

الجمهورية الجزائرية الديمقراطية الشعبية
République Algérienne démocratique et populaire

وزارة التعليم العالي والبحث العلمي
Ministère de l'enseignement supérieur et de la recherche scientifique

جامعة سعد دحلب البليدة
Université SAAD DAHLAB de BLIDA

كلية التكنولوجيا
Faculté de Technologie

قسم الإلكترونيك
Département d'Électronique

Réf : RT11



Mémoire de Master

Filière : Télécommunications
Spécialité : Réseaux & Télécommunications

Présenté par :
Zerrouk Meriem dina

L'implémentation d'un réseau IP-MPLS

Proposé par :
Promoteur : Dr. YAHIAOUI AMINA
Mr. GUENAB TARIK

Année Universitaire : 2024-2025

REMERCEMENTS

Avant tout, je veux remercier ALLAH pour m'avoir donné la patience pendant les moments difficiles, la force pour surmonter les défis, et la volonté de continuer jusqu'à la fin de ce projet.

Un grand merci à Madame Yahiaoui, ma directrice de mémoire, pour son soutien, sa disponibilité et ses bons conseils tout au long de cette aventure.

Je tiens aussi à remercier chaleureusement OOREDOO pour m'avoir accueilli lors de mon stage de fin d'études. Un merci spécial à Monsieur Tarik Guenab pour sa supervision, son soutien constant et son professionnalisme durant mon temps là-bas.

Je suis reconnaissante envers les membres du jury pour l'honneur qu'ils me font en évaluant ce travail et pour leurs remarques utiles.

Enfin, je veux exprimer ma gratitude à ma famille pour leur soutien moral et leur présence inébranlable, ainsi qu'à mes amis et camarades pour tous les bons moments partagés, les encouragements et la bonne humeur tout au long de ce parcours. Votre aide a été indispensable pour mener à bien ce travail.

DEDICACES

À ma chère mère Naïma et mon père Aliouat, merci pour tout l'amour et le soutien que vous m'avez donné au fil des ans. Ce mémoire est le résultat de vos encouragements et de votre présence à mes côtés. Que Dieu vous garde en bonne santé.

À mes sœurs Isra et Zineb, et à mon frère Moumen, je vous suis reconnaissante pour votre tendresse et votre soutien moral qui m'ont aidé dans les moments difficiles.

À toute ma famille, merci pour vos prières et votre affection.

À mes amies, je vous remercie pour votre présence et votre motivation tout au long de ce parcours.

Enfin, un grand merci à Madame Yahiaoui, ma super encadrante, pour son aide précieuse et ses conseils qui ont vraiment aidé à faire avancer ce travail.

Zerrouk Meriem dina

المُلخَص

تُعدّ تقنية تبديل التسمية متعددة البروتوكولات حلاً فعالاً لمواجهة المتطلبات المتزايدة من حيث السرعة، والموثوقية، وإدارة حركة البيانات. يعرض هذا البحث عملية تنفيذ وتحقيق من بنية شبكة متكاملة تدعم خدمات الشبكات الخاصة الافتراضية من الطبقة الثالثة، بهدف ضمان الترابط الآمن بين مواقع العملاء التي تتميز بمتطلبات عالية من حيث الأداء، والمرونة، وقابلية التوسع. يشمل العمل تفعيل وظائف تبديل التسمية متعددة البروتوكولات على أجهزة شبكات من نوع CISCO باستخدام المحاكى الرسومي للشبكات – الإصدار الثالث بالإضافة إلى إعداد شبكات الخاصة الافتراضية من المستوى الثالث والتحقق من أدائها من خلال سيناريوهات اختبار واقعية.

الكلمات المفتاحية : 3GNS, BGP-MP, VRF, MPLS/IP

Résumé

La technologie MPLS (Multiprotocol Label Switching) est apparue comme une solution efficace face aux exigences croissantes en termes de rapidité, de fiabilité et de gestion du trafic. Ce mémoire présente un processus d'implémentation et de validation d'une architecture IP/MPLS intégrant des services L3VPN, permettant d'assurer l'interconnexion sécurisée de sites clients avec une forte exigence de performance, de résilience et d'évolutivité, en mettant en œuvre les fonctionnalités MPLS sur des équipements réseau CISCO dans le simulateur GNS3, de configurer les VPN de niveau 3 et de vérifier leur bon fonctionnement à travers des scénarios de test concrets.

Mots clés : IP/MPLS, VRF, MP-BGP, GNS3.

Abstract

MPLS (Multiprotocol Label Switching) technology has emerged as an effective solution to meet the growing demands for speed, reliability, and traffic management. This work presents the implementation and validation process of an IP/MPLS architecture integrating Layer 3 VPN (L3VPN) services, aimed at ensuring secure interconnection of customer sites with high requirements for performance, resilience, and scalability. The project involves deploying MPLS functionalities on CISCO network devices within the GNS3 simulator, configuring Layer 3 VPNs, and verifying their correct operation through concrete test scenarios.

Key words : IP/MPLS, VRF, MP-BGP, GNS3.

Table des matières

Liste des Figures

Liste des Tableaux

Liste des Abréviations

INTRODUCTION GENERALE1

Chapitre I . Présentation de l'organisme d'accueil Erreur ! Signet non défini.

I . 1 Introduction Erreur ! Signet non défini.

I . 2 Historique d'Ooredoo Algérie Erreur ! Signet non défini.

I . 3 Présentation de Ooredoo Algérie Erreur ! Signet non défini.

I . 4 Le Réseau de Ooredoo Erreur ! Signet non défini.

I . 5 Les valeurs de Ooredoo..... Erreur ! Signet non défini.

I . 6 Evolution du logo Ooredoo..... Erreur ! Signet non défini.

I . 7 Conclusion Erreur ! Signet non défini.

Chapitre II . Fondements Théoriques sur les Réseaux Erreur ! Signet non défini.

II . 1 Introduction..... Erreur ! Signet non défini.

II . 2 Définition des réseaux Erreur ! Signet non défini.

II . 3 Types de réseau Erreur ! Signet non défini.

II . 3 .1 Réseau local (LAN Local Area Network) Erreur ! Signet non défini.

II . 3 .2 Réseau étendu (WAN Wide Area Network) Erreur ! Signet non défini.

II . 3 .3 Réseau métropolitain (MAN Metropolitan Area Network) Erreur ! Signet non défini.

II . 3 .4 Réseau personnel (PAN *Personal Area Network*) Erreur ! Signet non défini.

II . 4 Topologie des réseaux Erreur ! Signet non défini.

II . 4 .1 Topologie logique Erreur ! Signet non défini.

II . 4 .2 Topologie physique Erreur ! Signet non défini.

II . 5 Modèles de référence Erreur ! Signet non défini.

II . 5 .1 Modèle OSI Erreur ! Signet non défini.

II . 5 .2 Modèle de référence TCP/IP Erreur ! Signet non défini.

II . 6 Protocoles Réseaux Erreur ! Signet non défini.

II . 6 .1 Types de protocole Erreur ! Signet non défini.

II . 6 .2 Interaction entre les protocoles Erreur ! Signet non défini.

II . 6 .3 Protocole TCP/IP Erreur ! Signet non défini.

II . 6 .3 .a Couche application Erreur ! Signet non défini.

II . 6 .3 .b Couche transport Erreur ! Signet non défini.

II . 6 .3 .c Couche internet Erreur ! Signet non défini.

II . 6 .3 .d Couche accès réseau Erreur ! Signet non défini.

II . 7 Les équipements réseaux Erreur ! Signet non défini.

II . 7 .1	Le concentrateur (<i>Hub</i>)	Erreur ! Signet non défini.
II . 7 .2	Le commutateur (<i>Switch</i>)	Erreur ! Signet non défini.
II . 7 .3	Le routeur	Erreur ! Signet non défini.
II . 7 .4	Modem	Erreur ! Signet non défini.
II . 8	Routage	Erreur ! Signet non défini.
II . 8 .1	Le routage statique	Erreur ! Signet non défini.
II . 8 .2	Le routage dynamique	Erreur ! Signet non défini.
II . 8 .3	Protocoles de routage.....	Erreur ! Signet non défini.
II . 8 .3 .a	RIP (Routing Information Protocol)	Erreur ! Signet non défini.
II . 8 .3 .b	IGRP (Interior Gateway Routing Protocol)	Erreur ! Signet non défini.
II . 8 .3 .c	EIGRP (Enhanced Interior Gateway Routing Protocol)	Erreur ! Signet non défini.
II . 8 .3 .d	OSPF (Open Shortest Path First)	Erreur ! Signet non défini.
II . 8 .3 .e	IS-IS (Intermediate System to Intermediate System)	Erreur ! Signet non défini.
II . 8 .3 .f	EGP (Exterior Gateway Protocol)	Erreur ! Signet non défini.
II . 8 .3 .g	BGP (Border Gateway Protocol)	Erreur ! Signet non défini.
II . 8 .4	Table de routage	Erreur ! Signet non défini.
II . 9	Commutation	Erreur ! Signet non défini.
II . 10	Conclusion	Erreur ! Signet non défini.
Chapitre III .	Principe Réseau IP-MPLS	Erreur ! Signet non défini.
III . 1	Introduction	Erreur ! Signet non défini.
III . 2	Technologie MPLS	Erreur ! Signet non défini.
III . 3	Éléments et Terminologie du MPLS.....	Erreur ! Signet non défini.
III . 3 .1	Élément du MPLS	Erreur ! Signet non défini.
III . 3 .2	Terminologie utilisée	Erreur ! Signet non défini.
III . 4	Définition labels	Erreur ! Signet non défini.
III . 4 .1	Fonctionnement des labels dans MPLS	Erreur ! Signet non défini.
III . 4 .1 .a	Label Switching (commutation par label)	Erreur ! Signet non défini.
III . 4 .1 .b	Label forwarding (transmission par label)	Erreur ! Signet non défini.
III . 4 .1 .c	Label distribution (Distribution des labels)	Erreur ! Signet non défini.
III . 4 .2	Les protocoles de distribution label	Erreur ! Signet non défini.
III . 5	Format de l'en-tête MPLS	Erreur ! Signet non défini.
III . 6	Application de la technologie MPLS	Erreur ! Signet non défini.
III . 6 .1	L'ingénierie du trafic (<i>TE</i>)	Erreur ! Signet non défini.
III . 6 .1 .a	Types de tunnels dans MPLS-TE ..	Erreur ! Signet non défini.
III . 6 .2	Qualité de service (<i>QOS</i>)	Erreur ! Signet non défini.
III . 7	Réseau Privés Virtuels dans MPLS	Erreur ! Signet non défini.
III . 7 .1	Virtuel routage et forwarding VRF	Erreur ! Signet non défini.

III . 7 . 1 . a VRF Forwarding Table	Erreur ! Signet non défini.
III . 7 . 1 . b Propagation des informations de routage VPN	Erreur ! Signet non défini.
III . 7 . 2 Multiprotocol BGP (MP-BGP)	Erreur ! Signet non défini.
III . 7 . 3 Le Transfert des Paquets IP (IP Packet Forwarding).....	Erreur ! Signet non défini.
III . 8 Conclusion	Erreur ! Signet non défini.
Chapitre IV . Simulation d'un Réseau IP/MPLS sous GNS3.....	Erreur ! Signet non défini.
IV . 1 Introduction.....	Erreur ! Signet non défini.
IV . 2 Présentation de l'outil de simulation GNS3	Erreur ! Signet non défini.
IV . 3 Installation du logiciel GNS3	Erreur ! Signet non défini.
IV . 4 Topologie physique de notre plateforme	Erreur ! Signet non défini.
IV . 5 Configuration des interfaces des routeurs	Erreur ! Signet non défini.
IV . 6 Configuration de protocole de routage	Erreur ! Signet non défini.
IV . 7 Configuration de MPLS.....	Erreur ! Signet non défini.
IV . 8 Configuration de VRF	Erreur ! Signet non défini.
IV . 9 Configuration de MP-BGP	Erreur ! Signet non défini.
IV . 10 Conclusion	Erreur ! Signet non défini.
CONCLUSION GENERALE	60
Références	Erreur ! Signet non défini.
ANNEXE 1. Installation de GNS3	Erreur ! Signet non défini.

Liste des Figures

Chapitre I : Présentation de l'organisme d'accueil

FIGURE I-1: LOGO DE OOREDOO ALGERIE DE 2004 A 2009	41
FIGURE I-2: LOGO DE OOREDOO ALGERIE DE 2010 A 2013	ERREUR ! SIGNET NON DEFINI.
FIGURE I-3: LOGO DE OOREDOO ALGERIE DEPUIS 2014	ERREUR ! SIGNET NON DEFINI.

Chapitre II : Fondements Théoriques Des Réseaux

FIGURE II-1: LES PRINCIPAUX TYPES DE RESEAUX	ERREUR ! SIGNET NON DEFINI.
FIGURE II-2: TOPOLOGIES DE RESEAUX	ERREUR ! SIGNET NON DEFINI.
FIGURE II-3: MODELE DE REFERENCE	ERREUR ! SIGNET NON DEFINI.
FIGURE II-4: INTERACTION ENTRE LES PROTOCOLES	ERREUR ! SIGNET NON DEFINI.
FIGURE II-5: MODELÉ TCP /IP	30
FIGURE II-6: EQUIPMENT PHYSIQUE HUB	33
FIGURE II-7: EQUIPMENT PHYSIQUE SWITCH	34
FIGURE II-8: EQUIPMENT PHYSIQUE ROUTEUR.....	35

Chapitre III : Principe Réseau IP-MPLS

FIGURE III-1: MPLS TECHNOLOGIE DE COUCHE 2.5	42
FIGURE III-2: DETAILS D'UN LABEL MPLS	45

Chapitre IV : Simulation d'un réseau IP-MPLS sous GNS3

FIGURE IV-1: LA TOPOLOGIE SIMULEE	54
FIGURE IV-2: AFFICHAGE LES INTERFACES ET LES ADRESSES IP DU R1.	59
FIGURE IV-3: AFFICHAGE LES INTERFACES ET LES ADRESSES IP DU R3.....	60
FIGURE IV-4: AFFICHAGE LES INTERFACES ET LES ADRESSES IP DU R2.....	60
FIGURE IV-5: AFFICHAGE LES INTERFACES ET LES ADRESSES IP DU R8.....	61
FIGURE IV-6: TEST DE CONNECTIVITE ENTRE R3 ET DIFFERENTES ADRESSES (COMMANDE PING).	61
FIGURE IV-7: ETAT DES INTERFACES OSPF SUR R1	63
FIGURE IV-8: AFFICHAGE LES VOISINS OSPF DETECTES PAR R1	63
FIGURE IV-9: TABLE DE ROUTAGE DU R1	64
FIGURE IV-10: AFFICHAGE L'ETAT DES INTERFACES OSPF SUR R3.....	65
FIGURE IV-11: AFFICHAGE LES VOISINS OSPF DETECTES PAR R3.....	65
FIGURE IV-12: AFFICHAGE LE TABLE DE ROUTAGE DU R3	65
FIGURE IV-13 : VERIFICATION DE LA CONNECTIVITE ENTRE R3 ET L'INTERFACE LOOPBACK DE R1	66
FIGURE IV-14 : AFFICHAGE L'ETAT DES INTERFACES OSPF SUR R2.....	66
FIGURE IV-15 : AFFICHAGE LES VOISINS OSPF DETECTES PAR R2	67
FIGURE IV-16 : TABLE DE ROUTAGE DU R2	67
FIGURE IV-17 : AFFICHAGE L'ETAT DES INTERFACES OSPF SUR R8.....	68
FIGURE IV-18 : AFFICHAGE LES VOISINS OSPF DETECTES PAR R8	68
FIGURE IV-19 : TABLE DE ROUTAGE DU R8	68
FIGURE IV-20 : VISUALISE LES INTERFACES MPLS ACTIVEES SUR R1	70
FIGURE IV-21 : AFFICHAGE LES VOISINS LDP DETECTES PAR R1	70
FIGURE IV-22 : VISUALISE LES INTERFACES MPLS ACTIVEES SUR R3	71
FIGURE IV-23 : AFFICHAGE LES VOISINS LDP DETECTES PAR R3	72
FIGURE IV-24 : VISUALISE LES INTERFACES MPLS ACTIVEES SUR R2	73

FIGURE IV-25 : AFFICHAGE LES VOISINS LDP DETECTES PAR R2	73
FIGURE IV-26 : VISUALISE LES INTERFACES MPLS ACTIVEES SUR R8	74
FIGURE IV-27 : AFFICHAGE LES VOISINS LDP DETECTES PAR R8	74
FIGURE IV-28 : MONTRE LA CONFIGURATION DES VRF SUR R1	76
FIGURE IV-29 : MONTRE LA CONFIGURATION DES VRF SUR R3	77
FIGURE IV-30 : AFFICHAGE LA TABLE DE ROUTAGE ASSOCIEE A LA VRF « CLIENT » SUR R1....	79
FIGURE IV-31 : AFFICHAGE LA TABLE DE ROUTAGE ASSOCIEE A LA VRF « CLIENT » SUR R3....	80
FIGURE IV-32 : CONFIGURATION DU ROUTAGE STATIQUE SUR R4 ET TEST DE CONNECTIVITE	81
FIGURE IV-33 : CONFIGURATION DU ROUTAGE STATIQUE SUR R5 ET TEST DE CONNECTIVITE	ERREUR ! SIGNET NON DEFINI. 81

ANNEX : INSTALLATION DE GNS3

FIGURE ANNEX-1 : Interface principale de GNS3	87
FIGURE ANNEX-2 : MENU « Edit » OUVERT AVEC « Préférences » sélectionné	88
FIGURE ANNEX-3 : FENETRE « GENERAL PREFERENCES » DANS GNS3	88
FIGURE ANNEX-4 : DEBUT DU PROCESSUS D'AJOUT D'UN ROUTEUR IOS	89
FIGURE ANNEX-5 : SELECTION DU FICHIER IMAGE IOS DANS GNS3	90
FIGURE ANNEX-6 : IMAGE IOS SELECTIONNEE POUR CREER UN NOUVEAU MODELE DE ROUTEUR	90
FIGURE ANNEX-7 : PARAMETRES GENERAUX DU MODELE	91
FIGURE ANNEX-8 : AFFICHE LES DETAILS DU MODELE DANS LA FENETRE « IOS ROUTER TEMPLATE » DU GNS3	91
FIGURE ANNEX-9 : PRESENTE LA FENETRE « Router » DU GNS3 AVEC L'AJOUT DU ROUTEUR R1 DANS LE PROJET	92

Liste des Tableaux

TABLEAU IV -1 : ADRESSES DES INTERFACES DES ROUTEURS.....	41
---	----

Liste des Abréviations

ARP	Address Resolution Protocol
AS	Autonomous System
ATM	Asynchronous Transfer Mode
BGP	Border Gateway Protocol
CE	Customer Edge
CEF	Cisco Express Forwarding
DHCP	Dynamic Host Configuration Protocol
DNS	Domain Name System
DHCPv4	Dynamic Host Configuration Protocol for IPv4
DHCPv6	Dynamic Host Configuration Protocol for IPv6
EIGRP	Enhanced Interior Gateway Routing
EGP	Exterior Gateway Protocol
FTP	File Transfer Protocol
FAI	Fournisseur d'Accès à Internet
FIB	Forwarding Information Base
FEC	Forwarding Equivalence Class
GSM	Global System for Mobile Communications
GNS3	Graphical Network Simulator 3
GNS3 VM	GNS3 Virtual Machine
HSPA+	Evolved High Speed Packet Access
HTTP	Protocol HyperText Transfer Protocol
IP	Internet Protocol

ICMP	Internet Control Message
IPv4	Internet Protocol version 4
IPv6	Internet Protocol version 6
IMAP	Internet Message Access Protocol
ICMPv4	Internet Control Message Protocol version 4
ICMPv6	Internet Control Message Protocol version 6
ICMPv6 ND	ICMPv6 Neighbor Discovery
ISP	Internet Service Provider
IGP	Interior Gateway Protocol
IGRP	Interior Gateway Routing Protocol
IS-IS	Intermediate System to Intermediate System
LAN	Local Area Network
LSR	Label Switch Router
LER	Label Edge Router
LSP	Label Switched Path
LIB	Label Information Base
LFIB	Label Forwarding Information Base
LDP	Label Distribution Protocol
L2VPN	Layer 2 Virtual Private Network
L3VPN	Layer 3 Virtual Private Network
MPLS	Multiprotocol Label Switching
MAN	Metropolitan Area Network
MAC	Media Access Control Address
MP-BGP	Multiprotocol Border Gateway protocol
NAT	Network Address Translation

OSI	Open Systems Interconnection		
OSPF	Open Shortest Path First		
ONT	Optical Network Terminal		
PAN	Personal Area Network		
POP3	Post Office Protocol version 3		
PR	Provider Router		
PE	Provider Edge Router		
QOS	Qualité of Service		
REST	Representational State Transfer		
RJ45	Registered Jack 45		
RIP	Routing Information Protocol		
RSTP	Rapid Spanning Tree Protocol		
SSH	Secure Shell		
SSL	Secure Sockets Layer		
SMTP	Simple Mail Transfer Protocol		
SFTP	SSH File Transfer Protocol		
STP	Spanning Tree Protocol		
TCP	Transmission Control Protocol		
TLS	Transport Layer Security		
TFTP	Trivial File Transfer Protocol		
TTL	Time To Live		
TE	Traffic Engineering		
UDP		User	Datagram
VPN	Virtual Private Network		
VLAN	Virtual Local Area Network		

VRF	Virtual Routing and Forwarding
VPNv4	VPN version 4
WAN	Wide Area Network
WLAN	Wireless Local Area Network

INTRODUCTION GENERALE

INTRODUCTION GENERALE

Les Réseaux informatiques sont nés du besoin de faire communiquer des terminaux distants avec un site central puis des ordinateurs entre eux. Dans un premier temps, ces communications étaient limitées aux transports de données informatiques alors qu'aujourd'hui on se dirige plutôt vers des réseaux qui intègrent à la fois des données mais en plus, la parole, et la vidéo.

Avec l'augmentation rapide du volume des données échangées, les infrastructures classiques rencontrent des difficultés pour répondre aux exigences croissantes en termes de rapidité, de fiabilité et de gestion du trafic. Face à ces défis, la technologie MPLS (Multiprotocol Label Switching) est apparue comme une solution efficace. Contrairement aux méthodes traditionnelles de routage, MPLS utilise un système d'étiquetage qui permet d'accélérer la transmission des données et d'optimiser la gestion du trafic. Cette technologie apporte plusieurs avantages, notamment la qualité de service (QOS), la gestion avancée du trafic (Traffic Engineering) et la mise en place de réseaux privés virtuels (VPN).

L'usage des protocoles IP en télécommunications sert principalement à assurer la convergence des diverses technologies 2G 3G et 4G via le réseau IP-MPLS. Avec la croissance continue du trafic IP et la diversification des besoins des entreprises multisites, les opérateurs téléphoniques sont confrontés à la nécessité de fournir des services réseau performants, flexible et sécurisés. Les solutions classiques de routage ou de VPN GRE atteignent rapidement leurs limites en termes de scalabilité, d'isolation des clients et de gestion centralisée. Dans ce contexte, l'intégration d'une architecture IP-MPLS combinée aux VPN de niveau 3 (L3VPN) s'impose comme une réponse adaptée, spécialement dans les réseaux de fournisseurs de service nécessitant l'isolation des différents clients.

L'objectif de ce travail est de concevoir, implémenter et valider une architecture IP-MPLS intégrant des services L3VPN, permettant d'assurer l'interconnexion sécurisée de sites clients avec une forte exigence de performance, de résilience et d'évolutivité. Il s'agira de mettre en œuvre les fonctionnalités MPLS (label switching, PE/P, VRF et MP-BGP) sur des équipements réseau CISCO dans le simulateur GNS3, de configurer les VPN de niveau 3 et de vérifier leur bon fonctionnement à travers des scénarios de test concrets.

INTRODUCTION GENERALE

Pour atteindre nos objectifs, nous avons suivi une méthodologie structurée autour de quatre chapitres :

Chapitre 1 : Présentation de l'organisme d'accueil, dans lequel nous présentons l'entreprise d'accueil tout en soulignant son historique, ses valeurs et sa vision.

Chapitre 2 : Fondements Théoriques Des Réseaux, propose une présentation des bases théoriques sur les réseaux informatiques ainsi que les différents protocoles de routage.

Chapitre 3 : Réseau IP-MPLS, exploite en détail la technologie MPLS, en expliquant ses divers composants et ses applications variantes.

Chapitre 4 : Simulation d'un Réseau IP-MPLS sous GNS3, décrit le processus d'une implémentation pratique d'un réseau IP-MPLS, en expliquant la topologie choisie et les différentes configurations nécessaires pour un fonctionnement idéal.

Chapitre I. Présentation de l'organisme d'accueil

Chapitre I . Présentation de l'organisme d'accueil

I . 1 Introduction

Il est important de présenter l'organisme d'accueil, au sein duquel, nous menons notre travail. Dans ce chapitre, nous allons présenter l'entreprise Ooredoo Algérie, un grand nom dans le secteur des télécommunications en Algérie. Nous proposons un aperçu de son histoire, décrivons ses principales activités, et examinons son réseau, ses valeurs ainsi que l'évolution de son image de marque. L'objectif est de mieux comprendre ses fondements, ses missions, et son rôle stratégique dans le paysage numérique du pays.

I . 2 Historique d'Ooredoo Algérie

Ooredoo Algérie, avant appelée Nedjma, est un opérateur de télécommunications qui a vu le jour grâce à Wataniya Telecom Algérie, une branche du groupe koweïtien Wataniya Telecom. Ce groupe fait partie de KIPCO, l'un des grands groupes privés du Koweït. En 2004, Wataniya a obtenu la troisième licence de téléphonie mobile en Algérie pour une période de 15 ans, lançant ainsi ses activités sous la marque Nedjma [1].

La société s'est vite fait remarquer grâce à ses gros investissements dans le réseau et les services mobiles, ce qui l'a aidée à se faire un nom dans l'innovation et les offres multimédias accessibles. En 2007, l'opérateur qatari Qtel, maintenant connu sous le nom d'Ooredoo Group, a pris une part majoritaire dans Wataniya Telecom, obtenant ainsi une grande part dans Nedjma.

En novembre 2013, un moment important se produit à Alger lors d'une conférence de presse. La marque Nedjma devient officiellement Ooredoo, se rapprochant ainsi de l'identité du groupe international. Ce changement d'image arrive en même temps que le lancement de la 3G en Algérie, ce qui transforme les télécommunications mobiles pour l'opérateur [1].

Depuis tout ce temps, Ooredoo Algérie continue d'élargir ses services et s'investir dans les nouvelles technologies comme la 4G/5G, tout en cherchant à assurer une communication moderne et efficace à ces clients.

I . 3 Présentation de Ooredoo Algérie

Grace à son approche innovative, et son excellent service, Ooredoo Algérie a gagné vite sa place dans le marché Algérien. L'opérateur propose une variété de services, comme la téléphonie mobile, l'internet rapide, et des solutions numériques pour les particuliers et les entreprises. Sa bonne infrastructure, lui assure de bien répondre aux demandes croissantes des utilisateurs tout en assurant une couverture réseau sur l'ensemble du pays.

Ooredoo Algérie investit dans de nouvelles technologies pour aider à la transformation digitale du pays. L'entreprise se démarque par son engagement envers le développement des gens à travers la communication, tout en suivant la vision du groupe Ooredoo, qui est présent dans plusieurs pays du Moyen-Orient, d'Afrique du Nord et d'Asie du Sud-Est [2].

Ooredoo Algérie est souvent saluée pour la qualité de ses services et son rôle important dans le secteur des télécoms en Algérie, ce qui renforce sa place sur le marché.

I . 4 Le Réseau de Ooredoo

Ooredoo Algérie a mis en place un réseau solide qui couvre une bonne partie du pays. Ils utilisent surtout la technologie GSM sur les bandes 900 et 1800 MHz, ce qui permet de garantir un service de qualité aussi bien en ville qu'à la campagne. Avec l'arrivée récente de la troisième génération HSPA+, la vitesse et la fiabilité des connexions mobiles se sont améliorées, ce qui répond aux besoins croissants des utilisateurs pour des services Internet plus rapides [2]. Ce développement technologique aide Ooredoo à offrir une connectivité qui reste accessible à tous .

I . 5 Les valeurs de Ooredoo

Depuis son lancement officiel fin 2013, Ooredoo s'engage à renforcer les valeurs fondamentales qui guident sa stratégie et ses relations avec ses clients et partenaires [2]. Ces principes reflètent l'esprit d'innovation et de responsabilité sociale qui caractérise l'entreprise :

Caring : Mettre le respect et l'écoute au centre de chaque interaction, en mettant l'accent sur la satisfaction des clients et le bien-être des employés.

Connecter : Il s'agit de rassembler les personnes et les organisations en proposant des solutions technologiques qui aident à bâtir une communauté.

Défi : viser l'excellence en cherchant toujours à s'améliorer et à proposer de nouvelles idées, afin de pouvoir relever les défis du secteur des télécommunications.

Ces valeurs forment la base sur laquelle Ooredoo construit sa réputation en Algérie, tout en jouant un rôle actif dans le progrès social et économique grâce à ses services.

I . 6 Evolution du logo Ooredoo

Au fil des ans, Ooredoo Algérie a changé de logo trois fois, ainsi illustre dans les figures ci-dessous . De 2004 à 2009, quand l'entreprise s'appelait Nedjma, le logo avait une étoile sur un fond orange (figure I-1). Entre 2010 et 2013, ils ont modernisé le logo tout en gardant l'étoile (figure I-2). En 2013, l'entreprise a pris le nom de OOREDOO avec un nouveau logo rouge fait de cercles (figure I-3), montrant son lien avec le groupe international et son envie d'innover et de rester connecté [2].



Figure I-1: Logo de Ooredoo Algérie de 2004 a 2009 [2]



Figure I-2: Logo de Ooredoo Algérie de 2010 a 2013 [2]



Figure I-3: Logo de Ooredoo Algérie depuis 2014 [2]

I . 7 Conclusion

Ooredoo Algérie se démarque par son rôle dans les télécommunications et son engagement à innover, à offrir un bon service et à satisfaire ses clients. Dans ce chapitre nous avons présente un aperçu sur l'histoire, la structure, les valeurs ainsi que les aspects clés qui façonnent l'identité de l'opérateur a fin de mieux comprendre le contexte du stage et les missions qui ont été effectuées.

Chapitre II. Fondements Théoriques Des Réseaux

Chapitre II . Fondements Théoriques Des Réseaux

II . 1 Introduction

De nos jours, les réseaux informatiques sont devenus indispensables dans notre vie quotidienne. Que ce soit pour le travail, les études ou les loisirs, nous utilisons internet et les réseaux pour échanger des informations, communiquer et accéder à divers services.

Ce chapitre est un départ de connaissance des concepts élémentaires dans les réseaux informatiques. Il présente une lecture complète de la définition d'un réseau, son éventail de topologie, les modèles de référence (OSI et TCP/IP) ainsi que les principaux protocoles de communication. Il fait également suite des aspects physiques concernant un réseau, avant d'évoquer les composantes fondamentales du transport de donnée et d'échanges.

II . 2 Définition des réseaux

Un réseau informatique est un ensemble d'ordinateurs connectés entre eux afin d'échanger des données et des ressources (fichiers, imprimantes ou connexions Internet). Parmi ces ordinateurs, on trouve différents types d'ordinateurs, serveurs, ainsi que des équipements spécifiques tels que des routeurs et des commutateurs (switchs).

L'objectif d'un réseau est de faciliter la communication entre tous ces éléments, qu'ils soient situés au même endroit (réseau local) ou distants (réseau étendu).

II . 3 Types de réseau

Selon la distances qui sépare les ordinateurs, on distingue plusieurs catégories de réseaux :

II . 3 .1 Réseau local (LAN Local Area Network)

Est un type de réseau qui connecte plusieurs appareils (ordinateurs, imprimantes, smartphones, etc.) dans une zone géographique limitée, comme une maison, un bureau, une école ou une petite ou moyenne entreprise.

II . 3 .2 Réseau étendu (WAN Wide Area Network)

Est un réseau qui connecte plusieurs réseaux locaux (LAN) sur de vastes zones géographiques, comme des villes ou des pays différents.

II . 3 .3 Réseau métropolitain (MAN Metropolitan Area Network)

Il couvre une zone plus grande qu'un réseau local (LAN), mais plus petite qu'un réseau étendu (WAN), comme une ville.

II . 3 .4 Réseau personnel (PAN *Personal Area Network*)

Ce type de réseau est très petit, et connecte des appareils proches d'une seule personne, comme un téléphone avec un ordinateur [3].

La figure (II-1) illustre les trois principaux types de réseaux.

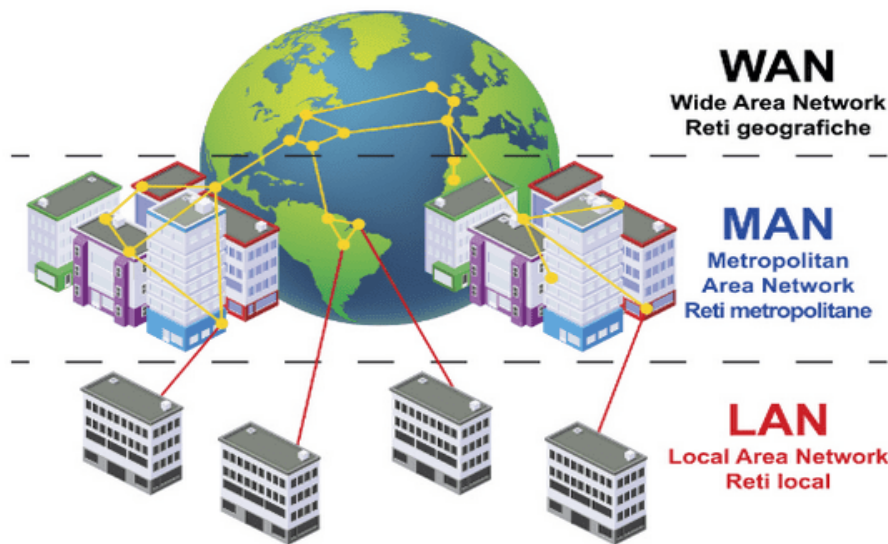


Figure II-1: Les principaux types de réseaux [3].

II . 4 Topologie des réseaux

Les diagrammes de topologie sont essentiels pour comprendre la structure d'un réseau. Ils montrent comment les appareils sont connectés et facilitent la gestion du réseau [4]. On distingue deux grands types de topologies :

II . 4 .1 Topologie logique

Représente la manière avec laquelle les données circulent sur les supports de transmission. Elle décrit la manière dont les appareils échangent des données et comment les adresses sont configurées.

II . 4 .2 Topologie physique

Représente la configuration spatiale visible du réseau. Elle indique l'emplacement des appareils, les câbles et les connexions réelles dans le réseau. On distingue plusieurs sous-types, ainsi illustrée dans la figure (II-2), a savoir :

- **Topologie en bus** : Tous les appareils partagent un seul câble central pour échanger les données.
- **Topologie en anneau** : Les équipements sont reliés en boucle fermée, les données circulent d'un appareil à l'autre.
- **Topologie en étoile** : Chaque appareil est connecté à un point central, ce qui simplifie la gestion du réseau.
- **Topologie maille** : Chaque appareil est relié à plusieurs autres, assurant une meilleure fiabilité et redondance.
- **Topologie arbre** : Combine les topologies en étoile et en bus, idéale pour les réseaux de grande taille.

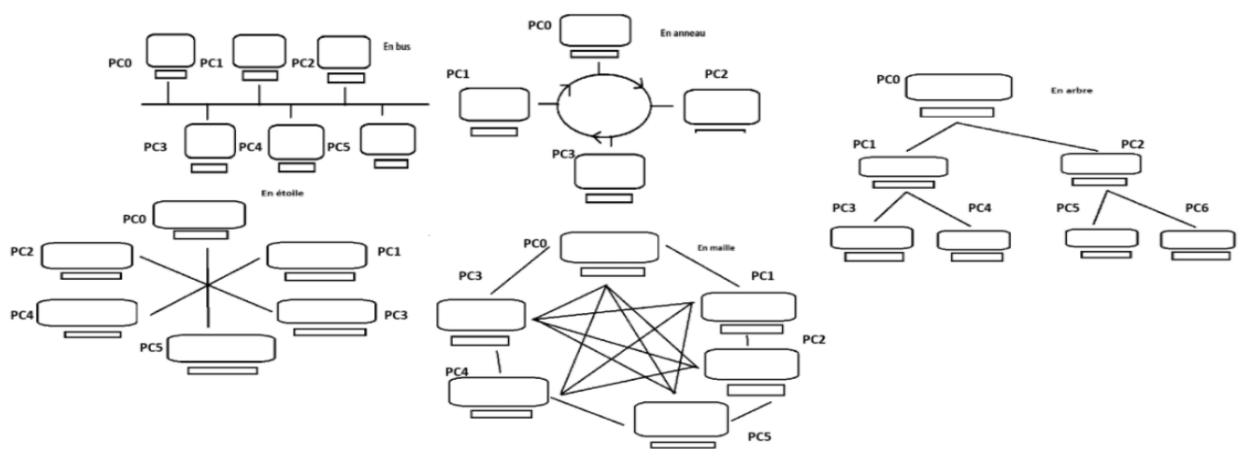


Figure II-1 Topologies de réseaux.

II . 5 Modèles de référence

A l'apparition des réseaux informatiques, chacun des constructeurs a conçu sa propre architecture de réseau et développé des protocoles propriétaires. Par conséquent, ces architectures fonctionnaient très bien dans l'environnement propre du constructeur mais la communication devenait très difficile dès qu'on veut lier des machines de constructeurs différents. Face à cette situation, en 1978, l'ISO (International Standards Organization) a

développé un modèle standard d'architecture de communication : le modèle de référence pour l'interconnexion des systèmes ouverts : Le modèle OSI (Open Systems Interconnexion). Il s'agit d'un modèle abstrait qui sert de cadre à la description des concepts utilisés et la démarche suivie pour l'interconnexion des systèmes ouverts [5].

Même si le modèle de référence OSI est universellement reconnu, historiquement et techniquement, la norme ouverte d'Internet est le protocole TCP/IP (pour Transmission Control Protocol/Internet Protocol). Le modèle de référence TCP/IP et la pile de protocoles TCP/IP rendent possible l'échange de données entre deux ordinateurs, partout dans le monde, à une vitesse quasi équivalente à celle de la lumière.

Donc, pour résumé, comme le montre la figure (II-3), il existe deux modèles en couche qui sont utilisés pour décrire les opérations réseau :

- ❖ Modèle de référence pour l'interconnexion des systèmes ouverts (OSI).
- ❖ Modèle de référence TCP/IP.

II . 5 .1 Modèle OSI

Le modèle OSI est une référence standard qui décrit le fonctionnement des réseaux en sept couches distinctes. Chaque couche du modèle joue un rôle spécifique :

- **Couche physique** : gère la transmission des signaux électriques ou optiques à travers les câbles ou les ondes
- **Couche liaison de données** : assure la communication entre deux machines directement connectées et détecte les erreurs de transmission
- **Couche réseau** : responsable du routage des paquets de données entre différents réseaux. Par exemple : IP (*Internet Protocol*), ICMP (*Internet Control Message Protocol*)
- **Couche transport** : garantit une transmission fiable des données entre les applications sur deux machines. Par exemple : TCP (*Transmission Control Protocol*), UDP (*User Datagram Protocol*).
- **Couche session** : gère l'établissement et la maintenance des connexions entre applications.
- **Couche présentation** : assure la conversion des données dans un format compréhensible par les applications (*cryptage, compression*)

➤ **Couche application** : fournit les services aux utilisateurs finaux, comme le courrier électronique et la navigation web [5].

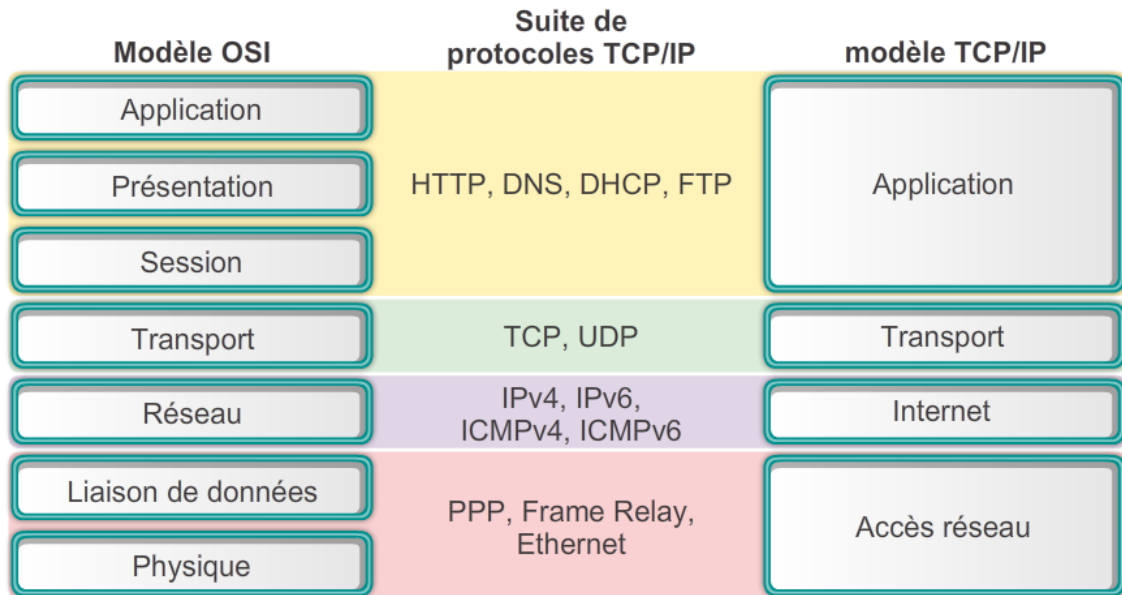


Figure 0-2 : Modèle de référence [5]

II . 5 .2 Modèle de référence TCP/IP

Ce modèle correspond étroitement à la structure d'une suite de protocoles spécifique. Le modèle TCP/IP, par exemple, décrit les fonctions exercées à chaque couche au sein de cette suite et sert également de modèle de référence.

- **Accès réseau** : contrôle les périphériques matériels et les supports qui constituent le réseau.
- **Internet** : détermine le meilleur chemin à travers le réseau.
- **Transport** : prend en charge la communication entre plusieurs périphériques à travers divers réseaux.
- **Application** : représente des données pour l'utilisateur, ainsi que du codage et un contrôle du dialogue [5].

II . 6 Protocoles Réseaux

Le protocole réseau est un ensemble de règles qui définissent comment les données sont échangées entre les appareils connectés. Grâce à ces règles, les équipements, même s'ils

proviennent de fabricants différents ou fonctionnent sous des systèmes d'exploitation varies, peuvent communiquer efficacement au sein d'un réseau [6].

II . 6 .1 Types de protocole

Les protocoles réseaux peuvent être regroupés en :

Protocole de communication de réseau : ce protocole permet aux appareils de communiquer sur différents réseaux. Exemples : IP, TCP, HTTP.

Protocole de sécurité des réseaux : il assure la sécurité des données par l'authentification, le chiffrement et l'intégrité. Exemple : SSH, SSL, TLS.

Protocole de routage : aide les routeurs à choisir le meilleur chemin pour envoyer les données. Exemple : OSPF, BGP.

Protocole de découverte de services : sert à détecter automatiquement les appareils ou services. Exemple : DHCP pour l'attribution des IP, DNS pour la traduction des noms en IP [6].

II . 6 .2 Interaction entre les protocoles

Lorsqu'un message envoyé sur un réseau informatique, il nécessite généralement l'utilisation de plusieurs protocoles, chacun avec ses propres fonctionnalités et format. La figure (II-4) montre certains protocoles réseau courants qui sont utilisés lorsqu'un périphérique envoie une demande à un serveur Web pour sa page Web [6].

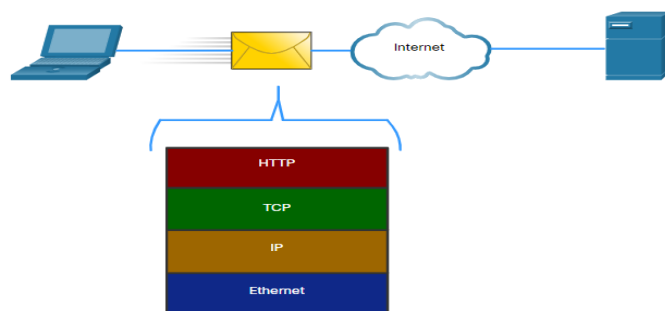


Figure II-3 Interaction entre les protocoles [6]

Les protocoles décrits dans la figure, sont :

❖ **Hypertext Transfer Protocol (HTTP)** : le protocole permet la communication entre un navigateur (*client*) et un serveur web. Il définit la structure des requêtes envoyées par le client ainsi que des réponses fournies par le serveur. Fonctionnant au niveau applicatif, il s'appuie sur d'autres protocoles pour assurer le transport des données.

❖ **Transmission Control Protocol (TCP)** : TCP est chargé d'établir et de maintenir une connexion fiable entre deux hôtes. Il garantit que les données sont transmises correctement, sans perte ni duplication, tout en régulant le flux de transmission.

❖ **Internet Protocol (IP)** : le protocole IP assure le routage des paquets entre les réseaux. Il permet de transmettre les messages depuis la source jusqu'au destinataire, même à travers plusieurs segments de réseau.

❖ **Ethernet** : Ethernet est un protocole utilisé dans les réseaux locaux (*LAN*) pour permettre l'échange de données entre dispositifs connectés au même réseau. Il définit le format des trames et la méthode d'accès au support physique [6].

II.6.3 Protocole TCP/IP

Aujourd'hui, la suite de protocoles TCP/IP inclut de nombreux protocoles et continue de s'adapter pour offrir de nouveaux services. Les protocoles les plus courants sont présentés dans la figure (II-5) ci-dessous :

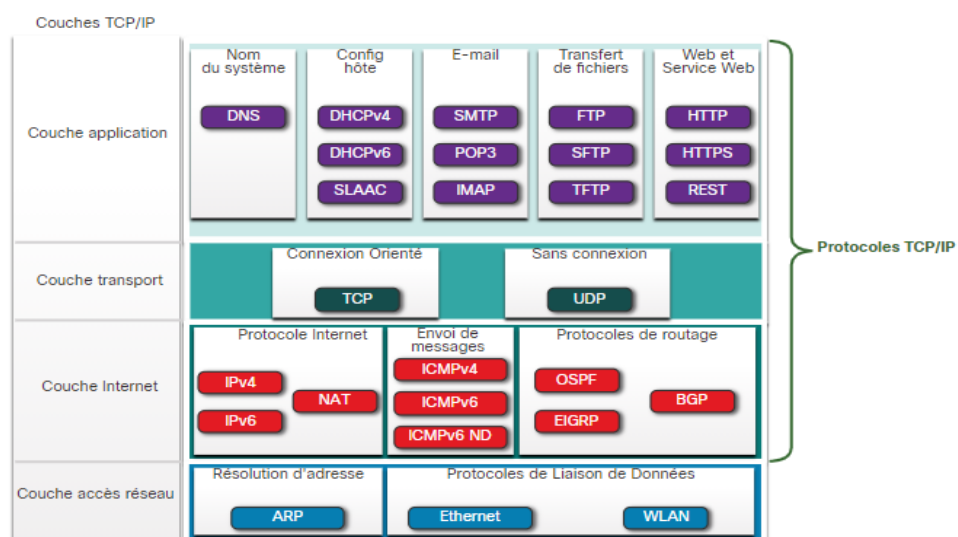


Figure II-0: modelé TCP /IP [6]

II .6 .3 .a Couche application

Nous trouvons ici les protocoles applicatifs. Ce sont des protocoles de haut niveau, destinés à permettre le dialogue entre applications serveurs et clientes. HTTP, FTP, POP et SMTP sont loin d'être les seuls. Ce sont cependant ceux que les internautes utilisent le plus souvent.

DNS : Un protocole qui traduit les noms de domaine, exemple : cisco.com, en adresses IP. Il aide à trouver des ressources sur le réseau.

DHCPv4 : Un protocole qui attribue des adresses IP et des réglages réseau aux appareils sur les réseaux IPv4.

DHCPv6 : Une version du DHCP pour IPv6 qui attribue,automatiquement, des adresses IPv6 aux appareils .

SLAAC : Une méthode qui permet à un appareil de définir son adresse IPv6 sans avoir besoin d'un serveur DHCPv6.

STMP : Utilisé pour envoyer des e-mails entre les serveurs de messagerie et les clients.

POP3 : Permet aux utilisateurs de récupérer leurs e-mails d'un serveur et de les télécharger sur leur appareil.

IMAP : Permet d'accéder aux e-mails sur un serveur tout en les gardant là.

FTP : Etablit des règles pour le transfert de fichiers entre ordinateurs sur un réseau.

SFTP : Une version sécurisée de FTP qui utilise SSH pour transférer des fichiers de manière chiffrée.

TFTP : Un moyen simple pour transférer des fichiers sans avoir besoin d'une connexion ou d'une confirmation de réception.

HTTP : Un protocole pour partager des textes, des images et des vidéos sur le web.

HTTPs : Une version sécurisée de HTTP qui chiffre les données pour mieux protéger vos informations en ligne.

REST : Création d'applications web avec des API et des requêtes HTTP.

II .6 .3 .b Couche transport

Ici, ce sont les protocoles orientés transport de données. UDP est dit "sans connexion" et TCP "est dit "avec connexion". Nous verrons plus loin ce que ceci veut dire. Ces protocoles permettent à ceux de la couche 4 de transporter leurs données de façon fiable.

TCP : C'est un protocole de communication qui assure que les données envoyées entre des systèmes éloignés arrivent bien, en vérifiant et en corrigeant d'éventuelles erreurs.

UDP : Ce protocole envoie des données rapidement entre deux systèmes, mais ne s'assure pas que les paquets arrivent ni ne propose un moyen de vérifier ou de corriger les erreurs.

II .6 .3 .c Couche internet

Ce sont ici des protocoles de haut niveau de la couche réseau : IP permet le routage des informations entre réseaux, c'est ici que l'adresse IP est utilisée, et ICMP qui est un protocole de "contrôle". Il met à disposition des outils de dépistage d'erreur et de signalisation.

IPv4 : Le protocole sert à envoyer des données sur le réseau en utilisant des adresses IP de 32 bits. Les messages sont divisés en paquets qui sont adressés, ce qui aide à assurer une bonne livraison d'un point à un autre.

IPv6 : Version améliorée d'IPv4, utilisant des adresses de 128 bit, offrant ainsi un nombre d'adresses beaucoup plus grand pour répondre aux besoins croissants des dispositifs connectés.

NAT : Technique permettant de convertir les adresses IP privées d'un réseau local en adresses publiques uniques, permettant ainsi à plusieurs appareils d'utiliser une seule adresse IP publique pour se connecter à internet.

ICMPv4 : Un protocole sert à envoyer des messages de contrôle et à donner des infos sur les erreurs de livraison de paquets dans un réseau IPv4.

ICMPv6 : Variante du protocole ICMP utilisée dans les réseaux IPv6, permettant d'informer sur les erreurs de transmission de paquets.

ICMPv6 ND : Protocole dans IPv6 utilisé pour la résolution d'adresses et la détection des doublons d'adresses, à travers plusieurs types de messages spécifiques.

OSPF : Protocole de routage interne basé sur l'état des liens, qui utilise une hiérarchie en zones pour déterminer les chemins les plus courts dans un réseau.

EIGRP : Cisco a son propre protocole de routage interne qui prend en compte plusieurs critères, comme la bande passante, le temps de réponse et la fiabilité, pour améliorer le routage.

BGP : Protocole de routage externe utilisé entre les fournisseurs de services internet (ISP) pour échanger des informations de routage, notamment pour gérer les routes entre les réseaux différents.

II .6 .3 .d Couche accès réseau

Protocole de plus bas niveau sur le réseau, il assure la bonne gestion du médium (détection de collisions) et permet l'acheminement des informations entre émetteur et destinataire au niveau des adresses MAC.

ARP : Un protocole pour résoudre les adresses. Il fait le lien entre une adresse IPv4 et une adresse physique.

Ethernet : Il fixe les règles pour le câblage et la signalisation dans la partie d'accès du réseau.

WLAN : Un réseau local sans fil fonctionne en utilisant des signaux sur les fréquences radio de 2,4 GHz et 5 GHz.

II . 7 Les équipements réseaux

II . 7 .1 Le concentrateur (*Hub*)

Le concentrateur, ou hub présenté dans la figure (II-6), est un dispositif réseau permettant de relier plusieurs machines dans une architecture en étoile. Il dispose de plusieurs ports (généralement 4, 8, 16 ou 32), il reçoit les données arrivant sur un port et les réplique sur tous les autres, sans distinction.

Étant donné qu'il diffuse à tous les équipements sans vérifier leur destination, les performances du réseau peuvent être affectées lorsque de nombreux appareils sont connectés. Pour cette raison, son utilisation a diminué au profit d'équipements plus intelligents comme le switch [7].



Figure II-4: Equipement physique Hub [7]

On distingue principalement deux types de concentrateurs :

➤ **Hub actif** : il est alimenté électriquement et joue un rôle similaire à celui d'un répéteur. Il régénère et amplifie le signal avant de le transmettre à tous les ports, assurant ainsi une meilleure qualité de transmission sur le réseau.

➤ **Hub passif** : il se contente de transmettre le signal tel quel, sans amplification ni régénération. Son rôle se limite à la diffusion des données vers tous les hôtes connectes.

II . 7 .2 Le commutateur (Switch)

Est un équipement réseau dote de plusieurs ports, permettant de connecter différents appareils via des câbles RJ45, ainsi illustré dans la figure (II-7). Contrairement au concentrateur qui diffuse les données à tous les hôtes connectes, le commutateur est capable d'identifier les destinataires et d'établir une communication directe entre les machines concernées. Cette capacité rend le switch plus performant, notamment lorsque le réseau devient dense et charge en trafic [7].



Figure II-7 : Equipement physique Switch [7]

Le switch analyse les adresses MAC source et destination des trames qu'il reçoit, et construit progressivement une table d'adresses MAC (appelée table de commutation). Cette table lui permet d'identifier précisément quel appareil est connecte à quel port. Ainsi, lorsqu'un message doit être envoyé, le switch le transmet uniquement sur le port correspondant à la machine destinataire, laissant les autres ports disponibles pour d'éventuelles communication simultanées.

Dans un réseau moderne, les commutateurs administrables permettent la création de VLANs (Virtual Local Area Networks). Un VLAN est une technologie qui permet de segmenter un réseau physique en plusieurs réseaux logiques indépendants. Cela améliore la sécurité, réduit les domaines de broadcast, et facilite la gestion du réseau. Grace aux VLANs, des utilisateurs connectes à différents ports du switch peuvent être regroupe dans un même réseau logique, même s'ils ne sont pas physiquement proches [8].

II . 7 .3 Le routeur

Est un dispositif réseau essentiel permettant l'interconnexion entre plusieurs réseaux, même s'ils sont de types différents. Il joue un rôle central dans l'acheminement des paquets en analysant les informations de la table de routage, qui associe adresses IP et ports. Grace à son système d'exploitation et au logiciel de routage intègre, il détermine le chemin optimal

que doivent suivre les données. Il constitue également un point de passage incontournable pour accéder à internet [9]. La figure (II-8) montre un exemple de routeur avec 16 ports.



Figure 0-5: Equipement physique routeur [8]

II . 7 .4 Modem

Le modem, ou modulateur-démodulateur, est un dispositif essentiel qui permet à un réseau local de se connecter à Internet. Sa fonction principale est de convertir les signaux numériques de l'ordinateur en signaux analogiques pouvant circuler sur des lignes téléphoniques ou câblées, et vice versa. Il agit donc comme un pont entre le réseau domestique et le fournisseur d'accès à Internet (FAI). Il existe plusieurs types de modems selon la technologie utilisée [9] :

- **Le modem ADSL** : utilise la ligne téléphonique pour offrir un accès Internet haut débit tout en permettant les appels.
- **Le modem câble** : utilise le réseau de télévision câblée pour transmettre les données.
- **Le modem 3G/4G/5G** : se connecte au réseau mobile pour fournir Internet, souvent utilisé dans les zones sans ligne fixe.
- **Le modem fibre optique (ONT)** : convertit les signaux lumineux de la fibre en signaux utilisables par les appareils domestiques.

Aujourd'hui, il est fréquent de trouver des modems intégrés à des routeurs, formant un seul appareil capable de se connecter à Internet et de distribuer le réseau local .

II . 8 Routage

Le routage est le principe de guider les paquets dans le réseau pour qu'ils arrivent à leur destination. Cela se fait grâce à des équipements appelés routeurs qui utilisent des règles pour choisir le bon chemin.

II . 8 .1 Le routage statique

Le routage statique repose sur une configuration manuelle des tables de routage. Une fois définies, ces tables ne changent pas, même si la topologie du réseau évolue. Ce type de routage présente plusieurs limites, notamment :

- La nécessité de mettre à jour chaque équipement manuellement
- La difficulté à gérer des routes redondantes
- Le risque de formation de boucles en cas de coupure de lien
- L'apparition possible de routages asymétriques

Malgré ces inconvénients, le routage statique reste adapté aux équipements situés en périphérie du réseau, notamment dans les systèmes autonomes, car il offre un meilleur contrôle et une sécurité renforcée [9].

II . 8 .2 Le routage dynamique

Dans le routage dynamique, les tables de routage sont mises à jour automatiquement grâce à des algorithmes de routage. Ce processus repose sur des protocoles appelés protocoles de routage, dont le rôle est de fournir les informations nécessaires pour permettre un acheminement efficace des paquets.

Pour mieux les échanges, le réseau est souvent divisé en sous-ensembles appelés systèmes autonomes (SA).

- A l'intérieur d'un même SA, on utilise des protocoles de routage interne (*Interior Gateway Protocols-IGP*)
- Entre différents SA, ce sont les protocoles de routage externe (*Exterior Gateway Protocols-EGP*) qui sont utilisés.

Ces protocoles ont été conçus pour faciliter le routage. Certains sont encore en usage, d'autres ont été remplacés ou abandonnés. Ils utilisent différents types d'algorithmes, basés soit sur la distance, soit sur l'état des liens, soit une combinaison des deux [9].

II . 8 .3 Protocoles de routage

Les protocoles de routage jouent un rôle essentiel dans le fonctionnement des réseaux. Ils permettent aux routeurs de choisir le chemin optimal pour transmettre les données vers leur destination. Ces protocoles assurent la communication entre routeurs afin de partager et mettre à jour les informations de routage. On distingue généralement deux grandes familles : les protocoles à vecteur de distance et ceux à état de lien, selon l'algorithme utilisé [10].

II .8 .3 .a RIP (Routing Information Protocol)

Le protocole RIP est l'un des plus anciens protocoles de routage. Il utilise le nombre de sauts (hops) comme métrique pour déterminer le chemin vers la destination. Bien qu'il soit simple à configurer, il présente des limitations importantes en termes de performance et de capacité, ce qui le rend adapté uniquement aux petits réseaux.

II .8 .3 .b IGRP (Interior Gateway Routing Protocol)

Développé par Cisco, IGRP a été conçu pour surmonter les limites de RIP. Il prend en compte plusieurs paramètres comme la bande passante, le délai, la fiabilité et la charge. Malgré son amélioration, IGRP est aujourd'hui considéré comme obsolète et n'est plus largement utilisé.

II .8 .3 .c EIGRP (Enhanced Interior Gateway Routing Protocol)

Également développé par Cisco, EIGRP est un protocole hybride qui combine les caractéristiques des protocoles à vecteur de distance et à état de lien. Il offre une convergence rapide et une gestion efficace du routage. Cependant, il reste propriétaire et ne fonctionne que sur les équipements Cisco.

II .8 .3 .d OSPF (Open Shortest Path First)

OSPF est un protocole à état de lien basé sur l'algorithme de Dijkstra. Il est largement utilisé dans les réseaux de taille moyenne à grande, grâce à sa capacité à calculer le chemin le plus court et à s'adapter rapidement aux changements topologiques. Il est également open source et interopérable entre différents fabricants.

II .8 .3 .e IS-IS (Intermediate System to Intermediate System)

IS-IS fonctionne de manière similaire à OSPF, utilisant aussi une approche par état de lien. Il est souvent utilisé dans les grands réseaux, notamment ceux des opérateurs de télécommunications. Moins connu que OSPF dans le monde académique, il est pourtant très robuste et évolutif.

II .8 .3 .f EGP (Exterior Gateway Protocol)

EGP est un protocole de routage externe utilisé historiquement pour l'échange d'informations entre systèmes autonomes. Il est aujourd'hui considéré comme obsolète et a été remplacé par des solutions plus modernes comme BGP.

II .8 .3 .g BGP (Border Gateway Protocol)

BGP est le protocole de routage principal de l'Internet. Il permet aux systèmes autonomes, tels que les fournisseurs d'accès, d'échanger des informations de routage. Contrairement aux autres protocoles, BGP utilise des politiques de routage basées sur des règles et des préférences, plutôt que sur des métriques simples comme la distance [11].

II . 8 .4 Table de routage

Chaque routeur a une table qui lui montre comment se connecter aux autres réseaux. Cette table peut être remplie à la main (statique) ou de manière automatique en utilisant des protocoles (dynamique).

II . 9 Commutation

C'est une technique qui permet d'envoyer des données efficacement entre différents appareils dans les réseaux informatiques. Elle utilise l'adresse MAC pour diriger chaque paquet de données vers sa destination. On la trouve surtout dans les réseaux locaux (LAN) avec un appareil appelé switch. Grâce à cette méthode, les échanges sont plus rapides et directs, ce qui évite les collisions qu'on avait avec les équipements plus anciens comme les hubs [12]. On distingue trois grands types de commutation, à savoir :

➤ **Commutation de circuit** : C'est une méthode où on crée une connexion continue entre deux points du réseau avant de commencer à échanger des données. Cette connexion reste ouverte jusqu'à ce qu'on termine la communication.

➤ **Commutation de paquets** : Consiste à diviser les données en petits paquets, qui sont ensuite transmis séparément à travers le réseau. Chaque paquet peut emprunter un chemin différent selon la disponibilité.

➤ **Commutation de cellules** : fonctionne un peu comme la commutation de paquets, mais les données sont envoyées en petites cellules de taille fixe. Utilisée surtout dans des technologies comme l'ATM [12,13].

Pour ces différentes catégories de commutation, on déroule deux types de protocoles décrits ainsi :

Protocole STP (Spanning Tree Protocol) : C'est utilisé pour éviter les boucles de commutation qui peuvent survenir quand il y a plusieurs chemins entre les équipements. Ça choisit tout seul un chemin actif et bloque les autres pour empêcher les problèmes.

Le RSTP (Rapid Spanning Tree Protocol) : C'est une version améliorée de STP, qui permet une réaction plus rapide si le réseau change, ce qui réduit le temps d'arrêt [12].

II . 10 Conclusion

Ce chapitre présente une vue d'ensemble sur les bases essentielles pour la compréhension des réseaux informatiques. Nous avons commencé par définir ce qu'est un réseau et explorer ses différentes topologies. Nous avons ensuite examiné les modèles de communication, notamment OSI et TCP/IP, ainsi que les protocoles utilisés pour garantir les échanges de données. Par la suite, nous avons étudié les éléments physiques qui composent un réseau, avant de nous pencher sur les mécanismes fondamentaux de routage et de commutation. Ces concepts posent les bases nécessaires pour comprendre les technologies plus avancées que nous détaillerons dans les chapitres suivants, en particulier les solutions de réseau telles que MPLS.

Chapitre III . Principe Réseau IP- MPLS

Chapitre III . Principe Réseau IP-MPLS

III . 1 Introduction

Avec l'augmentation des besoins en performance, flexibilité et sécurité des réseaux IP, la technologie MPLS (Multiprotocol Label Switching) se présente comme une option efficace et populaire. Ce chapitre vise à explorer les bases de MPLS, en commençant par son fonctionnement et ses principaux composants, tels que les labels et les protocoles de distribution. Nous aborderons également son intégration avec les VPN. Les principales applications de MPLS seront détaillées, notamment l'ingénierie du trafic et la qualité de service (QoS) et les types de tunnels. Enfin, nous examinerons des mécanismes importants tels que les tables VRF et le protocole MP-BGP, afin de mieux comprendre comment MPLS permet un acheminement efficace et sécurisé des paquets IP au sein des réseaux multi-sites.

III . 2 Technologie MPLS

MPLS, ou Multiprotocol Label Switching, est une méthode utilisée dans les réseaux pour accélérer le transfert des données. Contrairement au routage IP traditionnel, où chaque appareil doit examiner l'adresse IP de chaque paquet, MPLS attribue un label à chaque paquet quand il entre dans le réseau. Cela aide les routeurs à faire des choix rapides sans avoir à vérifier les tables de routage. On considère souvent MPLS comme une technologie de couche 2,5, car elle se situe entre la couche de liaison (couche 2) et la couche réseau (couche 3), ce qui lui permet de combiner la rapidité du niveau 2 avec les fonctions du niveau 3 (figure III-1).

Cette technologie est surtout utilisée par les grandes entreprises qui ont plusieurs sites éloignés. Elle offre une connexion fiable entre les bureaux et les centres de données. On rajoute à cela, sa capacité de donner la priorité à certains types de trafic, comme les appels ou les vidéos, assurant une bonne qualité de service tout en diminuant la latence et la perte de paquets.

Néanmoins, le problème majeur des MPLS est son coût. Difficile à mettre en place et n'est pas aussi flexible face à la montée des solutions cloud. Dans les systèmes classiques où

tout passe par un point central, les données doivent souvent passer par le siège avant d'atteindre le cloud, ce qui peut ralentir la transmission et demande plus de bande passante.

Quant à la sécurité, MPLS n'offre pas de protection des données de manière intégrée. Il est donc important de le compléter avec des outils comme le chiffrement, les pare-feu ou des systèmes pour détecter et prévenir les intrusions afin de garantir la confidentialité et la fiabilité des données qui circulent sur le réseau [14].

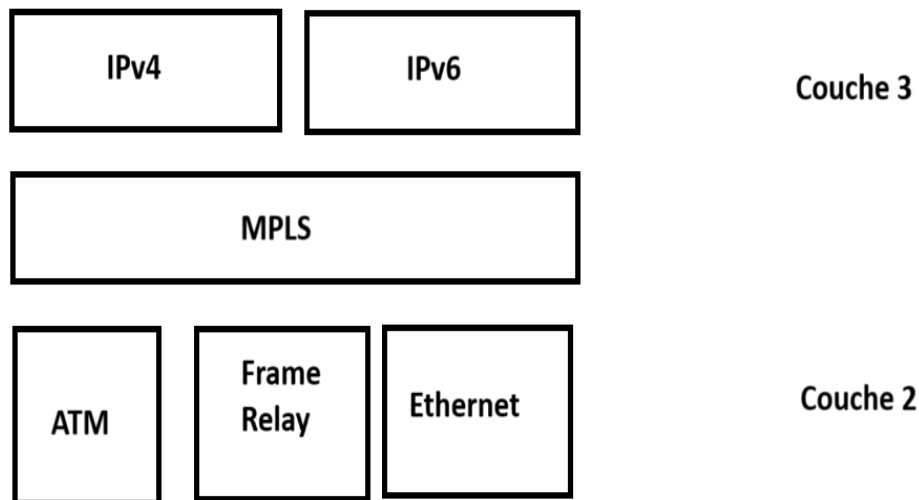


Figure III-1 technologie de couche 2.5

III . 3 Eléments et Terminologie du MPLS

III . 3 .1 Elément du MPLS

Dans les réseaux MPLS, les données passent par des dispositifs comme des routeurs ou des commutateurs qui gèrent et remplacent les étiquettes. Mais ces appareils ne peuvent pas analyser les en-têtes des paquets au niveau réseau. L'acheminement se fait via des étiquettes au lieu des adresses IP classiques. Les équipements principaux utilisés dans le cadre de l'architecture MPLS sont : LSR et LER.

LSR (Label Switching Router) : Il s'agit d'un routeur ou d'un commutateur qui fonctionne à l'intérieur du domaine MPLS. Ses rôles sont : échanger des informations réseau, échanger les étiquettes et acheminer les paquets en fonction de ces étiquettes.

LER (Label Edge Router) : Ce type de routeur agit comme une interface entre un domaine MPLS et l'extérieur. Il a deux fonctions principales : il ajoute des étiquettes aux paquets entrants (Ingress LER) et les supprime des paquets sortants (Egress LER) [15].

III . 3 .2 Terminologie utilisée

Dans un réseau MPLS, il existe quelques concepts de base qui aident à acheminer efficacement les données à l'aide de balises. Ci-dessous, un aperçu simple de ces termes :

LSP (Label Switch Path) : Il s'agit du chemin unidirectionnel emprunté par les paquets à travers les LSR, sans modifier leur contenu. Cette route permet un transport rapide et organisé.

Label : Il s'agit du numéro qu'un routeur LER place sur un paquet IP lorsqu'il entre dans un réseau MPLS. Ceci est utilisé pour acheminer les paquets via les LSR, sans avoir à analyser l'intégralité du contenu des en-têtes IP.

LIB (Label Information Base) : Un enregistrement contenant des informations sur les associations d'étiquettes. Aide à choisir la balise appropriée pour accéder au routeur dans un réseau MPLS.

LFIB (Label Forwarding Information Base) : Il associe chaque étiquette entrante à une action à effectuer (ajouter, supprimer ou remplacer l'étiquette) et à une sortie à effectuer.

FIB (Forwarding Information Base) : Il associe chaque destination réseau à une balise et à une interface de sortie, simplifiant ainsi le processus initial d'attribution de balise.

FEC (Forwarding Equivalence Class) : Un groupe de paquets qui ont des caractéristiques de routage similaires (telles que l'adresse de destination ou la classe de service). MPLS attribue la même étiquette à tous les paquets du même FEC [15].

III . 4 Définition labels

Un label est un petit identifiant numérique placé entre la couche 2 (liaison de données - MAC) et la couche 3 (réseau - IP) d'un paquet MPLS. Il permet d'acheminer les paquets sans avoir à lire leur contenu IP. Avec ce système, les routeurs qui changent d'indicateur n'ont plus besoin d'effectuer la recherche d'itinéraire habituelle. Ils consultent simplement le tableau d'échange d'étiquettes.

Ces labels peuvent être attribuées de manière dynamique à l'aide de protocoles tels que LDP (Label Distribution Protocol) ou de manière statique également dans certains cas. Chaque paquet suit un chemin déjà établi, appelé chemin de commutation d'étiquettes (LSP).

Cela contribue à rendre les réseaux MPLS plus efficaces, plus résilients et à offrir une meilleure qualité de service [15].

III . 4 .1 Fonctionnement des labels dans MPLS

III .4 .1 .a Label Switching (commutation par label)

Dans un réseau MPLS, chaque paquet se voit attribuer un petit label numérique, qui est placée entre la couche de liaison de données et la couche réseau. Cet indicateur aide les routeurs MPLS à choisir rapidement le chemin à suivre à l'aide d'une table basée sur ces indicateurs, sans avoir à vérifier les adresses IP. Contrairement au routage IP traditionnel, cette méthode accélère le traitement et facilite la transmission des paquets au sein du réseau.

III .4 .1 .b Label forwarding (transmission par label)

Dans un réseau MPLS, le transfert des paquets se fait grâce à un système de labels. Chaque routeur dans le réseau, qu'on appelle un LSR, utilise le label d'un paquet pour décider où l'envoyer. Quand un paquet arrive dans le réseau, le routeur d'entrée lui donne un label selon sa destination. Ensuite, à chaque saut entre les routeurs, le LSR lit le label, regarde dans sa table pour savoir quoi faire, change le label pour un nouveau, et envoie le paquet au prochain routeur sur le chemin. Ce système permet de transférer les paquets rapidement, car il ne faut pas passer du temps à vérifier les en-têtes IP à chaque étape. Le paquet suit un chemin prédéfini jusqu'au routeur de sortie, où le label est enlevé avant qu'il ne soit livré. Ce fonctionnement qui fait que MPLS est différent des réseaux IP classiques [16].

III .4 .1 .c Label distribution (Distribution des labels)

Dans les réseaux MPLS, la façon dont les labels sont distribués est super importante pour envoyer les paquets rapidement. Au lieu de se baser sur les adresses IP comme le fait le routage classique, MPLS utilise des petits labels pour suivre des chemins spécifiques, appelés chemins commutés par label. Ces labels sont échangés et donnés entre les routeurs MPLS avec des protocoles spéciaux, ce qui permet de faire une commutation efficace à travers tout le réseau.

III . 4 .2 Les protocoles de distribution label

LDP (Label Distribution Protocol) : est le protocole le plus courant pour l'attribution de labels de façon dynamique. Il aide les routeurs MPLS à partager des infos sur les préfixes IP et les labels qui leur sont liés, en utilisant des routes faites par des protocoles comme OSPF ou IS-IS. LDP est une façon simple de créer des chemins LSP, sans avoir à gérer des configurations compliquées pour le trafic [17].

RSVP-TE (Resource Reservation Protocol – Traffic Engineering) : est une version du protocole RSVP qui s'occupe de la gestion du trafic. Il sert à établir des chemins clairs en réservant les ressources dont on a besoin, comme la bande passante et la latence, pour assurer la qualité de service. Ce protocole est bien adapté aux réseaux qui ont besoin d'un contrôle précis du trafic et d'une bonne gestion des ressources [18]

MP-BGP (Multiprotocol Border Gateway Protocol) : une version du protocole BGP qui gère plusieurs types d'adresses, comme IPv4 ou VPNv4. Dans les réseaux MPLS VPN, il sert à transmettre les routes des clients entre les routeurs PE avec leurs labels. MP-BGP permet de garder les clients bien séparés tout en rendant le réseau plus facile à développer [19].

III . 5 Format de l'en-tête MPLS

L'en-tête MPLS, ou Multiprotocol Label Switching, fait 32 bits et se place entre l'en-tête de la couche liaison, comme Ethernet, et celui de la couche réseau, comme IP. Grâce à cet en-tête, les routeurs peuvent choisir comment acheminer les données en se basant sur des labels au lieu de devoir utiliser des adresses IP complètes [18].

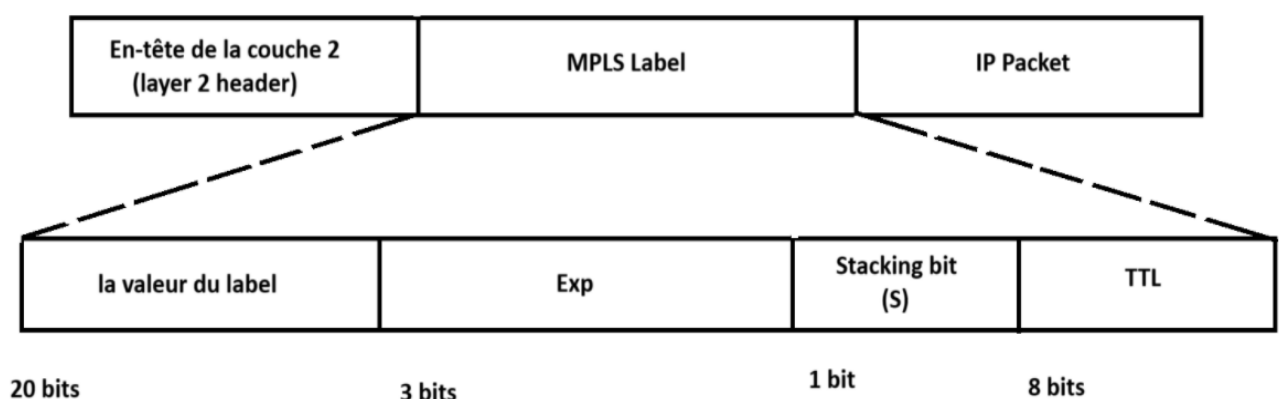


Figure Erreur ! Utilisez l'onglet Accueil pour appliquer Heading 1 au texte que vous souhaitez faire apparaître ici..2 : détails d'un label MPLS

Comme le montre la figure (III.2), on distingue quatre principaux champs :

Label (20bits) : C'est l'identifiant numérique principal du paquet. Il sert à le relier à un groupe de transfert, ce qui aide les routeurs MPLS à faire leur travail sans avoir besoin de regarder les détails IP.

Exp / Traffic Class (3bits) : Ce champ, qu'on appelait autrefois champ expérimental (Exp), est maintenant utilisé pour gérer la qualité de service. Il sert à codifier la priorité du trafic ou la classe de service, ce qui aide à différencier les flux dans le réseau.

Stacking bit / S (1bit) : Ce bit indique si le label en cours est le dernier de la pile de labels. Une valeur de 1 signifie que c'est le dernier label à lire ; une valeur de 0 indique qu'il y a d'autres labels empilés. Cela permet la hiérarchisation du routage (par exemple, pour VPN ou trafic multicast).

TTL (Time To Live) (8bits) : Fonctionne comme dans les paquets IP. Il limite la durée de vie d'un paquet dans le réseau MPLS en évitant les boucles de routage. À chaque saut, sa valeur est décrémentée.

III . 6 Application de la technologie MPLS

Les applications principales de la technologie MPLS incluent la création de réseaux privés virtuels (VPN), la garantie d'une qualité de service (QoS) variée et la gestion du trafic pour mieux utiliser les ressources du réseau. Ces caractéristiques font de MPLS un choix populaire dans les réseaux compliqués et exigeants.

III . 6 .1 L'ingénierie du trafic (TE)

Le trafic engineering (TE) est une façon efficace d'utiliser la technologie MPLS. L'idée est de mieux gérer les ressources du réseau en contrôlant les chemins que prennent les données, ce qui permet d'éviter les congestions et d'améliorer le routage.

Contrairement au routage IP classique, qui peut parfois surcharger certains points et laisser d'autres peu utilisés, MPLS-TE permet de choisir des chemins spécifiques tout en tenant compte de choses comme la bande passante, le temps de réponse ou la perte de paquets. Il s'appuie sur un algorithme qui choisit les routes non seulement en fonction de la distance, mais aussi de la qualité du service. Des protocoles comme RSVP-TE sont importants ici, car ils aident à réserver les ressources le long du chemin.

De plus, MPLS-TE offre des moyens de changer de route facilement, ce qui permet de passer à un nouveau chemin sans perturber le trafic, surtout en cas de problèmes ou quand un

meilleur chemin est disponible. Cela garantit une bonne disponibilité, un équilibre de charge et une réponse rapide aux besoins des applications qui exigent un temps de réponse rapide.

III .6 .1 .a Types de tunnels dans MPLS-TE

Dans la technologie MPLS Traffic Engineering, on trouve deux types de tunnels : les tunnels dynamiques et les tunnels explicites.

Tunnel dynamique : Dans le mode tunnel dynamique, chaque routeur décide de donner un label à une classe d'équivalence de routage. Ce label est ensuite partagé avec les routeurs voisins. Ça fonctionne un peu comme le routage IP traditionnel, où chaque nœud choisit lui-même quel chemin prendre pour un paquet, souvent en se basant sur le chemin le plus court selon les protocoles comme OSPF ou IS-IS.

Tunnel explicite : Un tunnel explicite, c'est quand le routeur à l'extrémité (Provider Edge, PE) indique une liste précise de routeurs par lesquels les données doivent passer. Même si ce chemin n'est pas toujours le plus court, il permet de réserver des ressources et d'assurer un certain niveau de qualité de service (QoS) durant le trajet. Ce genre de tunnel est souvent choisi quand il y a des contraintes particulières à respecter, comme la bande passante ou la latence [20].

III . 6 .2 Qualité de service (QoS)

Dans un réseau MPLS, la qualité de service (QoS) est super importante pour assurer que tout fonctionne bien, surtout pour les applications comme la voix et la vidéo qui doivent avoir la priorité. On peut gérer la QoS dans MPLS de deux manières principales : IntServ et DiffServ.

L'approche Integrated Services (IntServ) : se base sur la réservation des ressources à l'aide du protocole RSVP. Chaque flux de données est traité à part, avec sa propre classe de service. Même si cette méthode permet une gestion précise des données, elle ne fonctionne pas bien sur de grands réseaux car elle est compliquée, consomme beaucoup de ressources et ne s'adapte pas facilement à l'augmentation du trafic.

D'un autre côté, **l'approche Différentiâtes Services (DiffServ)** classe les flux de manière générale en marquant les paquets, comme avec les bits EXP dans l'en-tête MPLS. C'est une méthode plus simple et plus adaptable, car elle gère la qualité de service par groupes de trafic. Mais elle est moins précise pour les flux qui ont besoin de plus d'attention. En ajoutant la QoS

à MPLS, on peut diriger le trafic en fonction de ses besoins spécifiques. Ça aide à mieux répartir les ressources, que ce soit pour un seul opérateur ou plusieurs [20].

III . 7 Réseau Privés Virtuels dans MPLS

Dans une situation où les entreprises ont plusieurs sites éloignés, il est essentiel de les relier de manière sécurisée et efficace. Auparavant, cela impliquait des lignes coûteuses, mais avec les réseaux MPLS, on a maintenant une option plus économique et flexible : les VPN MPLS.

Un VPN, c'est un réseau privé virtuel qui connecte différents sites à travers l'infrastructure d'un fournisseur de services, tout en gardant une séparation entre les clients. Chaque site est lié à un ou plusieurs VPN, et pour que deux sites communiquent, ils doivent être dans le même VPN, ce qui assure la sécurité du trafic. On trouve deux types :

- Intranet VPN : c'est quand tous les sites font partie de la même entreprise.
- Extranet VPN : c'est quand plusieurs entreprises utilisent le même VPN pour travailler ensemble.

MPLS permet aux fournisseurs de proposer des services VPN de niveau 3 (L3VPN) ou de niveau 2 (L2VPN) en mettant en place des règles pour le routage, la sécurité et parfois la qualité de service. Cette structure aide les entreprises à s'agrandir facilement, à gérer leurs réseaux plus simplement et à garantir de bonnes performances [21].

VPN couche 2 (L2VPN) : est une solution qui permet de connecter directement les sites des clients en transférant des paquets de niveau 2, comme Ethernet ou Frame Relay, sur un réseau MPLS. Contrairement au L3VPN, le réseau du fournisseur ne se mêle pas du routage IP du client. Le L2VPN fait office de pont virtuel, ce qui permet aux clients de gérer leurs propres protocoles de routage. C'est un bon choix pour les entreprises qui veulent garder le contrôle sur leur réseau ou pour relier des sites utilisant des protocoles non-IP.

VPN couche 3 (L3VPN) : est le type de VPN qu'on voit le plus souvent dans les réseaux MPLS. Il permet à plusieurs sites de clients de partager des données en utilisant l'infrastructure IP de leur fournisseur de services. Dans ce système, les routeurs PE (Provider Edge) gèrent le routage IP du client en maintenant une table de routage séparée pour chaque client grâce aux VRF (Virtual Routing and Forwarding). On utilise souvent des protocoles de

routing comme BGP pour échanger des infos de routing entre les différents sites. Ce VPN est vraiment adapté pour les services entre différents sites à un niveau IP [21].

III . 7 .1 Virtuel routing et forwarding VRF

La technologie Virtual Routing and Forwarding (VRF) permet à un routeur de gérer plusieurs tables de routing indépendantes. Chaque instance VRF fonctionne comme un routeur à part entière, ce qui aide à garder le trafic réseau séparé pour différents clients ou services tout en utilisant la même infrastructure physique. Avec cette séparation, on peut utiliser les mêmes plages d'adresses IP dans divers réseaux privés sans craindre des conflits.

Cette séparation est très importante pour les environnements multi-clients, comme chez les opérateurs ou fournisseurs de services, où chaque client a besoin d'un espace de routing particulier. Chaque interface du routeur d'accès (PE) est reliée à une instance VRF spécifique, et le routing des paquets se fait selon la table de routing de cette VRF, pas selon celle du routeur en général [22].

III .7 .1 .a VRF Forwarding Table

Dans un réseau MPLS qui utilise la technologie VRF, chaque VRF a sa propre table de routing. Cette table garde les routes spécifiques à un client, ce qui permet de diriger le trafic sans interférer avec les autres tables de routing sur le même routeur.

Quand un routeur de bord reçoit un paquet IP d'un client, il ne regarde pas la table de routing globale. Il se réfère plutôt à la table qui correspond au VRF de l'interface qui a reçu le paquet. Cela permet de garder le trafic de chaque client bien séparé, même si les adresses IP se chevauchent. Chaque VRF a un nom défini sur le routeur PE, mais ce nom ne veut rien dire pour les autres équipements du réseau. Chaque interface client sur le routeur PE est associée à un VRF particulier, ce qui permet d'appliquer des règles de routing différentes pour chaque client.

Ce système est essentiel pour gérer le trafic en toute sécurité et efficacement dans un VPN MPLS. Il permet de séparer les réseaux et de personnaliser la connectivité et les règles de routing en fonction des besoins de chaque client [22].

III . 7 .1 .b Propagation des informations de routage VPN

Dans une infrastructure MPLS VPN, on trouve trois types de routeurs, chacun ayant un rôle unique pour gérer le trafic et séparer les réseaux des clients.

Routeur P (Provider) : Ce routeur, placé au centre du réseau du fournisseur, s'occupe seulement du transfert des paquets avec des étiquettes MPLS entre les routeurs PE. Il ne connaît pas les routes des clients et ne s'implique pas dans le partage des infos de routage VPN. Sa tâche se concentre donc sur le basculement rapide grâce aux labels MPLS, ce qui assure une bonne performance au sein du réseau du fournisseur.

Routeur PE (Provider Edge) : Situé à la limite du réseau du fournisseur, ce routeur se connecte directement aux routeurs CE des clients. Il a la tâche de gérer des instances VRF (Virtual Routing and Forwarding) séparées pour chaque client, ce qui permet de garder les réseaux bien isolés. Le PE ajoute ou enlève les labels MPLS en fonction du trafic, et il partage les routes VPN en utilisant des protocoles comme MP-BGP.

Routeur CE (Customer Edge) : Le CE est installé chez le client et il connecte son réseau privé au réseau MPLS du fournisseur. Il ne gère pas directement MPLS, mais il utilise des protocoles standards comme BGP, OSPF, RIP ou des routes statiques pour échanger les informations de routage avec le PE. Cette séparation des rôles aide à avoir une architecture qui peut évoluer, est sécurisée et peut offrir des services VPN à plusieurs clients [23].

III . 7 .2 Multiprotocol BGP (MP-BGP)

Le protocole MP-BGP (Multiprotocol Border Gateway Protocol) joue un rôle clé dans la gestion des infos de routage dans les réseaux MPLS VPN. À la différence du BGP classique, MP-BGP prend en charge plusieurs types d'adresses, y compris les adresses VPNv4 qu'on voit dans les VPN de niveau 3.

Dans ce contexte, les routeurs PE (Provider Edge) utilisent MP-BGP pour partager des routes VPN. Quand un routeur CE (Customer Edge) annonce une route IPv4 à son PE, celui-ci lui donne un identifiant unique qu'on appelle Route Distingasse (RD). Ce RD, qui fait 64 bits, ajouté à l'adresse IPv4 de 32 bits, crée un préfixe VPNv4 de 96 bits. Cela assure que les routes restent uniques, même si les clients ont des adresses qui se chevauchent.

Les préfixes VPNv4 créés sont ensuite partagés via MP-BGP avec les autres routeurs PE dans le réseau. Une fois qu'ils les ont reçus, ces routeurs PE peuvent enlever le RD et remettre l'adresse IPv4 d'origine dans la bonne table de routage VRF. Mais, l'ajout du RD ne

dit pas quelles routes doivent être acceptées ou rejetées par les VRF. C'est là que le concept de Route Target (RT) entre en jeu.

Le RT est une sorte de communauté BGP qui agit comme un filtre pour les routes VPN. Chaque VRF a ses propres RTs pour décider quelles routes VPNv4 peuvent être acceptées ou annoncées. Cela permet de créer des règles de routage pratiques entre les différents endroits d'un VPN [24].

En conclusion, MP-BGP, avec l'usage de RD et RT, permet de gérer efficacement les routes des clients dans les MPLS VPN, tout en gardant une séparation entre les clients.

III . 7 .3 Le Transfert des Paquets IP (IP Packet Forwarding)

Dans un réseau MPLS VPN, le transfert des paquets IP fonctionne grâce à un système d'encapsulation à deux niveaux, qu'on appelle empilement d'étiquettes. Quand un paquet IP part d'un routeur CE (Customer Edge) vers un autre site, le routeur PE (Provider Edge) à l'origine ajoute deux étiquettes MPLS. La première étiquette aide à envoyer le paquet à travers le backbone MPLS vers le routeur PE de destination, tandis que la deuxième étiquette indique quelle interface utiliser sur le PE de destination pour le routeur CE qui reçoit le paquet. Cette deuxième étiquette est apprise via des mises à jour du protocole MP-BGP.

Le transfert de données se fait en regardant les tables de routage et les tables CEF (Cisco Express Forwarding) des routeurs. Cela aide à choisir la meilleure interface de sortie. En gros, un routeur reçoit un paquet sur une interface d'entrée et le renvoie sur une interface de sortie selon ce qui est dans sa table de routage. L'efficacité de tout ça dépend de la capacité du routeur à choisir le meilleur chemin vers le réseau de destination [25].

III . 8 Conclusion

Les reseaux IP/MPLS sont omnipresent dans les infrastructures des operateurs. Leur valeur reside dans leur capacite a surmonter les limitations des protocoles de routages IP classiques en introduisant la communication de labels. De plus, ces reseaux offrent une solutions strayantes en integrant facilement de nouvelles technologies basees sur la virtualisation des services reseau afin de fournir des offres virtuelles aux clients et l'échange de routes via MP-BGP. En comprenant ses principes et ses fonctionnalités, nous pouvons, maintenant, créer des architectures réseau solides, flexibles, et qui répondent aux besoins

Chapitre III. Principes reseau IP/MPLS

actuels des entreprises, ce qui constituera le theme central de notre travail detaille dans le prochain chapitre.

Chapitre IV. Simulation d'un réseau IP-MPLS sous GNS3

Chapitre IV . Simulation d'un Réseau IP-MPLS sous GNS3

IV . 1 Introduction

Ce chapitre est consacré à la mise en œuvre d'une simulation pratique d'un réseau IP-MPLS, dans le but d'approfondir la compréhension des notions théoriques étudiées précédemment. A travers la conception et la configuration d'une topologie représentative, nous cherchons à illustrer le fonctionnement réel du MPLS, notamment le rôle des protocoles de routage comme OSPF, ainsi que le processus d'attribution et d'échange des labels via le protocole LDP. Cette approche permet de visualiser concrètement les mécanismes de commutation par labels au sein d'un réseau IP moderne.

IV . 2 Présentation de l'outil de simulation GNS3

Dans le cadre de cette étude, l'outil GNS3 (Graphical Network Simulator 3) a été choisi pour la modélisation et la simulation de l'architecture réseau. Il s'agit d'un simulateur puissant et largement utilisé dans le domaine des réseaux, qui permet de créer des topologie complexes en intégrant de véritables images de systèmes d'exploitation réseau (comme IOS de Cisco). Grâce à son interface graphique intuitive, GNS3 facilite l'interconnexion, la configuration et la supervision des équipements virtuels, tout en offrant un comportement très proche d'un environnement réel.

Ce simulateur est particulièrement adapté aux projets pédagogiques et professionnels, car il autorise l'émulation de divers protocoles et technologies de routage, notamment MPLS, ce qui en fait un outil de choix pour tester et valider des scénarios réseaux avant leur déploiement concert.

IV . 3 Installation du logiciel GNS3

Le logiciel GNS3 a été adopté pour la simulation du réseau IP-MPLS en raison de sa capacité à exécuter des images réelles d'équipements réseau. Nous l'avons installé sur un système Windows 10, avec l'activation de GNS3 VM pour de meilleures performances. Après l'installation, les images IOS des routeurs ont été ajoutées via les paramètres du programme,

ce qui nous a permis de créer un environnement de simulation réaliste pour tester différents scénarios réseau. Les détails de l'installation sont présentés en Annexe 1.

IV . 4 Topologie physique de notre plateforme

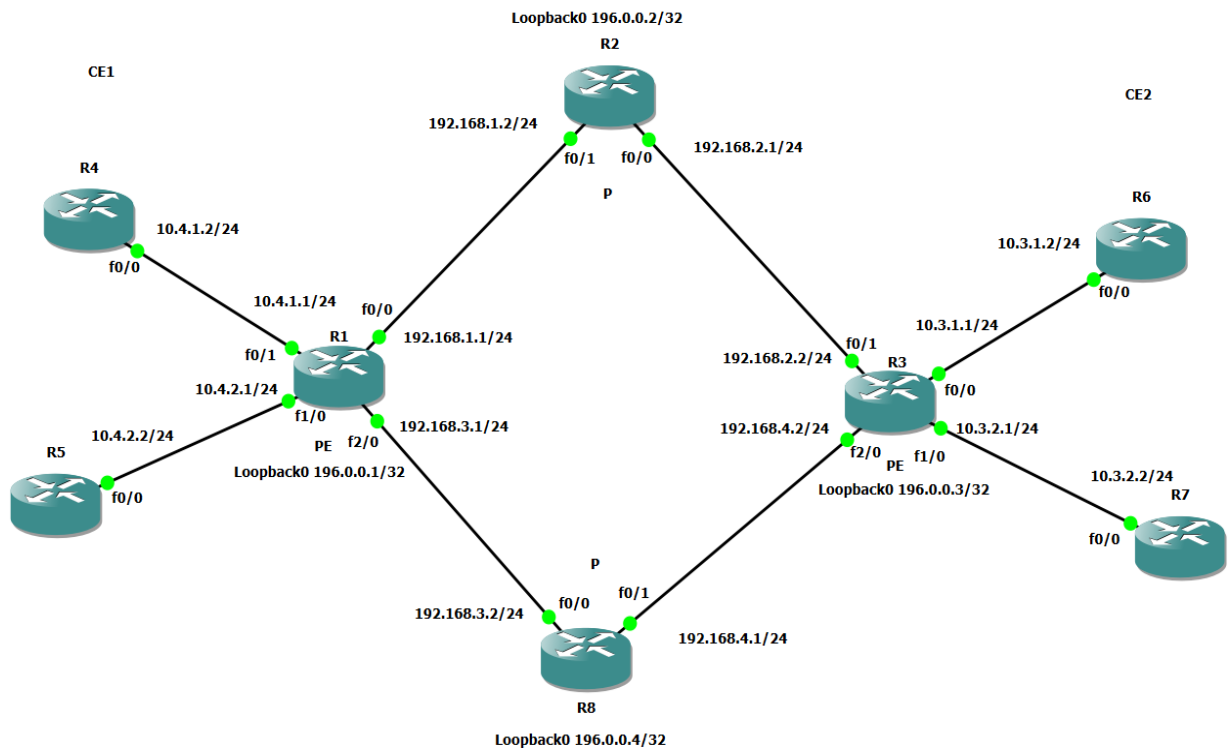


Figure IV-1 : La topologie simulée

Pour la mise en place de la plateforme, nous avons choisi les séries de routeurs C3725 en raison des avantages suivants :

- Support MPLS : Ils prennent en charge MPLS, ce qui permet un routage plus rapide et plus efficace via la commutation d'étiquettes.
- Accès VPN : Ils assurent des connexions VPN sécurisées, essentielles pour protéger les données sur les réseaux publics.
- Capacités multiservices : Ils supportent les services voix, fax et données, offrant une solution de communication complète.
- Commutation d'étiquettes via BGP : Ces routeurs intègrent la commutation d'étiquettes via BGP, améliorant ainsi la flexibilité du routage.

Le tableau (IV-1) présenté ci-dessous, décrit la configuration réseau des routeurs, en mettant l'accent sur leurs interfaces, adresses IP, masques de sous-réseau ainsi que les adresses Loopback.

Tableau Erreur ! Utilisez l'onglet Accueil pour appliquer Heading 1 au texte que vous souhaitez faire apparaître ici.-1

Routeur	Interface	Adresse et masque	Loopback
R1	FastEthernet0/0 FastEthernet0/1 FastEthernet1/0 FastEthernet2/0	192.168.1.1/24 10.4.1.1/24 10.4.2.1/24 192.168.3.1/24	196.0.0.1
R2	FastEthernet0/1 FastEthernet0/0	192.168.1.2/24 192.168.2.1/24	196.0.0.2
R3	FastEthernet0/1 FastEthernet0/0 FastEthernet1/0 FastEthernet2/0	192.168.2.2/24 10.3.1.1/24 10.3.2.1/24 192.168.4.2/24	196.0.0.3
R4	FastEthernet0/0	10.4.1.2/24	
R5	FastEthernet0/0	10.4.2.2/24	
R6	FastEthernet0/0	10.3.1.2/24	
R7	FastEthernet0/0	10.3.2.2/24	
R8	FastEthernet0/0 FastEthernet0/1	192.168.3.2/24 192.168.4.1/24	196.0.0.4

IV . 5 Configuration des interfaces des routeurs

Une fois la topologie, (figure IV-1), physique mise en place dans l'espace de travail, les routeurs sont démarrés à l'aide du bouton dédié. Lorsque les liaisons deviennent actives (passant du rouge au vert), nous passons à la configuration des adresses IP sur les interfaces de chaque routeur, ainsi qu'à l'attribution des adresses loopback via la console de configuration. Étant donné que la configuration est similaire pour l'ensemble des routeurs, seuls les exemples des routeurs R1 et R3 seront présentés.

Configuration de routeur R1

```
R1>enable
R1#configuration terminal
R1(config)#interface fastEthernet0/0
R1(config-if)#ip address 192.168.1.1 255.255.255.0
R1(config-if)#no shutdown
```

```
R1(config-if)#exit
R1(config)#interface fastEthernet0/1
R1(config-if)#ip address 10.4.1.1 255.255.255.0
R1(config-if)#no shutdown
R1(config-if)#exit
R1(config)#interface fastEthernet1/0
R1(config-if)#ip address 10.4.2.1 255.255.255.0
R1(config-if)#no shutdown
R1(config-if)#exit
R1(config)#interface fastEthernet2/0
R1(config-if)#ip address 192.168.3.1 255.255.255.0
R1(config-if)#no shutdown
R1(config-if)#exit
R1(config)#interface loopback0
R1(config-if)#ip address 196.0.0.1 255.255.255.255
R1(config-if)#end
```

Discussion

- R1>enable

Cette commande permet d'entrer dans le mode privilégié (EXEC privilégié) afin d'avoir accès à plus de commandes, y compris celles de configuration.

- R1#configuration terminal

Elle permet de passer en mode de configuration globale, où l'on peut modifier les paramètres du routeur.

- R1(config)#interface fastEthernet1/0

Cette commande sélectionne l'interface FastEthernet1/0 pour la configurer. Les interfaces sont les points d'interaction entre le routeur et le réseau.

- R1(config-if)#ip address 192.168.1.1 255.255.255.0

Elle attribue l'adresse IP 192.168.1.1 à l'interface FastEthernet1/0. Le masque de sous-réseau 255.255.255.0 (/24) indique un petit sous-réseau avec seulement deux adresses IP utilisables (idéal pour les liaisons point-à-point).

- R1(config-if)#no shutdown

Cette commande active l'interface (elle est désactivée par défaut).

- R1(config-if)#exit

Elle permet de sortir du mode de configuration de l'interface et de revenir au mode de configuration globale.

- R1(config)#interface loopback0

Cette commande sélectionne l'interface virtuelle Loopback0. Les interfaces loopback sont utilisées à des fins de gestion et de tests.

- R1(config-if)#ip address 196.0.0.1 255.255.255.255

Elle attribue l'adresse IP 196.0.0.1 à l'interface loopback. Le masque 255.255.255.255(/32) représente une seule adresse IP, ce qui est typique pour une loopback.

- R1(config-if)#end

Cette commande permet de quitter tous les modes de configuration et de revenir au mode EXEC privilégié. Dans ce qui suit, nous passons à la configuration de R3, R2 et R8 respectivement.

Configuration de routeur R3

```
R3>enable
R3#configuration terminal
R3(config)#interface fastEthernet0/1
R3(config-if)#ip address 192.168.2.2 255.255.255.0
R3(config-if)#no shutdown
R3(config-if)#exit
R3(config)#interface fastEthernet0/0
R3(config-if)#ip address 10.3.1.1 255.255.255.0
R3(config-if)#no shutdown
R3(config-if)#exit
R3(config)#interface fastEthernet1/0
R3(config-if)#ip address 10.3.2.1 255.255.255.0
R3(config-if)#no shutdown
R3(config-if)#exit
R3(config)#interface fastEthernet2/0
R3(config-if)#ip address 192.168.4.2 255.255.255.0
R3(config-if)#no shutdown
R3(config-if)#exit
R3(config)#interface loopback0
```

```
R3(config-if)#ip address 196.0.0.3 255.255.255.255
R3(config-if)#end
```

Configuration de routeur R2

```
R2>enable
R2#configuration terminal
R2(config)#interface fastEthernet0/0
R2(config-if)#ip address 192.168.1.2 255.255.255.0
R2(config-if)#no shutdown
R2(config-if)#exit
R2(config)#interface fastEthernet0/1
R2(config-if)#ip address 192.168.2.1 255.255.255.0
R2(config-if)#no shutdown
R2(config-if)#exit
R2(config)#interface loopback0
R2(config-if)#ip address 196.0.0.2 255.255.255.255
R2(config-if)#end
```

Configuration de routeur R8

```
R8>enable
R8#configuration terminal
R8(config)#interface fastEthernet0/0
R8(config-if)#ip address 192.168.3.2 255.255.255.0
R8(config-if)#no shutdown
R8(config-if)#exit
R8(config)#interface fastEthernet0/1
R8(config-if)#ip address 192.168.4.1 255.255.255.0
R8(config-if)#no shutdown
R8(config-if)#exit
R8(config)#interface loopback0
R8(config-if)#ip address 196.0.0.4 255.255.255.255
R8(config-if)#end
```

Vérification des interfaces

Après la configuration des adresses IP sur les interfaces, il est essentiel d'effectuer une vérification pour s'assurer que celles-ci sont en état UP. Cela permet de confirmer que les interfaces sont bien actives et prêtes à échanger des données. Cette étape garantit également que les paramètres appliqués ont été pris en compte correctement par le routeur. Les figures (IV-2) et (IV-3) et (IV-4) et (IV-5), montre un exemple d'affichage des interfaces ainsi que les adresses IP des routeurs R1 et R3 et R2 et R8, respectivement.

R1#Show ip interface brief

Interface	IP-Address	OK?	Method	Status	Protocol
FastEthernet0/0	192.168.1.1	YES	NVRAM	up	up
Serial0/0	unassigned	YES	NVRAM	administratively down	down
FastEthernet0/1	10.4.1.1	YES	NVRAM	up	up
Serial0/1	unassigned	YES	NVRAM	administratively down	down
Serial0/2	unassigned	YES	NVRAM	administratively down	down
Serial0/3	unassigned	YES	NVRAM	administratively down	down
Serial0/4	unassigned	YES	NVRAM	administratively down	down
Serial0/5	unassigned	YES	NVRAM	administratively down	down
FastEthernet1/0	10.4.2.1	YES	NVRAM	up	up
FastEthernet2/0	192.168.3.1	YES	NVRAM	up	up
Loopback0	196.0.0.1	YES	NVRAM	up	up

Figure IV-2 : Affichage les interfaces et les adresses IP du R1.

Interface	IP-Address	OK?	Method	Status	Protocol
FastEthernet0/0	10.3.1.1	YES	manual	up	up
Serial0/0	unassigned	YES	unset	administratively down	down
FastEthernet0/1	192.168.2.2	YES	manual	up	up
Serial0/1	unassigned	YES	unset	administratively down	down
Serial0/2	unassigned	YES	unset	administratively down	down
Serial0/3	unassigned	YES	unset	administratively down	down
Serial0/4	unassigned	YES	unset	administratively down	down
Serial0/5	unassigned	YES	unset	administratively down	down
FastEthernet1/0	10.3.2.1	YES	manual	up	up
FastEthernet2/0	192.168.4.2	YES	manual	up	up
Loopback0	196.0.0.3	YES	manual	up	up

Figure IV-3 : Affichage les interfaces et les adresses IP du R3

```
R2#sh ip int bri
```

Interface	IP-Address	OK?	Method	Status	Prot
FastEthernet0/0	192.168.2.1	YES	NVRAM	up	up
Serial0/0	unassigned	YES	NVRAM	administratively down	down
FastEthernet0/1	192.168.1.2	YES	NVRAM	up	up
Serial0/1	unassigned	YES	NVRAM	administratively down	down
Serial0/2	unassigned	YES	NVRAM	administratively down	down
Serial0/3	unassigned	YES	NVRAM	administratively down	down
Serial0/4	unassigned	YES	NVRAM	administratively down	down
Serial0/5	unassigned	YES	NVRAM	administratively down	down
FastEthernet1/0	unassigned	YES	NVRAM	administratively down	down
FastEthernet2/0	unassigned	YES	NVRAM	administratively down	down
Loopback0	196.0.0.2	YES	NVRAM	up	up

Figure IV-4 : Affichage les interfaces et les adresses IP du R2

```
R8#sh ip int bri
```

Interface	IP-Address	OK?	Method	Status	Prot
FastEthernet0/0	192.168.3.2	YES	NVRAM	up	up
Serial0/0	unassigned	YES	NVRAM	administratively down	down
FastEthernet0/1	192.168.4.1	YES	NVRAM	up	up
Serial0/1	unassigned	YES	NVRAM	administratively down	down
Serial0/2	unassigned	YES	NVRAM	administratively down	down
Serial0/3	unassigned	YES	NVRAM	administratively down	down
Serial0/4	unassigned	YES	NVRAM	administratively down	down
Serial0/5	unassigned	YES	NVRAM	administratively down	down
FastEthernet1/0	unassigned	YES	NVRAM	administratively down	down
FastEthernet2/0	unassigned	YES	NVRAM	administratively down	down
Loopback0	196.0.0.4	YES	NVRAM	up	up

SFigure IV-5 : Affichage les interfaces et les adresses IP du R8

Ping : La commande Ping, illustrée dans la figure (IV-6), suivie de l'adresse IP de la machine cible, permet de vérifier la connectivité entre deux équipements. Dans notre cas, elle est utilisée pour s'assurer que les routeurs peuvent communiquer entre eux après la configuration des interfaces.

```
R3#PING 192.168.1.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 192.168.1.1, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 36/60/76 ms
R3#ping 192.168.3.2
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 192.168.3.2, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 20/29/32 ms
R3#ping 192.168.2.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 192.168.2.1, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 24/29/32 ms
R3#
```

Figure IV-6 : Test de connectivité entre R3 et différentes adresses (commande Ping).

IV . 6 Configuration de protocole de routage

La configuration du protocole OSPF sur un routeur se fait en deux étapes :

1. Activer le processus de routage OSPF (par exemple, le processus OSPF 100).
2. Spécifier les interfaces associées à ce processus OSPF.

Configuration de routeur R1

```
R1>en
R1#conf t
R1(config)#router ospf 100
R1(config)#network 192.168.1.1 0.0.0.255 area 0
R1(config)#network 10.4.1.1 0.0.0.255 area 0
R1(config)#network 10.4.2.1 0.0.0.255 area 0
R1(config)#network 192.168.3.1 0.0.0.255 area 0
R1(config)#network 196.0.0.1 0.0.0.0 area0
R1(config)#end
```

Discussion

La configuration du protocole OSPF (*Open Shortest Path First*) commence par l'activation d'un processus OSPF avec un identifiant (*ici, le numéro 100*), à l'aide de la commande suivante

- R1(config)#router ospf 100

Cette commande permet de démarrer le processus de routage OSPF, un protocole à état de lien utilisé pour calculer dynamiquement le chemin le plus court vers chaque destination dans le réseau. Ensuite, on associe les interfaces concernées à une zone OSPF.

Dans notre cas, la zone 0 (*appelée zone backbone*) :

- R1(config-router)#network 192.168.1.1 0.0.0.255 area 0

Cette ligne indique que le réseau 192.168.1.0/24 (où 0.0.0.255 est le masque générique – wildcard mask – équivalent à 255.255.255.0) doit être inclus dans la zone 0.

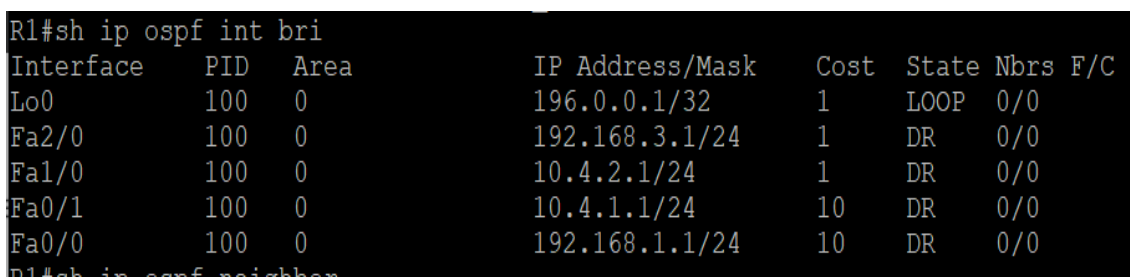
On ajoute également l'interface loopback à la même zone OSPF, avec la commande :

- R1(config-router) #network 196.0.0.1 0.0.0.0 area 0

Ici, le masque 0.0.0.0 permet d'indiquer une adresse IP précise, ce qui signifie que seule l'adresse 192.0.0.1 (*interface loopback*) sera prise en compte par OSPF dans la zone backbone.

Vérification de l'activation de l'OSPF

Pour s'assurer que le protocole OSPF est bien activé sur les interfaces du routeur, on utilise la commande suivante : « Show ip ospf interface brief ». Cette commande affiche un résumé concis de l'état du protocole OSPF sur chaque interface, figure (IV-7). Elle permet notamment de vérifier quelles interfaces participent au processus OSPF, leur état opérationnel, ainsi que leur association à une zone spécifique.



Interface	PID	Area	IP Address/Mask	Cost	State	Nbrs	F/C
Lo0	100	0	196.0.0.1/32	1	LOOP	0/0	
Fa2/0	100	0	192.168.3.1/24	1	DR	0/0	
Fa1/0	100	0	10.4.2.1/24	1	DR	0/0	
Fa0/1	100	0	10.4.1.1/24	10	DR	0/0	
Fa0/0	100	0	192.168.1.1/24	10	DR	0/0	

Figure IV-7 : Etat des interfaces OSPF sur R1

Sur le routeur R1, on peut constater que les trois interfaces sont associées au processus OSPF numéro 100 (colonne PID) et qu'elles appartiennent à la zone 0 (colonne Area). Cependant, il est également important de vérifier que les routeurs sont effectivement voisins. Pour cela, on utilise la commande : « show ip ospf neighbor », figure (IV-8).

```
R1#sh ip ospf neighbor

Neighbor ID      Pri   State           Dead Time   Address      Interface
196.0.0.2        1    FULL/BDR        00:00:34   192.168.1.2  FastEthernet0/0
```

Figure IV-8 :Affichage les voisins OSPF détectés par R1

Pour analyser les routes connues par le routeur R1, nous avons utilisé la commande suivante en mode privilégié,(figure IV-9) :

```
R1#sh ip route
Codes: 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

Gateway of last resort is not set

    196.0.0.0/32 is subnetted, 1 subnets
C       196.0.0.1 is directly connected, Loopback0
    10.0.0.0/24 is subnetted, 2 subnets
C       10.4.2.0 is directly connected, FastEthernet1/0
C       10.4.1.0 is directly connected, FastEthernet0/1
C     192.168.1.0/24 is directly connected, FastEthernet0/0
C     192.168.3.0/24 is directly connected, FastEthernet2/0
```

Figure IV-9 : Table de routage du R1

Pour R3, on aura :

Configuration de router R3

```
R3>en
R3#conf t
R3(config)#router ospf 100
R3(config)#network 192.168.2.2 0.0.0.255 area 0
R3(config)#network 10.3.1.1 0.0.0.255 area 0
R3(config)#network 10.3.2.1 0.0.0.255 area 0
R3(config)#network 192.168.4.2 0.0.0.255 area 0
R3(config)#network 196.0.0.3 0.0.0.0 area0
R3(config)#end
```

```
R3#sh ip ospf int bri
Interface      PID    Area      IP Address/Mask    Cost    State Nbrs F/C
Lo0            100    0          196.0.0.3/32       1       LOOP  0/0
Fa2/0          100    0          192.168.4.2/24     1       DR    0/0
Fa1/0          100    0          10.3.2.1/24        1       DR    0/0
Fa0/0          100    0          10.3.1.1/24        10      DR    0/0
Fa0/1          100    0          192.168.2.2/24     10      BDR   1/1
R3#sh ip ospf neighbor

Neighbor ID     Pri   State           Dead Time   Address      Interface
196.0.0.2       1     FULL/DR         00:00:38   192.168.2.1  FastEthernet0/1
```

Figure IV-10: Affichage l'état des interfaces OSPF sur R3

Pour vérifier l'établissement des relations de voisinage OSPF sur le routeur R3, nous avons utilisé la commande illustrée dans la figure (IV-11) suivante :

```
R3#sh ip ospf neighbor

Neighbor ID     Pri   State           Dead Time   Address      Interface
196.0.0.4       1     FULL/DR         00:00:37   192.168.4.1  FastEthernet2/0
196.0.0.2       1     FULL/BDR        00:00:36   192.168.2.1  FastEthernet0/1
```

Figure IV-11 : Affichage les voisins OSPF détectés par R3

Pour analyser les routes connues par le routeur R3, nous avons utilisé la commande suivante en mode privilégié, ainsi affichée dans la figure (IV-12) :

```
R3#sh ip route
Codes: 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

Gateway of last resort is not set

    196.0.0.0/32 is subnetted, 3 subnets
O       196.0.0.1 [110/21] via 192.168.2.1, 00:05:36, FastEthernet0/1
O       196.0.0.2 [110/11] via 192.168.2.1, 00:05:36, FastEthernet0/1
C       196.0.0.3 is directly connected, Loopback0
C       192.168.4.0/24 is directly connected, FastEthernet2/0
    10.0.0.0/24 is subnetted, 4 subnets
C       10.3.1.0 is directly connected, FastEthernet0/0
C       10.3.2.0 is directly connected, FastEthernet1/0
O       10.4.2.0 [110/21] via 192.168.2.1, 00:05:36, FastEthernet0/1
O       10.4.1.0 [110/30] via 192.168.2.1, 00:05:38, FastEthernet0/1
O       192.168.1.0/24 [110/20] via 192.168.2.1, 00:05:38, FastEthernet0/1
C       192.168.2.0/24 is directly connected, FastEthernet0/1
O       192.168.3.0/24 [110/21] via 192.168.2.1, 00:05:38, FastEthernet0/1
```

Figure IV-12: Affichage le table de routage du R3

Ping : Pour vérifier la connectivité entre le routeur R3 et l'interface Loopback de R1, nous avons utilisé la commande ping depuis R3 en ciblant l'adresse IP de la Loopback de R1, dans la figure (IV-13).

```
R3#ping 196.0.0.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 196.0.0.1, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 44/58/80 ms
```

Figure IV-13: Vérification de la connectivité entre R3 et l'interface Loopback de R1

Pour R2, nous avons la configuration suivante :

Configuration de routeur R2

```
R2>en
R2#conf t
R2(config)#router ospf 100
R2(config)#network 192.168.1.2 0.0.0.255 area 0
R2(config)#network 192.168.2.1 0.0.0.255 area 0
R2(config)#network 196.0.0.2 0.0.0.0 area0
R2(config)#end
```

Pour examiner l'état des interfaces configurées avec OSPF sur le routeur R2, nous avons utilisé la commande dans la figure (IV-14) suivante :

```
R2#sh ip ospf int bri
Interface    PID    Area      IP Address/Mask    Cost    State Nbrs F/C
Lo0          100    0         196.0.0.2/32       1       LOOP  0/0
Fa0/0        100    0         192.168.2.1/24     10      BDR   1/1
Fa0/1        100    0         192.168.1.2/24     10      DR    1/1
```

Figure IV-14 : Affichage l'état des interfaces OSPF sur R2

Pour vérifier l'établissement des relations de voisinage OSPF sur le routeur R2, nous avons utilisé la commande présentée dans la figure (IV-15) suivante :

```
R2#sh ip ospf neighbor
```

Neighbor ID	Pri	State	Dead Time	Address	Interface
196.0.0.3	1	FULL/DR	00:00:32	192.168.2.2	FastEthernet0/0
196.0.0.1	1	FULL/BDR	00:00:34	192.168.1.1	FastEthernet0/1

Figure IV-15 : Affichage les voisins OSPF détectés par R2

Pour analyser les routes connues par le routeur R2, nous avons utilisé la commande suivante en mode privilégié comme affiche la figure (IV-16) :

```
R2#sh ip route
Codes: 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

Gateway of last resort is not set

    196.0.0.0/32 is subnetted, 4 subnets
O       196.0.0.4 [110/12] via 192.168.2.2, 00:01:13, FastEthernet0/0
          [110/12] via 192.168.1.1, 00:01:13, FastEthernet0/1
O       196.0.0.1 [110/11] via 192.168.1.1, 00:01:23, FastEthernet0/1
C       196.0.0.2 is directly connected, Loopback0
O       196.0.0.3 [110/11] via 192.168.2.2, 00:01:13, FastEthernet0/0
O       192.168.4.0/24 [110/11] via 192.168.2.2, 00:01:13, FastEthernet0/0
C       192.168.1.0/24 is directly connected, FastEthernet0/1
C       192.168.2.0/24 is directly connected, FastEthernet0/0
O       192.168.3.0/24 [110/11] via 192.168.1.1, 00:01:15, FastEthernet0/1
```

Figure IV-16 : Table de routage du R2

Maintenant, même procédure pour R8 :

Configuration de routeur R8

```
R8>en
R8#conf t
R8(config)#router ospf 100
R8(config)#network 192.168.3.2 0.0.0.255 area 0
R8(config)#network 192.168.4.1 0.0.0.255 area 0
R8(config)#network 196.0.0.4 0.0.0.0 area 0
R8(config)#end
```

Pour examiner l'état des interfaces configurées avec OSPF sur le routeur R8, nous avons utilisé la commande affichée dans la figure (IV-17) :

```
R8#sh ip ospf int bri
Interface      PID    Area      IP Address/Mask    Cost    State  Nbrs  F/C
Lo0            100    0          196.0.0.4/32       1       LOOP   0/0
Fa0/1          100    0          192.168.4.1/24     10      DR     1/1
Fa0/0          100    0          192.168.3.2/24     10      DR     1/1
```

Figure IV-17 : Affichage l'état des interfaces OSPF sur R8

Pour vérifier l'établissement des relations de voisinage OSPF sur le routeur R8, nous avons utilisé la commande, figure (IV-18) :

```
R8#sh ip ospf neighbor
Neighbor ID     Pri   State           Dead Time   Address        Interface
196.0.0.3       1     FULL/BDR        00:00:36    192.168.4.2    FastEthernet0/1
196.0.0.1       1     FULL/BDR        00:00:39    192.168.3.1    FastEthernet0/0
```

Figure IV-18 : Affichage les voisins OSPF détectés par R8

Pour analyser les routes connues par le routeur R8, dans la figure (IV-19), nous avons utilisé la commande suivante en mode privilégié :

```
R8#sh ip route
Codes: 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

Gateway of last resort is not set

    196.0.0.0/32 is subnetted, 4 subnets
C       196.0.0.4 is directly connected, Loopback0
O       196.0.0.1 [110/11] via 192.168.3.1, 00:10:24, FastEthernet0/0
O       196.0.0.2 [110/21] via 192.168.4.2, 00:10:14, FastEthernet0/1
         [110/21] via 192.168.3.1, 00:10:24, FastEthernet0/0
O       196.0.0.3 [110/11] via 192.168.4.2, 00:10:24, FastEthernet0/1
C     192.168.4.0/24 is directly connected, FastEthernet0/1
O     192.168.1.0/24 [110/20] via 192.168.3.1, 00:10:24, FastEthernet0/0
O     192.168.2.0/24 [110/20] via 192.168.4.2, 00:10:26, FastEthernet0/1
C     192.168.3.0/24 is directly connected, FastEthernet0/0
```

Figure IV-19 : Table de routage du R8

IV . 7 Configuration de MPLS

Dans cette étape, nous procédons à l'implémentation de la technologie MPLS sur notre plateforme. Pour cela, nous activons tout d'abord le protocole de distribution des labels (LDP) sur chaque routeur, puis nous activons MPLS sur les interfaces concernées. Cette configuration permet aux routeurs de construire une table de labels et de préparer le réseau à l'acheminement des paquets selon le mécanisme de commutation d'étiquettes.

Configuration du routeur R1

```
R1>en
R1#conf t
R1(config)#ip cef
R1(config)#mpls ip
R1(config)#mpls label protocol ldp
R1(config)#mpls ldp router-id loopback0 force
R1(config)#int f0/0
R1(config-if)#mpls ip
R1(config-if)#exit
R1(config)#int f0/1
R1(config-if)#mpls ip
R1(config-if)#exit
R1(config)#int f1/0
R1(config-if)#mpls ip
R1(config-if)#exit
R1(config)#int f2/0
R1(config-if)#mpls ip
R1(config-if)#end
```

Discussion

- R1(config)#ip cef

Cette commande active Cisco Express Forwarding (*CEF*) sur le routeur. Le CEF est une technologie de commutation rapide qui optimise le transfert des paquets IP.

- R1(config)#mpls ip

Cette commande active MPLS (*Multiprotocol Label Switching*) globalement sur le routeur. Le MPLS permet de transférer les paquets selon des labels au lieu des adresses IP classiques.

- R1(config)#mpls label protocol ldp

Cette commande spécifie que le protocole de distribution des labels est LDP (Label Distribution Protocol), utilisé pour échanger les labels entre routeurs.

- R1(config)#mpls ldp router-id loopback 0 force

Cette commande définit l'identifiant du routeur pour LDP à partir de l'adresse de l'interface Loopback 0. On utilise souvent une interface loopback car elle reste toujours active. L'option force applique immédiatement cette configuration.

- R1(config)#int F1/0

Ceci est un raccourci pour sélectionner l'interface FastEthernet1/0 en mode configuration.

- R1(config-if)#mpls ip

Cette commande active MPLS sur l'interface sélectionnée, lui permettant de traiter les paquets labellisés.

Vérification de la configuration MPLS

Pour vérifier si le protocole MPLS est bien activé sur les interfaces du routeur, on utilise la commande suivante : `show mpls interfaces`, la figure (IV-20). Cette commande permet d'afficher les interfaces sur lesquelles MPLS est activé, ainsi que leur état de fonctionnement.

```
R1#sh mpls int
Interface          IP          Tunnel  Operational
FastEthernet0/0    Yes (ldp)   No       Yes
FastEthernet0/1    Yes (ldp)   No       Yes
FastEthernet1/0    Yes (ldp)   No       Yes
FastEthernet2/0    Yes (ldp)   No       Yes
```

Figure IV-20: Visualise les interfaces MPLS activées sur R1

Pour vérifier l'établissement des relations de voisinage via le protocole LDP sur le routeur R1, nous avons utilisé la commande de la figure (IV-21) :

```
R1#show mpls ldp neighbor
Peer LDP Ident: 196.0.0.4:0; Local LDP Ident 196.0.0.1:0
TCP connection: 196.0.0.4.22310 - 196.0.0.1.646
State: Oper; Msgs sent/rcvd: 12/12; Downstream
Up time: 00:01:25
LDP discovery sources:
  FastEthernet2/0, Src IP addr: 192.168.3.2
Addresses bound to peer LDP Ident:
  192.168.3.2      196.0.0.4      192.168.4.1
Peer LDP Ident: 196.0.0.2:0; Local LDP Ident 196.0.0.1:0
TCP connection: 196.0.0.2.22839 - 196.0.0.1.646
State: Oper; Msgs sent/rcvd: 12/12; Downstream
Up time: 00:01:17
LDP discovery sources:
  FastEthernet0/0, Src IP addr: 192.168.1.2
Addresses bound to peer LDP Ident:
  192.168.2.1      196.0.0.2      192.168.1.2
```

Figure IV-21 : Affichage les voisins LDP détectés par R1

On passe de R3

Configuration de routeur R3

```
R3>en
R3#conf t
R3(config)#ip cef
R3(config)#mpls ip
R3(config)#mpls label protocol ldp
R3(config)#mpls ldp router-id loopback0 force
R3(config)#int f0/0
R3(config-if)#mpls ip
R3(config-if)#exit
R3(config)#int f0/1
R3(config-if)#mpls ip
R3(config-if)#exit
R3(config)#int f1/0
R3(config-if)#mpls ip
R3(config-if)#exit
R3(config)#int f2/0
R3(config-if)#mpls ip
R3(config-if)#end
```

Pour identifier les interfaces sur lesquelles le protocole MPLS est activé sur le routeur R3, nous avons utilisé la commande présentée dans la figure (IV-22)suivante :

```
R3#sh mpls int
Interface          IP          Tunnel    Operational
FastEthernet0/0    Yes (ldp)   No        Yes
FastEthernet0/1    Yes (ldp)   No        Yes
FastEthernet1/0    Yes (ldp)   No        Yes
FastEthernet2/0    Yes (ldp)   No        Yes
```

Figure IV-22: Visualise les interfaces MPLS activées sur R3

Pour vérifier l'établissement des relations de voisinage via le protocole LDP sur le routeur R3, figure (IV-23).

```
R3#show mpls ldp neighbor
Peer LDP Ident: 196.0.0.2:0; Local LDP Ident 196.0.0.3:0
TCP connection: 196.0.0.2.646 - 196.0.0.3.60234
State: Oper; Msgs sent/rcvd: 11/11; Downstream
Up time: 00:00:26
LDP discovery sources:
FastEthernet0/1, Src IP addr: 192.168.2.1
Addresses bound to peer LDP Ident:
192.168.2.1    196.0.0.2    192.168.1.2
Peer LDP Ident: 196.0.0.4:0; Local LDP Ident 196.0.0.3:0
TCP connection: 196.0.0.4.14029 - 196.0.0.3.646
State: Oper; Msgs sent/rcvd: 11/11; Downstream
Up time: 00:00:14
LDP discovery sources:
FastEthernet2/0, Src IP addr: 192.168.4.1
Addresses bound to peer LDP Ident:
192.168.3.2    196.0.0.4    192.168.4.1
```

Figure IV-23: Affichage les voisins LDP détectés par R3

Pour R2 :

Configuration de routeur R2

```
R2>en
```

```
R2#conf t
```

```
R2(config)#ip cef
```

```
R2(config)#mpls ip
```

```
R2(config)#mpls label protocol ldp
```

```
R2(config)#mpls ldp router-id loopback0 force
```

```
R2(config)#int f0/0
```

```
R2(config-if)#mpls ip
```

```
R2(config-if)#exit
```

```
R2(config)#int f0/1
```

```
R2(config-if)#mpls ip
```

```
R2(config-if)#end
```

Pour identifier les interfaces sur lesquelles le protocole MPLS est activé sur le routeur R2, nous avons utilisé la commande suivante :

```
R2#sh mpls int
Interface          IP          Tunnel    Operational
FastEthernet0/0    Yes (ldp)   No        Yes
FastEthernet0/1    Yes (ldp)   No        Yes
```

Figure IV-24 : Visualise les interfaces MPLS activées sur R2

Pour vérifier l'établissement des relations de voisinage via le protocole LDP sur le routeur R2, nous avons utilisé la commande suivante :

```
R2#sh mpls ldp neighbor
Peer LDP Ident: 196.0.0.3:0; Local LDP Ident 196.0.0.2:0
TCP connection: 196.0.0.3.15617 - 196.0.0.2.646
State: Oper; Msgs sent/rcvd: 46/46; Downstream
Up time: 00:30:50
LDP discovery sources:
  FastEthernet0/0, Src IP addr: 192.168.2.2
Addresses bound to peer LDP Ident:
  192.168.2.2      196.0.0.3      192.168.4.2
Peer LDP Ident: 196.0.0.1:0; Local LDP Ident 196.0.0.2:0
TCP connection: 196.0.0.1.646 - 196.0.0.2.60670
State: Oper; Msgs sent/rcvd: 45/45; Downstream
Up time: 00:30:40
LDP discovery sources:
  FastEthernet0/1, Src IP addr: 192.168.1.1
Addresses bound to peer LDP Ident:
  192.168.1.1      196.0.0.1      192.168.3.1
```

Figure IV-25 : Affichage les voisins LDP détectés par R2

Pour R8 :

Configuration de routeur R8

```
R8>en
R8#conf t
R8(config)#ip cef
R8(config)#mpls ip
R8(config)#mpls label protocol ldp
R8(config)#mpls ldp router-id loopback0 force
R8(config)#int f0/0
R8(config-if)#mpls ip
R8(config-if)#exit
R8(config)#int f0/1
R8(config-if)#mpls ip
R8(config-if)#end
```

Pour identifier les interfaces sur lesquelles le protocole MPLS est activé sur le routeur R8, nous avons utilisé la commande suivante :

```
R8#sh mpls int
Interface          IP          Tunnel    Operational
FastEthernet0/0    Yes (ldp)   No        Yes
FastEthernet0/1    Yes (ldp)   No        Yes
R8#
```

Figure IV-26 : Visualise les interfaces MPLS activées sur R8

Pour vérifier l'établissement des relations de voisinage via le protocole LDP sur le routeur R8, nous avons utilisé la commande suivante :

```
R8#sh mpls ldp neighbor
Peer LDP Ident: 196.0.0.1:0; Local LDP Ident 196.0.0.4:0
TCP connection: 196.0.0.1.646 - 196.0.0.4.39103
State: Oper; Msgs sent/rcvd: 48/48; Downstream
Up time: 00:32:33
LDP discovery sources:
FastEthernet0/0, Src IP addr: 192.168.3.1
Addresses bound to peer LDP Ident:
192.168.1.1      196.0.0.1      192.168.3.1
Peer LDP Ident: 196.0.0.3:0; Local LDP Ident 196.0.0.4:0
TCP connection: 196.0.0.3.646 - 196.0.0.4.37412
State: Oper; Msgs sent/rcvd: 48/47; Downstream
Up time: 00:32:23
LDP discovery sources:
FastEthernet0/1, Src IP addr: 192.168.4.2
Addresses bound to peer LDP Ident:
192.168.2.2      196.0.0.3      192.168.4.2
```

Figure IV-27 : Affichage les voisins LDP détectés par R8

IV . 8 Configuration de VRF

Configuration de routeur R1

```
R1>en
R1#conf t
R1(config)#ip vrf client
R1(config-vrf)#rd 65000:1
R1(config-vrf)#route-target export 65000:1
R1(config-vrf)#route-target import 65000:1
R1(config-vrf)#exit
R1(config)#interface FastEthernet 0/1
R1(config)#ip vrf forwarding client
R1(config-vrf)#ip address 10.4.1.1 255.255.255.0
R1(config-if)#no shutdown
R1(config-vrf)#end
```

Discussion :

- R1(config)#ip vrf client

Cette commande permet de créer une instance VRF nommée client. La technologie VRF (*Virtual Routing and Forwarding*) permet la segmentation des tables de routage sur un seul

routeur. Cela signifie que le routeur peut gérer plusieurs domaines de routage totalement isolés les uns des autres.

- R1(config-vrf)#rd 65000:1

Cette commande définit le Route Distinguisher (*RD*) de l'instance VRF client avec la valeur 65000:1. Le RD sert à différencier les routes provenant de différentes VRFs, ce qui permet l'utilisation d'adresses IP qui peuvent être identiques mais dans des contextes VRF séparés.

- R1(config-vrf)#route -target export 65000:1

Cette commande définit le Route Target (*RT*) utilisé pour l'exportation des routes de la VRF client. Le RT permet de marquer les routes afin de contrôler leur distribution via le protocole MP-BGP, en spécifiant vers quelles autres VRFs ces routes peuvent être envoyées.

- R1(config-vrf)#route -target import 65000:1

Cette commande définit le Route Target (*RT*) utilisé pour l'importation des routes vers la VRF client. Ainsi, seules les routes marquées avec le RT 65000:1 seront acceptées dans cette VRF.

- R1(config-vrf)#exit

Cette commande permet de quitter le mode de configuration de la VRF pour revenir en mode de configuration global.

- R1(config)#interface FastEthernet 0/1

Cette commande sélectionne l'interface FastEthernet 0/0 pour lui appliquer une configuration.

- R1(config-if)#ip vrf forwarding client

Cette commande associe l'interface à la VRF client. Cela signifie que tout le trafic qui transite par cette interface utilisera exclusivement la table de routage spécifique à cette VRF.

- R1(config-if)#ip address 10.4.1.1 255.255.255.0

Cette commande attribue l'adresse IP 10.4.1.1/24 à l'interface. Cette adresse appartient désormais au domaine de routage de la VRF client.

Pour afficher les VRF déjà configurées sur le routeur PE1, ainsi que les interfaces auxquelles les clients doivent être connectés, on peut utiliser la commande `show ip vrf`. Cette commande permet de visualiser la liste des VRF actives ainsi que les interfaces qui leur sont associées, ce qui facilite l'identification des points de connexion client.

```
R1#sh ip vrf
Name                Default RD          Interfaces
client              65000:1            Fa0/1
                   Fa1/0
```

Figure IV-28: Montre la configuration des VRF sur R1

Pour R3 :

Configuration de routeur R3

```
R3>en
R3#conf t
R3(config)#ip vrf client
R3(config-vrf)#rd 65000:1
R3(config-vrf)#route-target export 65000:1
R3(config-vrf)#route-target import 65000:1
R3(config-vrf)#exit
R3(config)#interface FastEthernet 0/0
R3(config)#ip vrf forwarding client
R3(config-vrf)#ip address 10.3.1.1 255.255.255.0
R3(config-if)#no shutdown
R3(config-vrf)#end
```

Dans cette étape, nous affichons la configuration des instances VRF mises en place sur le routeur R3. Pour cela, nous avons utilisé la commande suivante :

```
R3#sh ip vrf
  Name                Default RD      Interfaces
  client              65000:1       Fa0/0
                   Fa1/0
```

Figure IV-29: Montre la configuration des VRF sur R3

IV . 9 Configuration de MP-BGP

Configuration de routeur R1

```
R1#enable
R1#conf t
R1(config)#router bgp 65000
R1(config-router)#neighbor 196.0.0.3 remote-as 65000
R1(config-router)#neighbor 196.0.0.3 update-source Loopback0
R1(config-router)#address-family vpnv4
R1(config-router-af)#neighbor 196.0.0.3 activate
R1(config-router-af)#neighbor 196.0.0.3 send-community extended
R1(config-router-af)#exit
R1(config)#router bgp 65000
R1(config-router)#address-family ipv4 vrf client
R1(config-router-af)#redistribute connected
R1(config-router-af)#exit
R1(config)#end
```

Discussion

- R1(config)#router bgp 65000

Cette commande permet d'activer le processus BGP sur le routeur avec l'AS numéro 65000.

- R1(config-router)#neighbor 196.0.0.3 remote-as 65000

Elle définit le voisin BGP ayant pour adresse 196.0.0.3, appartenant au même AS. Cela établit une session iBGP.

- R1(config-router)#neighbor 196.0.0.3 update-source Loopback0

Cette instruction force l'utilisation de l'interface Loopback0 comme source pour l'établissement de la session BGP, assurant ainsi une stabilité accrue.

- R1(config-router)#address-family vpnv4

Active le mode de configuration pour l'adresse familiale VPNv4, nécessaire pour transporter les routes MPLS entre les routeurs PE via MP-BGP.

- R1(config-router-af)#neighbor 196.0.0.3 activate

Active la session BGP pour l'adresse familiale VPNv4 avec le voisin spécifié.

- R1(config-router-af)#neighbor 196.0.0.3 send-community extended

Cette commande permet l'envoi des attributs communautaires étendus (*extended communities*), nécessaires notamment pour la gestion des Route Targets.

- R1(config-router)#address-family ipv4 vrf client

Permet d'accéder à la configuration BGP pour le VRF nommé "client", en utilisant l'adresse familiale IPv4.

- R1(config-router-af)#redistribute connected

Cette commande autorise l'injection des routes connectées dans le processus BGP pour le VRF concerné. Cela permet à BGP de propager ces routes vers les autres routeurs du réseau

En utilisant la commande `show ip route vrf client`, figure (IV-30), il est également possible d'afficher la table de routage spécifique au VRF nommé client

```

R1#sh ip route vrf client

Routing Table: client
Codes: 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

Gateway of last resort is not set

    10.0.0.0/24 is subnetted, 2 subnets
C       10.4.2.0 is directly connected, FastEthernet1/0
C       10.4.1.0 is directly connected, FastEthernet0/1

```

Figure IV-30: Affichage la table de routage associée a la VRF « client » sur R1

Configuration de routeur R3

```

R3#enable
R3#conf t
R3(config)#router bgp 65000
R3(config-router)#neighbor 196.0.0.1 remote-as 65000
R3(config-router)#neighbor 196.0.0.1 update-source Loopback0
R3(config-router)#address-family vpnv4
R3(config-router-af)#neighbor 196.0.0.1 activate
R3(config-router-af)#neighbor 196.0.0.1 send-community extended
R3(config-router-af)#exit
R3(config)#router bgp 65000
R3(config-router)#address-family ipv4 vrf client
R3(config-router-af)#redistribute connected
R3(config-router-af)#exit
R3(config)#end

```

Dans cette étape, nous procédons à l'affichage de la table de routage spécifique à la VRF nommée "client" sur le routeur R1. Pour cela, nous utilisons la commande de la figure (IV-31) suivante :

```
R3#sh ip route vrf client

Routing Table: client
Codes: 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

Gateway of last resort is not set

    10.0.0.0/24 is subnetted, 4 subnets
C       10.3.1.0 is directly connected, FastEthernet0/0
C       10.3.2.0 is directly connected, FastEthernet1/0
B       10.4.2.0 [200/0] via 196.0.0.1, 00:01:21
B       10.4.1.0 [200/0] via 196.0.0.1, 00:01:21
R3#
```

Figure IV-31: Affichage la table de routage associée a la VRF « client » sur R3

Ping : Dans cette étape, nous configurons un routage statique sur le routeur R4 afin de lui permettre d'atteindre le réseau [10.3.1.0/24](#), situé derrière un autre routeur, figure (IV-32) :

```
R4#conf t
Enter configuration commands, one per line.  End with CNTL/Z.
R4(config)#ip route 10.3.1.0 255.255.255.0 10.4.1.1
R4(config)#end
R4#
*Mar  1 01:12:58.391: %SYS-5-CONFIG_I: Configured from console by console
R4#PING 10.3.1.1

Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 10.3.1.1, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 64/88/96 ms
R4#PING 10.3.1.2

Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 10.3.1.2, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 96/121/140 ms
R4#
```

Figure IV-32: Configuration du routage statique sur R4 et test de connectivite

Dans cette étape, nous configurons un routage statique sur le routeur R5 afin de lui permettre d'atteindre le réseau [10.3.2.0/24](#), qui se trouve derrière un autre routeur, figure (IV-33) :

```
R5#conf t
Enter configuration commands, one per line. End with CNTL/Z.
R5(config)#ip route 10.3.2.0 255.255.255.0 10.4.2.1
R5(config)#END
R5#
*Mar  1 01:14:02.795: %SYS-5-CONFIG_I: Configured from console by console
R5#ping 10.3.2.1

Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 10.3.2.1, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 68/86/112 ms
R5#ping 10.3.2.2

Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 10.3.2.2, timeout is 2 seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 104/126/168 ms
R5#
```

Figure IV-33: Configuration du routage statique sur R5 et test de connectivité

IV . 10 Conclusion

Ce chapitre a été consacré à la mise en œuvre pratique d'un réseau IP-MPLS à l'aide de l'environnement de simulation GNS3. Après avoir installé et configuré l'outil, nous avons construit une topologie réseau représentative d'une architecture opérateur. Chaque étape a été réalisée avec rigueur : de l'attribution des adresses IP à la configuration des interfaces et des protocoles de routage, jusqu'à l'activation du protocole MPLS. La suite de la simulation a permis d'intégrer les technologies avancées telles que la segmentation du réseau via les instances VRF et l'échange de routes multi-VRF grâce au protocole MP-BGP et l'utilisation des adresses VPNv4. Ce travail a ainsi démontré notre capacité à modéliser un réseau de cœur moderne, à la fois fiable, évolutif et adapté aux besoins des fournisseurs de services VPN de niveau

3.

CONCLUSION GENERALE

Conclusion générale

Le mémoire présenté se concentre sur l'étude et la mise en œuvre d'un réseau IP/MPLS au sein du centre de l'opérateur mobile Ooredoo. Il s'agira de mettre en œuvre les fonctionnalités MPLS (label switching, PE/P, VRF et MP-BGP) sur des équipements réseau CISCO dans le simulateur GNS3, de configurer les VPN de niveau 3 et de vérifier leur bon fonctionnement à travers des scénarios de test concrets. Les solutions classiques de routage ou de VPN GRE atteignent rapidement leurs limites en termes de scalabilité, d'isolation des clients et de gestion centralisée. Dans ce contexte, l'intégration d'une architecture IP/MPLS combinée aux VPN de niveau 3 (L3VPN) s'impose comme une réponse adaptée, spécialement dans les réseaux de fournisseurs de service nécessitant l'isolation des différents clients.

Dans un premier temps, nous avons établi une petite étude introductive sur l'opérateur Ooredoo, en présentant son historique en Algérie ainsi que ses valeurs. Par la suite, nous avons abordé en détails les fondements théoriques principaux dans la compréhension du sujet, allant des topologies et architectures jusqu'aux protocoles et commutation.

Dans un deuxième temps, nous avons examiné les concepts du MPLS et constaté que l'application des VPN grâce à son mécanisme d'acheminement d'étiquettes a connu un grand succès.

Avec ces bases, on a pu s'attaquer à la partie technique de notre projet, surtout l'implémentation d'un réseau IP MPLS, qui est une technologie moderne et qui fait bien le job en termes de performance, de flexibilité et de sécurité.

Finalement, la simulation faite avec le logiciel GNS3 a confirmé les configurations mises en place et montré que MPLS fonctionne bien pour le transport de données dans un réseau organisé.

Ce mémoire nous a donné l'occasion de consolider ce que nous avons appris, de voir comment fonctionne une entreprise de télécommunications et de mettre en œuvre une véritable solution technique qui répond aux besoins des réseaux modernes.

Références

- [1] : Ooredoo. (n.d.). *Tout sur Ooredoo – Ooredoo en détails – Ooredoo Algérie*. Site officiel Ooredoo. Consulté le 24 mars 2025, sur <https://www.ooredoo.dz/>
- [2] : Bentchakal, R., & Djanane, I. *Contribution des méthodes prévisionnelles dans l'amélioration de la performance des entreprises : Ooredoo Algérie* [Mémoire de Master, Université M'hamed Bougara de Boumerdes], (2022).
- [3] : Tanenbaum, A. S., & Wetherall, D. J. *Computer Networks* (5e éd.). Pearson Education, (2011).
- [4] : Rziza, M. *Cours des réseaux informatiques*. [Document pédagogique],(2010–2011).
- [5] : Toumi, M. *Le modèle OSI et le modèle TCP/IP* [Cours universitaire],(2017).
- [6] : Cisco Networking Academy. *Introduction to Networks v7*. Cisco Press,(2022).
- [7] : Forouzan, B. A. *Data Communications and Networking* (5e éd.). McGraw-Hill,(2013).
- [8] : Toumi, M. *Cours n° 3 : Les équipements d'interconnexion* [Cours universitaire],(2013).
- [9] : Charifi, T. *Conception d'un modem PLC implémentation sur FPGA* [Mémoire d'ingénieur, École Nationale Polytechnique], (2006).
- [10] : Noui, H. (n.d.). *Chapitre 5 : Routage* [Cours, 3e année Licence Informatique].
- [11] : Hamouma, M. *Routage et interconnexion* [Cours, Université de Batna 2, Département d'informatique], (2022).
- [12] : Riahla. *Commutation LAN* [Cours, Université de Boumerdes], (2017).
- [13] : Hatira, N. (n.d.). *Cours de commutation* [Support de cours].
- [14] : Sehaba, K. (n.d.). *Cours de réseaux*. Université Lumière Lyon 2.
- [15] : Palo Alto Networks. (n.d.). *What is Multiprotocol Label Switching (MPLS)?* <https://www.paloaltonetworks.com/>
- [16] : Ould Lamara, S., & Takilt, M. *Implémentation du SDN dans une structure IP/MPLS* [Mémoire de Master, Université Mouloud Mammeri de Tizi-Ouzou], (2018).
- [17] : Atman, Y. A., & Khenafif, S. A. *Ingénierie de trafic d'un réseau VPN L3 IP/MPLS* [Mémoire de Master, Université Saad Dahleb Blida 1], (2024).
- [18] : Awduche, D., Berger, L., Gan, D., Li, T., & Swallow, G. (2001). *RFC 3209 – RSVP-TE: Extensions to RSVP for LSP Tunnels*. Internet Engineering Task Force (IETF). <https://datatracker.ietf.org/doc/html/rfc3209>
- [19] : Andersson, L., et al. (2007). *RFC 5036 – LDP Specification*. IETF. <https://datatracker.ietf.org/doc/html/rfc5036>
- [20] : Rosen, E. C., & Rekhter, Y. (2006). *RFC 4364 – BGP/MPLS IP Virtual Private Networks (VPNs)*. IETF. <https://datatracker.ietf.org/doc/html/rfc4364>

References

- [21] : Cisco Systems. *Configurer un réseau VPN MPLS de base*. Cisco Systems, (2022).
- [22] : Bhatia, A. (n.d.). *HVAC Variable Refrigerant Flow (VRF) Systems* (Course No: M03-014, Credit: 3 PDH). <https://www.cedengineering.com>
- [23] : Cisco Systems. (n.d.). *MP-BGP MPLS VPN Configuration Guide*. Cisco Systems. <https://www.cisco.com/>
- [24] : Cisco Systems. (n.d.). *Implementing MPLS Forwarding*. Cisco Systems. <https://www.cisco.com/>

ANNEXE 1. Installation de GNS3

Le logiciel GNS3 a été adopté pour la simulation du réseau IP-MPLS en raison de sa capacité à exécuter des images réelles d'équipements réseau. Nous l'avons installé sur un système Windows 10, avec l'activation de GNS3 VM pour de meilleures performances. Après l'installation, les images IOS des routeurs ont été ajoutées via les paramètres du programme, ce qui nous a permis de créer un environnement de simulation réaliste pour tester différents scénarios réseau.

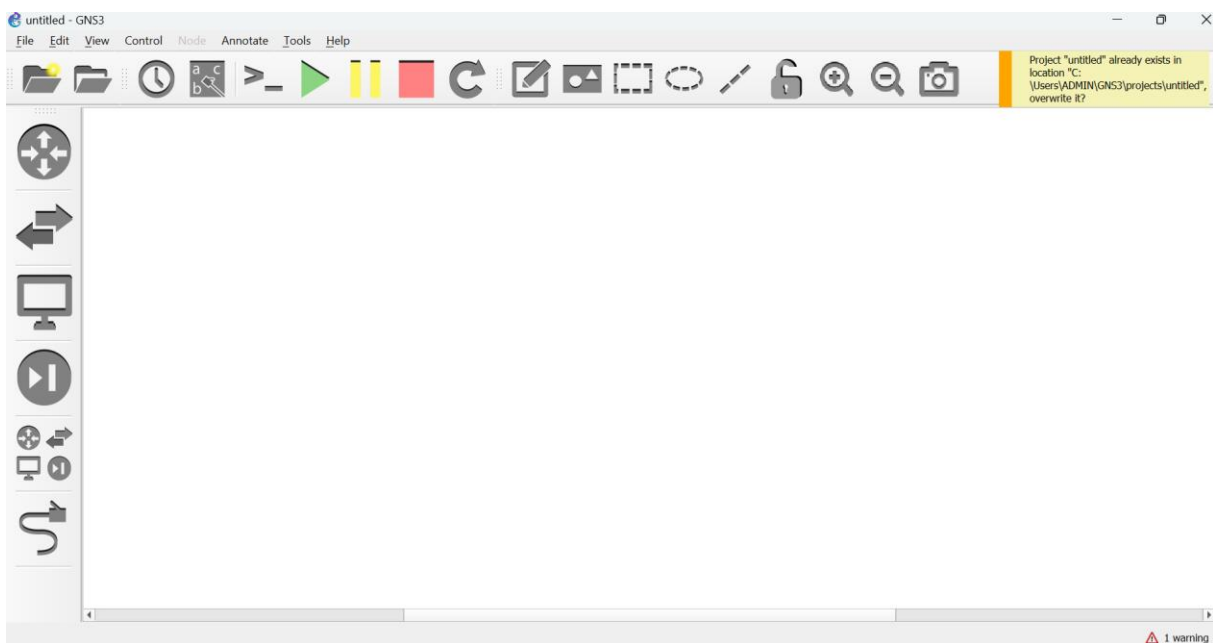


Figure Annexe-1 : Interface principale de GNS3

Voici les étapes à suivre :

1. Allez dans "Édition" puis sélectionnez "Préférences" dans la barre d'outils.
2. Dans la section des images IOS, ajoutez l'image que vous avez téléchargée.

Cette étape garantit que vous disposez des images IOS nécessaires pour configurer et simuler les équipements réseau dans GNS3, vous permettant ainsi de créer des environnements de test réalistes

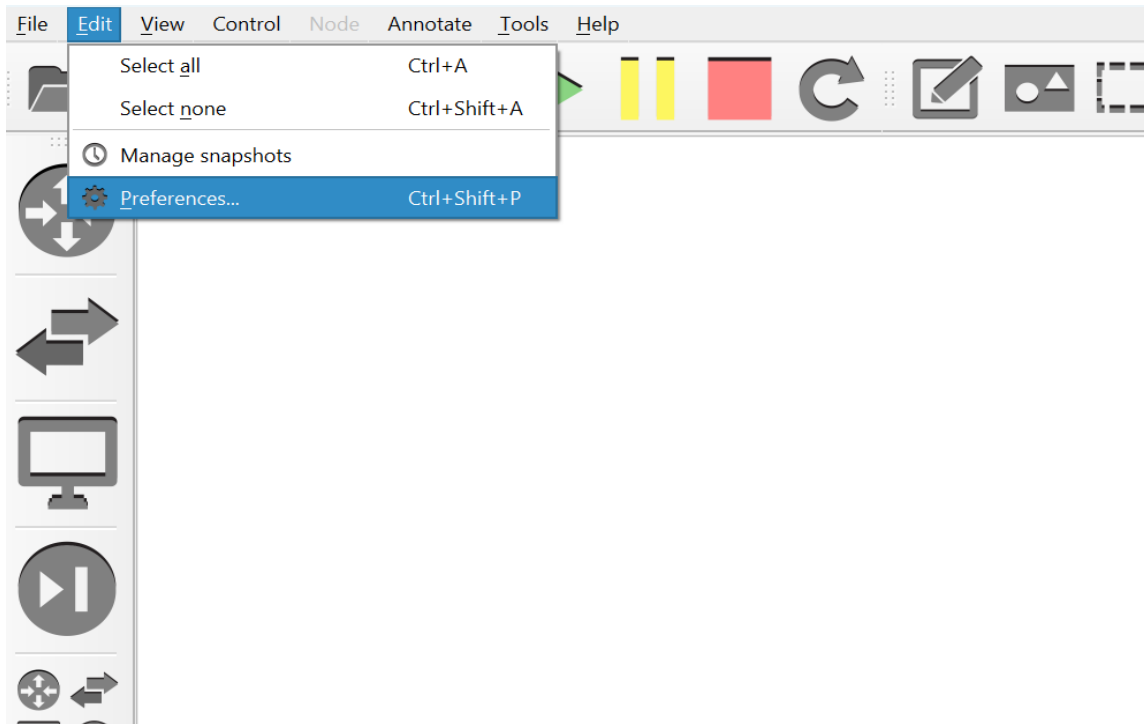


Figure Annexe-2 : Menu « Edit » ouvert avec "Préférences" sélectionné

Une fois dans les "Préférences", cliquez sur **IOS Routers**

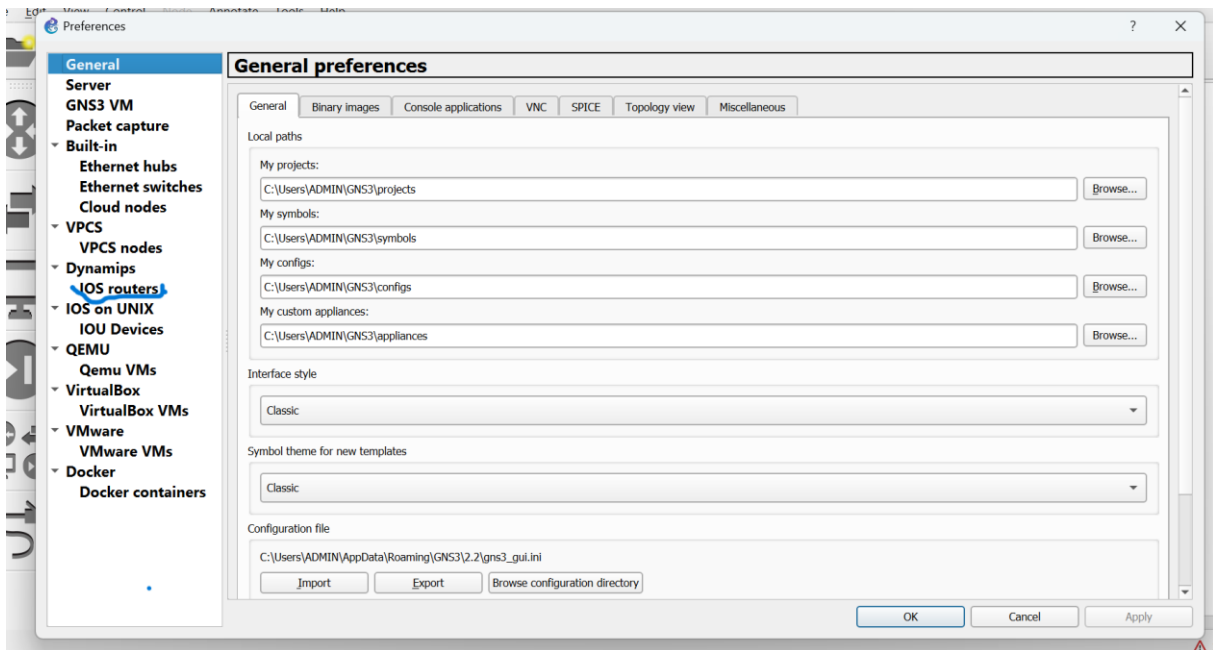


Figure Annexe-3 : Fenêtre « General Préférences » dans GNS3

Sur cette interface, cliquez sur le bouton « Browse ». Une fenêtre s'ouvre alors, vous permettant de sélectionner l'image souhaitée depuis son emplacement dans le répertoire.

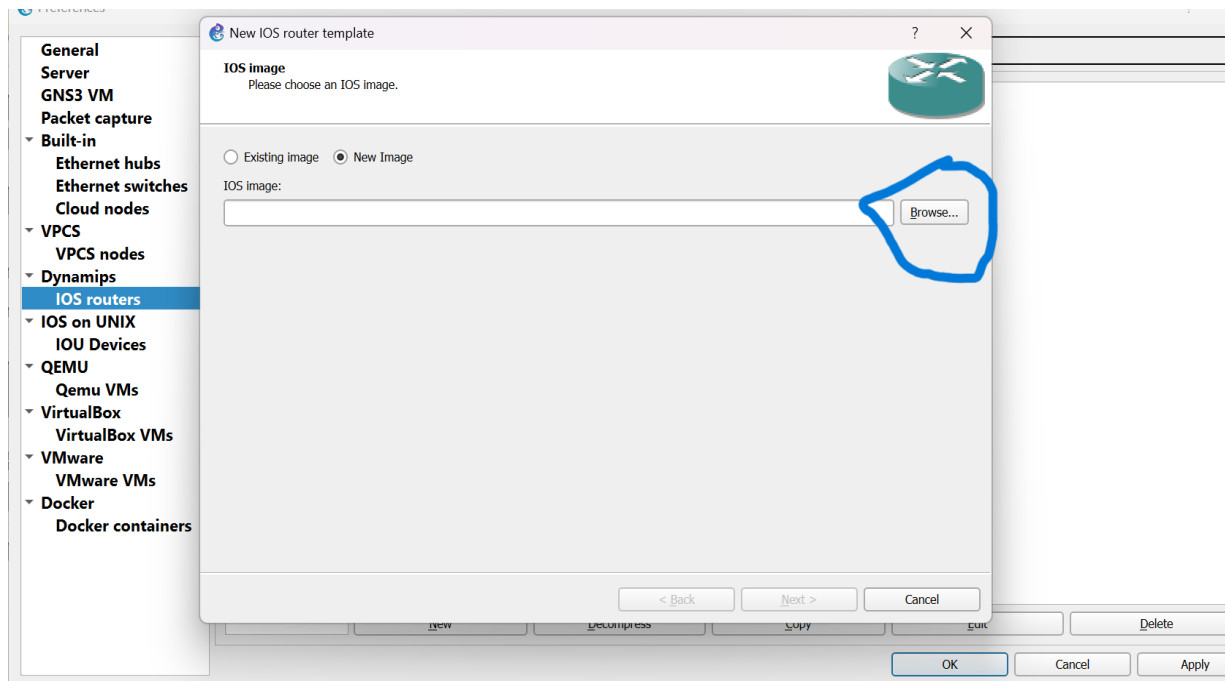


Figure Annexe-4 : Début du processus d'ajout d'un routeur IOS

Une fois les images sélectionnées, cliquez sur « Ouvrir » en bas de la fenêtre. Les images seront alors importées automatiquement.

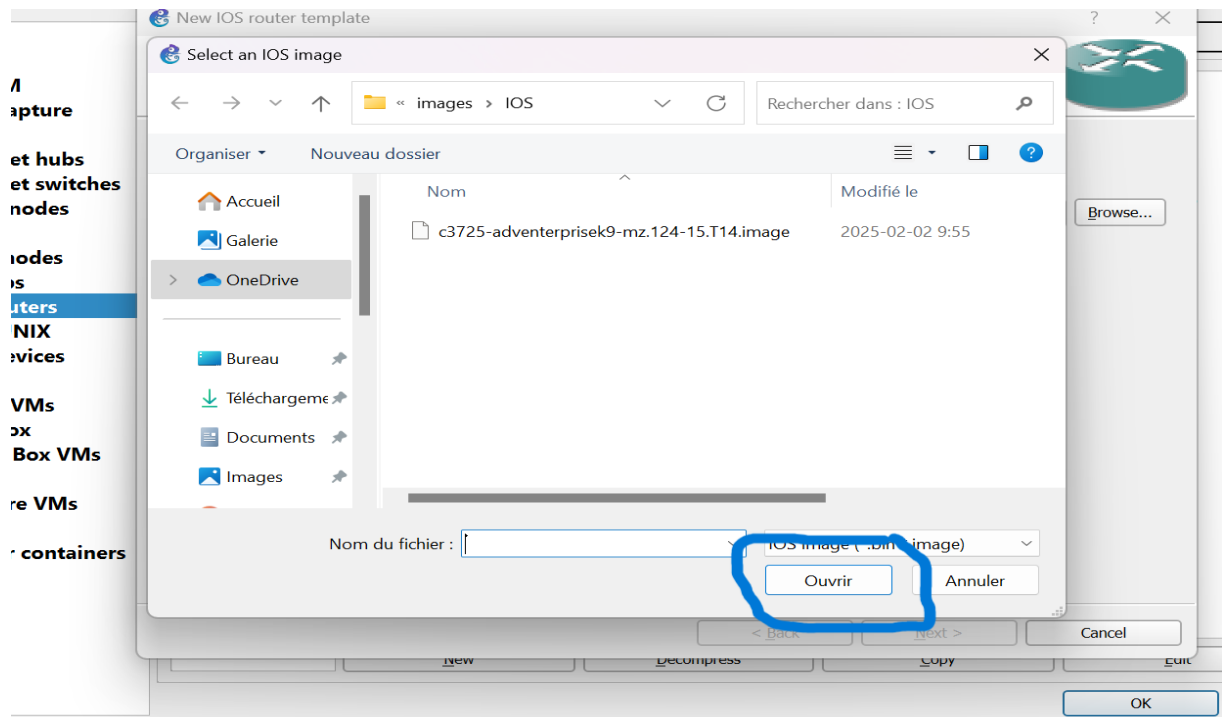


Figure Annexe-5 : Sélection du fichier image IOS dans GNS3

L'image sélectionnée sera configurée en cliquant sur « Next ».

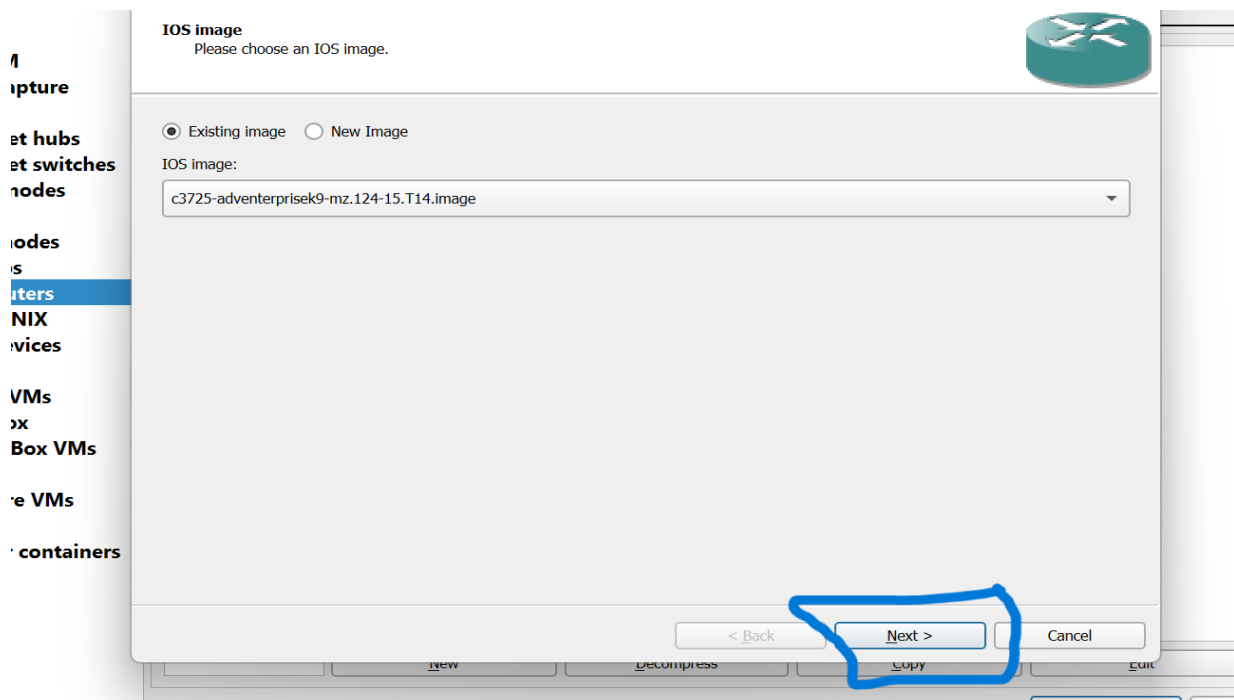


Figure Annexe-6 : Image IOS sélectionnée pour créer un nouveau modèle de routeur

ANNEXE

On obtient une fenêtre comme celle illustrée ci-dessous, où les champs de saisie sont remplis automatiquement avec le lien de l'image dans la section « fichier image ». Dans cet exemple, nous sélectionnons le routeur C3725.

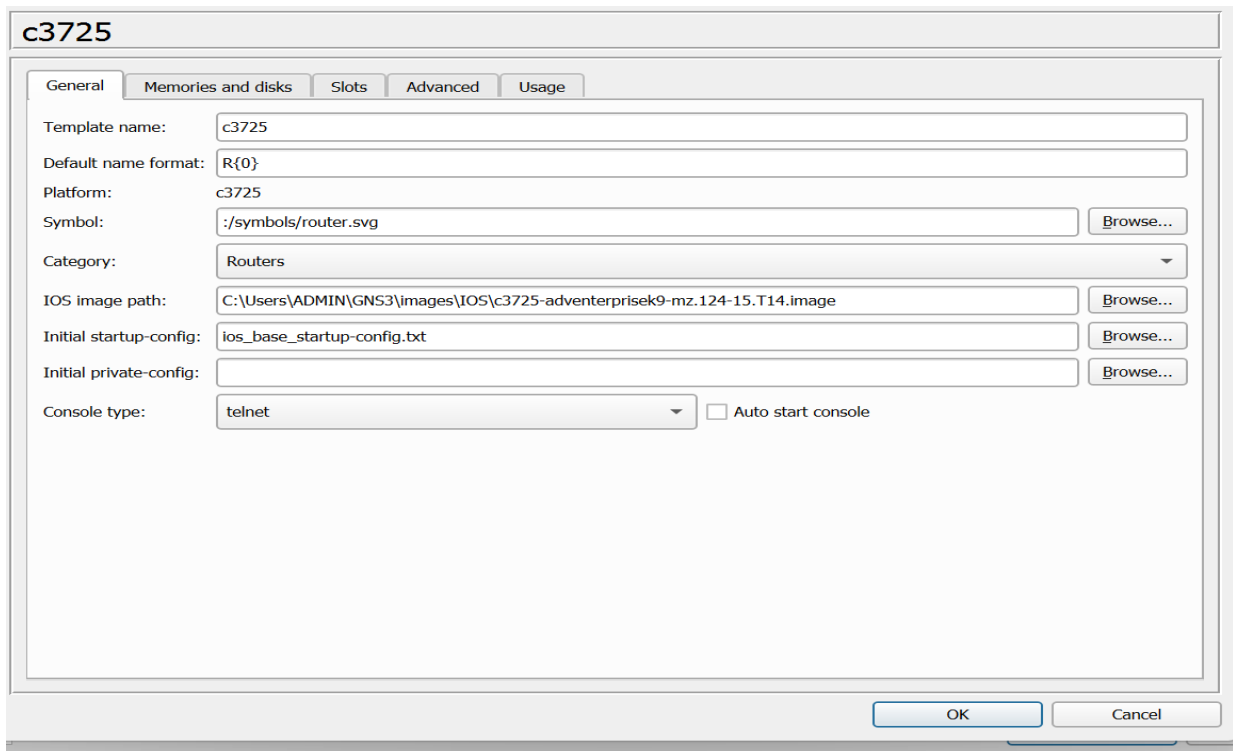


Figure Annexe-7 : Paramètres généraux du modèle

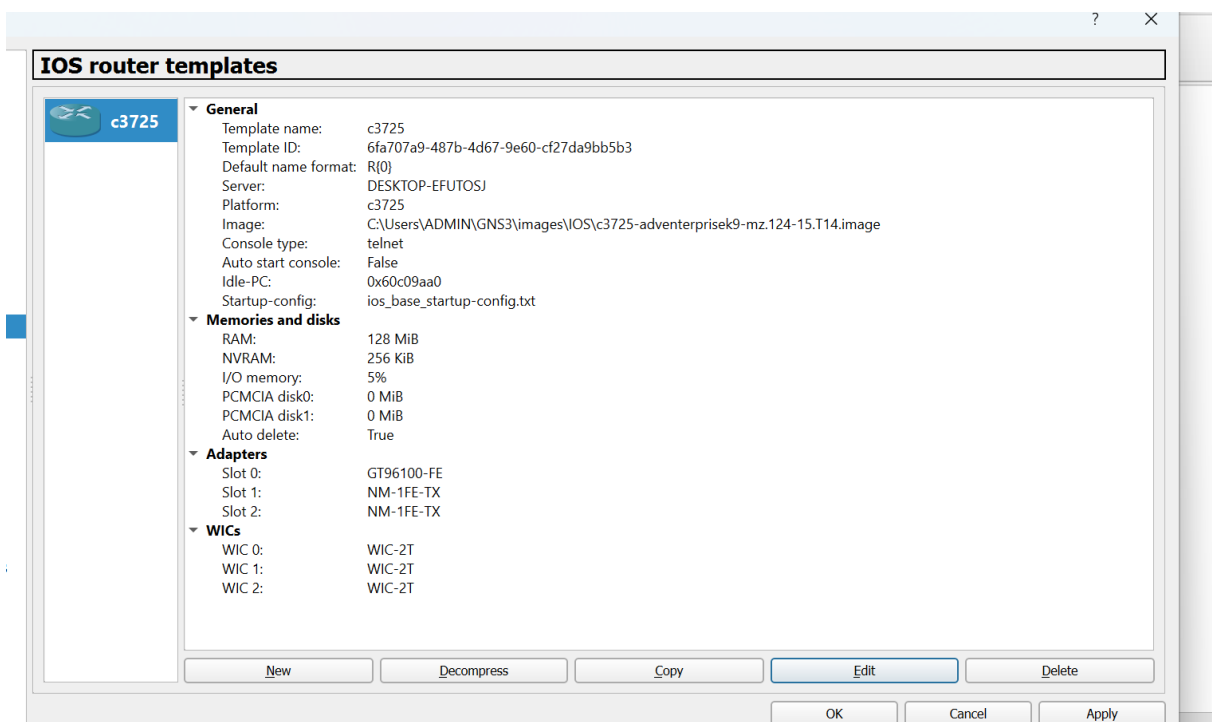


Figure Annexe-8 : Affiche les détails du modèle dans la fenêtre « IOS router Template » de GNS3

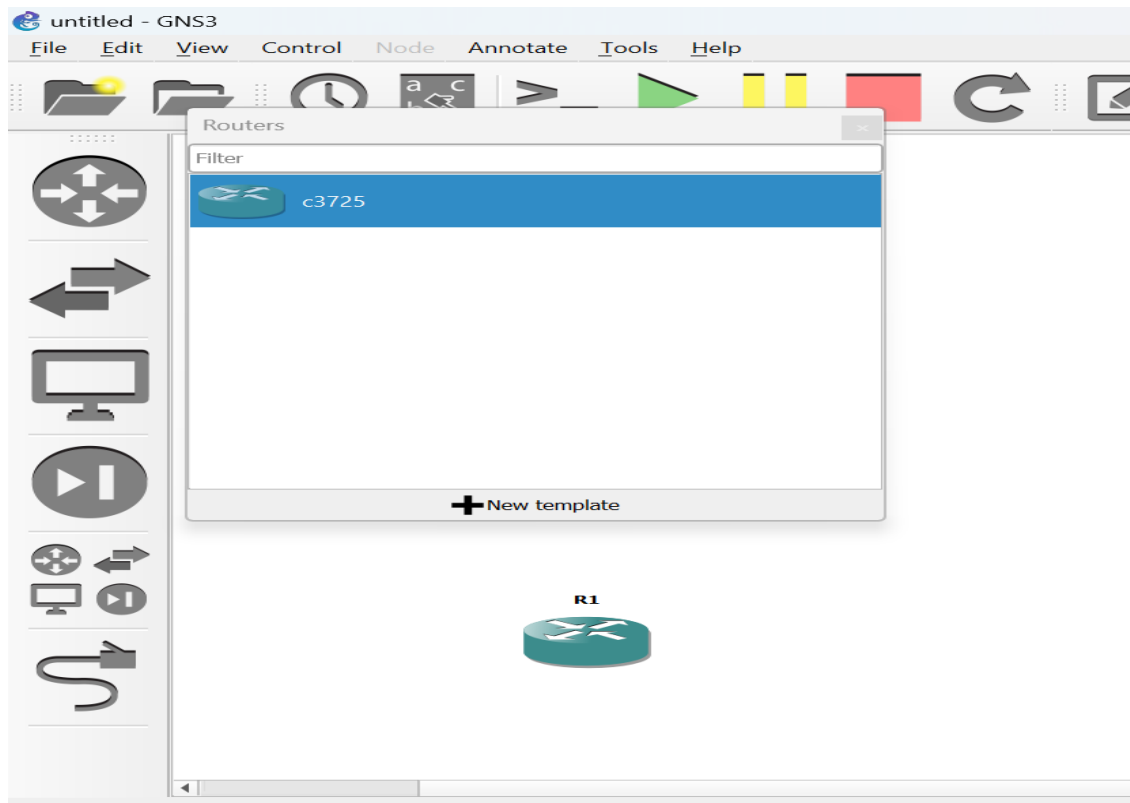


Figure Annexe-9 : Présente la fenêtre « Router » de GNS3 avec l'ajout du routeur R1 dans le projet