



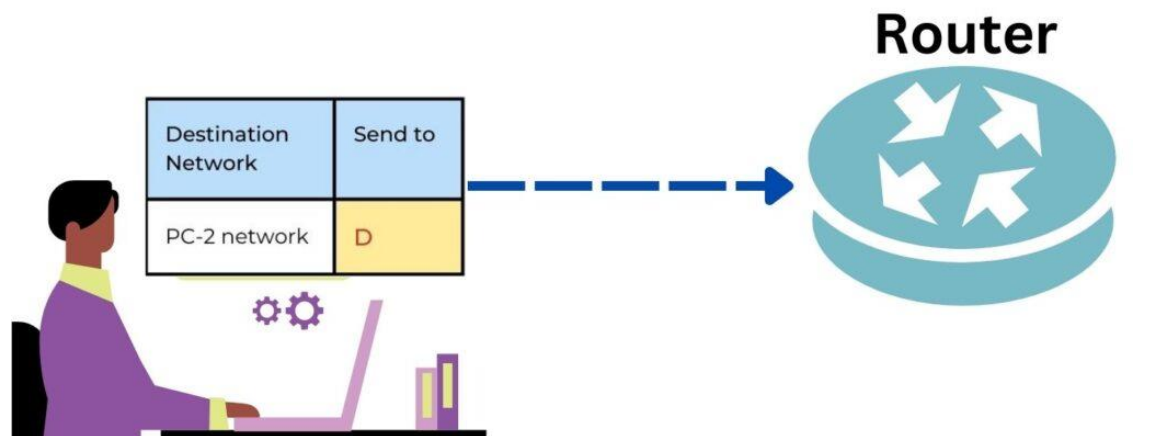
شبکه های کامپیوتری

(آمادگی برای آزمون Network+)

اسلاید دهم

(بخش دوم مسیریابی)

آشنایی با پروتکل های مسیریابی، RIP، IGRP، OSPF، IS-IS، EIGRP، BGP



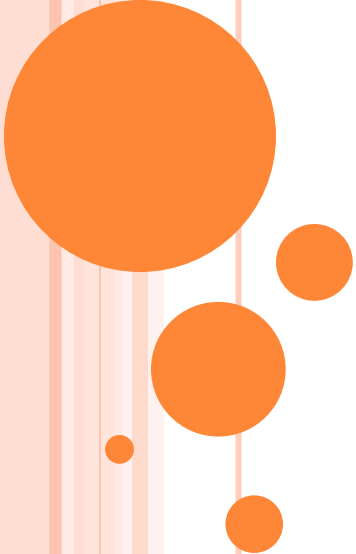
محمد سعید صفایی صادق

(استفاده از اسلایدها صرفا برای دانشجویان مجاز می باشد!)

۱۴۰۳

درس
شبکه های
کامپیوتری

۱۰



Distance Vector **RIP** Routing Information Protocol

پروتکل (RIP (Routing Information Protocol :

یک پروتکل Open Standard هست به اصطلاح میگویند Vendor Independent یعنی ازین پروتکل میتوان چه روی تجهیزات سیسکو و چه سایر تجهیزات مثل میکروتیک استفاده کرد.

همچنین این پروتکل یک پروتکل Metric هست یعنی مسیر ها را بر اساس تعداد Hop ها استفاده می شود. یعنی تعداد روتر های بین مبدا و مقصد.

نکته ۱: وقتی میگوییم $\text{Hop Count} = 0$ یعنی شبکه به صورت مستقیم یا directly به ما متصل است.

نکته ۲: بیشترین Hop Count بین مبدا و مقصد باید ۱۶ تا باشد و اگر بیشتر ازین باشد (مثلا ۱۶ تا) آنرا Infinity یا بینهایت در نظر میگیرند و عملاً دراپ خواهد شد.

پروتکل (RIP (Routing Information Protocol :

نکته ۳ : این پروتکل از نوع Distance Vector هست و AD آن برابر ۱۲۰ می باشد.

پیام های Advertise هر ۳۰ ثانیه یکبار آپدیت میکند.

اصلاحا به پروتکل RIP میگویند Routing on Rumours یعنی مسیر یابی بر اساس شایعه، به این معنا که من به همسایگان خود اعتماد دارم و همسایگان من هم به من اعتماد دارند (آپدیت جدول مسیریابی)

پس عملا هر اطلاعاتی که در رابطه با جدول مسیریابی دریافت میکند برایش مورد اعتماد است.

دارای قابلیت Passive Interface یعنی با توجه به اینکه پیام های آپدیت به صورت دوره ای از شبکه Lan ارسال می شود امکان دارد هکری در شبکه اسنیفینگ یا کپچرینگ انجام دهد، پس برای جلوگیری، این قابلیت را دارد هر

اینترفیسی را که میخواهد نگه دارد و یا آنرا Deny کند. شبکه های کامپیوتری - محمد سعید صفایی

پروتکل (RIP (Routing Information Protocol :

- ۱- معیار انتخاب مسیر: فقط به تعداد هاپ‌ها توجه می‌کند و کوتاه‌ترین مسیر (از نظر تعداد روترها) را انتخاب می‌کند.
- ۲- محدودیت هاپ‌ها: حداکثر تعداد هاپ‌ها ۱۵ است. اگر مسیر بیشتر از ۱۵ هاپ باشد، آن مسیر غیرقابل دسترسی در نظر گرفته می‌شود.
- ۳- به‌روزرسانی دوره‌ای: هر ۳۰ ثانیه، روترها اطلاعات جدول مسیریابی خود را به روترهای مجاور ارسال می‌کنند.
- ۴- پروتکل RIP : Distance Vector از الگوریتم Distance Vector استفاده می‌کند که در آن هر روتر فقط اطلاعاتی از روترهای مجاور خود دارد.
- ۵- بسته های RIP به سرعت در حال ردو بدل اطلاعات بین روتر ها هستند در صورت قطعی یکی از مسیرها، جدول مسیریابی را آپدیت میکند و مسیر جایگزین برای ارسال به دست می آید.

نسخه‌های پروتکل RIP

- ۱- **RIP نسخه ۱** **RIPv1**: اطلاعات مسیریابی را به صورت Broadcast به آدرس ۲۵۵.۲۵۵.۲۵۵.۲۵۵ ارسال می‌کند.
در این نسخه VLSM ساپورت نمی‌شود چون Classfull است.
- ۲- **RIP نسخه ۲** **RIPv2**: اطلاعات مسیریابی را به جای Broadcast، به صورت Multicast به آدرس ۲۲۴.۰.۰.۹ ارسال می‌کند، که باعث کاهش بار شبکه می‌شود.
- ۳- **RIPng (RIP Next Generation)** : **RIPng** برای شبکه‌هایی که از IPv6 استفاده می‌کنند طراحی شده است.
از آدرس Multicast FF02::9 برای ارسال اطلاعات مسیریابی استفاده می‌کند.

پروتکل (RIP (Routing Information Protocol :

فرض کنید سه روتر داریم:

Router A → Router B → Router C

مرحله ۱:

در ابتدا، هر روتر فقط از خودش خبر دارد.

جدول Router A: [(0 هاپ) A]

جدول Router B: [(0 هاپ) B]

جدول Router C: [(0 هاپ) C]

مرحله ۲:

Router A جدول خود را به Router B می‌فرستد و می‌گوید: "من به A دسترسی دارم (0 هاپ)".

Router B جدول خود را به Router A و C می‌فرستد.

جدول‌ها به‌روزرسانی می‌شوند:

جدول Router B: [(0 هاپ) B, (1 هاپ) A]

جدول Router C: [(0 هاپ) C, (1 هاپ) B]

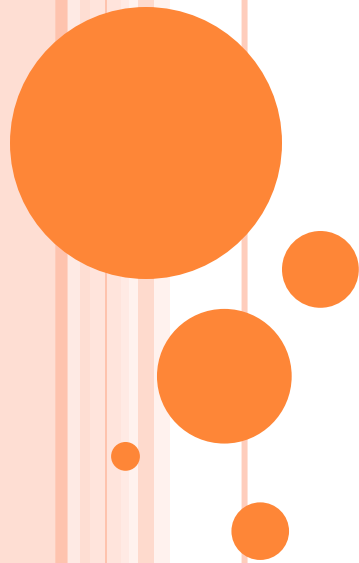
جدول Router A: [(0 هاپ) A, (1 هاپ) B]

مرحله ۳:

این روند ادامه پیدا می‌کند تا همه روترها اطلاعات مسیرهای مختلف را داشته باشند.

Distance Vector IGRP

Interior Gateway Routing Protocol

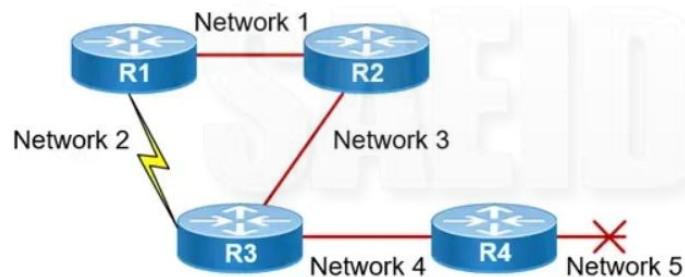


پروتکل (Interior Gateway Routing Protocol) IGRP :

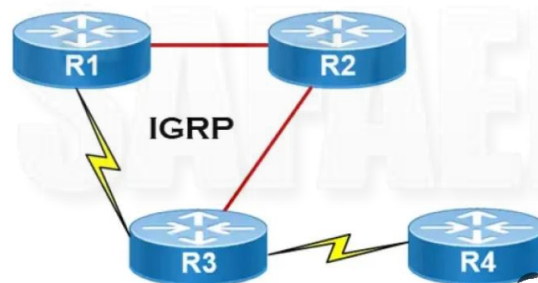
IGRP یک پروتکل مسیریابی Classfull است که توسط سیسکو طراحی شده است و برای شبکه‌های متوسط و بزرگ استفاده می‌شود.

این پروتکل مانند RIP از روش Distance Vector استفاده می‌کند اما نسبت به RIP پیشرفته‌تر است.

IGRP برای غلبه بر محدودیت‌های RIP مانند تعداد هاپ‌ها و نادیده گرفتن سرعت لینک طراحی شده و از معیارهای متنوع‌تری برای انتخاب بهترین مسیر استفاده می‌کند.



RIP Protocol



پروتکل (IGRP (Interior Gateway Routing Protocol :

- ۱- معیارهای مسیریابی متنوع: از معیارهایی مثل پهنای باند، تأخیر، قابلیت اطمینان و بار شبکه برای انتخاب مسیر استفاده می‌کند. برخلاف RIP که فقط تعداد هاپ‌ها را در نظر می‌گیرد.
- ۲- محدودیت تعداد هاپ‌ها: حداکثر تعداد هاپ‌ها ۲۵۵ (پیش‌فرض ۱۰۰) است، در حالی که RIP فقط تا ۱۵ هاپ را پشتیبانی می‌کند.
- ۳- زمان به‌روزرسانی یا Advertising: جداول مسیریابی هر ۹۰ ثانیه به‌روزرسانی می‌شوند (نسبت به RIP که ۳۰ ثانیه بود، کارآمدتر است).
- ۴- پشتیبانی از توپولوژی‌های پیچیده: برای شبکه‌هایی با چند مسیر به مقصد Multi-path مناسب است.

پروتکل (IGRP (Interior Gateway Routing Protocol :

نکته : اگر یک Update از یک روتر موجود در شبکه IGRP به مدت ۳ وهله زمانی یا ۲۷۰ ثانیه دریافت نشود روترهای شبکه Route مربوط به این روتر را نامعتبر یا Invalid می کنند ، اگر این فرآیند ادامه پیدا کند و بعد از ۷ وهله زمانی همچنان Routing Update از یک روتر در شبکه دیده نشود روتر مورد نظر بصورت کامل از Routing Table های موجود در IGRP شبکه حذف خواهد شد.

پروتکل (IGRP (Interior Gateway Routing Protocol :

بر خلاف RIP که تنها عامل متریک آن تعداد Hop ها بود، در پروتکل IGRP برای محاسبه Metric از پارامترهایی مثل **Delay**، **Bandwidth**، **Reliability** و **MTU** Load برای محاسبه بهترین مسیر برای هر Packet استفاده می کند.

بصورت پیشفرض الگوریتم مورد استفاده در IGRP فقط از **Bandwidth** و **Delay** برای تعیین Metric استفاده می کند اما سایر فاکتورهای محاسبه Metric هم می تواند فعال کرد.

پروتکل (IGRP (Interior Gateway Routing Protocol

فرض کنید سه روتر به این شکل به هم متصل هستند:

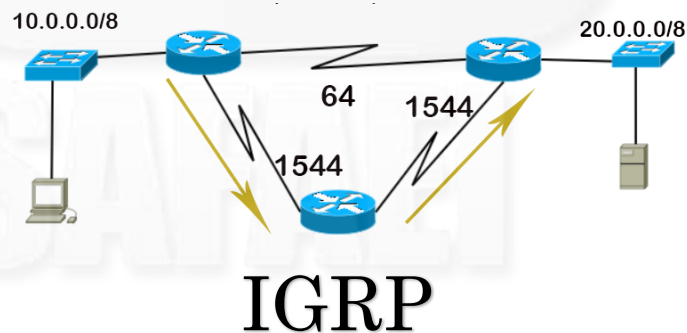
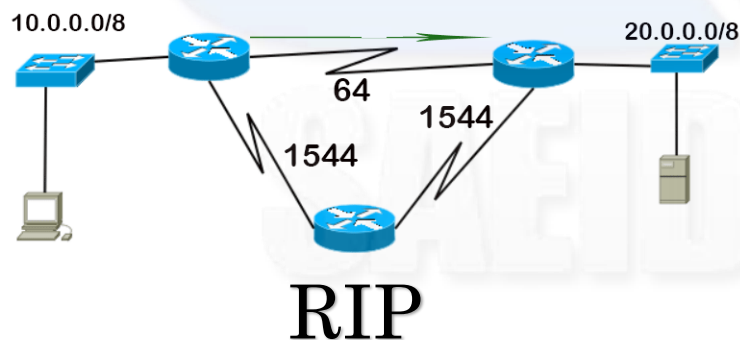
Router A ↔ Router B ↔ Router C

معیارهای لینکها:

لینک بین A و B: پهنای باند 64 گیگابیت/ثانیه، تأخیر 10 میلی ثانیه.

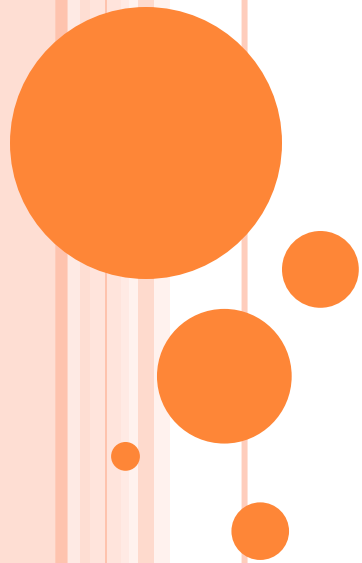
لینک بین B و C: پهنای باند 1544 مگابیت/ثانیه، تأخیر 10 میلی ثانیه.

لینک بین A و C: پهنای باند 1544 مگابیت/ثانیه، تأخیر 10 میلی ثانیه.



Link State OSPF

Open Short Path First

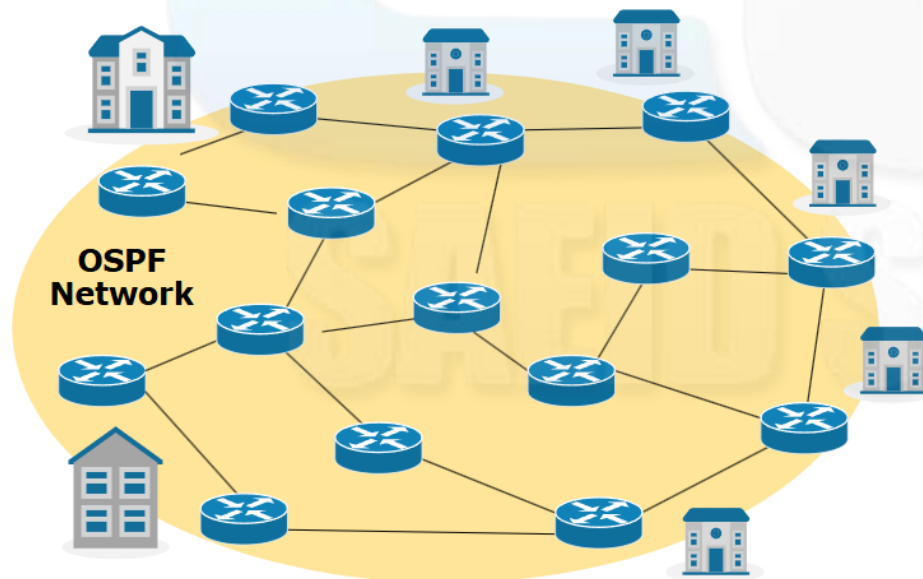


پروتکل (Open Shortest Path First) OSPF :

یک پروتکل مسیریابی Link State که از معیاری به نام هزینه Cost استفاده می کند.

این معیار بر اساس سرعت لینک بین دو روتر محاسبه می شود.

OSPF به دلیل ویژگی هایی مانند مقیاس پذیری، همگرایی سریع و سازگاری با محصولات مختلف تولید کنندگان، یکی از پروتکل های IGP محبوب است.



هزینه لینک Link Cost:

یکی از ویژگی‌های Link-State یک لینک در پروتکل OSPF، هزینه Cost نامیده می‌شود.

این مقدار عددی به هر لینک بین روترها اختصاص داده می‌شود و نشان‌دهنده "هزینه" یا "فاصله" ارسال بسته‌ها از طریق آن لینک است.

الگوریتم SPF از مقدار هزینه لینک برای محاسبه کوتاه‌ترین و کارآمدترین مسیر در شبکه استفاده می‌کند.

لینک‌هایی با مقادیر هزینه کمتر، ارجحیت بیشتری دارند و این موضوع به اتخاذ تصمیمات بهینه مسیریابی منجر می‌شود.

$$\text{Cost} = \frac{\text{Reference Bandwidth}}{\text{Link Bandwidth}}$$

اجزای فرمول:

Reference Bandwidth پهنای باند مرجع:

یک مقدار ثابت که به عنوان مرجع برای محاسبه هزینه استفاده می‌شود. معمولاً این مقدار برابر با ۱۰۰ Mbps یا ۱ Gbps است، اما می‌تواند توسط مدیر شبکه تغییر کند.

Link Bandwidth پهنای باند لینک:

پهنای باند واقعی لینک بین دو روتر که معمولاً بر حسب Mbps یا Gbps اندازه‌گیری می‌شود.

$$\text{Cost} = \frac{\text{Reference Bandwidth}}{\text{Link Bandwidth}}$$

اجزای فرمول:

Reference Bandwidth پهنای باند مرجع:

یک مقدار ثابت که به عنوان مرجع برای محاسبه هزینه استفاده می‌شود. معمولاً این مقدار برابر با ۱۰۰ Mbps یا ۱ Gbps است، اما می‌تواند توسط مدیر شبکه تغییر کند.

Link Bandwidth پهنای باند لینک:

پهنای باند واقعی لینک بین دو روتر که معمولاً بر حسب Mbps یا Gbps اندازه‌گیری می‌شود.

$$\text{Cost} = \frac{\text{Reference Bandwidth}}{\text{Link Bandwidth}}$$

نحوه محاسبه:

هزینه با تقسیم پهنای باند مرجع بر پهنای باند لینک محاسبه می‌شود.

نتیجه: لینک‌هایی با پهنای باند بیشتر (سرعت بالاتر) هزینه کمتری خواهند داشت، که نشان‌دهنده اولویت آن‌ها برای مسیریابی است.

مثال:

- اگر Reference Bandwidth برابر با 100 Mbps باشد و پهنای باند لینک برابر با 10 Mbps:

$$10 = \frac{100}{10} = Cost$$

- برای لینک 100 Mbps:

$$1 = \frac{100}{100} = Cost$$

این روش تضمین می‌کند که OSPF مسیرهایی با پهنای باند بیشتر را ترجیح دهد و مسیریابی بهینه انجام شود.

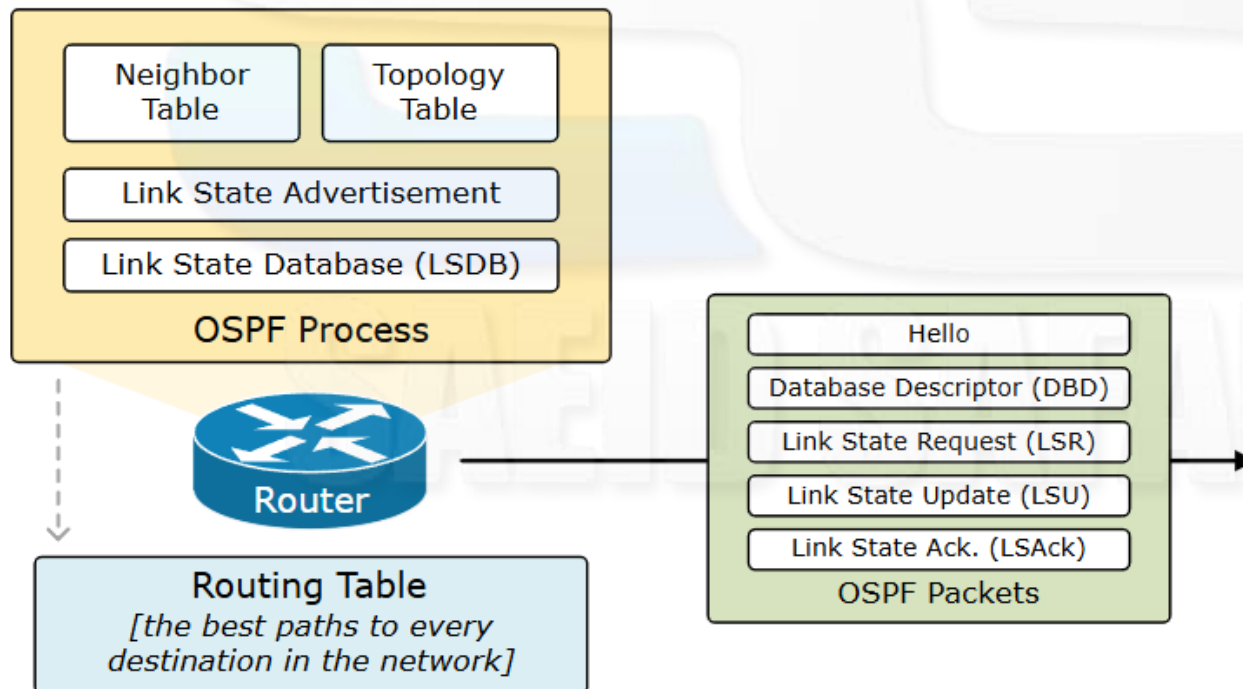
$$Cost = \frac{Reference\ Bandwidth}{Link\ Bandwidth}$$

OSPF فرآیند

از دیدگاه یک روتر، OSPF یک فرآیند است که همراه با سایر فرآیندها روی دستگاه اجرا می‌شود.

این فرآیند از مجموعه‌ای از جداول، پایگاه داده‌ها و پیام‌ها برای تبادل اطلاعات با سایر روترها استفاده می‌کند.

هدف نهایی این فرآیند مسیریابی، به‌روزرسانی جدول مسیریابی دستگاه با بهترین مسیرها است



فرآیند OSPF

به طور پیش فرض، فرآیند OSPF در تمامی دستگاه‌های Cisco IOS و IOS-XE غیرفعال است.

برای فعال‌سازی این فرآیند، از یک دستور در حالت پیکربندی کلی استفاده کرده و یک Process ID اختصاص می‌دهیم.

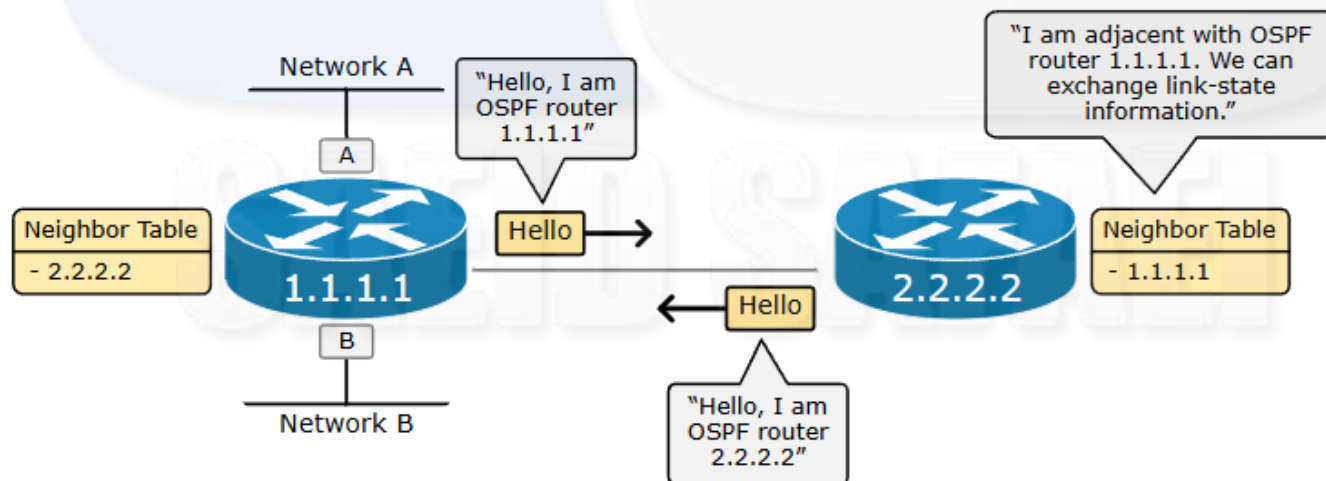
Process ID نیازی به مطابقت بین همه روترهای داخل شبکه ندارد. این شناسه یک شناسه محلی است که برای شناسایی نمونه خاصی از فرآیند مسیریابی OSPF که روی آن روتر اجرا می‌شود، استفاده می‌شود. این شناسه کاملاً محلی است و هیچ تأثیری خارج از روتر ندارد.



ایجاد همسایگی : Becoming Neighbors

زمانی که فرآیند OSPF روی یک روتر با موفقیت یک Router ID انتخاب کند، شروع به ارسال پیام‌های Hello روی تمامی اینترفیس‌های فعال OSPF می‌کند.

وقتی دو روتر روی یک VLAN، لینک سریال یا لینک WAN قرار داشته باشند، پیام‌های Hello یکدیگر را دریافت کرده و به همسایگان OSPF تبدیل می‌شوند، همان‌طور که در تصویر زیر نشان داده شده است.



در پروتکل OSPF، Hello Interval و Dead Interval دو پارامتر مهم برای برقراری و حفظ همسایگی بین روترها هستند:

۱. Hello Interval

تعریف: مدت زمان بین ارسال هر پیام Hello توسط روتر.
کاربرد: این پیام‌ها برای شناسایی و تأیید زنده بودن همسایه‌ها استفاده می‌شوند.
مقدار پیش فرض:

روی لینک‌های Point-to-Point و Broadcast : معمولاً ۱۰ ثانیه.
روی لینک‌های Non-Broadcast : معمولاً ۳۰ ثانیه.

۲- Dead Interval

تعریف: مدت زمانی که اگر طی آن هیچ پیام Hello از یک روتر دریافت نشود، همسایه مرده فرض می شود.

کاربرد: برای تشخیص قطع ارتباط یا خاموش شدن یک همسایه.

مقدار پیش فرض: معمولاً ۴ برابر Hello Interval است
نکته:

هر دو پارامتر باید بین روترهای همسایه یکسان باشند؛ در غیر این صورت، همسایگی OSPF برقرار نمی شود.

مسیر یابی

فرض کنید دو نفر در یک اتاق هستند و هر ۱۰ ثانیه Hello Interval با گفتن "سلام، من اینجا هستم!" به هم اطلاع می‌دهند که هنوز حضور دارند. اگر یکی از آنها برای ۴۰ ثانیه Dead Interval هیچ "سلامی" از دیگری نشنود، فرض می‌کند که فرد دیگر اتاق را ترک کرده یا اتفاقی افتاده است. در OSPF هم دقیقاً همین کار انجام می‌شود:

Hello Interval: زمانی است که روترها به طور منظم پیام Hello برای هم ارسال می‌کنند.

Dead Interval: اگر روتر در این مدت پاسخی دریافت نکند، فرض می‌کند که ارتباط با همسایه قطع شده است.

ایجاد همسایگی : Becoming Neighbors

مدل همسایگی دارای دو عملکرد اصلی است:

۱- کشف خودکار روترهای OSPF جدید:

این مدل به روترها اجازه می‌دهد تا به‌طور خودکار روترهای OSPF جدید را در یک بخش مشترک کشف کنند، بدون اینکه نیاز به پیکربندی دستی از سوی مدیر شبکه باشد.

۲- مبادله اطلاعات وضعیت لینک:

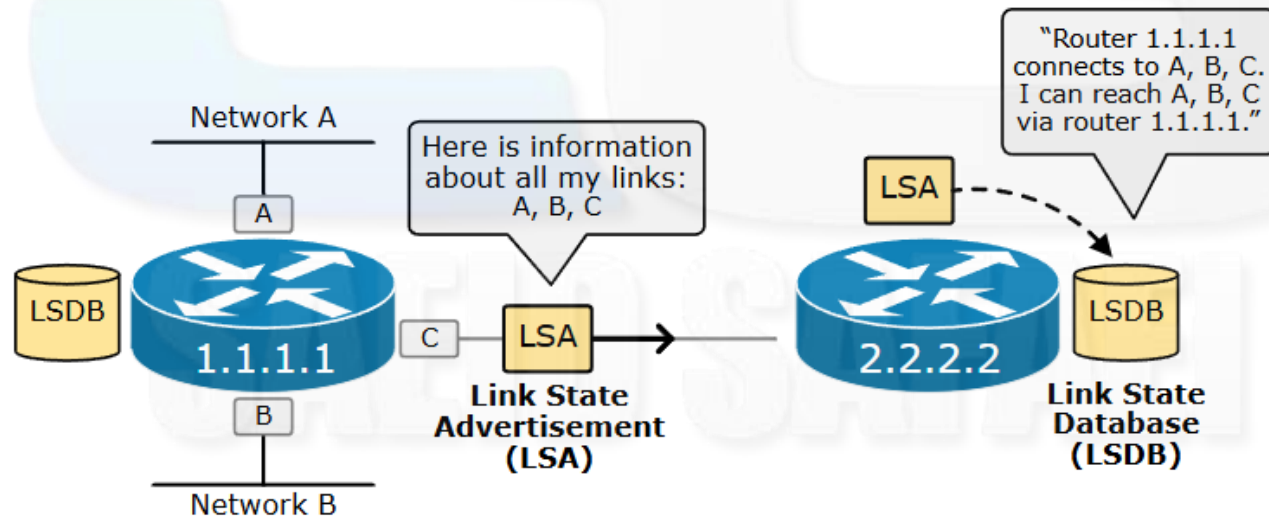
این مدل به روترها این امکان را می‌دهد که اطلاعات وضعیت لینک شبکه خود را با یکدیگر مبادله کنند.

مبادله اطلاعات پایگاه داده

زمانی که دو روتر همسایه شدند، Link-State Database را با هم تبادل می کنند.

۱- لینک یعنی هر اینترفیس روتر، مانند A یا B یا C.

۲- وضعیت لینک ویژگی های هر اینترفیس مانند آدرس IP، نوع اینترفیس و ارتباط با روترهای همسایه را توصیف می کند.



به عنوان مثال، روتر با آدرس ۱.۱.۱.۱ سه لینک دارد: A، B و C. با استفاده از یک Link-State Advertisement (LSA)، روتر اطلاعات وضعیت لینک های A، B و C را با روتر همسایه اش به آدرس ۲.۲.۲.۲ به اشتراک می گذارد.

پخش LSA (LSA Flooding) :

در فرآیند مبادله پایگاه داده در مقیاس بزرگ تر، می توانیم نحوه عملکرد فرآیند Link-State Advertisement (LSA) را مشاهده کنیم.

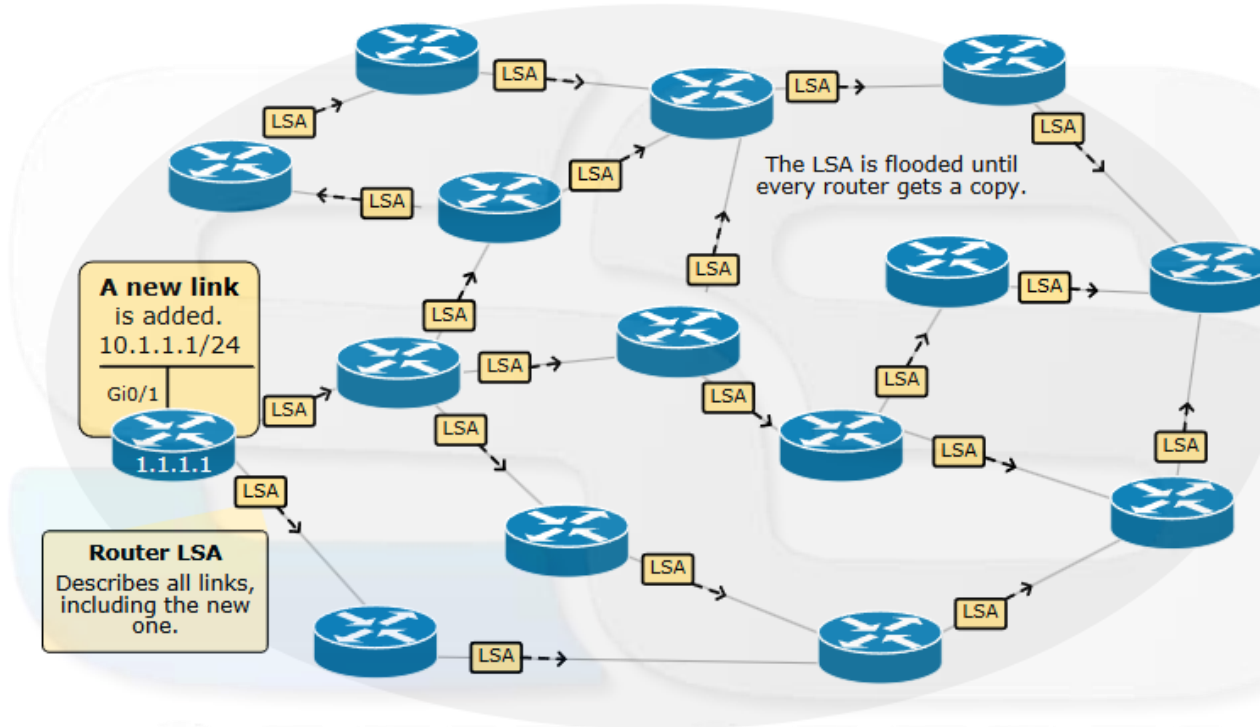
فرض کنید یک لینک جدید به روتر ۱.۱.۱.۱ اضافه کرده ایم، همان طور که در نمودار نشان داده شده است.

برای اطمینان از اینکه همه روترها از وجود یک لینک جدید با آدرس ۱۰.۱.۱.۱ که به شبکه ۱۰.۱.۱.۰/۲۴ متصل است مطلع شوند، روتر ۱.۱.۱.۱ یک به روزرسانی LSA جدید به تمام همسایگان خود ارسال می کند.

هر روتر سپس این LSA را به تمام همسایگان دیگر ارسال می کند تا زمانی که همه روترها یک نسخه از LSA را دریافت کنند.

این فرآیند پخش LSA (LSA Flooding) نامیده می شود.

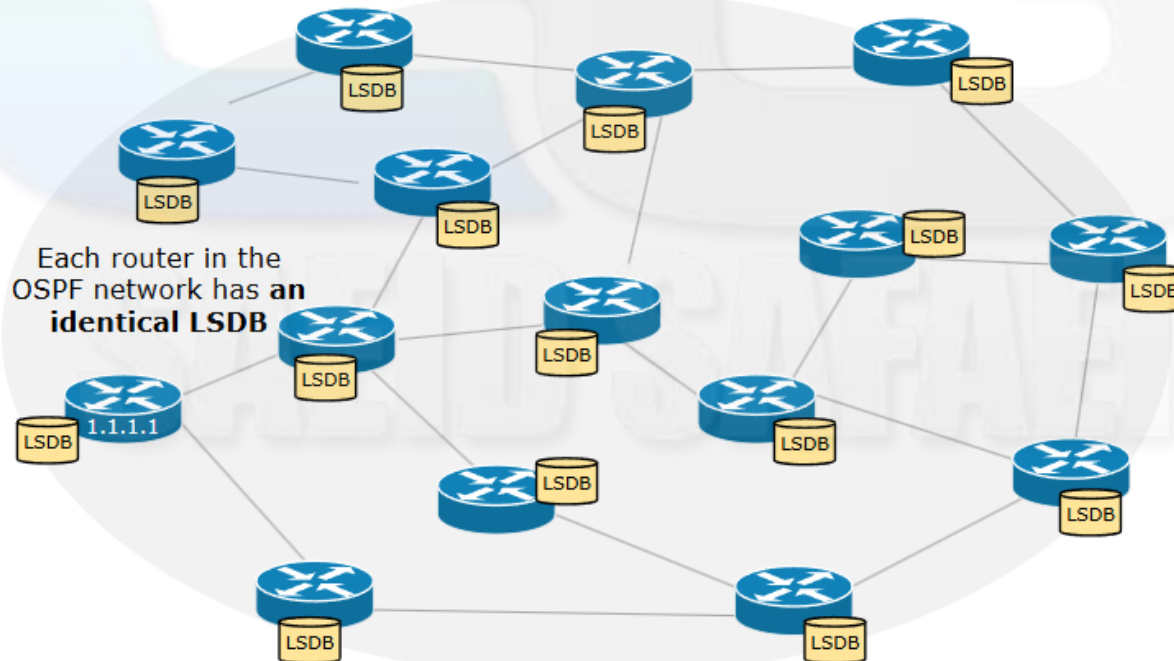
پخش LSA (LSA Flooding):



البته، مکانیزمی برای جلوگیری از حلقه وجود دارد که اطمینان می‌دهد یک به‌روزرسانی LSA به طور بی‌پایان در شبکه گردش نکند.

پایگاه داده LS (LSDB) :

هر روتر پیام‌های LSA دریافتی را در پایگاه داده‌ای به نام Link-State Database (LSDB) ذخیره می‌کند. زمانی که یک روتر یک LSA جدید دریافت می‌کند، پایگاه داده وضعیت لینک LSDB خود را به‌روزرسانی کرده و LSA را به همسایگان خود ارسال می‌کند. این فرآیند تضمین می‌کند که تمامی روترها در شبکه یک LSDB یکسان داشته باشند.



الگوریتم SPF

در پروتکل OSPF، پیام های LSA به تنهایی اطلاعاتی برای اضافه کردن مستقیم به جدول مسیریابی ندارند. هر LSA فقط بخشی از اطلاعات شبکه را نشان می دهد.

مجموعه این پیام ها یک پایگاه داده به نام LSDB را تشکیل می دهد که نقشه کامل شبکه را ذخیره می کند. اما LSDB خودش مسیرها را مشخص نمی کند، بلکه فقط اطلاعات خام را در خود دارد.

پس:

برای اینکه روتر بتواند از اطلاعات موجود در LSDB مسیرهای بهینه را پیدا کند، نیاز به پردازش توسط الگوریتم SPF (Shortest Path First) دارد.

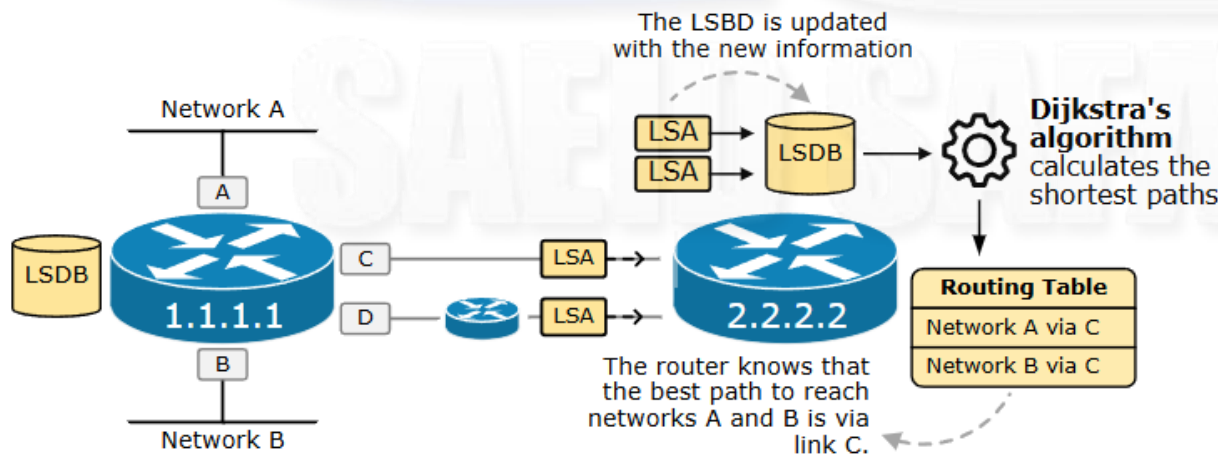
الگوریتم SPF

SPF از اطلاعات ذخیره شده در LSDB استفاده می کند تا توپولوژی کامل شبکه را بازسازی کند.

با استفاده از این توپولوژی، مسیر کوتاه ترین و بهینه ترین راهها را به هر مقصد محاسبه می کند.

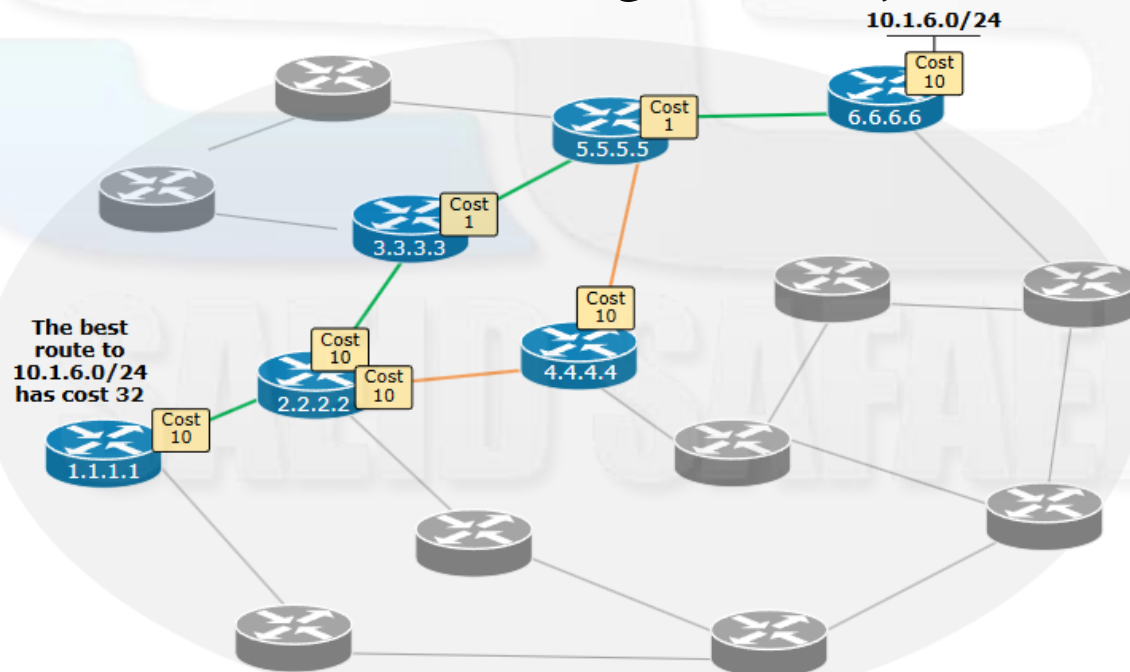
سپس این مسیرها را در جدول مسیریابی روتر ثبت می کند.

به زبان ساده: LSDB نقشه خام شبکه است، و الگوریتم SPF مثل یک مسیریاب عمل می کند که بهترین مسیرها را از روی این نقشه پیدا کرده و مشخص می کند.



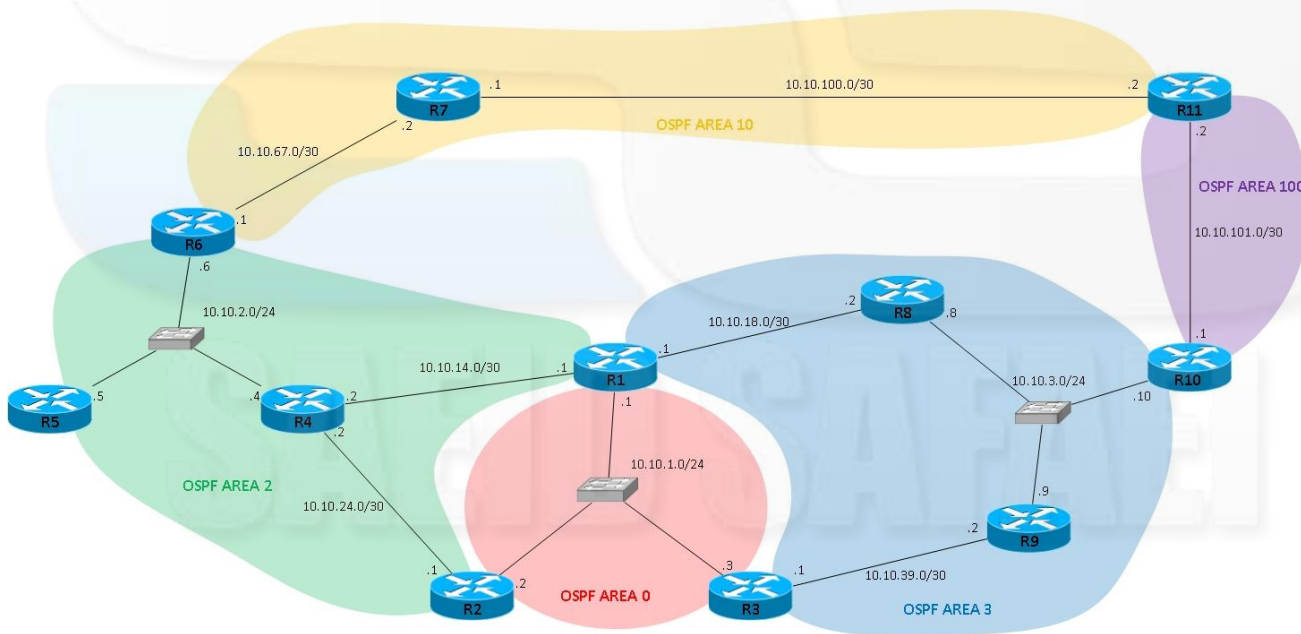
الگوریتم SPF

مهم است که بدانید LSDB روی همه دستگاهها یکسان است، اما الگوریتم SPF که به آن الگوریتم دایجسترا هم میگویند از این اطلاعات برای محاسبه بهترین مسیرها استفاده می کند. برای مثال، هر روتر بهترین مسیرهای متفاوتی به شبکه A دارد، چون هر کدام از روترها مسیرهای بهینه را بر اساس وضعیت شبکه از دید خود محاسبه می کنند.



الگوریتم SPF

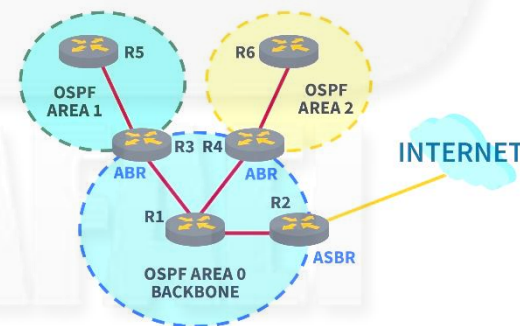
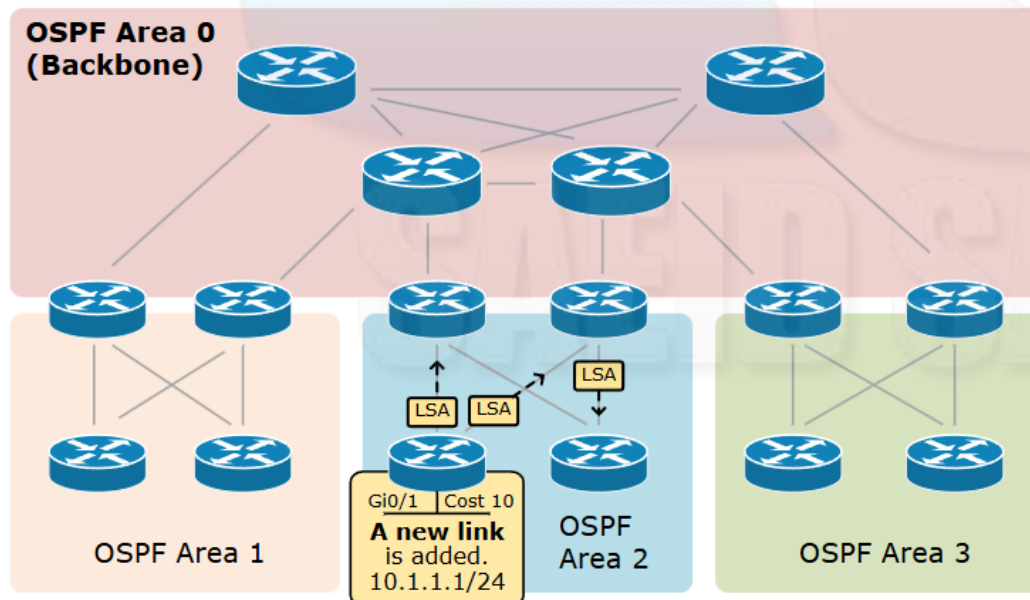
الگوریتم SPF از هزینه‌های لینک‌ها در LSDB برای ساخت یک درخت کوتاه‌ترین مسیر استفاده می‌کند. این الگوریتم هزینه لینک‌ها را در طول مسیرهای مختلف جمع کرده و مسیری را انتخاب می‌کند که کمترین هزینه کل را داشته باشد.



مناطق Areas

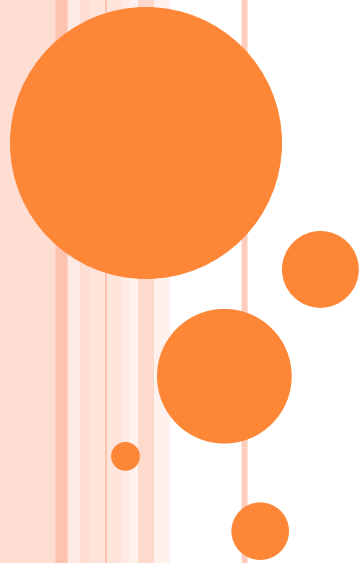
برای مقیاس‌پذیری بهتر، OSPF مفهوم منطقه Area را پیاده‌سازی می‌کند. هر منطقه بخشی از شبکه است که در آن روترها اطلاعات مسیریابی را با یکدیگر تبادل می‌کنند.

این کار به مقیاس‌پذیری شبکه‌های بزرگ کمک می‌کند، زیرا شبکه به بخش‌های کوچکتر تقسیم می‌شود و بدین ترتیب مقدار اطلاعات مسیریابی که هر دستگاه باید پردازش و ذخیره کند کاهش می‌یابد.



Link State IS-IS

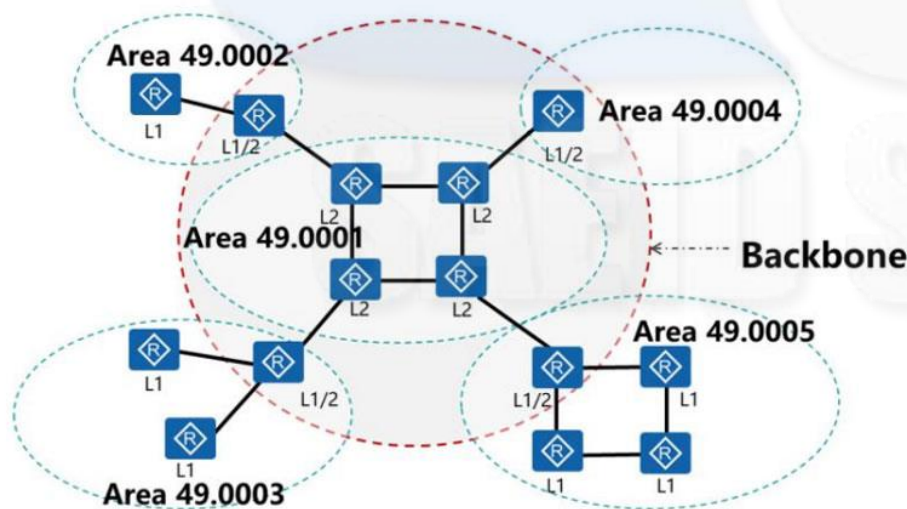
Intermediate System
to
Intermediate System



پروتکل IS-IS (Intermediate System to Intermediate System)

IS-IS یک پروتکل مسیریابی داخلی IGP و Link-State است که عملکردی مشابه با OSPF دارد.

این پروتکل ارتباط همسایگی برقرار می‌کند، از نواحی Areas استفاده می‌کند، بسته‌های Link-State Packets LSP را مبادله می‌کند، یک پایگاه داده Link-State Database LSDB می‌سازد و با استفاده از الگوریتم Dijkstra کوتاه‌ترین مسیر SPF را برای هر مقصد محاسبه کرده و در جدول مسیریابی نصب می‌کند.



پروتکل IS-IS (Intermediate System to Intermediate System)

زمانی که OSPF و IS-IS توسعه داده شدند، پروتکل IP هنوز به عنوان پروتکل غالب شناخته نمی شد.

در آن زمان، سازمان ISO چیزی مشابه با IP و UDP به نام

CLNP (Connectionless-mode Network Protocol)

و

CLNS (Connectionless-mode Network Service)

طراحی کرده بود. ISO همچنین از برخی اصطلاحات متفاوت استفاده می کند. برای مثال:

روتر (Router) معادل Intermediate System است.

میزبان (Host) معادل End System است.

معماری و نحوه عملکرد

- روترها در IS-IS به دو نوع تقسیم می‌شوند:

۱- **Intermediate System (IS)** : وظیفه مسیریابی دارد.

۲- **End System (ES)** : سیستم انتهایی (مانند کلاینت‌ها) است که به مسیریابی نیازی ندارد.

- پایگاه داده **LSDB** : مانند OSPF، هر روتر IS-IS اطلاعات وضعیت لینک‌های خود را به تمام روترهای دیگر در ناحیه ارسال می‌کند. این اطلاعات در یک LSDB ذخیره شده و با استفاده از الگوریتم Dijkstra کوتاه‌ترین مسیرها محاسبه می‌شوند.

- پروتکل‌های حمل و نقل: پیام‌های مسیریابی IS-IS از طریق پروتکل لایه پیوند **Layer 2** ارسال می‌شوند و نیازی به حمل و نقل با پروتکل‌های IP ندارد. این ویژگی موجب کاهش پیچیدگی و سربار در پروتکل می‌شود.

چرا IS-IS برای سرویس‌دهندگان مناسب است؟

پروتکل IS-IS بیشتر در سطح شبکه‌های بزرگ و پیچیده مانند شبکه‌های سرویس‌دهندگان اینترنتی Service Providers استفاده می‌شود.

دلیل این انتخاب، ویژگی‌های مقیاس‌پذیری بالا، پایداری و توانایی آن در مدیریت تعداد زیادی از روترها و مسیرها است.

۱- قابلیت مقیاس‌پذیری: می‌تواند تعداد زیادی از نواحی Areas و روترها را مدیریت کند.

۲- سادگی در طراحی: IS-IS بر روی لایه ۲ Data Link Layer کار می‌کند و نیازی به وابستگی به IP ندارد، که باعث کاهش پیچیدگی می‌شود.

۳- پشتیبانی از IPv6: به طور طبیعی می‌تواند IPv4 و IPv6 را مدیریت کند، که برای شبکه‌های امروزی ضروری است. نیازی به مهاجرت یا تغییر پروتکل نیست.

چرا IS-IS برای سرویس‌دهندگان مناسب است؟

در پروتکل IS-IS، بیشتر از **Traffic Engineering** استفاده می‌شود تا به بهینه‌سازی مسیرها و منابع شبکه پرداخته شود.

IS-IS می‌تواند با استفاده از روش‌هایی مانند MPLS و RSVP-TE برای هدایت ترافیک از مسیرهای خاص، به مهندسی کمک کند.

این پروتکل برای مدیریت بهتر جریان‌های ترافیکی در شبکه‌های پیچیده و بزرگ طراحی شده است.

Load Balancing هم در IS-IS به صورت غیرمستقیم می‌تواند انجام شود، اما بیشتر به کنترل و توزیع بار در لایه‌های بالاتر (مانند دروازه‌ها یا روترها) مربوط است تا خود پروتکل IS-IS.

تفاوت Traffic Engineering با Load Balancing:

Traffic Engineering : بهینه‌سازی استفاده از منابع شبکه برای بهبود کارایی کلی. مهندسی ترافیک به دنبال هدایت جریان‌های داده از طریق مسیرهای خاص است تا از ازدحام جلوگیری کرده و از ظرفیت شبکه به طور بهینه استفاده کند.

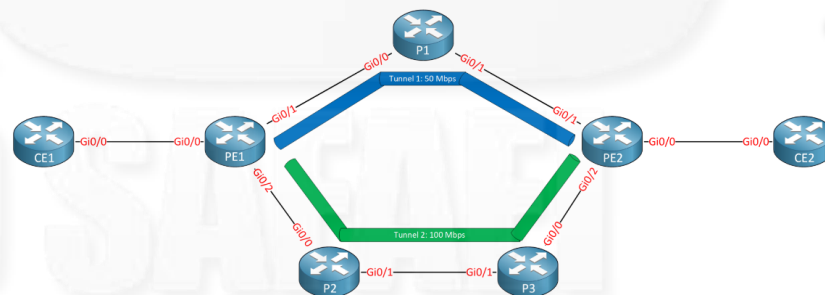
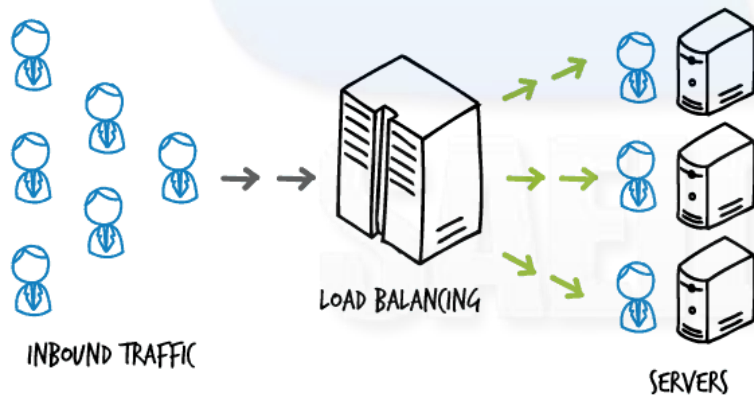
Load Balancing : توزیع متعادل بار ترافیکی بین منابع مختلف برای جلوگیری از بار زیاد روی یک منبع خاص.

ویژگی	Traffic Engineering	Load Balancing
هدف	استفاده بهینه از منابع شبکه	توزیع متعادل بار بین منابع
سطح کاربرد	کل شبکه (لینک‌ها و مسیرها)	منابع خاص (سرورها یا لینک‌ها)
مقیاس‌پذیری	در شبکه‌های بزرگ و پیچیده	بیشتر در شبکه‌های کوچک‌تر یا مراکز داده
ابزارها و پروتکل‌ها	الگوریتم‌های مسیریابی، MPLS، RSVP-TE	ها، الگوریتم‌های توزیع بار، Load Balancer
تمرکز	جلوگیری از گلوگاه در کل شبکه	جلوگیری از بار زیاد روی یک منبع خاص

تفاوت Traffic Engineering با Load Balancing:

Traffic Engineering : فرض کنید در یک شهر، چندین جاده وجود دارد. مهندسی ترافیک شبیه به مدیریت شهری است که مسیرهای خاصی برای خودروها تعیین می‌کند تا جاده‌های شلوغ آزاد شوند.

Load Balancing : شبیه به قرار دادن چندین صندوق در یک بانک است که مشتری‌ها بین آن‌ها توزیع می‌شوند تا از ازدحام در یک صندوق جلوگیری شود.



تفاوت Traffic Engineering با Load Balancing:

Traffic Engineering: در شبکه‌های بزرگ (مانند ISP ها)، هدف از مهندسی ترافیک هدایت جریان‌های داده به مسیرهای خاص است تا از ازدحام و ترافیک زیاد در یک لینک خاص جلوگیری شود. مثلا IS-IS، ترافیک به‌طور هوشمند به لینک‌های با ظرفیت‌های مختلف هدایت می‌شود.

Load Balancing: در یک محیط با چندین منبع (مثل سرورها)، هدف از توازن بار توزیع متعادل ترافیک است تا از بار زیاد روی یک منبع خاص جلوگیری شود. مثلا Load Balancing، ترافیک بین سرورها توزیع می‌شود تا همه سرورها به طور یکنواخت بار داشته باشند.

تفاوت OSPF با ISIS:

۱. پروتکل و معماری:

OSPF (Open Shortest Path First) یک پروتکل مسیریابی مبتنی بر Link-State است که در لایه ۳ Network Layer مدل OSI کار می‌کند و به طور گسترده در شبکه‌های IP استفاده می‌شود.

IS-IS (Intermediate System to Intermediate System) نیز یک پروتکل مسیریابی حالت-پیوندی است، اما در لایه ۲ Data Link کار می‌کند و برای شبکه‌های OSI طراحی شده است، ولی به راحتی می‌تواند با شبکه‌های IP نیز کار کند.

IS-IS برای مسیریابی در سطح لایه ۲ و تبادل اطلاعات وضعیت لینک‌ها طراحی شده است: به جای ارسال آدرس‌دهی IP، این پروتکل از آدرس‌های MAC برای شناسایی روترها استفاده می‌کند.

تفاوت OSPF با ISIS:

۲. سازگاری و استفاده:

OSPF بیشتر برای شبکه‌های IP و به ویژه شبکه‌های داخلی IGP طراحی شده است و به صورت گسترده در شبکه‌های IP استفاده می‌شود. IS-IS برای استفاده در شبکه‌های OSI طراحی شده بود ولی به سرعت در شبکه‌های IP نیز محبوب شد، به ویژه در سرویس‌دهندگان اینترنت ISP و شبکه‌های بزرگ، به دلیل مقیاس‌پذیری بهتر.

تفاوت OSPF با ISIS:

۳. مدیریت نواحی Areas :

OSPF از نواحی برای تقسیم‌بندی شبکه استفاده می‌کند و شبکه را به مناطق مختلف تقسیم می‌کند تا پیچیدگی شبکه کاهش یابد.

IS-IS نیز از نواحی استفاده می‌کند، ولی ساختار آن در این زمینه ساده‌تر و انعطاف‌پذیرتر است.

تفاوت OSPF با ISIS:

۴. ساختار پیام‌ها و پروتکل:

OSPF پیام‌ها و اطلاعات را در قالب LSA (Link-State Advertisements) ارسال می‌کند.

IS-IS از LSP (Link-State PDU) برای ارسال اطلاعات شبکه استفاده می‌کند.

۵. پشتیبانی از پروتکل‌ها:

OSPF فقط برای پروتکل IP طراحی شده است.

IS-IS می‌تواند هم برای IP و هم برای CLNP (Connectionless Network Protocol) استفاده شود و از پروتکل‌های مختلف پشتیبانی می‌کند.

تفاوت OSPF با ISIS:

۶. سازگاری و مقیاس پذیری:

IS-IS به دلیل طراحی ساده تر و استفاده از لایه ۲، مقیاس پذیری بهتری در شبکه های بزرگ (مثل ISP ها) دارد.

OSPF پیچیدگی بیشتری در مدیریت نواحی و مقیاس پذیری دارد، اما در شبکه های متوسط و کوچک محبوب تر است.

۷. رویکرد به همسایگی:

در OSPF، روترها باید در یک شبکه مشابه شبکه های broadcast یا point-to-point باشند تا ارتباط همسایگی برقرار شود.

IS-IS از یک روش مشابه ولی با انعطاف پذیری بیشتر استفاده می کند و می تواند با شبکه های مختلف کار کند.

تفاوت OSPF با ISIS:

خلاصه:

OSPF برای شبکه‌های IP و روترهای کوچک و متوسط مناسب است و در شبکه‌های مقیاس پذیر استفاده می‌شود. در حد سازمانی.

IS-IS به طور عمده در شبکه‌های بزرگ و ISP ها استفاده می‌شود و به دلیل پشتیبانی از لایه ۲ و مقیاس پذیری بالا در شبکه‌های پیچیده تر مناسب تر است. در حد ارائه دهنده سرویس Service Provider

تفاوت OSPF با ISIS:

دلیل استفاده از LSP:

OSPF که در لایه ۳ کار می‌کند از LSA برای ارسال اطلاعات وضعیت لینک‌ها استفاده می‌کند، زیرا این پروتکل به‌طور خاص برای شبکه‌های IP طراحی شده است.

IS-IS که در لایه ۲ کار می‌کند از LSP برای ارسال وضعیت لینک‌ها استفاده می‌کند، زیرا این پروتکل به‌طور اصلی برای شبکه‌های OSI طراحی شده بود. در این شبکه‌ها، LSP برای ارسال اطلاعات وضعیت لینک‌ها بهینه‌تر و سازگارتر است.

LSP در IS-IS ساختاری ساده‌تر دارد و اطلاعات وضعیت لینک‌ها را در قالب یک واحد داده PDU ارسال می‌کند.

Ethernet Header

IS-IS Header

IS-IS Data

FCS

تفاوت OSPF با ISIS:

دلیل استفاده از LSP:

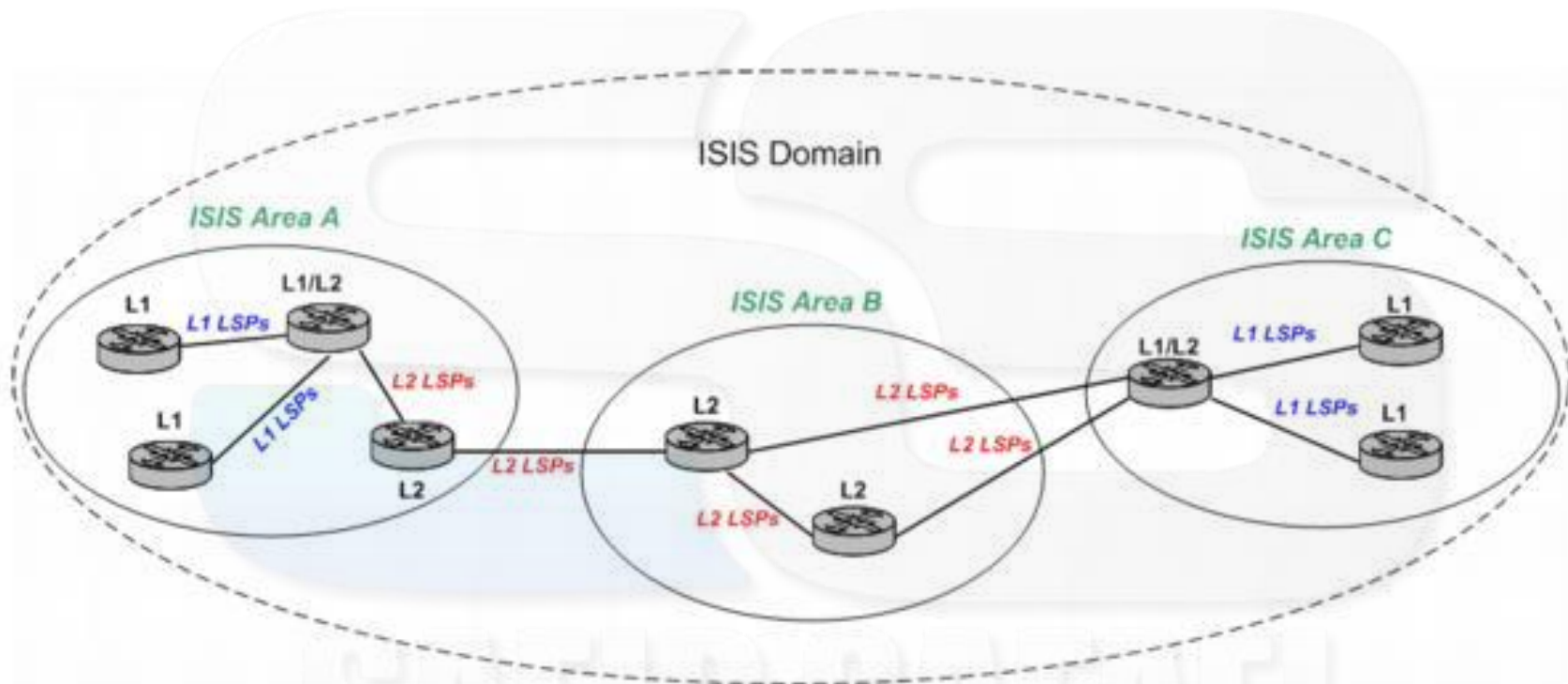
در IS-IS از پکت‌هایی به نام LSP (Link-State PDU) به جای پکت‌های IP استفاده می‌شود. (چون در لایه ۲ است)

در واقع، LSP اطلاعات وضعیت لینک‌ها را در شبکه به اشتراک می‌گذارد و این پکت‌ها برای به‌روزرسانی پایگاه داده وضعیت لینک Link-State Database در روترها و محاسبه بهترین مسیرها استفاده می‌شوند.

در IS-IS، برخلاف IP که آدرس‌دهی به مقصدها را مدیریت می‌کند، LSP نقش انتقال اطلاعات وضعیت لینک‌ها بین روترها را بر عهده دارد. این اطلاعات به روترها کمک می‌کند تا مسیرهای بهینه را برای انتقال داده‌ها پیدا کنند.

بنابراین، IS-IS از LSP برای برقراری ارتباط بین روترها و انجام مسیریابی استفاده می‌کند، نه از IP که مخصوص مسیریابی در شبکه‌های مبتنی بر پروتکل اینترنت است.

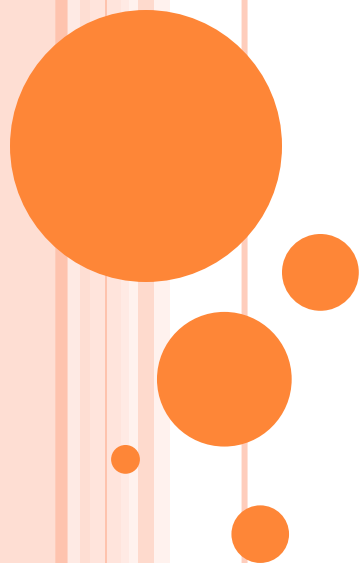
.ISIS



IGP (Hybrid)

EIGRP

Enhanced Interior Gateway
Routing Protocol



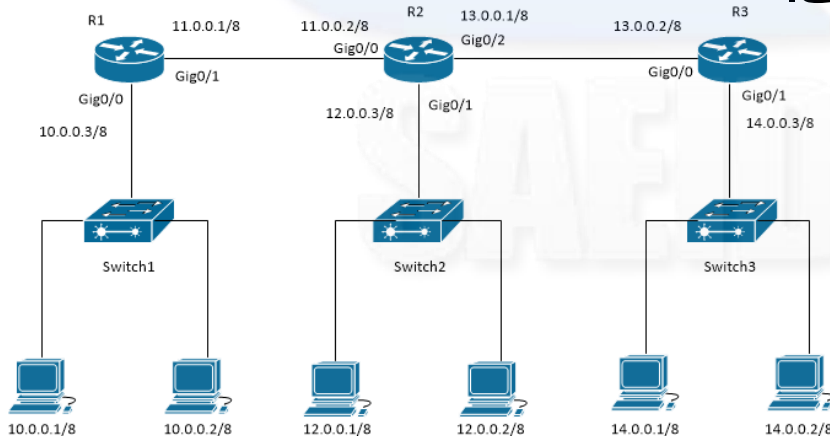
پروتکل EIGRP (Enhanced Interior Gateway Routing Protocol)

یک پروتکل اختصاصی شرکت Cisco است.

این پروتکل در شبکه‌های مخصوص سیسکو بسیار محبوب است، اما در محیط‌هایی با تجهیزات تولیدکنندگان مختلف کمتر مورد استفاده قرار می‌گیرد.

مشابه OSPF، EIGRP نیز یک پروتکل IGP است که دارای زمان همگرایی سریع و مقیاس‌پذیری بالاست.

طبقه‌بندی این پروتکل به عنوان یک پروتکل مسیریابی بردار فاصله (DV) یا حالت پیوند (LS) کمی چالش‌برانگیز است.



پروتکل EIGRP (Enhanced Interior Gateway Routing Protocol)

به طور پیش فرض، EIGRP در محاسبه متریک خود از پهنای باند و تأخیر Delay استفاده می کند، اما پارامترهای دیگری نیز می توانند در نظر گرفته شوند.

این پارامترهای اختیاری شامل قابلیت اطمینان Reliability، بار Load و اندازه واحد انتقال حداکثر MTU می شوند.

با استفاده از تأخیر به عنوان بخشی از متریک، EIGRP می تواند تأثیر زمان تأخیر ناشی از کندترین لینک ها در مسیر را در نظر بگیرد.

این جمله به این معناست که پروتکل EIGRP در محاسبه متریک (معیار انتخاب بهترین مسیر) خود، از تأخیر Delay به عنوان یکی از عوامل کلیدی استفاده می کند.

پروتکل EIGRP (Enhanced Interior Gateway Routing Protocol)

این ویژگی به EIGRP اجازه می‌دهد تا مسیرهایی را انتخاب کند که کمترین تأخیر را داشته باشند، حتی اگر یک یا چند لینک در مسیر نسبتاً کند باشند.

به عبارت دیگر:

وقتی یک مسیر شامل لینک‌های با سرعت‌های متفاوت است (مثلاً لینک‌های کند و سریع)، EIGRP می‌تواند تأثیر کندترین لینک را بر زمان کل تأخیر در نظر بگیرد.

این قابلیت کمک می‌کند تا مسیری که عملکرد بهتری از نظر زمان انتقال داده دارد، انتخاب شود، نه صرفاً مسیری با تعداد هاپ کمتر یا پهنای باند بالاتر.

این انعطاف‌پذیری، یکی از مزایای EIGRP نسبت به برخی از پروتکل‌های دیگر است.

MTU (Maximum Transmission Unit) چیست؟

MTU یعنی بیشترین اندازه یک بسته اطلاعاتی که می‌تواند از یک شبکه عبور کند.

یک مثال ساده:

فرض کن می‌خواهی یک نامه بزرگ را با پست بفرستی. اگر پاکت پستی خیلی کوچک باشد، مجبور می‌شوی نامه را به چند قسمت تقسیم کنی.

این همان چیزی است که در شبکه هم اتفاق می‌افتد:

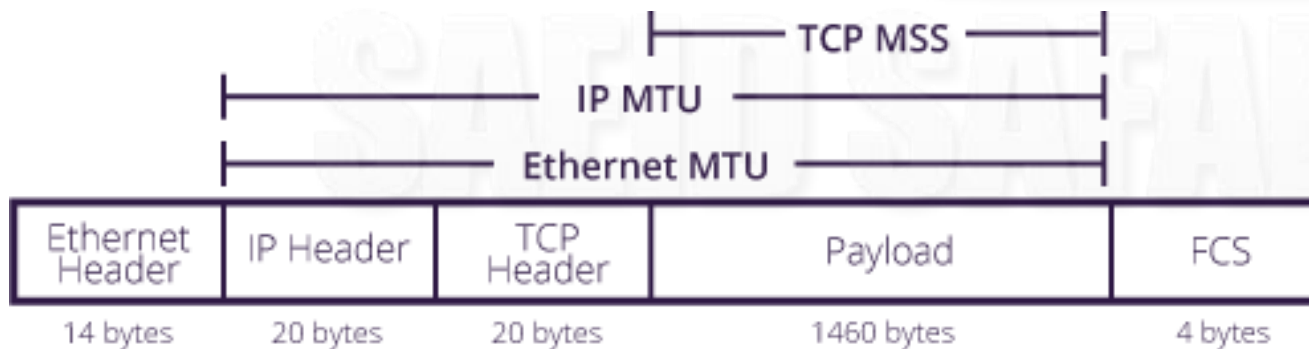
اگر اندازه داده‌های ارسالی بزرگ‌تر از MTU باشد، داده‌ها به تکه‌های کوچک‌تر تقسیم می‌شوند (که به آن Fragmentation می‌گویند). این تقسیم شدن باعث می‌شود ارسال و دریافت داده کمی کندتر شود.

MTU (Maximum Transmission Unit) چیست؟

اگر MTU خیلی کوچک باشد، بسته‌ها زیادتر می‌شوند و شبکه شلوغ می‌شود.
اگر MTU خیلی بزرگ باشد، ممکن است همه دستگاه‌ها نتوانند بسته‌های بزرگ را مدیریت کنند.

مقدار استاندارد: در شبکه‌های معمولی، اندازه MTU ۱۵۰۰ بایت است. یعنی هر بسته می‌تواند حداکثر ۱۵۰۰ بایت داده داشته باشد.

MTU مشخص می‌کند که داده‌ها در شبکه چه اندازه باشند تا به‌درستی و با کمترین مشکل ارسال شوند.



پروتکل EIGRP (Enhanced Interior Gateway Routing Protocol)

EIGRP به طور دقیق در دسته بندی پروتکل های مسیریابی بردار فاصله Distance Vector یا حالت پیوند Link State قرار نمی گیرد. اما می توان آن را به این صورت توصیف کرد:

بردار فاصله پیشرفته Advanced Distance Vector :

EIGRP مانند پروتکل های بردار فاصله از اطلاعات همسایگان خود برای تصمیم گیری در مورد بهترین مسیر استفاده می کند و مسافت و جهت Distance و Vector را در نظر می گیرد.

ویژگی های حالت پیوند:

برخلاف پروتکل های بردار فاصله کلاسیک، EIGRP یک پایگاه داده توپولوژی ایجاد می کند که شامل اطلاعات دقیق تری از کل شبکه است.

این ویژگی مشابه پروتکل های حالت پیوند است.

پروتکل EIGRP (Enhanced Interior Gateway Routing Protocol)

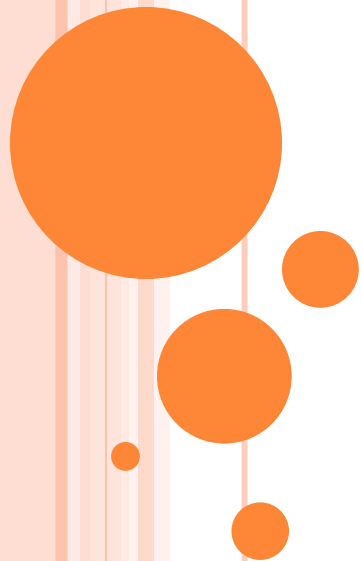
به همین دلیل، EIGRP را معمولاً پروتکل مسیریابی هیبریدی یا ترکیبی می‌نامند، زیرا ترکیبی از ویژگی‌های هر دو نوع پروتکل Distance Vector و Link State را دارد.

با این حال، به‌طور رسمی EIGRP یک پروتکل مسیریابی بردار فاصله پیشرفته محسوب می‌شود.

Exterior Gateway Protocol

BGP

Border Gateway
Protocol



پروتکل مسیریابی مرزی BGP

تنها EGP است که امروزه به طور گسترده مورد استفاده قرار می گیرد. در واقع، BGP به عنوان پروتکل مسیریابی که اینترنت را اجرا می کند، شناخته می شود؛ اینترنت یک اتصال میان چندین سیستم مستقل Autonomous Systems است.

اگرچه برخی منابع، BGP را به عنوان یک پروتکل مسیریابی بردار فاصله Distance-Vector طبقه بندی می کنند، اما توصیف دقیق تر آن، یک پروتکل مسیریابی بردار مسیر Path-Vector است.

این بدان معناست که BGP از تعداد هاپ های سیستم های مستقل AS Hops که باید برای رسیدن به یک شبکه مقصد طی شوند، به عنوان متریک خود استفاده می کند، برخلاف تعداد هاپ های روتر.

با این حال، انتخاب مسیر BGP تنها بر اساس تعداد هاپ های سیستم های مستقل نیست.

: Path Vector

Path-Vector یک روش مسیریابی است که شبیه به Distance-Vector عمل می‌کند، اما به جای فقط شمارش هاپ‌ها (مانند تعداد روترها)، مسیر دقیق عبوری را نیز ثبت و استفاده می‌کند.

توضیح ساده:

فرض کن می‌خواهی از شهر خودت به یک مقصد سفر کنی و چند مسیر مختلف وجود دارد.

در مسیریابی Distance-Vector، فقط به تو می‌گویند که "این مسیر ۳ ایستگاه دارد".

اما در مسیریابی Path-Vector، علاوه بر تعداد ایستگاه‌ها، به تو می‌گویند: مسیر این است:

مقصد $\rightarrow$ شهر C $\rightarrow$ شهر B $\rightarrow$ شهر A

: Path Vector

فرض کن می‌خواهی به خانه دوست بروی و در مسیر چند راه مختلف وجود دارد.

در روش Distance-Vector، به تو گفته می‌شود که باید ۳ مسیر مختلف را طی کنی. یعنی فقط تعداد مسیرها مهم است.

در روش Path-Vector، به تو گفته می‌شود که باید از مسیر خاصی بروی، مثلاً: "از خانه من به خیابان اصلی برو، بعد از آن به چهارراه چپ بپیچ، سپس به بزرگراه برو و از آنجا به خانه دوست برس."

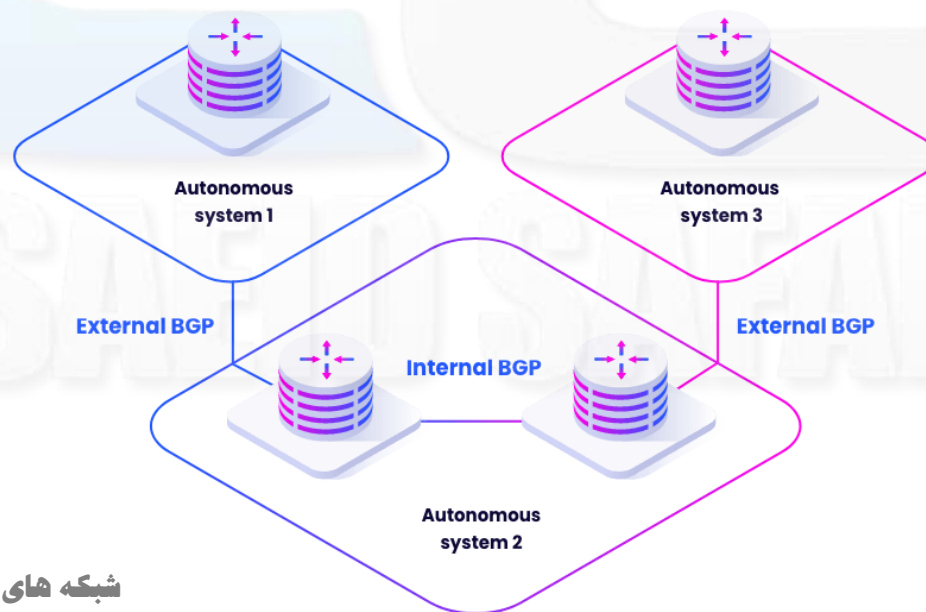
در Path-Vector، مسیر دقیق هر مرحله مشخص است، نه فقط تعداد مراحل.

: Path Vector

در BGP به همین روش عمل می‌کند.

وقتی یک روتر می‌خواهد داده‌ها را به مقصد بفرستد، به جای گفتن اینکه چند روتر را باید رد کند، مسیر دقیق و سیستم‌های مستقل AS که داده‌ها از آن‌ها عبور می‌کند، را ذخیره می‌کند و استفاده می‌کند.

خلاصه: در Path-Vector، مسیر دقیق داده‌ها مشخص است، نه فقط تعداد توقف‌ها.



مفهوم Route Redistribution :

یک شبکه می‌تواند به‌طور همزمان از چندین پروتکل مسیریابی از طریق فرآیند توزیع مجدد مسیرها Route Redistribution پشتیبانی کند.

برای مثال، یک روتر می‌تواند یکی از رابط‌های خود را در یک ناحیه OSPF از شبکه قرار دهد و رابط دیگری را در یک ناحیه EIGRP قرار دهد.

این روتر می‌تواند مسیرهایی که از طریق OSPF یاد گرفته است را دریافت کرده و آن‌ها را به فرآیند مسیریابی EIGRP وارد کند.

به همین ترتیب، مسیرهایی که از طریق EIGRP یاد گرفته شده‌اند نیز می‌توانند به فرآیند مسیریابی OSPF توزیع شوند.

توزیع مجدد مسیرها اجازه می‌دهد که مسیرهایی که توسط یک پروتکل مسیریابی یاد گرفته شده‌اند، به پروتکل مسیریابی دیگری منتقل شوند.



با تشکر از همراهی شما

محمد سعید صفایی صادق

