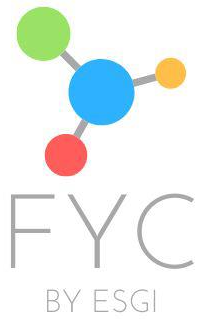


RAVIER Jules

BALLENS François

LANORE Gabin

DA SILVA Brice



FYC - Présentation des différents désign d'architecture réseaux Datacenter

Table des matières

Introduction	4
Vidéo de présentation de l'équipe (DRIVE).....	4
Vidéo de présentation de l'équipe (Youtube).....	4
Vidéo de présentation du cours (DRIVE)	4
Vidéo de présentation du cours (Youtube)	4
QCM de positionnement.....	4
Correction QCM de positionnement	5
I – Architecture traditionnelle « Core, Aggregation, Access ».....	7
A- Présentation du principe d'architecture	7
B- Avantages et inconvénients de l'architecture	9
C- Présentation des protocoles	11
Vidéo de Présentation de la maquette d'architecture (DRIVE)	11
Vidéo de Présentation de la maquette d'architecture (Youtube)	11
Télécharger le LAB	11
1. Qu'est-ce que les VLAN et les Private VLAN ?	11
QCM VLAN et PVLAN	14
Correction QCM VLAN et PVLAN.....	15
Vidéo de de mise en œuvre VLAN et PVLAN (DRIVE)	16
Vidéo de de mise en œuvre VLAN et PVLAN (Youtube)	16
Fiche TP VLAN et PVLAN	16
Télécharger Correction TP VLAN et PVLAN	20
2. Qu'est-ce qu'une ACL ?	20
QCM ACL	21
Correction QCM ACL.....	22
Vidéo de de mise en œuvre ACL (DRIVE)	23
Vidéo de de mise en œuvre ACL (YOUTUBE)	23
Fiche TP ACL	23
Télécharger Correction TP ACL.....	23
3. Qu'est-ce que PAgP et LACP ?	23
Télécharger le paquet Wireshark LACPDU	25
Télécharger le paquet Wireshark PAgPDU	28
QCM LACP et PAgP	29
Correction QCM LACP et PAgP	30
Vidéo de de mise en œuvre LACP et PAgP (DRIVE)	30

Vidéo de de mise en œuvre LACP et Pag (Youtube).....	30
Fiche TP LACP et PAgP	30
Télécharger Correction TP LACP et PAgP	31
4. Qu'est-ce que VRRP et GLBP ?	31
Télécharger le paquet Wireshark VRRP Advertisement	32
Télécharger le paquet Wireshark GLBP	34
QCM VRRP et GLBP	37
Correction QCM GLBP et VRRP	37
Vidéo de de mise en œuvre VRRP et GLBP (DRIVE)	38
Vidéo de de mise en œuvre VRRP et GLBP (Youtube)	38
Fiche TP VRRP et GLBP	38
Télécharger Correction TP VRRP et GLBP	39
5. Qu'est-ce que MSTP et RPVST+ ?.....	39
Télécharger le paquet Wireshark MSTP BPDU	44
Télécharger le paquet Wireshark RPVST+ BPDU	46
QCM MSTP et RPVST+	48
Correction QCM MSTP et RPVST+	49
Vidéo de mise en œuvre MSTP et RPVST+ (DRIVE)	49
Vidéo de mise en œuvre MSTP et RPVST+ (Youtube).....	49
Fiche TP MSTP et RPVST+	49
Télécharger Correction TP MSTP et RPVST+	51
6. Qu'est-ce que IS-IS ?	51
Télécharger le paquet Wireshark HELLO PDU	54
Télécharger le paquet Wireshark LSPDU	55
Télécharger le paquet Wireshark CSNPDU	55
Télécharger le paquet Wireshark PSNPDU	56
QCM IS-IS.....	58
Correction QCM IS-IS	59
Vidéo de de mise en œuvre IS-IS (DRIVE).....	59
Vidéo de de mise en œuvre IS-IS (Youtube).....	59
Fiche TP IS-IS	59
Télécharger Correction TP IS-IS	60
7. Qu'est-ce que BGP ?.....	61
Télécharger le paquet Wireshark OPEN	62
Télécharger le paquet Wireshark KEEPALIVE	63

Télécharger le paquet Wireshark UPDATE	63
Télécharger le paquet Wireshark NOTIFICATION	63
QCM BGP	67
Correction QCM BGP	68
Vidéo de de mise en œuvre BGP (DRIVE)	68
Vidéo de de mise en œuvre BGP (Youtube).....	68
Fiche TP BGP.....	68
Télécharger Correction TP BGP	69
II – Architecture « Spine Leaf »	70
A- Présentation du principe d'architecture	70
B- Avantages et inconvénients de l'architecture	71
C- Intégration des protocoles dans l'architecture « Spine Leaf ».....	72
III – Confrontation des deux architectures	74
VIDEO DE CONCLUSION (DRIVE)	76
VIDEO DE CONCLUSION (Youtube).....	76
QCM FINAL	76
Correction QCM FINAL	83
Annexe	90
Bibliographie	90
Glossaire	90

Introduction

[Vidéo de présentation de l'équipe \(DRIVE\)](#)

[Vidéo de présentation de l'équipe \(Youtube\)](#)

[Vidéo de présentation du cours \(DRIVE\)](#)

[Vidéo de présentation du cours \(Youtube\)](#)

[QCM de positionnement](#)

a) **Qu'est-ce qu'un VLAN ?**

- ☐ Matériel réseau utilisé pour augmenter la vitesse d'un réseau.
- ☐ Logiciel qui permet d'installer des applications sur plusieurs ordinateurs simultanément.
- ☐ Protocole de communication utilisé pour chiffrer les données sur internet.
- ☐ Réseau local virtuel qui permet de segmenter un réseau physique en plusieurs réseaux logiques distincts.

b) **Que peut-on dire sur le Spanning-Tree (2 réponses) ?**

- ☐ Le Spanning-Tree est un protocole de routage utilisé pour déterminer le chemin le plus court dans un réseau.
- ☐ Il permet de prévenir les boucles de réseau en créant une topologie arborescente à partir des connexions existantes.
- ☐ Il n'est pas nécessaire dans les réseaux modernes en raison de l'utilisation généralisée des commutateurs intelligents.
- ☐ Il peut automatiquement ajuster la topologie du réseau en fonction des changements de connexion.

c) **Que peut-on dire sur eBGP ?**

- ☐ eBGP est responsable de l'attribution d'adresses IP aux appareils sur un réseau local.
- ☐ eBGP est un protocole de routage utilisé pour échanger des informations de routage entre différents systèmes autonomes sur Internet.
- ☐ eBGP utilise un système de métriques pour choisir le meilleur chemin de routage en fonction de la bande passante disponible.
- ☐ eBGP est un protocole de routage utilisé pour échanger des informations de routage internes à un système autonome.

d) **Que peut-on dire sur LACP ?**

- ☐ Il permet de combiner plusieurs liaisons réseau en une seule liaison logique pour augmenter la bande passante et permettre la redondance.
- ☐ LACP est un protocole de routage qui détermine le chemin le plus court entre deux dispositifs sur un réseau.
- ☐ LACP nécessite que les appareils soient configurés pour utiliser des adresses IP statiques afin de fonctionner correctement.
- ☐ Il peut équilibrer le trafic entre les liaisons agrégées en fonction de différents critères, tels que l'adresse MAC ou l'adresse IP.

e) **IS-IS est principalement utilisé dans quel type de réseau ?**

- ☐ Réseaux domestiques.
- ☐ Réseaux de petite entreprise.
- ☐ Réseaux de grande envergure et fournisseurs de services.
- ☐ Réseaux personnels.

f) **Quel est l'objectif principal des ACL ?**

- ☐ Améliorer la bande passante.
- ☐ Contrôler l'accès au réseau et sécuriser les ressources.
- ☐ Augmenter la vitesse de connexion.
- ☐ Gérer les sessions utilisateurs.

g) **Quelle fonctionnalité permet de restreindre l'accès à un équipement au réseau ?**

- ☐ Sticky MAC.
- ☐ OSPF.
- ☐ QoS.
- ☐ NAT.

h) **Quelle est la fonction d'un message ARP Reply ?**

- ☐ Il envoie une demande pour résoudre une adresse IP.
- ☐ Il effectue le routage entre 2 réseaux.
- ☐ Il fournit l'adresse MAC correspondant à une adresse IP demandée.
- ☐ Il bloque les adresses IP non autorisées.

i) **Quel élément est généralement contenu dans une table de routage ?**

- ☐ Les adresses MAC.
- ☐ Les adresses IP des prochains sauts.
- ☐ Contenu des paquets de données.
- ☐ Historique des connexions internet.

j) **Que peut-on dire sur la QoS ?**

- ☐ Elle garantit une connexion internet pour tous les utilisateurs.
- ☐ Elle réduit le nombre d'appareils connectés au réseau.
- ☐ Elle est uniquement utilisée pour les réseaux sans fil.
- ☐ Elle permet de prioriser certains types de trafic pour assurer une performance optimale.

Correction QCM de positionnement

a) **Qu'est-ce qu'un VLAN ?**

- ☐ Matériel réseau utilisé pour augmenter la vitesse d'un réseau.
- ☐ Logiciel qui permet d'installer des applications sur plusieurs ordinateurs simultanément.
- ☐ Protocole de communication utilisé pour chiffrer les données sur internet.
- ☒ Réseau local virtuel qui permet de segmenter un réseau physique en plusieurs réseaux logiques distincts.

b) Que peut-on dire sur le Spanning-Tree (2 réponses) ?

- ☐ Le Spanning-Tree est un protocole de routage utilisé pour déterminer le chemin le plus court dans un réseau.
- ☒ **Il permet de prévenir les boucles de réseau en créant une topologie arborescente à partir des connexions existantes.**
- ☐ Il n'est pas nécessaire dans les réseaux modernes en raison de l'utilisation généralisée des commutateurs intelligents.
- ☒ **Il peut automatiquement ajuster la topologie du réseau en fonction des changements de connexion.**

c) Que peut-on dire sur eBGP ?

- ☐ eBGP est responsable de l'attribution d'adresses IP aux appareils sur un réseau local.
- ☒ **eBGP est un protocole de routage utilisé pour échanger des informations de routage entre différents systèmes autonomes sur Internet.**
- ☐ eBGP utilise un système de métriques pour choisir le meilleur chemin de routage en fonction de la bande passante disponible.
- ☐ eBGP est un protocole de routage utilisé pour échanger des informations de routage internes à un système autonome.

d) Que peut-on dire sur LACP ?

- ☒ **Il permet de combiner plusieurs liaisons réseau en une seule liaison logique pour augmenter la bande passante et permettre la redondance.**
- ☐ LACP est un protocole de routage qui détermine le chemin le plus court entre deux dispositifs sur un réseau.
- ☐ LACP nécessite que les appareils soient configurés pour utiliser des adresses IP statiques afin de fonctionner correctement.
- ☒ **Il peut équilibrer le trafic entre les liaisons agrégées en fonction de différents critères, tels que l'adresse MAC ou l'adresse IP.**

e) IS-IS est principalement utilisé dans quel type de réseau ?

- ☐ Réseaux domestiques.
- ☐ Réseaux de petite entreprise.
- ☒ **Réseaux de grande envergure et fournisseurs de services.**
- ☐ Réseaux personnels.

f) Quel est l'objectif principal des ACL ?

- ☐ Améliorer la bande passante.
- ☒ **Contrôler l'accès au réseau et sécuriser les ressources.**
- ☐ Augmenter la vitesse de connexion.
- ☐ Gérer les sessions utilisateurs.

g) Quelle fonctionnalité permet de restreindre l'accès à un équipement au réseau ?

- ☒ **Sticky MAC.**
- ☐ OSPF.
- ☐ QoS.
- ☐ NAT.

h) Quelle est la fonction d'un message ARP Reply ?

- ☐ Il envoie une demande pour résoudre une adresse IP.
- ☐ Il effectue le routage entre 2 réseaux.
- ☒ **Il fournit l'adresse MAC correspondant à une adresse IP demandée.**
- ☐ Il bloque les adresses IP non autorisées.

i) Quel élément est généralement contenu dans une table de routage ?

- ☐ Les adresses MAC.
- ☒ **Les adresses IP des prochains sauts.**
- ☐ Contenu des paquets de données.
- ☐ Historique des connexions internet.

j) Que peut-on dire sur la QoS ?

- ☐ Elle garantit une connexion internet pour tous les utilisateurs.
- ☐ Elle réduit le nombre d'appareils connectés au réseau.
- ☐ Elle est uniquement utilisée pour les réseaux sans fil.
- ☒ **Elle permet de prioriser certains types de trafic pour assurer une performance optimale.**

I – Architecture traditionnelle « Core, Aggregation, Access »

A- Présentation du principe d'architecture

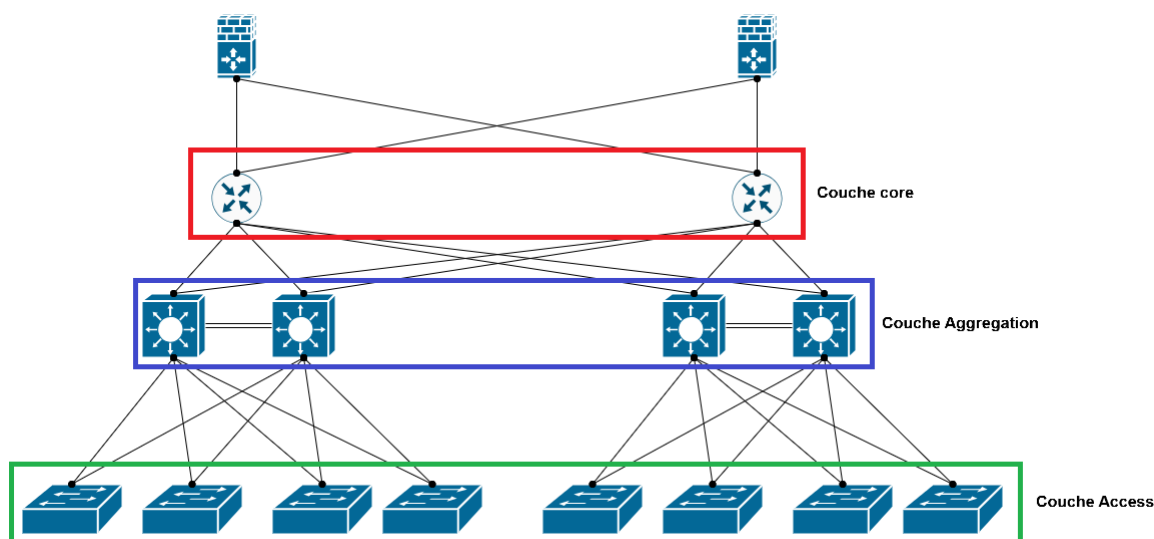
L'architecture « Core, Aggregation, Access » ou en français « Cœur, Agrégation, Accès » est la plus vieille architecture réseaux à 3 niveaux hautement disponibles qui est encore très utilisée de nos jours. La création de cette architecture ne s'est pas faite en un jour, elle est le fruit des diverses évolutions pour répondre aux besoins croissants des entreprises au cours du temps.

C'est dans les années 80 que les 2 premières couches « Core, Access » sont apparues. Les LAN (Local Area Network) des entreprises étaient à l'époque très peu segmentés et sécurisés. Ils servaient uniquement à interconnecter des postes de travail, des imprimantes et des serveurs avec une notion de routage inter réseaux quasiment inexistante.

L'apparition de la troisième couche « Aggregation » a eu lieu dans les années 90, avec l'augmentation du nombre d'utilisateurs et des services mis à disposition sur les réseaux. De plus, la volonté de segmenter les réseaux via des VLAN (Virtual Local Area Network), de les sécuriser via des ACL (Access Control List) mais également d'ajouter de la QoS (Quality of Service), a démontré l'intérêt d'ajouter une troisième couche au modèle d'architecture déjà existant.

Dans les années 2000, avec la démocratisation d'internet et la montée en puissance des infrastructures informatiques au sein des entreprises, l'architecture « Core, Aggregation, Access » a su s'imposer comme une solution redondante et évolutive offrant aussi bien des capacités de commutation que de routage.

Nous allons maintenant aborder le rôle, les fonctionnalités, les caractéristiques et les protocoles pouvant être utilisés pour chacune des couches de l'architecture afin que vous puissiez mieux comprendre leur utilité.



Couche Core :

La couche « Core » en assurant l'interconnexion des différentes zones d'agrégation entre elles mais également l'interconnexion avec les ressources externes (site secondaire, accès internet) constitue le point central de notre architecture. Cette couche est en quelque sorte l'autoroute de notre réseau et doit être capable d'acheminer le trafic avec une grande bande passante, une faible latence et de manière résiliente. C'est pourquoi, les équipements choisis auront généralement des ports optiques disposant d'une grande bande passante. La couche « Core » est, dans la majorité des cas, uniquement composée de routeurs, elle n'interagit donc qu'au niveau 3 du modèle OSI. Les protocoles les plus utilisés sur cette couche sont des protocoles de routage dynamique comme OSPF, EIGRP, IS-IS, BGP mais aussi les protocoles d'interconnexion comme IPSEC, MPLS et VXLAN. Il est à noter que cette couche ne prend pas en charge la sécurité et le filtrage des flux de l'infrastructure. Pour les flux internes, ce rôle est porté par la couche « Aggregation » et pour les flux externes par une couche de firewall située en amont. Cependant, un minimum de sécurité doit être implémenté, notamment l'authentification des protocoles de routage via des clés (MD5, SHA256, etc..) ou un mot de passe à défaut de compatibilité. Le but étant d'empêcher un acteur malveillant d'injecter de fausses routes dans les tables de routage.

Couche Aggregation :

Cette couche travaille aussi bien sur le niveau 2 que sur le niveau 3 du modèle OSI. Cela lui permet, d'une part, de segmenter le réseau vers la couche « Access », et d'autre part d'agréger les différents réseaux pour une redirection vers la couche « Core » dans le but d'accéder aux ressources externes. Les équipements utilisés pour la mise en œuvre de cette couche de l'infrastructure sont des MLS (Multilayer switch). La segmentation du réseau se fait sur cette couche via des VLAN qui seront également propagés sur la couche « Access ». Une défaillance de cette couche pouvant paralyser plusieurs switch « Access », la redondance des liaisons physiques est souvent implémentée via des protocoles comme LACP et/ou PAGP qui permettent également une augmentation de la bande passante. D'autres protocoles de niveau 2 et 3 comme RSTP et/ou VRRP permettent la tolérance aux pannes d'un équipement et ainsi, une continuité d'activité en service dégradé. C'est aussi au niveau de la couche « Aggregation » que le filtrage et la sécurité des flux se font au travers des ACL et de la QoS. Les ACL pourront par exemple être utilisés pour limiter et/ou interdire les communications inter-VLAN. Quant à la QoS, elle permet de prioriser certains VLAN par rapport à d'autres, comme le VLAN d'un client VIP. Si chaque zone

d'agrégation peut disposer de protocoles de niveau 2 différents, à l'inverse les protocoles de niveau 3 doivent être similaires entre la couche « Aggregation » et la couche « Core ».

Couche Access :

La couche « Access » est celle qui est au plus près des serveurs et des clients finaux de l'infrastructure. La proximité avec les périphériques clients va influencer le choix des équipements ; en effet, ils ne disposeront généralement pas de ports optiques, mais de ports cuivre compatibles POE (Power over Ethernet) pour le raccordement de téléphones IP et/ou de bornes WIFI. Cette couche interagit uniquement au niveau 2 du modèle OSI. Bien que l'utilisation de MLS soit possible sur cette couche, cela reste très rare, car de simples switch de couche 2 sont suffisants et moins onéreux. Dans la continuité de la couche « Aggregation », les réseaux sont segmentés via des VLAN et la redondance est assurée via les protocoles STP, RSTP, etc... L'utilisation des protocoles LACP et/ou PAgP permettant la redondance des liaisons physiques et l'augmentation de la bande passante, bien qu'elle soit possible, reste moins répandue qu'au niveau de la couche « Aggregation ». L'implémentation de la sécurité de manière plus granulaire est également possible via l'implémentation du protocole 802.1x, pour limiter l'accès aux réseaux en fonction de l'adresse MAC (Media Access Control) des clients et/ou des serveurs. Il est également courant de voir l'ajout du DHCP snooping pour prévenir la tentative de diffusion d'un DHCP illégitime par un client mal intentionné sur le réseau. Dans la continuité de la couche supérieure, une QoS plus fine peut également être appliquée sur cette couche.

Ce principe d'architecture très populaire a plusieurs fois été décliné et modifié pour s'adapter aux besoins et aux contraintes des entreprises. Il n'est pas si rare de voir des entreprises, avec des budgets limités, utiliser les MLS pour fusionner la couche « Aggregation » et « Core ». Ce type d'architecture a dû également ajouter des points d'accès WIFI sur la couche « Access » pour accueillir les utilisateurs avec des ordinateurs portables se déplaçant dans les locaux des entreprises.

B- Avantages et inconvénients de l'architecture

Maintenant que vous êtes familiarisés avec le fonctionnement des 3 couches de l'architecture, dans cette partie nous allons d'abord aborder les avantages, puis dans un second temps les inconvénients du modèle « Core, Aggregation, Access ».

Les avantages de l'architecture à 3 couches « Core, Aggregation, Access » sont les suivants :

- **Evolutivité** : Chacune des 3 couches de l'architecture peut accueillir des équipements supplémentaires pour augmenter les capacités de l'infrastructure en fonction des besoins. Par exemple, une entreprise qui connaît une forte croissance pourra ajouter des switch sur la couche « Access » pour augmenter le nombre d'utilisateurs connectés au réseau. Dans le cadre d'une construction d'un second bâtiment dans une entreprise, un routeur de couche « Core » et/ou un switch de couche « Aggregation » avec des switch « Access » pourront être ajoutés.
- **Cloisonnement des tâches** : Les tâches du réseau sont réparties sur chacune des couches de manière indépendante. Dans le cas d'une entreprise, la couche « Core » assure le transit général vers l'extérieur, la couche « Aggregation » assure le filtrage et le routage entre les différents services et la couche « Access » distribue le réseau aux clients finaux.

- **Optimisation** : Chacune des couches disposant de son propre rôle, le risque de surutilisation des performances des équipements est fortement réduit et le risque de latence est diminué. Cependant, il est à noter que ce gain de performance minime ne peut être constaté que sur de grosses infrastructures et à condition que les équipements ne soient pas sous-dimensionnés par l'architecte réseau. Dans le contexte d'un DATACENTER avec de gros volumes de flux à destination d'internet, la couche « Core » ne sera chargée que du transit sans se soucier du filtrage de sécurité, il y aura donc une faible latence de traitement.
- **Résilience** : Le design de cette infrastructure inclut implicitement la redondance des équipements qui sont quasiment toujours par paire pour la couche « Core » et « Aggregation ». La présence de deux équipements seulement ne suffit pas à la redondance, les protocoles mis en place doivent être adaptés. En ce qui concerne la couche « Aggregation », la présence de deux MLS reliés à un ou plusieurs switch « Access » ne suffit pas à redonder la Gateway. La mise en place du protocole VRRP ou GLBP est indispensable pour la mise en place d'une redondance de Gateway.
- **Gestion de la sécurité** : La mise en place de la sécurité de plus en plus granulaire à chacune des couches est un véritable atout. La couche « Access » prend en charge la sécurisation de l'accès au réseau par le biais du protocole 802.1x ou du « Port Security », la couche « Aggregation » assure le filtrage des flux entre les différents VLAN via la mise en place d'ACL. La couche « Core » n'applique aucune sécurité ; la sécurité en provenance d'un WAN (Wide Area Network) peut être gérée par la couche « Aggregation » cependant cela est de plus en plus rare. La sécurité des flux en provenance du WAN est aujourd'hui exclusivement confiée à une couche de firewall avant la couche « Core ».

Nous allons maintenant aborder ci-dessous les inconvénients d'une telle infrastructure :

- **Coût élevé** : La mise en place d'une infrastructure à 3 couches nécessite des achats de matériels pouvant être rapidement onéreux. Par exemple, une entreprise qui possède cinq sites, qui souhaite mettre en place une telle architecture sur chacun de ses sites, à raison d'au minimum deux équipements par couche, soit 6 équipements par site, cela représente un total de 30 équipements. Dans le cadre d'une très grosse infrastructure, il est également important de considérer le coût de fonctionnement énergétique.
- **Mise en œuvre complexe** : La mise en place d'une architecture à 3 couches nécessite l'utilisation d'un large panel de protocoles. Il faut également tenir compte de la compatibilité entre les différents constructeurs, mais également des compatibilités entre les protocoles. Par exemple, une entreprise utilisant du matériel CISCO et des protocoles propriétaires CISCO ne pourra pas aisément remplacer son matériel par des équipements provenant d'autres constructeurs.
- **Gestion complexe** : De la même manière que la mise en œuvre de tous ces protocoles peut être complexe, leur maintenance nécessite une maîtrise poussée. Cela implique une large contrainte de formation des équipes IT.

C- Présentation des protocoles

[Vidéo de Présentation de la maquette d'architecture \(DRIVE\)](#)

[Vidéo de Présentation de la maquette d'architecture \(Youtube\)](#)

[Télécharger le LAB](#)

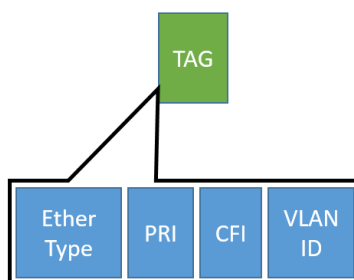
1. Qu'est-ce que les VLAN et les Private VLAN ?

Commençons par un rappel de base sur les VLAN (Virtuelle Local Area Network), le premier concept d'utilisation de VLAN est introduit par le protocole ISL (Inter-Switch-Link) de chez CISCO en 1994. C'est seulement en 1998 que l'IEEE proposera un standard 802.1Q qui sera largement adopté en raison de sa compatibilité inter-constructeur. Les principes de ISL et 802.1Q sont semblables, cependant de légères différences de fonctionnement existent.

Le rôle des VLAN est indispensable afin de segmenter les réseaux et ainsi limiter les zones de diffusion, en vue de réduire la congestion sur la topologie. De plus, cette segmentation permet de sécuriser les réseaux en limitant leur zone d'action. A noter que les VLAN sont cloisonnés entre eux (sauf utilisation d'un routeur ou switch de niveau 3), à l'inverse les hôtes au sein d'un même VLAN peuvent tous communiquer entre eux.

Pour ce faire, ISL ajoute une couche d'encapsulation à l'inverse de 802.1Q qui ajoute seulement une étiquette entre le champ « MAC sources » et « EtherType ». Le fonctionnement de 802.1Q est donc bien moins gourmand en bande passante et en ressources qu'ISL.

L'étiquette « TAG » de 802.1Q est composée de la manière suivante :



- **EtherType** : Indique sur 16 bits qu'il s'agit d'une trame 802.1Q.
- **PRI** : Indique sur 3 bits le niveau de QoS compris entre 0 et 7.
- **CFI** : Indique sur 1 bit s'il s'agit d'un paquet pour un réseau « Ethernet » ou « Token Ring ».
- **VLAN ID** : Indique sur 12 bits le numéro d'identifiant du VLAN auquel appartient la trame.

Dans le cadre de l'utilisation de VLAN, un switch dispose de deux modes de fonctionnalité d'interface :

- **Port Access** : Assigne l'interface à un seul VLAN à l'aide du PVID, ce mode de configuration est réservé aux clients finaux.
- **Port Trunk** : Permet le transit de plusieurs VLAN, ce mode de configuration est réservé pour l'interconnexion du switch avec un autre switch ou un routeur.

Le Port VLAN Identifier ou PVID, est défini sur les ports en mode « Access » pour transmettre les trames non taguées dans un VLAN spécifique. Par exemple, l'ordinateur connecté sur le port d'un switch qui a son PVID égal à 10, verra ses trames taguées sur le VLAN 10 de manière transparente pour lui.

Un autre concept des VLAN à maîtriser est le VLAN natif, ce dernier se configure sur un lien Trunk. Il a pour but d'attribuer les trames non taguées à un VLAN pour la communication sur le lien Trunk. Le natif VLAN doit être identique sur les deux ports de la liaison Trunk pour garantir une communication.

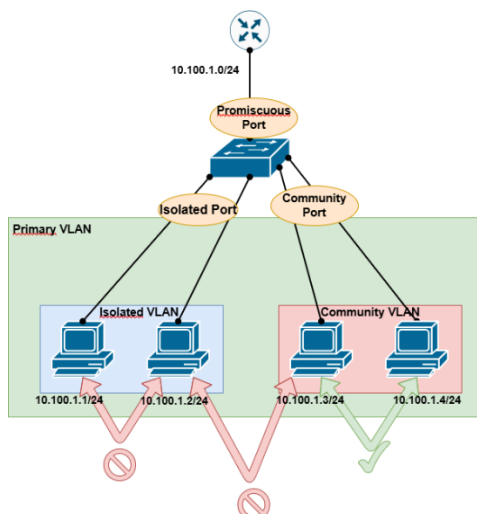
Les VLAN disposent d'une fonctionnalité avancée, les Private VLAN (PVLAN), qui permet de cloisonner les VLAN en sous-ensembles afin d'empêcher les hôtes d'un même VLAN de communiquer entre eux. Les PVLAN sont particulièrement appréciés dans les environnements nécessitant une confidentialité et une sécurité strictes, mais également dans les infrastructures mutualisées.

A l'inverse des VLAN, les PVLAN ne disposent pas de deux modes de fonctionnalités des interfaces, mais de trois :

- **Promiscuous Port** : Peut recevoir et transmettre des paquets avec tous les ports du VLAN. Ce mode est généralement réservé pour l'interconnexion avec un routeur ou un switch de niveau 3.
- **Isolated Port** : Peut échanger des paquets uniquement avec le ou les « Promiscuous Port », il est totalement isolé des autres ports, qu'ils soient au sein du même PVLAN ou non.
- **Community Port** : Peut communiquer avec les autres ports faisant partie de la même communauté de PVLAN, mais aussi avec le ou les « Promiscuous Port ». Cela permet de définir des communautés de ports isolées entre elles ainsi que des « Isolated Port ».

Le VLAN divisé en sous-ensembles se nomme le « Primary VLAN », il permet le transfert entre les interfaces en mode « Promiscuous Port » et les sous-ensembles « Secondary VLAN ». A noter qu'il existe deux types de « Secondary VLAN » :

- **Isolated VLAN** : Permet aux interfaces « Isolated Port » de communiquer avec le ou les « Promiscuous Port ».
- **Community VLAN** : Permet aux interfaces « Community Port » de communiquer entre elles au sein de la même communauté et de communiquer avec le ou les « Promiscuous Port ».



Les switch CISCO sont dotés de protocoles propriétaires complémentaires pour la configuration VLAN comme Dynamic Trunk Protocol « DTP » et VLAN Trunking Protocol « VTP ».

Le protocole DTP permet de négocier de manière automatique les liaisons Trunk entre deux switch CISCO, via l'échange de messages de négociation. Quant au protocole VTP, il permet à un switch dit « serveur » de propager les VLAN sur les switch dits « client ».

Bien que les VLAN soient une manière de sécuriser une topologie réseau, ils doivent eux-mêmes être sécurisés, cela passe aussi bien par une bonne configuration que par un contrôle d'accès.

L'utilisation de DTP, bien qu'utile dans certains cas en vue de simplifier les configurations, peut être une faille de sécurité. En effet, si DTP est activé sur un port, un utilisateur malintentionné connecté sur ce port pourrait envoyer de faux paquets DTP et ainsi faire passer ce port du switch en mode Trunk. Une fois le port en mode Trunk, il n'y a plus qu'à créer des sous-interfaces virtuelles taguées sur les VLAN auxquels on souhaite accéder. L'utilisateur pourra maintenant accéder à tous les VLAN qu'il souhaite sans restriction. Cette attaque peut être contrée en restreignant l'utilisation de DTP à un nombre limité de ports ou encore mieux, en désactivant totalement DTP.

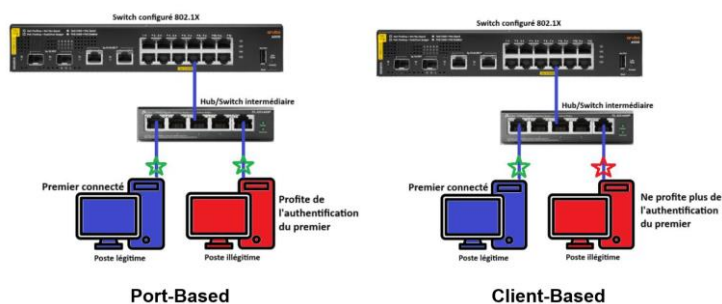
VTP fonctionne sur le principe du client serveur, c'est-à-dire qu'un switch qui porte le rôle du serveur peut créer des VLAN sur les switch clients, bien que cela permette de centraliser la configuration, cela représente une faille de sécurité. En effet, si le protocole VTP reste actif sur des switch clients, un attaquant pourrait simuler l'envoi de paquets VTP pour écraser la configuration des VLAN actuelle et déclarer d'autres VLAN. Le réseau se trouverait donc totalement paralysé.

L'une des bonnes pratiques de configuration est également la modification du VLAN natif sur les interfaces. En effet, par défaut le VLAN natif est le VLAN 1, le premier réflexe d'un attaquant sera donc d'envoyer du trafic tagué sur le VLAN 1 ; or, si le VLAN natif a été personnalisé, cela compliquera la tâche de l'attaquant.

La sécurisation des VLAN passe par un contrôle d'accès rigoureux des clients sur le réseau via l'utilisation du 802.1x et des Ports Security.

Le protocole 802.1x permet l'authentification des postes clients via un certificat, en vue d'identifier à quel VLAN ils seront assignés. Pour ce faire, une fois connecté au réseau, le poste client envoie ses informations d'authentification au switch qui les relaie au serveur d'authentification (Ex : RADIUS). Tant que le serveur n'a pas validé les informations d'authentification, le port du switch est maintenu en état « non authentifié », aucun trafic réseau n'est autorisé. Une fois que le client est authentifié par le serveur, le switch peut appliquer la politique réseau dédiée à ce client : l'assignation du VLAN et de la QoS. A noter que l'authentification peut se faire de deux manières :

- **Port-Based** : Dans ce cas, l'authentification est faite pour le port physique du switch peu importe le nombre de clients reliés à ce port. Si un hub ou un petit switch est connecté au switch principal, le premier client authentifié fait profiter de son authentification à tous les autres clients reliés sur ce port.
- **Client-Based** : Dans ce cas, l'authentification est faite pour chacun des clients de manière individuelle sans se soucier du port.



L'authentification « Client-Based » sera préférée, car elle permet une gestion des accès réseaux plus fine évitant ainsi qu'un client illégitime ne bénéficie de l'authentification d'un client légitime.

La mise en place du 802.1x c'est un véritable atout en complément des VLAN pour segmenter et sécuriser les réseaux.

Le « Port Security » est une fonctionnalité de sécurité qui permet au switch de filtrer et de restreindre les accès aux réseaux en fonction de l'adresse MAC des postes clients. L'apprentissage de la MAC des postes clients peut se faire des trois manières suivantes :

- **Statique** : Les MAC sont configurées de manière manuelle sur les interfaces.
- **Dynamique** : Les MAC sont apprises à la connexion du poste client sur le switch, le nombre de MAC pouvant être apprises est paramétré par l'administrateur. A noter que les MAC ne sont pas enregistrées dans la configuration et seront oubliées au redémarrage du matériel.
- **Sticky** : Les MAC sont également apprises à la connexion du poste client sur le switch. Cependant, les MAC sont enregistrées de manière permanente dans la configuration du switch.

Dans le cas d'une violation de la politique de sécurité, l'une des trois réactions ci-dessous sera mise en œuvre en fonction de la configuration du switch.

- **Protect** : Le port reste actif, mais le trafic est bloqué pour la MAC illégitime sans journalisation.
- **Restrict** : Le port reste actif, mais le trafic est bloqué pour la MAC illégitime avec une journalisation (SYSLOG, SNMP).
- **Shutdown** : Le port est désactivé avec une journalisation. A noter qu'une intervention humaine ou un redémarrage du switch sera nécessaire pour réactiver le port.

La gestion des accès via le « Port Security » est le minimum qui puisse être implémenté dans une topologie réseau. Cependant, l'utilisation du 802.1x lui sera préférée pour les grandes topologies nécessitant une gestion centralisée. De plus, il est aujourd'hui possible d'usurper une MAC : le Port Security, bien qu'une première couche de sécurité, n'est pas inviolable.

Maintenant que vous maîtrisez le concept de VLAN et PVLAN, voyons comment l'intégrer dans notre infrastructure à 3 niveaux. L'utilisation des VLAN est majoritairement mise en place sur les couches « Access » et « Aggregation ». En effet, la couche « Access » dispose de plusieurs VLAN avec des ports en mode « Access » pour connecter les clients finaux dans leurs VLAN respectifs et des ports en mode « Trunk » pour l'interconnexion avec la couche « Aggregation ». La couche « Aggregation » ne dispose que d'interfaces en mode « Trunk » connectées à la couche « Access » qui sont utilisées pour le routage entre les différents VLAN. La couche « Core », n'interagissant qu'au niveau 3, ne nécessite pas l'implémentation de VLAN.

Dans le cadre de notre infrastructure nous mettrons en place des VLAN sur nos couches « Access » et « Aggregation ».

QCM VLAN et PVLAN

a) Quel est le rôle principal d'un VLAN dans un réseau ?

- ☐ Améliorer la vitesse des liens physiques.
- ☐ Isoler les segments du réseau logiquement.
- ☐ Fournir une redondance automatique.
- ☐ Créer un routage entre les segments.

b) Quel type de PVLAN permet aux hôtes de communiquer uniquement avec le routeur ?

- ☐ Primary VLAN
- ☐ Community VLAN
- ☐ Isolated VLAN
- ☐ Secondary VLAN

c) Quel VLAN est utilisé par défaut sur un switch Cisco ?

- ☐ VLAN 1
- ☐ VLAN 10
- ☐ VLAN 100
- ☐ Aucun VLAN par défaut

d) Quel est l'objectif principal du VTP dans un réseau VLAN ?

- ☐ Configurer automatiquement les ports en mode Trunk.
- ☐ Propager les informations de VLAN entre les switch.
- ☐ Mettre à jour les versions du firmware des switch.
- ☐ Assurer la communication entre PVLAN.

e) Quel est le rôle principal de DTP ?

- ☐ Assurer la redondance des liens.
- ☐ Configurer dynamiquement les liens en mode Trunk ou Access.
- ☐ Synchroniser les VLAN entre les switch.
- ☐ Activer les PVLAN sur les interfaces.

Correction QCM VLAN et PVLAN

a) Quel est le rôle principal d'un VLAN dans un réseau ?

- ☐ Améliorer la vitesse des liens physiques.
- ☒ **Isoler les segments du réseau logiquement.**
- ☐ Fournir une redondance automatique.
- ☐ Créer un routage entre les segments.

b) Quel type de PVLAN permet aux hôtes de communiquer uniquement avec le routeur ?

- ☐ Primary VLAN
- ☐ Community VLAN
- ☒ **Isolated VLAN**
- ☐ Secondary VLAN

c) Quel VLAN est utilisé par défaut sur un switch Cisco ?

- ☒ **VLAN 1**
- ☐ VLAN 10
- ☐ VLAN 100
- ☐ Aucun VLAN par défaut

d) Quel est l'objectif principal du VTP dans un réseau VLAN ?

- ☐ Configurer automatiquement les ports en mode Trunk.
- ☒ **Propager les informations de VLAN entre les switch.**
- ☐ Mettre à jour les versions du firmware des switch.
- ☐ Assurer la communication entre PVLAN.

e) Quel est le rôle principal de DTP ?

- ☐ Assurer la redondance des liens.
- ☒ **Configurer dynamiquement les liens en mode Trunk ou Access.**
- ☐ Synchroniser les VLAN entre les switch.
- ☐ Activer les PVLAN sur les interfaces.

[Vidéo de de mise en œuvre VLAN et PVLAN \(DRIVE\)](#)

[Vidéo de de mise en œuvre VLAN et PVLAN \(Youtube\)](#)

Fiche TP VLAN et PVLAN

1. Configurer les noms de tous les équipements.
2. Désactiver la recherche DNS sur tous les équipements.
3. Configurer le nom de domaine « esgi.cloud » sur tous les équipements.
4. Activer le chiffrement des mots de passe sur tous les équipements.
5. Créer un compte « admin » avec le mot de passe « esgi » sur tous les équipements.
6. Définir un mot de passe pour la connexion au terminal et le mode d'exécution privilégié.
7. Configurer SSH version 2 sur tous les équipements.
8. Configurer les bannières de login sur tous les équipements.
9. Configurer les VLAN sur les équipements de la couche « Access » et la couche « Aggregation ».

Zone	VLAN	Description
IEEE	100	ADMIN
IEEE	110	CLT-110
IEEE	120	CLT-120
IEEE	130	CLT-130
IEEE	140	CLT-140
CISCO	200	ADMIN
CISCO	210	CLT-210
CISCO	220	CLT-220
CISCO	230	CLT-230
CISCO	240	CLT-240

10. Désactiver DTP sur tous les équipements de la couche « Access » et la couche « Aggregation ».
11. Configurer les interfaces entre la couche « Access » et la couche « Aggregation » en mode « Trunk » et n'autoriser que les VLAN appropriés.

Zone	IEEE	CISCO
VLAN autorisé	100,110,120,130,140	200,210,220,230,240

12. Configurer les interfaces clientes de la couche « Access » en mode « Access » pour y affecter un VLAN.
13. Configurer les interfaces clientes de la couche « Access » avec les paramètres « Port Sécurité » suivants :

Méthode d'apprentissage	Nombre d'apprentissage	Politique de violation
sticky	2	shutdown

14. Configurer les IP et les descriptions des interfaces pour tous les équipements (se référer au tableau ci-dessous).

DC1-CORE1			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.1.1/30	2001:db8:0:1::1/64	TO-DC1-DST1-G0/0
GigabitEthernet0/1	172.16.1.5/30	2001:db8:0:2::1/64	TO-DC1-DST2-G0/0
GigabitEthernet0/2	172.16.1.9/30	2001:db8:0:3::1/64	TO-DC1-DST3-G0/1
GigabitEthernet0/3	172.16.1.13/30	2001:db8:0:4::1/64	TO-DC1-DST4-G0/1
GigabitEthernet1/0	172.16.1.42/30	2001:db8:0:9::1/64	TO-DC1-FW1-P1
GigabitEthernet1/1	172.16.1.45/30	2001:db8:0:11::1/64	TO-DC1-FW2-P1

DC1-CORE2			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.1.25/30	2001:db8:0:5::1/64	TO-DC1-DST3-G0/0
GigabitEthernet0/1	172.16.1.29/30	2001:db8:0:6::1/64	TO-DC1-DST4-G0/0
GigabitEthernet0/2	172.16.1.21/30	2001:db8:0:7::1/64	TO-DC1-DST2-G0/1
GigabitEthernet0/3	172.16.1.17/30	2001:db8:0:8::1/64	TO-DC1-DST1-G0/1
GigabitEthernet1/0	172.16.1.50/30	2001:db8:0:10::1/64	TO-DC1-FW1-P2
GigabitEthernet1/1	172.16.1.53/30	2001:db8:0:12::1/64	TO-DC1-FW2-P2

DC1-DST1			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.1.2/30	2001:db8:0:1::2/64	TO-DC1-CORE1-G0/0
GigabitEthernet0/1	172.16.1.18/30	2001:db8:0:8::2/64	TO-DC1-CORE2-G0/3
interface vlan 100	192.168.100.2/24	N/A	VLAN-ADMIN
interface vlan 110	192.168.110.2/24	N/A	VLAN-CLT-110
interface vlan 120	192.168.120.2/24	N/A	VLAN-CLT-120
interface vlan 130	192.168.130.2/24	N/A	VLAN-CLT-130
interface vlan 140	192.168.140.2/24	N/A	VLAN-CLT-140

DC1-DST2			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.1.6/30	2001:db8:0:2::2/64	TO-DC1-CORE1-G0/1
GigabitEthernet0/1	172.16.1.22/30	2001:db8:0:7::2/64	TO-DC1-CORE2-G0/2
interface vlan 100	192.168.100.3/24	N/A	VLAN-ADMIN

interface vlan 110	192.168.110.3/24	N/A	VLAN-CLT-110
interface vlan 120	192.168.120.3/24	N/A	VLAN-CLT-120
interface vlan 130	192.168.130.3/24	N/A	VLAN-CLT-130
interface vlan 140	192.168.140.3/24	N/A	VLAN-CLT-140

DC1-DST3			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.1.26/30	2001:db8:0:5::2/64	TO-DC1-CORE2-G0/0
GigabitEthernet0/1	172.16.1.10/30	2001:db8:0:3::2/64	TO-DC1-CORE1-G0/2
interface vlan 200	192.168.200.2/24	N/A	VLAN-ADMIN
interface vlan 210	192.168.210.2/24	N/A	VLAN-CLT-210
interface vlan 220	192.168.220.2/24	N/A	VLAN-CLT-220
interface vlan 230	192.168.230.2/24	N/A	VLAN-CLT-230
interface vlan 240	192.168.240.2/24	N/A	VLAN-CLT-240
DC1-DST4			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.1.30/30	2001:db8:0:6::2/64	TO-DC1-CORE2-G0/1
GigabitEthernet0/1	172.16.1.14/30	2001:db8:0:4::2/64	TO-DC1-CORE1-G0/3
interface vlan 200	192.168.200.3/24	N/A	VLAN-ADMIN
interface vlan 210	192.168.210.3/24	N/A	VLAN-CLT-210
interface vlan 220	192.168.220.3/24	N/A	VLAN-CLT-220
interface vlan 230	192.168.230.3/24	N/A	VLAN-CLT-230
interface vlan 240	192.168.240.3/24	N/A	VLAN-CLT-240

DC1-ACCESS1			
Port	IPv4	IPv6	Description
VLAN100	192.168.100.11/24	N/A	VLAN-ADMIN

DC1-ACCESS2			
Port	IPv4	IPv6	Description
VLAN100	192.168.100.12/24	N/A	VLAN-ADMIN

DC1-ACCESS3			
Port	IPv4	IPv6	Description
VLAN100	192.168.100.13/24	N/A	VLAN-ADMIN

DC1-ACCESS4			
Port	IPv4	IPv6	Description
VLAN100	192.168.100.14/24	N/A	VLAN-ADMIN

DC1-ACCESS5			
Port	IPv4	IPv6	Description
VLAN100	192.168.200.15/24	N/A	VLAN-ADMIN

DC1-ACCESS6			
Port	IPv4	IPv6	Description

VLAN100	192.168.200.16/24	N/A	VLAN-ADMIN
---------	-------------------	-----	------------

DC1-ACCESS7			
Port	IPv4	IPv6	Description
VLAN100	192.168.200.17/24	N/A	VLAN-ADMIN

DC1-ACCESS8			
Port	IPv4	IPv6	Description
VLAN100	192.168.200.18/24	N/A	VLAN-ADMIN

RTR-ORANGE			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	10.1.1.2/30	N/A	TO-DC1-FW1-P4
GigabitEthernet0/1	10.1.1.6/30	N/A	TO-DC1-FW2-P4
GigabitEthernet0/3	10.1.1.9/30	N/A	TO-RTR-BOUYGUE-G0/3

RTR-BOUYGUE			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	10.1.1.21/30	N/A	TO-RTR-SFR
GigabitEthernet0/3	10.1.1.10/30	N/A	TO-RTR-ORANGE

RTR-SFR			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	10.1.1.25/30	N/A	TO-RTR-FRE
GigabitEthernet0/3	10.1.1.22/30	N/A	TO-RTR-BOUYGUE

RTR-FREE			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	10.1.1.30/30	N/A	TO-DC2-FW1-P4
GigabitEthernet0/1	10.1.1.33/30	N/A	TO-DC2-FW2-P4
GigabitEthernet0/3	10.1.1.26/30	N/A	TO-RTR-SFR

DC2-SPINE1			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.3.1/30	2001:db8:0:13::1/64	TO-DC2-LEAF1-G0/0
GigabitEthernet0/1	172.16.3.5/30	2001:db8:0:14::1/64	TO-DC2-LEAF2-G0/0
GigabitEthernet0/2	172.16.3.9/30	2001:db8:0:15::1/64	TO-DC2-LEAF3-G0/0
GigabitEthernet0/3	172.16.3.13/30	2001:db8:0:16::1/64	TO-DC2-LEAF4-G0/0
GigabitEthernet1/0	172.16.3.33/30	2001:db8:0:17::1/64	TO-DC2-FW1-P1
GigabitEthernet1/1	172.16.3.37/30	2001:db8:0:18::1/64	TO-DC2-FW2-P1

DC2-SPINE2			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.3.17/30	2001:db8:0:19::1/64	TO-DC2-LEAF1-G0/1
GigabitEthernet0/1	172.16.3.21/30	2001:db8:0:20::1/64	TO-DC2-LEAF2-G0/1

GigabitEthernet0/2	172.16.3.25/30	2001:db8:0:21::1/64	TO-DC2-LEAF3-G0/1
GigabitEthernet0/3	172.16.3.29/30	2001:db8:0:22::1/64	TO-DC2-LEAF4-G0/1
GigabitEthernet1/0	172.16.3.41/30	2001:db8:0:23::1/64	TO-DC2-FW1-P2
GigabitEthernet1/1	172.16.3.45/30	2001:db8:0:24::1/64	TO-DC2-FW2-P2

DC2-LEAF1			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.3.2/30	2001:db8:0:13::2/64	TO-DC2-SPINE1-G0/0
GigabitEthernet0/1	172.16.3.18/30	2001:db8:0:19::2/64	TO-DC2-SPINE2-G0/0

DC2-LEAF2			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.3.6/30	2001:db8:0:14::2/64	TO-DC2-SPINE1-G0/1
GigabitEthernet0/1	172.16.3.22/30	2001:db8:0:20::2/64	TO-DC2-SPINE2-G0/1

DC2-LEAF3			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.3.10/30	2001:db8:0:15::2/64	TO-DC2-SPINE1-G0/2
GigabitEthernet0/1	172.16.3.26/30	2001:db8:0:21::2/64	TO-DC2-SPINE2-G0/2

DC2-LEAF4			
Port	IPv4	IPv6	Description
GigabitEthernet0/0	172.16.3.14/30	2001:db8:0:16::2/64	TO-DC2-SPINE1-G0/3
GigabitEthernet0/1	172.16.3.30/30	2001:db8:0:22::2/64	TO-DC2-SPINE2-G0/3

[Télécharger Correction TP VLAN et PVLAN](#)

2. Qu'est-ce qu'une ACL ?

Les ACL pour « Access Control List », ne sont pas à proprement parler un protocole, mais plutôt un principe. Il existe des ACL pour divers types de ressources informatiques, comme les fichiers et dans notre cas les réseaux. Le principe d'une ACL est de limiter l'accès à certaines ressources pour certains utilisateurs. Les ACL réseaux sont apparues dans les années 80 avec les ACL standard et ont continué à se perfectionner avec l'apparition des ACL étendues dans les années 90.

Les ACL sont comparées aux paquets réseaux de manière séquentielle, c'est-à-dire que, dès qu'une règle peut être associée à un paquet, l'action de cette règle est appliquée et les règles suivantes sont ignorées. En l'absence de règle applicable aux paquets réseaux, une règle implicite de « Deny » comprenez de refus, est appliquée. Une ACL est toujours accompagnée d'une action, soit d'autorisation « Permit », soit d'interdiction « Deny ».

A noter que les ACL sont dites « Stateless », c'est-à-dire que ce qui est autorisé dans un sens, ne l'est pas dans l'autre.

Il existe deux types d'ACL présentés ci-dessous :

- **ACL standard** : Les ACL standard sont très limitées, se basent uniquement sur l'IP source pour autoriser ou non le paquet (Permit/Deny). A noter que les ACL standard sont placées au plus proche de la destination.
- **ACL étendues** : Les ACL étendues sont beaucoup plus performantes et peuvent se baser sur les IP sources et destinations, les MAC sources et destinations, les ports et les protocoles (TCP, UDP, ICMP). A noter que les ACL étendues sont placées au plus près de la source.

Les ACL qu'elles soient standard ou étendues peuvent être déclarées des deux manières suivantes :

- **ACL Numérotées** : Elles sont définies dans des plages, de 1 à 99 puis de 1300 à 1999 pour les ACL standard et de 100 à 199 puis de 2000 à 2699 pour les ACL étendues. Une fois créées, les ACL numérotées ne peuvent pas être modifiées, en cas de nécessité il faut supprimer et recréer la règle. Le principe de numérotation réduit la lisibilité des ACL, cela limite l'utilisation des ACL numérotées aux petites infrastructures réseaux peu complexes.
- **ACL Nommées** : Les ACL sont nommées de manière personnalisée par l'administrateur. De plus, contrairement aux ACL numérotées, il est possible d'éditer les ACL après leur création. Cette lisibilité et cette flexibilité font que les ACL nommées sont majoritairement préférées aujourd'hui, surtout dans les architectures denses et complexes.

Les ACL sont un mécanisme de sécurité et leur sécurisation ne dépend que de leur bonne configuration, il n'y a donc aucun mécanisme supplémentaire pour sécuriser une ACL.

Dans les architectures à 3 niveaux, l'utilisation des ACL se fait en grande majorité sur la couche « Aggregation » pour restreindre le routage entre les différents VLAN, afin de les rendre hermétiques et de s'assurer du respect de la confidentialité. Leurs utilisations plus minimales sur la couche « Core » permettent d'ignorer certaines routes indésirables et aussi l'interconnexion entre les DATACENTER via la mise en place de tunnels IPSEC.

Dans notre cas, nous mettrons en œuvre les ACL sur la couche « Aggregation » pour le cloisonnement des VLAN.

QCM ACL

a) Quel est le rôle principal d'une ACL dans un réseau ?

- ☐ Chiffrer le trafic réseau.
- ☐ Contrôler le flux de trafic basé sur des règles définies.
- ☐ Permettre le routage inter-VLAN.
- ☐ Configurer automatiquement les VLAN.

b) Quel numéro identifie une ACL standard dans un environnement Cisco ?

- ☐ De 1 à 99
- ☐ De 100 à 199
- ☐ De 2000 à 2699
- ☐ De 1000 à 1999

c) Où une ACL standard doit-elle être appliquée pour bloquer un trafic indésirable le plus efficacement possible ?

- ☐ Aussi proche que possible de la source.
- ☐ Aussi proche que possible de la destination.
- ☐ Toujours sur l'interface entrante.
- ☐ Toujours sur l'interface sortante.

d) Quel mode de restriction Port Security permet de désactiver un port si une violation est détectée ?

- ☐ Restrict
- ☐ Shutdown
- ☐ Protect
- ☐ Monitor

e) Quel est le rôle principal du Port Security sur un switch ?

- ☐ Bloquer les ports inutilisés.
- ☐ Limiter le nombre d'adresses MAC apprises sur un port.
- ☐ Configurer dynamiquement les ports en mode Access.
- ☐ Autoriser les VLAN spécifiques.

Correction QCM ACL

a) Quel est le rôle principal d'une ACL dans un réseau ?

- ☐ Chiffrer le trafic réseau.
- ☒ **Contrôler le flux de trafic basé sur des règles définies.**
- ☐ Permettre le routage inter-VLAN.
- ☐ Configurer automatiquement les VLAN.

b) Quel numéro identifie une ACL standard dans un environnement Cisco ?

- ☒ **De 1 à 99**
- ☐ De 100 à 199
- ☐ De 2000 à 2699
- ☐ De 1000 à 1999

c) Où une ACL standard doit-elle être appliquée pour bloquer un trafic indésirable le plus efficacement possible ?

- ☒ **Aussi proche que possible de la source.**
- ☐ Aussi proche que possible de la destination.
- ☐ Toujours sur l'interface entrante.
- ☐ Toujours sur l'interface sortante.

d) Quel mode de restriction Port Security permet de désactiver un port si une violation est détectée ?

- ☐ Restrict
- ☒ **Shutdown**
- ☐ Protect
- ☐ Monitor

e) Quel est le rôle principal du Port Security sur un switch ?

- ☐ Bloquer les ports inutilisés.
- ☒ **Limiter le nombre d'adresses MAC apprises sur un port.**
- ☐ Configurer dynamiquement les ports en mode Access.
- ☐ Autoriser les VLAN spécifiques.

[Vidéo de de mise en œuvre ACL \(DRIVE\)](#)
[Vidéo de de mise en œuvre ACL \(YOUTUBE\)](#)

[Fiche TP ACL](#)

1. Configurer les ACL sur les MLS « DC1-DST2 », « DC1-DST3 » et « DC1-DST4 ». Ces ACL doivent empêcher la communication entre les VLAN 110, 120, 130, 140 pour la zone « IEEE » et entre les VLAN 210, 220, 230, 240 pour la zone « CISCO ». Toutes les autres communications sont autorisées.

[Télécharger Correction TP ACL](#)

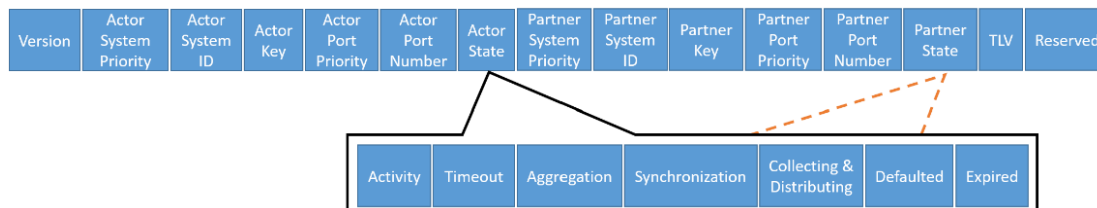
3. Qu'est-ce que PAgP et LACP ?

LACP et PAgP, connus sous le nom EtherChannel, sont deux protocoles semblables, qui permettent l'agrégation de plusieurs liaisons physiques en vue de créer une seule liaison logique. L'utilisation de ces protocoles offre une redondance des liens physiques, mais également une augmentation de la bande passante. Ces protocoles sont précieux afin de mettre en place une infrastructure résiliente avec de larges bandes passantes.

EtherChannel est apparu pour la première fois au début des années 90 avec le protocole PAgP proposé par CISCO. Il faudra attendre les années 2000 pour que l'IEEE propose une version standardisée 802.3ad, qui offrira une compatibilité inter-constructeur. Bien que ces deux protocoles ne soient pas totalement identiques dans leur fonctionnement, leur objectif reste le même.

LACP utilise des messages « Link Aggregation Control Data Units » ou LACPDU en vue d'établir et de maintenir un échange d'informations, indispensable à la négociation des liens agrégés.

Les LACPDU sont composés de la manière suivante :



- **Version** : Indique sur 1 octet la version du protocole LACP (aujourd'hui il n'existe qu'une version).
- **Actor System Priority** : Indique sur 2 octets la valeur du switch émetteur comprise entre 1 (priorité la plus élevée) et 65535 (priorité la plus basse), en vue de déterminer le switch acteur de l'agrégation. Le switch qui a la priorité la plus basse est le switch acteur, si les priorités sont égales l'adresse MAC la plus basse sera préférée. A noter que la valeur par défaut est « 32768 ».
- **Actor System ID** : Indique sur 6 octets l'identifiant du switch émetteur, généralement l'adresse MAC du switch.
- **Actor Key** : Indique sur 2 octets la valeur du switch émetteur comprise entre 1 et 65535 qui permet d'identifier le groupe d'agrégation. Cette valeur est obligatoirement définie manuellement durant la configuration.
- **Actor Port Priority** : Indique sur 2 octets la valeur du switch émetteur comprise entre 1 (priorité la plus élevée) et 65535 (priorité la plus basse) qui permet d'identifier le port actif de l'agrégation. En cas d'égalité le numéro de port le plus bas, ainsi que l'adresse MAC la plus basse seront préférés. A noter que par défaut cette valeur est « 32768 », cependant elle peut être personnalisée par l'administrateur.
- **Actor Port Number** : Indique sur 2 octets le numéro d'identifiant du port sur le switch émetteur. Ce numéro d'identification est attribué par le switch lui-même et ne peut être modifié.
- **Actor State** : Indique sur 1 octet les différents états du port émetteur au sein de l'agrégat, notamment :
 - **Activity** : Indique si le port est défini en mode actif ou passif.
 - **Timeout** : Indique la durée avant que le lien ne soit inactif.
 - **Aggregation** : Indique l'état de l'agrégation avec le partenaire.
 - **Synchronization** : Indique l'état de synchronisation avec le partenaire.
 - **Collecting & Distributing** : Indique la possibilité de recevoir et d'émettre les LACPDU.
 - **Defaulted** : Indique si les valeurs utilisées sont les valeurs par défaut ou configurées par l'administrateur.
 - **Expired** : Indique si l'agrégation a expiré ou non.
- **Partner System Priority** : Indique sur 2 octets la valeur du switch partenaire comprise entre 1 (priorité la plus élevée) et 65535 (priorité la plus basse), en vue de déterminer le switch acteur de l'agrégation. A noter que la valeur par défaut est « 32768 », cependant elle peut être personnalisée par l'administrateur.
- **Partner System ID** : Indique sur 6 octets l'identifiant du switch partenaire, généralement l'adresse MAC du switch.
- **Partner Key** : Indique sur 2 octets la valeur du switch partenaire comprise entre 1 et 65535 qui permet d'identifier les groupes d'agrégation.
- **Partner Port Priority** : Indique sur 2 octets la valeur du switch partenaire comprise entre 1 (priorité la plus élevée) et 65535 (priorité la plus basse) qui permet d'identifier le port actif de l'agrégation. En cas d'égalité le numéro de port le plus bas, ainsi que l'adresse MAC la plus

basse seront préférés. A noter que par défaut cette valeur est « 32768 », cependant elle peut être personnalisée par l'administrateur.

- **Partner Port Number** : Indique sur 2 octets le numéro d'identifiant du port sur le switch partenaire.
- **Partner State** : Indique sur 1 octet les différents états du port partenaire au sein de l'agrégat.
 - **Activity** : Indique si le port est défini en mode actif ou passif.
 - **Timeout** : Indique la durée avant que le lien ne soit inactif.
 - **Aggregation** : Indique l'état de l'agrégation avec le partenaire.
 - **Synchronization** : Indique l'état de synchronisation avec le partenaire.
 - **Collecting & Distributing** : Indique la possibilité de recevoir et d'émettre les LACPDU.
 - **Defaulted** : Indique si les valeurs utilisées sont les valeurs par défaut ou configurées par l'administrateur.
 - **Expired** : Indique si l'agrégation a expiré ou non.
- **TLV** : Indique les informations complémentaires si besoin.
- **Reserved** : Champ d'un octet réservé pour les extensions futures du protocole.

[Télécharger le paquet Wireshark LACPDU](#)

Les champs « Actor » sont remplis dès le premier échange LACPDU, contrairement aux champs « Partner » qui ne sont remplis qu'à partir du deuxième échange LACPDU. Les messages LACPDU sont ensuite échangés périodiquement et mis à jour en fonction des événements réseaux.

Les ports faisant partie de l'agrégation peuvent être configurés en deux modes :

- **Actif** : Dans ce mode le port envoie des LACPDU pour établir l'agrégation avec le switch voisin.
- **Passif** : Dans ce mode le port attend de recevoir un premier LACPDU avant d'établir l'agrégation avec le switch voisin.

Vous pouvez voir ci-dessous le tableau de négociation des agrégations de liens pour LACP :

Port équipement A	Port équipement B	Négociation
Actif	Actif	Agrégation LACP établie
Actif	Passif	Agrégation LACP établie
Passif	Actif	Agrégation LACP établie
Passif	Passif	Aucune agrégation

En plus du mode des ports, la fréquence d'envoi des messages LACPDU est un paramètre de configuration indispensable. Ce paramètre dispose de deux modes de configuration :

- **Vitesse rapide (Fast)** : Envoie des messages LACPDU toutes les 1 seconde, pour une détection rapide des incidents.
- **Vitesse lente (Slow)** : Envoi des messages LACPDU toutes les 30 secondes, pour une préférence de stabilité.

LACP offre deux grands modes de fonctionnalités avancées :

- **Failover** : intègre une gestion dynamique des pannes qui assure une stabilité du lien en retirant les liaisons défaillantes de l'agrégation et en ajustant la bande passante. En cas de disponibilité, le lien sera réintégré automatiquement dans l'agrégat.
- **Load Balancing** : Permet la répartition de charge de manière équilibrée sur plusieurs liaisons physiques. Cette répartition peut se faire au travers de différents algorithmes. Le choix de l'algorithme reste à la discrétion de l'administrateur en fonction de sa topologie réseaux et de ses besoins. Ces différents algorithmes s'appuient notamment sur :
 - **MAC source** : Cet algorithme se base sur l'adresse MAC source des paquets pour les répartir sur les liens. Cela peut être utile dans le cadre d'un réseau où un grand nombre de clients se connecte sur un nombre réduit de serveurs.
 - **MAC destination** : Cet algorithme se base sur l'adresse MAC de destination des paquets pour les répartir sur les liens. Dans le cadre d'un réseau où un nombre important de serveurs répondent à plusieurs clients, cet algorithme sera préféré.
 - **MAC source et destination** : Comme son nom l'indique, c'est une fusion des deux algorithmes MAC source et MAC destination. Il est privilégié dans les LAN avec une grande variété de combinaisons de sources et de destinations, pour répartir la charge au mieux.
 - **IP source** : Se base sur les IP sources, afin de répartir la charge entre les liens de l'agrégat. Il permet une répartition optimale dans un environnement où un grand nombre de clients cherche à atteindre un petit nombre de serveurs dans un autre réseau.
 - **IP destination** : Se base sur les IP destinations, afin de répartir la charge. Son utilisation est recommandée dans les topologies disposant de plusieurs sous-réseaux où un nombre élevé de serveurs communique avec un nombre restreint de clients.
 - **IP source et destination** : En se basant sur les deux IP, il permet de répartir la charge de manière stable dans un grand nombre de scénarios. Cela en fait l'algorithme utilisé par défaut dans la majorité des cas.

Pour s'assurer du bon fonctionnement de LACP, plusieurs conditions doivent être remplies. Les interfaces d'une agrégation de liens doivent toutes être dans le même mode de fonctionnement, soit « Access », soit « Trunk » et affectées au même VLAN. De plus, les interfaces doivent toutes disposer de la même configuration physique, à savoir la même vitesse de liens et le même mécanisme de multiplexage. Il est également intéressant de noter que l'IEEE limite le nombre de liens physiques à 16 dont 8 actifs par agrégation. Une fois ces conditions remplies et la configuration effectuée, les paquets LACPDU vont être échangés : si les configurations sont cohérentes, l'agrégation devrait être établie.

PAGP fonctionne lui aussi via un échange de paquets dans l'objectif d'assurer l'agrégation des liens physiques. Les PAGPDU sont composés de la manière suivante :

Version	Flag	Local Device ID	Local Learn Capability	Local Port Hot Standby Priority	Local Sent Port ifindex	Local Group Capability	Local Group ifindex	Partner Device ID	Partner Learn Capability	Partner Port Hot Standby Priority	Partner Sent Port ifindex	Partner Group Capability
Partner Group ifindex	Partner Count	Number of TLVs	TLV	Slow Hello	Auto Mode	Consistent State						

- **Version** : Indique sur 1 octet la version du protocole PAGP.
- **Flags** : Indique sur 1 octet les paramètres de négociation de l'agrégation :

- **Slow Hello** : Indique sur 1 bit la fréquence des paquets Hello.
- **Auto Mode** : Indique sur un 1 bit le mode de configuration du port, la valeur « 1 » représente le mode « Auto » et la valeur « 0 » le mode « Desirable ».
- **Consistent State** : Indique sur 1 bit si le port est actif (1) ou inactif (0).
- **Local Device ID** : Indique sur 6 octets l'identifiant du switch émetteur, généralement l'adresse MAC du switch.
- **Local Learn Capability** : Indique sur 1 octet la capacité du switch émetteur à apprendre ou non les adresses MAC, mais également les modalités d'apprentissage :
 - **0x00** : Le port ne peut pas apprendre les MAC.
 - **0x01** : Le port peut apprendre les MAC en mode standalone uniquement.
 - **0x02** : Le port peut apprendre les MAC en mode agrégé uniquement.
 - **0x03** : Le port peut apprendre les MAC en mode standalone et agrégé.
- **Local Port Hot Standby Priority** : Indique sur 1 octet la valeur du switch émetteur comprise entre 1 (priorité la plus élevée) et 255 (priorité la plus basse) qui permet d'identifier le port actif de l'agrégation. En cas d'égalité, le numéro de port le plus bas sera préféré. A noter que par défaut cette valeur est « 128 », cependant elle peut être personnalisée par l'administrateur.
- **Local Sent Port ifindex** : Indique sur 2 octets le numéro d'identifiant du port sur le switch émetteur. Ce numéro d'identification est attribué par le switch lui-même et ne peut être modifié.
- **Local Group Capability** : Indique sur 4 octets les types d'agrégation supportés par le switch émetteur, automatique, dynamique ou statique.
- **Local Group ifindex** : Indique sur 2 octets la valeur du switch émetteur comprise entre 1 et 65535 qui permet d'identifier le groupe d'agrégation.
- **Partner Device ID** : Indique sur 6 octets l'identifiant du switch partenaire, généralement l'adresse MAC du switch.
- **Partner Learn Capability** : Indique sur 1 octet la capacité du switch partenaire à apprendre ou non les adresses MAC, mais également les modalités d'apprentissage.
- **Partner Port Hot Standby Priority** : Indique sur 1 octet la valeur du switch partenaire comprise entre 1 (priorité la plus élevée) et 255 (priorité la plus basse) qui permet d'identifier le port actif de l'agrégation. En cas d'égalité, le numéro de port le plus bas sera préféré. A noter que par défaut cette valeur est « 128 », cependant elle peut être personnalisée par l'administrateur.
- **Partner Sent Port ifindex** : Indique sur 2 octets le numéro d'identifiant du port sur le switch partenaire. Ce numéro d'identification est attribué par le switch lui-même et ne peut être modifié.
- **Partner Group Capability** : Indique sur 4 octets les types d'agrégation supportés par le switch partenaire, automatique, dynamique ou statique.
- **Partner Group ifindex** : Indique sur 2 octets la valeur du switch partenaire comprise entre 1 et 65535 qui permet d'identifier le groupe d'agrégation.
- **Partner Count** : Indique sur 1 octet le nombre de ports utilisés par l'agrégation de liens.
- **Number of TLVs** : Indique sur 1 octet le nombre de champs TLV.
- **TLV Entry** : Indique des informations complémentaires si besoin.

Tout comme pour LACP, les champs « Local » (équivalent de Actor) sont remplis dès le premier échange PAgPDU, contrairement aux champs « Partner » qui ne sont remplis qu'à partir du deuxième échange PAgPPU. Les messages PAgPDU sont ensuite échangés périodiquement et mis à jour en fonction des événements réseaux.

Télécharger le paquet Wireshark PAgPDU

A la différence de LACP, pour PAgP les ports de l'agrégation peuvent être configurés en trois modes :

- **On** : Ce mode force l'échange de paquets PAgPDU pour l'établissement de l'agrégation de liens. Cette négociation entièrement manuelle, doit être cohérente des deux côtés.
- **Auto** : Ce mode attend le premier paquet PAgPDU d'un switch configuré en mode Desirable. Ce mode de fonctionnement est donc semblable au mode passif de LACP.
- **Desirable** : Ce mode envoie des paquets PAgPDU pour initier la négociation de l'agrégation, cela est semblable au mode Actif de LACP.

Vous pouvez voir ci-dessous le tableau de négociation des agrégations de liens pour PAgP :

Port équipement A	Port équipement B	Négociation
On	On	Agrégation forcée établie
On	Auto	Aucune agrégation
On	Desirable	Aucune agrégation
Auto	Auto	Aucune agrégation
Auto	Desirable	Agrégation PAgP établie
Desirable	Auto	Agrégation PAgP établie
Desirable	Desirable	Agrégation PAgP établie

De la même manière que LACP, PAgP dispose de deux fréquences d'envoi des messages, Fast qui envoie un message toutes les 1 seconde et Slow qui envoie un message toutes les 30 secondes.

Globalement, les fonctionnalités avancées de PAgP sont semblables à LACP, avec l'intégration du Failover et le Load Balancing.

De la même manière que LACP, le fonctionnement du protocole PAgP est conditionné par une correspondance du mécanisme de multiplexage et de la vitesse des interfaces. De même le mode de fonctionnement des interfaces doit être similaire, soit en mode « Access » soit en mode « Trunk » et affecté au même VLAN. De plus, le nombre de liens physiques est limité à 16 dont 8 actifs par agrégation. Une fois ces conditions remplies et la configuration effectuée, les paquets PAgPDU vont être échangés : si les configurations sont cohérentes, l'agrégation devrait être établie.

En dehors des modes de négociation plus nombreux sur PAgP, les deux protocoles permettent les mêmes fonctionnalités. La différence vient principalement du fait que LACP permet de profiter d'un renouvellement de matériel pour changer de fournisseur d'équipement de manière graduelle. En revanche, PAgP nécessitera un changement global en raison de sa compatibilité restreinte avec les équipements CISCO.

La sécurisation de ces deux protocoles présente des points communs, comme notamment le contrôle d'accès via des ACL et/ou la restriction « Port Security » basés sur les MAC. Afin de limiter les risques d'accès de périphériques non légitimes, leur sécurisation passe également par une bonne pratique de configuration. En effet, il faut toujours privilégier les modes Actif/Actif sur LACP et Desirable/Desirable sur PAgP. L'utilisation du mode Auto et/ou Passif, permet à un attaquant de simuler l'envoi de paquets Actif ou/et Desirable afin de perturber la stabilité et l'intégrité de l'agrégation. Cela peut impliquer une perturbation de l'algorithme de répartition de charge, mais également une perturbation de la disponibilité via la désactivation de liens. De plus, la mise en place de la QoS peut limiter l'impact en cas de saturation d'un lien de l'agrégat.

LACP et PAgP peuvent être implémentés sur toutes les couches de l'architecture à 3 niveaux en fonction des besoins et des problématiques. En effet, ils pourront être implémentés pour la redondance et l'augmentation de la bande passante. Cependant, sur les couches « Aggregation » et « Core », l'augmentation de la bande passante ne sera bien souvent pas la raison de son implémentation, du fait de l'utilisation de routeurs haut de gamme disposant de ports optiques 25, 50 ou 100 Go ne rendant pas nécessaire leur utilisation. La redondance de liaison reste intéressante sur toutes les couches de notre architecture.

Dans notre cas, nous allons mettre en œuvre ces deux protocoles pour interconnecter nos switch d'agrégation par paires, afin de s'assurer d'une redondance des liaisons physiques.

QCM LACP et PAgP

a) Lequel de ces 2 protocoles est propriétaire CISCO ?

- ☐ LACP
- ☐ PAgP

b) En plus du mode des ports, que faut-il configurer dans LACP ?

- ☐ La vitesse des messages LACPDU.
- ☐ La longueur des messages LACPDU.
- ☐ Le poids des messages LACPDU.
- ☐ L'origine des messages LACPDU.

c) Les interfaces d'une agrégation (LACP ou PAgP) de liens doivent toutes être dans le même mode de fonctionnement soit « Access » soit « Trunk ».

- ☐ Vrai
- ☐ Faux

d) Quelle est la limite de lien physique de LACP ?

- ☐ 14 dont 7 actifs par agrégation.
- ☐ 16 dont 8 actifs par agrégation.
- ☐ 18 dont 9 actifs par agrégation.
- ☐ 20 dont 10 actifs par agrégation.

e) Dans le cas de PAgP les ports de l'agrégation peuvent être configurés en trois modes, lesquels ?

Correction QCM LACP et PAgP

a) Lequel de ces 2 protocoles est propriétaire CISCO ?

- ☐ LACP
- ☒ PAgP

b) En plus du mode des ports, que faut-il configurer dans LACP ?

- ☒ La vitesse des messages LACPDU.
- ☐ La longueur des messages LACPDU.
- ☐ Le poids des messages LACPDU.
- ☐ L'origine des messages LACPDU.

c) Les interfaces d'une agrégation (LACP ou PAgP) de liens doivent toutes être dans le même mode de fonctionnement soit « Access » soit « Trunk ».

- ☒ Vrai
- ☐ Faux

d) Quelle est la limite de lien physique de LACP ?

- ☐ 14 dont 7 actifs par agrégation.
- ☒ 16 dont 8 actifs par agrégation.
- ☐ 18 dont 9 actifs par agrégation.
- ☐ 20 dont 10 actifs par agrégation.

e) Dans le cas de PAgP les ports de l'agrégation peuvent être configurés en trois modes, lesquels ?

On, Auto, Desirable

[Vidéo de de mise en œuvre LACP et PAgP \(DRIVE\)](#)

[Vidéo de de mise en œuvre LACP et Pag \(Youtube\)](#)

Fiche TP LACP et PAgP

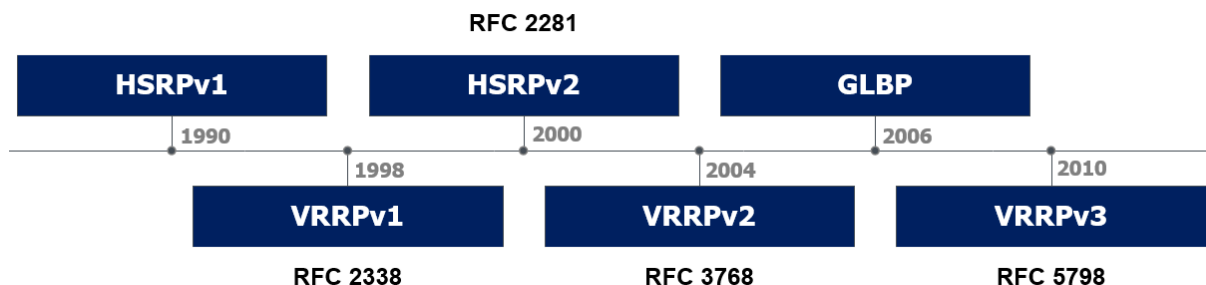
1. Dans la zone « IEEE » sur les switch MLS « DC1-DST1 » et « DC1-DST2 » créer une agrégation de lien LACP avec les interfaces « G1/2 » et « G1/3 » en mode « Active ».
2. Configurer les interfaces virtuelles de l'agrégation de lien LACP en mode « Trunk » pour les VLAN « 100,110,120,130,140 » uniquement, vous devez également désactiver DTP.
3. Configurer LACP pour utiliser l'algorithme de répartition MAC Source et Destination.
4. Dans la zone « CISCO » sur les switch MLS « DC1-DST3 » et « DC1-DST4 » créer une agrégation de lien PAgP avec les interfaces « G1/2 » et « G1/3 » en mode « Desirable ».

5. Activer le mode PAgP Fast.
6. Configurer les interfaces virtuelles de l'agrégation de lien PAgP en mode « Trunk » pour les VLAN « 200,210,220,230,240 » uniquement, vous devez également désactiver DTP.
7. Configurer PAgP pour utiliser l'algorithme de répartition MAC Source et Destination.

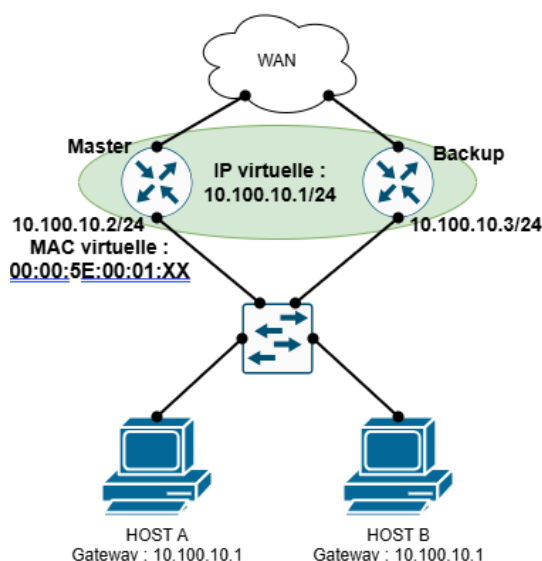
[Télécharger Correction TP LACP et PAgP](#)

4. Qu'est-ce que VRRP et GLBP ?

L'accès aux ressources externes étant indispensable, des protocoles de redondance de Gateway ont été développés dès le début des années 90 avec HSRPv1 (Hot Standby Router Protocol) propriétaire CISCO. Il faudra attendre 1998 pour qu'un standard de l'IETF soit publié dans la RFC 2338, il s'agit de VRRP (Virtual Router Redundancy Protocol). CISCO actualisera son protocole en 2000 avec HSRPv2 en ajoutant des fonctionnalités de sécurité et de gestion des VLAN. Une actualisation du standard VRRPv2 a été publiée dans la RFC 3768 en 2004 afin de prendre en charge de nouvelles fonctionnalités. GLBP (Gateway Load Balancing Protocol) a été proposé par CISCO en 2006 pour permettre une prise en charge du balancement de charge en plus de la redondance de Gateway. VRRP connaîtra sa dernière mise à jour en 2010 avec la RFC 5798 qui intègre la prise en charge d'IPv6.



Le protocole VRRP a pour but d'offrir une redondance de Gateway, afin de garantir une haute disponibilité en définissant un routeur « Master » et un « Backup » au sein d'un groupe de routeurs. Une IP et une MAC virtuelles seront associées au routeur « Master ». En cas de défaillance du « Master », le routeur « Backup » sera élu nouveau « Master » et se verra attribuer cette IP et cette MAC. Un échange de paquets VRRP Advertisement sur l'adresse multicast « 224.0.0.8 » périodique (avec un intervalle de 1 seconde) permet de procéder à l'élection du « Master », mais également de détecter une panne de ce dernier. A noter que la MAC virtuelle est normée sous la forme « 00:00:5E:00:01:XX », la valeur « XX » équivaut à l'identifiant du groupe VRRP en hexadécimal.



Vous trouverez le récapitulatif des états que peut prendre un routeur VRRP dans le tableau suivant :

Etat	Description
Initialize	Le routeur initialise sa configuration, ses timer et sa priorité, il ne participe pas encore activement au groupe VRRP.
Backup	Le routeur surveille le master, sans traiter aucun flux réseaux autre que les paquets VRRP Advertisement.
Master	Le routeur transmet activement le trafic et envoie des paquets VRRP Advertisement.

Nous allons maintenant procéder à l'analyse du paquet **VRRP Advertisement** :

Version	Type	Virtual Router ID (VRID)	Priority	Count IP Addresses	Authentication Type	Advertisement Interval	Checksum	Addresses	Authentication Data
---------	------	--------------------------	----------	--------------------	---------------------	------------------------	----------	-----------	---------------------

- **Version** : Indique sur 4 bits la version du protocole VRRP.
- **Type** : Indique sur 4 bits le type de paquet.
- **Virtual Router ID (VRID)** : Indique sur 8 bits l'identifiant unique du groupe VRRP.
- **Priority** : Indique sur 8 bits la priorité du routeur émetteur pour l'élection du « Master », la valeur par défaut est égale « 100 ». A noter qu'elle peut être personnalisée par l'administrateur entre « 0 » et « 255 ».
- **Count IP Addresses** : Indique sur 8 bits le nombre d'IP associées à ce routeur virtuel.
- **Authentication Type** : Indique sur 8 bits le type d'authentification utilisé, mot de passe, MD5.
- **Advertisement Interval** : Indique sur 8 bits l'intervalle maximal admissible en secondes entre deux messages en provenance du « Master ».
- **Checksum** : Indique sur 16 bits la somme de contrôle d'intégrité du paquet VRRP.
- **Addresses** : Indique la liste des adresses IP virtuelles du groupe de VRRP.
- **Authentication Data** : Indique les données d'authentification VRRP.

[Télécharger le paquet Wireshark VRRP Advertisement](#)

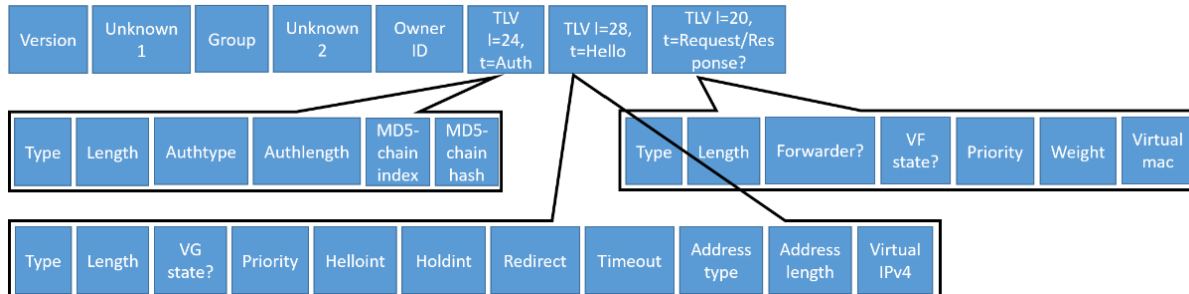
L'échange de ces paquets, permet au routeur de procéder à l'élection du routeur « Master » en se basant sur la priorité, plus elle est élevée plus le routeur sera préféré. En cas d'égalité de priorité, l'adresse IP la plus élevée sera utilisée pour départager les routeurs. Une fois le routeur « Master » élu, des paquets d'annonce sont envoyés toutes les 1 seconde en vue de détecter les pannes. En cas de perte de 3 paquets consécutifs, une nouvelle élection aura lieu. Le concept de préemption permet au routeur de redevenir « Master » après une défaillance temporaire. VRRP ne permet pas une répartition de charge native, contrairement à GLBP. Cependant, il existe un moyen détourné de créer un balancement de charge en créant deux groupes VRRP pour le même réseau IP avec une priorité inversée, puis de répartir les clients manuellement sur les Gateway virtuels.

Il est théoriquement possible de créer 255 groupes VRRP sur un routeur en fonction des ressources matériels de ce dernier. Chaque groupe VRRP peut prendre en charge 255 routeurs, soit 1 « Master » et 254 « Backup ».

Les protocoles HSRP et VRRP sont globalement semblables dans leurs fonctionnalités, les légères différences existantes sont les noms de mode des routeurs « Actif/Passif » au lieu de « Master/Backup », mais également sur les valeurs par défaut des timer.

GLBP est un protocole propriétaire CISCO, qui permet de mettre en œuvre un balancement de charge en plus d'une redondance. Pour garantir ce balancement de charge, le groupe GLBP procède à l'élection de l'AVG (Active Virtual Gateway) et de/des AVF (Active Virtual Forwarder). Cette élection intervient après la configuration du protocole via l'échange de paquets multicast via l'IP « 224.0.0.102 » sur le port UDP 3222 toutes les 3 secondes.

Nous allons maintenant analyser les différents champs qui composent le paquet GLBP :



- **Version** : Indique sur 1 octet la version du protocole GLBP, afin de garantir la compatibilité.
- **Unknown1** : Champ réservé pour une extension de fonctionnalité future du protocole.
- **Group** : Indique sur 1 octet le numéro de groupe GLBP, un groupe représente une seule IP virtuelle. Les valeurs admissibles sont comprises entre 0 et 1023.
- **Unknown2** : Champ réservé pour une extension de fonctionnalité future du protocole.
- **Owner ID** : Indique sur 6 octets l'adresse MAC du routeur AVG.
- **TLV l=24, t=Auth** : TLV regroupant les champs d'authentification ci-dessous.
 - **Type** : Indique sur 1 octet le type de TLV, ici authentification.
 - **Length** : Indique sur 2 octets la taille totale du champ TLV.
 - **Authtype** : Indique sur 1 octet la méthode d'authentification.
 - **Authlength** : Indique sur 1 octet la longueur du hash utilisé.
 - **MD5-chain index** : Indique sur 1 octet la position du hash dans la séquence.
 - **MD5-chain hash** : Indique sur 16 octets la valeur hash, utilisée pour la vérification d'intégrité et l'authentification des paquets GLBP.
- **TLV l=28, t=Hello** : TLV regroupant les champs d'annonce du groupe GLBP ci-dessous.
 - **Type** : Indique sur 1 octet le type de TLV, ici authentification.

- **Length** : Indique sur 2 octets la longueur totale du champ TLV.
- **VG state?** : Indique sur 1 octet l'état du routeur AVG, AVF et Listen.
- **Priority** : Indique sur 1 octet la priorité utilisée pour l'élection de l'AVG et l'AVF.
- **Helloint** : Indique sur 2 octets l'intervalle maximal admissible en millisecondes entre deux messages HELLO.
- **Holdint** : Indique sur 2 octets l'intervalle en millisecondes avant qu'un routeur ne soit considéré comme inactif en l'absence de réponse.
- **Redirect** : Indique sur 2 octets l'intervalle en secondes durant lequel un client peut continuer à requêter un AVF défaillant avant qu'il ne soit redirigé vers un autre AVF.
- **Timeout** : Indique sur 2 octets la durée en secondes pendant laquelle un AVF peut se retrouver en état inactif avant d'être supprimé de la topologie.
- **Address type** : Indique sur 1 octet la version d'IP utilisée, IPv4 ou IPv6.
- **Address length** : Indique sur 1 octet la longueur de l'adresse IP.
- **Virtual IPv4** : Indique sur 4 octets l'adresse IP virtuelle utilisée par les clients pour joindre le groupe GLBP.
- **TLV l=20, t=Request/Response?** : TLV regroupant les champs de coordination des rôles des routeurs du groupe GLBP ci-dessous.
 - **Type** : Indique sur 1 octet le type de TLV, ici coordination des rôles.
 - **Length** : Indique sur 2 octets la longueur totale du champ TLV.
 - **Forwarder?** : Indique sur 1 octet l'identifiant de l'AVF destinataire de la requête.
 - **VF state?** : Indique sur 1 octet l'état du routeur, les états possibles sont « Active », « Standby » et « Inactive ».
 - **Priority** : Indique sur 1 octet la priorité du routeur pour déterminer son rôle AVG ou AVF.
 - **Weight** : Indique sur 2 octets le poids du routeur qui sera utilisé par les algorithmes de balancement de charge.
 - **Virtualmac** : Indique sur 6 octets la MAC virtuelle du routeur qui sera utilisée pour recevoir le trafic.

[Télécharger le paquet Wireshark GLBP](#)

Maintenant que vous maîtrisez la composition des paquets GLBP, nous allons aborder les rôles des routeurs au sein d'un groupe GLBP :

- **AVG** : Le routeur qui porte ce rôle est le routeur principal, c'est lui qui est chargé de répondre aux requêtes ARP sur l'IP virtuelle du groupe GLBP. Cette réponse aux requêtes ARP permet de répartir la charge sur les différents AVF en tenant compte de l'algorithme sélectionné. L'AVG procède à la coordination et à l'attribution des rôles des AVF, tout en surveillant leur état de santé. A noter que l'AVG est également AVF. Dans le cadre d'une défaillance de l'AVG, l'AVG Standby prendra le relais, il agit également comme un AVF.
- **AVF** : Le ou les routeurs AVF servent de Gateway pour les clients qui ont été redirigés vers eux par l'AVG via leur MAC virtuelle. L'AVF répond de manière régulière à l'AVG pour confirmer qu'il est toujours opérationnel, cela permet à l'AVG de rediriger le trafic vers un autre AVF en cas de défaillance.

Comme vous l'avez compris, une fois la configuration effectuée, GLBP procède à l'élection de l'AVG. Pour cela les informations échangées via les paquets GLBP sont utilisées. Le routeur qui dispose de la plus grande priorité sera préféré, en cas d'égalité l'adresse IP la plus haute sera préférée à son tour. A

noter que la priorité peut être personnalisée par l'administrateur entre 1 et 255 (la valeur par défaut est 100). L'AVG Standby est élu selon le même principe que l'AVG avec la seconde priorité la plus haute. Les routeurs AVG et AVF passent par plusieurs états avant d'atteindre leur rôle.

Les routeurs AVG transitent par les 6 états suivants :

1. **Disabled** : Dispose d'une configuration GLBP partielle, il fait partie d'un groupe GLBP mais, ne dispose pas d'une IP pour participer au processus.
2. **Initial** : Dispose d'une configuration IP, cependant les timer et la priorité sont encore en cours d'initialisation.
3. **Listen** : Reçoit les paquets HELLO des autres routeurs AVG et/ou AVG standby, mais il n'envoie pas encore de paquets HELLO.
4. **Speak** : Envoie des paquets HELLO de manière active, pour participer à l'élection de l'AVG et l'AVG standby.
5. **Stanby** : Surveille l'AVG pour prendre le relais en cas de panne de ce dernier.
6. **Active** : Assume le rôle d'AVG de manière active et coordonne les AVF.

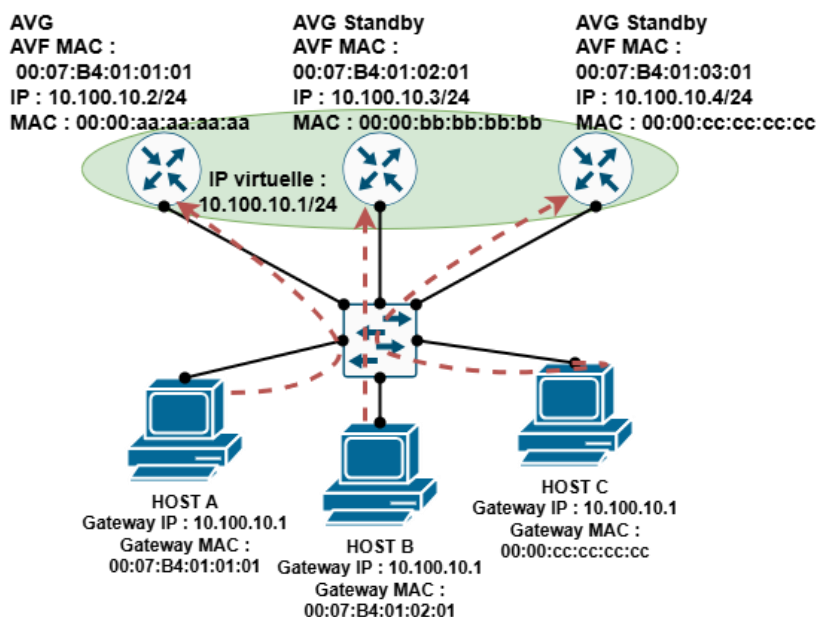
Les routeurs AVF transitent seulement par les 4 états ci-dessous :

1. **Disabled** : Dispose d'une configuration GLBP partielle, il fait partie d'un groupe GLBP mais, n'a pas encore de MAC virtuelle attribuée par l'AVG.
2. **Initial** : Dispose d'une MAC virtuelle et entre dans la phase d'initialisation finale de la configuration.
3. **Listen** : Reçoit les paquets HELLO de l'AVG et se prépare à prendre le rôle AVF à la demande de l'AVG.
4. **Active** : Agit comme un AVF actif et participe à l'équilibrage de charge en traitant le trafic du sous-ensemble de clients que l'AVG lui adresse.

Le protocole GLBP peut théoriquement supporter 1024 groupes, bien que certains routeurs d'entrée de gamme ne disposent pas des ressources matériels suffisantes pour gérer autant de groupes GLBP. Chaque groupe GLBP peuvent contenir 1024 routeurs, cependant seulement 4 routeurs peuvent être AVF actifs en même temps.

L'adresse IP virtuelle est définie par l'administrateur, contrairement à l'adresse MAC qui est générée de la manière suivante : « 00:07:B4:XX:YY:ZZ »

Partie	Description
00:07:B4	Préfix OUI (Organizationally Unique Identifier) attribué à CISCO pour GLBP.
XX	Numéro du groupe GLBP compris entre 0 et 1023.
YY	L'identifiant de l'AVF compris entre 0 et 4.
ZZ	Généré aléatoirement pour éviter les conflits sur le réseau.



Abordons maintenant les trois algorithmes de répartition de charge disponibles sur GLBP :

- **Round Robin** : Transmet les paquets de manière cyclique à chaque AVF à tour de rôle, cela assure une répartition totalement équitable. Cet algorithme est recommandé quand les AVF ont tous la même performance de routage.
- **Weighted** : Transmet les paquets en tenant compte du poids des AVF définis par l'administrateur (entre 1 et 255), plus le poids est élevé plus le routeur transmettra de trafic. A noter que la valeur « Weighted » par défaut est égale à 100. Cet algorithme est préféré quand les AVF n'ont pas tous les mêmes performances de routage, cela permet de maîtriser plus finement la charge envoyée à chaque routeur.
- **Host-Dependent** : Transmet le paquet aux AVF en fonction du client émetteur. Chaque client se voit attribuer un AVF via le hash de son IP. Cela permet de limiter les fluctuations de routage, ce qui sera préféré pour des services qui sont sensibles à ces fluctuations, comme la VOIP.

Comme tous les services et protocoles, de nos jours la sécurisation de ces protocoles est un point essentiel. La sécurisation de nos protocoles passe par deux grands axes : l'authentification et le filtrage. En effet, les protocoles GLBP et VRRP (jusqu'à la version 2 pour VRRP) permettent une authentification de leurs échanges via un mot de passe, le hachage MD5 et/ou SHA256. De plus, comme pour de nombreux protocoles, les ACL peuvent être implémentées afin de filtrer les paquets VRRP et GLBP. L'utilisation de ces deux méthodes permet de limiter grandement le risque d'usurpation d'un routeur illégitime, mais également la compromission des paquets.

Ces deux protocoles s'intègrent parfaitement sur la couche « Aggregation » de l'architecture à trois niveaux. En effet, la couche « Aggregation » étant constituée d'une paire de MLS qui assure le routage inter-VLAN pour les switch de la couche « Access », la mise en place de VRRP ou GLBP offrira une redondance de Gateway. Le choix entre ces deux protocoles se fait en fonction du fournisseur de matériels réseaux et du besoin de balancement de charge. Si vous ne souhaitez pas être dépendant d'un seul constructeur, VRRP sera un choix plus raisonnable. Au contraire, si la dépendance à CISCO ne vous gêne pas et que le balancement de charge est impératif pour vous, GLBP sera un très bon choix.

Dans le cadre de notre infrastructure, nous mettrons en œuvre GLBP sur une paire de MLS de la couche « Aggregation » et VRRP sur une autre paire de MLS de la couche « Aggregation », cela nous permettra d'opposer leur configuration.

QCM VRRP et GLBP

a) Le protocole VRRP a pour but d'offrir une redondance de Gateway ?

- ☐ Vrai
☐ Faux

b) Cochez l'adresse multicast qu'utilise VRRP pour ses paquets « Advertissement ».

- ☐ 224.0.0.8
☐ 225.0.0.8
☐ 224.0.0.7
☐ 225.0.0.8

c) Pour garantir ce balancement de charge, le groupe GLBP procède à l'élection de l'AVG (Active Virtual Gateway) et de ou des AVF (Active Virtual Forwarder). Cette élection intervient après la configuration du protocole via l'échange de paquets multicast. Sur quelle adresse et port sont diffusés ces paquets multicast ?

- ☐ 224.0.0.102 port UDP 3222
☐ 224.0.0.102 port TCP 3222
☐ 224.0.0.103 port TCP 3223
☐ 224.0.0.103 port UDP 3223

d) GLBP est un protocole propriétaire CISCO, qui permet de mettre en œuvre un balancement de charge en plus d'une redondance.

- ☐ Vrai
☐ Faux

e) L'état backup de VRRP se définit par la phrase suivante « Le routeur surveille le master, sans traiter aucun flux réseaux autre que les paquets VRRP Advertissement. »

- ☐ Vrai
☐ Faux

Correction QCM GLBP et VRRP

a) Le protocole VRRP a pour but d'offrir une redondance de Gateway ?

- ☒ **Vrai**
☐ Faux

b) Cochez l'adresse multicast qu'utilise VRRP pour ses paquets « Advertissement ».

- ☒ **224.0.0.8**
☐ 225.0.0.8
☐ 224.0.0.7
☐ 225.0.0.8

c) **Pour garantir ce balancement de charge, le groupe GLBP procède à l'élection de l'AVG (Active Virtual Gateway) et de ou des AVF (Active Virtual Forwarder). Cette élection intervient après la configuration du protocole via l'échange de paquets multicast. Sur quelle adresse et port sont diffusés ces paquets multicast ?**

☒ **224.0.0.102 port UDP 3222**

☐ 224.0.0.102 port TCP 3222

☐ 224.0.0.103 port TCP 3223

☐ 224.0.0.103 port UDP 3223

d) **GLBP est un protocole propriétaire CISCO, qui permet de mettre en œuvre un balancement de charge en plus d'une redondance.**

☒ **Vrai**

☐ Faux

e) **L'état backup de VRRP se définit par la phrase suivante « Le routeur surveille le master, sans traiter aucun flux réseaux autre que les paquets VRRP Advertisement. »**

☒ **Vrai**

☐ Faux

[Vidéo de de mise en œuvre VRRP et GLBP \(DRIVE\)](#)

[Vidéo de de mise en œuvre VRRP et GLBP \(Youtube\)](#)

Fiche TP VRRP et GLBP

1. Dans la zone « IEEE » sur les switch « DC1-DST1 » et « DC1-DST2 » créer et configurer les groupes VRRP suivants :

Groupe	Interface	IP virtuelle	Priorité DC1DST1	Priorité DC1DST2	Authentification	Préemption	Timer
100	Vlan 100	192.168.100.1	150	100	Mot de passe	Oui	1s
110	Vlan 110	192.168.110.1	150	100	Mot de passe	Oui	1s
120	Vlan 120	192.168.120.1	150	100	Mot de passe	Oui	1s
130	Vlan 130	192.168.130.1	100	150	Mot de passe	Oui	1s
140	Vlan 140	192.168.140.1	100	150	Mot de passe	Oui	1s

2. Dans la zone « CISCO » sur les switch « DC1-DST3 » et « DC1-DST4 » créer deux clés de chiffrement « GLBP-VLAN-ADM » et « GLBP-VLAN-CLT » en SHA256.

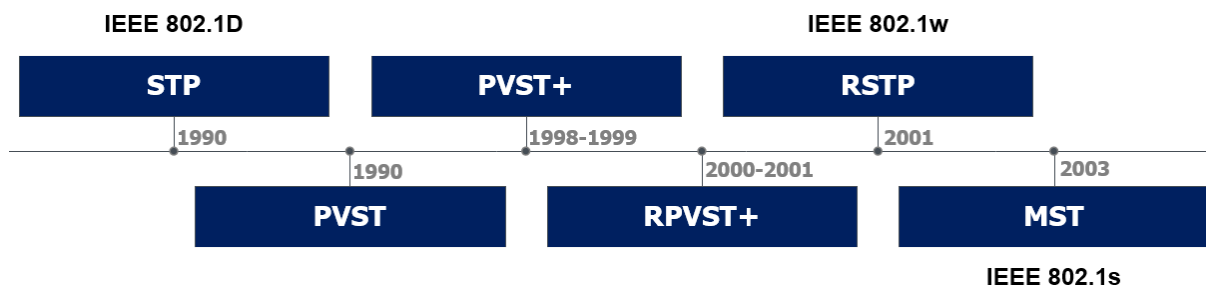
3. Dans la zone « CISCO » sur les switch « DC1-DST3 » et « DC1-DST4 » créer et configurer les groupes GLBP suivants :

Groupe	Interface	IP virtuelle	Priorité DC1DST3	Priorité DC1DST4	Authentification	Préemption	Timer	Algorithme
200	Vlan 200	192.168.200.1	150	100	Clé GLBP-VLAN-ADM	Oui	1s	round-robin
210	Vlan 210	192.168.210.1	150	100	Clé GLBP-VLAN-CLT	Oui	1s	round-robin
220	Vlan 220	192.168.220.1	150	100	Clé GLBP-VLAN-CLT	Oui	1s	round-robin
230	Vlan 230	192.168.230.1	150	100	Clé GLBP-VLAN-CLT	Oui	1s	round-robin
240	Vlan 240	192.168.240.1	150	100	Clé GLBP-VLAN-CLT	Oui	1s	round-robin

[Télécharger Correction TP VRRP et GLBP](#)

5. Qu'est-ce que MSTP et RPVST+ ?

RPVST+ et MSTP sont tous les deux des protocoles dérivés de STP (Spanning-Tree Protocol). En effet le protocole STP standardisé par l'IEEE en 1990 sous le 802.1D, a connu de nombreuses déclinaisons standardisées et propriétaires en vue d'offrir des fonctionnalités de plus en plus performantes. Parmi elles, nous retrouvons notamment PVST (propriétaire CISCO), PVST+ (propriétaire CISCO), RPVST+ (propriétaire CISCO), RSTP (IEEE 802.1w), MSTP (IEEE 802.1s). Afin de pouvoir comprendre les différents intérêts de toutes ces déclinaisons, il est indispensable au préalable de maîtriser STP.

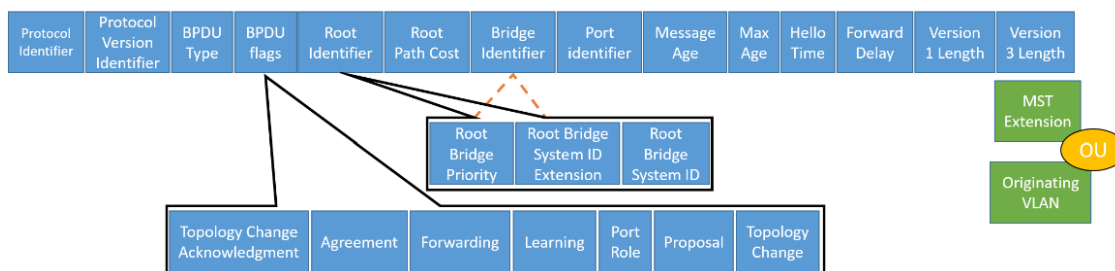


Ce protocole permet d'empêcher les boucles de manière dynamique, afin d'interconnecter des équipements avec plusieurs chemins possibles, le but étant d'assurer la redondance en cas de défaillance d'un lien ou d'un équipement, tout en évitant une boucle réseaux qui conduirait à une tempête de diffusion.

Une fois configuré Spanning-Tree va échanger des paquets BPDU (Bridge Protocol Data Units) afin de déterminer et maintenir à jour les différents chemins réseaux. Il existe deux types de BPDU :

- **Configuration BPDU**, qui crée et maintient la topologie avec tous les chemins possibles sans boucle.
- **Topology Change Notification BPDU**, qui met à jour la topologie en cas de changement.

Nous allons aborder et expliquer les différents champs qui composent un paquet BPDU Spanning-Tree :



- **Protocol Identifier** : Champ sur 2 octets identifiant le protocole utilisé.
- **Protocol Version Identifier** : Champ sur 2 octets identifiant la version du protocole utilisé, la valeur « 0 » pour SPT, la valeur « 2 » pour RSTP et la valeur « 3 » pour MSTP.
- **BPDU Type** : Champ sur 1 octet pour identifier le type de BPDU, soit les « Configuration BPDU », soit les « TCN BPDU ». Le paquet « Configuration BPDU » a toujours la valeur « 0x00 » peu importe la version du protocole utilisé, à l'inverse du « TCN BPDU » dont la valeur varie en fonction de la version du protocole. La valeur « 0x01 » pour STP et la valeur « 0x02 » pour RSTP et MSTP.
- **BPDU flags** : Champ sur 1 octet fournissant des informations sur l'état du port émetteur et la topologie.
 - **Topology Change Acknowledgment** : bit d'accusé de réception d'un changement topologique.
 - **Agreement** : bit de validation de synchronisation.
 - **Forwarding** : bit d'indication sur la politique de transfert de trame.
 - **Learning** : bit d'indication sur la politique d'apprentissage MAC.
 - **Port Role** : bit d'indication sur le rôle du port (Root, Designated, Blocking).
 - **Proposal** : bit d'indication sur les propositions de changement du rôle du port.
 - **Topology Change** : bit d'indication d'un changement de topologie.
- **Root Identifier** : Regroupe les trois champs ci-dessous :
 - **Root Bridge Priority** : Indique sur 2 octets la valeur de priorité du switch comprise entre 1 (priorité la plus élevée) et 65535 (priorité la plus basse) qui permet l'élection du « Root Bridge » (concept expliqué plus tard). Cette valeur est configurable par l'administrateur pour optimiser les flux, les valeurs doivent toujours être des multiples de 4096. A noter que la valeur par défaut est « 32768 ».
 - **Root Bridge System ID Extension** : Indique sur 1 octet, une valeur pour l'identification des instances dans le cadre de l'utilisation de MSTP, pour STP cette valeur est à « 0 ».
 - **Root Bridge System ID** : Indique sur 6 octets l'adresse MAC qui permet d'identifier le « Root Bridge ».
- **Root Path Cost** : Indique sur 4 octets le coût du chemin vers le « Root Bridge », dans le but de déterminer le chemin préféré.
- **Bridge Identifier** : Regroupe les trois champs ci-dessous :
 - **Bridge Priority** : Indique sur 2 octets la valeur de priorité du switch émetteur comprise entre 1 (priorité la plus élevée) et 65535 (priorité la plus basse) qui permet l'élection du « Root Bridge ». En cas d'égalité entre le switch émetteur et le « Root bridge » l'adresse MAC sera utilisée pour l'élection.

- **Bridge System ID Extension** : Comme pour le « Root Bridge » ce champ indique sur 1 octet, une valeur pour l'identification des instances dans le cadre de l'utilisation de MSTP, pour STP cette valeur est à « 0 ».
- **Bridge System ID** : Indique sur 6 octets l'adresse MAC qui permet d'identifier le switch émetteur.
- **Port identifier** : Indique sur 2 octets l'identifiant unique du port du switch émetteur. Cet identifiant est composé du numéro de port et de la priorité du port. A l'inverse du numéro de port, la priorité du port peut être personnalisée par l'administrateur de 0 à 240, sous condition d'être un multiple de 16. A noter que la valeur par défaut est « 128 ».
- **Message Age** : Indique sur 2 octets le temps en secondes depuis l'envoi du message par le « Root Bridge ».
- **Max Age** : Indique sur 2 octets la durée de validité du BPDU avant qu'il ne soit considéré comme obsolète.
- **Hello Time** : Indique sur 2 octets l'intervalle entre les BPDU, afin d'assurer que la topologie soit mise à jour régulièrement.
- **Forward Delay** : Indique sur 2 octets le délai en secondes durant lequel un port est en état « Listening » et « Learning » avant de passer en « Forwarding ». Cela permet d'éviter les boucles réseaux durant un changement topologique.
- **Version 1 Length** : Indique sur 1 octet la longueur des données spécifiques à STP et RSTP.
- **Version 3 Length** : Indique sur 1 octet la longueur des données spécifiques à MSTP.
- **« MST Extension » ou « Originating VLAN (PVID) »** : sont des champs avec plusieurs valeurs dédiées aux diverses déclinaisons de STP, ils seront détaillés plus tard dans ce cours.

Maintenant que vous connaissez mieux la composition des BPDU, nous allons aborder les mécanismes de fonctionnement de STP. Afin de créer une topologie sans boucle, Spanning-Tree doit en premier lieu sélectionner un switch qui sera le point central de notre topologie, le « Root Bridge ». L'élection du « Root Bridge » se fait via la sélection du switch qui dispose du plus faible « Bridge ID ». Cet ID est composé du « Bridge Priority » et du « Bridge System ID » dans le cadre de STP, le champ « Bridge System ID Extension » est utilisé uniquement dans le cadre de RSTP et MSTP. A noter que ces champs ont été précédemment détaillés dans la présentation du BPDU.

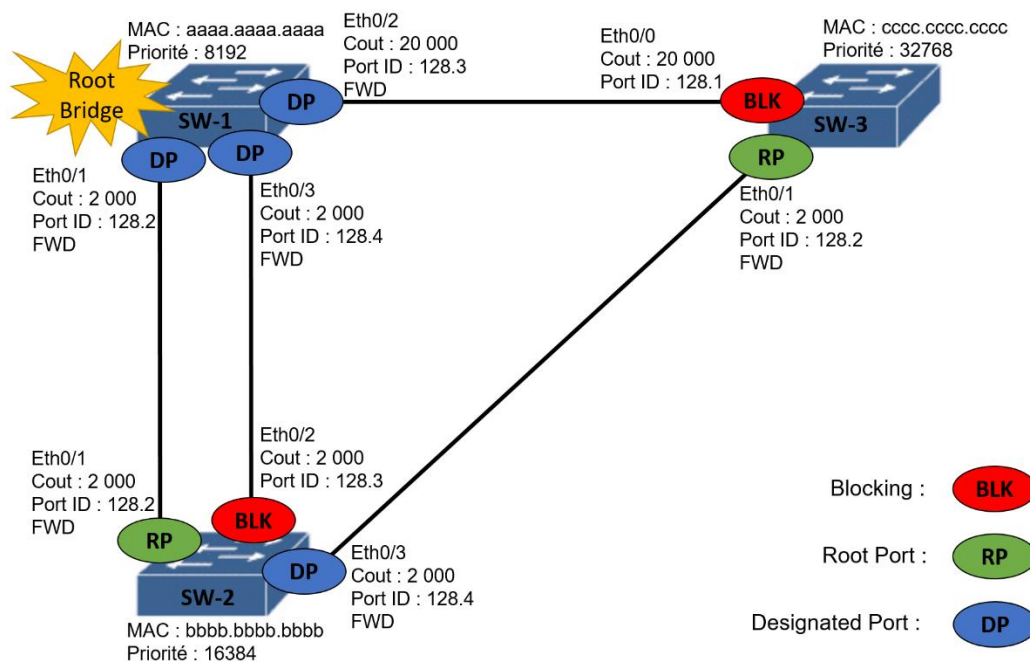
Une fois le « Root Bridge » élu, chaque switch (sauf le Root Bridge élu) procède à la sélection d'un « Root Port », ce port est celui qui dispose du chemin avec le coût le plus faible (le Root Path Cost) à destination du « Root Bridge ». Le « Root Path Cost » est incrémenté par le coût de la liaison réseaux traversée à chaque arrivée sur un nouveau switch. Le coût de la liaison peut être personnalisé par l'administrateur pour influencer le choix du « Root Port ». En l'absence de modification par l'administrateur, le coût dépend de la bande passante, vous pouvez trouver ci-dessous le tableau de correspondance défini par IEEE 802.1D en 2004 :

Bande passante	Coût de la liaison réseaux
10 Mbps	2 000 000
100 Mbps	200 000
1 Gbps	20 000
10 Gbps	2 000
100 Gbps	200
1 Tbps	20

Dans le cas d'une égalité de coût entre deux liaisons réseaux, le « Bridge ID » le plus faible sera préféré. Si les « Bridge ID » venaient eux aussi à être égaux, le numéro de port le plus faible sera préféré. A noter que le « Root Port » est toujours en état Forwarding.

Une fois le « Root Port » choisi, le « Designated Port » doit-être sélectionné sur chacune des liaisons réseaux. Tout comme le « Root Port », le port préféré pour être « Designated Port » est celui qui disposera du « Root Path Cost » le plus faible, en cas d'égalité le « Bridge ID » et le numéro de port pourront servir à les départager. A noter que le « Designated Port » est également toujours en état « Forwarding ».

Rôle du Port	État	Fonction
Root Port	Forwarding	Transmet le flux réseaux vers le Root Bridge.
Designated Port	Forwarding	Fait office de passerelle sur la liaison réseaux.



Comme vous l'avez sans doute compris, les ports peuvent se trouver dans plusieurs états. On dénombre cinq états de port que vous pouvez retrouver dans le tableau ci-dessous :

Etat du port	Transfert de données	Réception BPDU	Envoi BPDU	Apprentissage MAC	Etat
Blocking	NON	OUI	NON	NON	Stable
Listening	NON	OUI	OUI	NON	Temporaire
Learning	NON	OUI	OUI	OUI	Temporaire
Forwarding	OUI	OUI	OUI	OUI	Stable
Disabled	NON	NON	NON	NON	Stable

En cas de changement sur la topologie réseaux, la topologie Spanning-Tree est, elle aussi, impactée. Par exemple, dans le cas d'une défaillance d'un lien, le chemin vers le « Root Bridge » sera recalculé. Si nous nous référons au schéma ci-dessus, en cas défaillance du lien actif entre switch 1 et switch 2, le port « Eth0/2 » sur switch 2 prendra 20 secondes pour passer du mode « blocking » à « Listening ». Le port va ensuite rester durant 15 secondes dans cet état, recevoir les BPDU, les interpréter et mettre à jour la topologie du switch. Il entrera ensuite en mode « Learning » durant un second cycle de 15 secondes,

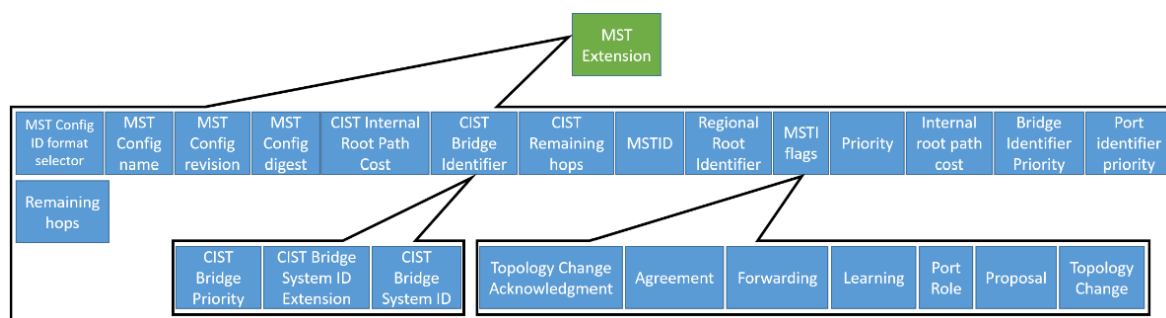
au cours duquel il apprendra les adresses MAC, pour finir par devenir le nouveau « Root Port » du switch 2 et entrer en mode « Forwarding ». Dans le cas d'une défaillance, STP prendra donc environ 50 secondes pour mettre en œuvre un nouveau chemin.

Ce processus de 50 secondes est également déclenché à la connexion d'un poste utilisateur sur la couche « Access ». C'est pourquoi il est fortement conseillé d'activer la fonctionnalité « PortFast » sur les interfaces destinées aux postes clients. Attention, cette option rend STP inopérant sur les ports concernés, l'option doit être uniquement utilisée sur les interfaces clients.

La vitesse de convergence de Spanning-Tree étant trop lente pour répondre aux besoins réseaux de nos jours, le protocole RSTP qui fonctionne de la même manière avec des temps de convergence plus rapide (maximum 6 secondes) sera préféré.

Maintenant que vous maîtrisez STP, nous allons aborder MSTP et RPVST+ qui ont un fonctionnement de base semblable à STP. Cependant, ils intègrent, une fonctionnalité de répartition de charge, mais également, tout comme RSTP, une vitesse de convergence plus rapide.

Je vous propose de continuer, par l'analyse des champs spécifiques au **BPDU de Multiple Spanning-Tree** ; en effet nous détaillerons ici uniquement les champs contenus dans « MST Extension », les autres champs étant communs avec STP :



- **MST Config ID format selector** : Indique sur 1 octet le format de présentation des informations de configuration MSTP, la valeur à « 0 » signifie une présentation standardisée.
- **MST Config name** : Indique le nom d'identification de la région MSTP, tous les switch de la région doivent avoir le même nom.
- **MST Config revision** : Indique sur 2 octets le numéro de version de configuration, ce numéro doit être le même sur tous les switch de la région. Il est à incrémenter de manière manuelle par l'administrateur au moment de la révision de la configuration.
- **MST Config digest** : Indique le Hachage des paramètres de configuration de la région MSTP, en vue de vérifier l'unicité de cette configuration sur les switch de la région.
- **CIST Internal Root Path Cost** : Indique sur 4 octets le coût du chemin vers le « Root Bridge » de l'arbre Spanning-Tree commun ou CIST (notion expliquée ci-dessous).
- **CIST Bridge Identifier** : Indique l'identifiant du « Bridge CIST », cet identifiant est composé des valeurs suivantes :
 - **CIST Bridge Priority** : Indique sur 2 octets la valeur de priorité du switch CIST comprise entre 1 (priorité la plus élevée) et 65535 (priorité la plus basse) qui permet l'élection du « Root Bridge CIST ».
 - **CIST Bridge Identifier System ID Extension** : Indique sur 1 octet, une valeur pour l'identification des instances dans le cadre de l'utilisation de MSTP.
 - **CIST Bridge Identifier System ID** : Indique sur 6 octets l'adresse MAC qui permet d'identifier le switch CIST.

- **CIST Remaining hops** : Indique sur 1 octet le nombre de sauts restant au sein du CIST avant que le BPDU ne soit abandonné, ce champ est décrémenté au fur et à mesure. Cela permet de limiter la propagation des BPDU, afin de maîtriser l'impact sur la charge réseau.
- **MSTID** : Indique sur 2 octets l'identifiant de l'instance MSTI.
- **Regional Root Identifier** : Indique sur 8 octets l'identifiant du « Root Bridge » régional de l'instance.
- **MSTI flags** : Champ sur 1 octet fournissant des informations sur l'état du port émetteur et la topologie au sein de l'instance.
 - **Topology Change Acknowledgment** : bit d'accusé de réception d'un changement topologique.
 - **Agreement** : bit de validation de synchronisation.
 - **Forwarding** : bit d'indication sur la politique de transfert de trame.
 - **Learning** : bit d'indication sur la politique d'apprentissage MAC.
 - **Port Role** : bit d'indication sur le rôle du port (Root, Designated, Blocking, Alternate, Backup, Master).
 - **Proposal** : bit d'indication sur les propositions de changement du rôle du port.
 - **Topology Change** : bit d'indication d'un changement de topologie.
- **Priority** : Indique sur 1 octet la priorité de cette instance MSTP.
- **Internal root path cost** : Indique sur 4 octets le coût du chemin à destination du « Root Bridge » régional pour cette instance MSTP.
- **Bridge Identifier Priority** : Indique la priorité du switch dans cette instance MSTP.
- **Port identifier priority** : Indique la priorité utilisée dans le choix du rôle du port au sein de l'instance MSTP.
- **Remaining hops** : indique sur 1 octet le nombre de sauts restant au sein de l'instance avant que le BPDU ne soit abandonné.

[Télécharger le paquet Wireshark MSTP BPDU](#)

Le BPDU MSTP, en dehors des champs « MST Extension », est composé exactement des mêmes champs que les BPDU de STP et RSTP, cela permet d'assurer une rétrocompatibilité.

Comme vous avez pu le constater en analysant les champs du BPDU, MSTP ajoute le concept d'instance (MSTI), ces instances sont configurées par l'administrateur pour regrouper un certain nombre de VLAN. A noter qu'en l'absence de configuration, tous les VLAN seront placés dans l'IST (Internal Spanning-Tree) qui est l'instance STP par défaut. Le protocole va calculer et déterminer un arbre Spanning-Tree par instance, l'administrateur pourra influencer de manière ciblée les topologies de chacune des instances pour répartir la charge réseaux. Les instances sont toutes déclarées dans un seul BPDU en vue de garantir la cohérence tout en limitant le nombre de paquets circulant sur les réseaux.

En plus du concept d'instance, MSTP inclut un concept de région. Une région regroupe tous les switch qui ont la même configuration pour les valeurs « MST Config name », « MST Config révision » et « MST Config digest ». Les instances MSTP sont attachées à une seule région et ne peuvent pas s'étendre au dehors. Les switch d'une région partagent donc la même vision de la topologie et de la répartition des VLAN. Chacune des régions dispose de sa topologie avec notamment un « CIST Regional Root » qui est le point d'interconnexion entre les régions MSTP et leurs diverses instances avec le CIST.

Le CIST ou « Common and Internal Spanning Tree » est la topologie Spanning par défaut qui permet l'interconnexion des topologies MSTP régionales entre elles, mais également la rétrocompatibilité avec les topologies STP et RSTP. Afin d'assurer cette interconnexion sans boucles réseaux, l'arbre CIST dispose également d'un « CIST Root Bridge » point central de la topologie. A noter que le CIST n'a pas de vision détaillée des topologies de chaque région, pour vulgariser, considérer que le CIST voit une région comme un seul switch.

A la différence de Spanning-Tree, MSTP procède à plusieurs élections. Pour commencer, le « CIST Root Bridge » est élu en fonction de « Priorité Bride Identifier » la plus faible et de l'adresse MAC la plus basse en cas d'égalité. Dans un deuxième temps, le « CIST Regional Root » est élu sur la base du chemin le plus court vers « CIST Root Bridge », en cas d'égalité le switch avec le « Bride Identifier » le plus faible sera sélectionné. La dernière election est celle du « Root Bridge » de chacune des instances MSTP, cette election se déroule de la même manière que l'élection du « Root Bridge » dans STP abordée précédemment.

MSTP inclut trois nouveaux rôles de port, de plus que STP, qui sont récapitulés ci-dessous :

Rôle du port	Etat	Fonction
Master Port	Forwarding	Transmet les flux entre la région MSTP et le CIST Root.
Bakcup Port	Blocking	Port de secours sur un même segment réseau pour transmettre les flux en cas de défaillance d'un Designated Port.
Alternate Port	Blocking	Port de secours pour transmettre les flux en cas de défaillance d'un Root Port.

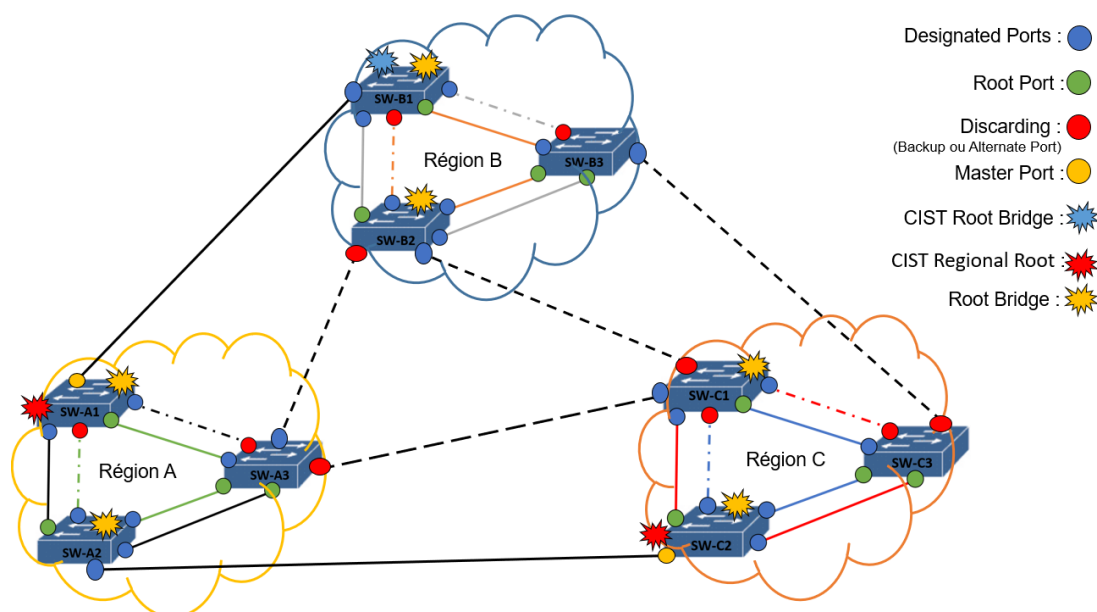
La sélection du « Master Port » repose sur le port qui a le coût de chemin vers le « CIST Root Bridge » le plus faible ; en cas d'égalité entre deux ports, le « Bridge Identifier » du switch connecté aux ports concurrents le plus faible sera préféré. Dans le cas d'une seconde égalité, c'est le switch avec le « Port Identifier » le plus bas qui sera utilisé pour les départager.

Pour ce qui est de la sélection du « Backup Port », il est choisi parmi les ports concurrents au « Designated port » sélectionnés sur un même segment réseau. Dans le cas où plusieurs ports peuvent prétendre à être « Backup Port », le « Bridge Identifier » le plus faible, ou à défaut le « Port Identifier » le plus faible sera préféré.

Le choix du « Alternate Port » se fait par le coût de chemin vers le « Root Bridge » le plus faible. Dans le cas de deux ports concurrents, le « Bridge Identifier » et le « Port Identifier » seront utilisés pour les départager.

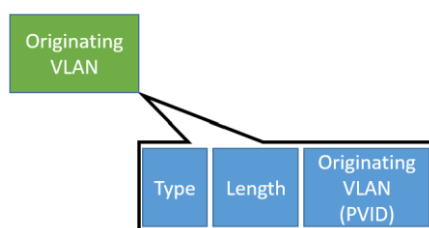
A noter que MSTP et RSTP ne disposent pas de l'état de port « Blocking » et « Listening », qui ont été fusionnés et remplacés par l'état « Discarding ». La fonctionnalité « PortFast » pour les interfaces reliant les clients existe également dans le cadre de MSTP.

Etat du port	Transfert de données	Réception BPDU	Envoi BPDU	Apprentissage MAC	Etat
Discarding	NON	OUI	OUI	NON	Stable



Je vous propose d'aborder la version propriétaire CISCO du protocole RPVST+.

Nous allons commencer par analyser les champs spécifiques au BPDU de RPVST+ :



- **Originating VLAN** : Contient les 3 valeurs suivantes.
 - **Type** : Indique sur 2 octets le type de valeur.
 - **Length** : Indique sur 2 octets la longueur de la valeur
 - **Originating VLAN (PVID)** : Indique sur 2 octets le numéro de VLAN émetteur de ce BPDU.

[Télécharger le paquet Wireshark RPVST+ BPDU](#)

A noter que les autres champs du paquet RPVST+ sont semblables au paquet STP présenté précédemment. La fonctionnalité « PortFast » pour les interfaces reliant les clients existe également dans le cadre de RPVST+.

Comme vous pouvez le constater, à la différence de MSTP, RPVST+ n'intègre pas le concept d'instance qui regroupe plusieurs VLAN, il crée un arbre Spanning-Tree par VLAN. La constitution de l'arborescence RPVST+ s'effectue de manière identique à l'arborescence STP, en intégrant la vitesse de convergence de RSTP.

Comme tous protocoles réseaux, la sécurisation de MSTP et RPVST+ est une étape majeure. Cette sécurisation passe par une sécurisation des protocoles en eux-mêmes, mais aussi par une gestion rigoureuse de l'accès aux réseaux.

Il existe plusieurs mécanismes de gestion des paquets BPDU permettant la mise en place d'un axe de sécurité.

Le BPDU Guard est l'un des mécanismes de sécurisation directement intégré dans les protocoles dérivés de STP (propriétaires ou non). Il est à configurer sur les interfaces clients qui sont reliées à des périphériques finaux, ces interfaces sont en principe en mode « Portfast ». Ce mécanisme désactive les interfaces sur lequel il est configuré en cas de réception d'un BPDU. Cependant, il est à noter que le port en mode BPDU Guard, qui n'est pas désactivé, continue d'émettre des BPDU comme tous les autres ports Spanning-Tree. Ce mécanisme offre une sécurisation contre les tentatives de détournement de l'arbre Spanning-Tree volontaires de la part d'un acteur malveillant mais également, accidentelles si un utilisateur décide de son propre chef de brancher un petit switch dans son bureau. Cela n'empêche en rien un acteur malveillant de procéder à une reconnaissance de la topologie en analysant les BPDU avec WireShark ou TCPDUMP.

BPDU Filter est un second mécanisme de sécurisation intégré à STP et ses variantes. Ce mode est activé par l'administrateur sur les interfaces mode « Portfast », soit au cas par cas, soit de manière globale. A noter que le comportement de BPDU Filter est impacté par la manière dont il a été configuré. Dans le cadre d'une configuration globale, l'envoi des BPDU est désactivé, cependant en cas de réception d'un BPDU, le BPDU Filter et le Portfast seront désactivés sur le port récepteur. Dans le cadre d'une configuration manuelle, les ports ne peuvent ni recevoir ni envoyer de BPDU. Tout comme BPDU Guard cela permet de protéger la topologie Spanning-Tree de toute modification, mais aussi d'empêcher la reconnaissance topologique. Cela réduit également le nombre de paquets inutiles sur le réseau. Il est important de souligner que cette couche de sécurité doit être implémentée avec précaution, en effet l'implémenter sur un lien reliant deux switch peut créer une boucle réseau.

Le Root Guard est un mécanisme qui est configuré par l'administrateur sur les ports souhaités. Une fois mis en place, ce mécanisme rejette les BPDU qui disposent d'une « Root Bridge Priority » favorable à l'émetteur, en vue d'empêcher un changement de « Root Bridge ». Le port sera donc placé dans le mode « root-inconsistent » afin de ne pas impacter la topologie. Cela est utile pour s'assurer qu'aucun acteur, malveillant ou non, n'usurpe les fonctions cruciales du « Root Bridge ».

Loop Guard permet de protéger la topologie d'une boucle en cas de défaillance Spanning-Tree. En effet, un port en mode « blocking » ou « Listening » qui ne reçoit plus de paquets BPDU pourrait passer en mode « Forwarding » et ainsi créer une boucle réseaux. L'implémentation de « Loop Guard » par l'administrateur permet de remédier à ce problème, quand un port en mode « Blocking » ou « Listening » ne reçoit plus de BPDU, il passe en état « loop-inconsistent ». Dans cet état, le port ne transmet plus aucun trafic, il sera automatiquement réintégré dans son rôle Spanning-Tree dès une nouvelle réception de BPDU. A noter que, de la même manière que les autres mécanismes de sécurité STP, Loop Guard peut être configuré de manière globale ou au cas par cas.

La gestion rigoureuse de l'accès aux réseaux est également un axe de sécurisation indirecte de Spanning-Tree et de ses variantes. Cette gestion des accès aux réseaux peut passer par l'implémentation de « Port Security » qui limite les clients pouvant se connecter aux réseaux via l'identification de leur MAC. Une gestion des accès réseaux peut passer par la mise en place du 802.1x pour limiter la connexion aux périphériques authentifiés par un certificat. Cela est bien plus sécurisé que le « Port Security » mais bien plus complexe à mettre en œuvre. Il est également possible d'ajouter des ACL pour affiner les conditions d'accès et de communication sur les réseaux.

Les possibilités d'intégration des variables du protocole Spanning-Tree sont nombreuses. Dans le cadre de l'architecture à trois niveaux, l'implémentation sera le plus souvent réservée aux couches « Aggregation » et « Access » ; en effet la couche « Core » interagissant au niveau 3 du modèle OSI, STP n'est pas adaptée. L'utilisation de STP est de nos jours considérée comme obsolète, RSTP lui sera

préférée pour sa vitesse de convergence. Cependant, dans un environnement nécessitant une répartition de charge réseaux, MSTP et RPVST+ seront préférés. RPVST+ étant uniquement compatible dans les environnements CISCO, il sera bien souvent délaissé au profit du standard MSTP. En plus de sa compatibilité universelle, MSTP permet de regrouper plusieurs VLAN afin de réduire le nombre d'arbres STP à prendre en compte durant la configuration. La gestion des BPDU de MSTP est également très avantageuse dans le cas d'un environnement avec un grand nombre de VLAN, car les informations sont regroupées dans un seul BPDU contrairement à RPVST+ qui émet un BPDU par VLAN. RPVST+ peut donc se montrer rapidement plus gourmand que MSTP en ressources. L'intégration multi-régions de MSTP, bien que complexe à mettre en œuvre, représente un véritable atout pour la gestion d'infrastructures très étendues.

Dans le cadre de notre maquette, nous mettrons en place MSTP et RPVST+ sur nos couches « Aggregation » et « Access ». Nous modifierons des valeurs telles que le « Bridge Priority » pour influencer les arbres STP. Nous créerons plusieurs instances MSTP pour la répartition de charge et mettrons en place certains mécanismes de sécurité présentés ci-dessus.

QCM MSTP et RPVST+

a) **RPVST+ et MSTP sont tous les deux des protocoles dérivés de ?**

- ☐ STP
- ☐ RPT
- ☐ MTP
- ☐ VTP

b) **PVST, PVST+, RPVST+ sont propriétaires CISCO.**

- ☐ Vrai
- ☐ Faux

c) **STP permet d'empêcher les boucles de manière dynamique.**

- ☐ Vrai
- ☐ Faux

d) **Donnez les deux types de BPDU Spanning-Tree.**

e) **Donnez la définition de « Root Path Cost » dans le paquet BPDU Spanning-Tree.**

Correction QCM MSTP et RPVST+

a) **RPVST+ et MSTP sont tous les deux des protocoles dérivés de ?**

- ☒ **STP**
- ☐ RPT
- ☐ MTP
- ☐ VTP

b) **PVST, PVST+, RPVST+ sont propriétaires CISCO.**

- ☒ **Vrai**
- ☐ Faux

c) **STP permet d'empêcher les boucles de manière dynamique.**

- ☒ **Vrai**
- ☐ Faux

d) **Donnez les deux types de BPDU Spanning-Tree.**

Configuration BPDU, qui crée et maintient la topologie avec tous les chemins possibles sans boucle.

Topology Change Notification BPDU, qui met à jour la topologie en cas de changement.

e) **Donnez la définition de « Root Path Cost » dans le paquet BPDU Spanning-Tree.**

Indique sur 4 octets le coût du chemin vers le “Root Bridge”, dans le but de déterminer le chemin préféré.

[Vidéo de mise en œuvre MSTP et RPVST+ \(DRIVE\)](#)

[Vidéo de mise en œuvre MSTP et RPVST+ \(Youtube\)](#)

Fiche TP MSTP et RPVST+

1. Dans la zone IEEE sur les switch « DC1-DST1 » et « DC1-DST2 » configurer le mode MSTP.
2. Dans la zone IEEE configurer sur les switch « DC1-DST1 » et « DC1-DST2 » la région « region-IEEE » et le numéro de révision de configuration.

3. Dans la zone IEEE sur le switch « DC1-DST1 » créer les instances de la manière suivante :

Instance	VLAN attaché	Priorité
instance 1	110,120	8192
instance 2	130,140	16384
instance 3	100	8192

4. Dans la zone IEEE sur le switch « DC1-DST2 » créer les instances de la manière suivante :

Instance	VLAN attaché	Priorité
instance 1	110,120	16384
instance 2	130,140	8192
instance 3	100	16384

5. Dans la zone IEEE sur les switch « DC1-ACCESS1 » à « DC1-ACCESS4 » procéder à l'activation du mode MSTP et à la déclaration de l'instance « 1 », « 2 » et « 3 » sans modifier leur priorité.
6. Dans la zone IEEE sur les switch « DC1-ACCESS1 » à « DC1-ACCESS4 » procéder à l'activation de « PortFast » et de « BPDUFiltering » sur toutes les interfaces clientes.
7. Dans la zone CISCO sur les switch « DC1-DST3 » et « DC1-DST4 » configurer le mode RPVST+.
8. Dans la zone CISCO sur le switch « DC1-DST3 » configurer les priorités de la manière suivante pour les VLAN :

VLAN	Priorité
200,210,220	8192
230,240	16384

9. Dans la zone CISCO sur le switch « DC1-DST4 » configurer les priorités de la manière suivante pour les VLAN :

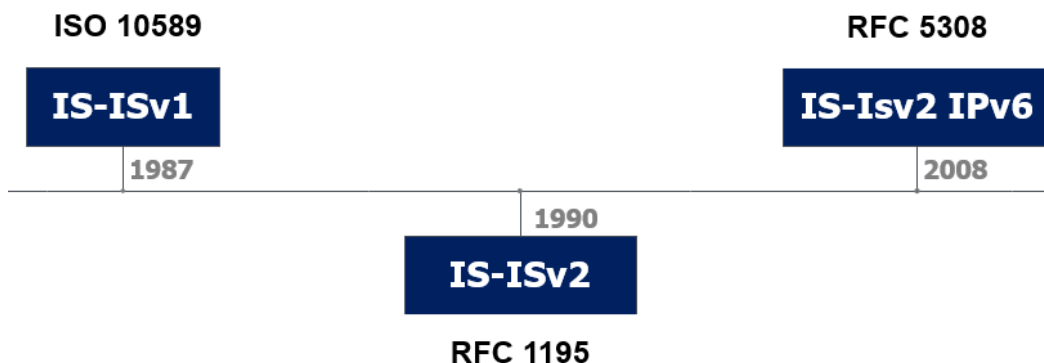
VLAN	Priorité
200,210,220	16384
230,240	8192

10. Dans la zone CISCO sur les switch « DC1-ACCESS5 » à « DC1-ACCESS8 » procéder à l'activation du mode RPVST+.
11. Dans la zone CISCO sur les switch « DC1-ACCESS5 » à « DC1-ACCESS8 » procéder à l'activation de « PortFast » et de « BPDUFiltering » sur toutes les interfaces clientes.

[Télécharger Correction TP MSTP et RVPST+](#)

6. Qu'est-ce que IS-IS ?

Créé en 1987 et défini dans l'ISO 10589 IS-IS (Intermediate System to Intermediate System), il sera par la suite standardisé pour TCP/IP en 1990 dans la RCF 1195 actualisé par le RFC-5308 en 2008, dans l'objectif de prendre en charge l'IPv6 et le concept de double pile.



IS-IS est un protocole de routage dynamique de la famille des IGP (Interior Gateway Protocol), tout comme OSPF et EIGRP. C'est-à-dire qu'il est utilisé pour le routage à l'intérieur d'une même AS, mais également pour le routage des très grandes infrastructures réseaux. Ce protocole dispose de fonctionnalités de routage avancées comme l'intégration au sein d'un MPLS, une vitesse de convergence élevée et la gestion de la double pile IPv4/IPv6. Il est donc particulièrement efficace et capable de maintenir une cohérence des tables de routage dans des infrastructures soumises à de fortes modifications. Il est à noter que IS-IS est également bien souvent plus performant qu'OSPF pour les infrastructures multi-zones très étendues. Cela en fait un protocole particulièrement apprécié par les opérateurs internet, mobiles et fournisseurs de services Cloud.

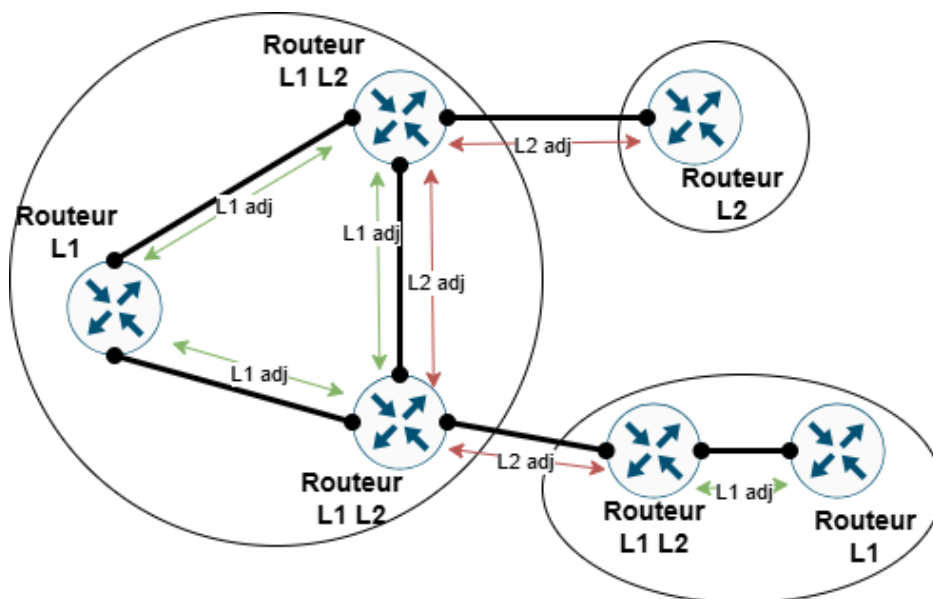
Les zones IS-IS sont une classification hiérarchique dans le routage des très grandes topologies. Elles permettent de diviser la topologie en sous-ensembles afin d'optimiser la gestion avec une réduction du trafic de mise à jour, l'augmentation de la vitesse de convergence, sans oublier la réduction de la taille de la LSDB (concept expliqué plus tard dans ce cours) et de la table de routage. Cela se traduit par un allègement important de l'utilisation des ressources de calculs. Les routeurs d'une zone forment une sous-topologie qui limite les informations échangées sur la topologie entre les zones. En l'absence de nécessité d'un routage inter-zones, les routeurs se contentent de communiquer uniquement au sein de leur zone.

Il existe deux niveaux de zones :

- **Zone niveau 1** : Elle ne connaît que l'information topologique interne à la zone, mais également les préfixes IP internes en vue de ne gérer que le routage interne à cette zone. En cas de routage inter-zones, le trafic sera routé vers les routeurs niveau 1 et 2.
- **Zone niveau 2** : Elle relie plusieurs zones entre elles et assure uniquement le routage inter-zones afin de constituer une « Backbone ». Elle n'interagit qu'avec des routeurs de niveau 2.

Les routeurs peuvent donc agir de trois manières différentes :

- **Routeur niveau 1** : Ces routeurs ont une visibilité uniquement sur la zone interne, les routes construites sont exclusivement à destination locale. Le trafic à destination d'une autre zone sera confié à un routeur niveau 1 et 2.
- **Routeur niveau 2** : Il a la visibilité d'accès à toutes les zones sans en connaître les détails internes, dans l'objectif de constituer la « Backbone ». Il effectue uniquement le transit des flux d'une zone à une autre, sans connaissance du chemin interne à la zone de destination.
- **Routeur niveau 1 et 2** : Ils sont connectés à des routeurs de niveau 1 et des routeurs de niveau 2, dans l'objectif de jouer le rôle de passerelle de zone. Pour cela, ces routeurs disposent de deux bases de données distinctes, une par niveau. Cela leur permet de connecter les routeurs niveau 1 d'une zone à une autre.



Les routeurs IS-IS vont passer par les 4 états suivants durant leur mise en fonction :

1. **Down** : Le protocole de routage est en cours de configuration, aucune adjacence n'est établie.
2. **Initializing** : Les paquets HELLO sont échangés pour établir des adjacences.
3. **Up (ou Two-Way)** : La relation d'adjacence bidirectionnelle est établie.
4. **Full** : Les LSDB sont échangées et synchronisées entre les voisins.

IS-IS est un protocole dit « à état de lien », c'est-à-dire qu'il envoie des messages à chaque fois qu'un changement de topologie se produit. Il utilise l'algorithme de Dijkstra tout comme OSPF, en vue de choisir la route avec le chemin le plus court en tenant compte de la topologie du réseau. Une fois configuré, IS-IS va échanger différents messages pour construire et mettre à jour sa topologie. Ces messages sont les suivants :

- **HELLO PDU** : Ce message périodique est utilisé pour créer et maintenir les relations de voisinage au sein de la topologie IS-IS, tout en vérifiant les paramètres de connectivité. Le message comporte plusieurs informations comme les timer d'actualisation. Il existe trois types de messages HELLO PDU en fonction des types de réseaux :
 - **P2P HELLO PDU** : Message HELLO PDU est utilisé spécifiquement sur les relations points à points, qui connectent directement deux routeurs ensemble. C'est le message HELLO PDU le plus courant.

- **LAN HELLO PDU** : Ce message est envoyé en multicast sur le segment réseau qui connecte plus de deux routeurs, cela permet à tous les routeurs de se connaître. Une fois que tous les routeurs se connaissent, ils vont procéder à l'élection du DIS « Designated Intermediate System » en fonction de la priorité la plus haute, en cas d'égalité l'identifiant de système le plus élevé sera préféré. Le DIS sera responsable de l'envoi des messages Link State PDU, afin de centraliser la LSDB.
- **Level-1 et Level-2 HELLO** : Ces deux types de messages HELLO fonctionnent de la même manière, pour établir et maintenir les relations de voisinage dans leurs zones respectives.
- **Link State PDU** : Les messages LSP contiennent les informations détaillées de la topologie comme l'état des liens, les métriques associées et les informations de routage des réseaux connus du routeur. Ces messages permettent de construire la Link State Database (LSDB) contenant l'intégralité de la topologie réseaux.
- **Complete Sequence Number PDU** : Les messages CSNP sont envoyés périodiquement pour synchroniser la LSDB entre les routeurs.
- **Partial Sequence Number PDU** : Ces messages sont utilisés comme accusés de réception des messages LSP, mais également pour solliciter un message LSP manquant, afin de compléter la LSDB.

Intéressons-nous maintenant à la composition des paquets des différents messages IS-IS :

Le paquet **HELLO PDU** est composé des éléments suivants :

Intradomain Routing Protocol Discriminator	Length Indicator	Version/Protocol ID Extension	ID Length	PDU Type	Version	Maximum Area Addresses	Circuit Type	Source ID
Holding Time	Local Circuit ID	Authentication	Point-to-Point Adjacency State	Protocols Supported	Area Address	IP Interface Address(es)	IPv6 Interface Address(es)	

- **Intradomain Routing Protocol Discriminator** : Indique sur 1 octet le protocole de routage utilisé, la valeur est « 0x83 » pour IS-IS.
- **Length Indicator** : Indique sur 1 octet la taille de l'en-tête IS-IS, cela permet de différencier l'en-tête du champ TLV.
- **Version/Protocol ID Extension** : Indique sur 1 octet la version du protocole.
- **ID Length** : Indique sur 1 octet la taille de l'identifiant de l'IS, une valeur à « 0 » signifie implicitement que la taille aura la valeur par défaut, soit 6 octets.
- **PDU Type** : Indique sur 1 octet le type de paquet, pour les paquets HELLO les valeurs peuvent être :
 - **1** : LAN Hello niveau 1
 - **2** : LAN Hello niveau 2
 - **3** : P2P Hello
 - **17** : P2P Hello (ancien format)
- **Version** : Indique sur 1 octet la version du protocole.
- **Maximum Area Address(es)** : Indique sur 1 octet le nombre maximal d'adresses de zone supporté (la valeur « 0 » signifie ce champ est ignoré).
- **Circuit Type** : Indique sur 1 octet les niveaux des routeurs traversés, niveau 1, niveau 2 et multi-niveau.
- **Source ID** : Indique sur 6 octets l'ID du routeur émetteur du paquet.
- **Holding Time** : Indique sur 2 octets le temps maximal admissible entre deux paquets, une fois ce temps dépassé la relation de voisinage est abandonnée.

- **PDU Length** : Indique sur 2 octets la longueur totale du paquet HELLO (en-tête et TLV).
- **Local Circuit ID** : Indique sur 1 octet l'identifiant unique de l'interface.
- **Authentication** : Indique la méthode d'authentification des communications du protocole.
- **Point-to-Point Adjacency State** : Indique l'état sur une interface P2P.
- **Protocols Supported** : Indique le ou les protocoles réseaux pris en charge, IPv4 et/ou IPv6.
- **Area Address(es)** : Indique l'ID de la zone à laquelle appartient l'émetteur.
- **IP Interface Address(es)** : Indique l'adresse IPv4 configurée sur l'interface.
- **IPv6 Interface Address(es)** : Indique l'adresse IPv6 configurée sur l'interface.

Télécharger le paquet Wireshark HELLO PDU

Le paquet **Link State PDU** est composé des éléments suivants :

Intradomain Routing Protocol Discriminator		Length Indicator	Version/Protocol ID Extension	ID Length	PDU Type	Version	Maximum Area Addresses	PDU Length	Remaining Lifetime
LSP ID	Sequence Number	Checksum	Attachment	Overload	Type of Intermediate System	Authentication	Area Address	Protocols Supported	
Hostname	Extended IS Reachability	IP Interface Address	Extended IP Reachability	IPv6 Interface Address	IPv6 Reachability				

- **Intradomain Routing Protocol Discriminator** : Indique sur 1 octet le protocole de routage utilisé, la valeur est « 0x83 » pour IS-IS.
- **Length Indicator** : Indique sur 1 octet la taille de l'en-tête IS-IS, cela permet de différencier l'en-tête du champ TLV.
- **Version/Protocol ID Extension** : Indique sur 1 octet la version du protocole.
- **ID Length** : Indique sur 1 octet la taille de l'identifiant de l'IS, une valeur à « 0 » signifie implicitement que la taille aura la valeur par défaut, soit 6 octets.
- **PDU Type** : Indique sur 1 octet le type de paquet, pour les paquets LSP les valeurs peuvent être :
 - 5 : LSP niveau 1
 - 6 : LSP niveau 2
- **Version** : Indique sur 1 octet la version du protocole.
- **Maximum Area Address(es)** : Indique sur 1 octet le nombre maximal d'adresses de zone supporté (la valeur "0" signifie que ce champ est ignoré).
- **PDU Length** : Indique sur 2 octets la longueur totale du paquet LSP (en-tête et TLV).
- **Remaining Lifetime** : Indique sur 8 octets la durée de validité du paquet LSP.
- **LSP ID** : Indique sur 2 octets l'identifiant unique du paquet LSP, qui est constitué de l'ID du routeur émetteur et du l'ID de l'instance LSP.
- **Sequence Number** : Indique sur 4 octets un ID unique, permettant de déterminer le LSP le plus récent.
- **Checksum** : Champ de contrôle d'intégrité.
- **Attachment** : Indique si le routeur émetteur est connecté à plusieurs zones.
- **Overload** : Indique si le routeur émetteur est surchargé.
- **Type of Intermediate System** : Indique le type du routeur émetteur.
- **Authentication** : Indique la méthode d'authentification des communications du protocole.
- **Area Address** : Indique l'ID de la zone à laquelle appartient l'émetteur.

- **Protocols Supported** : Indique le ou les protocoles réseaux pris en charge, IPv4 et/ou IPv6.
- **Hostname** : Indique le nom du routeur émetteur.
- **Extended IS Reachability** : Indique la liste des voisins IS-IS atteignables via le routeur émetteur.
- **IP Interface Address(es)** : Indique l'adresse IPv4 de l'interface du routeur émetteur.
- **Extended IP Reachability** : Indique les réseaux IPv4 joignables via le routeur émetteur.
- **IPv6 Interface Address(es)** : Indique l'adresse IPv6 de l'interface du routeur émetteur.
- **IPv6 Reachability** : Indique les réseaux IPv6 joignables via le routeur émetteur.

[Télécharger le paquet Wireshark LSPDU](#)

Le paquet **Complete Sequence Number PDU** est composé des éléments suivants :

Intradomain Routing Protocol Discriminator	Length Indicator	Version/Protocol ID Extension	ID Length	PDU Type	Version	Maximum Area Addresses	PDU Length	Source ID
Start LSP ID	End LSP ID	Authentication	LSP Entries					

- **Intradomain Routing Protocol Discriminator** : Indique sur 1 octet le protocole de routage utilisé, la valeur est « 0x83 » pour IS-IS.
- **Length Indicator** : Indique sur 1 octet la taille de l'en-tête IS-IS, cela permet de différencier l'en-tête du champ TLV.
- **Version/Protocol ID Extension** : Indique sur 1 octet la version du protocole.
- **ID Length** : Indique sur 1 octet la taille de l'identifiant de l'IS, une valeur à « 0 » signifie implicitement que la taille aura la valeur par défaut, soit 6 octets.
- **PDU Type** : Indique sur 1 octet le type de paquet, pour les paquets CNSP les valeurs peuvent être :
 - 7 : CNSP niveau 1
 - 8 : CNSP niveau 2
- **Version** : Indique sur 1 octet la version du protocole.
- **Maximum Area Address(es)** : Indique sur 1 octet le nombre maximal d'adresses de zone supporté (la valeur « 0 » signifie que ce champ est ignoré).
- **PDU Length** : Indique sur 2 octets la longueur totale du paquet CNSP (en-tête et TLV).
- **Source ID** : Indique sur 6 octets l'ID du routeur émetteur du paquet.
- **Start LSP ID** : Indique sur 6 octets l'identifiant unique du premier Link State de la plage synchronisée.
- **End LSP ID** : Indique sur 6 octets l'identifiant unique du dernier Link State de la plage synchronisée.
- **Authentication** : Indique la méthode d'authentification des communications du protocole.
- **LSP Entries** : Indique la liste de chaque paquet Link State accompagnée des différentes valeurs : LSP ID (sur 6 octets), Sequence Number (sur 4 octets), Lifetime (sur 2 octets) et Checksum (sur 2 octets).

[Télécharger le paquet Wireshark CNSPDU](#)

Le paquet **Partial Sequence Number PDU** est composé des éléments suivants :

Intradomain Routing Protocol Discriminator	Length Indicator	Version/Protocol ID Extension	ID Length	PDU Type	Version	Maximum Area Addresses	PDU Length	Source ID	Authentication	LSP Entries
--	---------------------	----------------------------------	--------------	-------------	---------	------------------------------	---------------	--------------	----------------	----------------

- **Intradomain Routing Protocol Discriminator** : Indique sur 1 octet le protocole de routage utilisé, la valeur est « 0x83 » pour IS-IS.
- **Length Indicator** : Indique sur 1 octet la taille de l'en-tête IS-IS, cela permet de différencier l'en-tête du champ TLV.
- **Version/Protocol ID Extension** : Indique sur 1 octet la version du protocole.
- **ID Length** : Indique sur 1 octet la taille de l'identifiant de l'IS, une valeur à « 0 » signifie implicitement que la taille aura la valeur par défaut, soit 6 octets.
- **PDU Type** : Indique sur 1 octet le type de paquet, pour les paquets PSNP les valeurs peuvent être :
 - **9** : PSNP niveau 1
 - **10** : PSNP niveau 2
- **Version** : Indique sur 1 octet la version du protocole.
- **Maximum Area Address(es)** : Indique sur 1 octet le nombre maximal d'adresses de zone supporté (la valeur « 0 » signifie que ce champ est ignoré).
- **PDU Length** : Indique sur 2 octets la longueur totale du paquet PSNP (en-tête et TLV).
- **Source ID** : Indique sur 6 octets l'ID du routeur émetteur du paquet.
- **Authentication** : Indique la méthode d'authentification des communications du protocole.
- **LSP Entries** : Indique la liste de chaque paquet Link State accompagnée des différentes valeurs : LSP ID (sur 6 octets), Sequence Number (sur 4 octets), Lifetime (sur 2 octets) et Checksum (sur 2 octets).

[Télécharger le paquet Wireshark PSNPDU](#)

Ces messages nous assurent de maintenir une LSDB à jour, qui à travers l'algorithme SPF, détermine les routes les plus courtes pour chaque destination en partant des routeurs. Bien que chaque routeur calcule ses routes de manière indépendante, la synchronisation périodique de la LSDB garantit leur cohérence.

Maintenant que la LSDB est constituée, le routeur doit procéder à la sélection des routes de la manière suivante :

1. **Hierarchie des niveaux** : La route annoncée par un routeur de niveau 1 sera préférée à la route annoncée par un routeur de niveau 2.
2. **Métrique** : Le second critère de sélection est le coût du chemin le plus faible. Le coût du chemin est incrémenté avec la métrique de chacune des interfaces traversées. A noter que la métrique d'une interface est par défaut égale à 10, cependant elle peut être personnalisée par l'administrateur.
3. **Route de secours** : En cas d'égalité entre deux routes, un équilibrage de charge natif sera effectué par le protocole de routage.
4. **Routes externes** : Les routes internes IS-IS sont préférées aux routes externes provenant d'un autre protocole. A noter qu'il existe deux types de routes externes :
 - **E1** : Route incluant la métrique du protocole de routage externe, mais également la métrique interne à IS-IS.

- **E2** : Ne tient compte que de la métrique du protocole de routage externe.
- 5. **Tiebreakers** : Si toutes les valeurs précédentes sont identiques, l'IP de l'interface de sortie et l'ID du routeur peuvent être utilisés pour départager deux routes.

En tant qu'IGP, IS-IS est moins exposé aux attaques que les EGP, cependant sa sécurisation reste indispensable. Pour cela, il est possible d'intégrer des sécurités pour empêcher l'injection de routes malveillantes.

La première sécurité pouvant être implémentée, est l'authentification des messages IS-IS afin d'empêcher l'injection de routes malveillantes. Cette authentification peut être assurée par divers moyens :

- L'authentification par mot de passe en clair est la méthode d'authentification la plus facile à mettre en œuvre mais considérée comme obsolète de nos jours.
- L'authentification par hachage MD5, qui repose sur un hachage des messages basé sur une clé préalablement pré-partagée.
- L'authentification par hachage SHA256 (voir SHA512 si compatible) ; ces niveaux de hachage ne sont généralement pas nativement pris en charge directement dans IS-IS. Cependant, un moyen détourné de mise en œuvre existe qui consiste en la génération préalable d'une clé hachée en SHA256, avant de la configurer comme clé d'authentification au sein IS-IS.

Ces authentifications simples à mettre en œuvre présentent des limites. Notamment le fait que le hachage en MD5 est de nos jours facilement réversible, mais également qu'en cas de fuite du mot passe, tous les routeurs de la topologie seront exposés. L'authentification via SHA256 est considérée comme acceptable.

TTL Security Check, consiste à limiter le nombre de sauts acceptables afin de réduire le risque d'injection de paquets en provenance d'un routeur extérieur aux réseaux légitimes. Bien que cette méthode puisse être efficace, elle présente des limites dues à la complexité en cas d'extension du réseau, le TTL devra être reconfiguré sur tous les routeurs. De plus, cette méthode ne doit pas être utilisée seule, mais en complément d'une authentification des messages IS-IS.

L'implémentation d'IS-IS au sein d'un tunnel IPSEC au-delà d'une authentification forte, permet de chiffrer les paquets au niveau de la couche transport. Cette méthode de sécurisation bien que très lourde à mettre en œuvre permet de se prémunir de bon nombre d'attaques, comme la reconnaissance, l'injection de données erronées, etc...

IS-IS peut être implémenté à différents niveaux dans l'architecture à trois couches. L'utilisation, la plus courante, est l'interconnexion de la couche « Aggregation » et « Core » dans les infrastructures utilisant la double pile IPV4 et IPV6 ou nécessitant une grande vitesse de convergence. Dans le cadre d'un fournisseur de service Cloud multi-sites, IS-IS peut être utilisé pour l'interconnexion des différentes couches « Core » des sites via la mise en place d'un MPLS.

Au sein de l'architecture que nous mettrons en œuvre durant les travaux pratiques de ce cours, nous utiliserons IS-IS pour l'interconnexion de la couche « Aggregation » et « Core » afin de bénéficier de la vitesse de convergence. IS-IS nous permettra également de profiter de la double pile IPv4/IPv6. Nous procéderons également à l'implémentation d'une authentification par hachage SHA256.

QCM IS-IS

a. **IS-IS est un protocole de routage dynamique de la famille des IGP tout comme ?
(2 bonnes réponses)**

- ☐ OSPF
- ☐ EIGRP
- ☐ SPB
- ☐ PVST+

b. **Combien de zones existe-t-il dans IS-IS ?**

- ☐ 1
- ☐ 2
- ☐ 3
- ☐ 4

c. **Par combien d'états vont passer les routeurs IS-IS durant leur mise en fonction ?**

- ☐ 4
- ☐ 3
- ☐ 2
- ☐ 1

d. **A quoi sert l'élément « Area Address » dans le paquet « hello PDU » de ISIS ?**

- ☐ Indique l'ID de la zone à laquelle appartient l'émetteur.
- ☐ Indique l'état sur une interface P2P.
- ☐ Indique sur 1 octet l'identifiant unique de l'interface.
- ☐ Indique sur 2 octets le temps maximal admissible entre deux paquets, une fois ce temps dépassé la relation de voisinage est abandonnée.

e. **Donnez les 4 états de mise en fonction de IS-IS ?**

Correction QCM IS-IS

a) **IS-IS est un protocole de routage dynamique de la famille des IGP tout comme ? (2 bonnes réponses)**

- ☒ **OSPF**
- ☒ **EIGRP**
- ☐ SPB
- ☐ PVST+

b) **Combien de zones existe-t-il dans IS-IS ?**

- ☐ 1
- ☒ **2**
- ☐ 3
- ☐ 4

c) **Par combien d'états vont passer les routeurs IS-IS durant leur mise en fonction ?**

- ☒ **4**
- ☐ 3
- ☐ 2
- ☐ 1

d) **A quoi sert l'élément « Area Address » dans le paquet « hello PDU » de ISIS ?**

- ☒ **Indique l'ID de la zone à laquelle appartient l'émetteur.**
- ☐ Indique l'état sur une interface P2P.
- ☐ Indique sur 1 octet l'identifiant unique de l'interface.
- ☐ Indique sur 2 octets le temps maximal admissible entre deux paquets, une fois ce temps dépassé la relation de voisinage est abandonnée.

e) **Donnez les 4 états de mise en fonction de IS-IS ?**

Down, Initializing, Up (ou Two-Way), Full

[Vidéo de de mise en œuvre IS-IS \(DRIVE\)](#)

[Vidéo de de mise en œuvre IS-IS \(Youtube\)](#)

Fiche TP IS-IS

1. Sur les routeur « DC1-CORE1 » et « DC-CORE2 » ainsi que les MLS « DC1-DST1 » à « DC1-DST4 », « DC2-SPINE1 » et « DC2-SPINE », puis de « DC2-LEAF1 » à « DC2-LEAF4 », vous devez configurer une clé en SHA256 pour la sécurisation de IS-IS.

2. Sur les routeur « DC1-CORE1 » et « DC-CORE2 » ainsi que les MLS « DC1-DST1 » à « DC1-DST4 », « DC2-SPINE1 » et « DC2-SPINE », puis de « DC2-LEAF1 » à « DC2-LEAF4 », vous devez activer IS-IS.
3. Dans le mode de configuration IS-IS vous devez définir les « Network Entity Titel » pour chacun des équipements suivant.

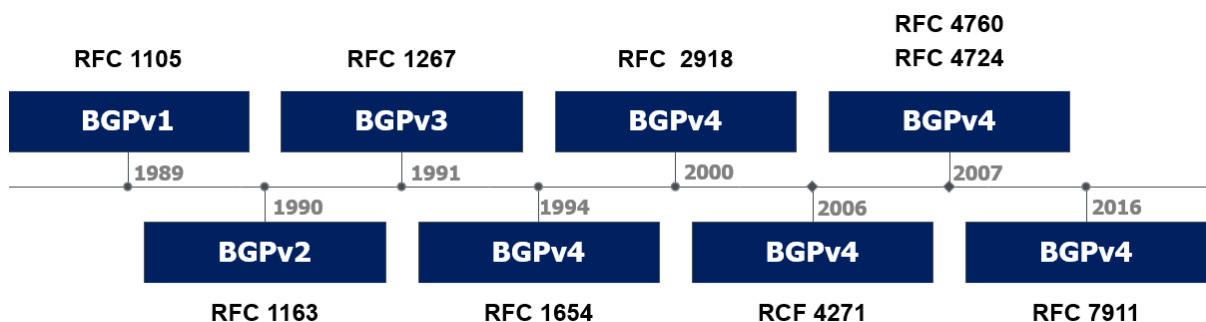
Equipement	Network Entity Titel
DC1-CORE1	49.0001.0000.0000.0001.00
DC1-CORE2	49.0001.0000.0000.0002.00
DC1-DST1	49.0001.0000.0000.0003.00
DC1-DST2	49.0001.0000.0000.0004.00
DC1-DST3	49.0001.0000.0000.0005.00
DC1-DST3	49.0001.0000.0000.0006.00
DC2-SPINE1	49.0002.0000.0000.0001.00
DC2-SPINE2	49.0002.0000.0000.0002.00
DC2-LEAF1	49.0002.0000.0000.0003.00
DC2-LEAF2	49.0002.0000.0000.0004.00
DC2-LEAF3	49.0002.0000.0000.0005.00
DC2-LEAF4	49.0002.0000.0000.0006.00

4. Sur tous les équipements sur lequel vous avez activé IS-IS définir le niveau hiérarchique, en mode niveau deux seulement.
5. Activer l'authentification IS-IS par clé, et utiliser la clé précédemment générée pour tous les équipements IS-IS.
6. Configurer le format de métrique « Wide » pour tous les équipements IS-IS.
7. Activer la redistribution des routes connectées pour tous les équipements IS-IS.
8. Sur les interfaces de tous les équipements IS-IS activer le routage IS-IS IPv4 et IPv6.
9. Sur les interfaces de tous les équipements IS-IS configurer le type de liaison en point-to-point, ainsi que le type d'adjacence en mode niveau deux uniquement.
10. Sur les interfaces de tous les équipements IS-IS configurer l'authentification IS-IS par clé, et utiliser la clé précédemment générée.

[Télécharger Correction TP IS-IS](#)

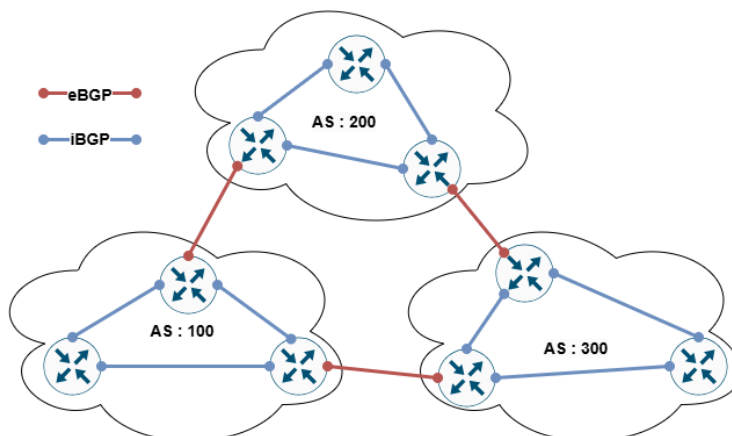
7. Qu'est-ce que BGP ?

BGP pour « Border Gateway Protocol » est un protocole de routage dynamique successeur du protocole EGP pour « Exterior Gateway Protocol » en raison des limitations de EGP quant à la gestion des boucles réseaux. La première version de BGP a été définie dans la RFC-1105 en 1989, qui a été déclinée en plusieurs versions jusqu'à la version actuelle BGPv4 définie pour la première fois en 1994 avec la RFC-1654 (réactualisée en 1995 avec la RFC-1771 et en 2006 avec la RFC-4271).



La version actuelle de BGP intègre la gestion du CIDR (Classless Inter-Domain Routing) afin de permettre une gestion efficace de l'agrégation des routes. A l'inverse des autres protocoles de routage dynamique (OSPF, EIGRP et IS-IS), BGP est conçu pour gérer le routage à travers plusieurs systèmes autonomes (AS). Il est également capable de prendre en charge des tables de routage avec plusieurs milliers de routes. Cela en fait un protocole de routage très utilisé par les fournisseurs d'accès internet et les fournisseurs de Cloud.

BGP est décomposé en deux branches iBGP (Interior BGP) et eBGP (Exterior BGP) ; iBGP est utilisé pour le routage au sein d'un même AS tandis que eBGP est utilisé pour le routage inter-AS. Plus concrètement, prenons l'exemple des opérateurs A et B disposant chacun de leur AS, le protocole utilisé sera eBGP pour communiquer de l'opérateur A à B ; à l'inverse pour communiquer entre deux routeurs du même opérateur, le protocole utilisé sera iBGP. Il est à noter que les routes iBGP ne sont pas redirigées vers d'autres routeurs iBGP, afin de ne pas créer de boucles réseaux. Pour ce qui est des routes eBGP, elles peuvent être aussi bien redirigées vers d'autres routeurs eBGP ou iBGP.



Pour rentrer un peu plus dans le détail, vous pouvez retenir les informations dans le tableau ci-dessous :

	Distance Administrative (AD) par défaut	TTL par défaut
iBGP	200	225
eBGP	20	1

Pour fonctionner, BGP doit établir des relations de voisinage appelées « paires », ces paires doivent être configurées manuellement. Une fois que les paires sont configurées sur les deux routeurs, la relation de voisinage est établie via des messages « OPEN » avec une connexion TCP sur le port 179. La connexion est ensuite maintenue active via le message « KEEPALIVE ». Les tables de routage sont échangées et mises à jour avec le message « UPDATE ». Le message « NOTIFICATION » est réservé pour les fins de session entre paires en cas d'erreur ou de reconfiguration. Rentrons plus en profondeur dans les paquets BGP :

Le paquet **OPEN** est composé des éléments suivants :

Marker	Length	Type	Version	My Autonomous System	Hold Time	BGP Identifier	Optional Parameters Length	Optional Parameters
--------	--------	------	---------	----------------------	-----------	----------------	----------------------------	---------------------

- **Marker** : Début du paquet contenant 16 octets remplis de « 1 » en temps normal dans le but de s'assurer que le paquet n'est pas corrompu.
- **Length** : Indique sur 2 octets la longueur totale du paquet.
- **Type** : Indique sur 1 octet le type de paquet, dans le cas du paquet OPEN la valeur est à « 1 ».
- **Version** : Indique sur 1 octet la version de BGP utilisée.
- **My Autonomous System** : Indique sur 2 octets le numéro de l'AS du routeur émetteur de ce paquet.
- **Hold Time** : Indique sur 2 octets un temps maximal admissible entre deux paquets KEEPALIVE ou UPDATE, une fois ce temps dépassé, la connexion entre les paires est fermée.
- **BGP Identifier** : Indique sur 4 octets l'identifiant unique du routeur émetteur, bien souvent cet identifiant est l'adresse de boucle locale (Loopback) ou à défaut l'adresse IP avec la valeur la plus élevée du routeur.
- **Optional Parameters Length** : Indique sur 1 octet la longueur du champ « Optional Parameters ».
- **Optional Parameters** : Indique les paramètres optionnels pouvant être négociés entre les deux routeurs.

[Télécharger le paquet Wireshark OPEN](#)

Le paquet **KEEPALIVE** est composé des éléments suivants :

Marker	Length	Type
--------	--------	------

- **Marker** : Début du paquet contenant 16 octets remplis de « 1 » en temps normal dans le but de s'assurer que le paquet n'est pas corrompu.
- **Length** : Indique sur 2 octets la longueur totale du paquet en octets.
- **Type** : Indique sur 1 octet le type de paquet, dans le cas du paquet KEEPALIVE la valeur est à « 4 ».

[Télécharger le paquet Wireshark KEEPALIVE](#)

Le paquet **UPDATE** est composé des éléments suivants :

Marker	Length	Type	Withdrawn Routes Length	Withdrawn Routes	Path Attributes Length	Path Attributes	Network Layer Reachability Information
--------	--------	------	-------------------------	------------------	------------------------	-----------------	--

- **Marker** : Début du paquet contenant 16 octets remplis de 1 en temps normal dans le but de s'assurer que le paquet n'est pas corrompu.
- **Length** : Indique sur 2 octets la longueur totale du paquet en octets.
- **Type** : Indique sur 1 octet le type de paquet, dans le cas du paquet UPDATE la valeur est à « 2 ».
- **Withdrawn Routes Length** : Indique sur 2 octets la longueur en octets des routes à supprimer de la table de routage, la valeur sera « 0 » en cas d'absence de routes à supprimer.
- **Withdrawn Routes** : Indique la liste des routes à supprimer de la table de routage.
- **Path Attributes Length** : Indique sur 2 octets la longueur en octets des attributs pour les routes nouvellement envoyées.
- **Path Attributes** : Indique la liste des attributs (les attributs seront expliqués ci-dessous) pour les routes nouvellement envoyées.
- **Network Layer Reachability Information** : Indique les adresses IP ainsi que leur masque de sous-réseaux.

[Télécharger le paquet Wireshark UPDATE](#)

Le paquet **NOTIFICATION** est composé des éléments suivants :

Marker	Length	Type	Error Code	Error Subcode	Data
--------	--------	------	------------	---------------	------

- **Marker** : Début du paquet contenant 16 octets remplis de 1 en temps normal dans le but de s'assurer que le paquet n'est pas corrompu.
- **Length** : Indique sur 2 octets la longueur totale du paquet en octets.
- **Type** : Indique sur 1 octet le type de paquet, dans le cas du paquet NOTIFICATION la valeur est à « 3 ».
- **Error Code** : Indique sur 1 octet le type d'erreur, erreur de configuration, erreur de session, etc...
- **Error Subcode** : Indique sur 1 octet la nature de l'erreur, par exemple l'incompatibilité entre deux versions BGP.
- **Data** : Indique les détails sur l'erreur rencontrée.

[Télécharger le paquet Wireshark NOTIFICATION](#)

Les routeurs BGP vont durant leur configuration passer par les 6 états suivants :

1. **Idle** : Dans cet état, le routeur est en cours de configuration et d'initialisation du timer « ConnectRetry ». Il n'a pas encore établi de connexion avec une paire. Cependant, il écoute et envoie activement des paquets pour établir une connexion. Dès réception d'un paquet, le routeur BGP passera en état « Connect ».
2. **Connect** : Dans cet état, le routeur parvient à un échange de paquets et attend que la connexion TCP soit complètement établie. En cas de succès, le routeur BGP passe en état « OpenSent », dans le cas d'un échec le routeur BGP passe en état « Active ». L'expiration du timer « ConnectRetry » relancera une tentative de connexion TCP.
3. **Active** : Dans cet état, le routeur BGP va tenter une nouvelle fois d'établir une connexion TCP. En cas de succès, le routeur BGP passe en état « OpenSent », dans le cas d'un échec ou d'expiration du timer « ConnectRetry » il retournera à l'état « Connect ». A noter qu'une réinitialisation manuelle du protocole renvoie à l'état « Idle ».
4. **OpenSent** : Dans cet état, le routeur BGP attend la réception du message « OPEN » du routeur voisin. Une fois le message reçu, il vérifie la cohérence de configuration, en cas d'incohérence ou de réinitialisation manuelle, il renvoie un message « NOTIFICATION » et retourne à l'état « Idle ». Dans le cas d'erreur de session TCP, il renvoie un message « NOTIFICATION » et retourne à l'état « Active ». Si aucune erreur de configuration ou de session TCP n'est détectée, l'échange des messages « KEEPALIVE » commence et le routeur BGP entre en état « OpenConfirm ».
5. **OpenConfirm** : Dans cet état, le routeur BGP attend la réponse « KEEPALIVE » de son voisin, une fois le message reçu le routeur BGP passera en état « Established ». Si un message « NOTIFICATION » est reçu le routeur retourne en état « Idle ». De même que si le message « KEEPALIVE » n'arrive pas dans le temps imparti par le timer, le routeur retournera dans l'état « Idle ».
6. **Established** : Dans cet état, le routeur BGP est maintenant prêt à échanger ses routes avec son voisin. L'échange périodique de messages « KEEPALIVE » permet de maintenir la connexion, en cas d'absence de message dans le temps imparti (90 secondes), le routeur retournera dans l'état « Idle ».

BGP est un protocole dit « à vecteur de chemin » c'est-à-dire qu'il transmet des informations détaillées en vue de choisir la meilleure route à emprunter pour atteindre une destination. Dans cet objectif, il s'appuie sur 10 attributs pour déterminer la route qui sera préférée. Les attributs sont traités dans l'ordre et de la manière suivante :

1. **Weight** : Propriétaire sur les équipements CISCO, est défini de manière locale sur le routeur, l'influence de cet attribut est uniquement locale. Plus le Weight est élevé, plus la route sera préférée.
2. **Local Preference** : Elle est définie au niveau de l'AS pour influencer la sélection de la route, plus elle est élevée plus la route sera prioritaire.
3. **Self Originated** : Les routes originaires du même routeur seront préférées.
4. **AS Path** : A chaque AS traversé par une route, le chemin « AS Path » est incrémenté, plus le « AS Path » est court, plus la route est préférée.
5. **Origine** : Les routes annoncées par iBGP sont privilégiées par rapport à celles annoncées par eBGP.
6. **Multi-Exit Discriminator** : Permet de sélectionner une priorité quand il y a plusieurs chemins vers une même AS, la MED la plus faible sera préférée.

7. **External** : Les routes priorisées sont celles qui disposent d'un chemin externe par rapport à un chemin interne. Cela permet une optimisation de performance, mais également de limiter le risque de boucles réseaux.
8. **IGP Cost** : La distance la plus petite vers le prochain saut, calculée par le protocole de routage interne (OSPF, IS-IS etc....) sera toujours préférée.
9. **eBGP peering** : En raison de sa large zone d'impact, BGP est conçu pour préférer la stabilité. C'est pour cette raison que la route la plus ancienne sera prioritaire.
10. **Routeur ID** : Le dernier attribut que BGP utilisera pour trancher entre deux possibilités de route est RID, la plus petite valeur sera préférée.

Ce processus de sélection complexe, les interactions à très grande échelle, ainsi que le besoin de stabilité font que BGP est le protocole de routage qui a la vitesse de convergence la plus lente.

Une fois configuré, BGP va maintenir à jour trois tables indispensables à son fonctionnement à savoir :

La **table de voisinage** qui contient chacune des paires avec leurs informations : état de la connexion, adresse IP, le numéro d'AS, la version du protocole, le nombre de messages envoyés et reçus, etc...

```
ORANGE#show ip bgp summary
BGP router identifier 10.1.1.9, local AS number 65004
BGP table version is 8, main routing table version 8
7 network entries using 1008 bytes of memory
8 path entries using 672 bytes of memory
4/4 BGP path/bestpath attribute entries using 640 bytes of memory
3 BGP AS-PATH entries using 72 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
BGP using 2392 total bytes of memory
BGP activity 7/0 prefixes, 8/0 paths, scan interval 60 secs

Neighbor      V      AS MsgRcvd MsgSent  TblVer  InQ  OutQ Up/Down  State/PfxRcd
10.1.1.1      4      65001    85     87      8     0    0 01:15:12      0
10.1.1.5      4      65001    85     88      8     0    0 01:15:13      0
10.1.1.10     4      65005    55     59      8     0    0 00:45:00      5
ORANGE#
```

La **table BGP** est constituée de toutes les routes possibles transmises par les paires, les routes sont accompagnées de leurs attributs (MED, LocPrf, WeiGGHT, AS_PATH etc....). C'est à partir de cette table que le routeur va choisir les meilleures routes qui seront utilisées en fonction de leurs attributs.

```
ORANGE#show ip bgp
BGP table version is 8, local router ID is 10.1.1.9
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale, m multipath, b backup-path, f RT-Filter,
               x best-external, a additional-path, c RIB-compressed,
               t secondary path,
Origin codes: i - IGP, e - EGP, ? - incomplete
RPKI validation codes: V valid, I invalid, N Not found

Network        Next Hop        Metric LocPrf Weight Path
*> 10.1.1.0/30   0.0.0.0          0           32768 i
*> 10.1.1.4/30   0.0.0.0          0           32768 i
* 10.1.1.8/30    10.1.1.10        0           0 65005 i
*> 10.1.1.12/30  0.0.0.0          0           32768 i
*> 10.1.1.12/30  10.1.1.10        0 65005 65006 65003 i
*> 10.1.1.16/30  10.1.1.10        0 65005 65006 65003 i
*> 10.1.1.20/30  10.1.1.10        0           0 65005 i
*> 10.1.1.24/30  10.1.1.10        0 65005 65006 i
ORANGE#
```

La **table routage IP** n'est pas uniquement utilisée par BGP, mais aussi par les autres protocoles de routage, elle contient donc toutes les routes utilisées par le routeur. Elle est mise à jour avec les meilleures routes contenues dans la table BGP, pour chacune des destinations.

```
ORANGE#show ip route
Codes: L - local, C - connected, S - static, R - RIP, M - mobile, B - BGP
        D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
        N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
        E1 - OSPF external type 1, E2 - OSPF external type 2
        i - IS-IS, su - IS-IS summary, L1 - IS-IS level-1, L2 - IS-IS level-2
        ia - IS-IS inter area, * - candidate default, U - per-user static route
        o - ODR, P - periodic downloaded static route, H - NHRP, l - LISP
        a - application route
        + - replicated route, % - next hop override, p - overrides from PfR

Gateway of last resort is not set

    10.0.0.0/8 is variably subnetted, 10 subnets, 2 masks
C       10.1.1.0/30 is directly connected, GigabitEthernet0/1
L       10.1.1.2/32 is directly connected, GigabitEthernet0/1
C       10.1.1.4/30 is directly connected, GigabitEthernet0/0
L       10.1.1.6/32 is directly connected, GigabitEthernet0/0
C       10.1.1.8/30 is directly connected, GigabitEthernet0/3
L       10.1.1.9/32 is directly connected, GigabitEthernet0/3
B       10.1.1.12/30 [20/0] via 10.1.1.10, 00:42:44
B       10.1.1.16/30 [20/0] via 10.1.1.10, 00:42:44
B       10.1.1.20/30 [20/0] via 10.1.1.10, 00:43:15
B       10.1.1.24/30 [20/0] via 10.1.1.10, 00:42:44
ORANGE#
```

Nous avons détaillé l'utilisation et le fonctionnement de BGP mais il reste un point crucial à aborder, sa sécurisation. En effet, BGP étant le protocole de routage le plus utilisé à l'échelle d'internet, son exposition est maximale et son bon fonctionnement est critique. Les grandes menaces sont l'usurpation de route (ou Hijacking), le spoofing et l'injection de routes malveillantes. Il existe aujourd'hui trois grandes manières de sécuriser l'utilisation de BGP, le BGPsec, le RPKI et l'authentification.

BGPsec repose sur l'utilisation d'une infrastructure de clés publiques et privées, dans l'objectif de permettre à chacun des AS traversés d'ajouter une signature numérique. Cela offre donc la possibilité de vérifier l'authenticité de bout en bout du chemin, mais aussi d'empêcher toute injection et/ou modification de routes non autorisées durant le transit. Cette solution de sécurité est extrêmement lourde à mettre en œuvre du fait que tous les AS traversés doivent être compatibles.

RPKI pour « Ressource Public Key Infrastructure », s'appuie sur des clés publiques qui mettent en relation des blocs d'IP publiques et les AS autorisés à leur utilisation. Quand une route est réceptionnée par un autre AS, elle est donc en mesure de contrôler via le préfixe IP que l'annonce est légitime. Cette sécurité permet de limiter grandement l'usurpation de routes. La complexité de mise en place est due à des problèmes de comptabilité pour la prise en charge des certificats et le niveau de gestion et de coordination des registres de blocs IP.

BGPsec et RPKI combinés représentent une véritable couche de sécurité renforcée qui, bien qu'applicable sur iBGP, ne présente pas le plus grand intérêt, sa mise en œuvre sera le plus souvent réservée à eBGP.

La sécurisation de iBGP, bien que fortement conseillée, est souvent moins répandue. L'authentification des paires BGP par mot de passe, MD5 (obsolète aujourd'hui) et SHA256 est le minimum requis.

Maintenant que BGP n'a plus aucun secret pour vous, nous allons voir comment nous pouvons l'intégrer dans notre architecture à trois niveaux. Dans le cadre d'infrastructures de très grande taille, iBGP peut être utilisé pour interconnecter la couche « Core » et la couche « Aggregation », tout comme d'autres IGP (IS-IS, EIGRP, OSPF) pourraient être utilisés. Pour eBGP, il sera exclusivement utilisé sur la couche « Core », afin d'effectuer la connexion avec l'AS de notre opérateur ou d'autres partenaires. Mais il pourrait également être utilisé dans le cadre d'une redondance de lien WAN avec plusieurs opérateurs pour l'optimisation des flux et de la disponibilité.

Dans l'infrastructure que nous allons mettre en place dans ce cours, nous utiliserons BGP pour la connexion avec l'opérateur.

QCM BGP

a) A quoi sert BGP ?

- ☐ Il permet de faire du NAT64.
- ☐ C'est un protocole de routage statique.
- ☐ Il gère le routage à travers plusieurs systèmes autonomes.
- ☐ Il permet de se connecter en SSH à un serveur.

b) Quelles sont les deux déclinaisons de BGP ?

- ☐ iBGP/eBGP
- ☐ xBGP/IBGP
- ☐ iBGP/uBGP
- ☐ eBGP/uBGP

c) Quelle est la Distance Administrative (AD) et le TTL par défaut de iBGP ?

- ☐ 200/225
- ☐ 210/220
- ☐ 210/230
- ☐ 200/210

d) Sur quel port sont envoyés les messages OPEN en TCP pour établir la relation de voisinage de BGP ?

- ☐ 178
- ☐ 179
- ☐ 176
- ☐ 177

e) Dans le paquet KEEPALIVE de BGP à quoi sert la composante « Network Layer Reachability Information » ?

- ☐ Indique les adresses IP ainsi que leur masque de sous-réseaux
- ☐ Indique la liste des attributs
- ☐ Indique la liste des routes à supprimer de la table de routage.
- ☐ La longueur en octets des attributs pour les routes nouvellement envoyées.

Correction QCM BGP

a) A quoi sert BGP ?

- ☐ Il permet de faire du NAT64.
- ☐ C'est un protocole de routage statique.
- ☒ **Il gère le routage à travers plusieurs systèmes autonomes.**
- ☐ Il permet de se connecter en SSH à un serveur.

b) Quelles sont les deux déclinaisons de BGP ?

- ☐ xBGP/IBGP
- ☒ **iBGP/eBGP**
- ☐ iBGP/uBGP
- ☐ eBGP/uBGP

c) Quelle est la Distance Administrative (AD) et le TTL par défaut de iBGP ?

- ☒ **200/225**
- ☐ 210/220
- ☐ 210/230
- ☐ 200/210

d) Sur quel port sont envoyés les messages OPEN en TCP pour établir la relation de voisinage de BGP ?

- ☐ 178
- ☒ **179**
- ☐ 176
- ☐ 177

e) Dans le paquet KEEPALIVE de BGP à quoi sert la composante « Network Layer Reachability Information » ?

- ☒ **Indique les adresses IP ainsi que leur masque de sous-réseaux.**
- ☐ Indique la liste des attributs.
- ☐ Indique la liste des routes à supprimer de la table de routage.
- ☐ La longueur en octets des attributs pour les routes nouvellement envoyées.

[Vidéo de de mise en œuvre BGP \(DRIVE\)](#)

[Vidéo de de mise en œuvre BGP \(Youtube\)](#)

Fiche TP BGP

1. Sur tous les routeurs opérateurs, configurer les numéros d'AS de la manière suivante :

Routeur	Numéro AS
Orange	65001
Bouygues	65002

SFR	65003
FREE	65004

2. Sur tous les routeurs opérateurs, configurer les réseaux adjacents de la manière suivante :

Routeur	Réseaux adjacents
Orange	- 10.1.1.0/30 - 10.1.1.4/30 - 10.1.1.8/30
Bouygues	- 10.1.1.8/30 - 10.1.1.20/30
SFR	- 10.1.1.20/30 - 10.1.1.24/30
FREE	- 10.1.1.24/30 - 10.1.1.28/30 - 10.1.1.32/30

3. Sur tous les routeurs opérateurs, configurer les paires de voisinage avec les informations suivantes :

Routeur	Voisin	AS distant	Description	Mot de passe	Interface source	Timer
Orange	10.1.1.10	65001	vers Bouygues	BouyguesBGP	Gi0/3	30 90
SFR	10.1.1.21	65003	vers Bouygues	SfrBGP	Gi0/0	30 90
SFR	10.1.1.26	65004	vers Free	FreeBGP	Gi0/3	30 90
Free	10.1.1.25	65003	vers SFR	FreeBGP	Gi0/3	30 90
Bouygues	10.1.1.9	65001	vers Orange	BouyguesBGP	Gi0/3	30 90
Bouygues	10.1.1.22	65003	vers SFR	SfrBGP	Gi0/0	30 90

4. Pour toutes les paires de voisinage, activer le mode « Reconfiguration souple ».

[Télécharger Correction TP BGP](#)

II – Architecture « Spine Leaf »

A- Présentation du principe d'architecture

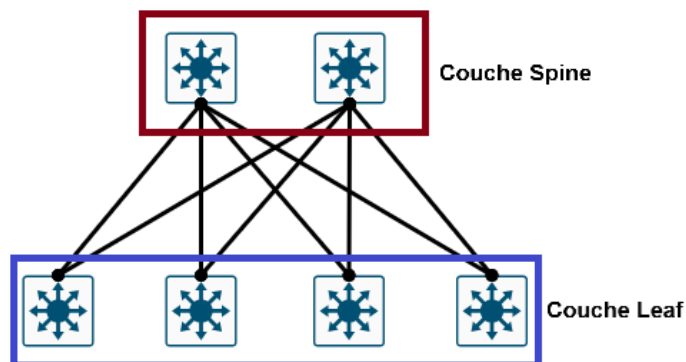
Durant de nombreuses années, l'architecture à 3 niveaux s'est imposée comme la seule architecture des DATACENTER. Cependant, avec l'avènement du Cloud Computing, de la virtualisation et de la conteneurisation, les besoins en scalabilité horizontale ont été accrus. L'architecture à 3 niveaux, bien que théoriquement scalable, est extrêmement complexe à faire évoluer en raison des nombreux protocoles utilisés pour sa mise en œuvre. De plus, le flux horizontal est confronté à des latences importantes. Les experts réseaux ont travaillé pour mettre en place un nouveau principe d'architecture réseau, c'est comme cela que le principe d'architecture « Spine Leaf » est apparu début 2010. Ce nouveau modèle d'architecture a su séduire les grandes entreprises comme Google, CISCO et Amazon dès son apparition.

En vue de simplifier l'architecture et de réduire la latence, le nombre de couches a été réduit à deux. La réduction du nombre de protocoles a été imaginée via la suppression de quasiment tous les protocoles de niveau 2, afin de ne garder qu'un protocole de routage dynamique capable d'acheminer le trafic avec une résilience et un balancement de charge. Les rares protocoles de niveau 2 qui sont utilisés au sein de ce modèle d'architecture sont les protocoles de redondance de lien comme PAgP et LACP. Cela peut être utile pour s'assurer de la résilience des liaisons entre la couche « Spine » et « Leaf ».

Chaque « Spine » étant connecté à tous les « Leaf » afin de former une topologie quasiment « full-mesh », cela requiert un grand nombre de ports. L'utilisation de PAgP et LACP est donc peu répandue du fait de la consommation significative de ports nécessaires pour leur mise en œuvre.

Les protocoles de niveau 3 s'approchant au plus près de clients finaux, cette architecture est donc uniquement réservée aux DATACENTER, contrairement à l'architecture « Core, Aggregation, Accès » qui peut aussi bien être utilisée pour des DATACENTER que pour des entreprises.

L'utilisation des protocoles de niveau 3, au plus près des serveurs, a grandement contribué à la démocratisation du SDN (Software Defined Networking). En effet, les hyperviseurs ont intégré des solutions SDN comme VmWare NSX où OpenStack Neutron afin d'interagir directement avec des protocoles de niveau 3, mais également de virtualiser complètement la couche 2. L'implémentation du SDN offre la possibilité d'automatiser les déploiements, avec l'utilisation de l'API de VmWare NSX. Prenons le cas plus concret d'un hébergement mutualisé en DATACENTER sur un hyperviseur VmWare. Ansible ou Terraform pourront être utilisés pour déployer une VM et son réseau virtuel en requêtant l'API de NSX.



Couche Spine :

La couche « Spine » gère l'interconnexion de notre infrastructure avec les ressources WAN, mais également l'interconnexion des « Leaf » entre eux à l'aide d'un protocole de routage dynamique (EIGRP, IS-IS, OSPF). Cette couche ne gère aucun mécanisme de sécurité à proprement parler, en dehors d'une isolation logique comme avec l'utilisation de VRF (Virtual Routing and Forwarding). La sécurisation des flux WAN se fait généralement par l'ajout d'une couche de firewall en amont.

Couche Leaf :

La couche « Leaf » gère l'accès des serveurs aux réseaux, mais également l'encapsulation et la désencapsulation des flux. En effet, l'utilisation exclusive de protocoles de niveau 3 nécessite la mise en œuvre d'interconnexion via VXLAN et EVPN, etc... La sécurité gérée sur cette couche se limite à la sécurisation des accès aux réseaux, via l'utilisation d'ACL et de « Port Sécurité ».

Cette architecture est modifiée par les entreprises pour répondre à leurs besoins, la modification la plus connue et utilisée est l'ajout d'une couche « Core », qui permet l'interconnexion de plusieurs domaines « Spine Leaf ». Cela est par exemple utilisé pour l'interconnexion de deux DATACENTER présents sur un même site.

B- Avantages et inconvénients de l'architecture

Les avantages de l'architecture « Spine Leaf » sont les suivants :

- **Simplification :** L'un des atouts majeurs de cette infrastructure est la simplification de mise en œuvre et de la maintenance, qui est possible par la limitation du nombre de protocoles utilisés. En effet, la simple maîtrise d'un protocole de routage dynamique suffit, là où dans le cas d'une infrastructure à 3 niveaux, une maîtrise d'un grand nombre de protocoles est requise.
- **Évolutivité :** Le design « Spine Leaf » offre la possibilité d'une évolutivité grâce la disposition de ses équipements, mais également via la simplification de l'architecture qui permet de bénéficier d'une évolutivité plus rapide.
- **Redondance et Résilience :** La disposition quasiment full-mesh des équipements, combinée au protocole de routage dynamique offre un re-routage des flux en cas d'une défaillance d'un équipement ou d'une liaison.
- **Performance et Balancement de charge :** L'utilisation d'un protocole de routage dynamique permet de mettre en œuvre l'ECMP. L'Equal-cost multi-path est une technique de gestion du routage, qui consiste à avoir plusieurs routes avec le même coût pour une même destination. Cela permet au protocole de routage dynamique de répartir les flux réseaux sur les différents

chemins afin de profiter de la bande passante de tous les liens et ainsi d'offrir des performances maximales.

- **Modernisation et Automatisation** : L'implémentation du SDN dans ce principe d'architecture a modernisé le déploiement et la maintenance via l'automatisation d'un certain nombre de tâches.
- **Flexibilité** : La disposition des équipements, l'utilisation du SDN et la réduction du nombre de protocoles rend l'infrastructure plus flexible pour s'adapter aux besoins du quotidien.

Nous allons maintenant aborder ci-dessous les inconvénients d'une telle infrastructure :

- **Coût Initial** : La mise en place d'une infrastructure « Spine Leaf » nécessitant l'utilisation de MLS haut de gamme comme des CISCO Nexus peut rapidement engendrer un coût élevé.
- **Migration et Compatibilité** : La mise en place de « Spine Leaf » nécessite que le matériel utilisé soit compatible avec les nouveaux protocoles plus modernes comme VXLAN (apparu en 2011). De plus, les grandes différences de protocoles et de câblage entre une infrastructure « Spine Leaf » et une architecture « Core, Aggregation, Access » rendent la migration complexe, voir quasiment impossible. C'est pourquoi, bien souvent la mise en place de « Spine Leaf » est réservée à un nouveau déploiement et très peu utilisée pour une migration d'infrastructure.

C- Intégration des protocoles dans l'architecture « Spine Leaf »

Comme nous l'avons abordé précédemment, le protocole le plus important dans une architecture « Spine Leaf » est le protocole de routage dynamique. Dans certains cas, un protocole sera un excellent choix et dans d'autre, il n'aura aucun intérêt et pourrait même nuire à l'infrastructure. A noter que les seuls protocoles de routage qui seront utilisés dans une architecture « Spin Leaf » sont des IGP.

Les IGP les plus utilisés au sein de « Spine Leaf » sont iBGP, OSPF et IS-IS, au contraire EIGRP et RIP ne sont que très rarement utilisés. RIP est de nos jours considéré comme obsolète du fait de sa vitesse de convergence très lente combinée à sa limitation du nombre de sauts, son implémentation dans « Spine Leaf » n'aurait absolument aucun sens. Le protocole EIGRP bien que toujours utilisé aujourd'hui, ne présente que très peu d'intérêt dans notre architecture. Cela vient du fait que EIGRP est propriétaire CISCO, son utilisation est donc restreinte aux infrastructures entièrement CISCO. De plus, même dans le cadre d'une infrastructure entièrement CISCO, les solutions SDN telle que VmWare NSX ne prennent pas en charge EIGRP.

Nous allons maintenant aborder les éléments à considérer pour faire votre choix entre les différents protocoles de routage.

OSPF, est un IGP à « état de lien », basé sur l'algorithme de Dijkstra qui a été standardisé par l'IETF en 1987. Ce protocole standardisé dispose d'une très large compatibilité, il est pris en charge par la totalité des constructeurs de matériels réseaux. OSPF dispose d'une convergence rapide du fait d'un échange périodique des paquets LSA, qui permet de maintenir la topologie à jour en cas de changement. De plus, il est basé sur un système hiérarchique constitué d'aires qui permettent la simplification de configuration. Cependant, malgré ses divers avantages, il existe également des inconvénients à l'utilisation d'OSPF.

L'évolutivité à très grande échelle d'OSPF peut être problématique car l'envoi périodique des LSA sur une grande infrastructure peut avoir un coût non négligeable, que ce soit sur la bande passante ou sur les ressources de calcul des routeurs. De plus, la gestion d'un grand nombre d'aires OSPF peut être complexe et ralentir le déploiement de l'infrastructure.

C'est pour toutes ces raisons qu'OSPF est préféré dans les petites et moyennes architectures « Spine Leaf » disposant d'une grande variété de constructeurs de réseaux, qui nécessitent une vitesse de convergence élevée.

IS-IS, présenté précédemment dans ce cours, dispose aussi de certains avantages et inconvénients qui le rendent plus ou moins utile en fonction de la taille et des besoins d'une infrastructure « Spine Leaf ». IS-IS est un protocole extrêmement flexible, du fait de son interopérabilité avec IPv4 et IPv6 plus efficace qu'OSPFv3 mais également du fait qu'il n'est pas soumis à une segmentation aussi rigide que les aires OSPF. Il est hiérarchisé par une simple classification à 2 niveaux. La gestion des paquets d'IS-IS lui offre une vitesse de convergence supérieure à OSPF, mais aussi une consommation en ressources matériels très légère. L'utilisation de ce protocole avec MPLS est fortement appréciée, en raison de son intégration native d'un champ TLV dédié.

Bien qu'il soit standardisé, IS-IS est moins répandu qu'OSPF de ce fait, il n'est pas pris en charge par tous les constructeurs, comme notamment Fortinet. Par ailleurs, les ressources et la documentation sont limitées. De plus, l'implémentation d'IS-IS peut être plus complexe au début, surtout pour des équipes d'ingénieurs habituées à OSPF.

C'est pourquoi IS-IS est un excellent choix dans le cas de très grandes infrastructures « Spine Leaf » dans lesquelles doivent cohabiter IPv4 et IPv6 avec une convergence élevée. Il faudra cependant s'assurer des compatibilités matériels, ainsi que de la formation des équipes.

iBGP, également présenté ci-dessus peut être lui aussi mis en œuvre au sein de « Spine Leaf ». Ce protocole dispose d'une très grande scalabilité, en raison des milliers de routes qu'il peut gérer dans sa table de routage, mais également grâce aux politiques de routage et aux routes map. Le grand atout de BGP est l'annonce des routes avec les AS des opérateurs et des partenaires. Cela permet aussi d'optimiser le transit à destination du WAN dans le cas où notre infrastructure « Spine Leaf » dispose de plusieurs liaisons opérateurs. En effet, BGP sera en mesure d'identifier le meilleur chemin parmi les différents AS opérateurs pour une destination donnée.

BGP dispose d'une convergence très lente ce qui peut être à la fois positif et négatif. En effet, cela sera utile si une infrastructure nécessite une stabilité de routage, au contraire cela sera un inconvénient si notre infrastructure est soumise à des changements réguliers et temporaires. Ce protocole bien que très efficace est gourmand en ressources et nécessite des routeurs haut de gamme disposant de grandes ressources de calcul. La mise en œuvre complexe de BGP nécessite aussi un haut niveau de maîtrise.

L'utilisation de BGP est réservée aux architectures connectées à plusieurs AS, qui préfèrent la stabilité à la vitesse de détection d'incident. Bien que majoritairement déployé dans de très grandes

infrastructures, il n'est pas choquant de voir BGP dans des moyennes et petites infrastructures qui sont connectées à plusieurs AS.

III – Confrontation des deux architectures

Les deux principes d'architecture que sont « Spine Leaf » et « Core, Aggregation, Access » disposent de leurs avantages et inconvénients. Le choix du principe d'architecture et des protocoles à implémenter, doit se faire sur la base d'une analyse préalable des besoins et contraintes de l'entreprise.

Les contraintes et besoins qui peuvent influencer la conception des architectures sont les suivants :

- **Taille du réseau :** Le nombre terminaux clients que l'on souhaite connecter influence le nombre d'équipements mis en œuvre au sein de l'architecture et les protocoles utilisés. Cependant cela n'influe pas particulièrement sur le choix du principe à proprement parler.
- **Nature des flux :** flux Est-Ouest (échange de flux entre différents serveurs et/ou micro-services) et flux Nord-Sud (échange de flux entre les utilisateurs et des serveurs). La Nature du trafic va nous faire choisir un principe architecture plutôt qu'un autre. En effet « Spine Leaf » sera préféré pour les flux Est-Ouest quant à « Core, Aggregation, Access » qui sera préféré pour les flux Nord-Sud.
- **Scalabilité :** Bien que les deux principes d'architecture soient scalables, « Spine Leaf » est plus facilement évolutif que « Core, Aggregation, Access ». Cela vient majoritairement du fait que le nombre de protocoles est plus restreint sur « Spine Leaf ».
- **Performance :** La latence étant bien souvent plus importante sur les architectures « Core, Aggregation Access », « Spine Leaf » sera préféré pour les infrastructures où les performances élevées sont une priorité.
- **Redondance :** La redondance est aussi bien présente sur l'infrastructure « Core, Aggregation, Access » que sur l'infrastructure « Spine Leaf ». Il est à noter que l'architecture « Spine Leaf » est bien plus résiliente que l'infrastructure « Core, Aggregation, Access », du fait qu'elle dispose de plus d'interconnexions entre les équipements.
- **Automatisation :** L'utilisation de solutions SDN comme Cisco ACI et VMware NSX, comme vu précédemment dans ce cours, est parfaitement bien intégrée dans « Spine Leaf ». Cela permet d'automatiser la gestion courante et les modifications d'infrastructure. La gestion courante de l'architecture « Core, Aggregation, Access » est plus archaïque, du fait de l'intégration plus rare des solutions SDN.
- **Complexité de mise en œuvre :** Le déploiement d'une architecture « Core, Aggregation, Access » reste aujourd'hui plus simple. Cela provient notamment du fait que plus d'équipes IT sont formées sur le principe d'architecture. « Spine Leaf » étant plus récente, il y a moins d'acteurs opérationnels sur le marché pour ce type d'architecture.
- **Maintenance :** La maintenance de « Spin Leaf » est bien plus simple au quotidien, car il y a beaucoup moins de protocoles à maîtriser que sur l'architecture « Core, Aggregation, Access ».

- **Coût financier** : Peu importe le principe d'architecture, l'une des préoccupations majeures des directions d'entreprises reste le coût financier. Que nous choisissons de mettre en œuvre « Spine Leaf » ou « Core, Aggregation, Access » le coût financier dépendra avant tout des besoins comme le nombre de terminaux à interconnecter, les bandes passantes minimum et de la qualité des équipements (CISCO vs Netgear...).

Vous trouverez ci-dessous un tableau benchmark :

	Spin Leaf	Core, Aggregation, Access
Taille du réseau	Grande échelle, milliers d'équipements	Moyenne à petite échelle, centaines d'équipements
Nature des flux	Est-Ouest (entre serveurs)	Nord-Sud (utilisateurs vers serveurs ou Internet)
Scalabilité	Très élevée	Moyenne
Performance	Elevée	Moyenne
Redondance	Très élevée	Moyenne
Automatisation	Nombreuses possibilités	Très limitée
Complexité de mise en œuvre	Complexe	Simple
Maintenance	Simple	Peut rapidement devenir complexe
Coût financier	Variable	Variable

Intéressons-nous maintenant à des cas d'utilisation plus concrets des principes d'infrastructure.

Une entreprise pour son siège social ou une succursale utilisera « Core, Aggregation, Access ». En effet, l'utilisation du réseau sera dans la plupart des cas l'accès à des ressources web (Linkedin, Outlook, ect...) mais également à un ERP ou autre logiciel métier hébergé en DATACENTER ou sur les serveurs de l'entreprise. L'utilisation quasiment exclusive d'un trafic Nord-Sud rend l'architecture « Spine Leaf » complètement inadaptée. Cependant l'utilisation de notre architecture à 3 couches offre un moyen simple de procéder à un découpage géographique. L'entreprise aura une couche « Core » par site, une couche « Aggregation » accompagnée de ses switch « Access » pour permettre la connexion des employés.

L'utilisation d'une telle infrastructure sera également préférée pour les mêmes raisons dans le contexte d'un campus universitaire, d'un hôpital ou bien d'une usine.

Pour ce qui est des fournisseurs cloud (AWS, GCP, Azure, OVH...) l'utilisation de « Spine Leaf » répond parfaitement aux contraintes. Une faible latence ainsi qu'une grande bande passante pour les flux Est-Ouest sont nécessaires pour les interactions entre les micro-services qui permettent le fonctionnement optimal de plateformes telles que OpenStak utilisée par les fournisseurs cloud. L'utilisation de « Spine Leaf » avec une solution SDN offre une intégration efficace des nouveaux clients en leur déployant un environnement réseau cloisonné de manière automatique. Sans oublier la simplification de la gestion quotidienne via l'automatisation et la limitation du nombre de protocoles.

Le déploiement de « Spine Leaf » est également apprécié pour les « Backbone » d'opérateurs internet, les DATACENTER et les supers calculateurs d'IA.

Comme vous avez pu le constater dans ce cours, les protocoles utilisés sur les réseaux existent pour la plupart depuis les années 80. Cependant, la manière de les utiliser au sein des infrastructures évolue de plus en plus vite en fonction des besoins. L'arrivée du SDN en 2010 et aujourd'hui de l'IA nous offre de nouvelles perspectives pour automatiser et dynamiser les infrastructures réseaux afin de répondre aux défis à venir. Nous pouvons nous demander à quoi ressembleront les infrastructures réseaux autogérées par l'IA ?

[VIDEO DE CONCLUSION \(DRIVE\)](#)

[VIDEO DE CONCLUSION \(Youtube\)](#)

[QCM FINAL](#)

1. **L'architecture « Spine Leaf » est plus ancienne que l'architecture « Core, Aggregation, Access ».**

- ☐ Vrai
- ☐ Faux

2. **Laquelle de ces couches de l'architecture « Core, Aggregation, Access » n'a généralement pas besoin de VLAN ?**

- ☐ Core
- ☐ Aggregation
- ☐ Access

3. **La couche « Core » dans une architecture traditionnelle peut être fusionnée avec la couche « Aggregation » pour des raisons budgétaires.**

- ☐ Vrai
- ☐ Faux

4. **Quel est l'objectif principal des Private VLAN (PVLAN) ?**

- ☐ Augmenter la bande passante disponible
- ☐ Sécuriser le trafic entre les VLAN
- ☐ Isoler les hôtes d'un même VLAN

5. **Le « Port Security » est une méthode d'authentification forte qui rend l'usurpation de MAC impossible.**

- ☐ Vrai
- ☐ Faux

6. **Quel type d'ACL est le plus souvent utilisé dans les infrastructures complexes ?**

- ☐ ACL standard
- ☐ ACL étendue
- ☐ ACL numérotées
- ☐ ACL nommée

7. **Les PVLAN sont principalement utilisés dans des environnements nécessitant une sécurité et une confidentialité élevées.**

- ☐ Vrai
- ☐ Faux

8. **Quel mode de configuration de port LACP permet d'envoyer activement des LACPDU pour établir l'agrégation ?**

- ☐ Actif
- ☐ Passif
- ☐ On
- ☐ Desirable

9. **Une ACL standard permet de filtrer le trafic en fonction des adresses MAC source et destination.**

- ☐ Vrai
- ☐ Faux

10. **Quel algorithme de répartition de charge LACP est recommandé pour un environnement LAN où un grand nombre de serveurs répondent à de multiples clients ?**

- ☐ MAC source
- ☐ MAC destination
- ☐ IP source
- ☐ IP destination

11. Le protocole VTP, bien que pratique pour la gestion des VLAN, peut présenter une faille de sécurité s'il est mal configuré.

- ☐ Vrai
- ☐ Faux

12. Quel protocole de redondance de Gateway permet un équilibrage de charge en utilisant un routeur « AVG » et des « Forwarder » ?

- ☐ VRRP
- ☐ GLBP
- ☐ HSRP
- ☐ Les trois

13. Dans une agrégation LACP, le switch ayant la priorité la plus élevée est élu switch acteur.

- ☐ Vrai
- ☐ Faux

14. Quel est l'état d'un port Spanning-Tree qui transmet activement les données et qui a été élu « Designated Port » ?

- ☐ Blocking
- ☐ Listening
- ☐ Forwarding
- ☐ Learning

15. Le protocole DTP, s'il est activé sur tous les ports d'un switch, peut permettre à un attaquant de prendre le contrôle du réseau.

- ☐ Vrai
- ☐ Faux

16. Quel protocole de Spanning-Tree est propriétaire CISCO et crée un arbre par VLAN ?

- ☐ RSTP
- ☐ MSTP
- ☐ RPVST+
- ☐ SPB

17. L'authentification 802.1x « Client-Based » offre une gestion des accès plus fine que l'authentification « Port-Based ».

- ☐ Vrai
- ☐ Faux

18. Quel type de routeur IS-IS fait le lien entre les routeurs de niveau 1 et les routeurs de niveau 2 ?

- ☐ Routeur de niveau 1
- ☐ Routeur de niveau 2
- ☐ Routeur de niveau 1 et 2
- ☐ Routeur de niveau 3

19. Un port Trunk ne peut pas être configuré avec un VLAN natif.

- ☐ Vrai
- ☐ Faux

20. Quel attribut BGP est utilisé pour indiquer la longueur du chemin parcouru par une route à travers différents AS ?

- ☐ Weight
- ☐ Local Preference
- ☐ AS Path
- ☐ MED

21. Quel est l'avantage principal de l'architecture « Spine Leaf » par rapport à l'architecture traditionnelle à 3 niveaux ?

- ☐ Coût initial plus faible
- ☐ Meilleure gestion de la sécurité
- ☐ Simplification de l'architecture et évolutivité accrue

22. GLBP est un protocole propriétaire CISCO qui permet d'assurer une redondance de Gateway et un équilibrage de charge.

- ☐ Vrai
- ☐ Faux

23. Quel type de protocole de routage est généralement utilisé dans une architecture « Spine Leaf » ?

- ☐ Un EGP
- ☐ Un IGP
- ☐ Les deux types de protocoles

24. VRRP et GLBP utilisent tous les deux des adresses MAC virtuelles pour rediriger le trafic en cas de panne.

- ☐ Vrai
- ☐ Faux

25. Lequel de ces protocoles de routage est le moins adapté à une architecture « Spine Leaf » ?

- ☐ OSPF
- ☐ IS-IS
- ☐ RIP
- ☐ iBGP

26. Le protocole MSTP est compatible avec les protocoles STP et RSTP.

- ☐ Vrai
- ☐ Faux

27. Quel protocole de routage est le plus adapté pour une grande infrastructure « Spine Leaf » utilisant IPv4 et IPv6 ?

- ☐ OSPF
- ☐ IS-IS
- ☐ iBGP
- ☐ EIGRP

28. Le « Root Path Cost » est calculé en fonction du nombre de sauts entre un switch et le « Root Bridge ».

- ☐ Vrai
- ☐ Faux

29. Quel est l'inconvénient majeur de BGP dans une architecture « Spine Leaf » ?

- ☐ Scalabilité limitée
- ☐ Convergence lente
- ☐ Manque de flexibilité

30. La fonctionnalité « PortFast » permet d'accélérer la convergence STP sur les ports connectés aux clients finaux.

- ☐ Vrai
- ☐ Faux

31. Quel mécanisme de sécurité STP permet de bloquer la réception de BPDU sur les ports connectés aux clients finaux ?

- ☐ BPDU Guard
- ☐ BPDU Filter
- ☐ Root Guard
- ☐ Loop Guard

32. RPVST+ est un protocole standardisé par l'IEEE qui permet de créer un arbre Spanning-Tree par VLAN.

- ☐ Vrai
- ☐ Faux

33. Quel mode de configuration PAgP est équivalent au mode « Actif » de LACP ?

- ☐ On
- ☐ Auto
- ☐ Desirable

34. IS-IS est un protocole de routage qui peut être utilisé pour l'interconnexion de la couche « Aggregation » et « Core » dans une architecture « Spine Leaf ».

- ☐ Vrai
- ☐ Faux

35. Quel est le type d'authentification le plus sécurisé pour les protocoles de routage comme IS-IS et BGP ?

- ☐ Mot de passe en clair
- ☐ Hachage MD5
- ☐ Hachage SHA256
- ☐ RPKI

36. BGP est un protocole de routage complexe qui nécessite une configuration manuelle des relations de voisinage.

- ☐ Vrai
- ☐ Faux

37. Quel est l'objectif principal de la couche « Leaf » dans une architecture « Spine Leaf » ?

- ☐ Interconnexion avec les ressources externes
- ☐ Gestion de la sécurité et du filtrage des flux
- ☐ Accès des serveurs aux réseaux et encapsulation des flux

38. L'architecture « Spine Leaf » est moins évolutive que l'architecture traditionnelle à 3 niveaux.

- ☐ Vrai
- ☐ Faux

39. Quel est le principal défi lors de la migration d'une architecture traditionnelle à 3 niveaux vers une architecture « Spine Leaf ».

- ☐ Le coût élevé du nouveau matériel
- ☐ La complexité de la migration et les problèmes de compatibilité
- ☐ La formation des équipes d'ingénieurs

40. Le SDN est souvent utilisé dans les architectures « Spine Leaf » pour automatiser certaines tâches et améliorer la flexibilité.

- ☐ Vrai
- ☐ Faux

Correction QCM FINAL

1. L'architecture « Spine Leaf » est plus ancienne que l'architecture « Core, Aggregation, Access ».

- ☐ Vrai
☒ Faux

2. Laquelle de ces couches de l'architecture « Core, Aggregation, Access » n'a généralement pas besoin de VLAN ?

- ☒ Core
☐ Aggregation
☐ Access

3. La couche « Core » dans une architecture traditionnelle peut être fusionnée avec la couche « Aggregation » pour des raisons budgétaires.

- ☒ Vrai
☐ Faux

4. Quel est l'objectif principal des Private VLAN (PVLAN) ?

- ☐ Augmenter la bande passante disponible
☐ Sécuriser le trafic entre les VLAN
☒ Isoler les hôtes d'un même VLAN

5. Le « Port Security » est une méthode d'authentification forte qui rend l'usurpation de MAC impossible.

- ☐ Vrai
☒ Faux

6. Quel type d'ACL est le plus souvent utilisé dans les infrastructures complexes ?

- ☐ ACL standard
☒ ACL étendue
☐ ACL numérotées
☐ ACL nommée

7. Les PVLAN sont principalement utilisés dans des environnements nécessitant une sécurité et une confidentialité élevées.

☒ **Vrai**

☐ Faux

8. Quel mode de configuration de port LACP permet d'envoyer activement des LACPDU pour établir l'agrégation ?

☒ **Actif**

☐ Passif

☐ On

☐ Desirable

9. Une ACL standard permet de filtrer le trafic en fonction des adresses MAC source et destination.

☐ Vrai

☒ **Faux**

10. Quel algorithme de répartition de charge LACP est recommandé pour un environnement LAN où un grand nombre de serveurs répondent à de multiples clients ?

☐ MAC source

☒ **MAC destination**

☐ IP source

☐ IP destination

11. Le protocole VTP, bien que pratique pour la gestion des VLAN, peut présenter une faille de sécurité s'il est mal configuré.

☒ **Vrai**

☐ Faux

12. Quel protocole de redondance de Gateway permet un équilibrage de charge en utilisant un routeur « AVG » et des « Forwarder » ?

☐ VRRP

☒ **GLBP**

☐ HSRP

☐ Les trois

13. Dans une agrégation LACP, le switch ayant la priorité la plus élevée est élu switch acteur.

- ☐ Vrai
- ☒ **Faux**

14. Quel est l'état d'un port Spanning-Tree qui transmet activement les données et qui a été élu « Designated Port » ?

- ☐ Blocking
- ☐ Listening
- ☒ **Forwarding**
- ☐ Learning

15. Le protocole DTP, s'il est activé sur tous les ports d'un switch, peut permettre à un attaquant de prendre le contrôle du réseau.

- ☒ **Vrai**
- ☐ Faux

16. Quel protocole de Spanning-Tree est propriétaire CISCO et crée un arbre par VLAN ?

- ☐ RSTP
- ☐ MSTP
- ☒ **RPVST+**
- ☐ SPB

17. L'authentification 802.1x « Client-Based » offre une gestion des accès plus fine que l'authentification « Port-Based ».

- ☒ **Vrai**
- ☐ Faux

18. Quel type de routeur IS-IS fait le lien entre les routeurs de niveau 1 et les routeurs de niveau 2 ?

- ☐ Routeur de niveau 1
- ☐ Routeur de niveau 2
- ☒ **Routeur de niveau 1 et 2**
- ☐ Routeur de niveau 3

19. Un port Trunk ne peut pas être configuré avec un VLAN natif.

- ☐ Vrai
- ☒ **Faux**

20. Quel attribut BGP est utilisé pour indiquer la longueur du chemin parcouru par une route à travers différents AS ?

- ☐ Weight
- ☐ Local Preference
- ☒ **AS Path**
- ☐ MED

21. Quel est l'avantage principal de l'architecture « Spine Leaf » par rapport à l'architecture traditionnelle à 3 niveaux ?

- ☐ Coût initial plus faible
- ☐ Meilleure gestion de la sécurité
- ☒ **Simplification de l'architecture et évolutivité accrue**

22. GLBP est un protocole propriétaire CISCO qui permet d'assurer une redondance de Gateway et un équilibrage de charge.

- ☒ **Vrai**
- ☐ Faux

23. Quel type de protocole de routage est généralement utilisé dans une architecture « Spine Leaf » ?

- ☐ Un EGP
- ☒ **Un IGP**
- ☐ Les deux types de protocoles

24. VRRP et GLBP utilisent tous les deux des adresses MAC virtuelles pour rediriger le trafic en cas de panne.

- ☒ **Vrai**
- ☐ Faux

25. Lequel de ces protocoles de routage est le moins adapté à une architecture « Spine Leaf » ?

- ☐ OSPF
- ☐ IS-IS
- ☒ **RIP**
- ☐ iBGP

26. Le protocole MSTP est compatible avec les protocoles STP et RSTP.

- ☒ **Vrai**
- ☐ Faux

27. Quel protocole de routage est le plus adapté pour une grande infrastructure « Spine Leaf » utilisant IPv4 et IPv6 ?

- ☐ OSPF
- ☒ **IS-IS**
- ☐ iBGP
- ☐ EIGRP

28. Le « Root Path Cost » est calculé en fonction du nombre de sauts entre un switch et le « Root Bridge ».

- ☐ Vrai
- ☒ **Faux**

29. Quel est l'inconvénient majeur de BGP dans une architecture « Spine Leaf » ?

- ☐ Scalabilité limitée
- ☒ **Convergence lente**
- ☐ Manque de flexibilité

30. La fonctionnalité « PortFast » permet d'accélérer la convergence STP sur les ports connectés aux clients finaux.

- ☒ **Vrai**
- ☐ Faux

31. Quel mécanisme de sécurité STP permet de bloquer la réception de BPDU sur les ports connectés aux clients finaux ?

- ☒ **BPDU Guard**
- ☐ BPDU Filter
- ☐ Root Guard
- ☐ Loop Guard

32. RPVST+ est un protocole standardisé par l'IEEE qui permet de créer un arbre Spanning-Tree par VLAN.

- ☐ Vrai
- ☒ **Faux**

33. Quel mode de configuration PAgP est équivalent au mode « Actif » de LACP ?

- ☐ On
- ☐ Auto
- ☒ **Desirable**

34. IS-IS est un protocole de routage qui peut être utilisé pour l'interconnexion de la couche « Aggregation » et « Core » dans une architecture « Spine Leaf ».

- ☒ **Vrai**
- ☐ Faux

35. Quel est le type d'authentification le plus sécurisé pour les protocoles de routage comme IS-IS et BGP ?

- ☐ Mot de passe en clair
- ☐ Hachage MD5
- ☒ **Hachage SHA256**
- ☐ RPKI

36. BGP est un protocole de routage complexe qui nécessite une configuration manuelle des relations de voisinage.

- ☒ **Vrai**
- ☐ Faux

37. Quel est l'objectif principal de la couche « Leaf » dans une architecture « Spine Leaf » ?

- ☐ Interconnexion avec les ressources externes
- ☐ Gestion de la sécurité et du filtrage des flux
- ☒ **Accès des serveurs aux réseaux et encapsulation des flux**

38. L'architecture « Spine Leaf » est moins évolutive que l'architecture traditionnelle à 3 niveaux.

- ☐ Vrai
- ☒ **Faux**

39. Quel est le principal défi lors de la migration d'une architecture traditionnelle à 3 niveaux vers une architecture « Spine Leaf ».

- ☐ Le coût élevé du nouveau matériel
- ☒ **La complexité de la migration et les problèmes de compatibilité**
- ☐ La formation des équipes d'ingénieurs

40. Le SDN est souvent utilisé dans les architectures « Spine Leaf » pour automatiser certaines tâches et améliorer la flexibilité.

- ☒ **Vrai**
- ☐ Faux

Annexe

Bibliographie

Site internet :

- Blog de cours réseau « networklessons.com » de Rene MOULENAAR
- Blog officiel FS « fs.com »
- Documentation officiel Cisco « cisco.com »
- Documentation officiel Nokia Network « documentation.nokia.com »
- Documentation officiel HP « support.hpe.com »

Ouvrage :

- Support de cours CCNA, CCNP et réseaux opérateur de Erwan GUILLEMOT
- CISCO CCNP Enterprise Design Network EENSLD 300-420 2nd Edition
- CISCO CCNP and CCIE Enterprise Core & CCNP Enterprise Advanced Routing Portable Commande Guide
- CISCO CCIE Routing and Switching v5.1 Foundation

Vidéo :

- Playlist « ENCOR 350-401 Complete Course » de Jeremy's IT Lab

Glossaire

Terme	Définition
LAN	Local Area Network
VLAN	Virtual Local Area Network
ACL	Access Control List
QoS	Quality of Service
OSI	Open Systems Interconnection
OSPF	Open Shortest Path First
EIGRP	Enhanced Interior Gateway Routing Protocol
IS-IS	Intermediate system to intermediate system
BGP	Border Gateway Protocol

IPSEC	Internet Protocol Security
MPLS	Multiprotocol Label Switching
VXLAN	Virtual Extensible LAN
MD5	Message-Digest Algorithm 5
SHA256	Secure Hash Algorithm 256
MLS	Multilayer Switch
LACP	Link Aggregation Control Protocol
PAgP	Port Aggregation Protocol
RSTP	Rapid Spanning Tree Protocol
VRRP	Virtual Router Redundancy Protocol
POE	Power over Ethernet
IP	Internet Protocol
WI-FI	Wireless Fidelity
STP	Spanning Tree Protocol
MAC	Media Access Control
DHCP	Dynamic Host Configuration Protocol
GLBP	Gateway Load Balancing Protocol
WAN	Wide Area Network
IT	Information Technology
ISL	Inter-Switch-Link
IEEE	Institute of Electrical and Electronics Engineers
PVID	Port VLAN ID
PVLAN	Private VLAN
DTP	Dynamic Trunk Protocol
VTP	VLAN Trunking Protocol
SYSLOG	System Logging Protocol
SNMP	Simple Network Management Protocol
TCP	Transmission Control Protocol

UDP	User Datagram Protocol
ICMP	Internet Control Message Protocol
LACPDU	Link Aggregation Control Protocol Data Unit
PAgPDU	Port Aggregation Protocol Data Unit
TLV	Type Length Value
HSRP	Hot Standby Router Protocol
IETF	Internet Engineering Task Force
RFC	Request for comments
VRID	Virtual Router ID
AVG	Active Virtual Gateway
AVF	Active Virtual Forwarder
ARP	Address Resolution Protocol
OUI	Organizationally Unique Identifier
VOIP	Voix sur IP
RPVST+	Rapid Per VLAN Spanning Tree+
MSTP	Multiple Spanning Tree Protocol
PVST	Per-VLAN Spanning Tree
BPDU	Bridge Protocol Data Units
TCN BPDU	Topology Change Notification BPDU
MSTI	Multiple Spanning Tree Protocol Instance
IST	Internal Spanning-Tree
CIST	Common and Internal Spanning Tree
IGP	Interior Gateway Protocol
LSDB	Link State Database
PDU	Protocol Data Unit
DIS	Designated Intermediate System
CSNP	Complete Sequence Number PDU
LSP	Link State PDU

PSNP	Partial Sequence Number PDU
P2P	Point-To-Point
AS	Autonomous System
EGP	Exterior Gateway Protocol
TTL	Time To Live
CIDR	Classless Inter-Domain Routing
iBGP	Interior BGP
eBGP	Exterior BGP
AD	Distance Administrative
MED	Multi-Exit Discriminator
RPIK	Resource Public Key Infrastructure
SDN	Software Defined Networking
VM	machine virtuelle
API	application programming interface
VRF	Virtual Routing and Forwarding
EVPN	Ethernet VPN
ECMP	Equal-cost multi-path
RIP	Routing Information Protocol
LSA	Link State Advertisement
ERP	Enterprise Resource Planning
IA	Intelligence artificielle