در تنظیم nat44 از inside و outside برای اینترفیس ها استفاده می شود ولی در nat64 فقط باید بر روی اینترفیس فعال گردد. ۲- در nat44 فقط مکان source یا destination تغییر می کند ولی در nat64 هدر کاملاً تغییر می کند. دو نوع NAT64 وجود دارد: stateless که گرفتن IP با استفاده از SLAAC است و stateful که گرفتن IP از DHCP است.

۱- State Less: معمولاً برای v6 client و v4 server استفاده می شود و می توان برای چند IPv6 از یک IPv4 استفاده نمود. IPv4 به صورت embedded در IPv6 قرار می گیرد.

۱. ابتدا IPv4 Valid Range مشخص می شود.

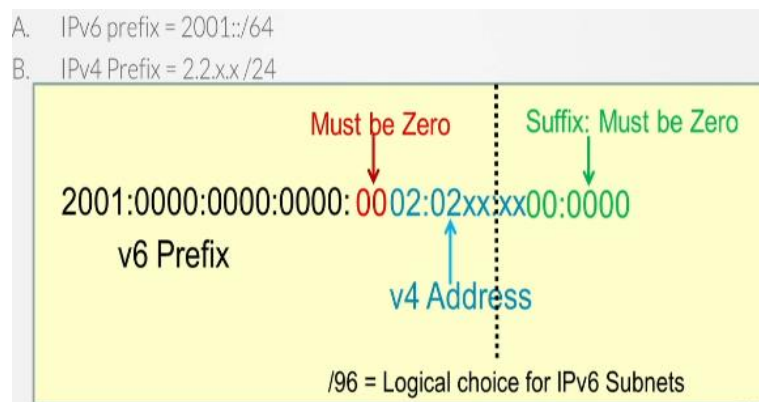
۲. با توجه به IPv4ها، Net-id های v4 در قسمت IPv6 مشخص می شود.
۳. حال Net-id های v6 نیز مشخص شود.
۴. IPv6 به صورت شکل زیر ساخته شود.
۵. به هر سیستم یک IP اختصاص داده شود.
۶. مسیریاب IPv6ها را از طریق IGP به مسیریابها ارسال می کند که از طریق من عبور داده شود.
۷. NAT64 بر روی مسیریاب NAT translator تنظیم می شود.

(Config-if)#nat64 enable

(Config)#nat64 prefix stateless <ipv6-prefix/32-/96> reach IPv4 Side

(Config)#nat64 route <ipv4-subnet/16 <if-ipv6side> reach ipv6 side

۸. حال DNS64 را می توان راه اندازی نمود.



۲- Stateful: به صورت خودکار است و sessionها را در جدول NAT ذخیره می کند. نیازی به IPv4 embedded نیست. شروع باید از سمت IPv6 باشد. همانند NAT44 که شروع باید از داخل به بیرون باشد. از حالت NSP برای IPv6 استفاده می نماید و توسط مسیریاب NAT64 با استفاده از IGP یا پروتکل های دیگر آن را اعلام می کند. بنابراین باید یک استخری از آدرس ها وجود داشته باشد و با استفاده از ACL اجازه استفاده داده شود. در آخر می توان NAT64 Stateful را تنظیم نمود:

(Config-if)#nat64 enable

(Config)#ipv6 access-list <name>

(Config-v6-acl)#permit ipv6 <prefix-/48> any

(Config)#nat64 prefix stateful <ipv6-prefix>

(Config)#nat64 v4 pool <name-pool> <ipv4-start-add> <ipv4-end-add>

(Config)#nat64 v6v4 list <name-of-ACL> pool <name-of-pool> overload

## ۸- VPN

VPN از دو بخش virtual network به معنای دسترسی دو یا چند شبکه به یکدیگر و Private network به معنای ایجاد policy برای ارتباط این شبکه‌ها، تشکیل می‌شود. این خود شامل route به معنای جدا کردن مسیرهای بسته‌های مسیریابی از دیگر محیط‌ها و data به معنای رمز گذاری و تونل کردن آنها است. باید توجه داشت که رمز گذاری اجباری نیست و فقط یک option است که عموماً استفاده می‌شود. قبل از این که از VPN استفاده شود، از PPP، Frame Relay و خط‌های ارتباطی T1 یا E1 استفاده می‌شد. این خود هزینه‌های زیادی به علت نصب خط‌های جداگانه و هزینه مدیریتی دارد. بعد از آن مسیریابی که به اینترنت متصل است، بسته‌ها را از طریق یک خط ارسال می‌کند. البته با ایجاد تونل آنها را به طرف دیگر منتقل می‌کند. به صورت کلی دو نوع VPN وجود دارد.

۱. Peer-to-peer: مسیریاب با ISP همسایه است و از طریق ISP به مسیریاب‌های دیگر دسترسی پیدا می‌کند.

مانند Layer 3 MPLS VPN.

۲. Overlay: مسیریاب ISP به عنوان همسایه نیست و مسیریاب طرف دیگر به عنوان همسایه است. مانند

DMVPN و Layer 2 MPLS VPN.

از ویژگی‌های VPN موارد زیر را می‌توان نام برد:

۱. مسیرها به صورت جداگانه یا ترکیب با مسیرهای دیگر، ارسال می‌شوند.

۲. می‌توان بسته‌ها را از مسیر مشخصی عبور داد.

۳. می‌توان با استفاده از IPSec امنیت آن را بالا برد.

## ۸-۱ IPSec Introduction

ویژگی‌های IPSec به شرح زیر است:

۱- اطلاعات را محرمانه یا رمز گذاری می‌نماید.

۲- جامعیت داده‌ها یا تغییر نکردن داده‌ها در بین مسیر را تضمین می‌کند.

۳- عدم تکرار بسته‌ها به معنای عدم ارسال دوباره بسته‌های قبلی را تضمین می‌کند.

اصطلاحاتی که در IPSec به کار برده می‌شود شامل موارد زیر است:

۱. Security Associate (SA): به معنای امن کردن تونل بین ابتدا و انتهای مسیر در صورت

فعال بودن IPSec در دو طرف است.

۲. Internet Key Exchange (IKE): استفاده از چندین key که ایجاد آن به صورت خود کار

یا دستی است. که این خود شامل دو مرحله می‌شود:

## I. Internet Security Associate Key Management (ISAKMP)

Protocol شروع IKE است. رسیدن و جابه جا نمودن Key از طرف دیگر است و

یک تونل امن بین دو طرف، ایجاد می نماید. ابتدا با یک Key جداگانه ایجاد و سپس

قابلیت ها بین دو طرف با استفاده از Transform Sets بررسی می شود.

## II. در مرحله بعد یک SA یک طرفه ایجاد می شود و طرف دیگر نیز یک تونل یک طرفه

دیگر با یک Key جداگانه ایجاد می نماید.

IPSec احراز اصالت، جلوگیری از تکرار و جامعیت بسته را ایجاد می کند. رمز گذاری در IPSec به صورت اختیاری

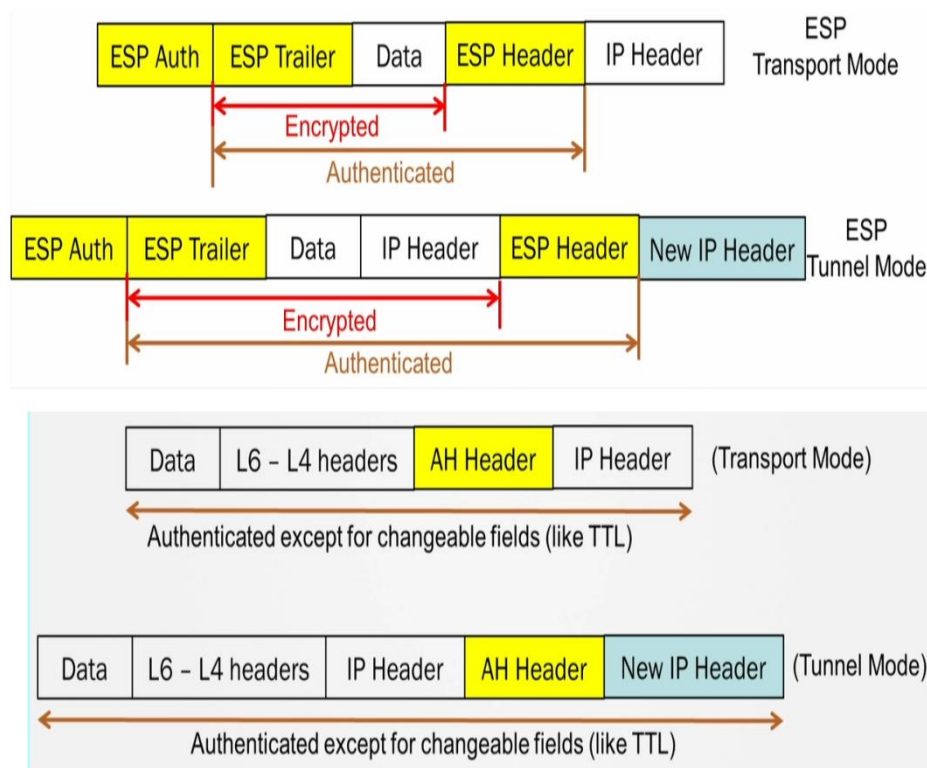
است. در IPSec هدری به نام AH یا Authentication Header وجود دارد که از IP protocol 51 استفاده

می کند. همچنین در IPSec بسته ها توسط ESP یا Encapsulation Security Payload که از ip protocol

50 استفاده می کند، بسته بندی می شوند. این دو برای احراز اصالت و جامعیت داده استفاده می شود. به یاد داشته باشید

که IP Protocol بین ۵۰ تا ۵۹ برای IP است و ۶۰ تا ۶۹ برای TCP و ۷۰ تا ۷۹ برای UDP است. تفاوت AH و

ESP در شکل زیر مشخص شده است.



AH: در لایه Transport جامعیت تمام هدرها، صحت مبدأ و مقصد را بررسی می کند. مبدأ و مقصد هدر

tunnel جدیدی که از تونل اینترنت فیس می آید را نیز بررسی می نماید.

ESP: برای رمزگذاری و در لایه transport برای endpoint استفاده می‌شود. در tunnel برای infrastructure استفاده می‌شود.

## ۸-۲- Dynamic Multipoint VPN (DMVPN)

Dynamic Multipoint VPN به صورت point-to-multipoint است که layer 3 Overlay vpn است. VPN همانند IPv6 Tunnel از tunnel interface استفاده و بسته را به این اینترنت ارسال می‌کند. این اینترنت نیز بسته را داخل تونل می‌فرستد و در طرف دیگر، آن را extract می‌کند. با استفاده از VPN مسیر به صورت virtual link یا تونل ایجاد می‌شود که مسیر را از دید دیگران پنهان و آن را به مکان دلخواه ارسال می‌کند. این عمل در لایه فیزیکی و data link انجام نمی‌گیرد بلکه در لایه ۳ یا network صورت می‌گیرد. DMVPN به صورت hub and spoke است و همچنین spoke-to-spoke را نیز پشتیبانی می‌کند. DMVPN از ابزارهایی مانند: multipoint gre tunnel، NHRP (Next Hop Resolution Protocol)، IPsec و Routing استفاده می‌کند. به دلایل زیر از DMVPN استفاده می‌شود:

۱. همه VPN ها سایت های دور از هم را به هم ارتباط می دهند.
۲. فقط نیاز به IP Connectivity است و نیازی به تنظیمات ISP ندارد.
۳. تنظیم routing policy توسط شرکت انجام می گیرد نه توسط ISP، زیرا Overlay است.
۴. بسیار مقیاس پذیر است چون تنظیمات روی Spoke انجام می گیرد و به تنظیمات spoke های دیگر ارتباطی ندارد.
۵. Hub routers باید در دسترس باشد. IP آن باید به صورت استاتیک و valid باشد.
۶. Spoke routers باید در دسترس باشد. IP آن باید به صورت استاتیک یا داینامیک و valid باشد.

## ۸-۳- GRE Introduction

از RFC 2784 استفاده می‌شود. یک Tunnel Mechanism است. به صورت پیش فرض برای cisco IOS Tunnel Encapsulation استفاده می‌شود. از چهار قسمت protocol type همان فیلد در قسمت data، version، Reserved0 و Checksum تشکیل می‌شود. در واقع حمل کننده یک پروتکل انتقال دهنده دیگر مانند IPv4 یا IPv6 یا apple talk است. دو حالت point-to-point و point-to-multipoint که هر دو نیاز به NHRP دارد را پشتیبانی می‌کند. از Tunnel key که برای بیشتر از یک gre tunnel استفاده می‌شود و برای جدا کردن آنها از یکدیگر است.

## ۸-۴- NHRP Introduction

از RFC 2332 استفاده می شود. مانند arp به صورت layer 2 resolution protocol and cache است. به صورت خود کار آدرس خصوصی و عمومی (public , private) نود مربوط به طرف دیگر تونل را پیدا می کند. مانند arp به صورت خود کار یا دستی آدرس ها را resolve می کند.

## ۸-۵ DMVPN Configuration

مراحل تنظیم و عملکرد DMVPN به شرح زیر است:

۱- یک tunnel interface ایجاد و تنظیمات عمومی آن انجام شود.

```
(Config-if)#tunnel source <public ip>
(Config-if)#tunnel mode gre multipoint
(Config-if)#tunnel ip add <private ip>
```

۲- تنظیمات NHS (Next Hop Server) به صورت دستی بر روی spoke در داخل NHRP Cache وارد شود.

NHRP NHS = hub private ip  
NHRP Cache : hub private ip -> hub public ip (static)

۳- IPSec بر روی tunnel interface ایجاد و فعال شود.

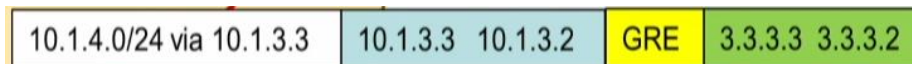
۴- Spoke درخواست NHRP registration request را برای Hub ارسال می کند.



۵- Hub پیام NHRP registration reply را ارسال می کند (به صورت داینامیک در cache ایجاد می شود).

۶- با استفاده از یک پروتکل مسیریابی IGP مسیریابی را بر روی tunnel interface فعال شود.

۷- در آخر اگر spoke-to-spoke نیاز باشد به صورت خود کار NHRP resolution request/reply به hub ارسال می شود تا تونلی بین spoke ها ایجاد شود.



پروتکل DMVPN چند فاز دارد: ۱- بین hub و spoke ۲- بین spoke و spoke ۳- modification of routing behavior: که به hub اجازه می دهد مسیرهای ارسال به spoke را خلاصه کند. که این در فاز دو نیست.

باید توجه داشته باشید که spoke-to-spoke به صورت on-demand است و زمانی که نیاز باشد ایجاد می شود. زمانی بین آنها به نام life time ایجاد می شود که بعد از هر ارسال بسته از صفر شروع می شود. به عبارت دیگر بستگی به ارسال ترافیک دارد. همچنین بر روی tunnel interface چندین پروتکل مسیریابی نیز می توان فعال نمود.

## ۹- VRF (Virtual Routing and Forwarding)

جدا کردن routing table شبکه ها از یکدیگر را VRF انجام می دهد. فرض شود اگر VRF وجود نداشت، دو مشتری که با آدرس های IP یکسان وجود دارند، ولی کاملاً از هم جدا هستند. حال مسیریاب بهترین مسیر را انتخاب می کند در صورتی که هر دو مسیر در شبکه خود صحیح هستند. بنابراین از VRF استفاده می شود تا جدول های مسیریابی از یکدیگر جدا شوند. این عمل هم بر روی ISP و هم بر روی مسیریاب مشتری انجام می گیرد. VRF بیشتر برای MPLS VPN استفاده می شود. این مطلب در CCNP گنجانده نشده است. در اینجا به VRF-Lite پرداخته می شود، که این عمل را پیاده سازی می کند.

```
(Config)#ip vrf <name-vrf>
```

```
(Config-if)#ip vrf forwarding <name-vrf>
```

با توجه به دستور بالا در واقع VRF بر روی interface اعمال می شود. قبل از این دستورات اینترفیس در global routing table بود. بعد از این به صورت خودکار از global جدا، آدرس آن حذف و به VRF table اضافه می شود. تنظیمات مربوط در داخل isp router به شرح زیر است:

```
(Config)#ip route vrf <name-vrf> <destination> <network> <next-hop>
```

```
(R1)#show ip vrf <name-vrf>
```

```
(R1)#show ip vrf interface <name-vrf>
```

```
(R1)#show ip route vrf <name-vrf>
```

```
(R1)#show ip protocol vrf <name-vrf>
```

این دستورات مخصوص RIP و EIGRP است که در تنظیم مسیریابی از address-family استفاده می شود.

```
(Config-router)#address-family ipv4 vrf <name-vrf>
```

```
r1#show ip <routing-protocol> vrf <name-vrf>
```

برای تنظیم OSPF از دستور زیر استفاده می شود.

```
(Config)#router ospf <process-number> vrf <name-vrf>
```

VRF این قابلیت را دارد که اگر OSPF وجود نداشته باشد به BGP نگاه می کند و از همان process-name استفاده می کند.

## ۱۰- Switch Introduction

عملکرد سوئیچ بر اساس آدرس MAC مقصد است. هنگام وارد شدن بسته بر اساس آدرس MAC بسته forward، discard و یا flood می‌شود. مانند مسیریاب که جدولی به نام routing دارد سوئیچ نیز جدولی به نام mac address دارد که به دو صورت تکمیل می‌شود. Static: در این جدول اطلاعاتی ذخیره شده است و همیشه بدون زمان‌بندی وجود دارد. مانند Cisco Discovery Protocol, Dynamic Trunking که مخصوص سوئیچ است. Dynamic: که در مرحله یادگیری سوئیچ در این جدول ذخیره شده و به صورت زمان‌بندی شده عمل می‌کند. هنگامی که بسته وارد سوئیچ می‌شود به قسمت src mac و vlan نگاه کرده و در جدول mac ذخیره می‌کند. زمانی به صورت پیش فرض ۵ دقیقه تنظیم می‌شود. حال اگر این پورت به حالت down برود به هر دلیلی از جدول mac حذف می‌شود.

```
(Config)#mac address-table aging-time <time-sec>
```

```
(Config)#mac address-table static <mac> vlan <vlan-id> interface <type/number>
```

```
(R1)#show mac address-table
```

```
(R1)#show mac address-table aging-time
```

استفاده از VLAN برای جدا کردن زیر شبکه‌ها از یکدیگر است. بین سوئیچ‌ها باید پورت در حالت Trunk باشد تا بتواند اطلاعات VLAN ها را جابه‌جا نماید. دو پروتکل برای VLAN وجود دارد. ISL: که مخصوص سیستم‌های سیسکو است 802.1q که تعریف شده IEEE است. در حالت trunk انواع VLAN وجود دارد

۱. Native: مانند vlan 1 که VLAN پیش فرض است. بیشترین پورت‌ها در این گروه است.

۲. Allowed: فقط اجازه عبور این vlan ها داده می‌شود.

۳. Encapsulation: بسته‌بندی کردن به فرمت 802.1q.

پروتکل Dynamic Trunking Protocol (DTP) مخصوص شرکت سیسکو است. این پروتکل وظیفه شنود خط را دارد و این که پورت در حالت access و یا trunk قرار گیرد را تشخیص می‌دهد. این پروتکل برای شنود دو حالت دارد. Auto یا Passive که خط را شنود و نوع VLAN را انتخاب می‌کند. Desirable یا Originate که با trunk صحبت را شروع می‌کند و بعد با استفاده از جواب آن طرف خط تصمیم می‌گیرد. در فرمت 802.1q که به صورت Tag VLAN عمل می‌کند شاخصه‌هایی دارد.

۱- Tag Protocol Identifier (TPID): به صورت کد برای شناسایی نوع Tag است.

۲- Priority: برای کیفیت سرویس استفاده می‌شود.

۳- Canonical Format Indicator (CFI): نوع شبکه Ethernet یا token ring را مشخص می‌کند.

۴- VLAN ID (VID): همان شماره VLAN است.

|    |    |            |          |      |     |
|----|----|------------|----------|------|-----|
| DA | SA | 802.1q Tag | Type/Len | Data | FCS |
|----|----|------------|----------|------|-----|

| 802.1q Tag     |                    |              |                |
|----------------|--------------------|--------------|----------------|
| TPID<br>16 Bit | Priority<br>3 Bits | CFI<br>1 Bit | VID<br>12 Bits |

به صورت پیش فرض تمامی VLAN ها اجازه عبور از روی trunk را دارند. برای محدود کردن آن از دستور زیر استفاده می شود.

(Config)#switchport trunk allow vlan {vlan-list | all | remove | add | except}

محدوده vlan به صورت 1 تا 4094 است که از 1 تا 1001 قابل استفاده، از 1001 تا 1005 رزرو برای token ring و از 1006 تا 4094 به عنوان extend-range که قابل استفاده در VTP Mode Transport یا VTPv3 است. VTPv3 در فایل VLAN.dat ذخیره نمی شود. اگر در حالت VTP تنظیم شود باید extend-VLAN ها به صورت دستی حذف شوند.

برای تنظیم normal vlan از دستور زیر استفاده می شود:

روش مدرن

(Config)#vlan <vlan-id>

(Config-vlan)#name <vlan-name>

روش قدیمی

(SW1)#vlan database

(SW1-vlan)#vlan <vlan-id>

برای تنظیم extend-vlan باید در حالت VTP Mode Transport باشد.

## ۱۱- Etherchannel

برای استفاده بیشتر از پهنای باند، استفاده از چند لینک به صورت هم زمان (load balance) و برای redundancy از این خصوصیت استفاده می شود. قبل از تنظیم این حالت سوئیچ این لینک ها را به عنوان حلقه می شناسد و STP آنها

را down می‌کند. برای رفع این مشکل باید این لینک‌ها را عضو port channel نمود تا سوئیچ این لینک‌ها را به عنوان یک لینک logical شناسایی کند. پیش تنظیمات آن به صورت زیر است:

۱. پورت باید trunk باشد.

۲. دو پروتکل وجود دارد که به صورت dynamic این کار را انجام می‌دهد. انتخاب یکی از آنها بنا به شرایط مورد نیاز استفاده می‌شود.

I. Port Aggregation Protocol (PAgP): به صورت desirable یا auto مانند trunk

desirable است. ابتدا صحبت را شروع و سپس نوع لینک را تشخیص می‌دهد.

II. Link Aggregation Control Protocol (LACP): این استاندارد ابتدا IEEE

802.3ad و سپس به 802.1ax تبدیل شد. دو حالت active و passive را پشتیبانی

می‌کند. این نیز همانند trunk است.

۳. ON: حالتی است که فقط Etherchannel فعال است و به صورت dynamic نیست.

## ۱۱-۱ PAgP vs LACP

تفاوت‌های دو پروتکل توضیح داده می‌شود. دقت داشته باشد که این مبحث یکی از پرکاربردترین و مباحث مهم CCNP است. اگر لینک به صورت half duplex باشد فقط PAgP پشتیبانی می‌شود. PAgP حداکثر ۸ لینک پشتیبانی می‌کند و با از بین رفتن یک لینک پهنای باند کاهش می‌یابد. LACP حداکثر ۸ لینک به صورت active و ۸ لینک به صورت standby را پشتیبانی می‌کند و اگر یک لینک از بین برود پهنای باند کاهش نمی‌یابد. در LACP می‌توان لینک‌هایی که عضو Etherchannel شوند را انتخاب و سپس آنها را تنظیم نمود. در پروتکل LACP چندین اصطلاح وجود دارد که توضیح داده می‌شود:

۱. System Priority: مقداری است که بر روی هر مسیریاب یا سوئیچ به صورت دستی یا خودکار تنظیم

می‌شود. پروتکل LACP با استفاده از System Priority و MAC دستگاه، System ID را ایجاد می‌کند

تا بتواند با دیگر دستگاه‌ها ارتباط برقرار کند. در صورت یکسان بودن System Priority، آدرس MAC با

کمترین عدد بهترین System ID است. در پروتکل LACP مقدار System Priority کمترین عدد بهترین

است که آن دستگاه master می‌شود.

۲. Port Priority: همانند SP مقداری است که بر روی هر مسیریاب یا سوئیچ به صورت دستی یا خودکار

تنظیم می‌شود. کمترین عدد بهترین است. پروتکل LACP با استفاده از Port Priority و Port Number

دستگاه، Port ID را ایجاد می‌کند. در صورت یکسان بودن port priority، شماره پورت با کمترین عدد

بهترین پورت است که عضو Etherchannel شود. این مقدار مشخص می کند که کدام پورت در حالت فعال و کدام پورت در حالت آماده باش قرار بگیرد.

در LACP یک سوئیچ Master و دیگری Slave می شود که بر اساس System ID انتخاب می شود. دستورات مورد استفاده برای مشاهده Etherchannel به شرح زیر است:

```
(SW1)#show etherchannel load-balance
```

```
(SW1)#test etherchannel load-balance interface port-channel <number> <ip | mac>
```

### ۱۱-۲- Etherchannel Layer

۱- Layer 2 OSI Etherchannel: تعداد ۲ تا ۸ لینک می تواند در etherchannel قرار گیرد و همه در یک broadcast domain قرار دارند. هیچ تفاوتی بین لینک ها نیست و همه به صورت trunk هستند. انتخاب مسیر بر اساس MAC Address است.

۲- Layer 3 OSI Etherchannel: load-balancing یا توازن بار بر اساس هر جریان است (per-flow) و هر جریان فقط از یک لینک استفاده می کند. اگر خواسته شود قسمت broadcast domain بخش مدیریتی و سرورها را از بخش end user جدا شود، از این مورد استفاده می شود. دلیل استفاده سرعت بالا و مدیریتی بودن آن است. برای تنظیم باید از حالت switchport خارج نموده و layer-3 port-channel را فعال نمود.

```
(Config)#interface port-channel <number>
```

```
(Config-if)#no switchport
```

### ۱۱-۳- Etherchannel Misconfiguration Guard

این خصوصیت خطای عدم تنظیم صحیح را به کاربر اعلام می کند. خطا ممکن است بدین شکل باشد که در یک طرف Etherchannel تنظیم شده است و طرف دیگر خیر. این خصوصیت از STP (Spanning Tree Protocol) بهره می گیرد. در STP بسته هایی به نام BPDUs (Bridge Protocol Data Unit) وجود دارد. در این بسته فیلدهایی به نام Port-ID و Port-Priority وجود دارد. در حالت عادی هر بسته STP BPDUs یک port-id دارد ولی در حالت etherchannel به تعداد لینک ها port-id در داخل این بسته قرار می گیرد. پس از این طریق Etherchannel قابل شناسایی است. سوئیچ های شبکه که بر روی آنها پروتکل STP اجرا شده برای اطلاع از احوال همدیگر از بسته های BPDUs استفاده می کنند. دو نوع بسته BPDUs در پروتکل STP وجود دارد. BPDUs Configuration: سوئیچ ها از این بسته ها برای محاسباتشان استفاده می کنند. TCN BPDUs: وقتی تغییری بر روی سوئیچی در شبکه رخ دهد برای اعلام تغییرات این مدل بسته را به سوئیچ های دیگر ارسال می کند تا بروز

تغییرات رو اعلام کند. برای جلوگیری از حلقه مابتدا باید پیش‌نیازهایی صورت گیرد که به ترتیب توضیح داده می‌شود. Root Bridge: یکی از سوئیچ‌ها در شبکه به عنوان سوئیچ Root Bridge انتخاب می‌شود. این سوئیچ تعیین کننده مسیر انتقال بسته‌ها در شبکه است. مبنای انتخاب Root Bridge در شبکه Bridge ID است. هر سوئیچ در شبکه یک Bridge ID منحصر به فرد دارد. سوئیچی که Bridge ID آن از همه کمتر باشد به عنوان Root Bridge انتخاب خواهد شد. Bridge ID یک مقدار ۸ بایتی است که از دو قسمت تشکیل شده است. Bridge Priority: این عدد (یک مقدار ۲ بایتی است) و Mac Address (یک مقدار ۶ بایتی است). Bridge Priority عددی بین ۰ تا ۶۵۵۳۵ است. سوئیچی که Bridge Priority آن از سوئیچ‌های دیگر کمتر باشد به عنوان Root Bridge انتخاب خواهد شد. در صورتی که Bridge Priority همه سوئیچ‌ها یکسان باشد، از آدرس Mac سوئیچ‌ها به عنوان الگوی انتخاب استفاده می‌شود. در این حالت سوئیچی که آدرس Mac آن از سوئیچ‌های دیگر کمتر باشد به عنوان Root Bridge انتخاب خواهد شد. به صورت پیش‌فرض Bridge Priority همه سوئیچ‌ها ۳۲۷۶۸ می‌باشد. (یعنی Bridge Priority همه سوئیچ‌ها به صورت پیش‌فرض یکسان است. به کمک دستور زیر می‌توان مقدار Bridge Priority رو در یک سوئیچ تغییر داد:

```
(Config)#interface <type> <number>
```

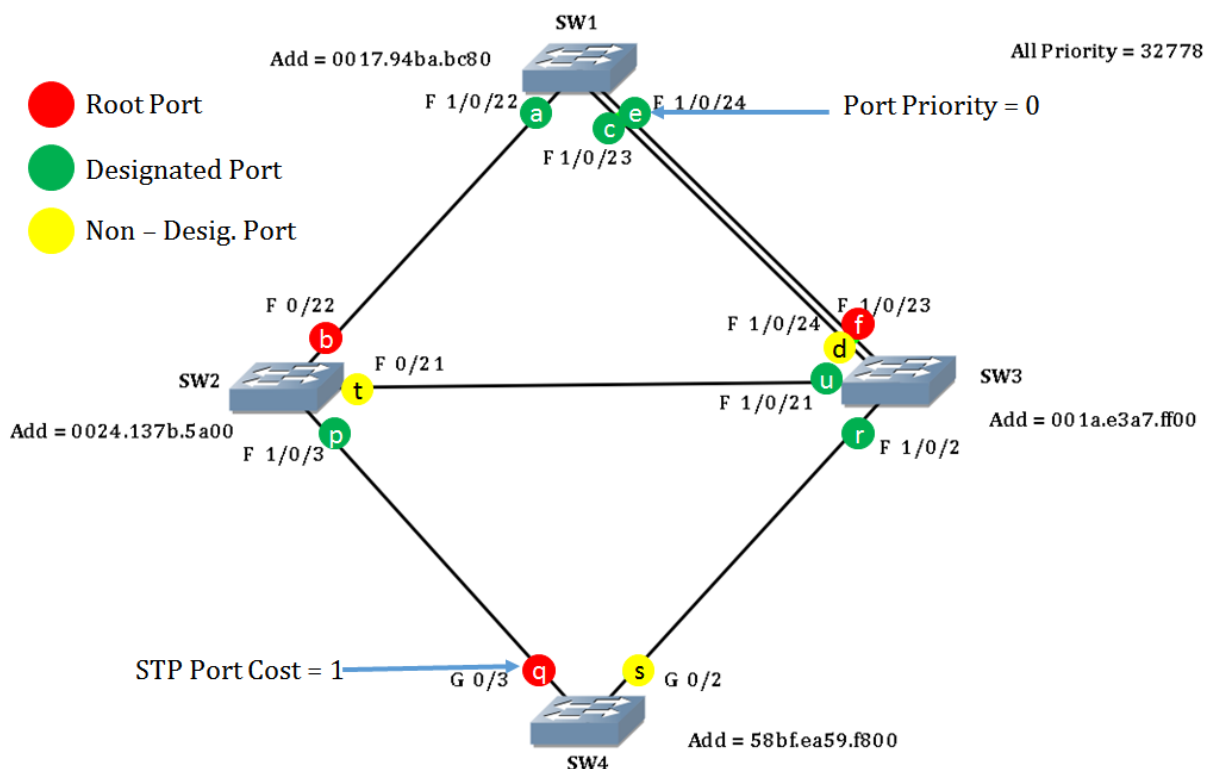
```
(Config-if)#spanning-tree port-priority <number>
```

بعد از اینکه Root Bridge انتخاب شد حال باید روی هر سوئیچ یکی از پورت‌ها به عنوان Root Port انتخاب شود. پورتهای به عنوان Root Port انتخاب می‌شود که Root Path Cost آن از همه پورت‌های دیگر کمتر باشد. Root Path Cost: هزینه یک مسیر از این سوئیچ تا سوئیچ Root Bridge است. Root Bridge و Root Port انتخاب شده‌اند اما هنوز حلقه در شبکه وجود دارد. در این مرحله نوبت به انتخاب Designated Port می‌رسد. بین هر لینک در شبکه یک پورت به عنوان Designated Port انتخاب می‌شود. فرمول انتخاب Designated Port هم طبق فرمولی است که Root Port انتخاب می‌شود. یعنی در هر لینک، پورتهای که کمترین هزینه به سمت Root Bridge رو داشته باشد به عنوان Designated Port در یک لینک انتخاب می‌شود.

به صورت کلی Etherchannel Misconfiguration Guard در حالت پیش‌فرض ON است. اگر طرف دیگر خطا وجود داشته باشد با استفاده از BPDUs شناسایی می‌کند. پس از شناسایی و وجود خطا پورت را در حالت err-disable می‌برد.

```
(Config)#spanning-tree etherchannel guard misconfig
```

```
(SW1)#show spanning-tree <active | summary>
```



## ۱۲- Virtual Trunking Protocol (VTP)

به معنای ایجاد، حذف و یا تغییراتی بر روی یک سوئیچ و اعمال آن یصورت خودکار بر روی سوئیچ‌های دیگر است. این پروتکل فقط مخصوص شرکت سیسکو و تجهیزات سوئیچ است و برای مسیریاب کارایی ندارد. پیش نیازهای آن برای تنظیمات شامل موارد زیر است:

۱. همه سوئیچ‌ها باید بر روی یک VTP Domain باشند و حروف کوچک و بزرگ اهمیت دارد (case sensitive).
۲. سوئیچ‌ها باید با لینک trunk به هم متصل شده باشند.
۳. پسورد انتخاب شده برای VTP باید یکسان باشد.
۴. نسخه VTP برای ارتباط باید یکسان باشد و در صورت عدم یکسان بودن، مشکلاتی از قبیل عدم ساخت VLAN بر روی سوئیچ‌های دیگر ایجاد می‌شود. اگر یکی از نسخه‌ها VTP Version 1 و VTP Version 2 Capable باشد و به یک همسایه با VTP Version 2 یا VTP Version 3 متصل باشد نسخه VTP به صورت خودکار به نسخه VTP Version 2 به روز می‌شود.

پروتکل VTP در نسخه‌های V1 & V2 & V3 با هم تفاوت دارند. در نسخه ۱ و ۲ از token ring پشتیبانی می‌کند. از ۱ به ۲ به صورت خودکار و به ۳ به صورت دستی به روز می‌شود. در نسخه ۱ و ۲ می‌توان نامی برای VTP Domain انتخاب نکرد ولی در نسخه ۳ انتخاب نام اجباری است.

VTP V3 دارای ویژگی‌های زیر است:

۱. با نسخه ۲ انطباق دارد.
۲. از همه VLAN ها (Normal & Extend) پشتیبانی می‌کند
۳. از تکثیر Private VLAN پشتیبانی می‌کند
۴. به صورت اختیاری می‌توان از پسورد به صورت clear text یا hidden یا hash استفاده نمود.
۵. تنظیمات (Multi Spanning Tree) 802.1s MST تکثیر می‌کند.
۶. می‌توان VTP را برای هر لینک و یا در سطح global خاموش و یا روشن نمود.
۷. احراز اصالت به صورت اختیاری است و در تمامی نسخه‌ها در running-config نشان داده نمی‌شود و با وارد نمودن دستور show vtp password نشان داده می‌شود. در کل سه روش normal، hidden و secret برای وارد نمودن پسورد وجود دارد.

```
(Config)#vtp password <clear-text>
```

```
(SW1)#vtp password <clear-text>
```

با دستور show vtp password نشان داده می‌شود.

```
(Config)#vtp password <clear-text> hidden
```

با دستور show vtp password نشان داده نمی‌شود.

```
(Config)#vtp password <hash-text> secret
```

با دستور show vtp password نشان داده نمی‌شود.

در نسخه ۱ و ۲ همه سوئیچ‌ها به صورت پیش فرض سرویس دهنده هستند و برای تنظیم vtp client و یا transport یک یا دو سوئیچ سرویس دهنده و بقیه سرویس گیرنده هستند. در نسخه ۳ یک سوئیچ Primary server به معنای اجازه داشتن برای ساخت و حذف VLAN و ارسال بسته به روز رسانی است و بقیه Secondary server هستند، به معنای اجازه نداشتن ساخت یا حذف VLAN و عدم اجازه برای ارسال بسته به روز رسانی VLAN Database به دیگران است. Primary server در هر VTP Domain تنها یکی وجود دارد.

تنظیمات مربوط به VTP V3 به شرح زیر است:

۱. ابتدا باید VTP Domain Name تنظیم شود.

۲. سپس نسخه آن به ۳ تغییر یابد.

۳. Primary server تنظیم شود.

۴. پسورد تنظیم شود.

```
(SW1)#vtp domain <domain-name>
(SW1)#vtp version {1 | 2 | 3}
(SW1)#vtp primary-server [VLAN | MST] [Force]
(SW1)#vtp password <password>
```

برای تنظیم primary server سه حالت وجود دارد ۱- Force: تداخل primary server را بررسی نکند. ۲- Multi Spanning Tree (MST): متصدی پایگاه داده MST است. ۳- VLAN: متصدی پایگاه داده VLAN است. ۴- <cr>: تداخل primary server را بررسی کند.

## ۱۲-۱ VTP Pruning

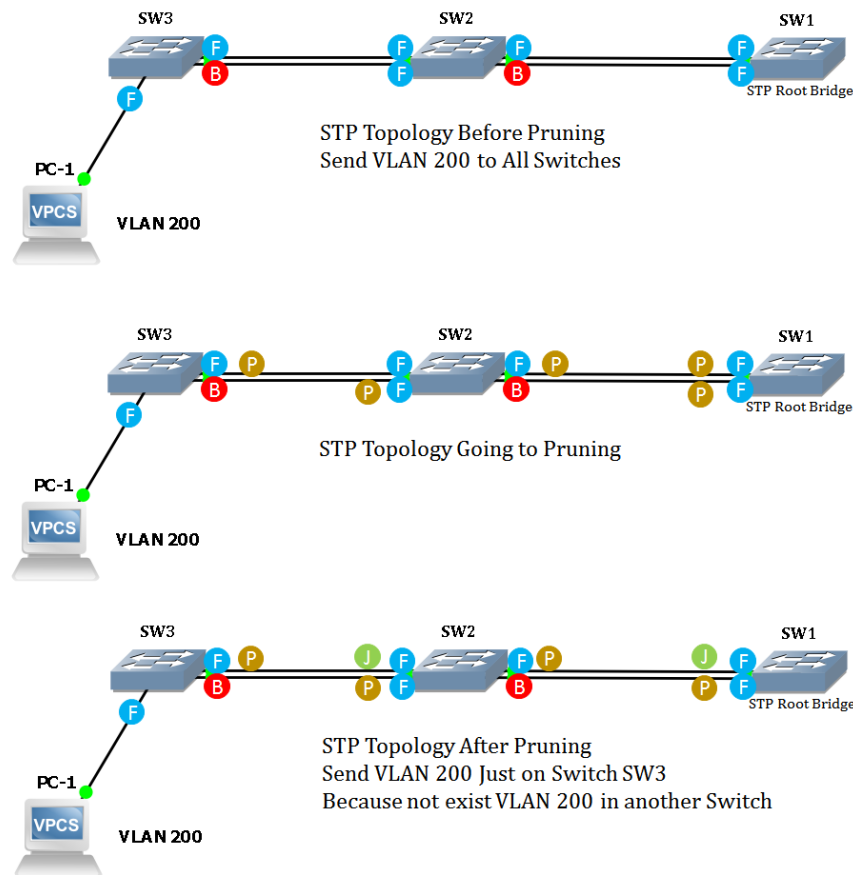
به معنای هرس کردن VLAN ها است. در واقع هرس کردن به معنای این است که اگر سوئیچی پورتی مربوط به یک VLAN را نداشت، اجازه عبور ترافیک آن VLAN از سوئیچ داده نمی شود. هرس کردن به شبکه کمک می کند تا ترافیک های غیر لازم کاهش یابند و پهنای باند شبکه افزایش بیابد. VTP Pruning فقط بر روی سرور قابل تنظیم است. مراحل آن شامل موارد زیر است:

۱. اگر سوئیچ هیچ access port در داخل VLAN نداشت، همه پورت های trunk مربوط به VLAN را هرس می کند.

۲. اگر VLAN در یک سوئیچ فعال شود، به معنای آنکه یک access port عضو VLAN داشته باشد. پیام Triggered Join به تمامی پورت های (non-root bridge (STP Root Ports و Root Bridge (all designated ports) ارسال می شود.

۳. سوئیچی که پیام Triggered Join را ارسال می کند، پورت های آن به حالت Prune رفته و اجازه عبور VLAN را نمی دهد (در صورتی اجازه می دهد که قبل از آن هاستی با همان VLAN به پورت آن وارد شود و مراحل Triggered Joining از قبل به این پورت وارد شده باشد و در حالت Join باشد). ولی سوئیچی که پیام را دریافت می کند و پورت آن در حالت forwarding است به حالت Join رفته و اجازه عبور VLAN را دارد.

در شکل زیر اگر vlan-x در سوئیچ سمت چپ فعال شود تمامی پورت‌های F به حالت J رفته و ارتباط بین همه سوئیچ‌ها برقرار می‌باشد.



### ۱۳- Multi-Layer Switch

سوئیچ‌های لایه ۲ فقط به MAC آدرس مقصد نگاه می‌کنند و تصمیم‌گیری انجام می‌دهد. همچنین دارای MAC آدرس‌های استاتیک پیش‌فرض به طور مثال CDP، STP، DTP و PAGP برای مدیریت هستند. در سوئیچ‌های لایه ۳ به هدر لایه ۲ و لایه ۳ نگاه می‌شود و با توجه به جدول‌های مسیریابی و MAC آدرس، مقصد را مشخص می‌کند. در سوئیچینگ مسیریابی نسبت به routing سریعتر انجام می‌گردد. در سوئیچ انواع کمتری از اینترفیس استفاده و پشتیبانی می‌شود ولی در router از ATM WAN T1 و انواع خاص دیگر استفاده می‌شود. زمانی که بر روی یک اینترفیس پروتکل مسیریابی فعال می‌شود IP Multicast به MAC Address Multicast گماشته می‌شود. در واقع سوئیچ‌های چند لایه به معنای پشتیبانی کردن یا بازرسی کردن بسته‌ها، بیشتر از یک لایه در مدل OSI است. این به معنای داشتن تمامی امکانات router نیست، بلکه هدف این است که هنگامی که بسته وارد سوئیچ چند لایه می‌شود تصمیم برای حذف، ارسال و پخش بسته انجام گیرد. بیشتر سوئیچ‌های چند لایه از TCP/UDP Header پشتیبانی

می‌کند. مثال: اگر چند کاربر در یک VLAN باشند و خواسته شود بین آنها ارتباطی نباشد، پروتکل مسیریابی نیز وجود نداشته باشد و از ACL هم استفاده نشود، به دلیل اینکه کاربران در یک VLAN و broadcast domain هستند از ابزار VACL (VLAN Access Control List) می‌توان استفاده نمود. سوئیچ‌های چند لایه VACL را می‌فهمند و بسته را حذف می‌کنند. اینترفیس سوئیچ‌های چند لایه فقط عمل switching و یا فقط عمل routing انجام می‌دهند.

Switch Virtual Interface (SVI): در واقع VLAN یک SVI است و برای استفاده از آن باید نام VLAN وجود داشته یا تنظیم شود و حداقل یک اینترفیس فیزیکی عضو VLAN باشد (vlan enable and usable). بعضی از سوئیچ‌ها برای vlanها یک mac address واحد و بعضی چند mac address دارند. روش عملکرد سوئیچ‌های چند لایه به این گونه است که جدولی به نام hardware table وجود دارد که شامل دو ردیف mac address و IP address است. قسمت layer 2 ASIC نتیجه را از جدول mac بر می‌دارد و تحلیل می‌کند. قسمت layer 3 ASIC نتیجه را از جدول IP بر می‌دارد و تحلیل می‌کند. این دو با هم در ارتباط بوده و در صورت پیدا کردن mac مقصد و مسیر آن بسته را ارسال می‌کنند. این جدول‌ها در حافظه TCAM (Ternary Content Addressable Memory) ذخیره می‌شوند. همچنین این حافظه VLAN security ACL و Routing را در خود جای می‌دهد. هنگام ورود بسته des mac و des ip به next hop تبدیل می‌شود.

| MAC Address                             | IP Address                        |
|-----------------------------------------|-----------------------------------|
| destination mac address via interface # | destination ip address via vlan # |
| Layer 2 Forwarding Engine               | Layer 3 Forwarding Engine         |
| Layer 2 ASIC (parse mac)                | Layer 3 ASIC (parse ip)           |

## ۱۴- Cisco Express Forwarding (CEF)

عمل switching به معنای گرفتن بسته از یک اینترفیس و فرستادن آن به یک اینترفیس دیگر است. این عمل هم در switch و هم در router صورت می‌گیرد. در CCNP ۳ روش برای switching وجود دارد:

1. Process-Base Switching: زمانی که بسته وارد سوئیچ می‌شود در CPU وقفه ایجاد می‌شود. طبیعی است که CPU مشغول انجام پردازش‌های دیگر است (مانند EIGRP و...). پس از رسیدن نوبت پروسه‌ای به نام "IP Input" را شروع می‌کند و با دستور show processes CPU می‌توان پروسه‌های CPU را مشاهده نمود. سیستم‌های امروزی به دلیل وقفه زیاد و عملکرد پایین از این روش استفاده نمی‌کنند.

۲. Fast Switching: اگر یک حافظه جدا کوچک برای هر اینترفیس قرار داده شود و سخت‌افزارها بتوانند به آن دسترسی داشته باشند، هنگامی که یک بسته وارد سوئیچ می‌شود نگاه به این حافظه شود و اگر تطابقی پیدا نمود آن را ارسال کند. اگر تطابقی نداشت آن‌گاه به CPU برای پیدا کردن مسیر ارسال شود و عملیات Process-Base را انجام دهد. حال اگر هیچ بسته‌ای مطابق بسته‌های موجود در حافظه در مدت زمان مشخصی ارسال نشود packet age آن به اتمام رسیده و از حافظه حذف می‌گردد.

۳- Cisco Express Forwarding: به این حالت Topology Driven Cache نیز گفته می‌شود. اگر حافظه‌ای که در مدل قبلی بحث شد همیشه به روز و مورد استفاده باشد، به عبارتی همه مسیرها در این حافظه قرار گیرد و یک کپی از routing و MAC باشد عملیات switching به سرعت انجام می‌شود. این حافظه از دو جدول تشکیل شده است.

۱. Forwarding Information Base (FIB): یک shadow copy از IP Routing Table

است و فقط آنهایی که نیاز است در این جدول قرار دارد. مانند: Prefix، Net mask Interface، Out

و Next Hop. این جدول همیشه با IP table در ارتباط است و به روز رسانی می‌شود.

```
(Config)#ip cef
```

```
(Config)#no ip cef
```

```
(R1)#show ip cef [ip] [mask] [detail]
```

۲. Adjacency Table: این جدول هم یک کپی از MAC address table است. در واقع جدول

FIB اطلاعات لایه ۳ را بر عهده دارد و این جدول اطلاعات لایه ۲ و اطلاعاتی از قبیل arp،

frame relay mapping table و ... را در خود جای می‌دهد.

```
(R1)#show adjacency <interface / number> [summary | detail]
```

```
(R1)#show adjacency vlan <vlan-id> detail
```

اگر لینک متصل باشد جدول FIB به روز شده و درخواست پروسه ARP انجام می‌گیرد و

adjacency table نیز به روز می‌شود. Adjacency انواع مختلفی دارد که بعضی از آنها

نمی‌تواند در CEF Switching شرکت کنند و باید به CPU ارسال شوند.

حالت‌هایی که در جدول CEF رخ می‌دهند به شرح زیر هستند:

۱. Glean: اگر آدرس مقصد با جدول تطابق داشت ولی نیاز به اطلاعات دیگری داشته باشد بسته به

CPU برای پردازش فرستاده می‌شود. به طور مثال نیاز به آدرس MAC باشد. در Glean نیاز به

پردازش لایه ۲ است و منتظر می‌ماند تا پاسخ دریافت کند. به طور مثال نیاز به ARP دارد. این حالت

قبل از حالت Attached قرار دارد و برای مدت بسیار کوتاهی است.

۲. Null: اگر بسته valid باشد و حذف گردد. مانند تعریف مسیر به Null0.
۳. Drop: اگر بسته‌ای valid باشد، ولی به گونه‌ای حذف گردد مانند encapsulation نا صحیح.
۴. Discard: بر اساس تعریف ACL بسته حذف گردد.
۵. Punt: شبیه به Glean است با این تفاوت که نیاز به پردازش لایه ۳ دارد. مانند route process.

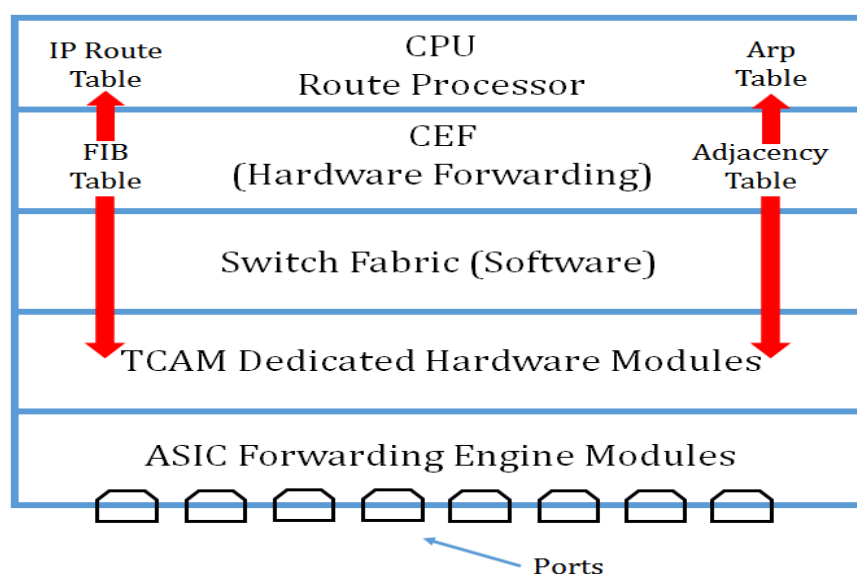
## ۱۴-۱- CEF vs TCAM

در واقع نتیجه عملیات CEF در حافظه TCAM قرار می‌گیرد. لایه ۳ حافظه TCAM نتیجه‌ای از جدول FIB و لایه ۲ آن، نتیجه Adjacency است. زمانی که سوئیچ روشن می‌شود CEF در پشت صحنه در حال اجراست. در سوئیچ‌های چند لایه به صورت پیش فرض روشن است و در بعضی دیگر باید روشن شود. بعضی از سوئیچ‌ها پایه و اساس آن CEF است مانند Cisco 6500 ولی بعضی خیر. بعضی از بسته‌ها نمی‌توانند در CEF Switching شرکت کنند. این بسته‌ها شامل موارد زیر هستند:

۱. ARP request باید در CPU پردازش شود.
  ۲. بسته‌هایی که نیازمند پاسخ CPU هستند مانند بررسی TTL، Fragmentation، IP Option not Empty و...، در CEF شرکت نمی‌کنند.
  ۳. بسته‌های Routing Protocol Traffic در CEF شرکت نمی‌کنند.
  ۴. بسته‌های Cisco Discovery Protocol در CEF شرکت نمی‌کنند.
  ۵. بسته‌های packet need encryption در CEF شرکت نمی‌کنند.
- باید توجه داشت که بیشتر استفاده switching از CEF است و کمتر از CPU استفاده می‌شود.

## ۱۵- Switch Software Introduction

در شکل زیر مؤلفه‌های یک سوئیچ نشان داده شده است. ASIC (Application Specific Integrated Circuit) به معنای مدارهای مجتمع با کاربرد خاص می‌باشد که یک سوئیچ می‌تواند یک (3750) یا چند (6500) ASIC داشته باشد. در قسمت Forwarding Engine بسته‌ها ذخیره و برای استخراج اطلاعات از آن نگهداری می‌شود و اطلاعات آن به TCAM ارسال می‌شود. سوئیچ‌ها می‌توانند یک یا چند Forwarding Engine داشته باشند به طور مثال TCAM: layer3، Security Stuff و QoS. مؤلفه CEF نیز برنامه‌ای برای تهیه و ذخیره اطلاعات در TCAM است. CPU عملیات‌های یادگیری مسیر، مسیریابی و ... را انجام می‌دهد که از جدول‌های IP و ARP استفاده می‌کند. اگر CEF در داخل CPU فعال باشد جدول‌های FIB و Adjacency اطلاعات را از IP و ARP استخراج می‌کند. TCAM یک کپی از FIB و Adjacency است و CEF آن را پر می‌کند. CEF فقط برای این ساخته نشده است و یک قسمت مهم و مفید برای کاهش بار CPU است.



## ۱۶- Switch Redundancy

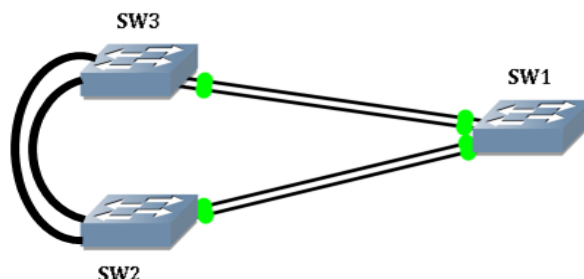
به هم متصل نمودن سوئیچ‌ها به صورت سخت‌افزاری است. به عبارتی stackwise می‌شوند که سوئیچ‌ها از طریق کابل مخصوصی به هم متصل می‌شوند و همه به عنوان یک سوئیچ منطقی شناخته می‌شوند. تعداد Stack شدن سوئیچ‌ها بستگی به نوع آنها دارد. حداکثر تعداد سوئیچ برای stack شدن در مدل‌های Catalyst برابر ۹ سوئیچ است و مدیریت آنها بر عهده یک سوئیچ Master است. در واقع یک IP برای مدیریت وجود دارد و همه به عنوان یک virtual switch شناخته می‌شوند. علت استفاده از آن برای مدیریت آسان و تحمل پذیری بالا در برابر خطا است. همه سوئیچ‌ها اطلاعات را از سوئیچ Master دریافت می‌کنند و در واقع همه دارای config file، MAC

Table و routing table یکسانی هستند. به غیر از این اطلاعات سوئیچ ها اطلاعات محلی نیز دارند. مرحله‌ای که برای انتخاب سوئیچ Master وجود دارد به شرح زیر است:

۱. User priority: در CCNP بحث نمی‌شود. این مؤلفه بین ۱ تا ۱۵ است. پیش فرض همه برابر ۱ است.
۲. Hardware and Software Priority: سوئیچی با قابلیت‌های بیشتر به عنوان Master انتخاب می‌شود.
۳. Default Configuration: سوئیچی که default configuration کمتری داشته باشد به عنوان master انتخاب می‌شود.
۴. Up Time: سوئیچی که زمان بیشتری روشن باشد به عنوان Master انتخاب می‌شود.
۵. MAC Address: سوئیچ با آدرس mac پایین تر Master انتخاب می‌شود.

پس از انتخاب سوئیچ Master اگر کابلی به پورت Console سوئیچ Slave متصل شود، Console به صورت خودکار به سوئیچ Master انتقال می‌یابد. در واقع مهم نیست به کدام سوئیچ متصل شد، در هر حال به سوئیچ Master با یک session اتصال برقرار می‌شود و CLI را کنترل می‌کند.

اگر یک سوئیچ به چند سوئیچ که با هم Stack شده‌اند، Ether Channel شوند به آن Cross-Channel Ether Channel گفته می‌شود.



در یک سیستم stack به صورت Online می‌توان سوئیچ‌ها را اضافه یا حذف نمود. اگر سوئیچ Master از بین برود مرحله انتخاب دوباره تکرار می‌شود و بقیه سوئیچ‌ها با استفاده از آخرین اطلاعاتی که از سوئیچ Master گرفته‌اند عملیات routing, Forwarding و Flow Control را انجام می‌دهند. وظیفه سوئیچ Master عملیات‌های نگهداری، به روزرسانی فایل تنظیمات، جدول مسیریابی و همچنین ارسال اطلاعات دیگر مربوط به Stack است. دیگر سوئیچ‌ها نیز اطلاعات جدول MAC محلی و اطلاعات STP و همچنین جدولی که از Master گرفته شده است را ذخیره می‌کنند. حلقه‌ای که کابل مخصوص stack ایجاد می‌کند در حالت block نمی‌رود. یک تفاوتی که در سیستم stack و حالت عادی است، مدیریت بر عهده یک سوئیچ و تنها با یک session است. در سیستم Stack دو نوع کابل Data و Power استفاده می‌شود.

## ۱۶-۱- Virtual Switching System (VSS)

برای فهمیدن این مطلب ابتدا تفاوت بین VSS و Stack شرح داده می شود.

۱. در سیستم stack سوئیچ ها با کابل مخصوصی به هم متصل می شوند که طول کوتاهی در حد یک rack یا حدود ۳ متر دارند ولی در virtual switching system از کابل های 10GE و یا FC می توان استفاده نمود که تا ۴۰ کیلومتر پشتیبانی می شود.

۲. در Stack حداکثر ۹ سوئیچ و در VSS تنها دو سوئیچ لازم است.

۳. در Stack یک سوئیچ master و بقیه slave هستند در vss یک سوئیچ Active و دیگری Standby است.

سوئیچ Active: اجرای پروتکل های کنترلی لایه ۲ و ۳ مانند پروتکل های مسیریابی و ...، Console interface و عملیات های مدیریتی بر عهده این سوئیچ است. سوئیچ Standby: ارسال بسته های کنترلی به سوئیچ Active بر عهده این سوئیچ است. هر دو سوئیچ وظیفه ارسال بسته ها به هاست های محلی را دارند. در VSS لینک بین دو سوئیچ که حداکثر ۸ لینک است را Virtual Switch Link (VSL) گویند و پروتکلی که برای ارتباط بین دو سوئیچ بر روی VSL اجرا می شود را Virtual Switch Link Protocol (VSLP) گویند. این پروتکل از دو زیر پروتکل Link Management Protocol (LMP) که وظیفه مدیریت لینک مانند جابه جایی عضوهای دامنه و پارامترهای مربوط به قابلیت های سوئیچ را بر عهده دارد و Role Resolution Protocol (RRP) که وظیفه انتخاب سوئیچ Active را بر عهده دارد، تشکیل می شود. اگر بسته ای وارد سوئیچ شود تا جایی که بتواند بسته را از لینک VSL عبور نمی دهد و سعی بر ارسال آن از طریق لینک های خود را دارد.

## ۱۷- DHCP Server

پروتکل Dynamic Host Configuration Protocol (DHCP) برای گرفتن آدرس IP از استخر آدرس ها به صورت خودکار است. این پروتکل در لایه Application قرار دارد و از پورت های 67, 68 UDP استفاده می نماید. این پروتکل ساختار کلاینت سروری دارد و مراحل آن بدین شکل است که ابتدا کلاینت درخواست DHCP Discover را broadcast می کند. این بدان معناست که به جست و جوی سرور برای دریافت آدرس می پردازد. سپس سرور پاسخ DHCP Offer را ارسال می کند. در واقع سرور به کلاینت پیشنهاد می دهد. بعد از آن کلاینت درخواست DHCP Request را به دو دلیل ارسال می کند، یکی برای مشخص کردن یک سرور از بین چند سرور و دیگری درخواست برای گرفتن آدرس IP. در آخر سرور پاسخ DHCP Acknowledge را ارسال می کند. در شبکه های کوچک ممکن است نیازی به DHCP Full-Scale نباشد. در این شبکه ها می توان از switch و یا

router استفاده نمود به دلیل اینکه IOS از همه ویژگی های DHCP پشتیبانی نمی کند. مراحل تنظیم به شکل زیر است:

```
(Config)#service dhcp
(Config)#ip dhcp pool <pool-name>
(Config-dhcp)#network <net-id>
(Config-dhcp)#default-router <default-gateway>
(Config-dhcp)#dns-server <dns-server-add>
(Config-dhcp)#lease <lease-time-duration>
(Config)#ip dhcp exclude-address <start-ip> <end-ip>
```

اگر کلاینت و سرور در یک VLAN نباشند و یا مسیریاب هایی در بین راه باشند مشکل مسیریابی ایجاد می گردد زیرا بسته ها به صورت broadcast است. پس باید به صورت unicast ارسال شود. با دستور زیر می توان این عیب را برطرف نمود.

```
(Config)#interface vlan <number>
(Config-if)#ip helper-address <ip-add>
```

Helper-address در فیلدی به نام GIADDR قرار دارد که مشخص می کند کدام VLAN بین چندین scope کدام آدرس را انتخاب کند.

```
(R1)#show ip dhcp binding
```

## ۱۸- Basic Manage Switch Port

دو ویژگی برای تنظیم راحت تر سوئیچ ها وجود دارد:

۱. Interface Range: مدیریت دو یا بیشتر پورت های سوئیچ از این طریق انجام می گیرد. همچنین می توان آنها را عضو یک گروه و هم زمان با هم تنظیم نمود. در این صورت از Macro می توان استفاده نمود و از یک اینترفیس به جای همه آنها برای مدیریت و تنظیم استفاده نمود.

```
(Config)#define interface-range <name-of-macro> <type-number> <number>
(Config)#interface range macro <name>
```

۲. Error-Disable: یک ابزاری است که هنگام وقوع یکسری عوامل پورت را به صورت خودکار غیر فعال می نماید. در این حالت پورت به حالت err-disable می رود. در واقع پورت از نظر فیزیکی مشکلی ندارد ولی بسته ها ارسال و دریافت نمی شوند. عواملی که ممکن است پورت را به این حالت ببرد با دستور زیر مشخص می شود:

```
(Config)#errdisable detect cause
```

عواملی که موجب رخ داد این چنین می شود ممکن است شامل BPDUGuard، ARP-Inspection، DTP-PaGP-Flap(etherchannel ports no longer identical) و... باشد. اگر قبل از دستور بالا "no" قرار گیرد عامل مشخص شده غیر فعال شده و پورت به حالت err-disable نمی رود. برای اینکه بعد از مدت مشخصی پورت به حالت فعال برود از دستورات زیر استفاده می شود:

```
(Config)#errdisable recovery cause <cause-name>
(Config)#errdisable recovery interval <time-to-recovery>
```

## ۱۹- Port Mirroring

کپی بسته ها از یک مبدأ به یک مقصد مشخص را Port Mirroring گویند. نام های دیگر آن Switched SPAN (Port Analyzer)، RSPAN(Remote SPAN) و Port Monitoring است. اگر بسته ای غیر عادی وارد شبکه شود و اختلالی به وجود آورد حل نمودن آن مشکل است. به همین دلیل واحدی با مشاهده و بررسی بسته ها این مشکلات را حل می نمایند. این عمل توسط Monitoring انجام می شود که یک کپی از بسته ها به یک مقصد مشخص ارسال می شود. برای تنظیم آن باید یک پورت مشخص به عنوان مبدأ و پورت دیگر به عنوان مقصد برای فرستادن کپی به آن انتخاب شود. دو مرحله برای تنظیم وجود دارد. ابتدا ایجاد یک Session و انتخاب مبدأ که بسته وارد آن می شود که نوع مبدأ در ادامه گفته می شود. سپس ارسال کپی از مبدأ به مقصد مشخص که همان ویژگی های مبدأ را دارد. سه نوع SPAN وجود دارد که در CCNP موارد ۱ و ۲ بررسی می گردد. ۱- Local or Port SPAN: مبدأ و مقصد بر روی یک سوئیچ باشند. ۲- RSPAN: اجازه می دهد که بسته ها از یک مبدأ بر روی یک سوئیچ به یک مقصد بر روی سوئیچ یا سوئیچ های دیگر ارسال شود. که لینک بین سوئیچ ها باید trunk باشد. ۳- ERSPAN(Enhanced RSPAN): کپی از یک مبدأ به یک مقصد در مکان دیگر. در موارد ۱ و ۲ فقط نیاز به دسترسی کامل به لایه ۲ وجود دارد ولی در مورد ۳ دسترسی کامل به لایه ۳ و ۲ نیاز است. که با GRE بسته بندی می شود. همچنین بیشتر سوئیچ ها از ERSPAN پشتیبانی نمی کنند. مبدأ می تواند شامل مواردی همچون: ۱- انتخاب یک یا چند پورت باشد. ۲- انتخاب طرف ارسال (RX, TX, Both) باشد. البته بعضی از سوئیچ ها فقط یک پورت را با TX ولی چند پورت را با RX پشتیبانی می کنند. ۳- انتخاب یک یا چند VLAN باشد. ۴- انتخاب remote vlan باشد. در صورتی که مقصد می تواند: ۱- انتخاب یک یا چند پورت باشد. ۲- انتخاب یک یا چند VLAN باشد. به صورت پیش فرض هنگام انتخاب پورت مقصد، پورت به حالت monitoring می رود و بسته ای ارسال نمی کند، تمام ویژگی ها غیر فعال می شود و فقط کپی بسته های دیگر را دریافت می کند. البته می توان با دستور ingress به حالت عادی برگرداند. تنظیمات مربوط به Local SPAN به شرح زیر است:

```
(Config)#monitor session <session-number> source <interface interface-id> <,->
<TX | RX | both> <vlan vlan-id> <remote vlan vlan-id>
```

```
(Config)#monitor session <session-number> destination <interface interface-id>
<,-> <encapsulation replicate | dot1q| ISL> <ingress> <remote vlan vlan-id>
```

در تنظیمات RSPAN این امکان به VLAN داده می شود تا RSPAN بین سوئیچ ها عبور کند. این VLAN فقط برای RSPAN استفاده شده و اگر پورت access در این VLAN باشد پورت به حالت غیر فعال می رود. در تنظیمات مربوط به RSPAN مبدأ مانند قبل است و مقصد به صورت زیر تنظیم می شود:

```
(Config)#monitor session <session-number> destination <remote vlan vlan-id>
<interface interface-id>
```

Reflector-port مکانیزمی است که بسته ها را به RSPAN VLAN کپی و ارسال می کند. در سوئیچ های 3750, 4500, 6500 نیاز به reflector-port نیست. این VLAN باید در تمامی سوئیچ ها تعریف شده باشد و remote-span بر روی آن فعال شود. مشکلی که در reflector-port ایجاد می شود این است که پورت غیر فعال، از حالت trunk خارج، LED پورت خاموش می شود و یک پورت دیگر باید به عنوان جابه جایی بین سوئیچ ها استفاده شود. در حالت SPAN نمی توان فیلتری انجام داد ولی ویژگی vlan access map این عمل را انجام می دهد.

## ۲۰- Common Spanning Tree 802.1d

به طور خلاصه این پروتکل برای تغییر وضعیت شبکه و جلوگیری از حلقه است. برای اولین بار شرکت Digital Equipment Corporation (DEC) Net آن را ثبت نمود و بعدها تبدیل به استاندارد IEEE 802.1d شد. علت اصلی ایجاد این پروتکل حل مشکل حلقه در شبکه است. سوئیچ ممکن است یک MAC را از چندین مسیر چندین بار حذف و یا بازنویسی نماید. زمانی که سوئیچ روشن می شود یا پورتی فعال می شود این پروتکل شروع به کار می کند و پیامی به نام BPDU ارسال می کند. که دارای فیلدهای زیر است:

| Protocol Identifier | Ver | Msg Type | Flags | Root ID | Root Path Cost | Bridge ID | Port ID | Msg Age | Max Age | Hello Time | Forward Delay |
|---------------------|-----|----------|-------|---------|----------------|-----------|---------|---------|---------|------------|---------------|
|---------------------|-----|----------|-------|---------|----------------|-----------|---------|---------|---------|------------|---------------|

۱- Protocol Identifier: نوع پروتکل STP را مشخص می کند که داخل 802.3 قرار می گیرد. اندازه این فیلد 2 Bytes است.

۲- Version: نسخه پروتکل را نمایش می دهد و اندازه آن 1 Bytes است.

۳- Message Type: نوع بسته را مشخص می کند که می تواند Configuration, regular BPDU, Topology Change و یا Topology Change Verification باشد. اندازه آن 1 Bytes است.

۴- Flags: نشان دهنده این است که در حال تغییر topology است یا خیر. اندازه آن 1 Bytes است.

۵- Priority:Root ID: مربوط به سوئیچ root است که به صورت پیش فرض 32768 است. انداز آن 8 Bytes است.

۶- Root Path Cost: بر پایه پهنای باند تنظیم می شود که به ترتیب  $10\text{Mb/s}=100$ ,  $100\text{Mb/s}=19$ ,  $1\text{Gb/s}=4$ ,  $10\text{GB/s}=2$  است. انداز آن 4 Bytes است.

۷- Bridge ID: شماره Priority سوئیچ ارسال کننده بسته است که به صورت پیش فرض 32768 است. انداز آن 8 Bytes است.

۸- Port ID: شماره پورت ارسال کننده بسته است. انداز آن 2 Bytes است.

۹- MSG Age: تعداد Hop های گذرانده شده است. انداز آن 2 Bytes است.

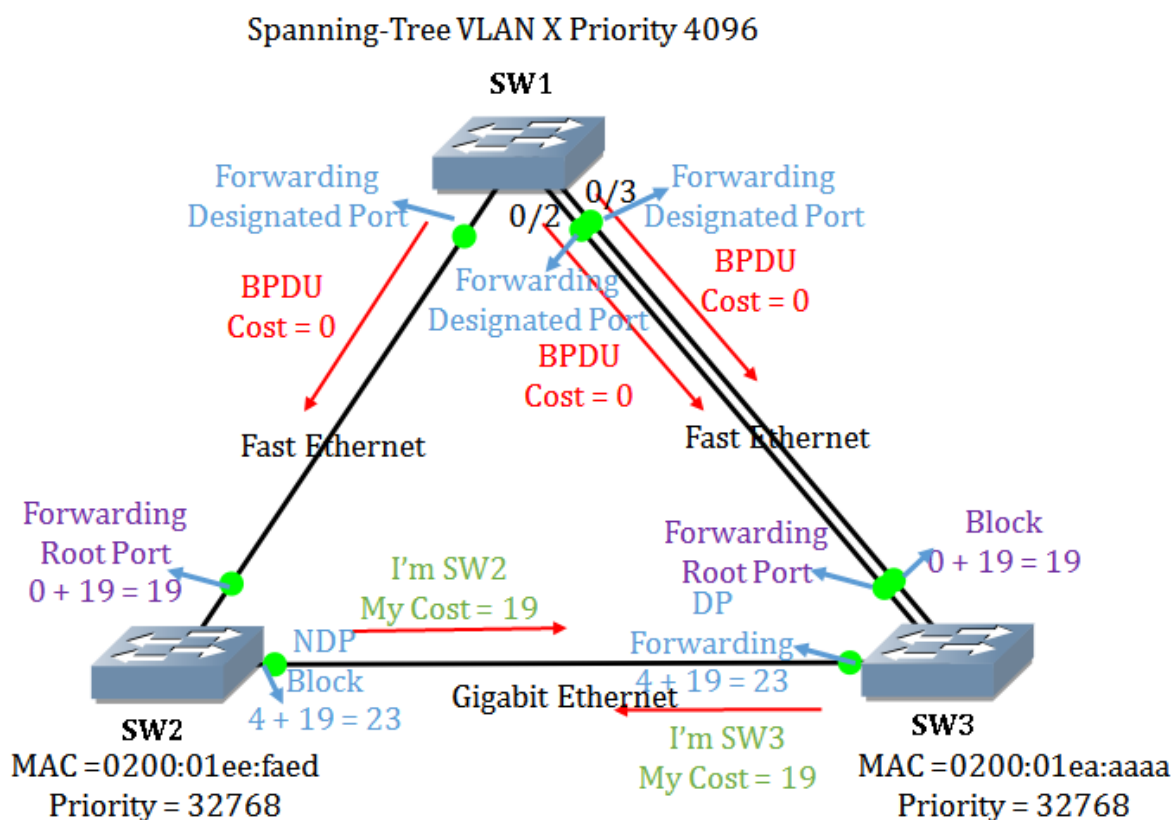
۱۰- MAX Age: حداکثر تعداد Hop است که پیش فرض برابر است با 20. انداز آن 2 Bytes است.

۱۱- Hello Time: هر چند ثانیه یکبار BPDUs ارسال شود که پیش فرض برابر است با 2 Sec. انداز آن 2 Bytes است.

۱۲- Forwarding Delay: زمانی که پورت فعال می شود به حالت Listening رفته و به خط گوش می دهد بعد به حالت learning می رود تا یاد بگیرد. پس از آن به حالت Forwarding می رود که پیش فرض برابر 15 Sec است و اطلاعات خود را ارسال می کند. انداز آن 2 Bytes است.

پس از ارسال بسته BPDUs باید سوئیچ root و نقش و حالت هر پورت مشخص شود. سوئیچ با پایین ترین Priority برنده و Root Bridge می شود و در صورت مساوی بودن کوچک ترین MAC به عنوان سوئیچ برنده انتخاب می شود. برای تنظیم دستی Priority نیز می توان از دستور `spanning-tree vlan <vlan-id> priority <number>` استفاده نمود و یا در آخر به جای عدد دستور `root priority` را وارد نمود، ولی به شرط آن که سوئیچ دیگری root نباشد. هر پورت دو نقش Designated و Root و سه حالت Learning و Forwarding و Disable دارد. این نقش ها و حالت ها در پروتکل های مختلف RSTP و MSTP و PVSTP تفاوت دارد. تمامی پورت های سوئیچ root به صورت forwarding و designated هستند. در هر collision domain حتماً نیاز به Designated Port است. وظیفه DP ارسال بسته BPDUs روی خط است. انتخاب Root Port بدین شکل است که کوتاه ترین مسیر برای رسیدن به سوئیچ root است. سوئیچ ها هزینه ها را برای یکدیگر ارسال می کنند و بهترین انتخاب می شود و اگر دو طرف لینک DP شد یکی غیر فعال می شود که بر اساس bridge id و MAC است. یک طرف در حالت DP و طرف دیگر در حالت NDP قرار می گیرد و Block می شود. اگر بین سوئیچ root و سوئیچ دیگری چند لینک باشد پورتهای فعال می ماند که از کمترین port-id سوئیچ root آمده باشد. که در بسته BPDUs مشخص شده است. بعد از زمان Learning پورت به یکی از سه حالت blocking, forwarding و

disable می‌رود که در این حالت پیامی روی root port ارسال می‌شود و تا زمانی که Acknowledge آن نیاید همچنان ارسال می‌شود. به این پیام Topology Change Notification(TCN) گفته می‌شود. اگر پورت در حالت port fast باشد این پیام ارسال نمی‌شود.



در شکل بالا پورت پورت سمت راست سوئیچ ۳ به دلیل اینکه به پورت ۳ سوئیچ ۱ متصل است به حالت Block می‌رود. زیرا شماره پورت سوئیچ ۱ بیشتر است. پورت سوئیچ ۳ که بین دو سوئیچ ۲ و ۳ است به دلیل MAC کوچک‌تر به حالت Forwarding می‌رود.

بسته TCN شامل فیلدهای زیر است:

|                    |             |         |                    |         |     |
|--------------------|-------------|---------|--------------------|---------|-----|
| 802.3/802.2 Header | Protocol ID | Version | Message Type(0x80) | Padding | FCS |
|--------------------|-------------|---------|--------------------|---------|-----|

هنگام دریافت بسته TCN توسط root عملیات‌های زیر صورت می‌گیرد:

۱. تغییر Flag به TC-Flag در بسته BPDU صورت می‌گیرد.
۲. این عمل را تا زمان  $\text{forwarding-delay} + \text{max-age} = 15 + 20 = 35$  ادامه می‌دهد.

۳. زمان جدول mac را به  $\text{forwarding-delay} = 15$  کاهش می دهد (aging time).

هنگام دریافت بسته TCN توسط دیگر سوئیچ ها عملیات های زیر رخ می دهد:

۱. زمان Content Addressable Memory (CAM) را به  $\text{forwarding-delay}$  کاهش می دهند.
  ۲. بعد از آن حافظه Media Access Control (MAC) را flush و زمان را دوباره به 5 MIN تغییر می دهند.
- زمانی که دستور `show spanning-tree detail` وارد شود قسمت `topology change flag set` در همه سوئیچ ها قابل مشاهده و به معنای تغییر است و `detect flag set` فقط در سوئیچ root قابل مشاهده و به معنای در حال تغییر `topology` است.

## ۲۰-۱-802.1d Optimize and Security

در این قسمت ۴ مورد برای تنظیمات بهتر `spanning-tree` گفته می شود:

۱. `Port-Fast`: تنظیم این دستور بر روی سوئیچ به معنای عدم ارسال BPDUs است. بیشتر برای `end device` ها یا مواردی است که از عدم وجود حلقه اطمینان باشد، استفاده می شود. این دستور عملیات های `Listening`, `Learning` را انجام نمی دهد و سریع پورت را به حالت UP می برد. این دستور هم در حالت `global` و هم بر روی `interface` تنظیم می شود. در حالت `global` به صورت پیش فرض برای همه پورت ها است.

```
(Config)#spanning-tree portfast
```

```
(Config-if)# spanning-tree portfast
```

۲. `UplinkFast`: اگر در یک سوئیچ دو پورت وجود داشته یکی `root port` و دیگری `designated port` باشد و `root port` به هر دلیلی `down` شود و DP قابلیت `root port` شدن را داشته باشد یعنی مسیری دیگر برای رسیدن به سوئیچ `root` وجود داشته باشد این پورت به حالت `root port` می رود ولی ۳۰ ثانیه طول می کشد ( $\text{listening} + \text{learning} = 15 + 15 = 30$ ). با وارد کردن این دستور در قسمت `global` پورت دیگر را جایگزین `root port` می کند. حال اگر سوئیچ های دیگر، از پورت قبلی آدرس های MAC را یاد گرفته باشند با ارسال بسته `Dummy Multicast` به سوئیچ های دیگر برای تغییر `mac address` ها به این لینک این مشکل حل می شود و بسته به شکل `src=mac address in table switch` و `des=multicast cisco address` می شود. این دستور بیشتر برای سوئیچ های `Access-Layer` طراحی شده است و بر روی سوئیچ های `Distributed` و `Core` کمتر دیده می شود، زیرا ارسال بسته چند پخش در شبکه باعث بالا رفتن `load` سوئیچ می شود و همچنین هزینه لینک و `bridge-id` را افزایش

می دهد (cost=3019 > cost=19-0) که خود این برای تصمیم گیری سوئیچ مشکل ایجاد می کند. زمان multicast به صورت پیش فرض 150 packet per second است.

```
(Config)#spanning-tree uplinkfast max-update-rate <pps>
```

```
(SW1)#debug spanning-tree uplinkfast
```

۳. BackboneFast: اگر سوئیچی دو لینک بین root و سوئیچ دیگر داشته باشد و پورت متصل به سوئیچ root به عنوان Root Port و دیگری به عنوان Designated Port باشد و پورت آن طرف DP و سوئیچ دیگر در حالت Block باشد. اگر RP به هر دلیلی down شود، سوئیچ شروع به ارسال بسته های BPDUs از طریق DP می کند. در حالت عادی سوئیچ دیگر ۲۰ ثانیه فرصت می دهد تا مشکل حل شود، بعد از آن ۳۰ ثانیه برای listening و Learning مصرف می شود و BPDUs صحیح را از root دریافت و ارسال می کند. در این ۴۵ ثانیه پورت به حالت down است. در حالت BackboneFast اگر از طرف DP به Block port بیاید و این بسته BPDUs صحیح و عادی باشد و اطلاع دهد که سوئیچ root است به معنای این است که یا سوئیچ جدید وارد شده است و یا RP به حالت down رفته است. سوئیچ دیگر با ارسال بسته Root Link Query (RLQ) به root و دریافت RLQ Response و ارسال به سوئیچ مشکل دار از زمان ۲۰ ثانیه انتظار می کاهد ولی هنوز ۳۰ ثانیه دیگر پورت down است. این تنظیم برای تشخیص خطا در root port است. دستور زیر بر روی همه سوئیچ ها باید وارد شود.

```
(Config)#spanning-tree backbonefast
```

۴. Loop Guard: در مورد قبل گفته شد اگر طرفی در حالت DP و طرف دیگر در حالت Block باشد در حالت عادی بعد از ۲۰ ثانیه به حالت listening و learning می رود. در حالت loop guard به پورت اجازه up شدن داده نمی شود و باید در همین حالت Block بماند. در حالت loop guard با ارسال BPDUs و عدم دریافت آن (سکوت طرف مقابل) پورت به حالت Block می رود. Uni-Directional Link به لینکی گفته می شود که BPDUs ارسال شود ولی جوابی دریافت نشود. در حالت Loop Guard اگر پروتکل routing و یا هر چیز دیگر بر روی لینک فعال باشد پورت down می ماند و تشخیص uni-directional link فقط با استفاده از BPDUs است. مکانیزم دیگری به نام UDLD (Uni-Directional Link Detection) وجود دارد که فقط به BPDUs بستگی ندارد و از اطلاعات دیگری نیز استفاده می کند. در هر دو سوئیچ این قابلیت باید فعال شود و با استفاده از بسته ها در لایه ۲ و همچنین لایه ۱ این تشخیص انجام می شود. بیشتر برای کابل های Fiber Optic استفاده می شود، چون یک لینک برای ارسال و یک لینک برای دریافت استفاده می شود. پورت بعد از تشخیص به حالت disable می رود و errdisable می دهد. UDLD مانند CDP است که اطلاعات همدیگر را برای یکدیگر ارسال می کنند. با دستور show udld

neighbors می‌توان همسایه‌های آن را مشاهده نمود. این بسته‌ها هر ۱۵ ثانیه ارسال می‌شود و با ارسال ۳ بسته (۴۵ ثانیه) و عدم دریافت آن پورت خاموش می‌شود. دو حالت aggressive و normal دارد که در حالت اول بعد از تشخیص پورت به حالت down می‌رود و در حالت normal پورت disable نمی‌شود.

```
(Config)#udld enable aggressive <Time-To-Block>
(Config)#udld enable
(SW1)#show udld <interface>
```

## ۲۰-۲- 802.1w Rapid Spanning Tree Protocol (RSTP)

قبل از این که سوئیچ طراحی شود پروتکل ST وجود داشته ولی در آن زمان سرعت لینک‌ها و این پروتکل بسیار کند بوده است و برای Bridge ها طراحی شده بود. بعدها که لینک‌هایی با سرعت و همگرایی بالا ایجاد شد پروتکل Rapid را طراحی کردند که پیشرفته پروتکل 802.1d است.

تفاوت‌هایی که پروتکل RSTP و 802.1d دارد به شرح زیر است:

۱. پروتکل RSTP سرعت بالاتری نسبت به 802.1d دارد.
۲. در 802.1d نوع لینک Full یا Half اهمیتی ندارد. در 802.1w باید حتماً Full باشد تا سرعت بیشتری داشته باشد. در این پروتکل Full Duplex به معنای p2p و Half Duplex به معنای share است.
۳. در 802.1d با قطع پورت متصل به root bridge دیگر سوئیچ‌های متصل به آن اجازه ارسال BPDU به بقیه را ندارند ولی در 802.1w اگر مشکلی ایجاد شود بقیه سوئیچ‌ها بسته BPDU را به عنوان Keep-Alive یا همان Hello ارسال می‌کنند. اگر بین دو سوئیچ چندین لینک وجود داشته باشد بسته‌های BPDU Keep-Alive، هم از پورت‌های Forwarding و هم از پورت‌های Blocking دریافت می‌شود. در 802.1d زمان Max Age=20 sec به معنای عدم دریافت ۱۰ بسته Hello است و پس از آن متوجه Neighbor Dead می‌شود. در 802.1w بعد از عدم دریافت ۳ بسته یا ۶ ثانیه متوجه Neighbor Dead می‌شود.

Roles یا نقش پورت‌ها در RSTP به شرح زیر است:

۱. Root Port: همانند پروتکل 802.1d است.
۲. Designated Port: همانند پروتکل 802.1d است.
۳. Backup Port: به معنای Backup DP است. زمانی که سوئیچی بسته BPDU خودش را از پورت دیگری که در حالت Block است دریافت کند (به طور مثال چندین پورت به Hub متصل باشد). در این صورت

Bridge-ID برابر خود سوئیچ و Port-ID یکی از پورت‌های سوئیچ که Non Block باشد است. در واقع پورت‌ها در یک Collision Domain هستند.

۴. Alternate Port: اگر پورت Block بسته BPDUs از سوئیچ دیگری دریافت کند به آن Alternate Port گویند.

۵. Edge Port: همان فعال نمودن port fast است. به دلیل استفاده در topology change با نام دیگری نام گذاری شده است.

پورت‌های RSTP در وضعیت Blocking و Learning آدرس‌های MAC را ذخیره نکرده و Discard می‌کنند. وضعیت پورت‌ها در RSTP به شرح زیر است:

۱. Forwarding: اطلاعات را ارسال و دریافت می‌کند.
  ۲. Learning: فقط آدرس‌های MAC را یاد می‌گیرد و اطلاعاتی ارسال و دریافت نمی‌کند.
  ۳. Discarding: ترکیبی از حالت‌های blocking, disable, learning در 802.1d است.
- ویژگی RSTP UpLinkFast به عنوان یک ویژگی دلخواه در 802.1w اضافه گردیده است و دیگر هزینه لینک، افزایش bridge-id و همچنین Dummy Multicast را ندارد.
- تفاوت 802.1d BPDUs و 802.1w BPDUs به شرح زیر است:

۱. Version field در 802.1d برابر ۰ و در 802.1w برابر ۲ است.
۲. فیلد Flag که ۸ بیتی است در 802.1d فقط از بیت ۱ و ۸ به عنوان topology change و topology change acknowledge استفاده می‌شود ولی در 802.1w به شکل زیر است.

| 0               | 1        | 2         | 3 | 4        | 5          | 6         | 7                   |
|-----------------|----------|-----------|---|----------|------------|-----------|---------------------|
| Topology Change | Proposal | Port Role |   | Learning | Forwarding | Agreement | Topology Change ACK |

Port Role: این قسمت شامل دو بیت است که چهار حالت را شامل می‌شود. 00 به معنای Unknown است. 01 به معنای Alternate/Backup است. 10 به معنای Root است. 11 به معنای Designated است.

Proposal و Agreement فقط برای لینک‌های full duplex است. Proposal به معنای خواستن برای root bridge شدن است که اول همه آن را تنظیم می‌کنند. سوئیچ با bridge-id بیشتر با پیغام Agreement به سوئیچ

دیگر اطلاع می دهد که به عنوان root انتخاب شده است، زیرا bridge-id پایین تری دارد. این عملیات تا زمانی که یک root واحد ایجاد شود ادامه می یابد. این عملیات بسیار سریع است چون حالت listening , learning ندارد و سوئیچ سریع تر به حالت forwarding و DP می رود و سوئیچ دیگر به حالت F و RP می رود.

Topology Change in RSTP: فقط زمانی اجرا می شود که یک پورت non-edge به حالت forwarding برود و در هیچ صورت دیگر اجرا نمی شود. ابتدا flag TC برابر ۱ می شود سپس همانند انتخاب root، سوئیچ بسته ای به تمام پورت های non-edge(DP,RP) ارسال می کند. بعد از آن حافظه CAM برای پورت هایی که این بسته را ارسال کرده است (نه دریافت کرده باشد) پاک می شود. بعد از ارسال بسته (BPDU TC) به مدت زمان 2 hello این عمل را تکرار و متوقف می شود.

### ۳-۲۰- Per Vlan Spanning Tree (PVST)

پروتکل 802.1d قبل از آن که VLAN وجود داشته باشد پیاده سازی و اجرا شده بود. در واقع تمامی سیستم های داخل شبکه در یک broadcast domain قرار داشتند. پس از پیاده سازی و اجرای VLAN این پروتکل تغییر نکرد و با vlan سازگار نشد. شرکت سیسکو نیاز به پروتکلی داشت که برای هر vlan بتوان ST را اجرا کند. پروتکل 802.1w نیز برای Common Spanning Tree(CST) طراحی شده بود که با VLAN هیچ ارتباطی نداشت و برای همه VLAN ها بود. شرکت سیسکو این دو پروتکل را برای هر VLAN سازگار کرد. پروتکل PVST که از 802.1d BPDU استفاده می نمود. پروتکل RPVST یا همان Rapid Per VLAN STP که از 802.1w BPDU استفاده می نمود. همچنین پروتکل Multiple Spanning Tree با استفاده از 802.1s BPDU را ارائه نمود. پروتکل RPVST با 802.1d نیز سازگار است. اگر سوئیچی 802.1w را پشتیبانی نکند می تواند با 802.1d اطلاعات را جابه جا کند. اگر سوئیچی با Bridge-ID کمتری وارد مدار شد و خواست به عنوان root باشد، ابتدا سوئیچ متصل به آن همه پورت های Non-Edge دیگر خود به غیر از لینک متصل به سوئیچ بهتر را در حالت Block می برد و این سوئیچ را Agreement می کند و سپس بسته BPDU را به تمامی همسایه های خود ارسال می کند. این عملیات بسیار سریع انجام می گیرد.

```
(Config)#spanning tree mode rapid-pvst
```

```
(Config-if)#spanning-tree link-type point-to-point
```

دستور بالا برای لینک های Half Duplex است.

```
(SW1)#show spanning-tree summary
```

## ۲۰-۴- Multiple Spanning Tree (MST) 802.1s

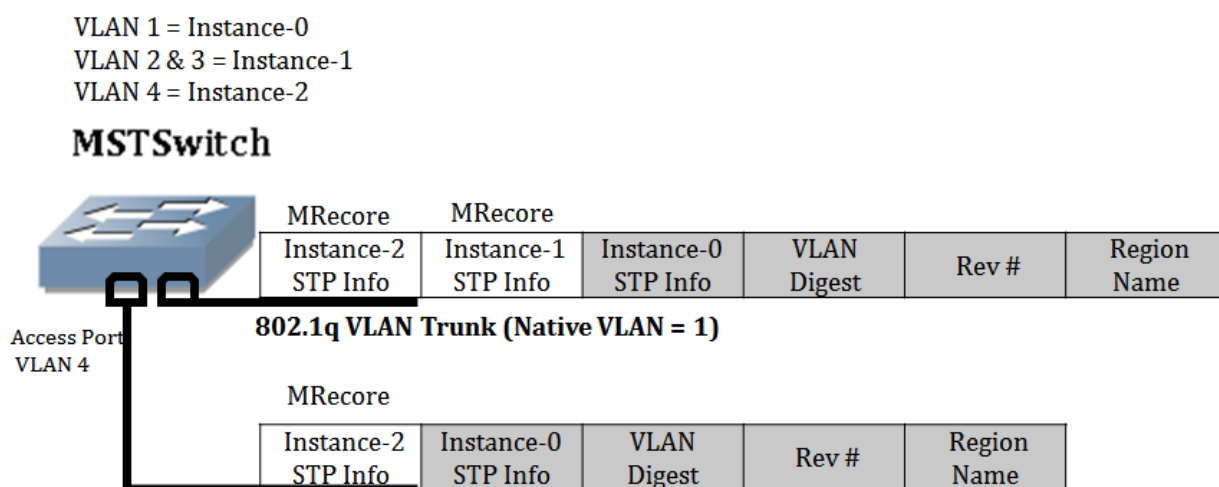
شرکت سیسکو از استاندارد IEEE 802.1s استفاده نموده و بر روی سیستم‌های خود مطابقت داده است. پروتکل مشابه دیگری به نام MIST (Multiple Instance Spanning Tree) وجود دارد که با MST کمی متفاوت است. در پروتکل‌های CST (Common ST) یک سوئیچ به عنوان root وجود دارد و بقیه سوئیچ‌ها خود را با آن مطابقت می‌دهند. مشکلی که وجود دارد این است که load balancing وجود ندارد. در پروتکل‌های Per-VLAN یک VLAN از یک پورت و VLAN دیگر از پورتی دیگر عبور می‌کند. IEEE تا قبل از MST هیچ پروتکل Per-VLAN ارائه نداده است ولی شرکت سیسکو این کار را انجام و با دستگاه‌های خود مطابقت داده است. حال اگر در بین راه سوئیچی غیر از سیسکو وجود داشت سوئیچ سیسکو به ازای هر VLAN یک BPDUD ارسال می‌کند ولی سوئیچ غیر سیسکو فقط BPDUDهای Native را عبور داده و بقیه را drop می‌کند چون از per-vlan پشتیبانی نمی‌کند. در نتیجه سوئیچ‌های سیسکو نیز load balancing انجام نمی‌دهند چون سوئیچ‌های بین راه غیر سیسکو اجازه یادگیری تمام نقشه شبکه را نمی‌دهند. شرکت سیسکو برای این مشکل نیز راه حلی ارائه داد. یک آدرس چند پخش‌ی رزرو شده انتخاب نمود و بسته‌های BPDUD را به این آدرس ارسال نمود. این سوئیچ‌های غیر سیسکو هنگام دریافت فقط آن را به تمامی همسایه‌ها ارسال می‌کنند که PVST+ نام دارد. این آدرس Cisco Proprietary Multicast Address نام دارد. بعدها IEEE پروتکلی طراحی نمود تا قابلیت load balancing بین VLAN‌ها و چندین ST topology وجود داشته باشد. یک پورت در هر topology ممکن است چندین وضعیت مختلف داشته باشد. این نکته را باید توجه داشت که هر PVST باری بر روی CPU دارد و ممکن است برای تعداد VLAN پایین مشکلی ایجاد نشود ولی اگر حدود 100 vlan و برای هر VLAN یک topology وجود داشته باشد، ممکن است سوئیچ دچار مشکل شود. سوئیچ‌های قوی‌تر در لایه‌های بالاتر از VLAN‌های بیشتر پشتیبانی می‌کنند ولی با این حال از تعداد محدودی PVST پشتیبانی می‌کند. به‌طور مثال 2950 SI از 128 vlan و 64 ST Topology پشتیبانی می‌کند. 2950 EI از 250 vlan و 64 ST Topology 3550, 3750 از 1005 vlan و 128 ST Topology پشتیبانی می‌کند. به هر یک از این topology‌ها یک Instance گفته می‌شود. اگر بیشتر از تعداد مجاز instance، topology ایجاد گردد، STP مختل شده و دیگر از ایجاد حلقه جلوگیری نمی‌کند. مشکلی که در پروتکل‌های per-vlan وجود داشت به ازای هر VLAN یک Instance وجود داشت ولی در MST نیازی نیست به ازای هر VLAN یک Instance ایجاد نمود. ممکن است به ازای چند VLAN یک Instance وجود داشته باشد. در واقع یک Instance ساخته می‌شود و vlan‌ها به آن assign می‌شوند و root bridge و port roles برای هر Instance تعریف می‌شود. می‌توان Instance را به عنوان Master و VLAN را به عنوان Slave نام گذاری نمود. Instance پیش فرض Instance-0 نام دارد و اگر تنظیمات پیش فرض تغییر نکند، هیچ گونه

load balancing وجود ندارد چون تمامی VLAN ها به صورت پیش فرض به این Instance تعلق دارند. اگر load balancing نیاز باشد باید Instance های دیگری ایجاد و VLAN ها را به آن assign نمود. Instance-0 را (Internal Spanning Tree) IST نامیده و بقیه Instance های ایجاد شده را MSTI (Multiple Spanning Tree Instance) می نامند. حداکثر تعداد Instance ایجاد شده 16 است. زمانی که MST فعال شود پروتکل RSTP نیز به صورت خود کار فعال شده و مانند VTP با همسایه خود ارتباط برقرار کرده و پارامترهایی که باید با هم برابر باشد را بررسی می کند و اگر برابر نباشد ارتباط برقرار نمی شود. پارامترهای MSTP شامل موارد زیر است:

۱. MST Region Name: مانند VTP Domain فقط یک نام است.
۲. MST Revision Number: در VTP با اضافه و حذف یک VLAN بسته VTP ارسال می شود و revision آن یکی اضافه می شود و کاربر به آن به صورت مستقیم دسترسی ندارد. در MST به صورت دستی تنظیم و باید در همه یکی باشد.
۳. VLAN to Instance Mapping: از Instance و Vlan ها، یک Hash با MD5 گرفته شده و در بسته BPDU قرار می گیرد. اگر این Hash برابر بود ارتباط برقرار می شود. در MST متغیرهای MST و هزینه ها به ازای هر Instance تعریف می شود نه به ازای هر VLAN و تمامی تعریف ها از per-vlan به Instance تغییر می یابد. هزینه در MST به سرعت لینک بستگی دارد. RSTP از مقدارهای قدیمی ولی MST از مقدارهای جدید استفاده می کنند.
- تفاوت BPDU در MST و PVST بدین شرح است که در PVST برای هر VLAN یک بسته BPDU ایجاد و ارسال می شود و برای Native VLAN ها تگ 802.1q ندارد و برای بقیه VLAN ها این Tag زده می شود. در MST برای تمامی Instance ها فقط یک بسته BPDU ارسال می شود. این بسته شامل موارد زیر است:

۱. Region Name
۲. Revision Number
۳. Vlan Digest که این سه قسمت باید در تمام سوئیچ ها یکسان باشد.
۴. Instance-0 Info: در تمامی بسته ها به صورت پیش فرض وجود دارد.
۵. MRecord: اطلاعات هر Instance در این قسمت است و به ازای هر Instance یک MRecord وجود دارد.

اطلاعات Instance شامل Root Bridge، Bridge-ID و Cost to Root است. در لینک‌های Trunk به صورت پیش فرض تمامی VLAN‌ها اجازه عبور دارند و همچنین تمامی Instance‌ها می‌توانند عبور کنند. Instance-0 به علت پیش فرض بودن و به علت وجود سوئیچ‌های Non-Cisco در بین راه، RSTP یا قابلیت‌های دیگر بر روی تمامی لینک‌ها چه Access و چه Trunk، اجازه عبور دارد.



بعضی از ویژگی‌های MST به شرح زیر است:

۱. پارامترهای زمانی فقط برای Instance-0 قابل تنظیم است و بقیه Instance‌ها از آن پیروی می‌کنند.
۲. Instance-0 MST برای پشتیبانی از RSTP و 8021.d است.
۳. Boundary Port: به پورتهایی گفته می‌شود که بعضی از متغیرها صحیح نباشد و MST روی آن فعال نشود. به طور مثال آن طرف لینک RSTP فعال باشد. در این حالت ارسال و دریافت به صورت Non-MST است و به صورت هماهنگ با آن طرف لینک انجام می‌گیرد. اگر چند Instance روی این نوع پورت عبور کند که آن طرف لینک RSTP یا CST باشد تمامی بر روی Instance-0 قرار می‌گیرد و load balancing روی این پورت‌ها انجام نمی‌گیرد. دستورات به صورت زیر است:

```
(Config)#spanning-tree mode mst
(Config)# spanning-tree mst configuration
(Config-mst)#name <name>
(Config-mst)#revision <number>
(Config-mst)#instance <number> vlan <vlan-id>
(Config)# spanning-tree mst <number>
(Config-mst)#[priority | root]
(Config-if)# spanning-tree mst <number> [cost | port-priority]
```

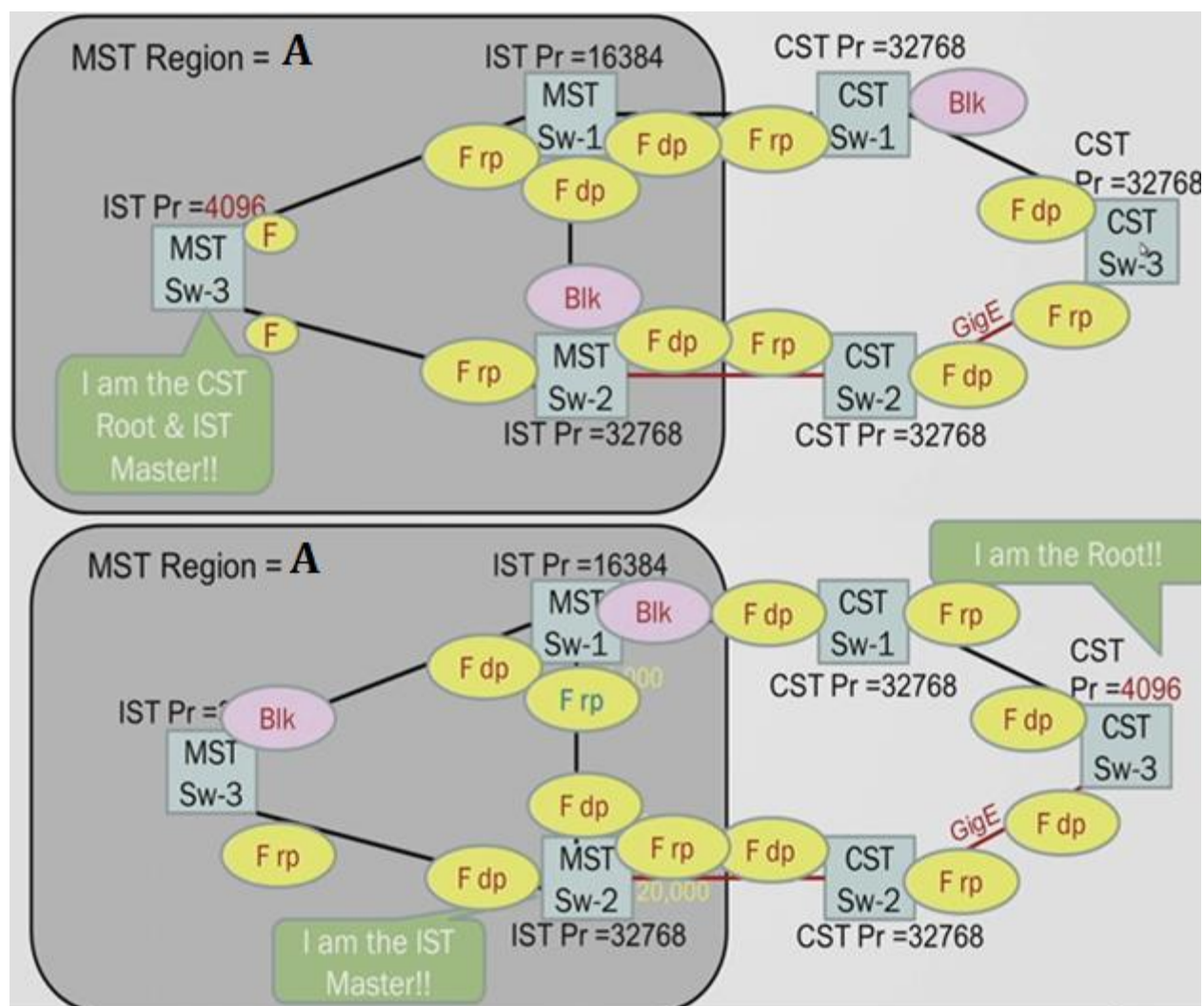
در Boundary Port هر تصمیمی که برای Instance-0 اتخاذ شود بر روی تمامی Instance های بر روی آن پورت اعمال می شود. زیرا فقط Instance-0 قابلیت صحبت با Non-MST ها را دارد. برای MSTI ها هیچ نگرانی بابت انتخاب Root Bridge نیست زیرا هیچ ارتباطی با محیط خارج یا همان Non-MST ها ندارند. ولی برای Instance-0 این مشکل وجود دارد. برای حل این مشکل دو Root Bridge انتخاب می شود یکی برای داخل MST و یکی برای محیط Non-MST. نحوه انتخاب به صورت زیر است:

1. Main Root Bridge: سوئیچی انتخاب می شود که در داخل و خارج کمترین Priority را داشته باشد.
2. CIST (Common Instance Spanning Tree) Regional Root (IST Master): سوئیچی با کمترین Priority داخل محیط MST انتخاب می شود و Root Port کمترین هزینه تا سوئیچ IST Master است. این دو سوئیچ دو حالت می توانند داشته باشند. یک کمترین Priority سوئیچ در داخل کل محیط MST باشد. در این صورت Main Root = IST Master. دو اگر کمترین Priority سوئیچ خارج از MST باشد آن گاه IST Master سوئیچی است با کمترین هزینه برای رسیدن به Main Root.

اگر زمانی قصد بر مهاجرت به MST باشد گاهی اوقات مشکل به وجود می آید که بعضی از سوئیچ ها MST را پشتیبانی نمی کنند. روش پیشنهادی این است که این سوئیچ ها در لایه Access قرار گیرند چون نیازی به دانستن تمامی VLAN ها نیستند، پس در محیط CST قرار می گیرند. سوئیچ های Core که از MST پشتیبانی می کنند و نیاز به دانستن تمامی VLAN ها دارند در محیط MST قرار می گیرند. تفاوتی برای سوئیچ های MST در نوع پورت (Boundary و Internal) برای ارسال هزینه وجود ندارد و در هر حال از استاندارد جدید (802.1t) استفاده می کنند و از نسخه قدیمی (RSTP) پشتیبانی نمی کنند. در Boundary Port فقط یک بسته BPDU ارسال می شود. در PVST به ازای هر Vlan یک بسته BPDU ارسال می شود و سوئیچ های Non-Cisco فقط یکی را دریافت و بقیه را حذف می کنند. Boundary Port ممکن است یک سوئیچ Non-Cisco باشد بنابراین شرکت سیسکو پروتکل MST را سازگار با PVST پیاده سازی نمود. کارهایی که سوئیچ بعد از تشخیص اتصال به PVST باید انجام دهد: به شرح زیر است:

1. پورت به عنوان Boundary علامت زده شود.
2. به ازای هر vlan که بر روی پورت عبور می کند یک بسته IST BPDU ارسال شود، زیرا فقط یک بسته می توان ارسال نمود.

۳. IST بر روی آن پورت به حالت Control Port State می‌رود. در زیر نحوه انتخاب دو حالت Root Bridge ترسیم شده است.



## ۲۱- High Availability

این موضوع برای حفظ افزونگی یا Redundancy و تحمل‌پذیری در برابر خطا یا Fault Tolerance برای مسیریاب‌ها است. در واقع زمانی که یک مسیریاب از بین رود مسیریاب دیگر وارد مدار شده و وظیفه آن را بر عهده می‌گیرد. تنظیمات مربوط به این پروتکل در قسمت Routing Interface و جزء خانواده First Hop Redundancy Protocol است. در معماری سه لایه‌ای سیسکو Access/Distribution/Core در لایه Access، VLAN‌ها و Broadcast Domain مشخص می‌شود. Broadcast Domain هنگامی پایان می‌یابد که به لایه Distribution می‌رسد که این یکی از وظایف routing است. در واقع در لایه Distribution سوئیچ‌های چندلایه استفاده می‌شود. باید توجه داشت فقط از router یا Swiych استفاده نمی‌شود بلکه از

سوئیچ‌های چند لایه استفاده می‌شود. در High Availability یک Virtual Interface ایجاد می‌شود که عملیات مسیریابی و Default Gateway را به PC و End Device‌ها ارائه می‌نماید. پروتکل‌های مربوط به High Availability ارائه دهنده در سوئیچ‌های چند لایه شامل HSRP، VRRP و GLBP می‌باشند.

## ۱-۲۱- Hot Standby Router Protocol (HSRP)

این پروتکل مخصوص سیستم‌های سیسکو است. از پروتکل UDP با پورت 1985 و آدرس چندپخشی 224.0.0.2 استفاده می‌کند. از این آدرس برای ارسال Hello به یکدیگر و صحبت با هم در پروتکل HSRP استفاده می‌شود و برای end device‌ها مفید نمی‌باشد. مسیریاب‌های HSRP دو نقش دارند یکی Active و دیگری Standby و بقیه سوئیچ‌ها در حالت Listening می‌مانند. End device‌ها بسته‌های خود را به سوئیچ Active ارسال کرده و این سوئیچ به آنها سرویس می‌دهد. سوئیچ Standby وضعیت را بررسی نموده و در صورت غیرفعال شدن Active جایگزین آن می‌شود. سوئیچ‌های دیگر در حالت Listening قرار می‌گیرند و به صحبت‌های Active و Standby گوش می‌دهند تا در صورت نیاز نقش خود را تغییر دهند. سوئیچ active سوئیچی است که بیشترین Priority را داشته باشد و در صورت مساوی بودن آن بیشترین IP انتخاب می‌شود. Default priority=100 است. هنگامی که HSRP بر روی interface تنظیم می‌شود یک group number باید وارد نمود و این به معنای قرار گرفتن در یک گروه است و در همه سوئیچ‌ها که در یک broadcast domain هستند باید یکسان باشد تا بتوانند با یکدیگر صحبت کنند. پس از تنظیم HSRP باید یک Virtual IP به عنوان default gateway وارد نمود که قاعدتاً باید در ردیف همان شبکه باشد. پس از تنظیم شماره گروه به صورت خودکار آدرس MAC به HSRP اختصاص داده می‌شود که به صورت 0000.0C07.AC(XX) است و در آن XX شماره گروه به صورت Hexadecimal است. اگر سوئیچ جدیدی وارد HSRP شد و تمایل به active شدن داشت، باید سوئیچ active جای خود را به این سوئیچ بدهد پس سوئیچ جدید باید priority بالاتری داشته باشد که به آن Preemption گویند. به صورت پیش فرض غیر فعال و قابل تنظیم است. ممکن است سوئیچی به اشتباه active شود که پس از آن سوئیچ متوجه می‌شود که اشتباه کرده است و سوئیچ active با استفاده از preempt مقدار priority خود را افزایش می‌دهد. در صورت فعال بودن preempt سوئیچ سعی می‌کند active شود ولی در غیر این صورت زمانی active می‌شود که سوئیچ active غیر فعال و در مشاخره بین سوئیچ‌ها پیروز شود. زمانی که سوئیچی برای اولین بار اجرا می‌شود و routing table آن خالی و یا ناقص است پیشنهاد می‌شود preempt غیر فعال باشد. HSRP و VRRP نمی‌توانند بین سوئیچ‌ها load sharing انجام دهند به معنای این که نمی‌توانند مشخص کنند default gateway بعضی از دستگاه‌ها یک سوئیچ و بعضی دیگر سوئیچ دیگر باشد. در حالتی که چندین گروه وجود داشته باشد و برای هر کاربر default gateway مشخصی اختصاص داده شود می‌توان از load sharing استفاده نمود

که اصطلاحاً به آن MHSRP (Multi Group HSRP) گویند. در این حالت می‌توان چند گروه برای یک interface معرفی نمود و با تغییر priority عمل load sharing را با استفاده از DHCP انجام داد.

```
(Config-if)#standby <group-id> ip <virtual-ip>
(Config-if)#standby <group-id> priority <priority-number>
```

می‌توان یک شماره گروه یکسان برای interface های متفاوت انتخاب نمود، به دلیل اینکه در broadcast domain های متفاوتی هستند. ولی نمی‌توان از virtual ip مشابهی استفاده نمود.

### HSRP State Machine -۱-۱-۲۱

زمانی که بر روی مسیریاب دستور show standby وارد شود، حالت‌های مختلف سوئیچ تا کامل شدن HSRP قابل مشاهده است. حالت‌های مختلف یا HSRP State Machine به شرح زیر است:

۱. Disable: هیچ عملی انجام نمی‌شود. ممکن اسن interface خاموش باشد.
۲. Initial (INIT): HSRP هنوز اجرا نشده است ولی interface در حال UP شدن است.
۳. Learn: ممکن است Virtual IP تنظیم نشده باشد و در حال یادگیری از بسته‌های hello باشد.
۴. Listen: Virtual IP را یاد گرفته ولی هیچ کدام از سوئیچ‌ها active و standby نیستند.
۵. Speak: سوئیچ‌ها در حال بررسی و صحبت برای active و standby شدن هستند. زمانی که یک سوئیچ جدید وارد شود و preempt بر روی آن فعال باشد با active صحبت می‌کند تا به این سوئیچ تغییر کند.
۶. Standby/Active: سوئیچ با بیشترین priority به عنوان active و بعدی به عنوان standby است. فقط سوئیچ‌های active و standby به یکدیگر بسته‌های hello با source MAC متفاوت ارسال می‌کنند. سوئیچ active از virtual mac address و سوئیچ standby از physical mac address استفاده می‌کند.

به صورت پیش فرض زمان hello سه ثانیه و زمان dead ده ثانیه حدود  $3 \times \text{hello}$  است. برای تغییر زمان یا HSRP Hello Time از دستور زیر استفاده می‌شود:

```
(Config-if)#standby <group-id> timers <hello-time> <hold-time>
(Config-if)#standby <group-id> timers msec <hello-time> msec <hold-time>
```

### HSRP Modify Preemption -۲-۱-۲۱

اگر سوئیچ جدیدی تازه روشن شود و بر روی آن پروتکل مسیریابی و HSRP فعال شود، HSRP زودتر از پروتکل مسیریابی کامل می‌شود. مشکلی که در این بین خواهد بود هنگام درخواست کاربر به دلیل کامل نشدن جدول مسیریابی پاسخ کاربر بی نتیجه می‌ماند. پس باید دقت داشت با active شدن سوئیچ مشکلی نخواهد بود فقط با فاصله

زمانی معینی این اتفاق رخ دهد تا جدول مسیریابی و سرویس‌های دیگر کامل اجرا شوند. برای تنظیم HSRP Modify Preemption از دستور زیر استفاده می‌شود:

```
(Config-if)#standby <group-id> preempt delay [minimum | reload | sync]
```

مؤلفه Minimum به معنای آن است که بعد از حداقل مدت مشخصی active شود.

مؤلفه Reload به معنای آن است که بعد از حداقل مدتی بعد از reload شدن سوئیچ active شود.

سرویس‌های دیگری از HSRP استفاده می‌کنند. بدین معنا که بعد از اجرای HSRP سرویس‌های دیگر اجرا شوند. این زمان یعنی بعد از مدت زمان مشخصی پس از sync بین IP Redundancy Client و HSRP سوئیچ به حالت active برود.

### ۳-۱-۲۱ HSRP Authentication

احراز اصالت HSRP به دو صورت plaintext و MD5 صورت می‌گیرد. در HSRP نیز می‌توان از key chain استفاده نمود که دستورات آن در زیر آمده است:

```
(Config-if)#standby <group-id> authentication <password>
(Config-if) #standby <group-id> authentication md5 key-string [0|7] <password>
(Config)#key chain <name-of-chain>
(Config-if) #standby <group-id> authentication md5 key-chain <name-of-chain>
```

### ۴-۱-۲۱ HSRP Object Tracking

اگر در سوئیچ active مشکلی پیش آید به طور مثال uplink آن قطع شود و یا هر مشکل دیگری برای interface رخ دهد، HSRP آن را با استفاده از object tracking تشخیص داده و سوئیچ standby به active تبدیل می‌شود. در واقع object همان مؤلفه‌ای است که در حال مشاهده است و track پیگیری کردن آن است. HSRP هنگام تشخیص آن priority سوئیچ active را به صورت پیش فرض 10 واحد کاهش می‌دهد.

```
(Config)#track <number> <object-to-track>
(Config-if)#standby <group-id> track <object-number> [decrement]
```

Decrement مقدار واحد کاهش priority است که پیش فرض 10 واحد است.

### ۲-۲۱ Virtual Router Redundancy Protocol (VRRP)

این پروتکل شبیه HSRP است ولی تفاوت‌ها و ویژگی‌های مخصوص به خود را دارد. VRRP یک پروتکل استاندارد اینترنتی است که با نام RFC 3768 و در ورژن جدید با نام RFC 5798 معرفی شده است. در صورتیکه HSRP یک پروتکل اختصاصی سیسکو است. پروتکل VRRP از پروتکل‌های UDP و TCP استفاده نمی‌کند بلکه از

پروتکل شماره 112 استفاده می کند. پروتکل VRRP از آدرس چند پخشی 224.0.0.18 استفاده می کند. در این پروتکل نقشی به نام Active وجود ندارد و این نقش تبدیل به Master شده است. این نقش به عنوان پاسخ دهنده به arp request و درخواست کاربران برای virtual IP شناخته شده است. در مسیریاب ها به صورت پیش فرض preemption فعال است و با همان priority پیش فرض 100 می باشند. انتخاب Master نیز مانند HSRP است. در این پروتکل زمان ها  $hello=1$  و  $dead=3.6$  ثانیه است. آدرس  $Virtual\ MAC=000.5e00.01xx$  و  $xx$  از group-id بدست می آید. مانند HSRP پروتکل VRRP نیز Load Balancing ندارد ولی با ایجاد چند گروه می توان این عمل را انجام داد. همانطور که گفته شد در پروتکل HSRP سه نقش وجود دارد. Active که پاسخ گو به درخواست های arp و کاربران است و بسته hello ارسال می کند. Standby که سرویسی به کاربران نمی دهد و بسته hello ارسال می کند. در آخر Listening که پس از خطا در active، انتخاب standby صورت می گیرد. در VRRP دو نقش بیشتر وجود ندارد. Master که بسته Hello ارسال می کند و به کاربران پاسخ می دهد. Listening که مسیریاب ها فقط گوش می دهند و هنگام خطا در active، انتخاب برای active بودن صورت می گیرد. به نظر می آید در HSRP هنگام خطای active همگرایی و active شدن به خاطر وجود standby زودتر به سرانجام برسد. در واقع در VRRP زودتر این عمل اتفاق افتاده و زمان کمتری قطعی وجود دارد. در VRRP هنگام خطای Master همه مسیریاب ها از یک آدرس MAC مبدأ استفاده می کنند و اگر همه به یک سوئیچ متصل باشند این سوئیچ دچار flip-flap می شود و پیغام MAC-Flap-Notification می دهد. دستورات تنظیم آن به صورت زیر است:

```
(Config-if)#vrrp <group-id> ip <virtual-ip>
(Config-if)#vrrp <group-id> priority <priority-number>
```

احراز هویت نیز همانند HSRP به دو صورت plaintext و md5 است.

بهینه کردن زمان در VRRP باید با دقت صورت گیرد. در یک گروه VRRP باید تمام مسیریاب ها زمان های مربوط به VRRP یکسانی داشته باشند. بیشتر پروتکل ها زمان Hello را به صورت خودکار از بسته ها یاد می گیرند ولی در VRRP باید به صورت دستی این عمل صورت گیرد. در VRRP دو حالت وجود دارد یا زمان Hello افزایش و یا کاهش داده می شود. در این پروتکل زمان بر حسب ثانیه و یک بایت در نظر گرفته شده است که 00000001 به معنای یک ثانیه است. حال اگر زمان بر حسب میلی ثانیه وارد شود نمی توان آن را اعلام نمود پس باید در همه مسیریاب ها این زمان تنظیم شود. ممکن است هدف از کاهش زمان سریع شدن همگرایی باشد. حال اگر هدف افزایش زمان باشد می توان در مسیریاب Master تنظیم نمود و در بقیه مسیریاب ها Learn را وارد نمود تا از Master یاد گیرند. نکته حائز اهمیت اینجاست که اگر زمان در مسیریاب غیر Master افزایش زمان تنظیم گردد این مسیریاب

فکر می کند Master است و خود را Master می کند. به عبارتی دو Master وجود دارد. شاید در نگاه اول خوب به نظر برسد که دو Master وجود دارد و می توان load balancing انجام داد ولی مشکلاتی پیش می آید که در زیر به آن توضیح داده می شود:

۱. همه Master از یک virtual mac address استفاده می کنند و زمانی که کاربر درخواست arp می دهد چندین پاسخ متفاوت می گیرد که بستگی به سیستم عامل کاربر دارد.

۲. همه Masterها از یک virtual mac address استفاده می کنند و هنگام اتصال به یک سوئیچ مشکل flip-flap به وجود می آید و سوئیچ این MAC را پاک و دوباره می نویسد و هر لحظه پاسخ متفاوتی به مسیریاب های مختلف داده می شود. مشکلی که ایجاد می گردد چند مسیر برای هر session ایجاد می کنند. در کل کاهش زمان باید در همه مسیریاب ها تنظیم گردد و با افزایش زمان بر روی Master در بقیه مسیریاب ها باید دستور Learn وارد شود. باید توجه داشت که VRRP نمی تواند میلی ثانیه را اعلام کند و این باید بر روی همه مسیریاب ها تنظیم شود.

```
(Config-if)#vrrp <group-id> timers advertise msec <value-time>
(Config-if)#vrrp <group-id> timers learn
(SW1)#show vrrp
(SW1)#show vrrp brief
```

## ۲۱-۳- Gateway Load Balancing Protocol (GLBP)

اگر هدف ایجاد load balancing برای تعداد بسیاری از کاربران که در یک گروه شبکه هستند و از مسیریاب ها مختلفی استفاده می کنند باشد، می توان از HSRP و یا VRRP استفاده نمود ولی باید با DHCP Server هماهنگی باشد. شرکت سیسکو برای اینکه مدیر شبکه با سرویس دهنده درگیر نشود پروتکلی ارائه نمود. در پروتکل های قبلی باید در DHCP Server تعریف شود که هر کاربر از کدام GW استفاده نماید. در پروتکل GLBP که شرکت سیسکو ارائه داده است همه کارها به صورت خودکار انجام می گردد. پروتکل GLBP مخصوص شرکت سیسکو است و دو عمل load balancing و افزایش اطمینان (Redundancy) مسیریاب ها را بالا برده است. اصطلاحاتی که در این پروتکل به کار برده می شود در ادامه توضیح داده می شود.

## ۲۱-۳-۱ Active Virtual Gateway (AVG)

این نقش شبیه active است که به arp و کاربران پاسخ می دهد و فقط یک مسیریاب AVG وجود دارد. ولی می توان هم زمان چندین مسیریاب active با virtual mac های متفاوتی داشت. در این پروتکل active به معنای مسیریابی است که بتواند پاسخ کاربران را بدهد و standby جایگزین AVG پس از خطای آن است. AVG می تواند بسته هایی که از سوی کاربران در یک گروه GLBP می آید را به مسیریاب های دیگر واگذار نماید تا آنها به کاربران خدمات

دهند که به Active Virtual Forwarder (AVF) معروف است. در واقع هنگام درخواست arp به AVG آدرس خود یا یکی از این AVF ها به کاربر داده می شود. Preemption به صورت پیش فرض برای نقش AVG غیر فعال است. تنها AVG می تواند پاسخ arp کاربران را بدهد. همانند پروتکل های قبلی انتخاب AVG مسیریاب با بیشترین اولویت یا آدرس IP است. که اولویت به صورت پیش فرض 100 است. حداکثر تعداد AVF ها در یک گروه 4 است که خود AVG نیز یک AVF محسوب می شود. Virtual MAC Address=0007.b4xx.xxyy که در آن xx.xx شماره گروه بر حسب هگزادسیمال است و yy شماره AVF است. AVF برای پیدا کردن شماره خود و virtual mac به AVG درخواست می دهد. تنها مسیریاب های AVG و AVF بسته Hello ارسال می کنند. این بسته برای آن است که AVG از وضعیت AVF ها با خبر باشد و در صورت خطای یکی از آنها Virtual MAC مربوط به AVF را به یک AVF دیگر داده تا پاسخ گوی کاربران AVF دچار خطا شده باشد. در واقع یک AVF دو آدرس MAC بر عهده دارد که انتخاب آن بستگی به اولویت و آدرس IP بزرگتر دارد. پس از برگشت AVF خراب این آدرس برگشت داده شده و دوباره به حالت عادی خود بر می گردد. مسیریابی که وظیفه AVF خراب را دارد به صورت پیش فرض تنها ۴ ساعت پاسخ گوی کاربران است و پس از آن کاربران قطع می شوند. این زمان در AVG ذخیره شده و پس از آن اگر کاربر جدیدی وارد شده آن را به یک AVF دیگر ارجاع می دهد. زمان Hello=3 ثانیه است. آدرس چند پخش 224.0.0.102 است و از پروتکل UDP با پورت 3222 استفاده می کند. دستور تنظیم آن به صورت زیر است:

```
(Config-if)#glbp <group-id> ip <virtual-ip>
(Config-if)#glbp <group-id> priority <priority-number>
```

الگوریتم هایی که GLBP از آن برای load Balancing هنگام ارجاع یک کاربر به AVF استفاده می کند یکی از موارد زیر است:

۱. Round Robin: به صورت پیش فرض از این الگوریتم استفاده می شود. به این صورت است که اولین کاربر از AVF اول و دومین کاربر از AVF دوم سرویس می گیرد و همین صورت ادامه می یابد و به صورت چرخشی است.

۲. Host Dependent: زمانی که یک کاربر درخواست می دهد یک AVF به آن اختصاص می یابد و بار دیگر اگر همان کاربر درخواست مجدد داد (به دلیل خاموش و روشن شدن سیستم کاربر، رفتن و بازگشت دوباره به شبکه و یا اتمام زمان cache) باز به همان AVF قبلی ارجاع داده می شود. در سیستم NAT بیشتر کاربرد دارد چون به یک MAC یکسان نیاز می شود.

۳. Weighted: به هر AVF با تناسب به میزان RAM و CPU و در کل معیارهای دیگر یک بر چسب زده می‌شود که به آن weight گفته می‌شود و هر چه شماره این برچسب بیشتر باشد به معنای ارجای کاربران بیشتر به این AVF است. دستورات تنظیم آن به صورت زیر است:

```
(Config-if)#glbp <group-id> load-balancing <weighted | round-robin | host-dependent>
```

دستور فوق در AVG وارد می‌شود.

```
(Config-if)#glbp <group-id> weighting <value> lower <value> upper <value>
```

دستور فوق در AVF وارد می‌شود.

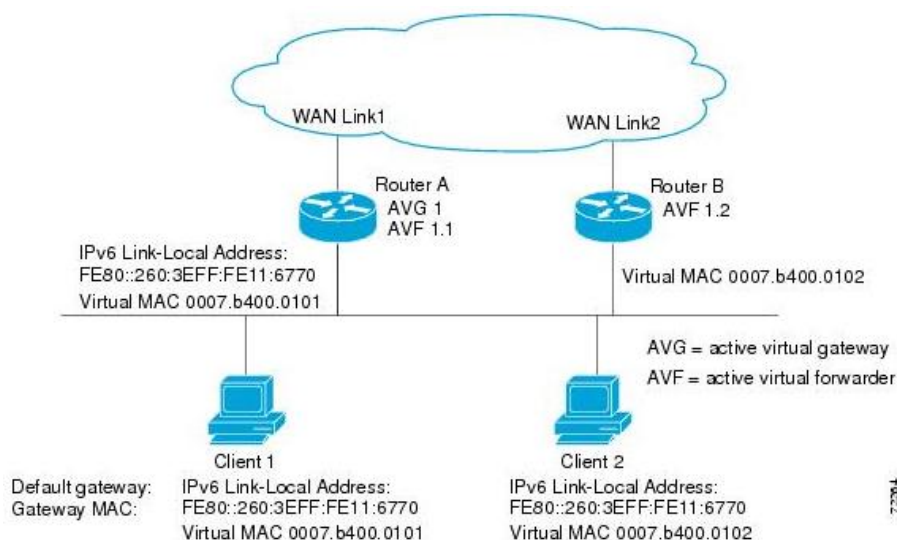
فرض کنید یک track تعریف شده است که باعث کاهش weighting می‌شود حال اگر این مقدار از lower بر اثر object tracking کمتر شد به معنای این است که AVF مناسبی نیست و باید از AVF ها حذف شود. اگر مقدار weight برابر یا بیشتر از upper شد به معنای بازگشت مسیریاب به AVF ها است. به صورت پیش فرض AVF Weighting=100 است.

```
(Config-if)#glbp <group-id> weighting track <track-number> decrement <weighting-decrement-unit>
```

زمان ها در GLBP: ۱- hello - ۲ hold time - ۳ redirect: هنگام fail شدن یک AVF بعد از چه مدت زمانی اگر کاربر درخواست داد آن را به AVF جدید ارجاع دهد. ۴- Time-out Interval: چه مدت زمانی AVF جدید به کاربران AVF خراب پاسخ دهد که به صورت پیش فرض ۴ ساعت است.

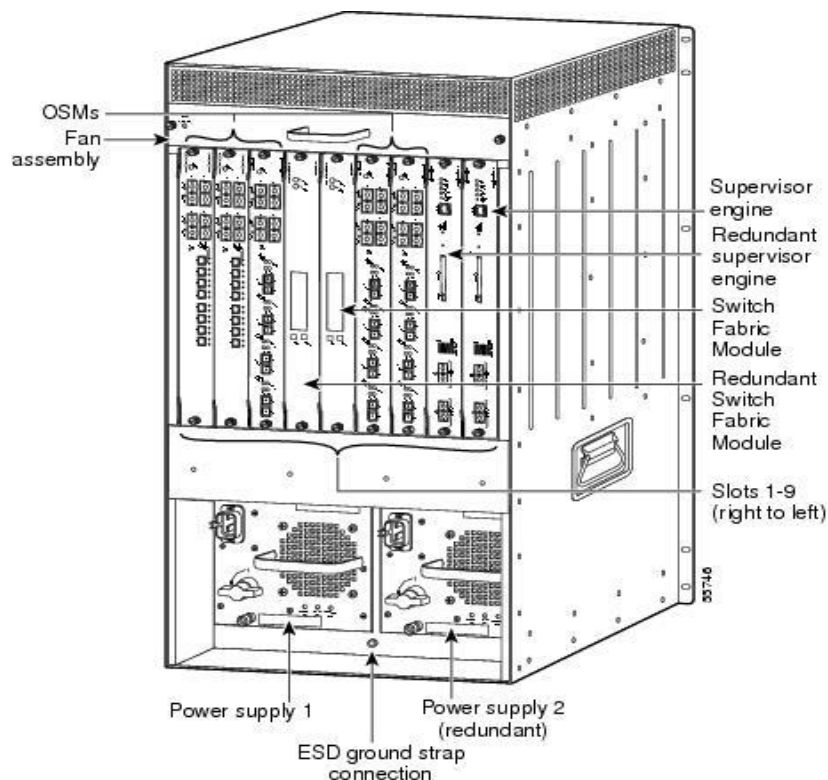
```
(Config-if)$glbp <group-id> timers <hello-time> <hold-time>
```

```
(Config-if)$glbp <group-id> timers redirect <value-redirect> <value-time-out>
```



## ۲۲- Supervisor & Route Processor Redundancy

بین مسیریاب‌ها و سوئیچ‌های تخت (Fixed) و غیر تخت (Modular) تفاوت‌هایی وجود دارد. در سوئیچ‌های Modular می‌توان اینترفیس‌های بیشتر با کاربردهای متفاوت اضافه نمود. به طور مثال در Nexus 7000 می‌توان اینترفیس‌ها با سرعت پایین را حذف و با سرعت بالاتری جایگزین نمود. در واقع این نوع سوئیچ‌ها یک شاسی یا بدنه دارند و بر روی آن Line Card های مورد نیاز نصب می‌شود. هر یک از این Line card ها یک Circuit Board دارد. درون این Circuit Board ها یک CAM Table و یک Forwarding Engine وجود دارد. این Line Card ها واحدی برای پردازش و دستور نیاز دارند. این واحد به عنوان مغز پردازشی سوئیچ شناخته می‌شود و به آن Supervisor گفته می‌شود. همچنین واحدی به نام Distributed Forwarding Card (DFC) وجود دارد که همه چیز را از Supervisor یاد می‌گیرد. در واقع Supervisor برد DFC را برنامه ریزی می‌کند. در بعضی از سوئیچ‌ها این Module به صورت On-Board روی Supervisor قرار دارد. واحد DFC وظیفه Forwarding، پیگیری عملیات بر روی هر Frame، پروتکل‌های مسیریابی، STP و همه وظیفه‌هایی که به عهده Supervisor است را بر عهده دارد. برای دسترس پذیری بالا بهتر است دو Supervisor استفاده شود. مکان Supervisor گاهی در وسط سوئیچ، گاهی در بالا و گاهی سمت راست قرار می‌گیرد که اهمیت زیادی ندارد. در حقیقت Supervisor برای سوئیچ‌های Modular به عنوان مغز متفکر و هوشمند است که واحدهای دیگر را برنامه ریزی می‌کند.



## ۲۲-۱- Route Processor Redundancy (RPR)

قبل از این که به توضیح این پروتکل پرداخته شود نیاز است پلی به گذشته زده شود. چندین سال قبل مدل catalyst 5000 وجود داشت. این سوئیچ از دو Supervisor پشتیبانی نمی کرد. بعد از آن مدل 6000 با کمی تفاوت نسبت به 5000 به بازار عرضه شد. در مدل 5000 سیستم عامل Catalyst OS نصب بود ولی در مدل 6000 از سیستم عامل IOS Cisco استفاده می شد. در واقع سیستم عامل مدل 6000 یا همان IOS برای مسیریابی و سیستم عامل Catalyst OS برای سوئیچینگ طراحی شده بود. اگر Supervisor قصد استفاده از تنظیمات سوئیچینگ داشت باید از دستورات catalyst استفاده می کرد و اگر مربوط به مسیریابی بود باید از دستورات IOS استفاده می نمود. همین امر باعث می شد که اگر Supervisor بر روی یک تیغه، هم سوئیچینگ و هم مسیریابی انجام دهد، باید به دو حالت مختلف جابه جا شود. به این حالت Hybrid گفته می شود. بعد از آن شرکت سیسکو با توجه به رضایت مشتری و دلایل دیگر تصمیم به تغییر سیستم عامل catalyst به IOS گرفت. در واقع به صورت بومی IOS برای مسیریابی استفاده می شد ولی در حال حاضر هیچ نگرانی نسبت به دستورات چندین سیستم عامل وجود ندارد و با یک دستور (switchport) می توان بین دو حالت سوئیچینگ و مسیریابی جابه جا شد. اگر دو Supervisor وجود داشته باشد، قابلیتی بین این دو برای جابه جایی اطلاعات نیاز است. به این قابلیت Route Processor Redundancy (RPR) گویند که وظیفه جابه جایی اطلاعات بین دو Supervisor را بر عهده دارد. شاید این سوال پیش آید که سخن از سوئیچ گفت می شود ولی نام این قابلیت Route Processor Redundancy است که جواب همان تاریخچه گفته شده است.

نقش های Supervisor در RPR به دو حالت Active و Standby تقسیم می شود. همه وظایف به عهده Active است و Standby فقط وظیفه Backup دارد و در هنگام وقوع خطا در Active وارد جریان می شود. باید توجه داشت که تمامی Uplink ها در Standby قابل استفاده است. Active همیشه در حال جابه جایی اطلاعات با Standby است که فقط Startup-config و Config-register را ارسال می کنند. در این حالت هر دو Supervisor اکیداً پیشنهاد می شود از نسخه IOS یکسانی استفاده نمایند ولی پیش نیاز نیست. ارتباط بین دو Supervisor از طریق Line Card Switch Fabric به نام Shared Bus است. فقط سوئیچ Active کامل Load می شود و Standby در حالت (Partially Boot) یا ناقص load شدن است. هر فریمی که از هر Line Card بر روی Shared Bus قرار می گیرد که یک کپی از همه واحدها است، قابل مشاهده است. یک Switch Fabric دیگری وجود دارد که Forwarding Engine پس از تصمیم گیری با ارسال پیامی بر روی آن نتیجه اینکه چه واحدی مسئول انجام چه کاری است را اعلام می کند. همه واحدها این پیام را مشاهده، کپی ها را حذف می کنند و پورت مسئول وظیفه خود را برای انتقال فریم انجام می دهد. در Standby واحد Forwarding Engine وجود

دارد ولی خاموش است و Uplink ها فعال توسط Active تصمیم گیری می شود. اگر Active دچار خطا شود حداقل دو دقیقه زمان برای تغییر Standby به Active نیاز است، زیرا تمامی line card ها باید دوباره load شوند، واحد مسیریابی و پروتکل های لایه دو آغاز به کار شوند.

```
(Config)#redundancy
(Config-red)#mode rpr
(SW1)#show redundancy state
```

## RPR+ -۲-۲۲

این قابلیت سریع تر از RPR است. برعکس RPR که اکیداً پیشنهاد می شود نسخه های Supervisor یکسان باشد، در RPR+ باید نسخه های IOS Supervisor یکسان باشند. Standby کاملاً load و تنظیم می شود. CPU بر خلاف RPR اجرایی است و Running-config نیز جابه جا می شود. در RPR حتماً تنظیمات باید ذخیره می شد تا بین Supervisor ها جابه جا شوند. بین ۳۰ تا ۶۰ ثانیه زمان برای Active شدن لازم است. ماژول های نصب شده نیازی به دوباره load شدن ندارند. در این حالت FIB Table به طور کامل پاک می شود و ترافیک های مربوط به مسیریابی drop می شوند اما مسیرهای Static باقی می ماند. همچنین Gig Uplink ها بر روی Standby همیشه فعال هستند.

برای تنظیمات می توان از دستورات زیر استفاده نمود:

```
(Config)#redundancy
(Config-red)#mode rpr-plus
```

## SSO (State SwitchOver) -۳-۲۲

این قابلیت شبیه به RPR+ است. تفاوت هایی بین RPR+ و SSO به شرح زیر است:

۱. تمامی اطلاعات مربوط به Stateful Feature مانند: FIB و Adjacency Table، STP و وضعیت پورت ها، نگهداری می شوند.
۲. به دلیل ارسال اطلاعات stateful سریع تر از RPR+ است و بین ۰ تا ۳ ثانیه زمان برای Active شدن لازم است.
۳. کمتر سوئیچی از این پروتکل پشتیبانی می کند.

اطلاعات مربوط به پروتکل های مسیریابی مانند همسایه ها و ... جابه جا نمی شوند. این اطلاعات بر روی قسمت route processor ماژول Multilayer Switch Feature Card (MSFC) قرار دارد که این ماژول یک Layer 3 Switching Engine است که قسمت جدا ناشدنی supervisor و به عنوان daughter card شناخته می شود.

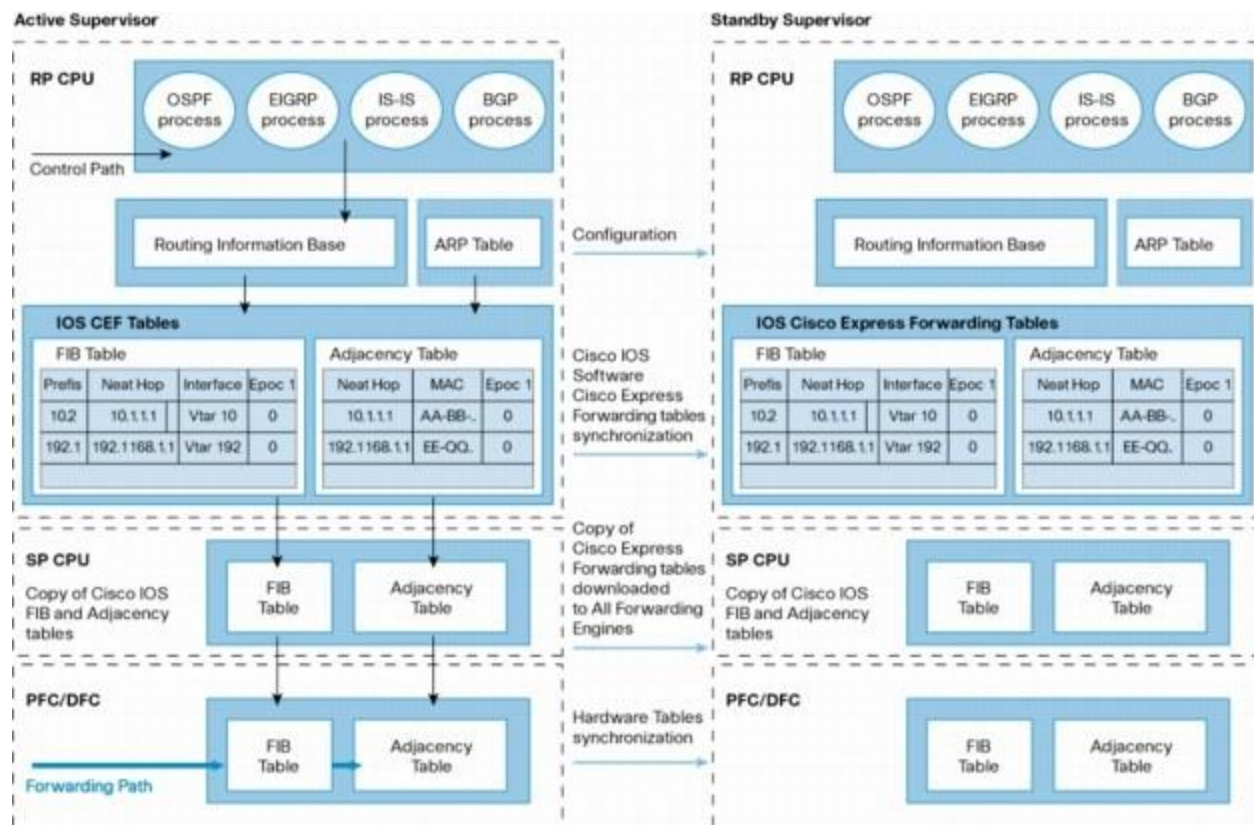
در واقع این ماژول با یک CPU وظیفه مسیریابی و با یک CPU دیگر وظیفه لایه دو را انجام می‌دهد. در واقع می‌توان گفت سیستم‌های شبکه دو قسمت دارند: یک قسمت Control Plane که وظیفه ارتباطات کنترلی و مدیریتی مانند پروتکل‌های مسیریابی و ... دارند. همچنین یک قسمت Data Plane که نتیجه Control Plane است و ارسال بسته‌ها به مقصد است. در این ماژول عملیات Control Plane انجام می‌شود مانند بررسی کردن پروتکل‌های مسیریابی لایه ۳ و نگهداری جدول مسیریابی و... در ادامه بیشتر توضیح داده می‌شود. سوئیچ Standby فقط تنظیمات مربوط به لایه ۳ را ذخیره می‌کند ولی ارتباط با همسایه‌ها و اطلاعات مربوط به آن را ندارد. این ارتباطات در مسیریاب‌ها در طرف دیگر لینک‌ها است و اگر Active دچار خطا شود بسته‌هایی که درخواست مسیریابی دارند حذف می‌شوند.

## ۲۲-۴- SSO with Non-Stop Forwarding (NSF)

تا این مرحله SSO بهترین پروتکل است ولی باز هم ۳ ثانیه زمان برای جابه‌جایی صرف می‌شود. شرکت سیسکو تصمیم گرفت قابلیت ارائه دهد که به هیچ وجه از ارسال بسته خودداری نشود (Non-Stop Forwarding). قبل از آن بهتر است مروری بر عملکرد مسیریابی و سوئیچینگ شود. ماژول MSFC در واقع تمام اطلاعات مسیریابی را در خود ذخیره می‌کند و پس از آن به Cisco Express Forwarding Information Base (FIB) که در قسمت route processor است و بعد از آن به قسمت switch processor انتقال می‌دهد. بعد از آن به قسمت Hardware Table یا Hardware Application Specific Integrated Circuits (ASIC) که بر روی Policy Feature Card (PFC) و Distributed Forwarding Card (DFC) قرار دارد، می‌دهد و در حافظه TCAM ذخیره می‌شود. در واقع forwarding engine از این جدول استفاده می‌کند. این حافظه TCAM در Standby کپی می‌شود و زمانی که Active دچار خطا شود جدول FIB و TCAM باقی می‌ماند. زمانی که درخواست مسیریابی داده شود از این جدول‌ها استفاده می‌شود و در همین هنگام ارتباطات همسایگی خود را نیز با دیگران برقرار می‌کند. در واقع قسمت Control Plane در حال برقراری ارتباط با همسایگان و پروتکل‌های مسیریابی است و قسمت Data Plane در حال استفاده از جدول‌های ذکر شده برای ارسال بسته‌ها است. NSF نیاز به همکاری با پروتکل‌های مسیریابی دارد تا بین Supervisor و مسیریاب که در طرف دیگر لینک است هنگام خطا ارتباط برقرار بماند و همگرایی سریع‌تر رخ دهد. مسیریاب این قضیه را می‌فهمد و هنگام عدم دریافت Hello ارتباط را نگه می‌دارد تا Supervisor دیگر ارتباط را دوباره بازسازی کند. دو اصطلاح در این میان استفاده می‌شود. NSF-Capable که همان سوئیچ می‌شود. NSF-Aware همان مسیریاب می‌شود که در صورت خطا در Active ارتباط را نگه می‌دارد.

BGP:(Config-router)#bgp graceful restart

OSPF:(Config-router)#nsf  
EIGRP:(Config-router)#nsf



## ۲۳- POE (Power over Ethernet)

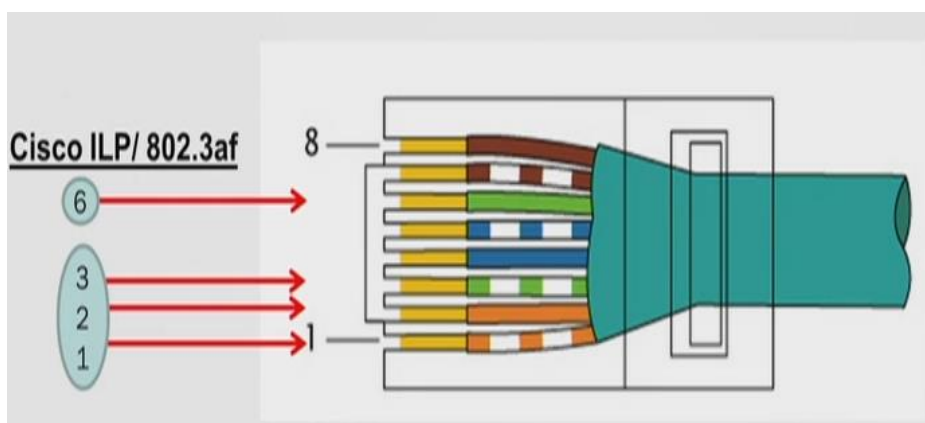
بعضی از تجهیزات شبکه نیاز به برق تقریباً کمی دارند و برق رسانی به آنها سخت یا نشدنی است. یک آنتن بر روی دکل را تصور کنید اگر بتوان فقط با یک کابل هم اطلاعات و هم برق رسانی را انجام داد کار ساده تر می شود. کابل شبکه به سوئیچ متصل شده و برق را از سوئیچ گرفته و به دستگاه مورد نظر انتقال می دهد. IP Phone را در نظر گرفته، این دستگاه نیاز به برق با ولتاژ بالا ندارد و بر روی میز قرار می گیرد. اگر این دستگاه از طریق همان کابل شبکه برق رسانی انجام دهد، شلوغی روی میز و بعضی از مشکلات دیگر کاهش می یابد. چند اصطلاح در زمینه ارسال برق بر روی کابل شبکه وجود دارد:

۱. PD (Powered Device): دستگاهی که نیازمند برق است.
۲. PSE (Power Source Equipment): دستگاهی که منبع برق است. دو نوع PSE وجود دارد:
  - I. End Span: دستگاه شبکه که برق را انتقال می دهد و در آخرین نقطه قرار دارد مانند سوئیچ.
  - II. MidSpan: دستگاهی بین سوئیچ و PD است مانند Patch Panel که بر روی آن POE قرار دارد.

روش‌ها و استانداردهای POE شامل موارد زیر است:

### ۱-۲۳ Cisco Inline Power (ILP)

این استاندارد مخصوص شرکت سیسکو است و توسعه یافته آن تبدیل به IEEE 802.3af شده است. در ادامه به شرح توضیح داده می‌شود. حداکثر توان قابل ارائه برای هر پورت 6.3 W است. توان مورد نیاز به این صورت محاسبه می‌شود که ابتدا بر روی پورت موجی با 340KHz بر روی TX ارسال می‌شود که هیچ آسیبی به دستگاه مقصد نمی‌رساند. پس از دریافت توسط PD این دستگاه آن را بر روی RX قرار می‌دهد و به سمت سوئیچ بر می‌گرداند که سوئیچ مقدار ولتاژ مورد نیاز این دستگاه را محاسبه می‌کند. همچنین سوئیچ انتقال دهنده باید منبع برقی با قدرت 714 W داشته باشد.



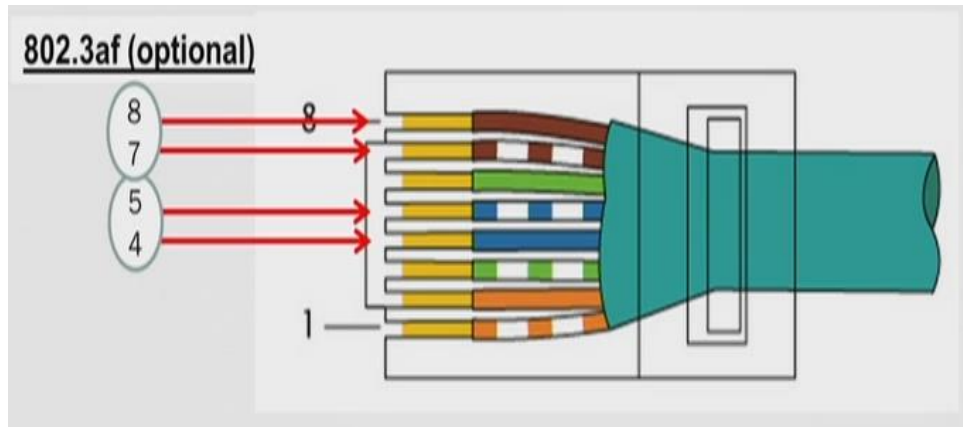
### ۲-۲۳ IEEE 802.3af

حداکثر توان برای هر پورت 15.4 W است و سوئیچ انتقال دهنده باید منبع برقی با قدرت 714 W داشته باشد. هنگامی که پورت فعال می‌شود ولتاژ کمی بر روی TX/RX ارسال می‌شود و مقاومت اندازه‌گیری می‌شود. اگر مقاومت برابر 25 K $\Omega$  باشد این دستگاه یک PD است و این روال تا چند بار صورت می‌گیرد و مقاومت اندازه‌گیری می‌شود. با توجه به مقاومت‌های بدست آمده مقدار ولتاژ مورد نیاز با توجه به کلاس بندی استاندارد بدست می‌آید. در این استاندارد بدلیل ولتاژ اولیه‌ای که ارسال می‌شود ممکن است به دستگاه‌های Non-PD آسیب برسد.

|                                           |                               |
|-------------------------------------------|-------------------------------|
| 0 (default by IEEE when no class is used) | 15.4 W                        |
| 1                                         | 4.0 W                         |
| 2                                         | 7.0 W                         |
| 3                                         | 15.4 W                        |
| 4                                         | Up to 50 W (Optional 802.3at) |

### IEEE 802.3at-۳-۲۳

در این استاندارد حداکثر توان قابل ارائه 25.5 W است. کمتر سوئیچی از این استاندارد پشتیبانی می کند. منبع برق سوئیچ باید توان 1418 W داشته باشد.



### Cisco Universal Power over Ethernet (UPOE)-۴-۲۳

این استاندارد تا 60 W برای هر پورت را فراهم می کند و سوئیچ هایی از قبیل 4500e و 3815 از آن پشتیبانی می کنند. یکی از تفاوت های بین استانداردها در روش تشخیص برق مورد نیاز PD است. CDP نیز می تواند در تشخیص برق مورد نیاز موثر باشد. به طور مثال هنگامی که از IP Phone سیسکو استفاده می شود این دستگاه مقدار برق مورد نیاز خود را به سوئیچ اطلاع می دهد.

```
(Config-if)#power inline {auto [max <milli-wats>] | static [max <milli-wats>] | never [max <milli-wats>]}
```

به صورت پیش فرض در حالت auto است.

دلایل قطعی POE ممکن است یکی از موارد زیر باشد:

۱. پورت توسط مدیر خاموش شده باشد.
۲. پورت در حالت err-disable باشد.
۳. "Power inline never" تنظیم شده باشد.
۴. پورت patch panel معیوب باشد.
۵. کابل شبکه معیوب باشد.
۶. طول کابل زیاد باشد (باید کمتر از ۱۰۰ متر باشد).

۷. سوئیچ توانی برای ارائه بر روی پورت نداشته باشد (پورت‌ها بیشتر از حد مجاز باشند).

۸. POE مربوط به پورت در سوئیچ خراب باشد.

۹. POE در قسمت پاور سوئیچ خراب باشد.

۱۰. خطای سخت افزاری وجود داشته باشد.

## ۲۴- Voice Vlan

برای جدا سازی صدا از اطلاعات دیگر از این ویژگی استفاده می‌شود. به طور مثال در اکثر تلفن‌های IP Phone یک سوکت ورودی و یک سوکت خروجی وجود دارد که سوکت ورودی متصل به سوئیچ و سوکت خروجی متصل به کامپیوتر دیگری می‌شود. اطلاعات مربوط به کامپیوتر و IP Phone هر دو به یک پورت سوئیچ متصل هستند. برای این که بتوان سیاست‌های متفاوتی بر روی بسته‌ها اعمال نمود باید این دو نوع ترافیک از یکدیگر جدا شوند. راه کارهای جداسازی بسته‌ها از یکدیگر شامل موارد زیر است:

۱. تگ 802.1q برای جداسازی استفاده شود. به عبارتی پورت سوئیچ دو VLAN داشته باشد یکی برای اطلاعات و دیگری برای صدا که vlan صدا با 802.1q tag جدا شده و اطلاعات به صورت untagged یا native vlan باشد. IP Phone در این جا نقش سوئیچ میانی بین کامپیوتر و سوئیچ اصلی را بازی می‌کند و 802.1q بین IP Phone و سوئیچ استفاده می‌شود.

۲. می‌توان با تنظیم Quality of Service (QoS) اولویت صدا را نسبت به اطلاعات افزایش داد.

سوئیچ در صورتیکه IP Phone Cisco استفاده شده باشد از طریق CDP مطلع شود که به IP Phone متصل است. اگر دستگاهی غیر سیسکو باشد از طریق option-156 پروتکل DHCP دستگاه مطلع می‌شود.

`(Config-if)#switchport voice vlan <vlan-id | dot1p | untagged | none>`

سوئیچ با استفاده از روش‌های زیر بسته را اولویت‌بندی می‌نماید:

۱. VLAN-ID: سوئیچ بسته CDP را با استفاده از option ارسال می‌کند و IP Phone از این قابلیت استفاده و ترافیک را جدا می‌کند.

۲. 802.1p: این هدر بسته در دستگاه‌هایی مانند Host یا End Device استفاده می‌شود. در VLAN در قسمت option باید عددی بین ۱ تا ۴۰۹۴ استفاده شود که به عنوان VLAN-ID شناخته می‌شود و ۱۲ بیتی است. در هدر 802.1p این مقدار ۱۲ بیتی، ۹ بیت آن ۰ و ۳ بیت آخر آن به عنوان بیت‌های اولویت استفاده می‌شود و IP Phone با دریافت این بسته فقط آن را forward می‌کند و تغییری بر روی آن نمی‌دهد. این اولویت در لایه ۲ صورت می‌گیرد.

۳. Untagged: در بسته CDP از VLAN-ID با شماره ۴۰۹۵ استفاده می‌شود که وجود ندارد. این مقدار رزرو شده برای صدا است و اطلاعات به صورت بدون tag است. این VLAN به Dead VLAN نیز معروف است و به معنای آن است که هیچ تفاوتی بین صدا و اطلاعات نیست.

۴. None: هیچ VLAN برای صدا اختصاص داده نمی‌شود.

می‌توان از دستور زیر برای نشان دادن وضعیت پورت سوئیچ استفاده نمود.

```
(SW1)#show interface <interface> switchport
```

## Voice QoS – ۱-۲۴

کیفیت سرویس از اهمیت خاصی برخوردار است، به این دلیل که ترافیک‌های پر اهمیت را از سایر ترافیک‌های داده دیگر جدا می‌کند. این عمل هم در سوئیچ‌ها و هم در مسیریاب‌ها انجام می‌گیرد. مزیت آن هنگامی نمایان می‌شود که در موقع وجود ازدحام و پر شدن حافظه Buffer، بسته‌ها با تأخیر ارسال و یا حذف می‌شوند. با زدن برچسب بر روی بسته‌ها اهمیت آنها مشخص می‌شود و بسته‌ها با اهمیت و اولویت بالاتر زودتر از سایر بسته‌ها ارسال می‌شوند و یا این که اصلاً حذف نگردند. در اکثر پروتکل‌ها از TCP/IP استفاده می‌گردد و با حذف تعدادی از بسته‌ها مشکلی ایجاد نمی‌گردد، زیرا TCP/IP مبتنی بر اطمینان است و بسته‌ها دوباره ارسال می‌شوند. در بسته‌های صوتی از پروتکل UDP استفاده می‌شود و تأخیر و ازدحام قابل تحمل نیست. در صورتی که سوئیچ به دلایلی همچون ازدحام، تصمیم بر حذف بسته‌ها بگیرد، بسته‌هایی را حذف می‌کند که اولویت پایین‌تری داشته باشند پس می‌توان به بسته‌های صوتی اولویت بالاتری داد. فواید کیفیت سرویس در صدا شامل پیش‌بینی ترافیک صوت و جلوگیری از Delay و Jitter است. صوت نمی‌تواند Delay و Jitter را بیش از مقدار معینی تحمل کند. ابتدا به تفاوت بین این دو پرداخته می‌شود. در Delay فضای بین ارسال بسته‌ها یا همان تأخیر برابر است و صدا قابل فهم نیست اما در Jitter فضای ارسال متغیر است و مانند قبل صدا قابل فهم نیست. به طور مثال هنگام ارسال بوق ممتد در delay فواصلی ایجاد می‌شود و بوق منقطع مساوی شنیده می‌شود ولی در jitter بوق منقطع و نامساوی شنیده می‌شود. در صوت هنگامی که یک سخن کوتاه گفته می‌شود به هزاران بسته تبدیل می‌شود که به نحوه نمونه‌گیری بستگی دارد. در QoS قسمتی از فضا برای بسته‌های صوت رزرو می‌شود و به این قسمت زودتر از سایر بسته‌ها سرویس داده می‌شود. بنابر این از Delay و Jitter جلوگیری می‌شود. کیفیت سرویس باید سرتاسر شبکه تنظیم شود، زیرا ممکن است یک سوئیچ crash کند یا اتفاقات دیگری رخ دهد. در نتیجه سایر سوئیچ‌ها باید آمادگی داشته باشند تا این مشکل را جبران کنند. ۲۰۰ میلی ثانیه حداکثر زمانی است که برای تأخیر ابتدا تا انتها بسته صوتی قابل تحمل است. توضیح داده شد که کیفیت سرویس Delay، Jitter و Lost را کاهش می‌دهد ولی باید توجه داشت که همه جا قابل اجرا نیست. به طور مثال اگر سرعت

شبکه همیشه پایین باشد و یا سوئیچی همیشه خراب باشد این عمل قابل استفاده نیست، زیرا طراحی شبکه مشکل دارد و باید راه حل هایی مانند etherchannel و یا دیگر موارد را استفاده نمود.

دو نوع Delay وجود دارد که در ادامه توضیح داده می شود.

۱. Fixed Delay تأخیری است که همیشه رخ می دهد مانند مشکل سخت افزاری و یا سرعت پایین Lookup Engine. این مشکل با QoS قابل حل نیست.

۲. Variable Delay در بعضی از موارد با QoS قابل حل نیست و در دیگر موارد QoS بر آن تاثیر می گذارد و یا آن را کاهش می دهد.

گاهی اوقات fixed delay مفید است. به طور مثال اگر یک سوئیچ ۵۰ میلی ثانیه تأخیر در forwarding engine داشته باشد و سایر سوئیچ ها ۲۰ میلی ثانیه می توان تأخیر را با استفاده از این fixed delay در سرتاسر شبکه کاهش داد. تأخیر موجود برای صوت می تواند شامل موارد زیر باشد.

۱. Coder Delay: زمان فشرده سازی و خارج کردن از فشرده سازی نمونه های PCM (الگوریتم برای تبدیل صوت به بیت کد) است.

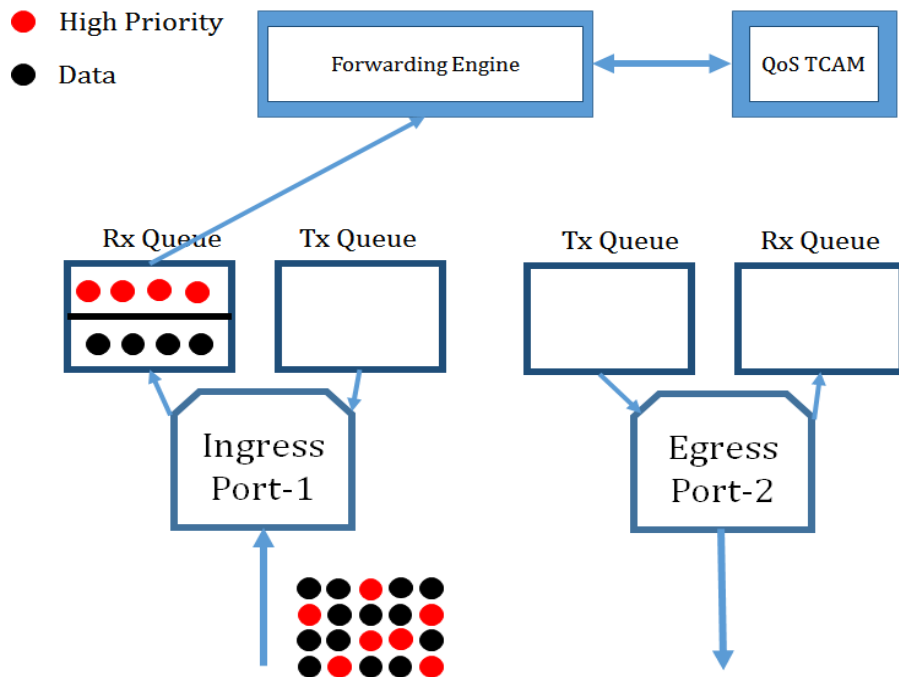
۲. Packetization Delay: زمانی که صرف پر کردن هدر بسته با encoding و فشرده سازی صوت می شود. به عبارتی ساخت بسته IP است.

۳. Serialization Delay: مدت زمانی که بسته از صف بر روی کابل قرار می گیرد. برای هر نوع کابل و پورت مانند پورت Serial و یا Ethernet متفاوت است. اگر بر روی پورت فیزیکی سریال، یک سطح ولتاژ ایجاد شود، به مدت زمانی که صرف قرار دادن بسته بر روی کابل می شود serialization گویند.

Out Put Queuing Delay: یک نوع Variable Delay است که قبل از serialization انجام می گیرد و با استفاده از QoS و کلاس بندی می توان آن را کاهش داد و به بسته اولویت بالاتری داد. دو delay دیگر وجود دارد که Network Switching Delay و De-Jitter Delay است.

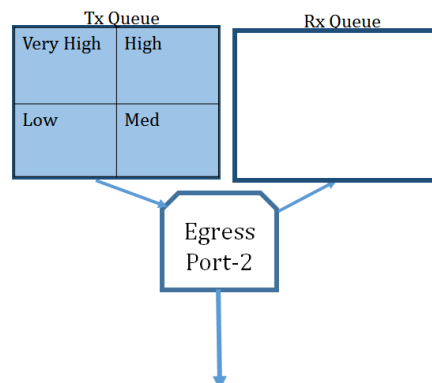
## ۲-۲۴- ابزارهای QoS

به صورت پیش فرض QoS غیر فعال است و از روش FIFO (First In First Out) برای ارسال به Forwarding engine و Transit Queue استفاده می شود. اگر QoS در قسمت Global فعال شود قسمتی از حافظه برای بسته ها رزرو می شود و بسته ها با اولویت بالاتر قبل از سایر بسته ها ارسال می شوند.



نحوه اولویت بندی بسته ها به شرح زیر است:

۱. Classification: در این روش کلاس بندی در قسمت Transit Queue است و از High به سمت Low ارسال می شود. نحوه انتخاب بسته ها و جایگذاری در کلاس ها توسط مدیر شبکه مشخص می شود.



ACL نمونه‌ای از کلاس‌بندی است و هنگامی که با آن مطابقت داشته باشد عملیاتی بر روی آن انجام می‌گیرد. در کلاس‌بندی فقط بسته‌ها داده، تشخیص و در کلاس‌های مربوط به خود جای داده می‌شوند.

۲. Queuing: با استفاده از کلاس‌بندی بسته‌ها را در صف‌های خروجی قرار می‌دهد. به عبارتی مشخص می‌کند کدام کلاس ابتدا ارسال شود. همچنین اندازه صف و نرخ ارسال کلاس به صف در Queuing قابل تغییر است.

۳. Policy/Shaping: در قسمت Policy می‌توان قوانین و سیاست‌هایی را اعمال نمود. به طور مثال ایجاد یک Threshold برای بسته‌های یک کلاس هنگام ارسال بسته‌ها به forwarding engine و هنگامی که از این حداکثر به سمت TX عبور کرد، بسته‌ها حذف شوند. Policy از پر شدن صف‌ها و عبور از threshold جلوگیری می‌کند. در shaping قسمت forwarding engine با تمام سرعت بسته‌ها را ارسال می‌کند و کنترل بسته‌ها در TX انجام می‌گیرد تا مطابق با سرعت کابل شود. تفاوت بین Policy و Shaping این است که policy در forwarding engine صورت می‌گیرد ولی shaping در TX queue و در policy اگر ترافیک ارسالی از نرخ تعریف شده تجاوز کند بسته‌ها حذف می‌شوند ولی در shaping فقط کنترل سرعت انجام می‌گیرد و بسته‌ها حذف نمی‌شوند و ممکن است صف‌ها پر شوند.

۴. Scheduling: در زمان‌بندی چگونگی ورود و یا خروج ترافیک در RX و TX مشخص می‌شود. به طور مثال ابتدا صف با اولویت بالاتر کاملاً خالی شود و سپس صف‌های دیگر ارسال شوند.

## ۲۴-۳- QoS Models

الگوریتم‌هایی که در کیفیت سرویس استفاده می‌شود شامل:

۱. Best Effort Delivery: در واقع مانند غیر فعال بودن QoS است و از FIFO استفاده می‌کند. این روش به صورت پیش فرض است. در واقع فقط بسته‌ها را ارسال می‌کند و همه چیز به شبکه بستگی دارد.
۲. Integrated Services: قبل از شروع سرویس‌ها در برنامه مبدأ، منابعی را برای کیفیت رزرو می‌کند. این روش از پروتکل RSVP (Resource Reservation Protocol) استفاده می‌کند تا تأخیر پهنای باند معینی را رزرو کند. حداقل و حداکثر پهنای باند شبکه بررسی می‌شود که آیا تأخیر برآورده می‌شود یا خیر. پس از تأیید، منابع مورد نیاز را تأمین و رزرو می‌کند. تمامی سوئیچ‌هایی که از ابتدا تا انتها وجود دارند باید پروتکل RSVP را پشتیبانی کنند. در واقع مبدأ درخواست تأمین منابع به صورت end-to-end ارسال می‌کند و منتظر apply از تمامی سوئیچ‌ها می‌ماند.

۳. Differentiated Service: این روش بر اساس سوئیچ به سوئیچ یا per-hope behavior است، به معنای این که هر بسته به هر سوئیچی که می‌رسد با توجه به تنظیمات مربوط به خود در classification عملکرد خاص خود را دارد. ممکن است یک بسته در یک سوئیچ صف اول و در سوئیچ دیگر صف دوم باشد. به معنای دیگر هر سوئیچ تنظیمات و مکانیزم مربوط به خود را در مورد مهمترین بسته‌ها دارد و با توجه به Flag هدر بسته‌ها تصمیمات QoS را اتخاذ می‌کند.

در تمامی موارد بالا Flag در دو لایه ۲ و ۳ انجام می‌گیرد که در زیر شرح داده می‌شود:

۱. Priority Tag in Layer-2: بسته‌های عادی Ethernet فیلد اولویت ندارند. هدر 802.1q/1p این هدف را پیاده‌سازی می‌کند. زمانی که یک tag در هدر مشخص می‌شود، به خودی خود مفهومی ندارد و در واقع مدیر شبکه باید به آن معنا دهد. مدیر شبکه مشخص می‌کند این tag چه عملی را انجام می‌دهد. 802.1q دارای ۴ بیت است که تنها ۳ بیت آن مربوط به priority است به نام Cost of Service (CoS) که بین ۰ و ۷ مقدار داده می‌شود. این فیلد برای تشخیص QoS است و به سوئیچ می‌فهماند که به سراغ این فیلد برود و ۷ را به عنوان بالاتر و صفر را به عنوان پایین‌ترین اولویت قرار دهد. مسأله‌ای که در اینجا اتفاق می‌افتد، تفاوت بسته ورودی با خروجی از سوئیچ است. اگر پورت خروجی access باشد هدر 802.1q حذف می‌شود و اگر trunk باشد هدر قبلی حذف و هدر جدیدی ایجاد می‌گردد و هدر اصلی از بین می‌رود. پس لایه ۲ قابل اطمینان نیست.

۲. Priority Tag in Layer-3: همانطور که گفته شد لایه ۲ قابل اطمینان نیست، پس در لایه ۳ از فیلدی به نام Type of Service (ToS) استفاده می‌شود. این فیلد end-to-end و همچنین رزرو شده است. حال اگر QoS بر روی سوئیچی تنظیم شود باید هدر لایه ۲ و ۳ قسمت CoS و ToS بسته را نگاه کند در غیر این صورت این مقادارها بی اهمیت است. استانداردهای ToS به دو نوع IP Precedency و DCSP است.

| Type of Service with IP Precedence |    |    |    |    |    |    |      |
|------------------------------------|----|----|----|----|----|----|------|
| P2                                 | P1 | P0 | T3 | T2 | T1 | T0 | Zero |

|                   |     |                 |                 |                 |
|-------------------|-----|-----------------|-----------------|-----------------|
| Ver               | IHL | Type of Service | Total Lenght    |                 |
| Identification    |     |                 | Flags           | Fragment Offset |
| Time To Live      |     | Protocol        | Header Checksum |                 |
| Source IP Address |     |                 |                 |                 |

|                        |
|------------------------|
| Destination IP Address |
| IP Options (if any)    |

ToS: IP Precedency دارای ۸ بیت است که ۳ بیت اول از سمت چپ همان CoS در لایه ۲ است و ۴ بیت بعدی برای درخواست منابع QoS است و بیت آخر آن همیشه صفر است. T3: بیت مربوط به Delay است. در صورت یک بودن درخواست منابع با کمترین Delay ارسال می‌شود. T2: بیت مربوط به Throughput است. در صورت یک بودن درخواست منابع با بیشترین Throughput ارسال می‌شود. T1: بیت مربوط به Reliability است. در صورت یک بودن درخواست منابع با بیشترین Reliability ارسال می‌شود. T0: رزرو شده برای آینده است. شبکه‌های کمتری به بیت‌های T توجه می‌کنند. همچنین ترتیب این بیت‌ها نام گذاری شده و حائز اهمیت هستند.

|     |           |     |                      |
|-----|-----------|-----|----------------------|
| 000 | Routine   | 100 | Flash Override       |
| 001 | Priority  | 101 | Critical             |
| 010 | Immediate | 110 | InterNetwork Control |
| 011 | Flash     | 111 | Network Control      |

استاندارد دیگر ToS به نام DSCP (Diffrentiated Services Code Point) است که استاندارد بعد از IP Precedency و اصلاح شده آن است و به جای ۳ بیت از ۶ بیت برای priority استفاده می‌کند. ۳ بیت اول آن از چپ به عنوان CS (Class Selector) و ۳ بیت بعدی به عنوان DP (Precedence Drop) و ۲ بیت بعدی ECN (Explicit Congestion Notification) نامیده می‌شود. برای تبدیل آن به IP Precedency، ۳ بیت CS همان CoS می‌شود. در DSCP هر ۳ CS زیر مجموعه دارد که به نام Assured Forwarding (AF) است و اهمیت حذف بسته را تعیین می‌کند. به عبارتی هر کلاس به سه دسته حذف کمتر، حذف متوسط و حذف زیاد تقسیم می‌شود. نحوه نام گذاری آن به صورت زیر است:

۱. فیلد AF از شماره ۰ تا ۶۳ مقدار دهی می‌شود.
۲. مقادیر خاصی از AF11 تا AF43 و CS1 تا CS7 می‌باشند و بقیه شماره‌ها با همان نام DSCP شناخته می‌شوند.
۳. مقادیر AFها به این صورت هستند AF## که دو عدد بعد از آن، اولی از چپ نشان دهنده شماره کلاس و عدد دومی نشان دهنده میزان حذف بسته‌ها است.
۴. در واقع CSها به همان AFها به صورت AF#0 هستند که رقم دوم آنها صفر است.

۵. مقدار DSCP46 نیز رزرو شده برای صوت است.

| Type of Service with DSCP |     |     |     |     |     |      |      |
|---------------------------|-----|-----|-----|-----|-----|------|------|
| DS5                       | DS4 | DS3 | DS2 | DS1 | DS0 | ECN1 | ECN0 |

|                       | Class 1           | Class 2           | Class 3           | Class 4           |
|-----------------------|-------------------|-------------------|-------------------|-------------------|
| Low drop probability  | AF11<br>(DSCP 10) | AF21<br>(DSCP 18) | AF31<br>(DSCP 26) | AF41<br>(DSCP 34) |
| Med drop probability  | AF12<br>(DSCP 12) | AF22<br>(DSCP 20) | AF32<br>(DSCP 28) | AF42<br>(DSCP 36) |
| High drop probability | AF13<br>(DSCP 14) | AF23<br>(DSCP 22) | AF33<br>(DSCP 30) | AF43<br>(DSCP 38) |

Class Selector 2 = IP Precedence 2

|      |                       |               |          |           |
|------|-----------------------|---------------|----------|-----------|
| AF42 | 100<br>Class Selector | 10<br>DS2/DS1 | 0<br>DS0 | 00<br>ECN |
|------|-----------------------|---------------|----------|-----------|

QoS Trust Boundary: زمانی که QoS در Globall وارد شود، بیت‌های اولویت لایه ۲ و ۳ بسته‌هایی که وارد می‌شوند به صفر تبدیل می‌شوند زیرا به عنوان بسته‌های نامطمئن هستند. برای رفع آن باید داخل پورت به صورت دستی Trust Port را تنظیم نمود تا مقدار QoS را تغییر ندهد.

## Configure QoS – ۴-۲۴

برای اختصاص اولویت از دستور زیر استفاده می‌شود:

```
(Config)#class-map <name>
```

دستور بالا مانند route-map در مسیریابی است.

```
(Config-map)#match ip <dscp | precedence > <value>
```

```
(Config)#mls qos
```

اگر فقط دستور بالا در Global تنظیم شود، به صورت Untrusted است و از FIFO استفاده می‌شود.

```
(Config-if)#mls qos trus <cos | ipp | dscp >
```

```
(Config-if)#mls qos trust device cisco-phone -> understand from CDP
```

```
(Config-if)#switchport pririty extend [cos <value | trust]
```

Trust در خط فرمان بالا برای 802.1p است.

802.1p شیه 802.1q است با این تفاوت که VLAN-ID برابر صفر است و بیت priority دارد. اگر بسته از سمت PC یا end-device بیاید می توان بیت CoS برای آن تنظیم و در سوئیچ دستور trust را وارد نمود تا بدون تغییر بسته را forward کند.

## ۲۴-۵ Auto QoS

این مفهوم به صورت macro و از پیش تعریف شده است. به این صورت عمل می کند که ابتدا باید QoS فعال باشد. سپس CoS را به DSCP تبدیل می کند. در واقع map کردن اولویت لایه ۲ به ۳ است. پس از آن صف های ورودی و خروجی را تنظیم می کند. این قابلیت صف اولویت پایدار محکمی برای ترافیک خروجی صوت و بر روی اینترنتی Trust Boundary QoS پایداری دارد.

```
(Config-if)#auto qos voip <cisco-phone | cisco-softphone | trust>  
(SW1)#show mls qos interface <interface-number>
```

## ۲۵- Securing Switch Access

در این قسمت به نحوه کار پروتکل های احراز اصالت و بالا بردن امنیت سوئیچ تا حدودی پرداخته شده است.

## ۲۵-۱ Authentication, Authorization and Accounting (AAA)

AAA مخفف Authentication، Authorization و Accounting است. در مواقعی استفاده می شود که کاربر نیاز به دسترسی به CLI سوئیچ یا مسیریاب یا سیستم دیگری داشته باشد. اگر در شبکه کاربران بسیار زیادی وجود داشته باشد که هر کدام نام و رمز عبور مختص به خود داشته باشند، ایجاد و تغییر دادن این مشخصات به صورت محلی دشوار است. برای آسانی می توان بر روی سرور AAA که خود دارای پایگاه داده است این مشخصات را ایجاد نمود. این سرور با سیستم های دیگر در ارتباط است و یک Credential با همه سیستم ها دارد و نیاز به تعریف به صورت محلی نیست. همچنین در این سرور می توان ACL یا Policy های مربوط به کاربران را Upload یا Download نمود. به طور مثال پس از تأیید هویت کاربر، میزان سطح دسترسی آن به سوئیچ را به صورت پویا می توان اضافه نمود. همچنین با استفاده از این سرویس می توان سطح دسترسی کاربر به شبکه را تغییر داد. به طور مثال هنگام اتصال به شبکه بیسیم ابتدا کاربر تأیید و سپس به آن سرویس دهی شود. پروتکل 802.1x نمونه ای از این پروتکل است. به دلیل وجود ۳ مؤلفه Client، Network Access Server (NAS) و Server دارای ۳ مرحله پردازی است که در زیر به توضیح هر یک پرداخته می شود:

۱. Client: دستگاهی است که نیاز به دسترسی شبکه دارد و می‌خواهد از یک سرویس دهنده سرویس دریافت نماید.

۲. Network Access Server (NAS): اولین دستگاهی است که Client به آن می‌رسد و احراز اصالت آن را بررسی می‌کند که خود به صورت از راه دور به server متصل می‌شود.

۳. AAA Server: در دستگاه NAS نیاز به ذخیره و نگهداری اطلاعات کاربر نیست. در واقع نیاز به سرویس دهنده مرکزی است تا اطلاعات کاربر در آن ذخیره شود و تمامی دستگاه‌های NAS به آن مراجعه کند. این سرور دارای تمامی Credentialهای دستگاه‌های سرویس دهنده است.

وظیفه سرویس دهنده AAA در زیر توضیح داده شده است:

۱- Authentication: این مؤلفه به تأیید احراز اصالت client می‌پردازد. وجود کاربر در پایگاه داده و تأیید اطلاعات را بررسی می‌کند. روش‌های مورد استفاده شامل نام کاربری و رمز عبور، کارت الکترونیکی، MAC Address و ... است.

۲- Authorization: سطح دسترسی کاربر، اجازه عبور و این که کدام کاربر به کدام سطح مجوز داشته باشد را بررسی می‌کند. سطح‌های دسترسی مورد استفاده در سوئیچ می‌تواند شامل: ۱- Basic Network Access ۲- CLI Availability ۳- VLAN Assignment ۴- Dynamic QoS Policy ۵- Dynamic ACL و ... باشد.

۳- Accounting: این مؤلفه اطلاعات آماری در رابطه با پیام‌های ورود و خروج، دسترسی‌های غیر مجاز، نوع سرویس و ... را جمع‌آوری می‌کند. بیشتر افراد این سیستم را پیاده‌سازی نمی‌کنند و از دو مؤلفه پیش جدا می‌دانند.

## Relationship between AAA Server and NAS ۱-۱-۲۵

برای ارتباط و همچنین تبادل اطلاعات بین client و AAA Server نیاز به پروتکل ارتباطی احساس می‌شود. نمونه‌ای از این پروتکل در لایه دو 802.1x است. در این جا client می‌تواند سوئیچ یا مسیریاب باشد، زیرا خود از AAA Server سرویس می‌گیرد و به عنوان NAS عمل می‌کند. پروتکل‌هایی که از طریق راه دور، عملیات احراز اصالت و خدمات مربوط به کنترل دسترسی شبکه را از طریق یک سیستم مرکزی انجام می‌دهند به شرح زیر است:

۱. TACACS+ (Terminal Access Controller Access Control System): این پروتکل مخصوص شرکت سیسکو است و برای تبادل اطلاعات AAA طراحی شده است. همچنین سه مؤلفه Authentication، Authorization و Accounting را به صورت جداگانه ارائه می‌دهد. به عبارت

دیگر می‌توان از پروتکل دیگری (مانند Kerberos) برای Authentication و از TACACS+ برای Authorization و Accounting استفاده نمود.

۲. RADIUS (Remote Authentication Dial In User Service): همانند پروتکل قبل برای تبادل اطلاعات AAA طراحی شده است با این تفاوت که در استاندارد IEEE با RFC 2058 تعریف و بعد از آن چندین بار به روز شده است. تفاوت دیگر آن این است که در پروتکل قبل سه مؤلفه جدا از هم بوده و می‌توان از پروتکل‌های جایگزین استفاده نمود ولی در این پروتکل Authentication و Authorization در قالب یک بسته بوده و از پروتکل سومی نمی‌توان استفاده نمود.

### AAA IOS Configuration ۲-۱-۲۵

برای تنظیم AAA بر روی تجهیزات سیسکو، ابتدا باید سیستم کاربر تنظیم شود و مشخص نمود کدام interface به AAA متصل است. پس از آن باید در سوئیچ یا مسیریاب AAA Server تنظیم شود و مشخص نمود کدام استاندارد استفاده می‌شود. AAA به صورت پیش فرض غیر فعال است. همانطور که گفته شد NAS خود یک سرویس گیرنده است. برای تنظیم ارتباط باید از کلمه عبور برای احراز بین AAA Server و NAS استفاده نمود. واضح است که بین این دو باید ارتباط IP وجود داشته باشد.

```
(Config)#aaa new-model
```

```
(Config)#tacacs-server host <aaa-server> key <key-value>
```

یا

```
(Config)#radius-server host <aaa-server> key <key-value>
```

برای دسترسی پذیری بالا می‌توان از چندین سرویس دهنده استفاده نمود که سرویس دهنده دومی پشتیبان اولی باشد.

```
(Config)#radius-server host <aaa-server1>
```

```
(Config)#radius-server host <aaa-server2>
```

```
(Config)#radius-server key <key-value>
```

در خط فرمان پایین مشخص می‌شود چه موقع باید از AAA استفاده نمود.

```
(Config)#aaa authentication login default <enable | group | krb5 | ... >
```

```
(Config)#aaa authentication login default group radius
```

در خط فرمان زیر در صورت قطع سرور، از local خود استفاده نماید.

```
(Config)#aaa authentication login default group radius group tacacs+ local
```

همچنین می‌توان زمانی تنظیم نمود که بعد از چه مدت از local استفاده نماید.

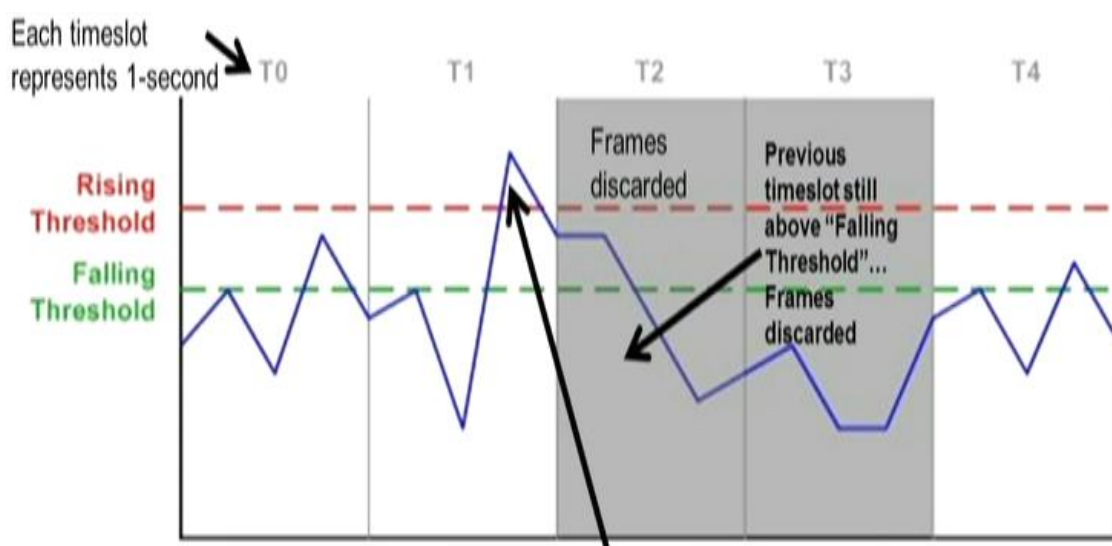
```
R1(config)#aaa authentication ?
  arap          Set authentication lists for arap.
  attempts      Set the maximum number of authentication attempts
  banner        Message to use when starting login/authentication.
  dot1x         Set authentication lists for IEEE 802.1x.
  enable        Set authentication list for enable.
  eou           Set authentication lists for EAPoUDP
  fail-message   Message to use for failed login/authentication.
  login         Set authentication lists for logins.
  onep          Set authentication lists for ONEP
  password-prompt Text to use when prompting for a password
  ppp           Set authentication lists for ppp.
  sgbp          Set authentication lists for sgbp.
  suppress      Do not send access request for a specific type of user.
  username-prompt Text to use when prompting for a username
```

## ۲-۲۵- Securing Switch Access (Storm Control)

گاهی اوقات در شبکه سیلی از بسته‌ها دیده می‌شود که ممکن است به دلیل ارسال بسته‌های Broadcast، Multicast و یا Unknown Unicast باشد. این سیل بسته‌ها باعث افزایش بار CPU به ۱۰۰ درصد و کاهش عملکرد شبکه و یا Host می‌شود. برای جلوگیری می‌توان از Storm Control استفاده نمود. این ابزار بر روی هر interface اجرا می‌شود و به صورت global نیست. قبل از پیاده‌سازی باید مشخص شود از کدام interface سیل بسته‌ها وارد می‌شود و یا این که بر روی تمام interface‌ها اجرا شود. برای فهمیدن این مطلب ابتدا باید storm را جست‌وجو و برای مشاهده Broadcast، Multicast و Unicast یا چند مورد را به صورت هم‌زمان مشخص نمود. به طور مثال در unicast storm بسته باز می‌شود و اگر آدرس MAC آن نا معتبر باشد آن را به عنوان storm شناسایی می‌کند. بعد از شناسایی storm باید از آن جلوگیری نمود. در هنگام جلوگیری باید دقت تمامی بسته‌ها را نمی‌توان فیلتر نمود. به طور مثال تمام بسته‌های Broadcast را نمی‌توان فیلتر نمود زیرا بسته‌های مدیریتی نیز جزئی از این بسته‌ها هستند. بنابر این باید یک مرزی تعیین شود که بعد از آن مرز بسته‌ها فیلتر شوند. این مرز یا Threshold می‌تواند یکی از موارد زیر باشد:

۱. درصد استفاده از پهنای باند: اگر پورتی از یک درصد بیشتر مورد استفاده قرار گرفت بسته‌ها حذف شوند.
۲. ترافیک بسته بر اساس Packet per Second: اگر ترافیک دریافتی بر اساس بسته بر ثانیه از مرز مشخصی عبور کرد بسته‌ها حذف شوند.
۳. ترافیک بسته بر اساس Bit per Second: اگر ترافیک دریافتی بر اساس بیت بر ثانیه از مرز مشخصی عبور کرد بسته‌ها حذف شوند.

۴. ترافیک بسته بر اساس Packet per Second و مورد استفاده برای small frames: این مرز می تواند برای هر interface و یا به صورت global تعریف شود. در بعضی از دستگاه ها دو مرز قابل تعیین است. Rising Threshold: اگر ترافیک از این مرز بالاتر رود بسته ها حذف می شوند. باید توجه داشت که ترافیک در یک بازه زمانی بررسی می شود و اگر در بازه زمانی قبلی از این مرز عبور کرد در بازه زمانی بعدی بسته ها حذف می شوند. Falling Threshold: اگر ترافیک از این مرز پایین نیاید همچنان بسته ها در بازه زمانی بعدی حذف می شوند. در شکل زیر نمودار بر اساس نرخ بسته و بازه زمانی، مورد بررسی قرار گرفته است. در بازه زمانی T1 نرخ بسته بالاتر از حد مجاز است و در بازه T2 بسته ها حذف می شوند.



دستورات مربوط به کنترل سیل یا Storm Control به صورت زیر است:

```
(Config-if)#storm-control broadcast level <rising-threshold-of-bandwidth>
```

در بعضی از platform ها به صورت زیر است:

```
(Config-if)#storm-control <broadcast | multicast | unicast> level <rising-threshold>  
<falling-threshold>
```

دستور زیر interface را به حالت errdisable می برد:

```
(Config-if)#storm-control action shutdown  
(SW1)#show storm-control
```

هنگامی که پورت به حالت errdisable می رود باید به صورت دستی فعال نمود و یا از دستور زیر استفاده نمود:

```
(Config)#errdisable recovery cause storm-control  
(Config)#errdisable recovery interval <after-second-for-enable>
```

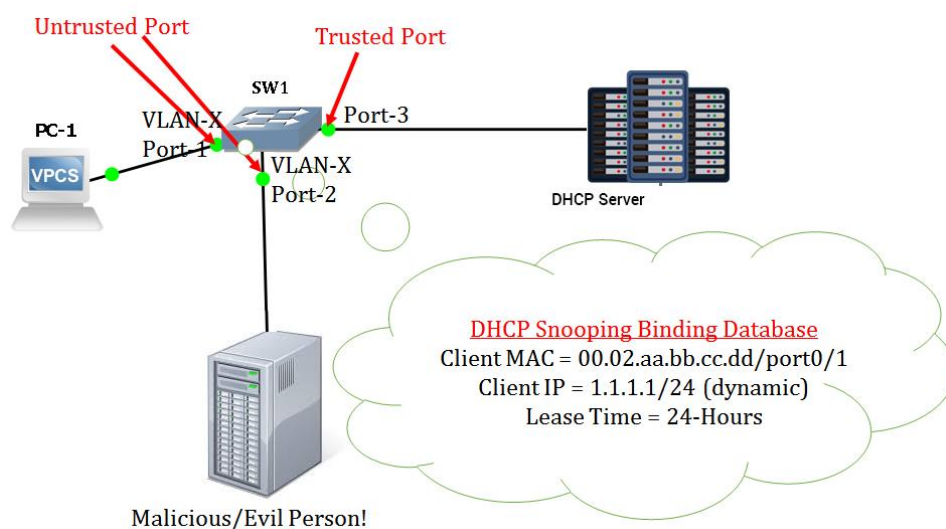
## DHCP Snooping – ۳-۲۵

زمانی که یک کاربر بسته DHCP Discover را ارسال می کند، همه افراد آن را می بینند زیرا این بسته به صورت Broadcast است. شخص خرابکار می تواند بسته DHCP Response را شبیه سازی و به کاربر ارسال کند. در همین هنگام DHCP Server نیز پاسخ می دهد ولی به دلیل ارسال چندین بسته از خرابکار و یا ارسال زودتر، کاربر بسته خرابکار را قبول می کند. بسته DHCP Offer که از سمت خرابکار ارسال می شود ممکن است شامل اطلاعات اشتباهی مانند default gateway و DNS باشد که پیامدهای نا مطلوبی را در پی خواهد داشت. ابزاری در سوئیچ وجود دارد که می توان از این پیامد جلوگیری نمود. اصطلاحاتی که در این زمینه وجود دارد شامل موارد زیر است:

۱. Untrusted Port: زمانی که DHCP Snooping بر روی سوئیچ فعال می شود، تمامی پورت های سوئیچ به حالت untrusted می رود و به صورت دستی باید پورت ها را به حالت Trust برد. این عمل شبیه QoS است.

۲. DHCP Snooping Binding Database: زمانی که بسته DHCP Discover از پورت Untrusted وارد سوئیچ می شود، جدولی ایجاد می شود و اطلاعات MAC، پورت کاربر و ... را به پورت Trusted ارسال می کند. تمامی بسته های DHCP که از پورت Untrusted وارد سوئیچ می شود پس از کامل شدن Binding Database به پورت Trusted ارسال می شود. جدول DHCP Binding Database فقط برای DHCP Snooping نیست و برای ابزارهای امنیتی دیگر مانند Dynamic ARP Inspection نیز استفاده می شود.

روند ارسال بسته ها از پورت Untrusted به صورت DHCP Discover، DHCP Request/Inform، DHCP Release، Decline و از پورت Trusted به صورت DHCP Offer، DHCP ACK، DHCP NACK.



است. ابزار DHCP Snooping مانند Uplink Fast برای سوئیچ‌های Access استفاده می‌شود زیرا بیشتر سیستم‌هایی که نیاز به DHCP دارند کاربران End Device هستند که در لایه Access قرار دارند. همچنین اگر در لایه‌های دیگر مانند Core یا Distributed باشد ممکن است دچار اشکال شود، زیرا اگر در بین راه مسیریابی باشد از IP Helper-address استفاده کند و سوئیچ DHCP Snooping به دلیل این که این بسته Unicast است نه Broadcast بسته را حذف می‌کند.

```
(Config)#ip dhcp snooping
```

```
(Config)#ip dhcp snooping vlan <vlan-id>
```

دستورات زیر برای اینترفیس است و دستور اول جهت محدود کردن میزان بار است که برای جلوگیری از crash و همچنین DOS Attack (Maximum Rate Packet Per Second) مورد استفاده می‌شود.

```
(Config-if)#ip dhcp snooping limit rate <1-2048>
```

```
(Config-if)#ip dhcp snooping trust
```

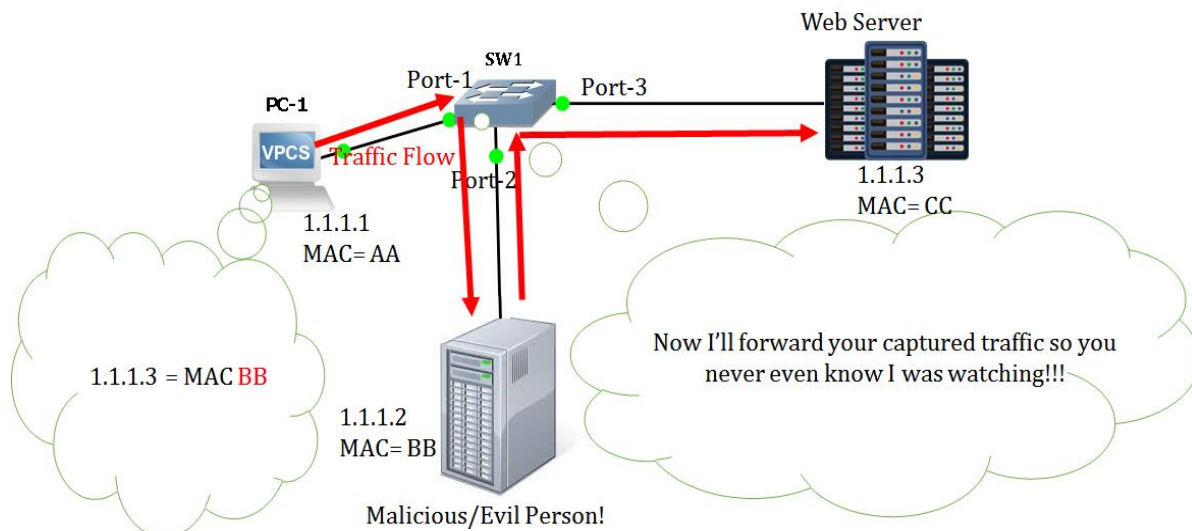
این دستور از گذاشتن اطلاعات کاربر و ارسال آن به DHCP Server با کد ۸۲ جلوگیری می‌کند. به صورت پیش فرض از آن استفاده می‌شود و بیشتر پروتکل‌های امنیتی DHCP این بسته را Drop می‌کنند.

```
(Config)#no ip dhcp snooping information option
```

```
(SW1)#show ip drop snooping binding
```

## ۲۵-۴- Dynamic Arp Inspection

زمانی که یک کاربر بسته ARP Request را ارسال می‌کند، به دلیل ارسال بسته به صورت Broadcast، فرد خرابکار این بسته را مشاهده می‌کند. فرد خرابکار پاسخ آن را با IP جعلی می‌دهد که این باعث می‌شود بسته‌ها از سمت کاربر به سمت فرد خرابکار و بعد از آن به سرویس دهنده ارسال گردد.



در واقع فرد خرابکار هر آدرسی را که بخواهد می تواند جعل کند. DAI با استفاده از جدول DHCP Snooping Binding می تواند ARP Request/Reply را بررسی نماید و در صورت صحیح نبودن بسته ARP، پیغام (log message) دهد. زمانی که سیستمی برای اولین بار Boot می شود نیاز به IP دارد که بسته DHCP و ARP برای تنظیم default gateway ارسال می شود. اگر در سوئیچ DHCP Snooping تنظیم شده باشد جدول DHCP Snooping Binding که شامل IP، MAC، SRC Port و VLAN است ایجاد می شود. این جدول کمک می کند تا از جعل آدرس خرابکار جلوگیری شود. مشکلی که در این هنگام پیش می آید آدرس های ایستا هستند مانند default gateway که در این جدول قرار نمی گیرند و سوئیچ این بسته ها را حذف می کند. در نتیجه هنگام ارسال بسته ARP Request کاربر مشکلی پیش نمی آید ولی هنگام ارسال بسته ARP Response سرویس دهنده، این بسته حذف می شود، زیرا در جدول DHCP Snooping Binding وجود ندارد. برای جلوگیری از حذف بسته ها با IP ایستا می توان از Static ARP ACL برای اطلاعات مسیریاب استفاده نمود و یا از ابزار DAI Trusted/Untrusted Port استفاده نمود.

دستورات زیر برای تنظیم DAI و Trust است.

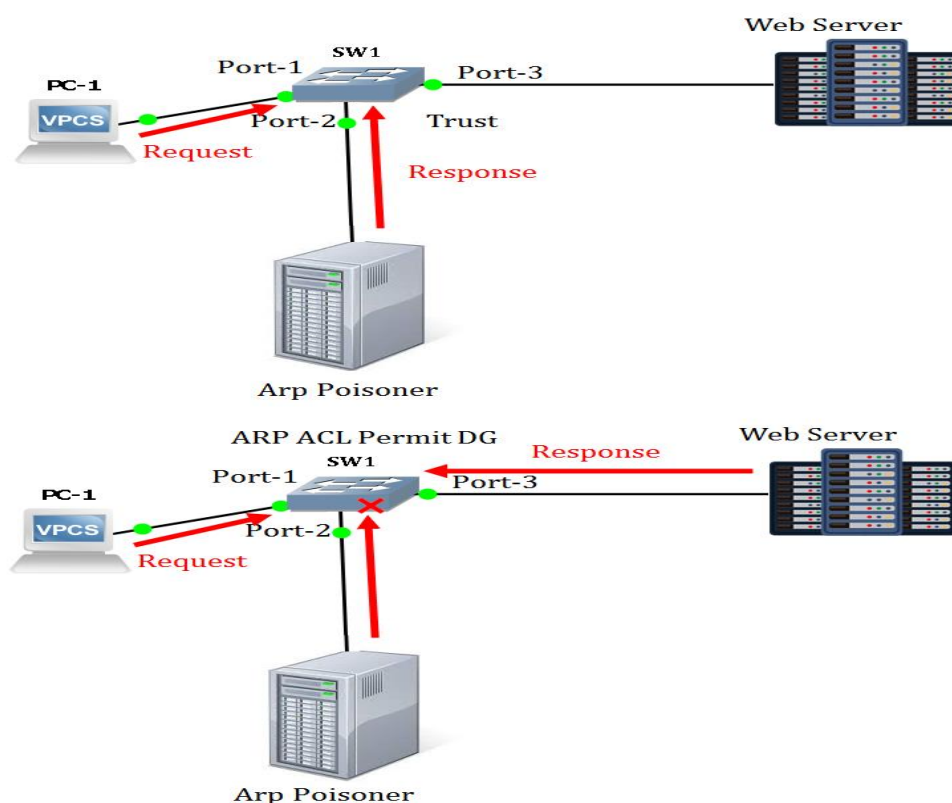
```
(Config)#ip arp inspection vlan <vlan-id>
(Config-if)#ip arp inspection trust
```

دستورات زیر برای ایجاد ARP ACL است.

```
(Config)#arp access-list <name>
(Config-ac)#permit ip host <ip-add> mac host <mac-add>
(Config)#ip arp inspection filter <arp-acl-name> vlan <vlan-id> [static]
```

دستور static دلخواه است ولی قدرت دستور را بالا می‌برد. در حالت عادی هنگام دریافت بسته ARP ابتدا static و سپس جدول DHCP Snooping Binding بررسی می‌شود ولی هنگام استفاده از static دیگر جدول DHCP Snooping Binding بررسی نمی‌شود.

زمانی که ARP Poisoner سمت سرور یا default gateway باشد باید از ARP ACL استفاده شود زیرا اگر از trusted port استفاده شود چون از یک interface عبور می‌کند و این پورت trust شده است می‌تواند آدرس را جعل کند.



هنگامی که Dynamic ARP Inspection (DAI) فعال می‌شود فیلدهای MAC، IP و VLAN مبدأ بررسی می‌شود. اگر مدیر سیستم بخواهد اطلاعات مقصد را نیز بررسی کند باید از دستور زیر استفاده نماید.

(Config)#ip arp inspection validate <[src-mac] [des-mac] [ip-target] >

| ARP Frame Format |             |               |                          |                        |          |               |         |            |               |            |               |
|------------------|-------------|---------------|--------------------------|------------------------|----------|---------------|---------|------------|---------------|------------|---------------|
| Dest MAC Add     | Src Mac Add | Type (0x0806) | H/W Type (0x01) Ethernet | Protocol Type (0x0800) | H/W Size | Protocol Size | Op Code | Sender MAC | Sender IP Add | Target MAC | Target IP Add |

ارسال بیش از حد درخواست‌های ARP نوعی حمله DoS است که ممکن است باعث کندی یا اختلال در سوئیچ و یا سرور شود. ابزار DAI قابلیت محدودیت ارسال بسته را نیز دارد. به صورت پیش فرض 15 بسته در هر ثانیه می‌توان ارسال نمود. این قابلیت و محدودیت بر روی پورت Trust اعمال نمی‌شود. برای محدود کردن ارسال بسته بر ثانیه از دستور زیر می‌توان استفاده نمود.

```
(Config-if)#ip arp inspection limit rate <0-2048>
(SW1)#show ip arp inspection
```

## ۲۶- Security VLAN

راه کارهایی برای امن نمودن VLAN، ارتباط بین آنها و استفاده در مسیریابی وجود دارد که در ادامه به توضیح برخی از آنها پرداخته می‌شود.

## ۲۶-۱- VLAN Access List

نام‌های دیگر آن VACL و VLAN Access-map است که شامل دو نوع زیر است:

۱- Routed access-list: بر روی اینترفیس لایه ۳ فعال می‌شود حال چه VLAN و چه routing باشد. فقط

برای ترافیک‌های مسیریابی استفاده می‌شود. اجازه و یا عدم اجازه ارسال ترافیک از یک شبکه خاص و مسیریابی به شبکه دیگر را فراهم می‌کند.

۲- VLAN access-list: می‌توان برای ترافیک مسیریابی استفاده نمود ولی قدرت و طراحی آن برای

ترافیک‌های Bridge است. اجازه و یا عدم اجازه ارتباط هاست‌ها بین VLANها را فراهم می‌کند. بر روی اینترفیس‌های لایه ۳ قابل تنظیم نیست و فقط بر روی VLAN Interface تنظیم می‌شود. هر ترافیکی ورودی و خروجی از VLAN در صورت مطابقت با VACL می‌تواند عبور و یا متوقف شود. مانند rout-map است که در این جا access-map نامیده می‌شود.

VACL می‌تواند مؤلفه‌های IP، IPX و MAC را بررسی نماید. در صورت تنظیم VACL برای MAC address باید از mac access-list استفاده شود. اولین مؤلفه‌ای که VACL بررسی می‌کند Ethernet type code است. در صورتی که mac access-list تنظیم شده باشد و Ethernet type code ترافیک مورد بررسی، شامل IP یا IPX باشد، از mac access-list گذر کرده و به بررسی ip access-list می‌پردازد. عملیات‌هایی که VACL انجام می‌دهد شامل Forward، Drop، Redirect و Capture است. عمل Redirect در سوئیچ به این شکل است که سوئیچ برای ارسال کردن (Forward) ترافیک جدول‌های MAC یا CAM را مشاهده می‌کند، در این حالت می‌توان یک اینترفیس خروجی را تعیین نمود تا ترافیک از آن اینترفیس ارسال (redirect) شود. در همه سوئیچ‌ها

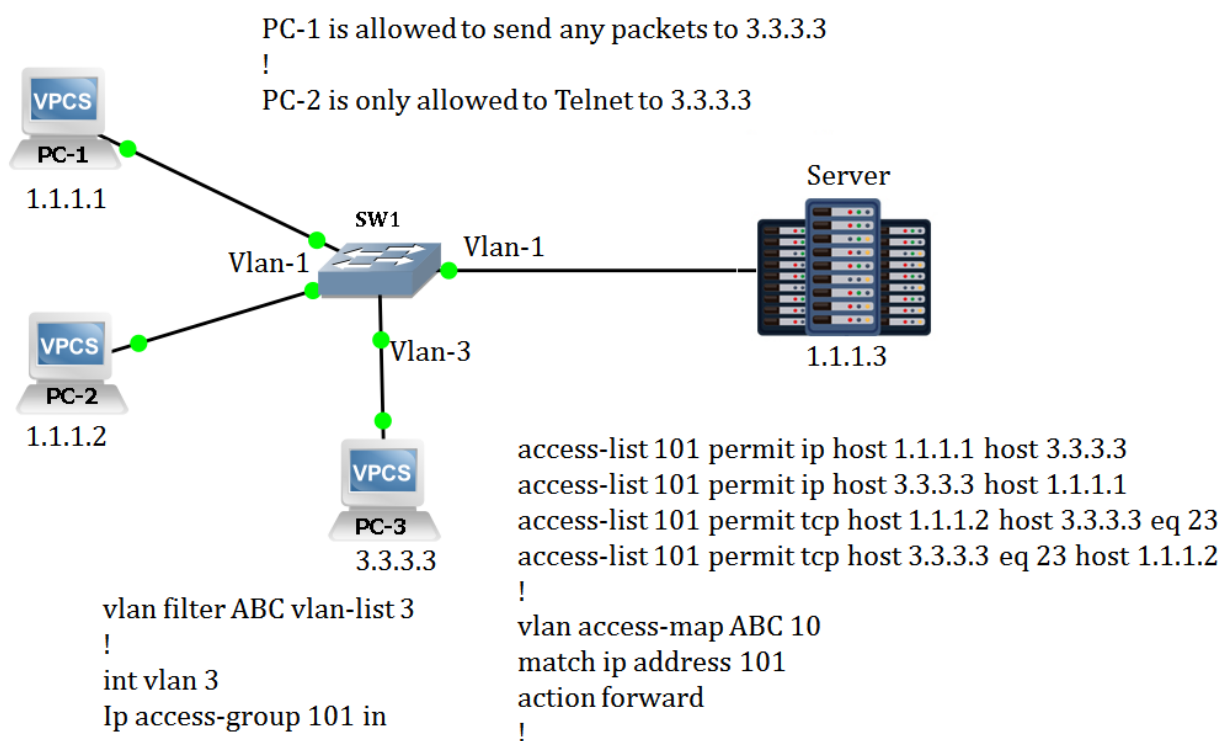
قابلیت Capture وجود ندارد. این قابلیت اجازه کپی ترافیک و ارسال آن به یک اینترفیس دیگر برای مشاهده ترافیک را می‌دهد. دستورات زیر تنظیم VACL را نشان می‌دهد.

```
(Config)#vlan access-map <map-name> [sequence-number]
(Config-access-map)#match <ip|ipx address> <acl-number | acl-name>
(Config-access-map)#action <drop | forward [capture] | redirect <type-if/number>
(Config)#vlan filter <map-name> vlan-list <vlan-id>
```

Permit یا Deny در access-list به معنای استفاده و یا عدم استفاده در access-map است.

تفاوت بین VACL و RACL به شرح زیر است:

۱. RACL بر روی Routed Interface انجام می‌گیرد اما VACL به صورت global تنظیم می‌شود.
۲. VACL هم برای RACL استفاده می‌شود و هم برای Bridge Traffic اما RACL فقط برای routed interface است.



## Private VLAN – ۲-۲۶

دو دلیل اصلی برای ساخت Private VLAN مطرح است. دلیل اول اینکه تنها ۴۰۹۶ VLAN می‌توان ایجاد کرد که این باعث محدودیت تعداد VLAN است. البته از access-list نیز می‌توان استفاده نمود ولی اگر تعداد کاربران بسیار زیاد باشد این راه آسان نیست و سربار مدیریتی به همراه دارد. دلیل دوم اینکه اگر تعدادی از کاربران در VLAN‌های

متفاوتی باشند و همه آنها دارای Policy های یکسانی باشند می توان از Private VLAN استفاده نمود. یا به عبارت دیگر کاربران در یک VLAN و یک Subnet باشند ولی بین آنها یک Security Policy باشد. Private VLAN در واقع ترکیبی از دو VLAN است. Primary VLAN که وظیفه کنترل دسترسی های IP آن شبکه را برعهده دارد. همه کاربران دارای یک سیاست خاص در آن قرار می گیرند. Secoundary VLAN که وظیفه کنترل سیاست های امنیتی بین کاربران عضو Primary VLAN را بر عهده دارد یا به عبارتی چه کاربری می تواند با کاربر دیگر در ارتباط باشد یا نباشد. Secoundary VLAN خود دارای دو نوع است.

۱. Community: تمام کاربرانی که در این نوع هستند می توانند با یکدیگر در ارتباط باشند. همه کاربران در یک شبکه IP هستند که به عنوان primary VLAN شناخته می شود. همه کاربران در یک Broadcast domain هستند. هیچ کاربری نمی تواند با یک Secoundary VLAN دیگر ارتباط برقرار کند.

۲. Isolate: هیچ کاربری نمی تواند با کاربری که در Isolate یا Community است ارتباط برقرار کند. مانند Community همه کاربران در یک شبکه IP هستند که به عنوان primary VLAN شناخته می شود. یک Private VLAN می تواند چندین Community VLAN ولی تنها یک Isolated داشته باشد.

با تعریف مفاهیم قبل ارتباط کاربران با یکدیگر مشخص شد. حال ارتباط کاربران با Default Gateway و سایر مسیریاب ها باید مشخص شود. این ارتباط توسط پورتی به نام Promiscuos انجام می شود. این پورت می تواند یک اینترفیس فیزیکی باشد که به یک مسیریاب یا سوئیچ چندلایه یا virtual interface متصل باشد.

پیش فرض های تنظیم PVLAN به شرح زیر است:

۱. سوئیچ باید در حالت VTP Transparent باشد. البته تنها VTP v3 از آن پشتیبانی می کند.  
۲. VLAN مورد استفاده برای Primary و یا Secoundary نباید قبلاً مورد استفاده شده باشد که به هر دو نمی توان اختصاص داد.

۳. برای عبور PVLAN باید Trunking PVLAN فعال باشد.

۴. برای تنظیم Etherchannel نباید PVLAN بر روی پورت ها فعال باشد. همانطور که مشخص است STP فقط برای Primary VLAN می باشد و بر روی PVLAN ها به دلیل منحصر به فرد بودن Broadcast Domain و نبودن حلقه معنایی ندارد.

برای تنظیم PVLAN به ترتیب Secoundary PVLAN، بعد از آن سه Primary VLAN ساخته می شود. پس از ساخت، باید این دو را به یکدیگر متصل نمود. بعد از آن باید تنظیمات مربوط به اینترفیس را انجام داد. پس از تنظیمات

ایترنیس Promiscuos تنظیم می‌شود. در آخر اتصال Primary/Secoundary Privat VLAN به Promiscuos را باید انجام داد.

```
(Config)#vlan <vlan-id>  
(Config-vlan)#private-vlan <community | isolate>
```

```
(Config)#vlan <vlan-id>  
(Config-vlan)#private-vlan primary
```

```
(Config)#vlan <primary-vlan-id>
```

```
(Config-vlan)#private-vlan accositation <add | remove> <secondary-vlan>
```

```
(Config-if)#switchport mode private-vlan host  
(Config-if)#switchport private-vlan host accosiation <primary-pvlan>  
<secondary-pvlan>
```

```
(Config-if)#switchport mode private-vlan promiscuos  
(Config-if)#switchport private-vlan mapping <primary-pvlan> <secondary-pvlan>
```

دستور زیر زمانی استفاده می‌شود که Virtual Interface به عنوان Primary قرار گیرد.

```
(Config-if)#private-vlan mapping <secondary-pvlan>
```