

**PONTIFICIA UNIVERSIDAD CATÓLICA DEL
ECUADOR
FACULTAD DE INGENIERÍA
MAESTRÍA EN REDES DE COMUNICACIONES**

**OPTIMIZACIÓN DE LA RED DE ACCESO IP PARA
INTERCONECTAR NODOS LTE (IP RAN) HACIA EL
CORE DE SERVICIOS DE LA PLATAFORMA DE
DATOS MÓVILES**

PAREDES VALENCIA JOHN DAVID

**PROYECTO PREVIO A LA OBTENCIÓN DEL TÍTULO
DE MAGISTER EN REDES DE COMUNICACIONES**

**Directora:
MSc. María Soledad Jiménez**

Quito, Enero del 2016

AGRADECIMIENTO

Agradezco en primer lugar a Dios por sus dones y luz, a la Ing. María Soledad Jiménez MSc por su incondicional apoyo y acertada dirección, a mi esposa Silvita por los sacrificios que representaron la base de este trabajo, a mis padres Juan y Rosita por su ejemplo y amor, a mis amigos que constituyeron un pilar importante en la recolección de información valiosa para este desarrollo y a la Pontificia Universidad Católica del Ecuador de la mano de Gustavo Chafra PhD por brindar un espacio de crecimiento profesional.

DEDICATORIA

El presente trabajo se lo dedico a mi esposa Silvita y a mi hija Gabriela Estefanía, por el incondicional apoyo, amor y por ser la inspiración de mi vida.

A mis padres Juan y Rosita, por su ejemplo, por brindarme sus sacrificios amistad y confianza.

John Paredes

ÍNDICE GENERAL

AGRADECIMIENTO	I
DEDICATORIA	II
ÍNDICE DE FIGURAS	VIII
ÍNDICE DE TABLAS	XI
CAPÍTULO 1 MARCO TEÓRICO	1
1.1 ANTECEDENTES	1
1.1.1 CRECIMIENTO DEL CONSUMO DE INTERNET MÓVIL	2
1.1.2 NUEVOS REQUERIMIENTOS DE LAS REDES DE TRANSPORTE PARA LAS REDES MÓVILES DE UN PROVEEDOR DE SERVICIOS	5
1.1.3 EVOLUCIÓN DE LAS REDES MÓVILES	8
1.1.4 CONCEPTOS BÁSICOS DE REDES MÓVILES LTE	13
1.1.4.1 ARQUITECTURA DE LAS REDES MÓVILES LTE COMO UNA EVOLUCIÓN DE LAS TECNOLOGÍAS UMTS Y GSM ^[7]	14
1.1.4.2 ARQUITECTURA DE LTE [7]	15
1.1.4.2.1 DISEÑO DE ALTO NIVEL DE LTE [7]	16
1.1.4.2.2 INTERCONEXIÓN DEL USUARIO MÓVIL Y LOS SERVICIOS LTE [7]	18
1.2 BENEFICIOS DE UNA IMPLEMENTACIÓN IP PARA EL ACCESO DE RADIO DE TECNLOGÍA LTE, IP-RAN	22
1.2.1 GENERALIDADES DE UNA RED DE ACCESO DE RADIO	22
1.2.2 LAYER 3 VS LAYER 2 (EJEMPLOS DE CONEXIONES DE BACKHAUL)	23
1.2.3 RECOMENDACIONES	25

1.2.4	GENERALIDADES DE LA CONECTIVIDAD LTE	27
1.3	UTILIZACIÓN DE LAS REDES MPLS	28
1.3.1	APLICACIONES DE REDES MPLS, L2VPN Y L3VPN	36
1.3.1.1	L3VPN	37
1.3.1.2	L2VPN	42
1.3.1.2.1	MODELOS DE SERVICIO L2VPN	43
1.3.1.2.2	ARQUITECTURA DEL SERVICIO L2VPN	43
1.3.1.2.3	L2VPN PLANO DE CONTROL.....	46
1.3.1.2.4	PLANO DE DATOS L2VPN PARA UNA RED IP, L2TPv3	47
1.3.1.2.5	PLANO DE DATOS L2VPN PARA UNA RED MPLS	48
1.3.1.2.6	TIPOS DE SERVICIOS DE UNA L2VPN.	49
1.3.2	ARQUITECTURAS DE REDES DE ACCESO IP-RAN PARA SERVICIOS LTE UTILIZANDO UNA RED MPLS COMO INTERCONEXIÓN AL CORE DE DATOS MÓVILES	51
1.3.3	SOLUCIONES DE L2VPN	52
1.3.4	SOLUCIONES DE L3VPN[5].....	55
1.3.5	SOLUCIONES CON INTER-AS VPN	57
1.3.5.1	MPLS ENTRE DIFERENTES DOMINIOS O INTER-AS MPLS.....	58
1.3.5.1.1	OPCIÓN A: CONEXIONES VRF-TO-VRF ENTRE LOS ASBR	58
1.3.5.1.2	OPCIÓN B: EBGp REDISTRIBUCIÓN DE RUTAS VPNv4 ETIQUETADAS.	60
1.3.5.1.3	OPCIÓN C: REDISTRIBUCIÓN DE RUTAS ETIQUETADAS VPNv4 POR UNA SESIÓN EBGp ENTRE LOS SISTEMAS AUTÓNOMOS, CON REDISTRIBUCIÓN DE ETIQUETAS IPv4 AL RESPECTIVO AS	63
1.3.5.2	SEAMLESS MPLS ^[14]	64
1.3.5.2.1	REDES DE ACCESO ETHERNET	65
1.3.5.2.2	TERMINOLOGÍA GENERAL DE REDES DE ACCESO MPLS	69

1.3.5.3	MECANISMOS DE TRANSPORTE CON SEAMLESS MPLS	71
1.3.5.3.1	LDP SIN SEGMENTACIÓN EN LA RED DE CORE Y AGREGACIÓN.....	75
1.3.5.3.2	SEGMENTACIÓN DEL IGP DE LA RED DE ACCESO Y UNIFICACIÓN DE LSP CON EL CORE Y AGREGACIÓN, MEDIANTE ETIQUETAS A TRAVÉS DE BGP	76
1.3.5.3.3	REDISTRIBUCIÓN DE ETIQUETAS DE BGP EN EL IGP DE LA RED DE ACCESO (OPCIONAL LDP DoD)	78
1.3.5.3.4	SEGMENTACIÓN DEL IGP DE LA RED DE AGREGACIÓN Y UNIFICACIÓN DE LSP CON EL CORE, MEDIANTE ETIQUETAS A TRAVÉS DE BGP	78
1.3.5.3.5	SEGMENTACIÓN DEL IGP DE LA RED DE ACCESO, AGREGACIÓN Y UNIFICACIÓN DE LSP CON EL CORE, MEDIANTE ETIQUETAS A TRAVÉS DE BGP	79
1.3.5.3.6	REDISTRIBUCIÓN DE ETIQUETAS DE BGP EN EL IGP AL INTERIOR DE LA RED DE ACCESO (OPCIONAL LDP DoD)	80
1.3.5.4	COMPARACIÓN DE LAS SOLUCIONES DE INTER-AS Y SEAMLESS MPLS	81
CAPÍTULO 2	ANÁLISIS DE LA RED ACTUAL IP- RAN DE LA CNT E.P. EN QUITO	85
2.1	DESCRIPCIÓN DE LA RED IP- RAN DE LA CNT E.P. EN QUITO	85
2.1.1	CONFIGURACIÓN DEL IS-IS	89
2.1.2	CONFIGURACIÓN DE MPLS LDP.....	90
2.1.3	CONFIGURACIÓN BGP	91
2.1.4	CONFIGURACIÓN DE MPLS VPN	93
2.1.5	CONFIGURACIÓN DE INTER-AS VPN	95
2.1.6	CONFIGURACIÓN DE QOS	99
2.1.7	CONFIGURACIÓN DEL SINCRONISMO DE LA RED LTE.....	102
2.2	ANÁLISIS DE LA RED IP- RAN DE LA CNT E.P. EN QUITO	103

2.2.1	CRECIMIENTO DE CLIENTES CON CRECIMIENTO DE CONSUMO DE ANCHO DE BANDA, TANTO DE DATOS COMO DE INTERNET.	103
2.2.2	CAPACIDADES DE LA SOLUCIÓN ACTUAL	104
2.2.2.1	EQUIPOS CSG Y ASG	106
2.2.2.2	EQUIPOS ASBR	107
2.2.2.2.1	CX-600 DE HUAWEI	108
2.2.2.2.2	CISCO ASR903	110
2.2.3	CONCLUSIONES DEL ANÁLISIS DE LA SOLUCIÓN IP-RAN DE LA CNT E.P.....	111
 CAPÍTULO 3 DISEÑO DE LA SOLUCIÓN DE REINGENIERÍA DE LA RED IP-RAN DE LA CNT E.P EN QUITO 116		
3.1	ANÁLISIS DE REQUERIMIENTOS	116
3.1.1	ANÁLISIS DE REQUERIMIENTOS CON RESPECTO AL DISEÑO DE INTER-AS	117
3.1.2	CRITERIOS DE DISEÑO	119
3.1.2.1	ESCALABILIDAD Y CAPACIDAD	120
3.1.2.2	RESILIENCIA.....	120
3.1.2.3	ARQUITECTURA DE LOS SERVICIOS LTE	121
3.2	DISEÑO DE LA NUEVA SOLUCIÓN	121
3.2.1	DISEÑO DE ALTO NIVEL	122
3.2.2	DISEÑO DE BAJO NIVEL.....	128
3.2.2.1	DISEÑO DE CONEXIONES FÍSICAS	128
3.2.2.2	DISEÑO DEL PROTOCOLO DE ENRUTAMIENTO IS-IS.....	130
3.2.2.3	DISEÑO DEL PROTOCOLO DE ENRUTAMIENTO BGP	131
3.2.2.4	DISEÑO DE LA L3 VPN PARA LOS SERVICIOS LTE	134
3.2.2.5	RESILIENCIA EN LOS SERVICIOS LTE	137

3.2.2.6	CALIDAD DE SERVICIO EN LOS SERVICIOS LTE	145
3.2.2.7	SINCRONISMO EN LOS SERVICIOS LTE	148
3.3	PLAN DE MIGRACIÓN DE FUNCIONALIDADES A LA NUEVA SOLUCIÓN ^[45]	149
3.3.1	PREPARACIÓN DE LA CAPA DE CORE Y ROUTE-REFLECTOR	152
3.3.2	PREPARACIÓN DE LA CAPA DE PRE-AGREGACIÓN.....	160
3.3.3	PREPARACIÓN DE LA CAPA DE ACCESO	162
3.3.4	PREPARACIÓN DE LA MPLS VPN LTE.....	163
3.3.4.1	PREPARACIÓN DE LA MPLS VPN LTE, SERVICIOS S1 Y X2.....	163
3.3.4.2	PREPARACIÓN DE LA MPLS VPN LTE, PLANO DE CONTROL	165
3.3.5	PREPARACIÓN DE LA RED PARA LA SINCRONÍA	169
3.3.5.1	PREPARACIÓN DE LA RED PARA ALTA DISPONIBILIDAD	172
3.3.5.2	LFA FRR CON BFD	172
3.3.5.3	BGP PIC O BGP FRR	175
3.3.6	PREPARACIÓN DE LA RED PARA CALIDAD DE SERVICIO (QOS).....	176
3.4	ANÁLISIS DEL DISEÑO PRESENTADO	184
CAPÍTULO 4	CONCLUSIONES Y RECOMENDACIONES	187
4.1	CONCLUSIONES.....	187
4.2	RECOMENDACIONES.....	194
	REFERENCIAS BIBLIOGRÁFICAS	196

ÍNDICE DE FIGURAS

Figura 1.1 Crecimiento del consumo Global de datos de las redes móviles ^[1]	3
Figura 1.2 Crecimiento de dispositivos móviles y nuevas conexiones ^[1]	4
Figura 1.3 Crecimiento de Tráfico de Video en redes móviles ^[1]	4
Figura 1.4 Crecimiento de conexiones 4G ^[1]	5
Figura 1.5 Evolución de las redes móviles ^[23]	9
Figura 1.6 Arquitectura general de una red UMTS,GSM ^[7]	14
Figura 1.7 Evolución de los sistemas UMTS y GSM en LTE ^[7]	16
Figura 1.8 Diseño de Alto Nivel de LTE ^[7]	16
Figura 1.9 Esquema de un equipo móvil ^[9]	19
Figura 1.10 Arquitectura de la E-UTRAN ^[7]	20
Figura 1.11 Principales componentes del Evolved packet Core ^[7]	21
Figura 1.12 Ejemplos de Red de Acceso movil (Mobile Backhaul) ^[12] y ^[24]	24
Figura 1.13 Red móvil de Interconexión ("Mobile backhaul") ^[12]	25
Figura 1.14 Cabecera MPLS ^[12]	30
Figura 1.15 Funciones de los equipos de MPLS ^[18]	32
Figura 1.16 Operación de la red MPLS ^[18]	35
Figura 1.17 Modelo de MPLS VPN L3 ^[20]	40
Figura 1.18 Ejemplo de configuración de Route Target VPNA y VPNB ^[20]	41
Figura 1.19 Propagación de rutas en una MPLS VPN ^[20]	42
Figura 1.20 Modelos de L2VPN, VPWS y VPLS ^[21]	44
Figura 1.21 Arquitectura del servicio L2VPN ^[21]	45
Figura 1.22 Plano de Control de L2VPN ^[21]	47
Figura 1.23 Plano de datos para un red no MPLS ^[21]	48
Figura 1.24 Plano de Datos L2VPN en redes MPLS ^[21]	49

Figura 1.25 L2VPN para interconectar LTE y EPC ^[5]	53
Figura 1.26 Modelo de MPLS L3VPN para interconectar LTE y EPC ^[5]	56
Figura 1.27 Inter-AS opción A: Conexión para a par de las VRFs en los ASBRs ^[14]	60
Figura 1.28 Opción B: Redistribución de rutas VPNv4 por E-BGP ^[14]	61
Figura 1.29 Opción C: Redistribución de etiquetas VPNv4 entre AS de origen y destino, con redistribución de etiquetas de IPV4 con E-BGP entre sistemas autónomos ^[14]	64
Figura 1.30 Terminología General de una red de Acceso ^[14]	69
Figura 1.31 Ejemplo de servicio implementado en una Red MPLS de CORE Y Red MPLS en el Acceso (Inter-AS) ^[14]	71
Figura 1.32 Ejemplo de un servicio implementado en una Red MPLS unificada o Seamless MPLS ^[14]	72
Figura 1.33 Arquitectura de <i>Seamless</i> MPLS ^[28]	75
Figura 1.34 LSP sin segmentación en el Core y agregación ^[27]	76
Figura 1.35 Segmentación del LSP y acceso mediante etiquetas por BGP ^[27]	77
Figura 1.36 Redistribución de etiquetas de BGP en el IGP o LDP de la red de Acceso (opcional LDP DoD) ^[27]	77
Figura 1.37 Segmentación del IGP de la red de Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP ^[27]	79
Figura 1.38 Segmentación del IGP de la red de Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP ^[27]	80
Figura 1.39 Redistribución de etiquetas de BGP en el IGP o LDP (opcional LDP DoD) en la red de Acceso desde la red de Agregación ^[27]	80

Figura 1.40 Servicios de L3VPN sobre pseudowires terminados en el nodo de transporte TN3, solución Inter-AS ^[14]	82
Figura 1.41 Servicios de L3VPN sobre pseudowires terminados en el nodo de servicio SN, solución Seamless MPLS ^[14]	82
Figura 2.1 Esquema general de la red de <i>backhaul</i> de CNT ^[13]	86
Figura 2.2 Arquitectura física de la red IP-RAN ^[38]	87
Figura 2.3 Equipos de borde de la red IP RAN y la red MPLS existente ^[13]	87
Figura 2.4 Segmentación del dominio IS-IS Core IP RAN varios procesos ^[38]	88
Figura 2.5 ASBR y sus subinterfaces para los procesos de IS-IS ^[38]	88
Figura 2.6 Arquitectura de Route Reflector ^[13]	92
Figura 2.7 Flujo de tráfico en interfaces S1 y X2 de LTE ^[38]	93
Figura 2.8 División de grupos BGP para los RR (Route Reflectors) ^[38]	94
Figura 2.9 Inter-AS Opción B en Core IP RAN ^[13]	95
Figura 2.10 Topología Inter-AS VPN de las Solución de CNT ^[38]	97
Figura 2.11 Balanceo de carga en el tráfico del Inter-AS ^[38]	98
Figura 2.12 Definición de la zona de confianza de QoS en el Core IP RAN ^[13]	99
Figura 2.13 IP Clock mediante 1588v2 ACR ^[13]	104
Figura 2.14 Router modelo ATN950- Huawei ^[33]	107
Figura 2.15 Router modelo CX-600- Huawei ^[35]	108
Figura 2.16 Router modelo ASR 903 ^[36]	110
Figura 3.1 Diseño de alto nivel del Tráfico LTE con L3VPN ^[41]	122
Figura 3.2 Diseño de Alto Nivel de la L3 VPN de servicios LTE ^[41]	124
Figura 3.3 Diseño de Alto Nivel de la L3 VPN de servicios LTE ^[41]	126
Figura 3.4 Diseño de bajo nivel de la IP-RAN de servicios LTE de Quito.....	130
Figura 3.5 Segmentación de dominios IS-IS de la IP-RAN de LTE de Quito ^[27]	132

Figura 3.6 Modelo del diseño de las L3VPN de LTE ^[27]	137
Figura 3.7 Tráfico X2 en caso de regiones contiguas L3VPN de LTE ^[27]	138
Figura 3.8 Funcionalidad de la RFC 5286 LFA FRR ^[43]	139
Figura 3.9 Funcionalidad de la RFC 7490 Remoto LFA FRR ^[43]	140
Figura 3.10 Funcionalidad de la RFC 7490 Remoto LFA FRR ^[43]	143
Figura 3.11 Esquema de mecanismo de resiliencia ^[27]	144
Figura 3.12 Modelo de Calidad de servicio del tráfico de subida ^[30]	147
Figura 3.13 Modelo de Calidad de servicio del tráfico de bajada ^[30]	149
Figura 3.14 Esquema de sincronismo ^[27]	150
Figura 3.15 Modelo de Topología de Migración del proyecto de reingeniería de la IP-RAN	151
Figura 3.16 Sesiones iBGP del Route-Reflector del Core	159
Figura 3.17 VRF movdcn de servicio y movoyrn_trust de gestión con su RTs	163
Figura 3.18 Sesiones BGP VPNV4, plano de control de la VRF movdcn	169
Figura 3.19 Esquema de calidad de servicio QoS ^[45]	177

ÍNDICE DE TABLAS

Tabla 1.1 Generaciones de Comunicaciones Celulares ^[8]	12
Tabla 1.2 Comparación de transporte de <i>Switch</i> Ethernet con redes de acceso MPLS ^[14]	69
Tabla 1.3 Modelos de red de transporte con <i>Seamless</i> MPLS ^[28]	74
Tabla 2.1 Definición del tercer octeto en la interface loopback 100 ^[13]	89
Tabla 2.2 Formato de la dirección NET del protocolo IS-IS ^[13]	89

Tabla 2.3 Representación del convenio para la dirección NET en los equipos de core [13]	90
Tabla 2.4 Representación del convenio para las direcciones NET en los equipos core IP-RAN [13]	90
Tabla 2.5 LDP vs RSVP-TE	91
Tabla 2.6 MPLS VPN para gestión y servicios LTE [13]	94
Tabla 2.7 Requerimientos de QoS en LTE [13]	100
Tabla 2.8 Correspondencia DSCP , CoS y EXP [13]	100
Tabla 2.9 Técnica de encolamiento por clases [13]	100
Tabla 2.10 Limites para WRED [13]	102
Tabla 2.11 Número de nodos de acceso por cada nodo de Agregación [39]	106
Tabla 2.12 Especificaciones técnicas del router ATN-950B correspondiente a la tarjeta Controladora instalada [33]	108
Tabla 2.13 Especificaciones técnicas del router CX-600 [35]	109
Tabla 2.14 Anuncio de soporte de la versión CX-600 V600R006	109
Tabla 2.15 Especificaciones técnicas del router ASR-903 [36]	110
Tabla 2.16 Anuncio de futuro soporte del router ASR-903 de Cisco [37]	111
Tabla 3.1 Distribución de los nodos de pre-agregación en tres zonas de agregación	129
Tabla 3.2 Comunidades BGP para optimizar la publicación de las interfaces <i>loopback</i>	133
Tabla 3.3 RT de la VRF de servicios.....	134
Tabla 3.4 Valores de temporizadores para BFD	142
Tabla 3.5 Modelo de QoS para el servicio LTE [27]	146

CAPÍTULO 1

MARCO TEÓRICO

1.1 ANTECEDENTES

En el crecimiento de la empresa de telecomunicaciones CNT E.P (Corporación Nacional de Telecomunicaciones Empresa Pública) la red MPLS (*Multiprotocol Label Switching*) ha sido una respuesta escalable y eficiente, que ha permitido brindar múltiples servicios, como voz sobre IP, Internet, datos, telefonía fija y móvil. La CNT E.P ha adquirido una red de acceso que agrupa en IP (*Internet Protocol*) a las radio bases del servicio de datos móviles LTE (*Long Term Evolution*), a esta infraestructura se la conoce como *backhaul IP* de radio bases, por sus siglas en inglés IP-RAN (*Internet Protocol Radio Access Network*), esta red tiene el objetivo de extender los servicios de banda ancha donde no se disponga de cobertura de redes fijas como ADSL (*Asymmetric Digital Subscriber Line*) o GPON (*Gigabit-capable Passive Optical Network*). La reciente adquisición se ha interconectado con la red MPLS, observándose en la actualidad algunos inconvenientes de redundancia y futura escalabilidad cuando haya superado el incremento de clientes esperados. El presente trabajo presenta un plan de acciones a considerar para fortalecer la nueva red adquirida aprovechando las cualidades de toda la red MPLS y, de esta manera asegurar la escalabilidad y calidad de servicio para los clientes de la CNT E.P.

1.1.1 CRECIMIENTO DEL CONSUMO DE INTERNET MÓVIL

En la industria de servicios móviles se ha podido confirmar¹ un amplio crecimiento en el consumo de servicios móviles, tanto para datos como para distribuir la conectividad a Internet, con mejoras tecnológicas en los equipos terminales, tales como el **iPhone y iPad** de Apple y sus competidores con sistema **Android** han determinado una creciente demanda de ancho de banda, para tener una medida de este crecimiento y también una proyección del mismo se ha tomado los resultados de un estudio, que anualmente lo actualiza la empresa Cisco Systems, denominado *Cisco Visual Networking Index: Global Mobile Data Traffic Forecast Update, 2013–2018*, documento del cual se extrae los siguientes datos:

- ✓ **El tráfico global de datos móviles presenta un crecimiento del 81% en el 2013**, ya que se alcanzó 1.5 exabytes por mes al término del 2013, con respecto 820 peta bytes por mes al término del 2012, como se indica en la figura 1.1
- ✓ **El año 2012 el tráfico de datos móviles fue 18 veces lo que fue el tráfico global de internet en el año 2000**. Un exabyte de tráfico se medía en el año 2000 y en el 2013 18 exabytes corresponden al tráfico de datos móviles, como se puede observar en la figura 1.1.
- ✓ **El tráfico de Video de datos móviles excedió el 50% en el 2012 por primera vez**. Convirtiéndose en un 53 % el tráfico de video en datos móviles al final del 2013, como se muestra en la figura 1.3.

¹ Datos tomados de la referencia [1].

- ✓ **Sobre medio millón de equipos móviles y nuevas conexiones fueron incluidas en el 2013.** Los equipos móviles globales y nuevas conexiones crecieron en 7 billones con respecto a 6.5 billones en el año 2012, como se puede confirmar en la figura 1.2. El crecimiento de los equipos como el smartphone llegó a 77% de este crecimiento de equipos móviles, con 406 millones de nuevas redes adicionales en el 2013.
- ✓ **En el 2013 las conexiones 4G han generado tráfico 14.5 veces mas en promedio que las que nos son 4G (3G y todavía 2.5G).** Aunque las conexiones 4G representan un 2.9 % de las conexiones móviles al 2013 este consumo representa el 30% del tráfico, como se muestra en la figura 1.4.

Luego de confirmar con estos estudios ^[1] la importancia que adquiere el tráfico de datos móvil en los siguientes cinco años, lo que representa una de las razones fundamentales para potenciar el crecimiento escalable de las redes móviles ya instaladas.

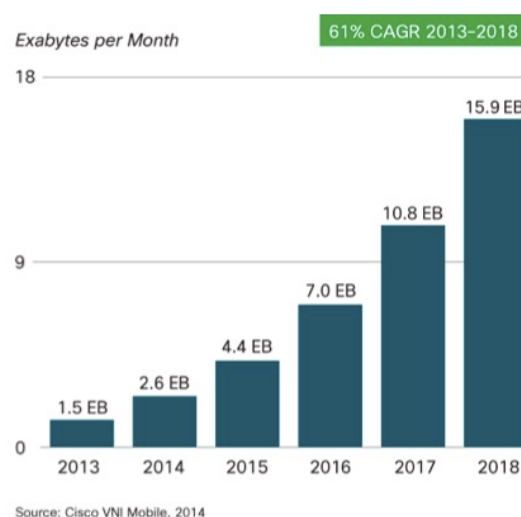


Figura 1.1 Crecimiento del consumo Global de datos de las redes móviles ^[1]

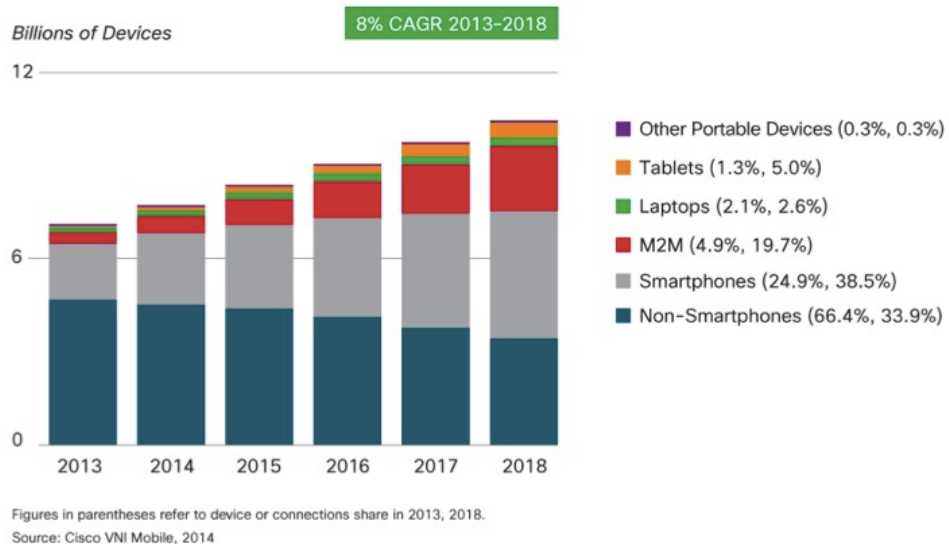


Figura 1.2 Crecimiento de dispositivos móviles y nuevas conexiones ^[1]

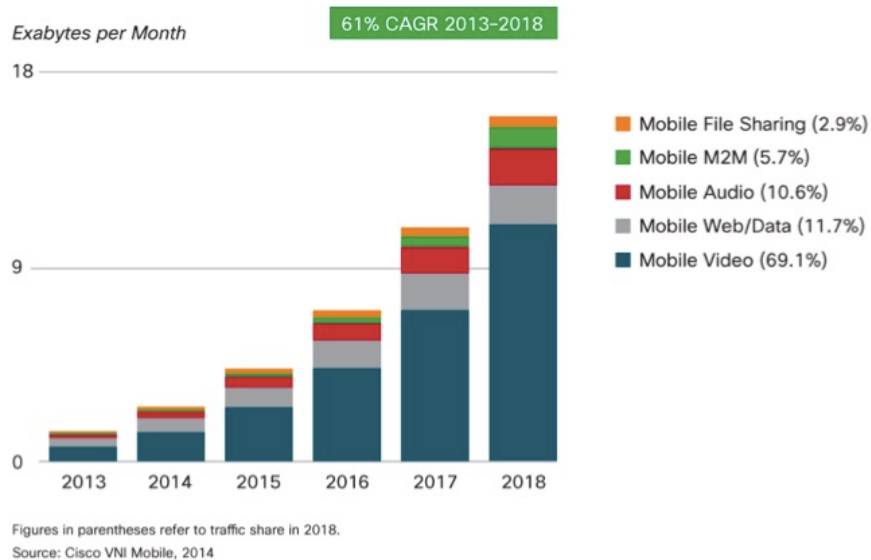


Figura 1.3 Crecimiento de Tráfico de Video en redes móviles

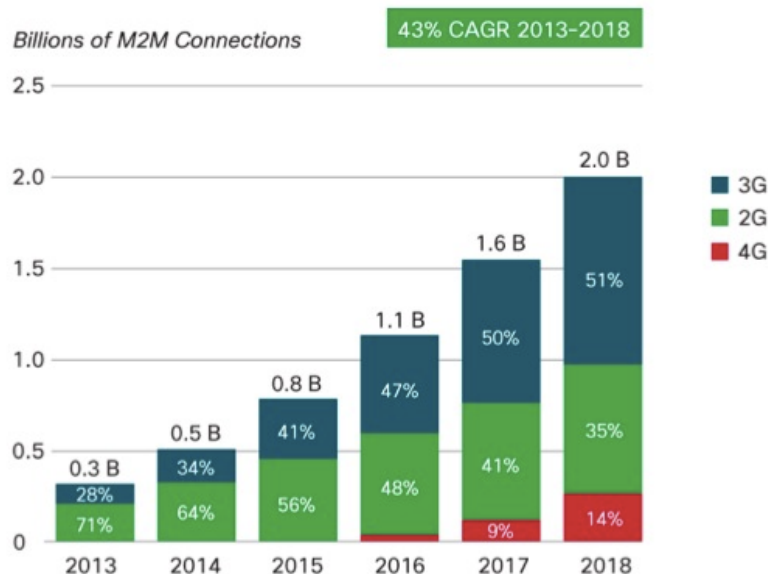


Figura 1.4 Crecimiento de conexiones 4G ^[1]

1.1.2 NUEVOS REQUERIMIENTOS DE LAS REDES DE TRANSPORTE PARA LAS REDES MÓVILES DE UN PROVEEDOR DE SERVICIOS

En el el servicio de datos móviles, las radio bases han presentado en el tiempo cambios en los tipos de tecnología para interconectarse con el Core de servicios, en las primeras generaciones con servicios de datos móviles como 2G y 3G ^[5], este transporte operaba tradicionalmente con enlaces TDM (SDH) o ATM, en las proyecciones previstas que se han revisado en el punto anterior ^[1] y confirmando la permanente evolución de las redes de datos, se observan varios tipos de soluciones, a este tipo de servicios de interconexión entre las radio bases y el Core de servicios móviles se denomina *backhaul*, para los casos en que esta conectividad comprende

enlaces IP se utiliza el término *Backhaul* IP-RAN y consiste en la esencia del presente estudio.

En esta sección se ha querido mencionar las principales características que se deben considerar para este nuevo tipo de interconexión, que en el constante crecimiento de la demanda esperada constituye uno de los medios que permitirá un adecuado aprovechamiento de las nuevas tecnologías de acceso a implementarse, como es el caso de redes LTE y LTE-Advanced.

A continuación se han revisado las siguientes características tomadas de las referencias [3] y [4], que se pueden considerar como buenas prácticas de la industria:

a. Escalabilidad para un crecimiento de 10 a 100 veces mas que el tráfico actual

En cuanto a escalabilidad existen algunos aspectos a tomar en cuenta, los proveedores de servicios deben como buena práctica prever su crecimiento de tráfico entre 1 y 100 veces su tráfico actual tomando en cuenta su demanda, ya que la mayor parte de operadoras o proveedores de servicios móviles diseñan sus redes dependiendo también de la densidad de abonados, es decir la cobertura en la zona rural es diferente a la que existe en la zona urbana. Otro aspecto a tomar en cuenta es que un diseño escalable de la red IP-RAN permitirá mantener la topología instalada y que su crecimiento en el acceso de radio no sea producto de cambios radicales, sino mas bien el reemplazo de los equipos de radio de la capa final de acceso; sin duda que estos cambios a futuro deben de ir de la mano de su respectivo estudio de retorno de inversión.

b. Simplificación y unificación de la operación

Debido que los operadores de redes móviles en su red de acceso de radio (RAN, *Radio Access Network*) presentan constantemente innovaciones o cambios tecnológicos, como el paso de una red HSPA+ (3G) hacia LTE o LTE – Advanced, la red de *backhaul* IP-RAN debe permitir unificar estas tecnologías y simplificar su migración, ésta es una de las principales ventajas sobre una red de *backhaul* tradicional TDM o ATM, especialmente cuando en las redes se tiene ya presente la implementación de LTE que opera completamente en IP.

c. Alta disponibilidad y calidad de servicio

En el mercado un proveedor de servicios móviles recibe diferentes requerimientos de calidad de servicio, dependiendo inclusive de la capacidad adquisitiva del cliente, por lo cual este operador móvil debe preparar su red para presentar diferentes soluciones a las diferentes necesidades de priorización de tráfico de sus clientes y cumplir con su promesa de valor, que generalmente se firma en su contrato, a esto se lo suele conocer con el nombre de SLA (*Service Level Agreement*). Tomando en cuenta estas premisas un *backhaul* de IP RAN debe permitir que el tráfico prioritario esté protegido tanto en las épocas de sobresuscripción como en la presencia de fallas de interconexión física y entregar el servicio de acuerdo al SLA comprometido con el cliente final.

d. Operación simultánea con tecnologías 2G, 3G y 4G de múltiples proveedores de acceso

Debido a la constante evolución de las redes de acceso inalámbricas, un operador está de la misma manera instalando estos nuevos componentes acorde a como se presentan nuevos desarrollos y es muy común el tener varias marcas de estos

equipos de radio; es por esto que un *backhaul* de IP-RAN debe permitir la integración de estas diversas tecnologías hacia un mismo Core de servicios. Gracias a que en la industria el protocolo IP es considerado cada vez más como parte de estos equipamientos, es en si una gran ventaja ya que es precisamente lo que permite esta integración, además esta facilidad adquiere mayor importancia en una migración de tecnología de 3G a LTE, ya que estas dos tecnologías trabajan sin problemas con *backhaul* IP.

e. La plataforma debería soportar tráfico multimedia

Nuevamente volviendo a los estudios de proyecciones de tráfico ^[1], como se pueden ver en la figura 1.2 y figura 1.3, las proyecciones del comportamiento del tráfico añadido a la evolución del LTE implica un crecimiento de los requerimientos de ancho de banda y del tipo de tráfico multimedia, especialmente el componente de video se hace cada vez mas importante frente al resto de tráfico de datos móviles. Acorde a lo mencionado la red de *backhaul* IP-RAN debe permitir el paso prioritario de este tipo de tráfico sin disminuir la promesa de valor de otros servicios considerados en los SLAs de cada cliente.

1.1.3 EVOLUCIÓN DE LAS REDES MÓVILES

En las últimas décadas la tecnología de datos por la red celular ha experimentado grandes crecimientos, en esta primera parte se desea mostrar un breve resumen de dicha evolución, para comprender como el requerimiento del ancho de banda del cliente final ha sido la principal razón de mejoramiento y constatación de evolución de estas redes.

En la figura 1.5, se tiene un resumen de cómo ha ido evolucionando las tecnologías y a la par cómo se ha tenido un incremento progresivo también del ancho de banda disponible para los usuarios finales.

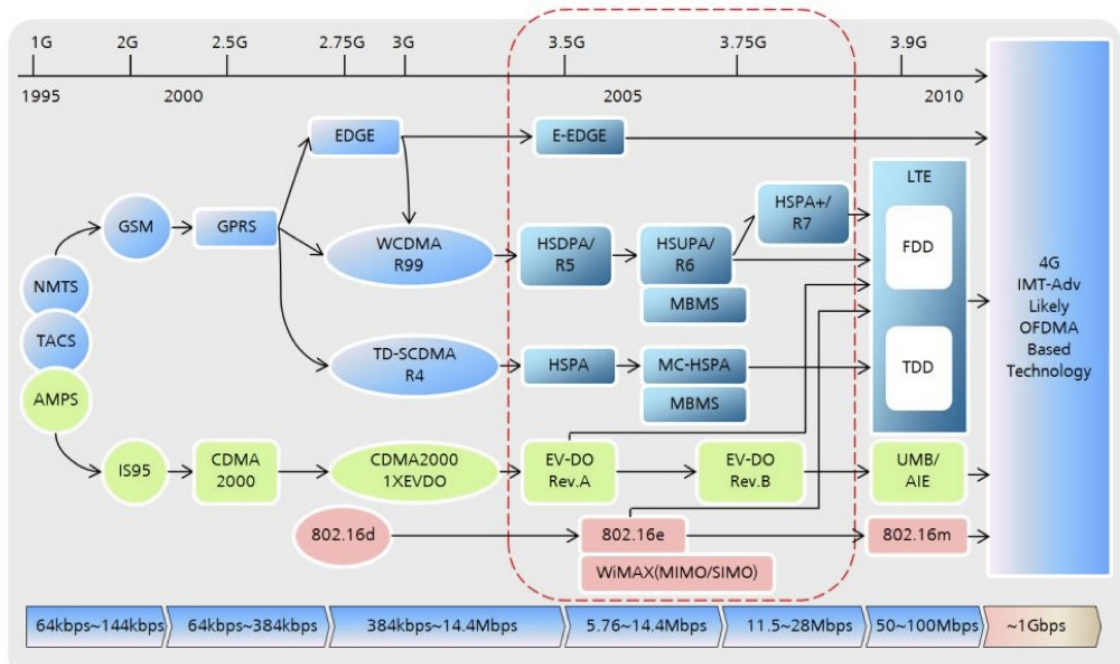


Figura 1.5 Evolución de las redes móviles ^[23]

A nivel comercial¹ en los años 1980s empiezan los primeros sistemas análogos de telefonía celular, denominados ahora como **1G**, es decir de primera generación, a cargo de empresas como *Nordic Mobile Telephone* (NMT) en Saudi Arabia, C-Netz en Alemania, Portugal y South África, *Total Access Communications System* (TACS) en el Reino Unido y el sistema conocido como AMPS (*Analog Advanced Mobile Phone System*) en todo el continente Americano. Este sistema fue conocido como analógico ya que no utilizó tecnología digital.

¹ Este resumen fue tomado de la parte “1.2 History of Mobile Telecommunications Systems” de la referencia [7]

En los años 1990s , introduciendo tecnología digital empieza lo que se conoce como 2G, la misma que consigue mejoras en la calidad de transmisión de voz, capacidad soportada así como también la nuevas aplicaciones como el SMS (*Short Message Service*). La CEPT (*The European Conference of Postal and Telecommunications Administrations*) decide a partir de 1982 crear un proyecto de segunda generación de comunicación móvil que se denominó 2G, esto fue el comienzo del conocido GSM (*Global System for Mobile Communications*) estándar que fue implementado en el año 1991. El principal objetivo de GSM fue el soporte de telefonía móvil y se incorpora el servicio de roaming¹, este estándar se basa en un mix entre TDMA (*Time Division Multiple Access*) y FDMA (*Frequency Division Multiple Access*) en contraste con los sistemas 1G fundamentados únicamente en FDMA. En paralelo a GSM otros estándares de 2G fueron desarrollados como IS-136, mejor conocido como D-AMPS (*Digital-AMPS*), IS-95A también conocido como CDMAOne (el primer *Code Division Multiple Access*) usado principalmente en el continente Americano y finalmente se conoce otro sistema de 2G denominado PDC (*Personal Digital Cellular*) empleado exclusivamente en el Japón.

Con la introducción de servicios de datos (*packet-switched services*) sobre la red de voz 2G y realizando una modificación en la interfaz inalámbrica para que se pueda intercambiar tanto datos como voz, a ésta evolución se denominó 2.5G. Las técnicas y conceptos que permitieron ésta evolución son el GPRS (*General Packet Radio Services*), que es una extensión del sistema GSM y el IS-95B que vino del IS-95A o

¹ En telefonía móvil la itinerancia o roaming es la capacidad de enviar y recibir llamadas en redes móviles fuera del área de servicio local de la propia compañía, es decir, dentro de la zona de servicio de otra empresa del mismo país, o bien durante una estancia en otro país diferente, con la red de una empresa extranjera.

CDMAOne. Al mismo tiempo que se incrementaba el tráfico de datos progresivamente, la industria mejoró la interfaz de radio con técnicas como EDGE (*Enhanced Data Rates for GSM Evolution*) y empezó a presentar las primeras alternativas para llegar a sistemas de tercera generación 3G.

Luego de entrar en operación los sistemas 2G, los miembros de la industria de telefonía celular se reunieron para discutir cuál son los próximos objetivos para un futuro desarrollo. En enero de 1998, la tecnología CDMA tuvo dos variantes, el WCDMA (*Wideband CDMA*) y la TD-SCDMA (*Time Division Synchronous Code Division Multiple Access*) que fueron adoptados por la ETSI (*European Telecommunications Standards Institute*) como estándares de UMTS (*Universal Mobile Telecommunication System*) para la tercera generación de comunicaciones móviles o 3G, también denominado IMT-2000 (*International Mobile Telecommunications 2000*). Estos dos nuevos estándares fueron utilizados por la 3GPP (*Third Generation Partnership Project*) para desarrollar su proyecto de tercera generación denominado UMTS que benefició a los países que venían trabajando con GSM, mientras que por otro lado el CDMA-2000 podría ser implementado donde se venía trabajando con IS-95.

Se desarrollaron nuevas especificaciones en el marco del grupo 3GPP conocidos como la Evolución de la 3G, se propusieron dos nuevos tipo de redes de acceso de radio; el primero que se denominó HSPA (*High Speed Packet Access*) que se lo calificó como 3.5G. HSPA se desarrolló en una primera fase conocida como *Release 5* como la tecnología HSDPA (*High Speed Downlink Packet Access*) y HSUPA (*High Speed Uplink Packet Access*) agregado en una segunda fase conocida como *Release 6* (Rel- 6) del proyecto UMTS; en estos estándares se consiguieron velocidades de 14.6 Mbps de *Downstream* y 5.76 Mbps de *Upstream*

respectivamente. Paralelamente en los sistemas 3G también se presentan alternativas como 1xEV-DO (2003) and 1xEV-DV (2006) que son el resultado del desarrollo del CDMA 2000.

	1G	2G/2.5G	3G	4G
Inicio/ Implementación	1970/ 1984	1980/ 1999	1990/ 2002	2000/ 2006
Tasa de Transmisión	2 kbps	14.4-64kbps	2Mbps	200Mbps hasta 1 Gbps para baja movilidad
Estándares usados	AMPS	2G:TDMA,CDMA, GSM 2.5G: GPRS, EDGE, 1xRTT	WCDMA, CDMA-2000	LTE-A
Tecnología	Analógica	Digital	Banda ancha CDMA, tecnología IP	Tecnología IP en combinación de Banda ancha, LAN/WAN/PAN y WLAN
Servicio	Telefonía Móvil (Voz)	2G:Voz (Digital) y SMS (mensajes) 2.5G: Alta Capacidad de Datos	Integración de Datos, Voz y Video	Información Dinámica Acceso de información de cualquier tipo
Multiplexación	FDMA	TDMA, CDMA	CDMA	CDMA
Conmutación de Datos	Conmutación de Circuitos	2G: Conmutación de Circuitos 2.5G: Conmutación de Circuitos para la red de acceso; Conmutación de Paquetes para el Core de datos	Conmutación de Paquetes, Excepto para la interfaz de aire (Conmutación de circuitos)	Todo es Conmutación de Paquetes
Core usado	PSTN	PSTN	Red de Paquetes	Internet
Handoff	Horizontal	Horizontal	Horizontal	Horizontal y Vertical

Tabla 1.1 Generaciones de Comunicaciones Celulares ^[8]

En el siguiente paso en la evolución de los sistemas UMTS se encuentran dos trabajos fundamentales realizados por la 3GPP, en cuanto al acceso de radio de los usuarios se denominó LTE (*Long Term Evolution*) y el Core de servicios también tuvo importantes modificaciones, a lo que se le nombró EPC (*Evolved Packet Core*)

dentro de los trabajos denominados SAE (*System Architecture Evolution*). El objetivo principal de LTE es tener un mejor desempeño (mas ancho de banda) a menor costo en el acceso, este mejor desempeño tiene como objetivos adicionales: una arquitectura plana basada en el protocolo IP para reducción de costos, incrementar la eficiencia de uso del espectro radioeléctrico que es un recurso limitado, incrementar la capacidad de transmisión del usuario final atado a cada celda de la red de acceso (*throughput*). En un principio con LTE se consiguieron velocidades de hasta 326 Mbps y ya con su evolución en LTE-Advanced o IMT-Advanced se obtendrá teóricamente hasta 1.6 Gbps por usuario.

En la tabla 1.1, se ha resumido brevemente las características de la evolución de la tecnología celular, donde los aspectos con mayor relevancia han sido el incremento de ancho de banda hacia el usuario final y el cambio de la conmutación de circuitos hacia la conmutación de paquetes, característica de la tendencia actual de las redes de datos, esto como énfasis en que la red que une a la parte de acceso de las celdas con el Core de datos también es en la 4G totalmente protocolo IP.

1.1.4 CONCEPTOS BÁSICOS DE REDES MÓVILES LTE

El proyecto de LTE fue diseñado en un ambiente de colaboración nacional y regional de los equipos que aportan para la realización de estándares de la industria, este equipo se denominó 3GPP (*Third Generation Partnership Project*). LTE fue la evolución de un sistema conocido como UMTS (*Universal Mobile Telecommunications*), el cual ha evolucionado del sistema conocido como GSM (*Global System for Mobile Communications*). Considerando este antecedente se hace importante mencionar conceptos de UMTS para entrar en contexto y explicar de mejor manera la funcionalidad de la Arquitectura de LTE.

1.1.4.1 Arquitectura de las redes móviles LTE como una evolución de las tecnologías UMTS y GSM ^[7]

Una red pública celular móvil con las que se encuentran actualmente operando en nuestro País, en cuanto a UMTS (3G) está constituida, como se indica en la Figura 1.6, por un equipo de usuario o teléfono celular, una red de acceso de radio móvil y una red de Core, que permite la interconexión con los sistemas de servicio.

La red que se denomina Core, está formada por dos dominios, uno que permite la interconexión con los circuitos de telefonía tanto fija, PSTN (*Public switched Telephone Network*) como móvil de otros proveedores, a este dominio se lo conoce como dominio de conmutación de circuitos o CS (*Circuit Switched*); por otro lado se tiene al dominio de conmutación de paquetes que se denomina PS (*Packet Switched*), para la transmisión de datos entre el equipo de usuario y una red externa de paquetes de datos o PDNs (*Packet Data Networks*), por ejemplo Internet.

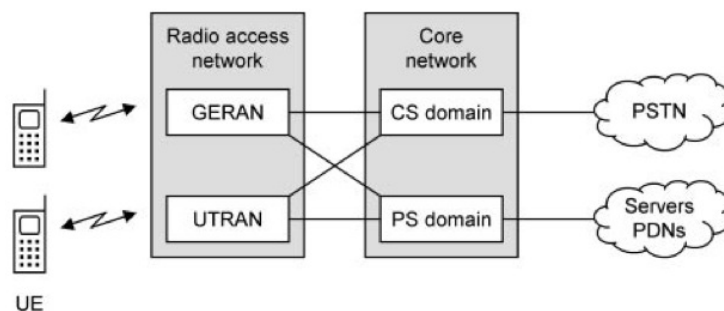


Figura 1.6 Arquitectura general de una red UMTS,GSM^[7]

En la red de acceso de radio se puede indicar que su propósito consiste en interconectar al usuario con el Core de servicios, para el caso de las redes GSM este acceso se conoce como GERAN (*GSM EDGE Radio Access Network*) y en las redes

UMTS se conoce como UTRAN (*UMTS Terrestrial Radio Access Network*). Estas dos tecnologías incluso pueden coexistir compartiendo el mismo Core de servicios.

En el 2004 empezaron las investigaciones dentro del 3GPP para lo que se conoce como LTE (*Long Term Evolution of UMTS*), en la figura 1.7 se puede observar cómo el servicio del usuario tiene dos partes fundamentales, la primera que es Core de servicios y la otra el acceso de radio, 3GPP se concentró por esto en dos grandes proyectos que se impulsaron, uno de ellos dedicado a la red de acceso de radio, lo que se conoce ahora como E-UTRAN (*Evolved UMTS Terrestrial Radio Access Network*) donde se ha conseguido mejorar notablemente el ancho de banda hacia el usuario, a este desarrollo de radio se lo ha denominado LTE; por otro lado se tiene el proyecto de evolución de la red de Core de servicios a lo que se denominó SAE (*System Architecture Evolution*) ahora todas las funciones son a través de una red conmutación de paquetes, es decir el EPC (*Evolved Packet Core*), en este nuevo sistema no se considera el servicio de voz tradicional en conmutación de circuitos. El proyecto integral de evolución de UMTS se lo conoce como EPS (*Evolved Packet System*), en el argot común se le conoce a todo el proyecto por el nombre del proyecto de radio acceso, es decir LTE.

1.1.4.2 Arquitectura de LTE [7]

Para aclarar conceptos acerca de la arquitectura de la red LTE se han tomado en cuenta dos aspectos, el diseño a alto nivel que nos proporciona una visión global y como el usuario se encuentra interconectado a los servicios que brinda el Core de datos

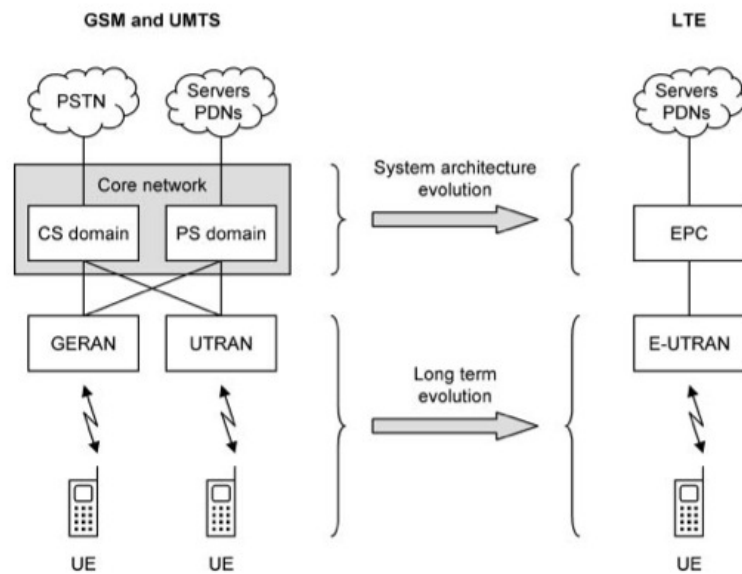


Figura 1.7 Evolución de los sistemas UMTS y GSM en LTE ^[7]

1.1.4.2.1 Diseño de alto nivel de LTE [7]

A continuación se presenta el diseño de alto nivel o lo que se conoce como HLD (*High Level Design*) de la arquitectura de las redes LTE ^[7], con el objetivo de describir la importancia de la interconexión entre el acceso de radio y el Core o centro de servicios para las redes LTE.

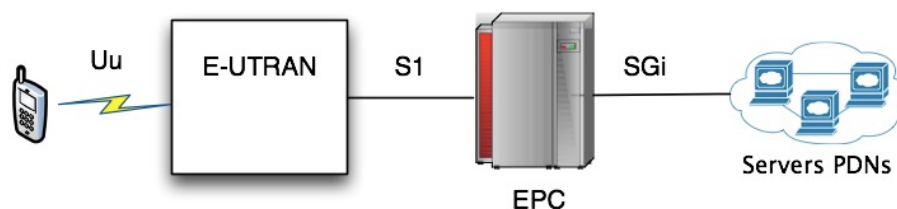


Figura 1.8 Diseño de Alto Nivel de LTE ^[7]

En la figura 1.8, se tiene la representación del HLD de LTE, aquí se puede apreciar a grandes rasgos el funcionamiento en general del sistema de comunicación. Se tiene entonces el equipo de usuario UE (*User Equipment*), el E-UTRAN (*Evolved UMTS Terrestrial Access Network*) y el EPC (*Evolved Packet Core*). Las interfaces de comunicación del EPC se denotan de la siguiente manera:

Uu: entre el equipo de usuario y la red de acceso móvil.

S1: transmisión entre las radio bases y el EPC.

SGi: comunicación entre el EPC y la red de servicios.

Para entregar el servicio S1, en la industria se han presentado varias soluciones tecnológicas, las que comprenden una red de transporte que se ha denominado en la práctica como *mobile backhaul*, cuando la principal tecnología involucrada en la solución es el protocolo IP a esta red de interconexión se le ha denominado IP-RAN (*IP Network Radio Access Network*). Es en esta red donde se concentra el estudio del presente trabajo, es decir se realizará el análisis de la solución IP-RAN de CNT y alternativas de optimización.

El EPC ahora transporta paquetes desde la E-UTRAN hacia la red de datos de servicios o PDNs (*Packet Data Networks*) a través de protocolos IP, utilizando para esto IP4, IPV6 o *dual stack*¹, una diferencia con GSM y UMTS es que ahora el Core mantiene una conexión permanente con la red de servicios y no una conexión temporal presente hasta la detección de falta de actividad desde el usuario.

¹ “dual stack” es una propiedad de software y hardware de los equipos de comunicaciones en una red de datos, la cual permite trabajar al mismo tiempo con paquetes de IPV4 y de IPV6 por las mismas interfaces en los mismos equipos.

El EPC^[7] se comporta como un canal de datos que transporta datos desde el usuario hacia sus correspondientes aplicaciones de forma transparente, esto hace que LTE se diferencie de los otros sistemas móviles, donde la información de voz era una parte integral y tenía un tratamiento directo, para LTE los paquetes de voz son manejados por una entidad externa denominada IMS (*IP Multimedia Subsystem*) y estos paquetes son transportados de la misma manera que cualquier otro tipo de información generada por el usuario. Esta forma de transmisión es similar a Internet, con la diferencia que el EPC contiene mecanismos específicos de control sobre la tasa de transmisión, la tasa de error y el retardo de un flujo de datos. No existe una definición específica del retardo de un paquete que pase por el EPC, pero existen recomendaciones de que se mantenga menor a 10 ms en esquemas que no tiene roaming y 50 ms con roaming¹

1.1.4.2.2 Interconexión del usuario móvil y los servicios LTE [7]

Es necesario revisar a un nivel general cómo se da la comunicación del usuario móvil, a través de la E-UTRAN para obtener los servicios que se encuentran centralizados en lo que se denomina EPC, esto pondrá en evidencia ciertas características de la red de *backhaul* que se quiere analizar en el presente trabajo. Como se puede ver en la figura 1.9, un equipo de usuario tiene tres componentes funcionales^[9] , un terminal móvil o MT (*Mobil Terminal*), encargado de la comunicación de radio con la E-UTRAN, un terminal que presenta los datos al usuario o TE (*terminal equipment*) y

¹ Este roaming^[7], se refiere al servicio que recibe un usuario a través de la red de otro proveedor.

una tarjeta denominada UICC (*Universal Integrated Circuit Card*). En estos circuitos se encuentra un programa que almacena características propias del usuario, como su número telefónico y la empresa operadora que le provee el servicio. En los primeros desarrollos de GSM se denominó SIM (*Subscriber Identity Module*) y en UMTS se denominó USIM (*Universal Subscriber Identity Module*); por estas razones este componente se lo conoce en el argot popular técnico como SIM card, cabe a notar que para el LTE las primeras versiones de esta tarjeta no son compatibles.

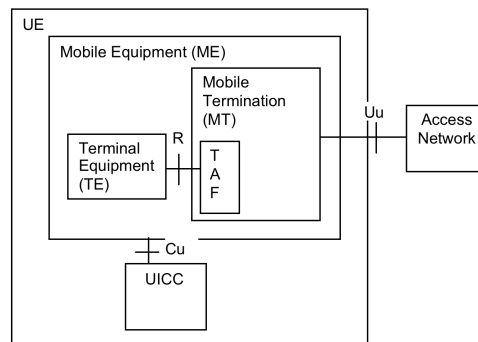


Figura 1.9 Esquema de un equipo móvil^[9]

El UE dependiendo de su marca y modelo se interconecta a la red de acceso de radio E-UTRAN con una variedad de capacidades de radio, las cuales le permitirán al usuario tener mas o menos ancho de banda para datos. Como se puede apreciar en la figura 1.10, la E-UTRAN es la encargada de la conexión entre el usuario móvil UE y el eNB (*Evolved Node B*), además es responsable por la conexión desde el eNB hasta el Core de servicios EPC. Se conoce como Uu al enlace entre el UE y el eNB, los enlaces hacia el EPC se conoce como S1, además la E-UTRAN se encarga de manejar apropiadamente la operación con radio bases adjuntas u otros eNB para lo que es el *handover*, a esta conexión se la ha nombrado como X2, su utilización no

es obligatoria y puede pasar su función a la conexión S1, pero con más retardo que cuando es directa, la X2 es una comunicación directa entre nodos cercanos.

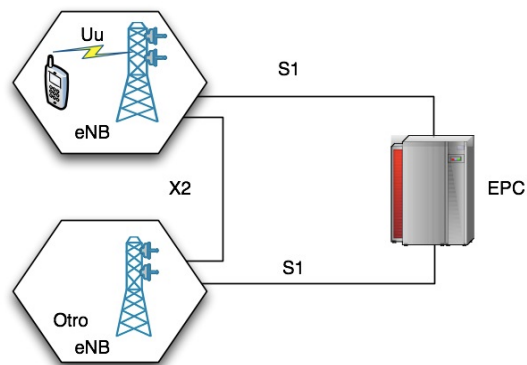


Figura 1.10 Arquitectura de la E-UTRAN ^[1]

A continuación se exponen los principales servicios que se encuentran en el EPC, como se puede apreciar en la figura 1.11, se tiene la base de datos que contiene toda la información del cliente, a la que se la denomina HSS (*Home Subscriber Server*) esta base se puede considerar la misma que se tiene en UMTS y en GSM.

El P-GW (*Packet Data Network Gateway*) constituye el punto de contacto con el mundo exterior de datos o internet, a través de una interfaz denominada SGi, cada paquete de datos es identificado por un nombre de punto de acceso o APN (*Access Point Name*), cada proveedor trabaja con varios números de APNs, por ejemplo un APN para sus datos internos y otro para internet. Un usuario móvil es asignado un APN por defecto, cuando enciende su equipo, esto se lo hace con el fin de tener siempre conectividad hacia fuera, como internet por ejemplo y si luego desea acceder a otra red de datos privados se le asigna otro APN, todos los APNs se mantienen constantes durante la conexión de datos.

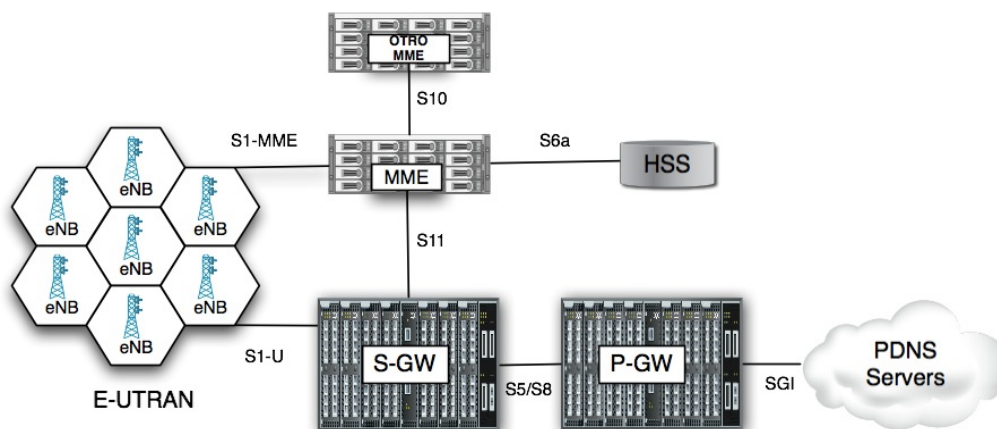


Figura 1.11 Principales componentes del Evolved packet Core ^[7]

El S-GW (*Serving Gateway*) actúa como un router llevando la data que proviene de las radio bases con sus respectivo PDN gateway, las operadoras suelen tener un grupo de S-GW en varios sitios de la red y los usuarios móviles se encuentran distribuidos en cada uno de ellos, es posible que un usuario cambie de S-GW si se ubica lo suficiente lejos para ser registrado por otro S-GW.

El control del servicio móvil se lo hace en una capa operativa superior denominada MME (*Mobility Management Entity*), donde se encuentran elementos como la seguridad, la distribución de los flujos de datos que no tienen nada que ver con la data propia del radio enlace. De igual forma que los S-GW, una operadora puede disponer de un grupo de MMEs distribuidos geográficamente que poseen la información de control de una parte de los usuarios y un usuario cambiaría de MME solamente en el caso de alejarse lo suficiente del sitio donde se encuentra.

El EPC tiene otros componentes que no han sido representados en la figura 1.11. Como por ejemplo el CBC (*Cell Broadcast Centre*) y el CBS (*Cell Broadcast Service*).

1.2 BENEFICIOS DE UNA IMPLEMENTACIÓN IP PARA EL ACCESO DE RADIO DE TECNOLOGÍA LTE, IP-RAN

Para las redes GSM y UMTS, el proporcionar la conectividad desde las radio bases hasta las RNC (*Radio Network Controller*) se convirtió en una necesidad cotidiana al momento de realizar la instalación de una nueva celda celular y con ello expandir la cobertura de nuevos clientes, cuanto mas ahora las redes LTE que tienen que observar niveles de servicio más exigentes especialmente para cumplir con las características de mayor ancho de banda, por lo que el requerimiento de un *backhaul* apropiado adquiere más importancia y debe mantener características técnicas adecuadas para alcanzar las metas propuestas.

1.2.1 GENERALIDADES DE UNA RED DE ACCESO DE RADIO

Como se indica en la figura 1.12, se puede observar que entre el sitio central de servicios móviles y las celdas celulares existe un tipo de interconexión, que siendo parte de una red de transporte se debe acoplar a las necesidades de los servicios móviles para brindar los servicios de señalización, voz y datos. Por lo que se puede considerar que este tipo de interconexión, se puede tratar en si como una red transporte especializada, a la cual se la conoce como "*Mobile Backhaul*" y en el caso de que la red de transporte sea por medio de una red de conmutación de paquetes IP, se le conoce como IP-RAN.

Desde que inició el servicio móvil existió la necesidad de un “*Mobile backhaul*” que permita distribuir los servicios desde un sitio central , durante este tiempo ha existido una evolución, acorde a las posibilidades que ofrecía en cada momento la tecnología de telecomunicaciones, por lo que se han dado varias soluciones como son TDM, ATM o IP sobre varios medios físicos de líneas dedicadas eléctricas y ópticas o con sistemas de radio. Como se expuso en la sección 1.1.1, el tráfico del servicio celular ha ido incrementándose, lo que implica el aumento de mas celdas y un aumento de la cantidad de enlaces de *backhaul*, a propósito de este comportamiento la tecnología de conmutación de paquetes permite disminuir los costos haciendo más eficiente ésta interconexión y de allí la importancia de la denominada IP-RAN.

1.2.2 LAYER 3 VS LAYER 2 (EJEMPLOS DE CONEXIONES DE BACKHAUL)

Las soluciones de *backhaul* para las redes móviles se plantean acorde a la situación en particular de cada operadora, tomando en cuenta su plan de negocios y posibilidades económicas, así como también de la tecnología que se encuentre al alcance. Es así como en el tiempo para las redes móviles, 2G, 3G y ahora LTE se han ido planteando varias soluciones, a continuación se describen algunos ejemplos de este tipo de redes.

El despliegue de las celdas celulares de las operadoras crece conforme su demanda lo indica, en un inicio el *backhaul* eran enlaces de baja capacidad TDM (*Time Division Multiplexing*) o de microondas PDH (*Plesyochronous Digital Hierarchy*) sin redundancia y en estrella, posteriormente nuevas tecnologías fueron apareciendo en el mercado, como es el caso de ATM (*Asynchronous Transfer Mode*) y ahora se tiene

el caso de MPLS (*MultiProtocol Label Switching*), estas tecnologías de transmisión usan como interconexión entre sus nodos SDH (*Synchronous Digital Hierarchy*) y ahora DWDM (*Dense Wave Division Multiplexing*). En la figura 1.13 se representa como ha ido evolucionando las soluciones de *backhaul*¹ que han cambiado acorde a las redes móviles 2.5G, 3G y para el caso de LTE se observa que es únicamente IP en el transporte.

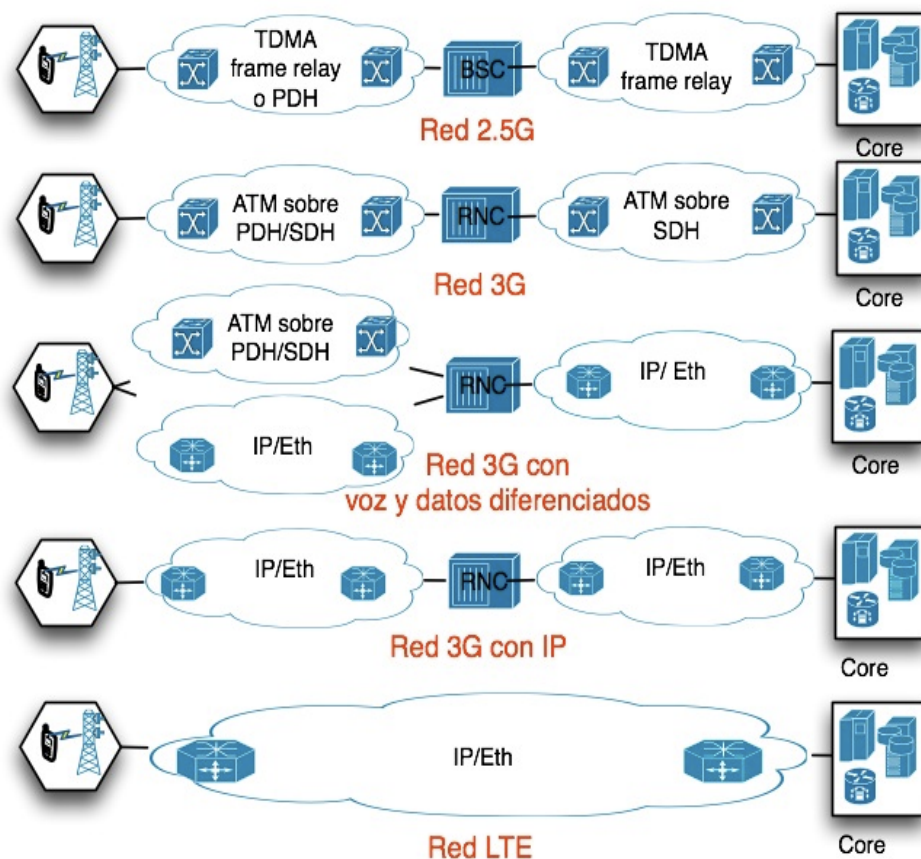


Figura 1.12 Ejemplos de Red de Acceso móvil (Mobile Backhaul) ^{[12] y [24]}

¹ Esta figura es la composición de dos referencias, la parte esencial se tomo de la referencia [12] y se complementó con la referencia [24] en cuanto a la red de transporte para 2.5G, para dar una referencia histórica de la evolución de esta red que une a las radio bases celulares con el core de servicios.

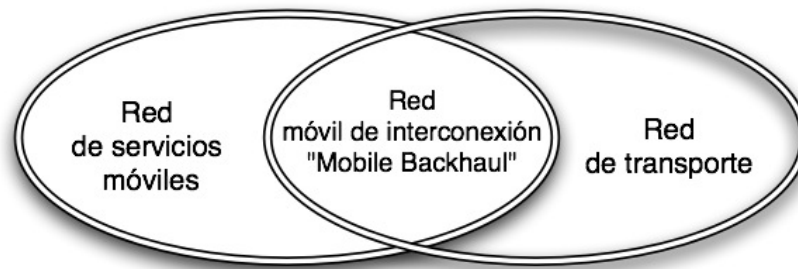


Figura 1.13 Red móvil de Interconexión ("Mobile backhaul") ^[12]

1.2.3 RECOMENDACIONES

En esta sección se quiere sustentar la razón por la que bajo otras opciones presentes en la implementación de redes de *backhaul* o IP-RAN, la tendencia ha sido utilizar las redes IP-MPLS.

Las operadoras como se indica en la figura 1.13, han tenido que resolver sus problemas de conectividad con lo que la tecnología les permitía en ese momento, como se pudo observar en la secciones 1.1.1 y 1.1.2 se concluye que una tecnología basada en conmutación de paquetes y dentro de esta clasificación las redes MPLS proporcionan mejores características que las primeras soluciones de *backhaul*. Entre las que se han consideraron más importantes :

- **Mayor exigencia de tráfico.** Las redes de conmutación de paquetes pueden manejar mayor cantidad de tráfico, más eficientemente, el cual ahora es característica de las nuevas aplicaciones y exigencias como se lo presentó en la sección 1.1.1.

- **Solución de *backhaul* de bajo costo¹.**
- **QoS.** Como se lo presentó en la sección 1.1.1, la diversidad de servicios y aplicaciones demandan de una diferenciación adecuada del tráfico y las redes IP cumplen con esta capacidad.
- **Escalabilidad.** Debido al constante crecimiento por la demanda de servicios móviles, que se observó en la sección 1.1.1, se requiere que el *backhaul* cumpla con la característica de incrementar fácilmente y ordenadamente la capacidad de interconexión.
- **Resiliencia.** Nuevamente los SLAs con los clientes demandan de la red que responda apropiadamente a las fallas que puedan presentarse, es decir que ofrezca redundancia o rápido restablecimiento en los casos extremos.
- **Sincronización.** En el caso de las radio bases, para los servicios móviles es muy importante la sincronización especialmente en el caso de LTE, razón por lo cual se deben observar los mecanismos necesarios para que las redes de transporte consideren esta recomendación.

Una de las soluciones de *backhaul* que cumple con estas recomendaciones son las redes MPLS, razón por la cual se utilizó en el diseño de la red de CNT E.P y por lo cual se estudiará la manera en que se encuentra aplicada esta tecnología y posteriormente cómo optimizarla.

¹ En el capítulo 2 de la referencia [12] se exponen las razones a considerar para afirmar que una red de conmutación de paquetes IP, como lo es la red MPLS, cumple con ser una tecnología de menor costo que sus predecesoras TDM y ATM.

1.2.4 GENERALIDADES DE LA CONECTIVIDAD LTE

Se ha presentado en la sección 1.2.3 las recomendaciones acerca de las redes IP-RAN o de *backhaul*, adicionalmente es necesario considerar requerimientos propios de las redes LTE que hacen que se tomen en cuenta especialmente en el diseño^[5].

Tipos de Tráfico

Como se observó en la figura 1.10, para LTE se tienen los siguientes tipos de tráfico:

- S1-u Tráfico destinado para el SGW. Se ha denominado de esta manera al tráfico que pasa al gateway de aplicaciones dentro del EPC.
- S1-c Tráfico destinado para el MME. Este tráfico es el que comprueba las características del tipo de usuario de la red y qué servicios pueden autorizarse.
- X2-u y X2-c tráfico destinado para otro e-NB. Esto es particular de LTE, nunca antes los nodos o celdas celulares requerían comunicación.

Seguridad

La EPC, al tratarse de una red IP tiene las mismas necesidades de seguridad propias de este tipo de redes, especialmente porque ahora un usuario final que también es IP, accede a los servicios que se encuentran en el EPC, para lo cual se necesitan esquemas de autenticación y protección en todas las fases de la conexión.

QoS

En LTE se tienen algunos nuevos conceptos que tomar en cuenta sobre el tráfico y su prioridad:

- QCI (*QoS Class Identifier*): Un valor del 1 al 9 que corresponde a los niveles de servicio para los tipos de aplicaciones del usuario final, como voz, video, etc.
- GBR (*Guaranteed Bit Rate*): velocidad en bits por segundo que se espera provisionar.
- MBR (*Maximun Bit Rate*): Limites de ancho de banda que espera provisionar.
- ARP (*Allocation and Retention Priority*): Se trata del control de los requerimientos que pueden ser reservados cuando los recursos se vuelven limitados.

Cada QCI corresponde a los diferentes flujos de datos del usuario final, como por ejemplo (voz, video, etc.) con sus correspondientes tipos de recurso de red con o sin GBR. En LTE se identifican los diferentes flujos de datos, definiendo también la prioridad e importancia en caso de congestión o disminución de recursos.

Sincronización

En las primeras redes móviles, la red de *backhaul* como por ejemplo TDM, tiene un esquema de sincronización desde la misma red de transporte, la cual en su momento tomó su respectivo desarrollo. En cuanto a la sincronización a través de una red de conmutación de paquetes, actualmente existen dos opciones: *Synchronous Ethernet* y PTP (*Precision Time Protocol*)

1.3 UTILIZACIÓN DE LAS REDES MPLS

Las redes MPLS (*Multiprotocol Label Switching*) han sido una respuesta escalable y eficiente, que han permitido brindar múltiples servicios, como voz sobre IP, Internet, datos, telefonía fija y móvil; la CNT E.P ha adquirido una red para *backhaul* que

permite interconectar los servicios de datos móviles con la red de acceso LTE, utilizando la tecnología MPLS, integrando esta red a la red de datos nacional que también está basada en MPLS.

En relación al modelo OSI, MPLS se encuentra ubicado entre la capa 3 y la capa 2 y se suele decir que es 2.5; esta consideración es debido a que los paquetes son transportados no por su dirección IP como ocurre en la capa 3 o plano de control, sino por la definición de una etiqueta (*labels*) que este protocolo asigna, creando así una nueva tabla considerada como plano de data (*data plane*); es precisamente con esta nueva tabla que los paquetes son despachados a su destino sin subir al plano de control haciendo mas eficiente su transporte disminuyendo notablemente el *delay*. En la figura 1.14 se muestra la cabecera MPLS que se inserta entre la cabecera de capa dos y el paquete IP, se tiene entonces lo siguiente:

- 20 bits Label, los paquetes son enviados gracias a la información de este campo que se lo conoce como etiqueta.
- Traffic Class se conoce el campo de 3 bits, que en mucha literatura se lo conoce como Experimental bits. Se utiliza esta etiqueta para definir la estrategia de calidad de servicio.
- *Bottom of stack bit* (S-bit), este bit identifica si el paquete es o no el último de la pila de etiquetas, con lo cual podemos indicar que para los servicios se puede trabajar con pilas o *stacks* de etiquetas, esto se entenderá mas adelante.
- 8 bits de *Time to Live* utilizado de la misma manera que en el protocolo IP, para evitar los lazos de enrutamiento, el paquete ingresa con un valor de TTL y se va disminuyendo en uno a lo largo del camino.

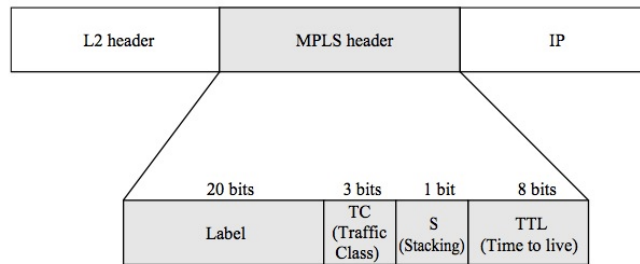


Figura 1.14 Cabecera MPLS^[12]

En el año 1997 ¹ empezó el grupo de IETF a tratar los temas de esta nueva tecnología que con sus características hicieron que una red de transporte permita hacer realidad muchas oportunidades de negocio en las redes de datos, a continuación podemos señalar las que se han considerado importantes:

1. **Escalabilidad de enrutamiento.** Al utilizar etiquetas para el tratamiento de los paquetes se agrega eficiencia frente al enrutamiento tradicional, gracias a la información adicional que estas etiquetas representan, a esta información se la conoce como *forwarding information* o FEC (*Forwarding Equivalence Class*), mientras tanto se trabaja en jerarquías de enrutamiento.
2. **Gran flexibilidad en la entrega de servicios de enrutamiento.** Debido a que el tráfico de cada paquete tiene etiquetas y se puede identificar un flujo de tráfico de todos los flujos presentes, es posible dar un tratamiento de calidad de servicio y conseguir diferenciar los tipos de servicio de la red.
3. **Mejora de desempeño.** Aunque la mayoría de routers modernos, tengan circuitos especializados que han logrado alcanzar tasas de transferencia de datos comparables a una red MPLS, dentro de la red MPLS existe la opción

¹ Tomado del Capítulo 1 foundations de la referencia [14] parte introductoria de MPLS con el apoyo del material didáctico [18] y [19]

de generar túneles de ingeniería de tráfico que permitan la optimización del uso de los enlaces de la red.

4. **Facilidad de integración.** Gracias a las funcionalidades de la red MPLS, se puede integrar con tecnologías como ATM o *frame relay* para aquellos proveedores de servicios que operan todavía con esta tecnología, por supuesto estas redes han sido reemplazadas poco a poco.

A continuación se presenta un conjunto de conceptos básicos¹ sobre la arquitectura de las redes MPLS, realizando una descripción de los servicios que se pueden implementar. En la figura 1.15² se ha querido representar la componentes de las redes MPLS: en la parte medular o en inglés *Core* se encuentran los equipos denominados LSR (*Label Switch Router*) o también conocidos como P (*Provider router*), luego se encuentran los equipos de frontera o en inglés *Edge* que se los conoce como LER (*Label Edge Router*) o PE (*Provider Edge*). Se crea un conjunto de túneles unidireccionales, conocidos con el nombre de LSP (*Label Switch Path*), para que los paquetes del cliente que entran por el PE de ingreso lleguen al apropiado PE de salida o egreso. Cuando los paquetes ingresan a un red MPLS, el router PE de ingreso determina a que FEC (*Forwarding Equivalence Class*)³ pertenece, por lo que todos los paquetes que requieran el mismo destino desde este PE tienen el mismo FEC, y comparten el mismo camino (LSP).

¹ Se utiliza las laminas didácticas de la referencia [18] y [19], las referencias [12] Capítulo 4.7 y [14] Capitulo 1 contienen el fundamento teórico de esta sección.

² Gráfica tomada de la referencia [18]

³ FEC, es un conjunto de paquetes que requieren llegar a un mismo destino y que son tratados de la misma manera por el mismo camino.

El papel principal del PE de ingreso es determinar a que equipo de salida se debe llegar y que LSP (camino) le corresponde, el cual está asociado al correspondiente FEC. MPLS tiene la propiedad de “multiplexar” por medio del mismo LSP varios tipos de tráfico, el operador de la red puede enviar a través del mismo LSP, para un mismo PE de ingreso todo el tráfico (por ejemplo L3VPN, túneles L2VPN que son los grupos de servicios más importantes que las redes MPLS ofrecen a sus clientes). Los equipos de Core LSR o P de la red en cambio realizan el paso del tráfico basados únicamente en las etiquetas¹ de la cabecera MPLS² y de esta manera este equipo no necesita conocer la información de enrutamiento de los clientes (tanto de servicios L3VPN o L2VPN) solo entrega los paquetes a su destino .

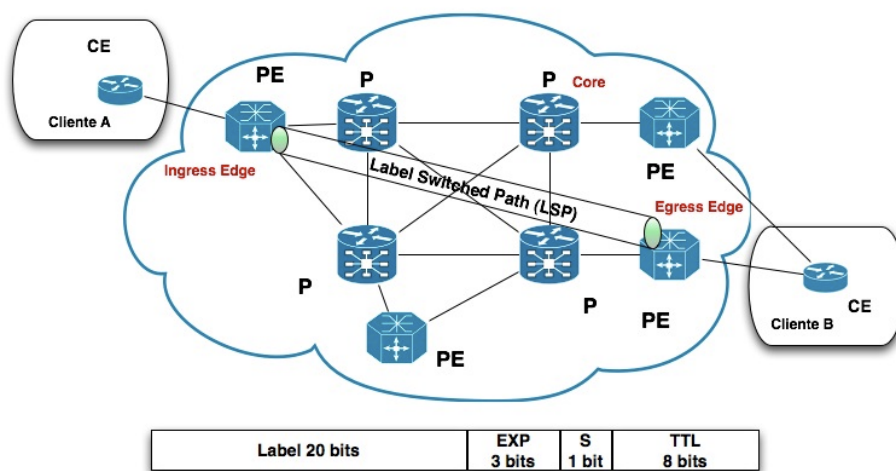


Figura 1.15 Funciones de los equipos de MPLS^[18]

¹ A esto se le suele conocer como *swapping of labels* o intercambio de etiquetas.

² Tanto en la figura 14 como la 15 se encuentra representado la cabecera MPLS, el autor de la referencia [18] considera todavía a los tres bits como experimental bits, en la referencia [12] se lo denomina ya como Traffic Class.

Los routers PE y los P se encuentran interconectados entre si por varios medios¹, que pueden ser ópticos de fibra oscura o DWDM; en algunos casos esta transmisión puede llegar por un puerto eléctrico (RJ45) a tecnologías de radio, NG-SDH (*Next Generation Synchronous Digital Hierarchy*) e incluso transmisión satelital. En la transmisión se debe asegurar que los paquetes IP sean mayores a los 1500 bits de la trama ethernet común ya que por las etiquetas y otras aplicaciones se requiere de paquetes de mayor tamaño para que los datos no se vean afectados². Además para la interconexión entre los equipos se utiliza un protocolo de enrutamiento, que para el caso de proveedores de servicio se encuentra OSPF o ISIS, la CNT E.P a preferido trabajar con ISIS en lugar de OSPF debido a su escalabilidad y sencilla configuración.

Para los PE la construcción de la FEC de cada destino depende de dos componentes muy importantes:

- Plano de Control, responsable por mantener la información de las etiquetas correctamente, existen tres mecanismos para que los equipos adquieran información de etiquetas³: LDP (*Label Distribution Protocol*), RSVP (*Resource Reservation Protocol*) o BGP (*Border Gateway Protocol*).
- Plano de datos, luego de que la información del plano de control es procesada, a nivel de hardware en los PEs y Ps de la red se tiene la equivalencia de cada destino con su correspondiente interface de salida⁴

¹ Esta nota se refiere a las interconexiones que los proveedores de servicio tienen en nuestro país.

² En algunos equipos se requiere activar una opción llamada "jumbo frames" para poder pasar tramas de mas de 1500.

³ A esto se lo conoce como LIB (*Label Information Base*).

⁴ Esta es la definición de LFIB (*Label Forwarding Information Base*).

para que el paquete se entregue lo antes posible sin pasar por el plano de control.

En la figura 1.16 se ilustra la operación¹ o el tratamiento de un paquete de datos al ingresar a una red MPLS, representado con la línea de color celeste se tiene la tabla de enrutamiento (protocolos OSPF o IS-IS) en la cual se conoce como llegar a las redes de destino, por otro lado se encuentra el protocolo que distribuye la información de la etiquetas (representado con color naranja), que en la mayoría de casos es LDP o RSVP y en otros que si veremos mas tarde BGP; en el plano de datos o el plano de “*forwarding*” se tiene la asignación que cada router hace de etiquetas por cada red de destino, esta información contiene también la interface de salida para alcanzar el destino. El proceso entonces es así:

1. El paquete de datos del cliente ingresa al primer PE buscando su destino, con la información que posee de las etiquetas² se asigna al paquete una etiqueta correspondiente a la red destino³.
2. El PE entonces envía el paquete por la interface que debe salir.
3. El siguiente router P, lee la etiqueta y asigna la que le corresponde según su tabla de destinos y envía el paquete, a esto se le suele denominar *swapping of labels* o intercambio de etiquetas.

¹ Este ejemplo ilustra el tratamiento de un paquete ipv4 sobre una red MPLS, el cual es el más simple.

² Esta información de las etiquetas con su correspondencia de red destino se lo conoce como LIB (*Label information Base*) y pertenece al plano de control, la tabla que añade la información de la interface de salida se la denomina LFIB (*Label Forwarding Informatio Base*) y corresponde al plano de datos.

³ Esta asignación de etiquetas es local en cada router, es decir para cada red de destino se tienen un etiqueta, luego mediante LDP se comparten los routers esta información entre si.

4. Finalmente el último router retira la etiqueta impuesta y entrega el paquete a su destino¹

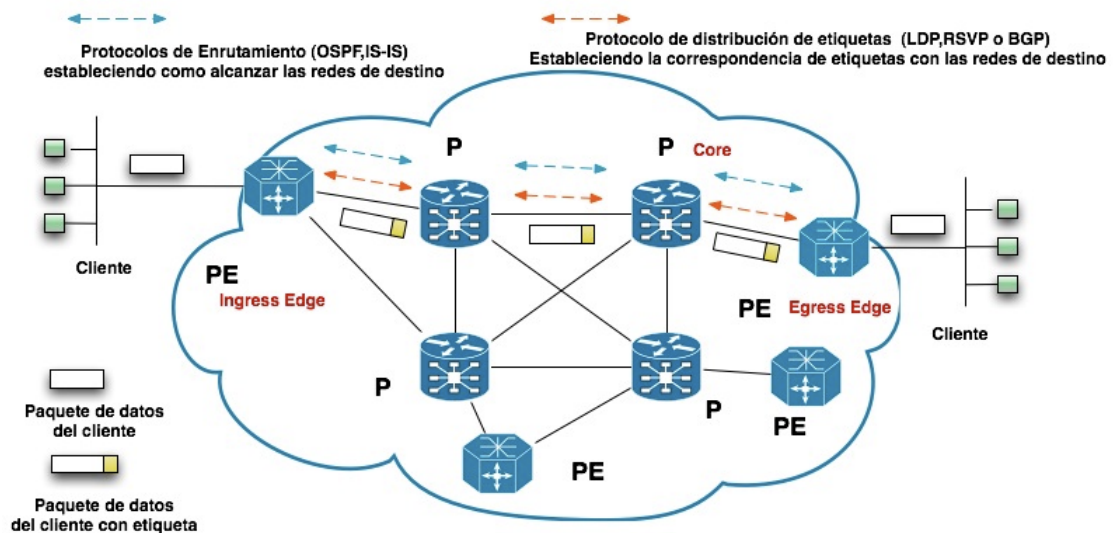


Figura 1.16 Operación de la red MPLS ^[18]

Como² se puede ver el paquete de datos ingresado es tratado en la red por sus etiquetas ya no mas por su información IP, es decir en cada salto la cabecera IP ya no es procesada sino su información de etiqueta. En el PE de ingreso se realiza la imposición de las etiquetas sobre el paquete de datos del cliente, como se había indicado al mencionar la cabecera MPLS, el paquete de datos puede tener no solo una única etiqueta sino un conjunto o una “pila de etiquetas” con el propósito de brindar servicios como por ejemplo:

¹ Por razones de optimización de tiempo los router al conocer el camino que se debe seguir y por lo tanto se conoce cual es el último router, entonces en un paso anterior retira la etiqueta y lo entrega así al ultimo router, quien entrega simplemente el paquete a su destino, a esto se le denomina PHP (*penultimun hope popping*)

² En la referencia [12] se realizan estas conclusiones en el Capítulo de MPLS Architecture.

- VPN de capa 3 (*Virtual Private Networks*, redes de clientes que tienen su propio enrutamiento privado)
- VPNs de capa dos, en este servicio la red MPLS trabaja como un *switch* de capa 2 sin ver direcciones *IP* sino proporcionando conectividad como una LAN extendida. Estos servicios son de manera general, punto a punto o multipunto.
- Túneles de Ingeniería de tráfico, que sirven para manipular un grupo de servicios a través de un LSP determinado por el operador.

1.3.1 APLICACIONES DE REDES MPLS, L2VPN y L3VPN

Los servicios de L2VPN y L3VPN son los que se utilizan para una red de *backhaul* o IP-RAN en cuanto a una red MPLS, es por esto que se requiere hacer una exposición del funcionamiento de estos servicios.

Una VPN¹ (*Virtual Private Network*) es una red que se comporta como una red privada dentro de una infraestructura compartida de un proveedor de servicios con otros clientes, como tal puede, según el modelo OSI, prestar servicios como una red de capa 2 (L2) o una red de capa 3 (L3). La conectividad en todos los sitios es el requerimiento mínimo que los clientes de una VPN esperan, de igual forma de acuerdo a la demanda de servicios existen varios servicios adicionales dentro del VPN como: Servicio de Internet, interconexión con otras VPNs, enrutamiento dinámico, soluciones punto a punto o punto multipunto.

¹ Estos conceptos básicos son tomados del capítulo 7 de la referencia [20]

1.3.1.1 L3VPN

Como¹ se presenta en la figura 1.17, en MPLS se nombra al equipo del cliente como CE (*Costumer Edge*) , a los equipos de Core como P (*Provider*) y los equipos MPLS que reciben la conexión del cliente PE(*Provider Edge*); en esta figura se puede apreciar como un mismo proveedor de servicios, con una misma red puede proporcionar varias VPNs sin que las redes del cliente A interfieran con el cliente B, incluso pueden tener los dos las mismas direcciones IP. Para explicar mejor como se llega a brindar estos servicios es necesario comprender algunos conceptos básicos como: VRF (*Virtual Routing Forwarding*) , *route distinguisher* (RD), *route targets* (RT), propagación de rutas por medio de MP-BGP (*Multiprotocol-Border Gateway Protocol*) y como se distribuyen los paquetes con etiqueta en la red MPLS.

Se denomina VPN L3 porque entre los CEs y el router PE existe una conexión en capa 3 y un protocolo de enrutamiento como: rutas estáticas, eigrp, ospf, rip, bgp o isis. Los routers P no están involucrados en este enrutamiento ellos solo realizan el “swaping” de etiquetas.

Los routers PE que trabajan con MPLS tienen una propiedad de virtualización, es decir dentro de un mismo router se puede tener varias instancias (routers virtuales) de enrutamiento que son independientes la una de la otra y ninguna se ve con la tabla de enrutamiento global, que es la comúnmente conocida y es donde el equipo se comunica con los otros routers. A estas instancias se las conoce como VRFs

¹ Tomado de la sección MPLS VPN Model de la referencia [20]

(*Virtual Routing Forwarding*), dentro de las VRFs también puede existir enrutamiento estático o dinámico y puede crearse una VRF por cada cliente; en la figura 1.17 se ha representado las tablas de enrutamiento del PE de ingreso, donde se pueden ver las tablas de las VRFs independientes.

Se utiliza MP-BGP para transportar los prefijos de cada VPN de los clientes por la red MPLS, esta información es importante en el PE de ingreso y en PE de destino. En el caso de que se utilizara BGP únicamente en ipv4 para transportar los prefijos del cliente, el momento que otro cliente tenga el mismo direccionamiento habría inconsistencias; es por esta razón que se crea el concepto de RD (*route distinguisher*) para resolver estas inconsistencias y de esta manera el prefijo del cliente asociado a este identificador único hace posible que se distinga a que cliente pertenece ese prefijo. Esta combinación del RD que es un campo de 64-bits mas el prefijo ipv4 del cliente se ha denominado prefijo vpnv4 y por esto MP-BGP tiene un *address-family* VPNV4 para transportar este tipo de información entre los PEs que lo requieran. El tamaño de un prefijo vpnv4 es de 96 bits y su formato puede ser ASN:nn o IP-address:nn ¹, donde nn es un número que ayuda a identificar a un cliente, gracias a esto las direcciones IP de un cliente pueden ser las mismas que otro y no existen problemas de comunicación al transportarlas por la red MPLS, como ejemplo se puede decir que un cliente tiene este prefijo ipv4: 10.1.1.0/24 y un RD 1:1 entonces su prefijo vpnv4 es 1:1:10.1.1.0/24.

¹ **ASN**, *autonomous system number* se conoce como el sistema autonomo, por lo general es el número que identifica a una entidad o proveedor de servicios de internet y cuyos prefijos llevan este identificador en BGP para indicar que pertenecen a un mismo grupo de administración. **IP-address**, es la dirección ipv4 de la red del cliente de la VPN.

Hasta el momento con el concepto del RD¹ no existiría la posibilidad de comunicación entre VPNs, en la práctica esta comunicación entre clientes existe gracias a otro concepto el RT (*Route Target*). El RT se lo define aprovechando la utilidad de un atributo del protocolo BGP conocido como *extended community* y se lo emplea para controlar que rutas se importan y exportan dentro de la VRF a través de las rutas vpnv4. El formato es el que se definió para el RD, ASN:nn o IP-address:nn y de la misma manera los proveedores prefieren emplear la primera opción; en el lugar donde se define los datos principales de la VRF se define el RT de importación y de exportación que pueden ser configurados a conveniencia, es decir que pueden ser varios de importación y exportación dependiendo de lo que se requiera. Para una VPN que solo tenga comunicación con sus sitios remotos sin involucrar sitios de otra VPN, el RT debe ser el definido para esta VPN correspondiendo el de exportación de un sitio con el de importación del otro; comúnmente lo que se suele hacer es definir un RT para esta VPN y configurar tanto para exportación como para importación.

En la figura 1.18 se puede observar en primera instancia, la VPNA y la VPNB que tienen dos sitios que se ven cada una entre si, a esto se lo denomina “intranet” y en el ejemplo se ha definido el RT de importación y exportación 1:1 para la VPNA y 1:2 para la VPNB. Ha pedido de los clientes se ha solicitado que solamente en el sitio 1 se vean las redes de la VPNA con la VPNB y para los otros dos sitios las redes entre los dos clientes continúen independientes. Para cumplir con este objetivo se ha definido un RT de exportación e importación de 100:1, el cual permite que las dos

¹ Este concepto de VPN con sus RDs y RTs son tomados del Capítulo 7 de la referencia [20], así también las figuras que sirven de ejemplo para estos conceptos.

VPNS comparten sus rutas únicamente¹ en el sitio 1, a esta compartición de rutas se la conoce como “extranet”

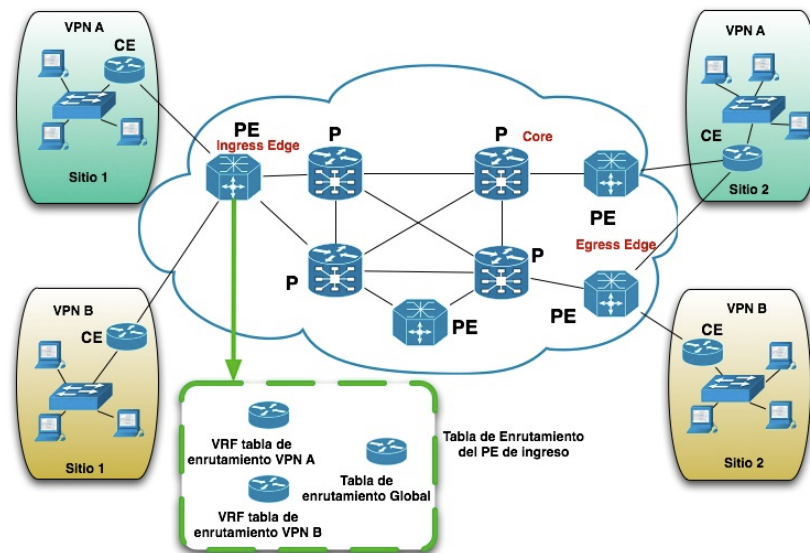


Figura 1.17 Modelo de MPLS VPN L3 [20]

Hasta el momento² se ha definido el concepto de VRF para separar las rutas de clientes en cada PE, los mecanismos necesarios en cada PE para generar los prefijos, así como también como intercambiar rutas entre VRFs de distintos clientes, restando tomar en cuenta cómo transportar estas rutas a través de la red MPLS. Las rutas de los clientes de MPLS VPN pueden llegar ser un número considerable y el protocolo que se caracteriza por su escalabilidad es BGP; se ha definido dentro del protocolo una nueva familia denominada VPNv4, las rutas ipv4 con su respectivo RD que identifica a que VPN son convertidas en rutas VPNv4 y mediante sesiones BGP

¹ A la configuración de este RT 100:1 se le puede añadir una política o filtro restrictivo para hacer más precisa la redistribución de rutas de una VPN a la otra, para evitar el cargar rutas innecesarias o que puedan causar inconsistencias.

² En esta parte se expone a cerca de la propagación de rutas vpn4, tomado de la sección VPNv4 Route Propagation in the MPLS VPN Network de la referencia [20].

internas¹ entre PEs se realiza el intercambio de rutas, además se debe tomar en cuenta que estas rutas se originan en un PE como directamente conectadas, como por enrutamiento estático o un IGP entre PE y CE dentro de la VRF, también se acostumbra en algunos casos a utilizar eBGP entre el CE y el PE. En la figura 1.19 se muestra el intercambio de rutas VPNv4.

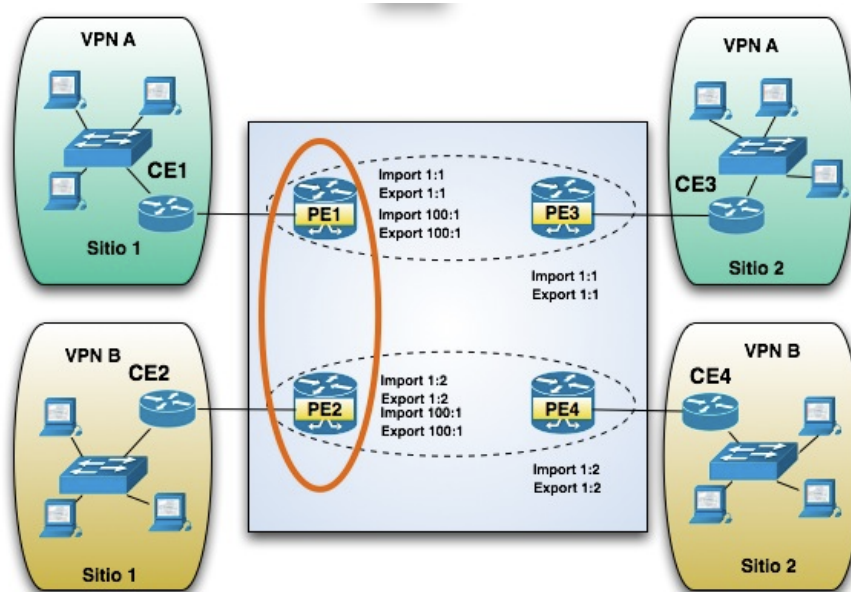


Figura 1.18 Ejemplo de configuración de Route Target VPNA y VPNB ^[20]

Es importante también mencionar que cada PE publica una etiqueta por cada prefijo VPNv4 y es entonces que los paquetes pueden pasar desde un sitio a otro del cliente a través de la red MPLS. Para redes con una cantidad considerable de PEs se tiene múltiples sesiones iBGP, lo cual no suele ser escalable, los proveedores utilizan el concepto de *route-reflector*.

¹ Estas sesiones BGP se las denomina también MP-BGP ya que existen varias address family y entonces se dice que se trata de Multiprotocol BGP

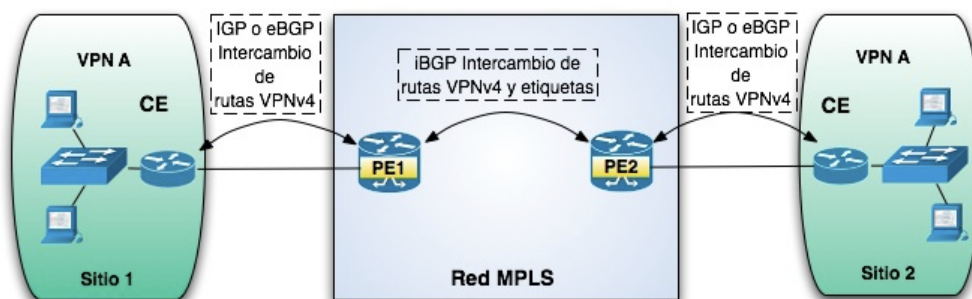


Figura 1.19 Propagación de rutas en una MPLS VPN ^[20]

1.3.1.2 L2VPN

Los servicios de L2VPN en MPLS surgieron como respuesta para que los proveedores de servicios unificaran en una red de Core convergente todos los servicios que tradicionalmente se brindó en un principio, como por ejemplo frame-relay, ATM y el tradicional IP¹; adicionalmente el cliente, con este servicio, puede mantener su propio enrutamiento, políticas de calidad de servicio, mecanismos de seguridad independientes del proveedor de servicios que se convierte en transporte transparente y finalmente este tipo de servicios interactúan con equipos de varios proveedores al ser de arquitectura abierta.

Las L2VPNs en general proveen conectividad punto a punto o punto multipunto entre los nodos de operación del cliente donde se le entrega el servicio, esta conectividad

¹ Este análisis y conceptos de L2VPN son tomados de la referencia [21] y algunas ilustraciones de la referencia [19]

es a nivel de capa 2, pudiendo transportarse los protocolos como: ATM, Frame-relay, Ethernet, PPP, HDLC, etc.

1.3.1.2.1 Modelos de servicio L2VPN

Los servicios de L2VPN puede darse en una red MPLS con lo que se denomina AToM(*Any Transport over MPLS*) o en una red IP con el protocolo L2TPv3¹, los servicios se suele clasificar en tres grupos: de conexión local o *Local Switching*, VPWS (*Virtual Private Wire Service*) que provee conectividad punto a punto y VPLS(*Virtual Private LAN Service*) que provee conectividad multipunto, como se ilustra en la figura 1.20. En esta figura se puede observar los modelos en general en una red IP y MPLS, además se menciona los tipos de interfaces que el cliente puede tener de su lado y las posibilidades de conectividad que se le pueden transportar por la red IP y la red MPLS. Los servicios punto a punto involucran una variedad de interfaces, pero únicamente la interface ethernet permite la conectividad multipunto.

1.3.1.2.2 Arquitectura del servicio L2VPN

El servicio de L2VPN² tiene tres elementos en su arquitectura, los AC (*attachment circuits*), el *pseudowire* y la tecnología de transporte

¹ Con redes IP solo se pueden entregar servicios punto a punto y con la misma tecnología al ingreso del punto de presencia del cliente.

² Esta definición de L2VPN es tomada de la referencia [21] Layer 2 VPN Architecture

Un AC es el circuito o enlace directamente conectado a un PE que es virtualmente extendido a otro PE a través de la red del proveedor que puede ser IP o MPLS.

Un *Pseudowire* es una conexión punto a punto a través de una red de conmutación de paquetes; el pseudowire emula la operación de un cable transparente que une dos AC de dos diferentes PEs.

El tercer elemento de esta arquitectura es la tecnología de transporte, el cual puede ser MPLS y sus interfaces capaces de soportar cualquier protocolo de capa 2 (punto a punto, punto a multipunto, o multipunto a multipunto), o una red IP con L2TPv3 en la cual solo se soportan conexiones punto a punto.

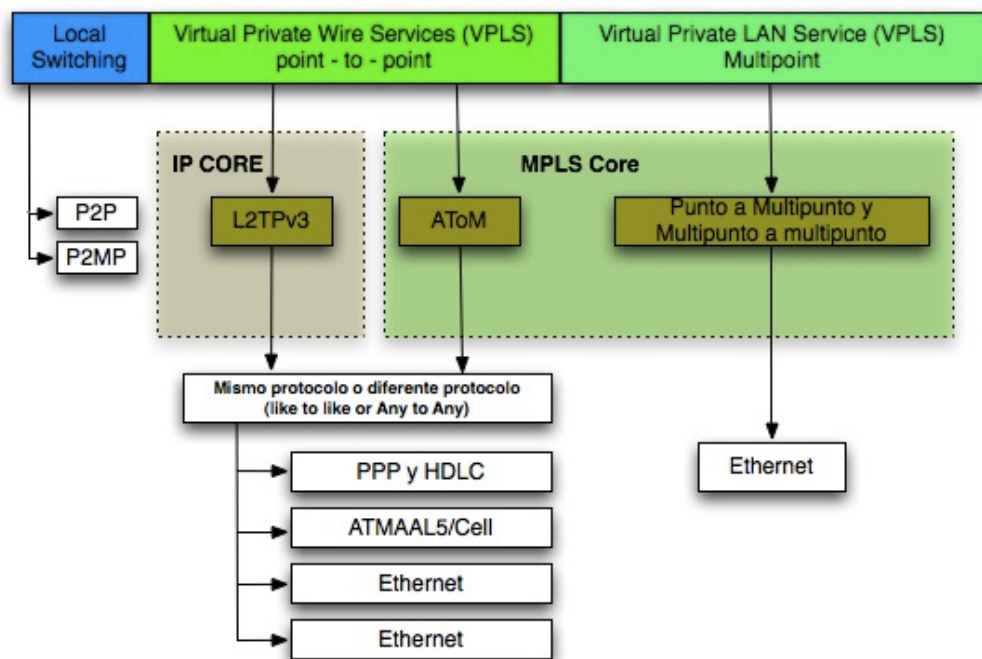


Figura 1.20 Modelos de L2VPN, VPWS y VPLS ^[21]

diferenciar en el paquete a que VPN pertenece y entregar el paquete. Esta segunda etiqueta se denomina también etiqueta de VC (*virtual circuit*) y significa que se hace la correspondencia entre el paquete y un VC o interface de salida.

1.3.1.2.3 L2VPN Plano de Control

Los PEs que se puede apreciar en la figura 1.22¹, requieren de un protocolo que permita intercambiar la información de los VCs. En el caso de las redes MPLS, este protocolo es LDP, luego del aprovisionamiento respectivo² una adyacencia directa entre los dos PEs es establecida³, el PE de salida o de egreso es el que le indica el valor de la etiqueta para determina la FEC de un VC⁴, es entonces el PE de ingreso quién asigna una segunda etiqueta que se apila⁵ en la trama que completa la FEC del VC. En el caso de que la red de transporte sea solo IP es el protocolo L2TPv3 el cual es encargado de transmitir los parámetros de establecimiento del túnel y no etiquetas como en el caso de MPLS.

En la figura 1.22 se ilustra la sesión directa que se establece entre los PEs de ingreso y egreso que es empleada para intercambiar la información de las etiquetas del VC. Es importante notar que la sesión LDP debe establecerse tanto en el PE de ingreso como en el PE de salida. Esta sesión constituye el plano de control de una L2VPN y

¹ Esta sección es tomada de la referencia [21] del tema Layer 2 VPN Control Plane

² En esta sección se está presentando el concepto global y fundamental de la solución, el cual se aplica a varias marcas cuyas sintáxis de configuración varían de acuerdo a cada plataforma.

³ De la experiencia cuando un circuito L2VPN deja de operar, una buena práctica es confirmar que esta sesión entre PEs siga existiendo ya que es fundamental su presencia.

⁴ Al inicio de la sección se expuso el concepto de FEC.

⁵ En inglés se entiende *label stack*

es responsable de proveer los elementos de interconexión con los servicios nativos, tales como LMI en frame relay, DLCI, etc.

El establecimiento de esta sesión de control se la puede hacer manualmente o por auto descubrimiento usando las extensiones MP-BGP o por LDP.

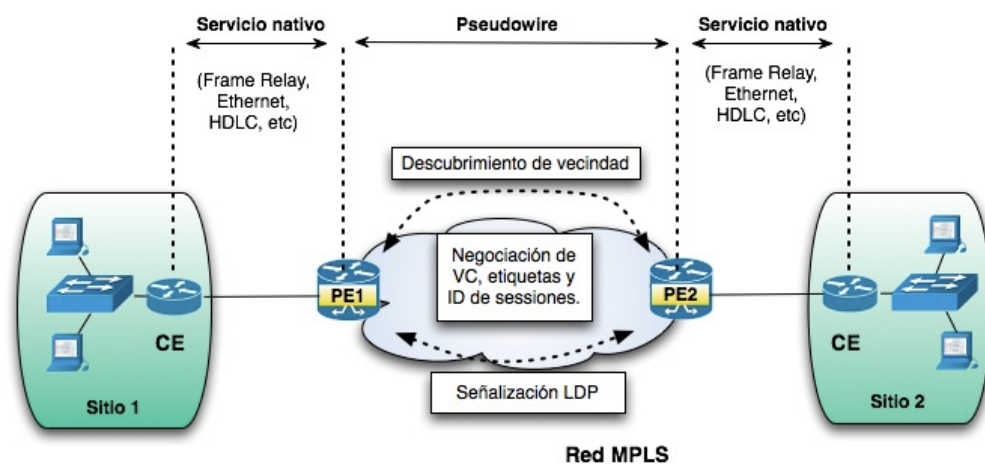


Figura 1.22 Plano de Control de L2VPN ^[21]

1.3.1.2.4 Plano de Datos L2VPN para una red IP, L2TPv3

En la figura 1.23 se puede observar la trama del protocolo L2TPv3 que es utilizada para transportar tramas de capa 2 sobre una red IP pura¹.

Todo el paquete L2TP, incluyendo la cabecera y la data se envían un datagrama UDP. Tradicionalmente el uso de L2TP ha sido transportar sesiones PPP con un túnel de este protocolo, adicionalmente la versión 3 provee seguridad, mejoras de encapsulación y la facultad de transportar no solo PPP, sino también otros protocolos como Frame Relay, Ethernet, ATM, etc.

¹ Con este término se quiere decir una red IP que no es MPLS.

La encapsulación L2TP incluye una parte de cabecera IP de transporte (20 bytes) y una parte variable de L2TP cuya parte fija es 4 bytes de Sesión ID. Los campos adicionales son 8 bytes de cookie y la control word. Finalmente el resto de la trama es el paquete original de L2.

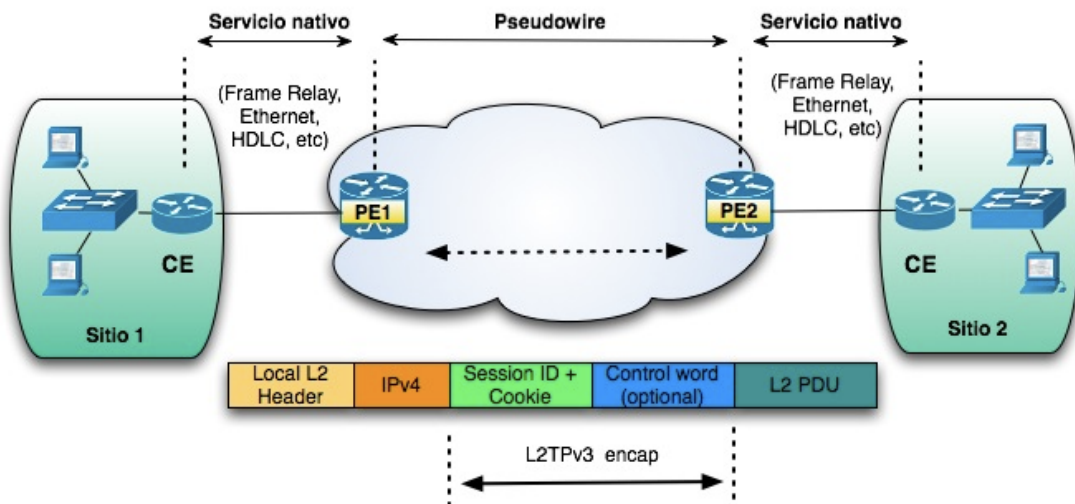


Figura 1.23 Plano de datos para un red no MPLS ^[21]

1.3.1.2.5 Plano de Datos L2VPN para una red MPLS

El mecanismo de transporte esta basado en el mismo mecanismo que se examinó para las L3VPN, salto a salto el protocolo LDP señala un camino unidireccional hasta el PE de salida, y una vecindad directa intercambia las etiquetas VC, la cual se encuentra como etiqueta interior en el paquete MPLS como se muestra en la figura 1.24. En general el uso la palabra de control es opcional y configurable, esta opción es requerida para algunos tipos de protocolos de transporte de la conexión con los clientes AC, como por ejemplo Frame Relay. Antes de la encapsulación MPLS algunos campos de la trama original del cliente son extraídos, dependiendo de la trama de capa 2 que se encapsula, como por ejemplo el campo checksums.

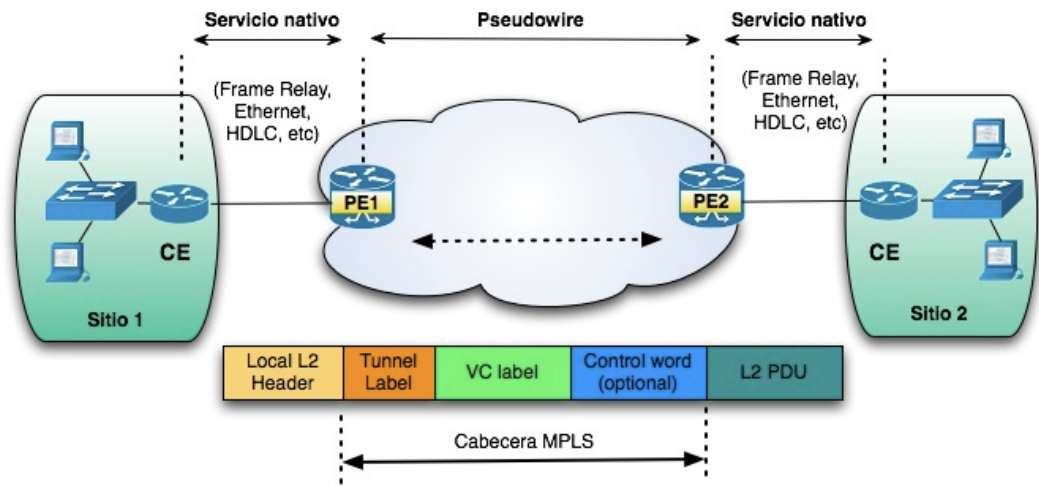


Figura 1.24 Plano de Datos L2VPN en redes MPLS ^[21]

1.3.1.2.6 Tipos de Servicios de una L2VPN.

Existen dos grandes grupos para clasificar los tipos de servicios de una L2VPN, estos son VPWS (*Virtual Private Wire Service*) y VPLS (*Virtual Private LAN Service*).

a) VPWS

Es el tipo de servicio utilizado para provisionar conexiones L2 punto a punto, no tiene la capacidad de aprender direcciones MAC¹ y posee dos métodos de transporte: L2TPv3 y AToM (*Any Transport over MPLS*).

b) VPLS

Es la tecnología que soporta conexiones L2 multipunto, agrupando una colección de *pseudowires* que terminan en un PE en una VFI (*Virtual Forwarding Interface*), Una

¹ Se refiere a las direcciones las direcciones físicas de una interface o dispositivo ethernet.

VFI representa una extensión virtual de un circuito físico agregado a un PE, se puede también indicar que tiene la capacidad de aprender direcciones MAC de los clientes. Una VPLS puede interconectar a varios PEs a través de una VLAN, lo cual le hace una solución escalable.

Adicionalmente se tiene identificado tres tipos de servicios: E-Line, E-LAN y E-Tree, bajo la clasificación de los servicios ethernet por parte de la MEF (*Metro Ethernet Forum*) y se definen a continuación:

- **E-Line**, se refiere a un servicio que interconecta dos clientes de puertos ethernet a través de una WAN, este servicio es subdivido en EPLs (*Ethernet Private Lines*), EVPLS (*Ethernet virtual private lines*) y EIA (*Ethernet Internet Acces*). Estos servicios son conocidos como de punto a punto.
- **E-LAN**, es una conexión multipunto de varios puntos de presencia del cliente a través del proveedor, de tal manera que cada punto tiene la experiencia de estar conectado a un bridge ethernet¹
- **E-Tree**, un servicio multipunto conectado uno o varios punto raíz, previniendo la comunicación entre los puntos extremos entre² si; este servicio es conocido también conocido como *rooted multipoint* . Este tipo de servicio es considerado como ideal para *Mobile backhaul*³ y para infraestructura *tripe-play*⁴

¹ El concepto de bridge es refiriendose a los equipos que unen dos puntos ethernet, se puede recordar que los puertos de un switch se considera cada uno como un bridge.

² Por aclarar que el punto raíz se comunica con todos los puntos finales pero ellos no entre si.

³ Este concepto fue tratado en la sección 1.2.

⁴ Se trata del paquete de servicios que ofrecen las empresas de telecomunicaciones como Voz, internet y Televisión.

Finalmente se puede asociar los servicios E-LINE a los servicios indicados de VPWS y los servicios tanto de E-LAN como de E-Tree a los de VPLS.

1.3.2 ARQUITECTURAS DE REDES DE ACCESO IP-RAN PARA SERVICIOS LTE UTILIZANDO UNA RED MPLS COMO INTERCONEXIÓN AL CORE DE DATOS MÓVILES

En la sección 1.1.4.2.2, en la figura 1.10 se puede observar la arquitectura de la red de acceso LTE y los tipos de tráfico siguientes:

- S1-u Tráfico destinado para el SGW, Se ha denominado de esta manera al tráfico que pasa al gateway de aplicaciones dentro del EPC.
- S1-c Tráfico destinado para el MME. Este tráfico es el que comprueba las características del tipo de usuario de la red y que servicios pueden autorizarse.
- X2-u y X2-c tráfico destinado para otro e-NB, esto es particular de LTE, nunca antes los nodos o celdas celulares requerían comunicación.

Para conseguir esta conectividad para los tipos de tráfico LTE en el la industria se han planteado varias opciones que se revisan a continuación¹.

¹ En la referencia [5] se ha extraído de la secciones Layer 2 VPN Model for LTE/EPC Deployments y Layer 3 VPN Model for LTE/EPC Deployments.

1.3.3 SOLUCIONES DE L2VPN

Utilizando aplicaciones de L2VPN para el *backhaul* de tráfico LTE se presenta una solución esquematizada en la figura 1.25, donde se puede ver que para el servicio de los e-NB, agregados en un *cell-site*¹, se tienen dos VLANs, una para alcanzar los servicios del Core o EPC y la otra para el tráfico X2. La VLAN de Core necesita operar bajo el concepto de punto a punto, entonces se extiende un servicio de E-LINE (se entiende Ethernet pseudowire) que se extiende desde los *cell-sites* hasta la capa de agregación pasando por la de pre-agregación. La VLAN para el servicio de X2 utiliza el servicio de E-LAN (Se entiende VPLS) que será como los eNB se comunicarán directamente para realizar el *Handover*, se recuerda que esta comunicación es crítica el momento del *Handover*.

Se deben tomar las siguientes consideraciones:

- Esta solución E-LAN para el tráfico X2 presenta un problema a resolver, debido a que un usuario móvil (el equipo del usuario) se conectará a diferentes *cell-sites*, los cuales deben tener conectividad entre si. Aunque el número de celdas vecinas sea bajo (se considera de 10 a 15 como bajo), el caso es que esta lista de celdas contiguas cambia constantemente debido a que el equipo del usuario se mueve de celda en celda. Se identifican dos factores a considerar:

¹ Para las soluciones de backhaul, cell site es el equipamiento que agrega los e-nodoB

- a. Los dominios E-LAN no pueden ser muy extensos ya que esto implica un gran dominio de *broadcast* y por lo tanto un alto riesgo de seguridad.
- b. Los dominios E-LAN deben comunicarse entre si para alcanzar el *Handover*, por lo cual se debe tener cierto grado de zonificación para el tráfico X2 y de esta manera interconectarse en la zona de servicios de pre agregación (jerarquía de los servicios E-LAN). Esta zonificación debe estar constituida de tal manera que los *cell-sites* sean alcanzables para que siempre el *Handover* sea posible complicando en cierto grado el diseño

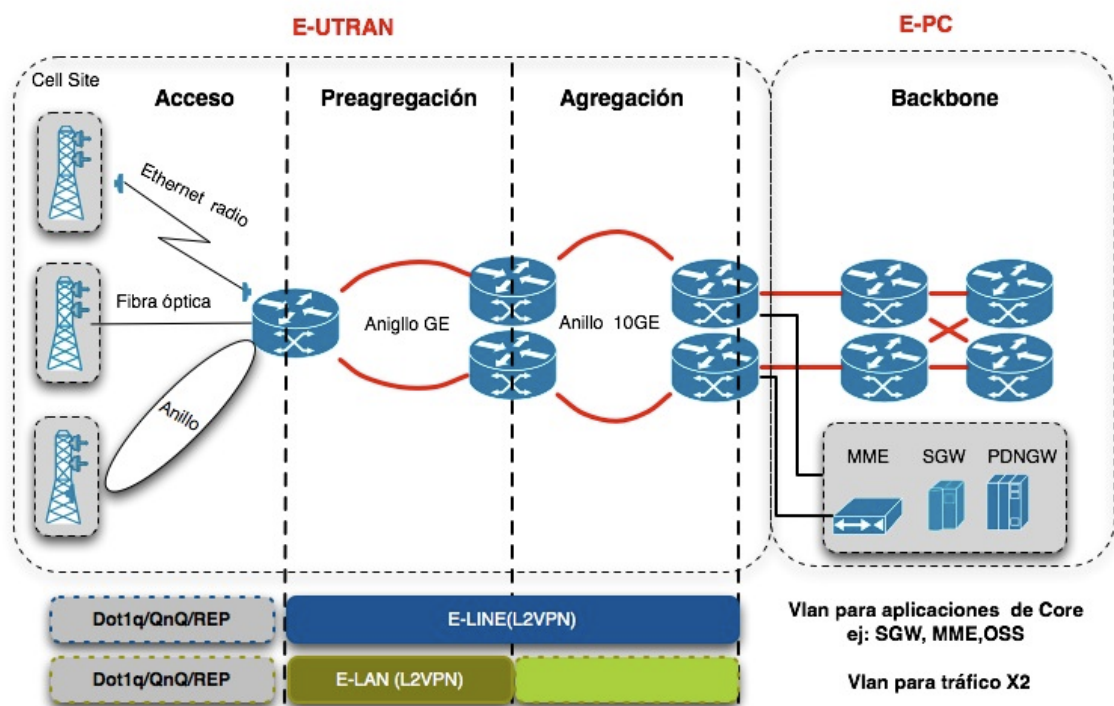


Figura 1.25 L2VPN para interconectar LTE y EPC ^[9]

- Esta solución de servicios E-LAN puede resultar en un gran dominio de *broadcast* lo que representa un mayor riesgo de seguridad, debido a que

todos los e-NB podrían ser sometidos a un ataque de DDOS (*Distributed Denial of Service*). Además para el propósito del tipo de tráfico X2 los e-NB vecinos deben verse, por lo tanto si todos están en un mismo dominio habría que crear listas de acceso de direcciones MAC, lo que hace al diseño poco escalable.

- La topología de E-LAN también dificulta el diseño en cuanto a la redundancia para los servicios ya se deben tomar en cuenta protocolos de capa 2.
- Para la implementación de un gateway de IPSec se recomienda que este localizado en la capa de pre-agregación o agregación, lo que demanda tener capa 3 sobre el diseño de capa2, lo cual dificulta aún mas la configuración.
- En estudio se encuentra consideraciones para mejorar el ancho de banda de consumo de los usuarios distribuyendo el gateway PDN, lo cual también tiene sus repercusiones en el diseño de capa 2 para proveer este servicio.
- En algunos casos se ha propuesto para mejorar la seguridad mecanismos como 802.1x, el cual tiene algunos inconvenientes ya que se trata de una solución con algunos saltos y dominios de broadcast que se conectan entre sí.
- Para estos servicios de E-LINE se han observado propuestas en modelos de conexión centralizada, lo que resulta en un denominado enrutamiento subóptimo afectando ciertos servicios como la VoIP o servicios de colaboración. Esta solución entonces rompe el esquema planteado por la 3GPP de tener una plena conectividad entre los e-NB.
- Una única VLAN dedicada para varios servicios en esta solución de capa 2 también representa un problema ya que la identificación de tráfico para el servicio respectivo complica el diseño a nivel del Core de servicios.

En base a lo presentado se puede concluir que esta solución no ha sido implementada a gran escala y por lo tanto cada vez menos considerada por los proveedores de servicio.

1.3.4 SOLUCIONES DE L3VPN[5]

A continuación se analiza una propuesta para proveer este *backhaul* y se encuentra representado en la figura 1.26, en donde se tiene una VLAN para los servicios de Core y otra para el tráfico X2. La VLAN de aplicación de Core debe tener conectividad hasta la capa de agregación para lo cual se tiene una MPLS VPN de capa 3, en muchos casos se ha visto que los equipos que se encuentran en los *cell-sites* son equipos que no tienen tantos recursos de memoria y procesamiento, para lo cual entre las soluciones propuestas existe un modelo denominado MPLS VPN *half-duplex*, ya que la idea es que para los servicios del EPC no es necesario que los *cell-sites* se vean entre sí, sino que todo tráfico de servicios pase por el sitio central que es el EPC.

Se tienen los siguientes aspectos a considerar:

- Este modelo presenta una arquitectura flexible y escalable, es decir sin mayores cambios se adapta a nuevos requerimientos, como por ejemplo seguridad con soporte de IPsec centralizado o distribuido, soluciones distribuidas de Core de servicios (SGW o PGW).
- Al tratarse de una solución MPLS VPN de capa 3, se puede identificar el tráfico y darle un tratamiento diferenciado, es decir calidad de servicio (QoS).

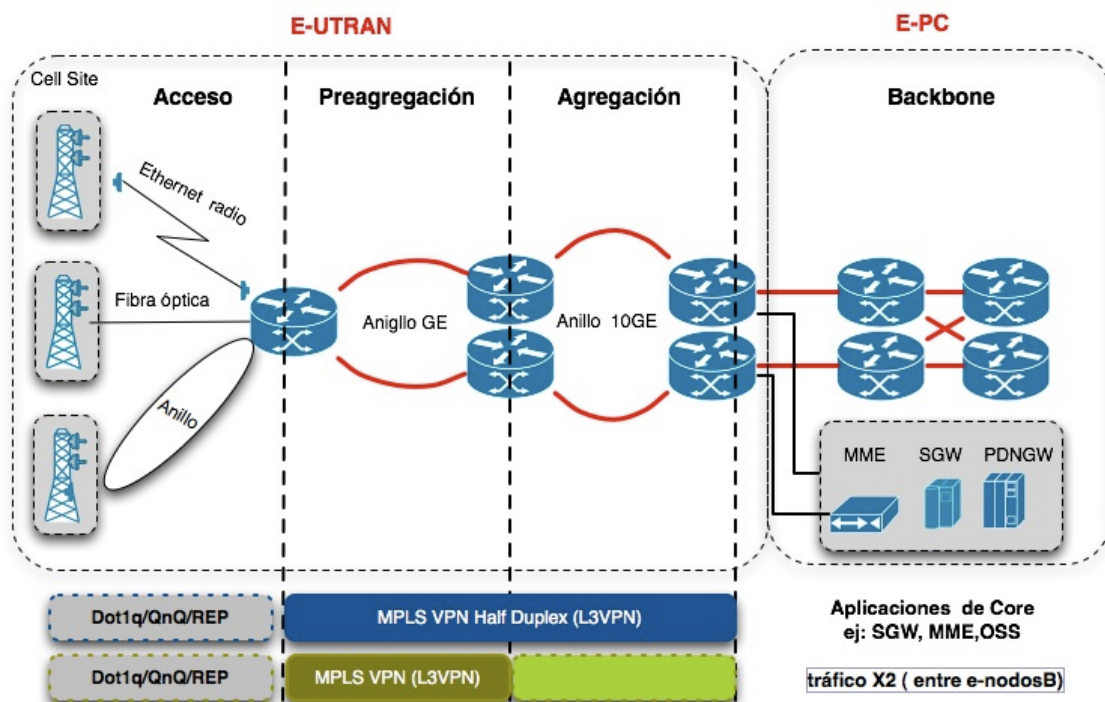


Figura 1.26 Modelo de MPLS L3VPN para interconectar LTE y EPC ^[5]

- Otra característica de esta solución es la resiliencia¹, lo que se consigue con configuraciones menos complicadas que las soluciones L2VPN².
- Esta solución no presenta enrutamiento subóptimo, ya que el tráfico de Core llega al Core y el tráfico de X2 se encuentran en comunicación directa.
- Aplicaciones colaborativas (MIMO) pueden exigir valores deseables de latencia, *jitter* y *delay*.

¹ La capacidad de continuar con el servicio a pesar de la presencia de eventos en la red.

² Al trabajar en Capa 2 existen siempre consideraciones de spanning tree y cualquier configuración debe evitar que se formen lazos ya que esto desencadena en inundación de tráfico y alto procesamiento innecesarios.

- Para implementaciones de IPsec con las soluciones de MPLS VPN *hub and spoke* se tendría enrutamiento subóptimo añadiendo un *delay* innecesario, aspecto que se debe tomar en cuenta.
- Existe una tendencia en los proveedores de servicios, que es extender el dominio MPLS lo más posible en el *backhaul*.

1.3.5 SOLUCIONES CON INTER-AS VPN

En implementaciones de redes MPLS de un número considerable de nodos, se presentan en los equipos una demanda de requerimientos de memoria y procesamiento, tanto para la tabla de enrutamiento producto del IGP, como de las rutas VPNv4 de las MPLS VPN L3; por lo que en el crecimiento de una red MPLS hace que, para cada nodo, el equipamiento deba cumplir con ciertas exigencias de hardware y software para ser considerado como PE de una nueva localidad, con el propósito de simplificar y por ende disminuir costos se han buscado alternativas¹, en la industria se para el caso de extender nuevos servicios como la telefonía móvil se ha tomado la decisión de separar la red MPLS con el objetivo de asegurar que el servicio este aislado del resto de servicios de la red y disminuir los requerimientos del nuevo equipamiento a instalarse. La solución de *backhaul* de la CNT E.P se basa en una nueva implementación dedicada a los servicios móviles dentro de un dominio diferente, es por eso que en esta sección se presentará el fundamento de lo que se

¹ Existe un draft en la IETF de lo que se considera como Seamless MPLS, esto está en la referencia [22]

conoce como Inter-AS o también conocido como Multi-AS *Backbone*, además se presentará un concepto fundamental de optimización denominado *Seamless MPLS*.

1.3.5.1 MPLS entre diferentes dominios o Inter-AS MPLS

Se tienen servicios entre varios sistemas autónomos o dominios MPLS diferentes¹, no necesariamente estos sistemas autónomos están relacionados como el mismo proveedor, por lo cual se presentan tres soluciones que permiten extender el servicio de extremo a extremo. Muchas veces el mismo proveedor decide dividir sus operaciones en varios sistemas autónomos para diversificar los servicios que entrega.

En el estándar RFC 4364 [25] se describen las opciones de solución a este tipo de conexiones y se ha denominado opción A, B y C. Los equipos de borde entre los sistemas autónomos son denominados ASBR (*Autonomous System Border Router*) y son los encargados de llevar la información necesaria para que los servicios pasen de un lado a otro.

1.3.5.1.1 Opción A: Conexiones VRF-to-VRF entre los ASBR

El método mas simple para la comunicación entre MPLS VPNs de dos sistemas autónomos diferentes es interconectar los equipos de borde ASBRs directamente a

¹ Esta sección se basa en el capítulo MULTI-AS BACKBONES de la referencia [14] y también fundamentado con la referencia [25]

través de varias subinterfaces, como se muestra en la figura 1.27 entre estas subinterfaces se establecen sesiones E-BGP; cada ASBR asocia un subinterface por cada VPN que requiere comunicación entre los dos sistemas autónomos, utilizando E-BGP para publicar las redes de la VRF en cada sistema autónomo.

Esta solución no requiere configurar MPLS en los equipos de borde, en lo que al router ASBR2 respecta, ve para cada VRF al router ASBR1 como que fuera un CE más. De la misma manera el router ASBR1 ve al router ASBR2 como un CE por cada VRF. El trabajar de esta manera con los ASBRs tienen las siguientes consecuencias:

- En los ASBRs se requiere una subinterface por cada VRF.
- Por razones de aprovisionamiento se requiere asignar los recursos de red suficientes para cada subinterface y su respectivo monitoreo.
- Cada ASBR requiere intercambiar la información de enrutamiento de cada VPN con los equipos del otro sistema autónomo.
- Por cada cliente VPN se hace necesario crear una nueva sesión E-BGP.

Debido a las características del servicio se puede inducir que, para un bajo número de clientes VPNs, con tablas de enrutamiento no muy grandes la solución es sencilla y beneficiosa. Entonces la opción A es fácil de entender e implementar, involucra a los PEs que contienen el servicio; además la interconexión entre los proveedores de servicio es simplemente una subinterface, donde se tiene también un punto de aplicaciones de las políticas del servicio (ancho de banda y filtros de acceso).

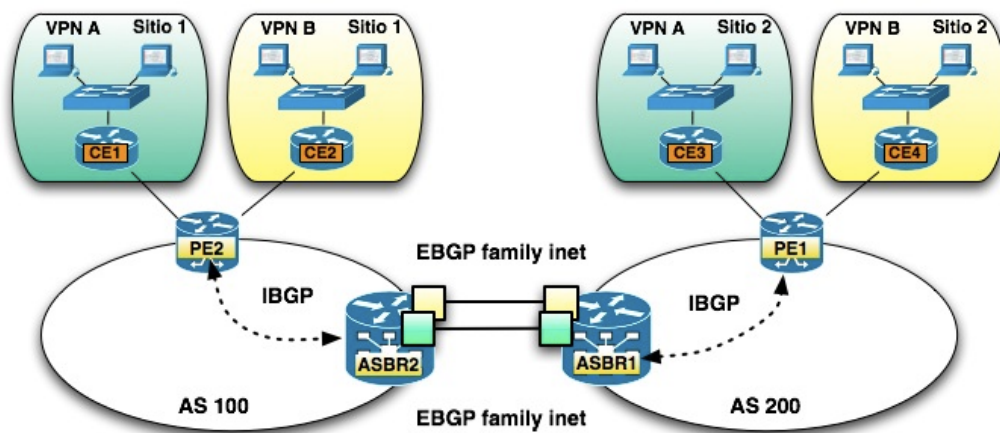


Figura 1.27 Inter-AS opción A: Conexión para a par de las VRFs en los ASBRs ^[14].

1.3.5.1.2 Opción B: EBGP redistribución de rutas VPNv4 etiquetadas.

En una red MPLS para activar un nuevo servicio de VPN basta con configurarlo en el PE de la respectiva localidad donde el cliente tiene sus oficinas. En la opción A, se presenta una situación que se vuelve no escalable, debido a que por cada nuevo servicio se hace necesario una nueva subinterface en los ASBR, las cuales son parte de la VPN del cliente y también se configura la VPN en su correspondiente PE. En la opción B se consigue evitar esta configuración repetitiva utilizando una sola sesión E-BGP entre los ASBRs sin importar la cantidad de VPNs que se requieran, en esta sesión E-BGP se propagan con etiquetas los prefijos VPNv4, los ASBR reciben estas rutas por medio una sesión IBGP con el PE respectivo en su AS y las intercambian entre los dos sistemas autónomos.

En la figura 1.28 se muestra el concepto de esta solución, además se recomienda que esta sesión EBGP entre los ASBRs sea asegurada en la configuración y evitar que se establezcan otras sesiones no autorizadas, de esta manera se consigue

proteger a las redes del cliente dentro de cada sistema autónomo. Tanto el PE1 y PE2 reciben las rutas de las VPNs a través de esta sesión E-BGP y al ser parte de dos dominios MPLS tienen el LSP definido hasta cada borde, es decir hasta el ASBR; sin embargo para que el tráfico entre el sitio 1 y el sitio 2 pueda pasar es necesario que se construya el camino de etiquetas que permita la conexión entre los dos sistemas autónomos; el protocolo que permite el transporte de las etiquetas entre los dos ASs es BGP y es lo que permitirá que se construya el LSP desde el sitio 1 al sitio2.

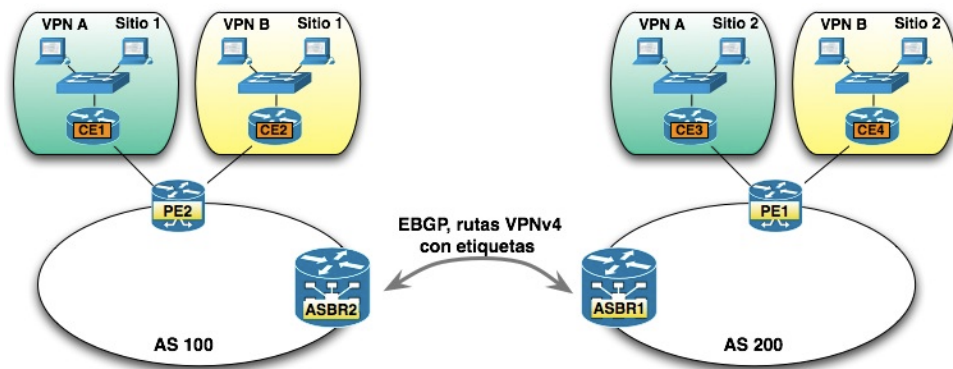


Figura 1.28 Opción B: Redistribución de rutas VPNv4 por E-BGP ^[14]

Como se indica en la figura 1.28, cuando el ASBR1 recibe una ruta VPNv4 del PE1, con la etiqueta L1, el ASBR1 le impone la etiqueta L2 para este prefijo y lo publica al ASBR2 como el mismo (ASBR1) como *next-hop*¹. Al mismo tiempo se instala en el plano de datos la equivalencia que permite el intercambio de etiquetas L1 a L2 o L2 a L1 según la dirección del paquete del cliente. El ASBR2 publica con etiqueta L3 las rutas de la VPN en su AS e instala en el plano de datos la equivalencia que permite

¹ *next-hop* es una funcionalidad de BGP que permite a un router de borde enseñar rutas, que están dentro de su red, a sus vecinos y que los mismos alcancen esas redes como que si se originaran en el equipo de borde, sin necesidad de conocer la topología del otro sistema autónomo.

el cambio entre L3 y L2, de la misma manera que se describió para el ASRB1. Cuando el tráfico es enviado desde el router PE2, para el prefijo de la VPN del cliente, este es etiquetado con L3, en el ASBR2 esta etiqueta se intercambia por la L2 y dentro del ASBR1 nuevamente se intercambia por la etiqueta L1 publicada por PE1. Nótese la diferencia de cuando se tiene todo en un mismo dominio donde la etiqueta de PE1 no hubiese cambiado.

Por seguridad el ASBR debería solo pasar tráfico con etiquetas que proviene de las interfaces por las cuales se encuentra publicando etiquetas, para evitar que un etiqueta enviada por otro equipo desvíe el tráfico correspondiente. Otra potencial falla de seguridad puede provenir de un peering no autorizado que puede capturar tráfico de una VPN en particular; en este caso los proveedores de servicio de los dos sistemas autónomos deben ponerse de acuerdo en que rutas intercambiar y que RTs utilizar.

En resumen esta solución presenta los siguientes puntos a destacar:

- No se requiere de subinterfaces ni recursos de red en la configuración para cada VPN en los ASBRs.
- Los ASBRs mantienen su estatus para todas las rutas de las L3VPN y L2VPN.
- No es sencillo a nivel de IP filtrar tráfico para los paquetes que cruzan los equipos de borde, por lo que no existen nada que actualizar el momento de crear nuevos servicios.
- Una única sesión EBGp es necesaria entre los dos sistemas autónomos.
- Esta solución requiere de un LSP entre el PE de origen y el PE de destino

1.3.5.1.3 Opción C: Redistribución de rutas etiquetadas VPNv4 por una sesión EBGp entre los sistemas autónomos, con redistribución de etiquetas IPv4 al respectivo AS

La opción B requiere que las rutas VPNv4 sean publicadas por los ASBRs, esto hace que la solución no sea escalable especialmente cuando existe una gran cantidad de rutas de VPNs. La opción C previene este problema, como se muestra en la figura 1.29 el cliente utiliza una sesión *multihop*¹ entre PE1 y PE2 para transportar los prefijos VPNv4 con etiquetas y los proveedores de servicio proporcionan conectividad hacia las direcciones de las interfaces *loopbacks* de los PEs publicando estas redes como prefijos IPv4 con etiquetas del un sistema autónomo al otro; de esta manera los routers ASBRs no deben transportar ninguna ruta de las VPNS. De las tres opciones ésta constituye la más escalable.

Se debe tomar en cuenta que no se debe cambiar el *next-hop*² de las rutas VPNv4 que se publican en las sesiones EBGp *multihop*. En la interface de ingreso de la VPN se impondrán tres etiquetas, a menos que las direcciones de los PE sean conocidas por los routers P en cada AS, en cuyo caso sólo se impondrán dos etiquetas en pila.

¹ Por lo general las sesiones BGP se establecen entre equipos directamente conectados, pero puede existir una conexión entre equipos que no lo están, a este tipo de sesiones se las conoce como *multihop*.

² Es parte de una ruta BGP y significa que equipo ha originado la ruta.

En cuanto a la seguridad un nuevo escenario se añade, el hecho que direcciones IP que pertenecen a un AS se redistribuyan en el otro, revelando de esta manera el plan de direccionamiento.

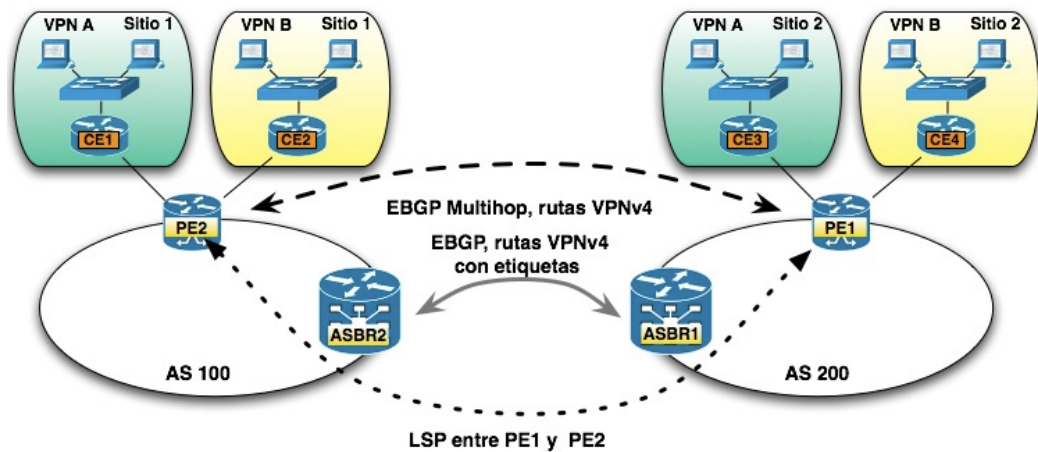


Figura 1.29 Opción C: Redistribución de etiquetas VPNv4 entre AS de origen y destino, con redistribución de etiquetas de IPV4 con E-BGP entre sistemas autónomos ^[14]

1.3.5.2 Seamless MPLS ^[14]

Hasta hace poco tiempo, la industria ha trabajado con redes MPLS en el Core y en en el acceso redes metro-ethernet. Debido al crecimiento de la demanda y la posibilidad de expandir la oferta de nuevos servicios, se han visto varias ventajas el extender el dominio del MPLS hasta la capa de acceso y de esta manera evitar inconvenientes que se pueden presentar en el acceso ethernet. Se analiza a continuación desventajas que las redes de acceso presentan y que serían superadas con un acceso *Seamless MPLS*.

1.3.5.2.1 Redes de Acceso Ethernet

En un inicio las redes de acceso IP de los proveedores de servicio entregaban la conectividad a sus clientes utilizando un *switch* ethernet de la misma manera que se lo hace para soluciones empresariales con sus respectivos inconvenientes como son los descritos a continuación:

1. **Inundación de Tráfico** (o *Flooding*). Ethernet utiliza el mecanismo de descubrimiento de direcciones *MAC* basado en el plano de datos. Si una trama ethernet tiene una *MAC address* que desconoce, entonces inunda de tráfico por todos los puertos del *switch*, a excepción del puerto de origen del requerimiento de tráfico. Este tipo de tráfico en ocasiones puede volverse un problema al incrementarse la cantidad de clientes y saturar los enlaces troncales.
2. **Aprendizaje de direcciones MAC**. Las direcciones *MAC* de los clientes deben almacenarse en la memoria de los equipos de acceso, por cambios en la topología de la red los equipos switch deben de tiempo en tiempo borrar sus tablas de direcciones y volverlas a aprender, esto hace que los tiempos de convergencia del servicio se vea afectado.
3. **Lazos de inundación de tráfico**. Los efectos pueden ser muy destructivos debido a que Ethernet no tiene un mecanismo como IP el campo de TTL (*Time-To-Live*). Lo que significa que un lazo de inundación de tráfico puede permanecer indefinidamente, con secuencias de saturación de enlaces

troncales y alto procesamiento de los equipos afectados con sus consecuente interrupción del servicio. Para solventar esto se diseñó el protocolo STP (*Spanning Tree Protocol*) que dentro de los enlaces redundantes crea caminos libres de lazos, lamentablemente este protocolo no es infalible y existen casos como por ejemplo que un enlace entre dos *switch* se interrumpa de un solo lado de la comunicación provocando que el protocolo deje de actuar y se formen los lazos, en cuyo caso la solución suele ser el apagado de varias interfaces, con la consecuente interrupción de servicios.

4. **Tiempos de Convergencia.** Al utilizar el protocolo STP implica tiempos de convergencia altos (decenas de segundos), si bien es cierto se tienen ya implementaciones con RSTP (*Rapid Spanning Tree Protocol*) que involucra mejores tiempos de convergencia, en el orden de segundos, siguen siendo altos para aplicaciones como video.
5. **Balanceo de carga.** Con STP se tienen enlaces redundantes que no tienen tráfico y que no pueden ser utilizados, para una implementación local no es de mucha importancia, pero un proveedor de servicio estos enlaces pueden ser de considerables distancia, complejidad y costo.
6. **Ingeniería de tráfico y control de Ancho de banda.** En equipos como los *switch* ethernet no existe un método para controlar el ancho de banda de un usuario y el único esquema para diferenciar tipos de tráfico es el método

basado en los valores de los bits IEEE 802.1p que se encuentran en la cabecera de la trama ethernet.

7. **Operación y Mantenimiento.** Ethernet trabaja en el inicio con métodos de descubrimiento de direcciones para crear la tabla de MAC address, lo que puede causar inconvenientes el momento de la operación. El momento de realizar el seguimiento a un evento de un cliente en particular el operador realiza un análisis del plano de control y luego del plano de datos en los switch, teniendo muchas veces que reiniciar puertos o los equipos involucrados con la consecuente pérdida de servicio para los otros clientes.

En cambio al acercar lo mas posible el uso de MPLS hacia el cliente final trae ventajas como las siguientes:

1. MPLS es una tecnología desarrollada para convergencia de redes en el Core del proveedor de servicios debido a su cumplimiento de los requerimientos que esta convergencia demanda y que se indican a continuación.
 - a. Un común plano de control y un plano de datos escalable para una gran variedad de tipos de tráfico.
 - b. Para un determinado tipo de tráfico, baja tasa de pérdida de paquetes en caso de fallas en la transmisión.
 - c. La capacidad de elegir el camino dentro de la red por donde se requiere pasar un servicio, así como también la posibilidad de balancear el tráfico por diferentes caminos.

- d. Para ciertos tipos de tráfico se puede reservar ancho de banda para ser utilizados en caso de congestión, protegiendo el SLA del cliente.
2. Mantener un mismo esquema en el acceso y en el Core facilita la operación, debido a la unificación de procedimientos. Además se vuelve más flexible la interconexión entre el Core de la red y los accesos.

En este modelo de trabajo los equipos *switch* MPLS son ethernet y se utilizan para interconectar los nodos de acceso y ciertos clientes directamente. En algunos casos ciertos nodos tiene la funcionalidad de MPLS, esto hace que los enlaces entre estos equipos y la red MPLS sean de gran importancia, con velocidades de transmisión oportunas (1,10,40 y 100Gbps); de esta manera entonces ethernet se vuelve una tecnología de conexión y no una red completa de transporte, lo que se resume en la tabla 1.2

Se considera que los *switch* que tienen MPLS son más caros que los *switch* que se tiene en los accesos tradicionales, pero se tienen las siguientes características y que los *switch* ethernet no poseen:

- Soporta una amplia cantidad de tipos de interfaces para la conexión.
- Se puede atender una cantidad considerable de solicitudes de búsqueda de un IP nueva a velocidad alta ¹.

¹ En la industria se tiene el término *wire rate* para las operaciones que son ejecutadas por los procesadores ASIC de los puertos de switch y son lo más rápido que puede ese equipo ejecutar una tarea.

- Capacidad de ofrecer servicios tales como *firewall* y NAT (*Network Address Translation*).
- Alta disponibilidad con *hardware* redundante.
- Plano de control extremo a extremo en el tráfico.

Switch Ethernet		MPLS
Esquema de transporte	MAC address	Etiquetas MPLS
Plano de Control	Limitado plano de control: Protocolo <i>Spanning Tree</i> para crear caminos libre de lazos. Pero no hay plano control para descubrimiento de direcciones.	IGP y RSVP, con LDP o BGP
Tecnología de enlace	Ethernet	Ethernet

Tabla 1.2 Comparación de transporte de *Switch Ethernet* con redes de acceso MPLS ^[14]

1.3.5.2 Terminología General de redes de acceso MPLS

Antes de continuar se requiere de la definición de cierta terminología que ayudará a describir de mejor manera el concepto de MPLS en redes de acceso.

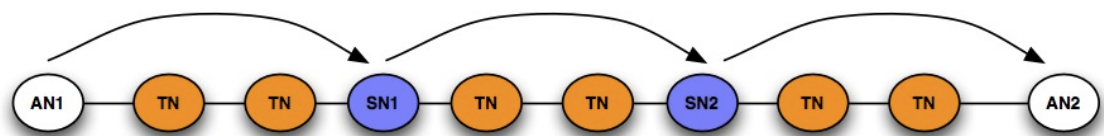


Figura 1.30 Terminología General de una red de Acceso ^[14]

En la figura 1.30 se pueden observar los siguientes elementos:

- AN (*Access Network*), es un término genérico en el cual se incluyen varios equipos de acceso, tales como DSLAMs (*Digital Subscriber Line Access*

Multiplexer), MSAN(*Multiservice Access Node*), GPON(*Gigabit-capable Passive Optical Network*), etc.

- SN (*Service Node*) es un término genérico que describe a varios tipos de nodo, en los cuales se aprovisionan servicios a los clientes, como por ejemplo:
 - a. En el caso de clientes DSL el SN puede ser la conexión de ancho de banda a un BRAS (Broadband Remote Access Server). En algunos casos el cliente puede tener más servicios, como Voz, datos o video.
 - b. En el caso de servicios MPLS VPN a clientes corporativos el SN está en el PE router, por ejemplo un L3VPN, las rutas contenidas en la VRF del cliente son parte del SN, así como también las políticas de ancho de banda
- TN (*Transport Node*) se designa este termino a un nodo MPLS que sirve de conexión entre un AN y un SN.

En la figura 1.30 se puede observar el camino que sigue un paquete de datos, desde su ingreso en la red en el AN1, hacia un SN principal, SN1 por medio de uno o varios TNs; SN1 relaciona el paquete con su contexto apropiado, por ejemplo en el caso de un L3VPN el SN1 se determina el nodo de servicio de destino SN2, el cual es un servicio específico para el caso del ejemplo puede ser un servicio de verificación de IP destino (*IP lookup*), luego de lo cual el servicio SN2 determina que el paquete debe salir por el nodo de acceso AN2.

Existen dos métodos de implementar MPLS en el Core y en el acceso, el primero, representado en la figura 1.31 donde se tiene una red MPLS en Core y otra red

MPLS en el acceso y su comunicación se denomina servicios de Inter-AS¹, con los métodos que se detallaron en la sección 1.3.5 y el segundo se lo conoce como *Seamless* MPLS o MPLS unificado que quiere decir sin bordes o fronteras, el cual se ha representado en la figura 1.32 y se lo estudiará mas adelante. En la figura 1.31 se tiene el ejemplo de un método de interconexión de dos redes MPLS con dominios diferentes, es decir Inter-AS opción A. Para el caso de la figura 1.32 se tiene una red unificada o *Seamless* MPLS, donde se ejemplifica la facilidad de reubicación de los servicios de la red.

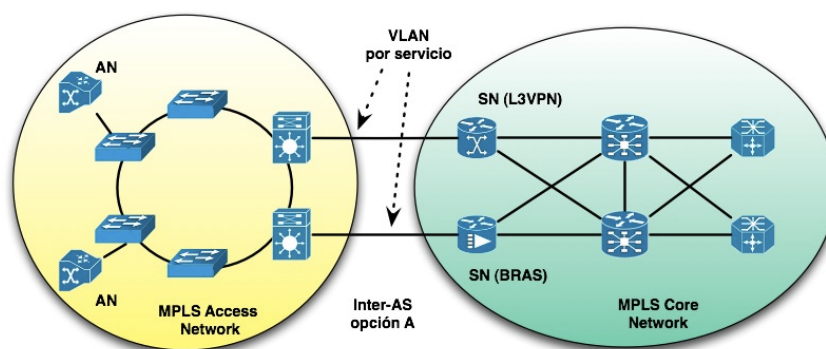


Figura 1.31 Ejemplo de servicio implementado en una Red MPLS de CORE Y Red MPLS en el Acceso (Inter-AS) ^[14]

1.3.5.3 Mecanismos de transporte con Seamless MPLS

El sistema *Seamless* provee de una arquitectura sobre la cual se pueden dar servicios de una manera escalable² a diferencia de la solución con Inter-AS, se trata

¹ En este ejemplo se ilustra la opción A, como uno de los métodos de implementar Inter-AS.

² Del Capítulo 4 de la referencia [26] se ha tomado los conceptos de los modelos de Seamless MPLS, además se ha complementado también con una actualización de Cisco Systems cuya referencia es la [27]

de un solo dominio unificado de MPLS en el cual no se tiene una separación entre el Core y el acceso, se puede presentar entre la red de Core y la red de acceso diferentes áreas de su IGP pero desde el punto de vista de la red MPLS, el tráfico de un cliente es transportado de un extremo a otro de un sólo dominio.

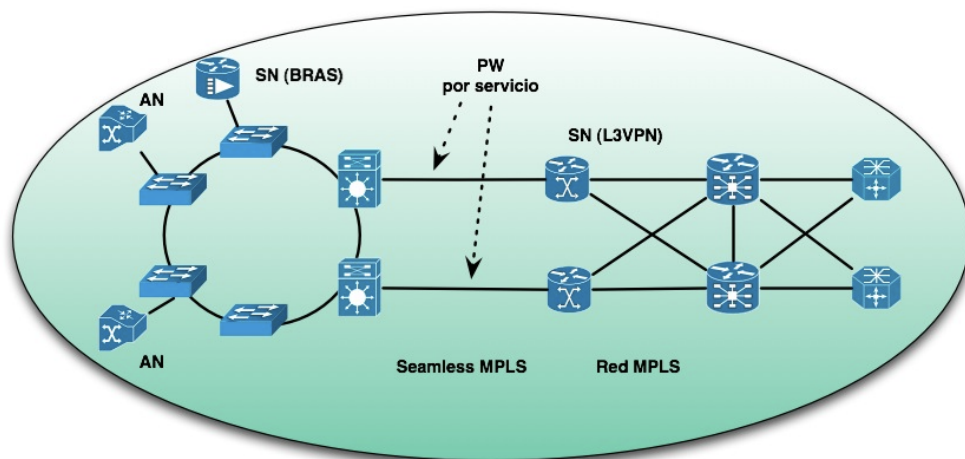


Figura 1.32 Ejemplo de un servicio implementado en una Red MPLS unificada o Seamless MPLS ^[14]

Antes de continuar con el análisis de los modelos de transporte propuestos en la industria de telecomunicaciones, es necesario introducir una terminología que ayudará a comprender de manera más sencilla componentes que se harán comunes en las soluciones y cuyas funciones o configuraciones harán que se vayan adaptando al tipo de transporte estudiado; Algunos de los términos ya se definieron en la sección anterior, pero se los volverá a ver sobre todo para tener una mejor idea de la definición de la solución *Seamless MPLS*.

Como se puede ver en la figura 1.33, en una red con *Seamless MPLS* se pueden identificar los siguientes componentes, para un servicio de extremo a extremo.

- AN (*Access Node*) es un término genérico en el cual se incluyen varios equipos de acceso, tales como DSLAMs (*Digital Subscriber Line Access Multiplexer*), MSAN (*Multiservice Access Node*), GPON(*Gigabit-capable Passive Optical Network*), CSGs (*Cell Site Gateway*¹) etc.
- SN (*Service Node*) es un término genérico que describe a varios tipos de nodo, en los cuales se aprovisionan servicios a los clientes, como por ejemplo: servicios L2VPN, L3VPN, BNGs (*Broadband Network Gateways*²), servidores de video, controladores de radio bases 3G conocido por las siglas RNC (*Radio Network Controller*), servidores de telefonía (*media gateways*), servicios de Core celular, etc.
- TN (*Transport Node*) se designa este término a un nodo MPLS que sirve de conexión entre un AN y un SN.
- BN (*Border Node*) son equipos que facilitan la interconexión entre regiones diferentes, también se los suele nombrar como ABR (*Area Border Router*).
- SH (*Service Helper*) habilita o hace escalable el plano de control de los servicios. SH no transporta datos del cliente final. Ejemplos de SH puede ser el servicio de *route-reflector*, control de políticas de ancho de banda del servicio, servicios como RADIUS y AAA.

Un equipo que puede tener mas de un rol como los descritos, eso va a depender de la implementación y de los recursos del proveedor de servicios.

¹ Este termino se añade ya que es el tipo de acceso en la red IP-RAN, también se lo denomina *cellsite*

² Este equipo también es conocido como BRAS (*Broadband Remote Access Server*)

Existen en la industria varios modelos de transporte con el sistema *Seamless MPLS*, en la referencia [27] se tiene una clasificación de estos modelos para redes pequeñas o para redes grandes¹, lo que se resume en la tabla 1.3.

Redes pequeñas	Redes grandes
Modelo 1.1 LSP sin segmentación en la red de Core y Agregación.	Modelo 2.1 Segmentación del IGP de la red de Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP
Modelo 1.2 Segmentación del IGP de la red de Acceso y unificación con el Core y Agregación mediante el LSP con etiquetas a través de BGP de manera jerárquica.	Modelo 2.2 Segmentación del IGP de la red de Acceso, Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP.
Modelo 1.3 Redistribución de etiquetas de BGP en el IGP o LDP de la red de Acceso (opcional LDP DoD ²)	Modelo 2.3 Redistribución de etiquetas de BGP en el IGP en la red de Acceso

Tabla 1.3 Modelos de red de transporte con *Seamless MPLS* [28]

Los modelos que a continuación se describen son soluciones que el fabricante ya a probado previamente en sus laboratorios y se han adaptado a una generalización de los casos particulares de redes implementadas en las empresas proveedoras de servicios de telecomunicaciones. Por lo cual se debe analizar el estado actual de la red MPLS que un proveedor tiene para adoptar la solución a su medida y condiciones de crecimiento.

¹ Una red se puede considera grande si tiene amplia cobertura (disponibilidad de quipos de acceso) en una zona geográfica extensa y cuenta por esta condición de un core de transporte con una topología de IGP que permita la segmentación para propósitos del *Seamless MPLS*.

² Actualmente lo del protocolo LDP DoD esta ya en una RFC 7032

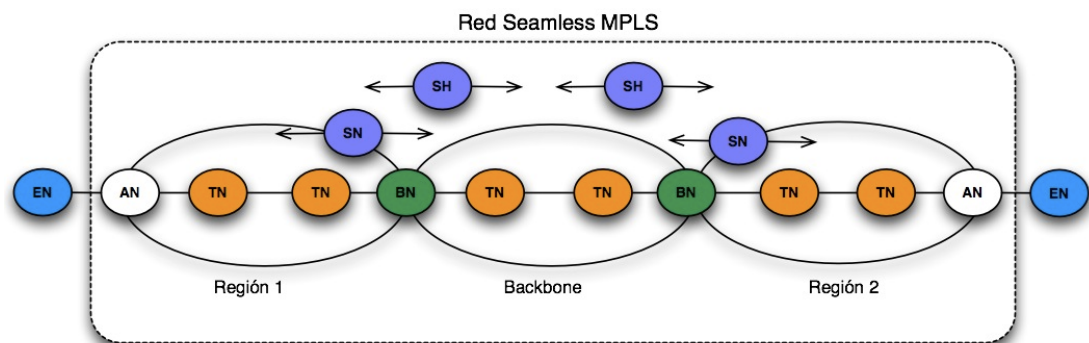


Figura 1.33 Arquitectura de *Seamless MPLS* [28]

1.3.5.3.1 LDP sin segmentación en la red de Core y Agregación

Para redes con un área geográfica no muy extensa como se muestra en la figura 1.33 y con implementaciones que permitan tener en el acceso MPLS directamente , se puede contar con un PE dedicado a prestar servicios específicos; por lo que no se requiere realizar trabajos adicionales para conseguir el objetivo de llegar más cerca al usuario con MPLS¹. Entonces se provisionan los servicios requeridos de L2VPN y L3VPN como se los trató en la sección 1.3.3 y 1.3.4, se debe considerar que dependiendo de la tecnología del proveedor de servicio puede tener como última milla ethernet con sus respectivas conversiones a microonda, TDM o SDH para servicios tanto móviles como fijos.

¹ Reflexión sobre el concepto de la referencia [27] ya que se requiere afirmar que el objetivo del Seamless es llegar con MPLS hasta donde sea posible en el acceso.

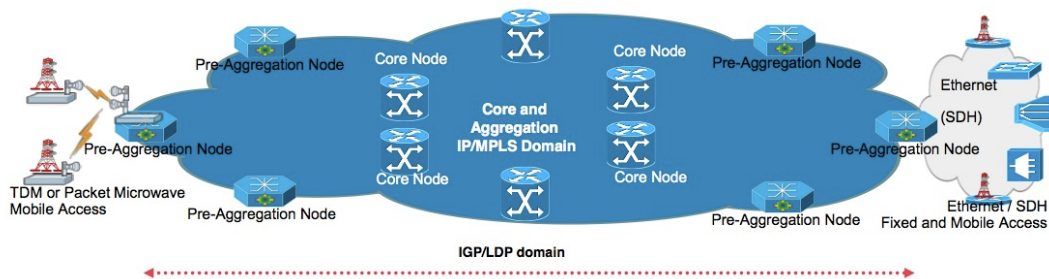


Figura 1.34 LSP sin segmentación en el Core y agregación [27]

1.3.5.3.2 Segmentación del IGP de la red de Acceso y unificación de LSP con el Core y Agregación, mediante etiquetas a través de BGP

Con el crecimiento de la red en cuanto a nodos y principalmente con el aumento de suscriptores o usuarios, se tiene un importante incremento de la tabla de enrutamiento tanto en el IGP así como también en las rutas VPNV4, lo que tiene consecuencias en el desempeño de la red ya que se tiene mayor exigencia de los PEs, los que enfrentan falta de capacidad de recursos, tanto de hardware como de red. Tomando en cuenta lo expuesto y debido a que se trata todavía de una red no muy grande se plantea la segmentación de la red de acceso final, para esta técnica se crea otra área del IGP¹ u otro proceso de enrutamiento con la finalidad de disminuir la tablas de enrutamiento y por ende minimizar los requerimientos de la red.

¹ Por lo general otra área cuando se trata de OSPF, en el caso de IS-IS se tiene L2 en la red de core y para el acceso se configura L1, siendo los equipos de agregación L1/L2 simultaneamente.

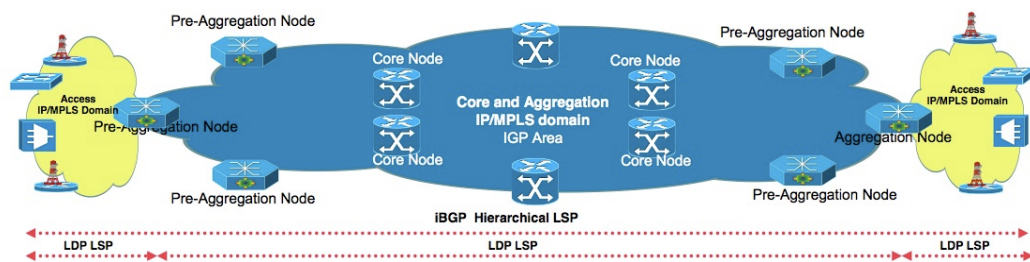


Figura 1.35 Segmentación del LSP y acceso mediante etiquetas por BGP ^[27]

En la figura 1.35 se puede ver la red de acceso en otra área de IGP, en cuanto a los servicios MPLS dependen de que el protocolo LDP termine el camino de extremo a extremo o LSP. Para resolver la interrupción del LSP provocada por la división del IGP, se utiliza BGP más *label*¹, por lo que en los PE de pre agregación (BN o ABR) se levantan sesiones con los ANs del segmento con iBGP en el *address-family* *VPNv4* y para enviar etiquetas por iBGP, además se agrega con cuidado los filtros necesarios para aprender exclusivamente las rutas de interés de los servicios requeridos.

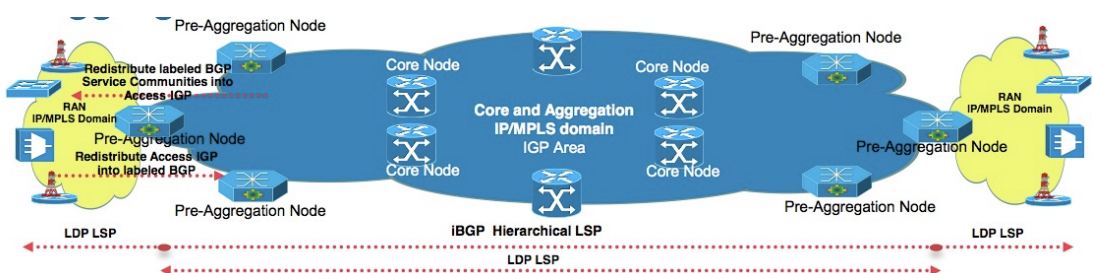


Figura 1.36 Redistribución de etiquetas de BGP en el IGP o LDP de la red de Acceso (opcional LDP DoD) ^[27]

¹ Esto se encuentra descrito en el RFC 3107

1.3.5.3.3 Redistribución de etiquetas de BGP en el IGP de la red de Acceso (opcional LDP DoD¹)

Se trata de una variación a la solución planteada en la sección anterior, como se ve en la figura 1.36 nuevamente los equipos de pre agregación hacen el papel de BN, hacia el lado de Core se interconecta con todos los servicios, pero hacia el lado del acceso tiene sesiones iBGP distribuyendo etiquetas para completar el LSP y redistribuye utilizando el atributo de *community* los servicios de interés al IGP/LDP. Alternativamente, se puede configurar la red de acceso y el BN para trabajar con LDP DoD². Cuando se utiliza estos nodos como un *backhaul* móvil y al mismo tiempo se requiere provisionar servicios fijos, se recomienda el uso de LDP DoD.

1.3.5.3.4 Segmentación del IGP de la red de Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP

En este caso como se puede ver en la figura 1.37, la red de acceso está directamente conectada a puertos de la red de agregación. La red de agregación se encuentra en otro segmento de IGP, en consecuencia se requiere completar el LSP mediante el protocolo de BGP con etiquetas, como ya se ha tratado con el RFC 3107.

¹ Actualmente lo del protocolo LDP DoD esta ya en una RFC 7032

² Siempre y cuando los equipos del segmento y el BN soporten el RFC 7032

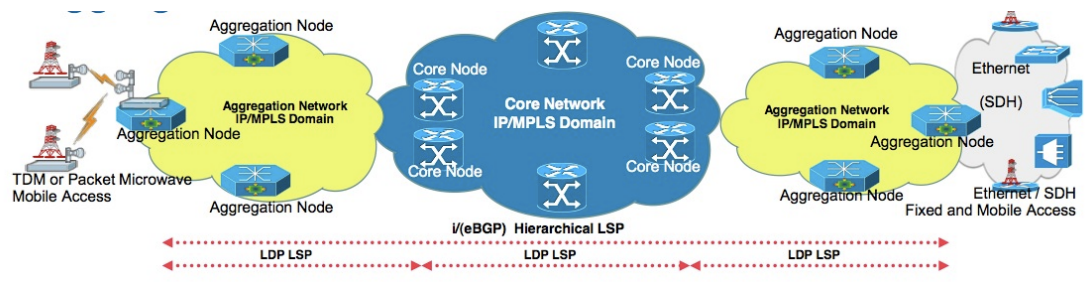


Figura 1.37 Segmentación del IGP de la red de Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP^[27]

1.3.5.3.5 Segmentación del IGP de la red de Acceso, Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP

En este caso se tiene, tal como se muestra en la figura 1.38 que la red de Acceso es MPLS y se encuentra en un segmento de IGP diferente a la red de agregación y la red de Core se encuentra en otro segmento de IGP. El LSP es entonces formado por segmentos, en una primera instancia los equipos de Core toman el rol de BN de la red de agregación para mantener sesiones iBGP¹ con etiquetas, de la misma forma los equipos de agregación actúan como BN para la red de Acceso. Se controla mediante filtros en los equipos BN la información de etiquetas necesaria para los servicios que se entregan en el acceso, además se asegura que estos recursos también pasen por los filtros configurados en la red de agregación.

¹ En la referencia [27] se indica que también se puede tener la red de Agregación en otro sistema autónomo, pero aquí ya se trataría de inter-as y no de seamless MPLS.

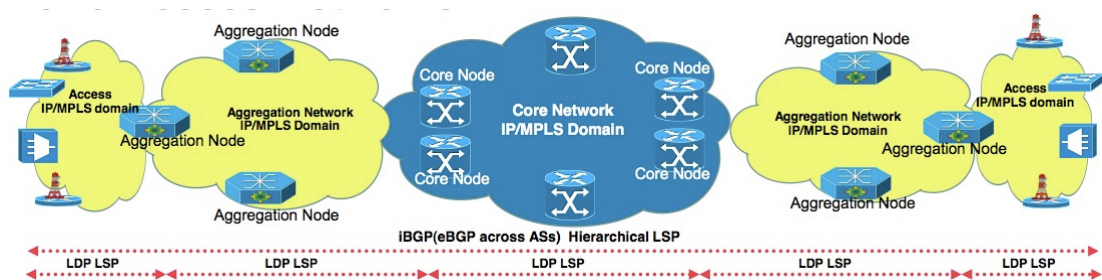


Figura 1.38 Segmentación del IGP de la red de Agregación y unificación de LSP con el Core, mediante etiquetas a través de BGP ^[27]

1.3.5.3.6 Redistribución de etiquetas de BGP en el IGP al interior de la red de Acceso (opcional LDP DoD)

Como se puede ver en la figura 1.39, se tiene la segmentación del IGP en el acceso, Core y agregación, para completar la información de etiquetas se utiliza sesiones iBGP entre los BN que son los equipos de Core y los de pre-agregación. Se redistribuye utilizando el atributo de community los servicios de interés al IGP/LDP de la red de acceso desde la red de agregación. Alternativamente, se puede configurar la red de acceso y el BN para trabajar con LDP DoD¹.

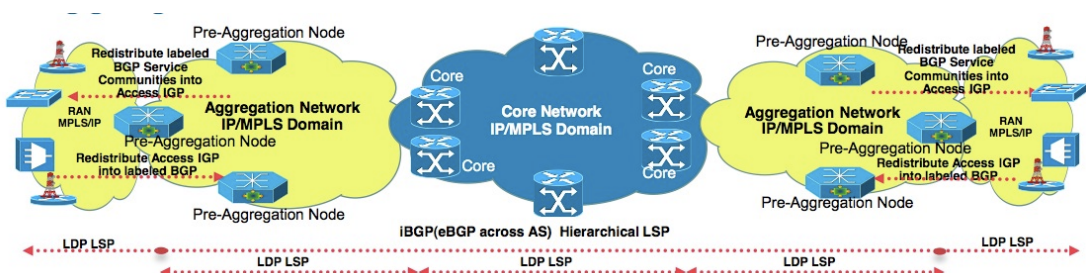


Figura 1.39 Redistribución de etiquetas de BGP en el IGP o LDP (opcional LDP DoD) en la red de Acceso desde la red de Agregación ^[27]

¹ Siempre y cuando los equipos del segmento y el BN soporten el RFC 7032

1.3.5.4 Comparación de las soluciones de Inter-AS y Seamless MPLS

Con el objetivo¹ de analizar ventajas y desventajas del *Seamless* MPLS sobre el Inter-AS, se ha tomado el servicio de L3VPN desde el nodo de acceso al nodo de servicio utilizando L2VPN como transporte, tal como se indican en la figuras 1.41 para Inter-AS y 1.42. para *Seamless* MPLS.

Una ventaja del Inter-AS es que si la red de acceso y la de core se encuentran administradas por departamentos diferentes, al tener MPLS en el acceso mediante su propio sistema autónomo, puede entonces seguir administrado de esta manera. La desventaja se presenta en la falta de flexibilidad de reubicación de servicios (nodos SN), por ejemplo en la figura 1.41 se tienen servicios masivos como clientes de GPON masivo o DSL masivo que tiene el servicio de PPPoE por medio de un BNG o BRAS ubicado en el Core; por el incremento de demanda se puede requerir que el servicio de BNG se distribuya a la red de acceso y esta red no esté en capacidad de brindar estos servicios.

Otra diferencia entre Inter-AS y *Seamless* se puede ver en la manera de entregar o provisionar el servicio de un lado a otro, por ejemplo para el caso de la figura 1.41 el servicio se debe entregar al SN en varias VLANs desde el TN3, mientras como se puede ver en la figura 1.42, el servicio en *Seamless* MPLS puede ser terminado

¹ La comparación de la referencia [14] se estima de importancia para fines de la reingeniería que se planteará mas adelante.

directamente en el nodo de servicio, a continuación se presenta un análisis del tráfico desde el AN hasta el SN.

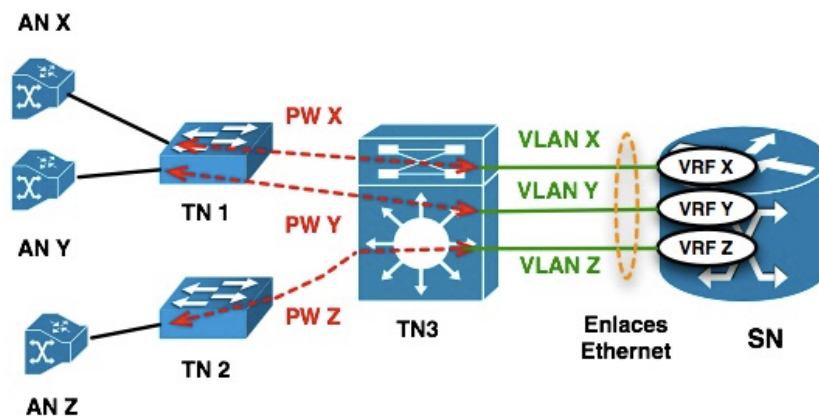


Figura 1.40 Servicios de L3VPN sobre pseudowires terminados en el nodo de transporte TN3, solución Inter-AS ^[14]

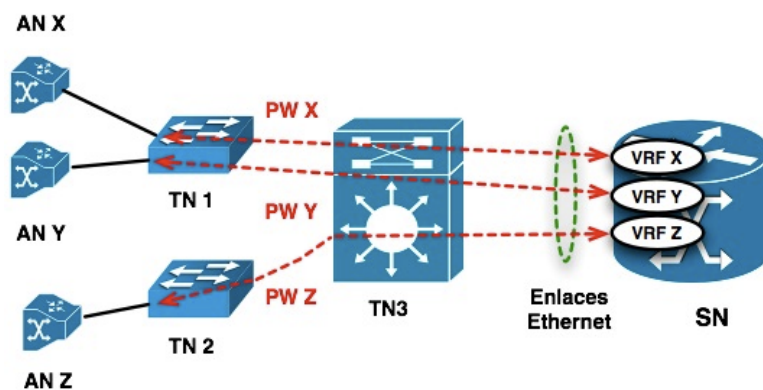


Figura 1.41 Servicios de L3VPN sobre pseudowires terminados en el nodo de servicio SN, solución Seamless MPLS ^[14]

En la figura 1.41 se puede ver que no existe una conectividad considerada como MPLS entre el nodo SN y la red de acceso. Por lo tanto el PE donde reside el servicio es el nodo TN3, los circuitos terminan en una VLAN entre TN3 y el SN, y desde el punto de vista del pseudowire el AN es un CE. En el SN, la VLAN es atada al servicio que se desea entregar, por ejemplo una L3VPN, una VPLS, una L2VPN o un servicio de control del equipo de usuario. En la figura 1.41 se tiene tres clientes X, Y y Z, los

cuales tienen L3VPN como servicio, sus correspondientes nodos de acceso son AN X, AN Y y AN Z. El SN es responsable por proveer el servicio de L3VPN, la VRF respectiva para cada usuario está configurada en el SN. Los AN X y AN Y se encuentran conectados al nodo de transporte TN1, AN Z se encuentra conectado al nodo TN2. El *pseudowire*¹ PW X opera entre TN1 y TN3, en donde el servicio se agrega en la VLAN X del TN3, esta VLAN X finalmente termina el servicio en la VRF X del PE SN; de manera similar el PW Y y PW Z terminan en las VRF Y y VRF Z.

En la figura 1.42, es decir en *Seamless MPLS*, la función de PE para los *pseudowires* está en el SN, de esta manera el TN3 actúa, en MPLS como un router P. Esto significa que no existe una interfaz física para los *pseudowires* en el SN, estos *pseudowires* se provisionan directamente con L3VPN VRF, una L2VPN o VPLS según la necesidad del usuario. En la figura 1.42 se tiene PW X, PW Y y PW Z que van directamente al router SN en el servicio de VRF sin requerir de una interfaz VLAN entre TN3 y SN. De la misma manera que los *pseudowires* se han provisionado en las VRFs se puede hacer para otros servicios.

Una ventaja del Inter-AS es que su configuración es muy familiar ya que se acerca a como siempre se ha venido tradicionalmente aprovisionando en las redes MPLS en los PEs, empleando con vista hacia el cliente con VLAN (y en algún momento con ATM PVC o DLCI de *frame relay*). En cambio cuando el MPLS es configurado en el acceso (*Seamless MPLS*) en lugar del tradicional acceso metro-ethernet, no existen cambios en el router PE, este tipo de implementación proporciona un punto de

¹ El *pseudowire* es el servicio de L2VPN que interconecta dos puntos, en esta parte solo lo ha nombrado como PW X para el cliente X.

demarcación mas claro entre la red de core y la red de acceso. Por otro lado *Seamless* MPLS tiene la ventaja, como en el ejemplo que en TN3 no se requiere configuración alguna ya que éste actúa como un P de la red; en contraste con el Inter-AS, que cada nuevo servicio que se desea aprovisionar, una nueva VLAN debe ser creada tanto en TN3 como en el router SN y su *pseudowire* hacia el nodo AN correspondiente¹.

Muchas operadoras implementan soluciones inicialmente con Inter-AS, por la facilidad de la operación tradicional y posteriormente migran a *Seamless* MPLS, la ventaja de que las dos sean tecnologías MPLS facilitan el trabajo de cambio de configuración en ventanas de mantenimiento.

¹ En este caso la referencia [14] sólo esta tomando en cuenta la opción A del inter-AS para ejemplificar las ventajas y desventajas, pero también es cierto que en las otras opciones también se deben realizar configuraciones adicionales cada que se incrementa un nuevo servicio. El propósito en este caso es didáctico.

CAPÍTULO 2

ANÁLISIS DE LA RED ACTUAL IP- RAN DE LA CNT E.P.

EN QUITO

2.1 DESCRIPCIÓN DE LA RED IP- RAN DE LA CNT E.P. EN QUITO

CNT E.P. con el objetivo de brindar servicios LTE en la ciudad de Quito, adquirió un sistema de comunicaciones, que permite interconectar a nuevos usuarios con equipos terminales LTE a su core de servicios de datos, telefonía e internet.

La figura 2.1 muestra la jerarquía de todo el sistema LTE, el diseño se divide en 3 capas:

- Acceso: dicha capa consiste de los equipos CSG (*Cell Site Gateway*), los cuales permiten la conexión de los nodos móviles de acceso NodoB y eNB.
- Agregación: consiste en los equipos ASG (*Agregation Site Gateway*), los cuales agregan o concentran el tráfico proveniente de los nodos de acceso.
- Core: consiste en los equipos RSG (*Radio Site Gateway*), los cuales permiten integrar la comunicación desde y hacia el acceso con sus respectivas plataformas de gestión y servidores externos de datos. Actualmente las plataformas de Core se encuentran conectadas en la red MPLS CISCO existente.

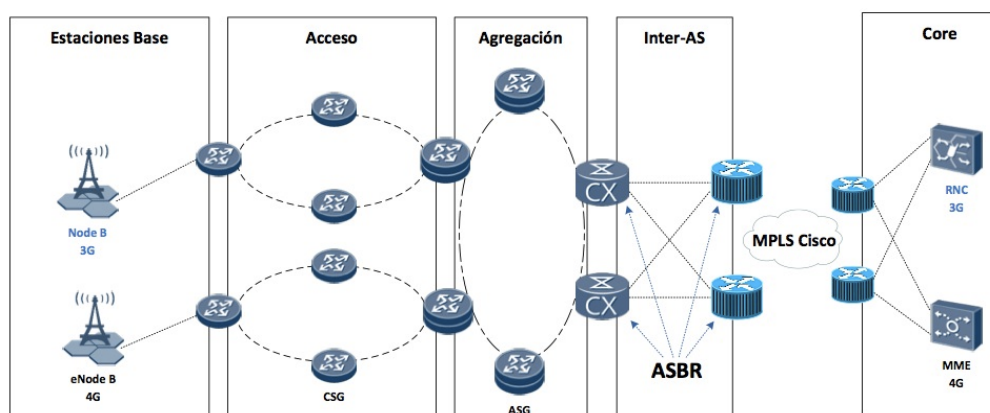


Figura 2.1 Esquema general de la red de *backhaul* de CNT ^[13]

En la figura 2.2 se muestra en detalle la ingeniería de diseño e implementación del Core IP RAN para el despliegue de la red de acceso móvil 4G LTE para el operador CNT EP; en este esquema se omite representar la agregación de los nodos finales o eNBs, debido a que se trata de información confidencial del giro del negocio de CNT E.P.

Para la interconexión con la red MPLS existente se tiene los equipos de borde tanto de la red IP RAN como de la red MPLS de CNT, tal como se representa en la figura 2.3. Estos equipos de borde se encuentran trabajando en lo que se denomina Inter-AS opción C, que se describe en la sección 1.3.5, para lo cual se tienen dos equipos CX-600 como borde de la red IP RAN y dos equipos ASR903 como bordes de la red MPLS existente.

El protocolo de enrutamiento interno o IGP (*Interior Gateway Protocol*) implementado en la solución es ISIS, por razones de escalabilidad y además para guardar coherencia con lo que se encuentra operando en la red MPLS a nivel nacional.

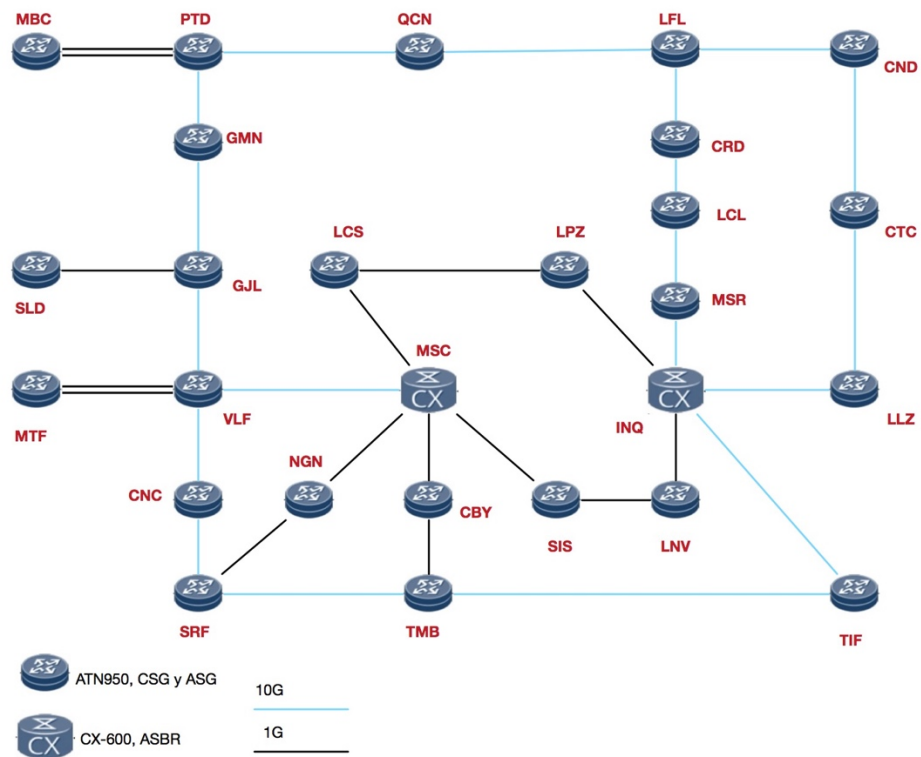


Figura 2.2 Arquitectura física de la red IP-RAN ^[38]

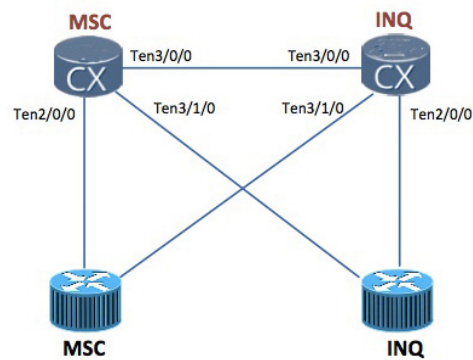


Figura 2.3 Equipos de borde de la red IP RAN y la red MPLS existente ^[13]

Con el fin de establecer la implementación de manera jerárquica, IS-IS fue diseñado en un esquema Multi-Proceso, el cual consiste en configurar cada anillo en un proceso distinto de IS-IS (como se puede ver en la figura 2.2); los puntos comunes a los anillos son los routers de borde de las dos plataformas, por lo tanto dichos equipos

deben formar parte de todos los procesos IS-IS existentes con el fin de cerrar los anillos a nivel de IGP.

Los equipos ASBR disponen de una interfaz física entre ellos, y sabiendo que ambos equipos deben cerrar los anillos de IS-IS, se configurarán sub-interfaces como se representa en la figura 2.5 cada una corresponderá a un proceso de IS-IS.

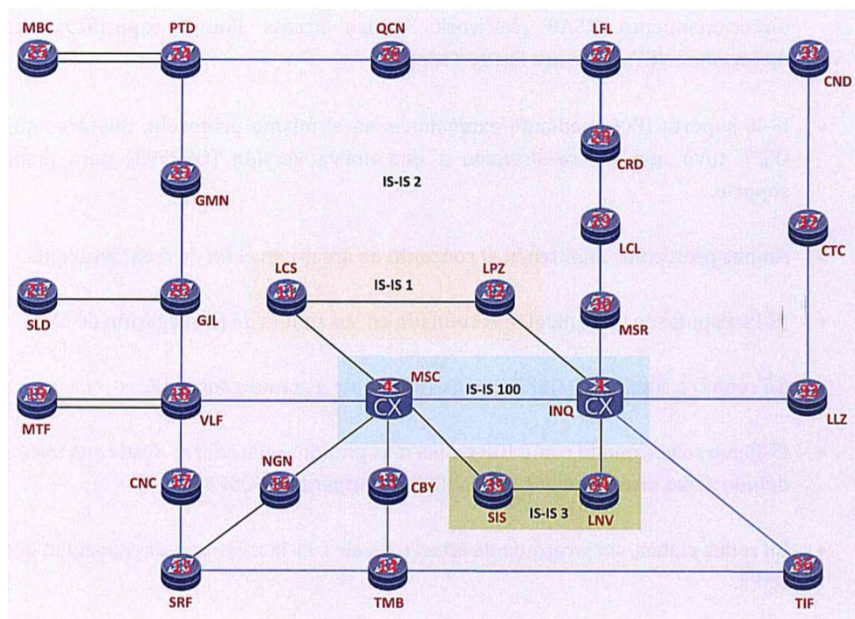


Figura 2.4 Segmentación del dominio IS-IS Core IP RAN varios procesos ^[38]

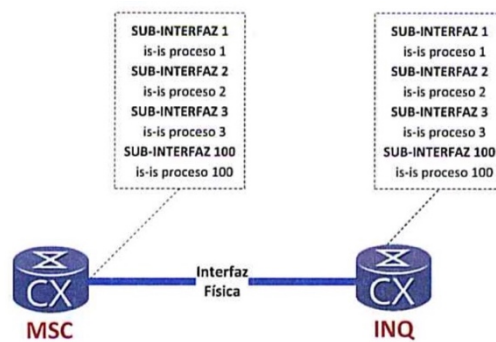


Figura 2.5 ASBR y sus subinterfaces para los procesos de IS-IS ^[38]

CNT EP ha provisto el segmento 10.20.A.B/16 para ser utilizado en el direccionamiento global del Core IP RAN; el direccionamiento provisto es asignado a las interfaces que conectan hacia otros equipos del Core, ejemplo CX-ATN y de igual manera a las interfaces *Loopback* las mismas que servirán como Router-ID de los protocolos que lo requieran. Con el fin de establecer un esquema de direccionamiento escalable, entendible y jerárquico, se han considerado los siguientes puntos:

- Enlaces punto a punto directamente conectados (P2P) tendrán longitud de prefijo /30.
- Interfaces *Loopback* tendrán longitud de prefijo /32. Dichas interfaces serán las “Loop100” en todos los equipos y en base a ello el esquema de direccionamiento es 10.20.X.100/32

Tercer Octeto (ID)	Equipos
1-10	ASBR
11-100	ASG
101-255	CSG

Tabla 2.1 Definición del tercer octeto en la interface loopback 100 ^[13]

2.1.1 CONFIGURACIÓN DEL IS-IS

En vista de que IS-IS maneja la identificación de router utilizando una dirección tipo NET, se debe definir previamente dicho direccionamiento para todos los equipos de la red. El formato básico de una dirección NET se muestra en la tabla 2.2:

AFI <-- 1 byte -->	AREA ID <-- 0-12 bytes -->	SYS ID <-- 6 bytes -->	SEL <-- 1 byte -->
-----------------------	-------------------------------	---------------------------	-----------------------

Tabla 2.2 Formato de la dirección NET del protocolo IS-IS ^[13]

Cada campo debe ser establecido en base a las siguientes reglas:

- AFI (*Authority and Format Identifier*), se usa el valor 49 para indicar que la dirección es local y privada (similar al rango privado de IPv4).
- AREA ID, para identificar a un grupo de routers bajo una misma área o nivel, es un campo opcional.
- SYS ID, para identificar al equipo dentro de un mismo dominio; existen 2 opciones para establecer este valor, mediante MAC-ADDRESS o a su vez mapeando los 32 bits de una dirección IPv4; la segunda opción es la más recomendada.
- SEL (NSAP Selector), generalmente se establece en 0x00.

El direccionamiento NET que actualmente se utiliza en las redes de Core de CNT EP, obedece al formato de la tabla 2.3:

AFI	AREA ID	SYS ID	SEL
49	0008	representación de la Ipv4 Loopback 100	00

Tabla 2.3 Representación del convenio para la dirección NET en los equipos de core ^[13]

Para el Core IP RAN, se sugiere utilizar como AREA ID el valor 0034 (3G y 4G) por lo tanto el formato sería el siguiente:

AFI	AREA ID	SYS ID	SEL
49	0034	representación de la Ipv4 Loopback 100	00

Tabla 2.4 Representación del convenio para las direcciones NET en los equipos core IP-RAN ^[13]

2.1.2 CONFIGURACIÓN DE MPLS LDP

Para el transporte de servicios sobre el core IP RAN, dos protocolos son analizados: LDP y RSVP-TE, a continuación se presentan sus principales características y diferencias, en la tabla 2-5.

Debido a los requerimientos de crecimiento dinámico y escalable se tomó la decisión de utilizar LDP en lugar de RSVP-TE.

Ítem	LDP	RSVP	Resumen
Implementación General	Configuración fácil, no existe necesidad de túneles	Configuración compleja, se requiere de túneles	LDP tiene mayor ventaja que RSVP-TE
Servicios	Tráfico X2 sigue el camino más corto	Tráfico X2 no sigue el camino más corto.	El tráfico X2 representa aproximadamente entre el 3% y 5% del tráfico S1
Nivel de protección	Bajo, depende absolutamente del IGP	Alto, incluye mecanismos propios de protección.	RSVP-TE ofrece mayor protección
Escalabilidad	Dinámica y simple	Manual y Compleja	LDP escala de mejor manera
Control de <i>path</i>	Bajo, depende del IGP	Alto, mayor control	RSVP-TE ofrece mayor control en el <i>path</i>

Tabla 2.5 LDP vs RSVP-TE

2.1.3 CONFIGURACIÓN BGP

Dentro del Core IP RAN, se utiliza BGP internamente con el fin de compartir información de prefijos en las distintas VPN de servicios (3G y 4G), adicionalmente para el intercambio de prefijos fuera del dominio IP RAN con la red MPLS existente, lo que comúnmente se conoce como Inter-AS VPN.

Con el fin de evitar un *full-mesh* y a la vez brindar un esquema jerárquico en BGP, los equipos ASBR serán habilitados como RR (*route-reflector*) en el *address-family* VPNv4 dentro del Core IP RAN, lo que implica que cada equipo de acceso CSG levantará sesiones MP- iBGP únicamente contra ambos RR (como se muestra en la figura 2.6). En un futuro, en caso de requerirse IPv6, se deberán levantar las mismas sesiones pero en el *address-family* VPNv6. Con el fin de mantener tablas de BGP y routing pequeñas en los CSG, los RR únicamente advertirán el prefijo *default* “0.0.0.0/0” en la VPNv4 ; los clientes de RR, es decir los CSG en cambio advierten todos los prefijos que se consideren necesarios hacia los RR; únicamente los RR

tendrán conocimiento de todos los prefijos VPNv4 de todo el Core IP RAN. Aplicando entonces esta configuración de RR a los tipos de tráfico de LTE S1 y X2, se tiene un comportamiento como el representado en la figura 2.7.

- Tráfico S1, conexión VPNv4 desde eNB hacia EPC (S1-MME es *control plane* y S1-SGW para *user-plane*), el tráfico se transporta desde cada CSG hacia los ASBR.
- Tráfico X2, conexión entre eNB para *handover*, el tráfico se transporta desde el CSG hacia el anillo de ASG para finalmente llegar al otro CSG. El volumen de tráfico X2 no es significativo.

En la figura 2.8 se representa el diseño de balanceo de carga de BGP o “*load-balance*”, se utiliza el atributo *Local-Preference* (LP) de BGP separando a los CSG en 2 grupos, un grupo tendrá a RR1 como *next-hop* BGP principal, el otro grupo tendrá a RR2 como *next-hop* BGP principal, ambos RR estarán de respaldo en ambos grupos cuando no operen como principal.

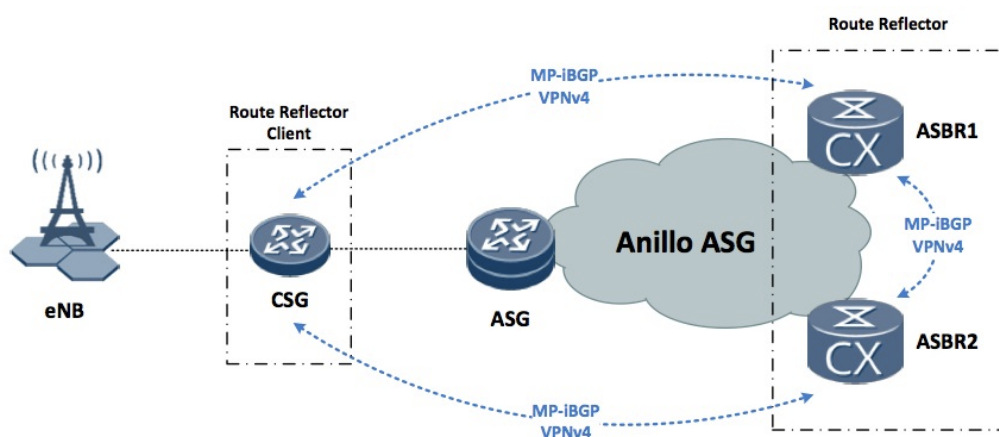


Figura 2.6 Arquitectura de Route Reflector ^[1-3]

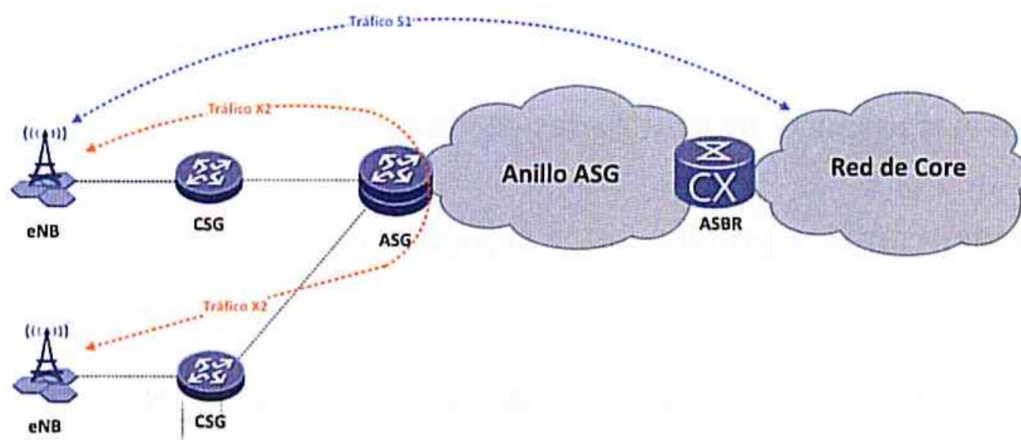


Figura 2.7 Flujo de tráfico en interfaces S1 y X2 de LTE ^[38]

Para la configuración de BGP el sistema autónomo asignado por CNT EP es 64986, bajo estas configuraciones el balanceo de carga se reparte en dos grupos, el grupo “RR Cliente 1” recibe dos prefijos *default*, uno por cada RR; ASBR1 es el principal de dicho grupo por lo tanto envía el prefijo con LP200, mientras que ASBR2 por ser respaldo envía con 100. El grupo “RR Cliente 2” recibe de igual manera dos prefijos *default*; ASBR2 es el principal de dicho grupo por lo tanto envía el prefijo con LP200, mientras que ASBR1 por ser respaldo envía con 100. Los grupos en BGP se dividen de la siguiente manera como se muestra en la figura 2.8.

2.1.4 CONFIGURACIÓN DE MPLS VPN

Los servicios a transportar sobre el Core IP RAN son 3G y 4G, ambos están configurados dentro de instancias VPNv4 o comúnmente conocidas como VRF. Las instancias de VPNv4 para LTE configuradas se resumen en la tabla 2.6.

Descripción	Nombre VRF	Subredes	VLAN CSG	RD	RT
Servicios LTE (EPC)	serlte	10.41.80.0/22	11	64986:605001	64986:605001
Gestión O&M LTE	geslte	10.64.20.0/22	10	64986:602003	64986:602003

Tabla 2.6 MPLS VPN para gestión y servicios LTE ^[13]

Los equipos CSG únicamente reciben la ruta VPNv4 *default* “0.0.0.0/0” desde cada ASBR a través de MP-iBGP puesto que dichos equipos se encuentran establecidos como *route-reflector*, mientras que los ASBR conocerán en detalle todos los prefijos por cada CSG.

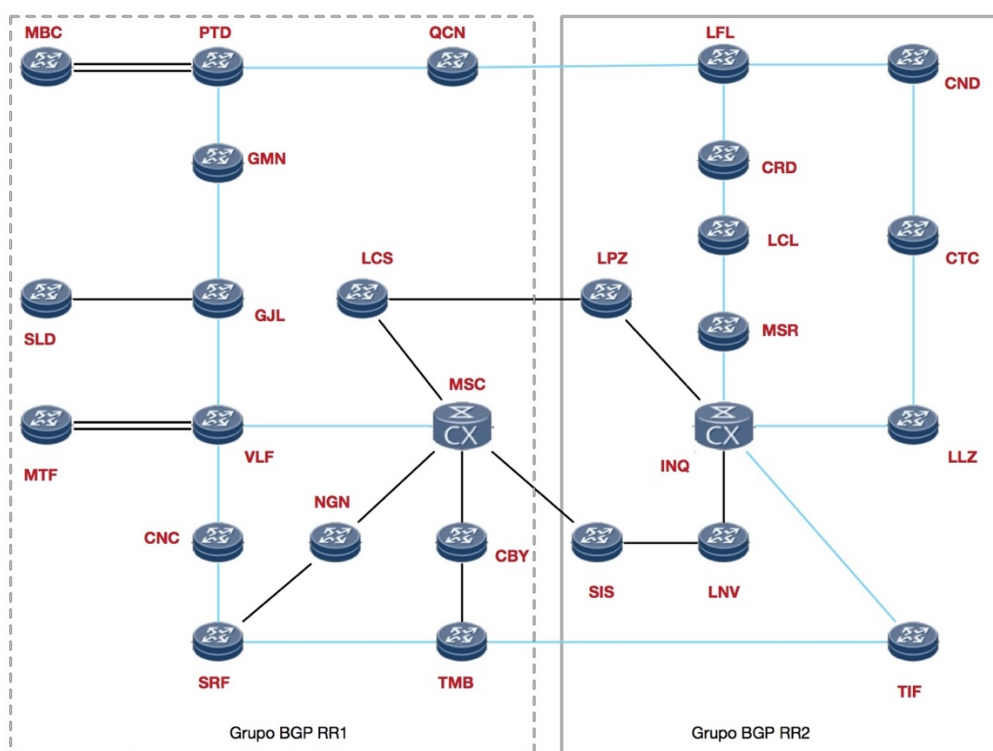


Figura 2.8 División de grupos BGP para los RR (Route Reflectors) ^[38]

2.1.5 CONFIGURACIÓN DE INTER-AS VPN

Con el fin de proveer comunicación a una VPN o servicio que cruza varios sistemas autónomos, se requiere establecer un escenario Inter-AS entre los distintos dominios por los cuales atraviesa la VPN. La técnica Inter-AS más viable para conectar ambos dominios MPLS es la opción B, como se muestra en la figura 2.9, la que se describe en la sección 1.3.5.

Los paquetes IP provenientes del e-NB, son convertidos en prefijos VPNv4 por el CSG y a su vez se transportan sobre un LSP establecido mediante RSVP-TE hacia el next-hop ASBR1, quien a su vez advierte los prefijos sobre la sesión MP-eBGP hacia el dominio MPLS existente y transporta los paquetes sobre el LSP establecido entre ambos dominios. El ASBR2 recibe los paquetes y los transporta sobre un LSP establecido bien sea por LDP o RSVP-TE hasta el next-hop PE en el cual se encuentran directamente conectados a las plataformas de Core Móvil o

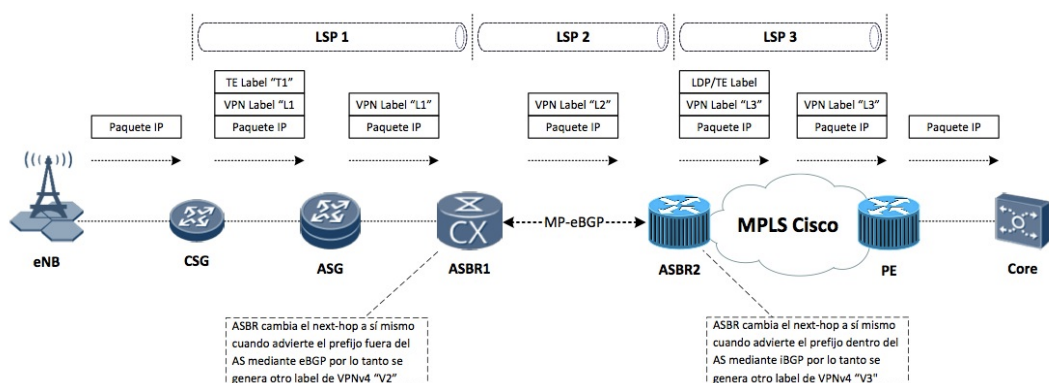


Figura 2.9 Inter-AS Opción B en Core IP RAN ^[13]

EPC. En el sentido inverso, es decir para los paquetes que fluyen del EPC hacia el e-NB, los prefijos de la EPC se advierten por VPNv4 dentro del AS de la IP-RAN mediante eBGP desde el ASBR2; debido a que los equipos CSG son clientes de los *route-reflectors* (ASBR del IP-RAN) conocen entonces el LSP o camino de comunicación desde el EPC hacia los e-NB.

Los siguientes son los servicios con los que cada e-nodo B debe tener conectividad, por lo tanto son los prefijos que deben ser advertidos por MP-eBGP desde la red MPLS existente:

- Red MME (maneja *control plane* y señalización de eNB): 10.52.1.25/32
- Red S-GW (maneja *data plane* de equipo terminal): 10.52.1.33/32

A continuación se detallan los prefijos correspondientes a servicios y O&M de eNB que el Core IP RAN y que deben ser alcanzables desde el core de servicios, que se encuentra en la red MPLS existente

- Servicios: 10.41.80.0/20
- O&M: 10.64.128.0/20

La topología¹ del Inter-AS entre ambos dominios MPLS se representa en la figura 2.10. Para evitar que el AS 64986 del Core IP RAN se convierta en tránsito del 28006 por los múltiples enlaces existentes entre ambos dominios, únicamente serán advertidos los prefijos originados localmente en el IP RAN mediante un filtro de AS-Path "*empty*" (^\$). Se crean entonces filtros en las sesiones eBGP permiten la salida de los prefijos únicamente de las VRFs de servicio y gestión y el ingreso de las rutas

¹ Se tiene la versión 1.7 del documento de diseño de la Solución para CNT en la cual se han introducido algunos cambios que se reflejan en la presente tesis, estos se denotan con la referencia [38], nuevamente se indica que este documento es confidencial.

por default de estas VRFs, también se crea una configuración de BGP de tal manera que se prefieren las rutas originadas en el enlace principal en lugar del que se aprende por enlace de respaldo (*backup*).

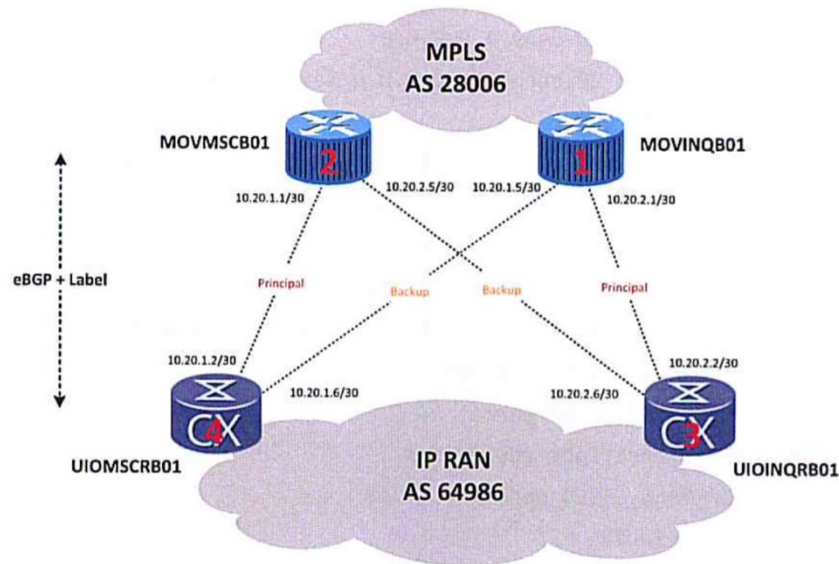


Figura 2.10 Topología Inter-AS VPN de las Solución de CNT
[38]

En ambos ASBRS se suman las redes de servicios y gestión con el fin de ser advertidas sobre el Inter-AS. Adicionalmente se habilita la cualidad de BGP denominada “*Dampening*” para minimizar el impacto en los servicios en caso de presentarse intermitencias sobre los prefijos recibidos en el Inter-AS.

Para seleccionar el enlace principal desde cada ASBR hacia la red MPLS CISCO, se utiliza un atributo de BGP propietario de Huawei y que tiene el significado local en el equipo (*apply preferred-value*), su nombre es PrefVal (similar al atributo Weight de CISCO).

En vista que el número de prefijos tiene un impacto directo en la utilización de los recursos físicos de los routers se ha limitado el número de prefijos en las sesiones

eBGP del Inter-AS para evitar que los equipos *cell-site* que son los más limitados puedan caer en un estado que afecte la entrega de los servicios ya mencionados.

BGP envía de manera periódica paquetes tipo *Keepalive* con el fin de detectar y confirmar el estado de sus vecinos (*neighbors*), este mecanismo se encuentra en el rango de segundos, lo que puede afectar directamente al tráfico de servicios. Habilitando a BGP como cliente del protocolo BFD (*Bidirectional Forwarding Detection*), se optimiza el tiempo de detección en el orden de los 50ms lo que permite la rápida convergencia de prefijos y así minimizar la pérdida de tráfico.

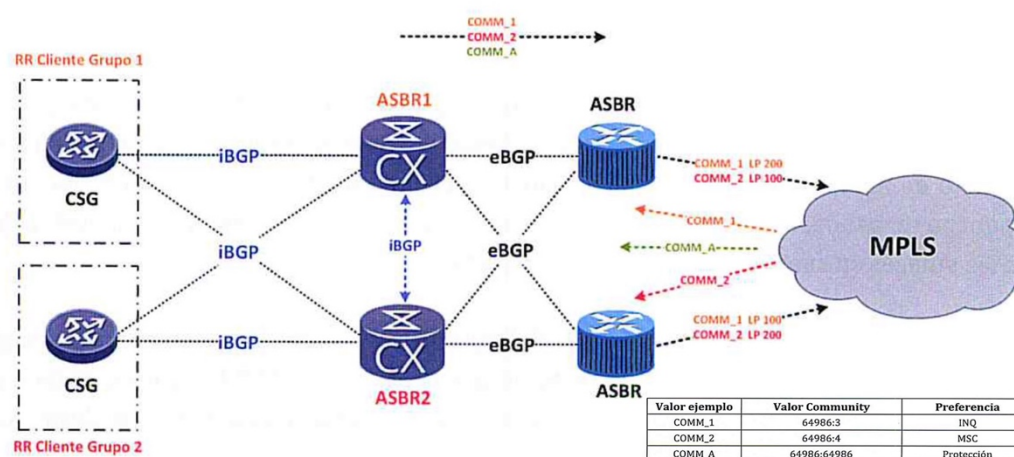


Figura 2.11 Balanceo de carga en el tráfico del Inter-AS [38]

Para el balanceo de carga en el sentido red MPLS CISCO hacia la red IP RAN, se utilizan atributos de community agregados a los prefijos advertidos sobre el Inter-AS para que la red MPLS pueda configurar el atributo Local-Preference dependiendo del valor de community recibido según se muestra en la figura 2.11, cada prefijo llevará dos communities, uno correspondiente al nodo principal y uno de protección, para conseguir distribuir el tráfico y obtener la redundancia en caso de fallos.

2.1.6 CONFIGURACIÓN DE QoS

La calidad de servicio o QoS es una herramienta que permite asegurar que el tráfico marcado como prioridad alta no sea descartado en caso de existir congestión en la red, minimizando el impacto de manera temporal; se debe tomar en cuenta que QoS no solventa problemas de ancho de banda, la única solución a la congestión continua es ampliar la capacidad de los enlaces en cuanto a ancho de banda se refiere.

Para el despliegue de QoS, es importante definir una zona de confianza en la cual los *tags* de QoS serán confiables; en base a la topología de la figura 2.12 y considerando el sentido de tráfico de Acceso hacia Core, la zona de confianza será establecida desde los eNB hasta los ASBR:

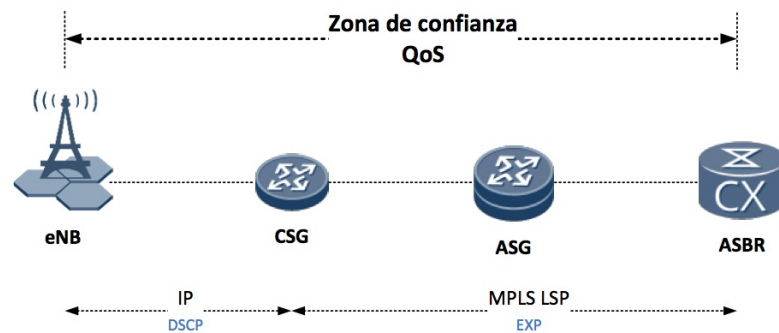


Figura 2.12 Definición de la zona de confianza de QoS en el Core IP RAN ^[13]

Para el flujo de tráfico inverso, proveniente de las plataformas de Core, los ASBR realizan la clasificación respectiva de los paquetes que se marcan con *tags* de QoS en base al diseño establecido dentro del Core IP RAN. En la tabla 2.7 se muestra los requerimientos de QoS en los servicios LTE y el respectivo tipo de encolamiento a usar en cada clase:

Los e-NB envían sus paquetes con los *tags* de DSCP mientras que los CSG mapean dichos valores a su correspondiente valor MPLS EXP. La tabla 2.8 muestra el mapeo entre valores DSCP y EXP por defecto

Servicio	Sub-servicio	QCI	DSCP (CoS)	EXP	QoS
Wireless	Voz en tiempo real	1	46(EF)	5	PQ
	Voz y video, ambos en tiempo real	2	26(AF31)	3	PQ+WFQ
	Juegos en tiempo real	3	34(AF41)	4	PQ+WFQ
	Datos en tiempo real	4	26(AF31)	3	PQ+WFQ
	Señalización IMS	5	46(EF)	6	PQ
	Video que no demanda tiempo real	6	18(AF21)	2	PQ+WFQ
	Voz, video y juegos que no demandan tiempo real	7	18(AF21)	2	PQ+WFQ
Control	-	STCP	46(EF)	5	PQ
O&M	-	MML	46(EF)	5	PQ
	-	FTP	10(AF11)	1	PQ+WFQ
IP Clock	-	-	46(EF)	5	PQ

Tabla 2.7 Requerimientos de QoS en LTE ^[13]

DSCP	CoS	EXP
46	EF	5
26	AF3	3
34	AF4	4
18	AF2	2
10	AF1	1
0	BE	0

Tabla 2.8 Correspondencia DSCP , CoS y EXP ^[13]

Los equipos Huawei disponen de 8 colas en sus interfaces, cada cola corresponde a una clase y en base a ello reciben la prioridad; en la tabla 2.9 se puede ver la distribución de dichas clases y sus respectivas prioridades por defecto asignadas en cuanto a la técnica de encolamiento se refiere:

Clases	Tipo Encolamiento	Valor (weight)
BE	WQF	10
AF1	WQF	10
AF2	WQF	10
AF3	WQF	15
AF4	WQF	15
EF	PQ	N/A
CS6	PQ	N/A
CS7	PQ	N/A

Tabla 2.9 Técnica de encolamiento por clases ^[13]

Para las clases de mayor prioridad (EF, CS6 y CS7), los paquetes son enviados a las colas PQ (*Priority Queuing*), las mismas que son atendidas con prioridad en caso de existir congestión y así disminuir *delay*, *jitter* y minimizar el impacto de la pérdida de paquetes; PQ clasifica las colas en niveles de mayor a menor prioridad, los paquetes de la cola de mayor prioridad serán enviados hasta que dicha cola se encuentre vacía, una vez vacía se da paso a la siguiente cola de menor prioridad y así sucesivamente hasta llegar a la última. En las clases establecidas para PQ no se limita al tráfico en términos de ancho de banda ni tampoco se aplican políticas de descarte.

Para el resto de clases, se utiliza WFQ (*Weighted Fair Queuing*) el cual asigna un peso a cada clase lo cual corresponde a un porcentaje de ancho de banda.

Inicialmente Huawei Technologies recomienda no aplicar políticas de servicio a las clases, es decir, limitar sus colas en porcentajes de *bandwidth*, especialmente a las colas PQ, debido a que dicho tráfico por su alta prioridad no debe ser limitado ya que puede afectar directamente a la señalización y servicios como la voz en casos de existir picos, dichos picos o disparos en el tráfico generalmente no son predecibles puesto que dependen de factores como por ejemplo las fechas en las cuales los operadores brindan promociones a sus usuarios, despliegue inicial del servicio, días festivos, etc.

Huawei Technologies a su vez recomienda analizar el comportamiento del tráfico conforme aumente el consumo, con el fin de optimizar los porcentajes de ancho de banda a las clases WFQ en caso de ser necesario.

En caso de existir congestión, WRED (*Weighted Random Early Detection*) es un algoritmo que permite descartar paquetes en base a prioridades o colores (verde,

amarillo y rojo); los paquetes con menor prioridad (rojo) son los primeros en ser descartados cuando existe congestión. Dichos colores son asignados únicamente a las clases WFQ en base a los umbrales establecidos; Huawei Technologies recomienda utilizar los siguientes valores de umbrales para WRED según la tabla 2.10:

Color	Threshold WRED inferior	Threshold WRED superior
Verde	80	100
Amarillo	60	80
Rojo	40	60

Tabla 2.10 Limites para WRED ^[13]

2.1.7 CONFIGURACIÓN DEL SINCRONISMO DE LA RED LTE

Actualmente la fuente del IP *Clock* se encuentra conectado en la red MPLS existente, por lo tanto el único requerimiento es proveer conectividad IP mediante una instancia VPNv4. Los equipos eNB llevarán la configuración necesaria para apuntar a los IP *Clock*, uno será establecido como MASTER y otro como BACKUP.

A nivel de QoS, los siguientes son los requerimientos que deben ser establecidos: un *delay* menor a los 20 ms y porcentaje de paquetes perdidos inferior al 1%.

Para la implementación de sincronización mediante el protocolo 1588v2, se utiliza el algoritmo ACR (*Adaptiva Clock Recovery*), el cual permite que los CSG se sincronicen con los ASBR, dichos ASBR tendrán las fuentes de *clock* y actuarán como SERVIDOR (Master/Slave por temas de protección y redundancia), mientras que los CSG serán CLIENTES tal como se muestra en la figura 2.13

2.2 ANÁLISIS DE LA RED IP- RAN DE LA CNT E.P. EN QUITO

La red IP-RAN de CNT E.P se encuentra al momento operando y a travesando un periodo de expansión, en el periodo de utilización de la red no se han conocido eventos extraordinarios más se encuentra proyectado un crecimiento notable, además las proyecciones analizadas en el punto 1.1.1 y las referencias [1] y [2] razones que fundamentan y que colocan las bases para el análisis de esta arquitectura con el objetivo de optimizarla. Por parte de la expansión se encuentran en proyección nuevos servicios sobre la red LTE y además los nodos 3G también van a ser migrados a esta red, otra razón para fortalecer el desempeño y operación de la IP-RAN de la CNT E.P. en la ciudad de Quito.

2.2.1 CRECIMIENTO DE CLIENTES CON CRECIMIENTO DE CONSUMO DE ANCHO DE BANDA, TANTO DE DATOS COMO DE INTERNET.

La CNT E.P, como resultado del análisis de la Gerencia de Planificación Estratégica corporativa¹ ha determinado el crecimiento hasta el 2019 de los servicios de 2G, 3G y 4G que se encuentra entregado a sus clientes, este estudio presenta proyecciones de internet basado en incremento de ventas y planes comerciales. De este estudio se desprende que a nivel nacional se esperaba un consumo total de internet móvil en Quito en el 2015 de 1.4 Gbps, lo que se ha cumplido en este año, por otra parte

¹ Debido a la alta confidencialidad del documento de planificación de la referencia [32], se mencionarán únicamente resultados que convienen para el análisis.

para el 2019, con el crecimiento de nuevos terminales de usuario se espera un consumo total acumulado de 13 Gbps, en este caso se requiere revisar el diseño actual para confirmar las capacidades del equipamiento y comprobar oportunidades de mejora, en el caso de una reingeniería.

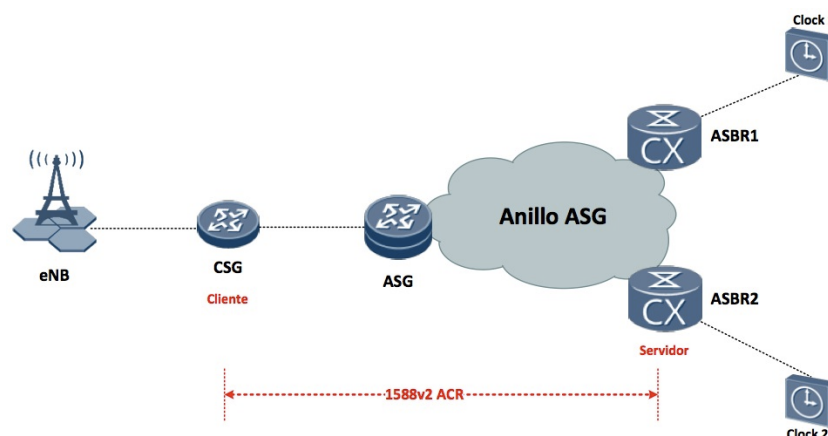


Figura 2.13 IP Clock mediante 1588v2 ACR ^[13]

2.2.2 CAPACIDADES DE LA SOLUCIÓN ACTUAL

En la figura 2.2 se puede observar la topología de la red IP-RAN implementada, en este esquema se ha puesto en detalle a los nodos¹ de agregación ASG, que agrupan a los CSG. Para analizar la capacidad de estos nodos se elaboró la tabla 2-11 basada en en la referencia [39], al tratarse de información confidencial, únicamente se analizará el número de nodos CSG que cada ASG transporta y de esta manera confirmar como se encuentra la capacidad instalada para cada nodo de agregación. Entonces del análisis de la tabla 2-11 se tienen las siguientes observaciones:

¹ Por temas de confidencialidad no se explicará el significado de las siglas que representan los nodos de CNT, esta información consta detallada en las referencias [13] y [38], también se tiene un diagrama detallado de la solución en la referencia [39]

1. El nodo que concentra la mayor parte del tráfico es el MSC, y se puede confirmar además que es uno de los ASBRs.
2. Los nodos que siguen en importancia son QCN, TMB, INQ, MSR, CRD y CTC, en cuanto a la capacidad de ancho de banda, todos tienen interfaces de 10Gbps por lo cual se garantizaría la tasa de transmisión entregada a cada nodo, más la capacidad de los ATN950 asociados, como se confirmará en la sección 2.2.2.1 debe pensarse en una mejora tecnológica para evitar que al crecer el número de clientes y su tráfico no afecten la calidad del servicio.
3. Los nodos PTD, GJL, VLF son nodos que tienen nodos dependientes, por lo que a futuro también se debe pensar en una mejora tecnológica para los equipos ATN950 que sirven en estas localidades.
4. Los nodos MBC, SLD,MTF, LCS, NGN, CBY, LPZ,SIS y LNV son nodos con baja capacidad de conexión hacia los ASBRs, al momento cumplen con los requerimientos de los clientes pero se debe revisar estas capacidades para evitar saturación y migrar los enlaces a tiempo.
5. Los equipos ASBRs, son los que concentran la mayor cantidad de tráfico y son los equipos que se han convertido en fundamentales de la solución, por lo cual deben estar dotados de una adecuada capacidad de conmutación y confiabilidad.

Tomando en cuenta además que la solución actual tiene equipos de dos tipos del proveedor Huawei, los ATN950 que están trabajando en la función de CSG y ASG, es decir como *cell-site* y como equipos de la capa de agregación, así como también se tiene el equipo CX-600, también de Huawei trabajando como ASBR. Por otra parte las interconexiones entre los nodos de agregación y los *cell-sites* tienen múltiples

soluciones con sus respectivas capacidades; a continuación se realizará un análisis de las capacidades de los equipos ya que las interconexiones no serán parte de este estudio, debido a que pueden ir ampliándose bajo demanda gracias a que CNT E.P tiene sus propios sistemas como la red DWDM metropolitana.

Nodo ASG	No de CSG	No de CSG tránsito	Capacidad Máxima Gbps
MBC	2	0	2
SLD	2	0	1
MTF	4	0	2
CNC	1	0	10
SRF	4	0	10
PTD	3	2	10
GMN	4	0	10
GJL	2	2	10
VLF	3	4	10
QCN	12	0	10
LCS	3	0	1
NGN	1	0	1
CBY	4	0	1
TMB	7	0	10
LPZ	2	0	1
MSC	4	60	10
INQ	10	18	10
SIS	1	0	1
LNV	2	0	1
MSR	6	0	10
LCL	3	0	10
CRD	6	0	10
LFL	1	0	10
CND	1	0	10
CTC	8	0	10
LLZ	1	0	10
TIF	1	0	10

Tabla 2.11 Número de nodos de acceso por cada nodo de Agregación ^[39]

2.2.2.1 Equipos CSG y ASG

Para la solución de la ciudad de Quito, la CNT E.P adquirió equipamiento Huawei, en el caso de los *cell-sites* y de los nodos de agregación se trabaja con el mismo equipo, se trata del router ATN950B, representado en la figura 2.14 el cual presenta la siguientes características:



Figura 2.14 Router modelo ATN950- Huawei ^[33]

Según la tabla 2-12, para desempeñar las funciones de *cell-site* este equipo cumple con las necesidades de tráfico, más al momento de agregar otros nodos su baja memoria se convierte en una desventaja. A continuación se de confirma la versión del sistema operativo de Huawei, que al momento no tiene notificaciones de fuera de tiempo de venta ni de soporte registrados en la página oficial de Huawei¹.

Para obtener la información de la versión del sistema operativo se debe ejecutar en el modo terminal no privilegiado el comando display versión.

```
display version
Huawei Versatile Routing Platform Software
VRP (R) software, Version 5.120 (ATN950B V200R002C00SPC300)
Copyright (C) 2011-2013 Huawei Technologies Co., Ltd.
HUAWEI ATN950B uptime is 723 days, 7 hours, 33 minutes
ATN950B version information:
```

2.2.2.2 Equipos ASBR

En los equipos que se encuentran realizando las funciones de ASBR se tienen dos tipos de equipo, el Cisco ASR903 en la red MPLS de Cisco y el router Huawei CX-600 en la red IP-RAN, a continuación se describe la información técnica de estos equipos.

¹ La página oficial de Huawei es www.huawei.com, la información se encuentra restringida para los usuarios que no tienen relación comercial con la compañía.

Especificación Técnica	Descripción
Capacidad de conmutación	AND1CXPA: 88 G bit/s (44 G bit/s for the upstream traffic and 44 G bit/s for the downstream traffic) AND1CXPB: 112 G bit/s (56 G bit/s for the upstream traffic and 56 G bit/s for the downstream traffic) AND2CXPA: 88 G bit/s (44 G bit/s for the upstream traffic and 44 G bit/s for the downstream traffic) AND2CXPB/AND2CXPE: 112 G bit/s (56 G bit/s for the upstream traffic and 56 G bit/s for the downstream traffic)
Capacidad de procesamiento de paquetes	AND1CXPA: 65.48M pps AND1CXPB: 83.33M pps AND2CXPA: 65.48M pps AND2CXPB/AND2CXPE: 65.48M pps
SDRAM	1024M Byte
Flash	128 M Byte
CF card	512M Byte

Tabla 2.12 Especificaciones técnicas del router ATN-950B correspondiente a la tarjeta Controladora instalada ^[33]

2.2.2.2.1 CX-600 de Huawei

De la información disponible para clientes en la página web oficial de Huawei se tiene las siguientes características técnicas que permiten analizar las capacidades de estos routers.

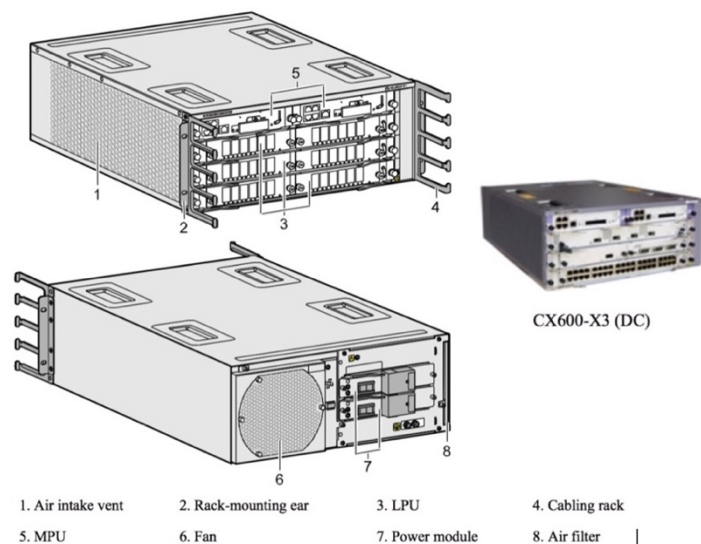


Figura 2.15 Router modelo CX-600- Huawei ^[33]

Especificación Técnica	Descripción
Capacidad de conmutación	1.08 Tbit/s (bidirectional)
Capacidad de procesamiento de paquetes	300 Mpps
SDRAM	2 GB
Flash	32 MB
CF card	Two 1-GB CF cards on each MPU

Tabla 2.13 Especificaciones técnicas del router CX-600 ^[35]

Se tiene a continuación que confirmar el futuro soporte del equipo, se debe recordar que es un ASBR del Inter-AS. Para este propósito hay que realizar la verificación de la versión del equipo que hace la función de borde en la red Huawei.

<UIOINQRB01>**disp version**

Huawei Versatile Routing Platform Software

VRP (R) software, Version 5.120 (CX600 **V600R006**C00SPC300)

Copyright (C) 2000-2013 Huawei Technologies Co., Ltd.

HUAWEI CX600-X3 uptime is 775 days, 3 hours, 46 minutes

CX600-X3 version information:

Se confirma en el sitio web del proveedor la información a cerca de la versión del router obteniéndose la información de la tabla 2-13

Milestone	Definition	Date
EOM	End of Marketing. The EOM date is the date from which the acceptance of the POs for new deployments and capacity expansions will be rejected. The product is not sold any longer after the date.	December 31, 2015
EOFS	End of Full Support. After the EOFS, R&D will not develop patches for the software version.	December 31, 2016(Planned)
EOS	End of Service and Support. After the EOS, Huawei does not provide software problem analysis services.	June 30, 2017(Planned)

Tabla 2.14 Anuncio de soporte de la versión CX-600 V600R006¹

¹ Tabla tomada del sitio <http://www.huawei.com/en/ProductsLifecycle/index.htm>, del producto Metro Service Platform, CX600

Del anuncio de soporte se confirma que se debe pensar en una migración tecnológica a mediano plazo.

2.2.2.2.2 Cisco ASR903

El equipo que desempeña las funciones de ASBR del lado de la red IP-MPLS Cisco, es un equipo de modelo ASR-903 con una procesadora A903-RSP1B-55, representado en la figura 2.15 y presenta las características indicadas en la tabla 2.15 y 2.16. De las características se deduce que por configuración física y soporte se confirma en la necesidad de una migración tecnológica a corto plazo, debido a que este equipamiento está concentrando todo el tráfico de la red IP-RAN, en caso de ser utilizado en una sección su vida útil todavía podría prolongarse..



Figura 2.16 Router modelo ASR 903 ^[36]

Especificación Técnica	Descripción
Capacidad de conmutación	55 Gbps
Capacidad de procesamiento de paquetes	65 Mpps
SDRAM	4G
Flash	External USB flash memory
CF card	N/A

Tabla 2.15 Especificaciones técnicas del router ASR-903 ^[36]

Milestone	Definition	Date
End-of-Life Announcement Date	The date the document that announces the end-of-sale and end-of-life of a product is distributed to the general public.	March 29, 2013
End-of-Sale Date	The last date to order the product through Cisco point-of-sale mechanisms. The product is no longer for sale after this date.	September 27, 2013
Last ShipDate: OS SW	The last possible ship date that can be requested of Cisco and/or its contract manufacturers. Actual ship date is dependent on lead time.	September 27, 2013
End of SW Maintenance Releases Date: OS SW	The last date that Cisco Engineering may release any final software maintenance releases or bug fixes. After this date, Cisco Engineering will no longer develop, repair, maintain, or test the product software.	March 28, 2014
Last Date of Support:	The last date to receive applicable service and support for the product as entitled by active service contracts or by warranty terms and conditions. After this date, all support services for the product are unavailable, and the product becomes obsolete.	September 30, 2018

HW = Hardware OS SW = Operating System Software App. SW = Application Software

Tabla 2.16 Anuncio de futuro soporte del router ASR-903 de Cisco ^[37]

2.2.3 CONCLUSIONES DEL ANÁLISIS DE LA SOLUCIÓN IP-RAN DE LA CNT E.P.

La solución de IP-RAN de la CNT E.P. está basada en Inter-AS opción B y tiene las desventajas que se estudiaron en su momento en la sección 1.3.5.4, entre ellas se tiene la discontinuidad en el servicio, es decir para aprovisionar un nuevo servicio se deben realizar configuraciones adicionales tanto en los nodos como en los ASBRs, esto incrementa la posibilidad de errores al realizar estas configuraciones, como se vio en la sección 2.1.5 existen varias consideraciones en los filtros de enrutamiento dentro de los ASBRs que incrementan la complejidad de la operación. Esto puede optimizarse con la implementación de un esquema de *Seamless MPLS*.

Luego de haber revisado el diseño de la red IP-RAN de Quito de la CNT E.P se tienen las siguientes conclusiones.

a) Falta de Flexibilidad del servicio

En la sección 1.3.5, también se observó que la red actual es poco flexible, en el caso de migrar servicios del un dominio al otro, de hecho se está estudiando la posibilidad de acercar el servicio de DNSs de una manera descentralizada, para llevar a a cabo esta operación en un esquema *Seamless* implica menos complejidad de cómo actualmente se opera con Inter-AS

b) Necesidad de renovación tecnológica de los nodos de agregación.

Como se pudo apreciar en la sección 2.2.2 se han descrito los nodos que requieren esta renovación tecnológica debido a que el mismo equipo que se selecciono como *cell-site* cumple con las funciones de agregación, esto limita el canal de agregación por las características que se analizaron en la sección 2.2.2.1 (la baja memoria para almacenar enrutamiento) y no asegura capacidad para el futuro.

c) Equipos de borde de bajo *backplane* dentro de la solución.

En la sección 2.2.2 se analizaron los nodos INQ y MSC que son los principales y que trabajan como equipos de borde de la solución es decir los equipos ASBRs, como se pudo confirmar en la sección 2.2.2.2 estos equipos trabajan dentro de los parámetros hasta el momento, pero dentro de la solución no aseguran una estabilidad y redundancia necesarios para el crecimiento de la red, además dentro de la solución de Inter-AS se tienen cuatro equipos que trabajan como

ASBRs y se ha confirmado que la redundancia al momento de faltar uno de los equipos ASR903 el tiempo de convergencias esta sobre los 3 minutos, provocando así la pérdida notable del servicio, por esto se debe fortalecer esta solución, en el caso de *Seamless* MPLS se puede aprovechar los equipos de core de la red CISCO MPLS que son CRS-3, con la capacidad suficiente para cumplir con las funciones de P en la red y de core dentro de una solución *Seamless* MPLS.

d) Todos los equipos de la red IP-RAN, están listos para utilizar el protocolo de sincronismo estándar IEEE 1588v2

En los temas de la sincronía lo que si se puede destacar es que se puede trabajar con PTP (IEEE 1588v2) o con syncE o una variación de los dos. Por el momento se encuentra operando los equipos ASRB como servidores principales del protocolo PTP, los cuales se encuentran integrados a dos equipos Huawei conocidos como IP-clock. Los equipos adquiridos están listos para operar con este protocolo, por lo tanto la migración tecnológica y crecimiento deben tomar en cuenta este requerimiento.

e) Equipos que deben sufrir una migración tecnológica por obsolescencia.

A consecuencia de lo presentado en las secciones 2.2.2.1 y 2.2.2.2 se recomienda realizar una mejora tecnológica para asegurar el crecimiento y continuidad del servicio.

f) La configuración e-BGP del Iner-AS debe ser continuamente manipulada.

En la configuración del e-BGP debe implementar mecanismos de control del número de prefijos aprendidos en la red, debido a que esta cantidad tiene directo impacto sobre el consumo de recursos (CPU) de los equipos que cumplen las funciones de *cell-sites* y agregadores. Por lo tanto rutinariamente se deben estar revisando estas sesiones para evitar que los equipos rebasen sus capacidades físicas y afecten de esta manera el servicio.

g) Balanceo de carga entre los ASBRs

Para el caso del Inter-AS, se hace necesario configurar una solución con Local preference para balancear el tráfico entre los dos equipos ASBR y evitar de esta manera saturación, cuya consecuencia es la intermitencia en los servicios. El diseño actual prevé un balanceo, pero se debe estar en constante monitoreo por medio de rutinas de mantenimiento para evitar el desbalance del consumo entre los enlaces de los ASBRs hacia el core del MPLS.

h) Escalabilidad

La Solución de Inter-AS, actualmente implementada tiene la desventaja de que al expandir la red LTE en el resto de ciudades importantes dentro del Ecuador, debe entonces implementarse como una red paralela a la red MPLS nacional, lo que implica duplicidad de recursos y en consecuencia una configuración de *Seamless* MPLS proporcionará una optimización de dichos recursos y ayudará a que a nivel Nacional se obtenga una solución escalable.

Otra desventaja de escalabilidad es que para la implementación de un nuevo nodo de agregación se debe pensar en la conectividad hacia los ASBRs sin

aprovechar los nodos existentes de la red de CNT E.P, duplicando recursos de transmisión.

i) Resultados

Por el análisis aquí expuesto se puede concluir que es necesario plantear una reingeniería que sea escalable y permita mejorar las falencias encontradas en el diseño actual, tanto de capacidad como de operación.

CAPÍTULO 3

DISEÑO DE LA SOLUCIÓN DE REINGENIERÍA DE LA RED IP-RAN DE LA CNT E.P EN QUITO

La arquitectura de las redes *Seamless* MPLS son utilizadas para obtener soluciones considerablemente escalables, las que reducen los puntos para el aprovisionamiento de servicios, reduce la complejidad y permite la flexibilidad de nuevos servicios en la red de los proveedores de servicios¹. Tomando en cuenta las conclusiones de la sección 2.2.3 se toma la decisión de realizar una mejora tecnológica que permitirá obtener una red escalable y sencilla de operar, a continuación se presenta un nuevo rediseño que para cambiar el esquema actual de Inter-AS por *Seamless* MPLS, se definirá el diseño de alto nivel y estrategia de migración.

3.1 ANÁLISIS DE REQUERIMIENTOS

Tomando como base lo expuesto en la sección 2.2.3 y lo que los requerimientos generales de la CNT E.P. en cuanto a crecimiento futuro se tiene a continuación características que nos permitirá elegir una de las formas de conexión del *Seamless* MPLS que conviene para la rediseñar el servicio de IP-RAN y aprovechar de las ventajas de esta evolución de la conectividad MPLS.

¹ Tomado de la referencia [40]

3.1.1 ANÁLISIS DE REQUERIMIENTOS CON RESPECTO AL DISEÑO DE INTER-AS

A continuación se definen los requerimientos que son el resultado de observar las desventajas de la configuración actual de la IP-RAN, es decir el modelo de Inter-AS que se encuentra aplicado y confirmar que en el nuevo diseño se superen éstas características.

a) Discontinuidad de extremo a extremo del servicio

Las soluciones de *Seamless* terminan con esta discontinuidad, planteando el servicio con las ventajas mencionadas en la sección 1.3.5.4. La reingeniería consiste entonces en plantear la nueva solución orientada a la topología sugerida, de tal manera de tener las ventajas del *Seamless* MPLS, en este caso de tener un servicio continuo.

b) Falta de Flexibilidad del servicio

Una de las principales características de las redes *Seamless* MPLS es la facilidad de crear nuevos servicios y descentralizarlos, en el caso del tráfico de datos de usuario para su autenticación y tarificación o servicios de datos protegidos, etc., que pueden mejorar el desempeño acercándose al usuario, o también servicios de streaming. Con el modelo de Inter-AS tal como se encuentra al momento trabajando la red LTE de Quito, se deben realizar muchos cambios para provisionar servicios que al momento están centralizados o no se encuentran en el portafolio actual.

c) Necesidad de renovación tecnológica de los nodos de agregación.

El alcance del rediseño debe contemplar el reemplazo por mejora tecnológica de los equipos que hacen agregación ASG, CSG y los ASBRs en el core. Se debe tomar en cuenta, tanto las capacidades de los equipos actuales, como las funcionalidades que se desea para la nueva solución.

d) Equipos de borde de bajo *backplane* dentro de la solución.

Actualmente los equipos que se encuentran operando como ASBRs tienen limitaciones en el crecimiento de tráfico, por una parte porque concentran la conectividad de toda la red y por que se encuentran en lista de limite de soporte en el corto plazo. El Nuevo diseño debe contemplar la mejora tecnológica necesaria en este aspecto.

e) Todos los equipos de la red IP-RAN, están listos para utilizar el protocolo de sincronismo estándar IEEE 1588v2

La solución actual de sincronismo para la red IP-RAN depende de dos equipos propietarios de Huawei denominados IP *clock*, los cuales toman la señal de reloj de la tradicional red de transmisión de la CNT E.P. Es por medio del protocolo IEEE 1588v2 o PTP que se hace la distribución de la señal de reloj en el resto de la red. En la solución planteada se trabaja de la misma manera, solo que los equipos de core están interconectados a los relojes principales de la red de transmisión y en los nodos de pre-agregación se puede reforzar la señal de reloj conectándolos a la señal de reloj de la central telefónica donde se encuentra

f) La configuración e-BGP del Iner-AS debe ser continuamente manipulada.

En la solución actual, la configuración de los ASBRs del Inter-AS debe ser siempre revisada, cuando se añade un nuevo nodo, para un servicio que ya está establecido, en cambio en la solución Seamless los equipos ASBR de core no se manipulan para agregar nuevos nodos luego de haber establecido un servicio. Las configuraciones del equipamiento se revisarán más adelante..

g) Balanceo de carga entre los ASBRs

En la operación actual se consideró crítico, por crecimiento de tráfico de los clientes el tener un balanceo de carga en los ASBR (ver figura 2.10), ya que estos equipos concentran todo el tráfico LTE de Quito; en cambio en la nueva solución propuesta esto no se vuelve crítico porque, se trata de una solución distribuida, tanto los nodos de pre-agregación como de Core tienen la capacidad necesaria y la posibilidad de añadir a la operación más interfaces de 10Gbps para crecimientos futuros.

3.1.2 CRITERIOS DE DISEÑO

Como resultado del análisis anterior se tienen un conjunto de criterios que permitirán llegar a un rediseño de la red IP-RAN, con el propósito de conseguir una red escalable acorde al crecimiento de usuarios y sus respectivos requerimientos de ancho de banda que demandan las nuevas aplicaciones de internet y dispositivos que el usuario dispone para utilizar la red LTE de la CNT E.P.

3.1.2.1 Escalabilidad y Capacidad

Debido a la topología que se encuentra actualmente operando, cuyo esquema se observa en la figura 2.2 y tabla 2-11; en la que se tiene: una capa previa de agregación de nodos *cell-site* MPLS, una etapa de pre agregación y dos equipos ASBRs se ha determinado por afinidad y escalabilidad que el modelo que conviene para el rediseño es la opción de la sección 1.3.5.3.6. En cuanto a la capacidad de transmisión, como se puede confirmar en la figura 2.2, la sección de pre-agregación se encuentra con enlaces de 10Gbps, que se encuentra dentro de los planificado en la demanda según la referencia [32] y la tabla 2-11, por lo tanto en el diseño se aprovecharán estos enlaces. Existe un aspecto muy importante a considerar, se trata de la creación de una capa de agregación en la topología de la red, debido a que a la topología de la red MPLS de CNT en Quito (referencia [42]) se encuentra la capa de acceso directamente conectada a la capa de CORE, la razón de este diseño, es que no existía la visión de aplicaciones que aprovechen una red de agregación, al tratarse de una arquitectura plana de MPLS.

3.1.2.2 Resiliencia

En la topología de la figura 2.2 se puede observar que existen anillos de la capa de agregación de los nodos ASG y se puede entonces confirmar que la protección de los servicios se da gracias a estos anillos, pero que de todas maneras el anillo depende de un solo equipo ASBR, tomando muchas veces un único punto de redundancia. Entonces el nuevo diseño debe plantear una topología con mejor redundancia, tanto de enlaces como de equipamiento. Para alcanzar este objetivo se

ha solicitado ayuda a la gerencia de transmisiones y gracias a la experiencia de los años en la jefatura de plataforma IP-MPLS se pudo obtener información de la capacidad de transmisión que nos ayudará a plantear una mejor arquitectura para la red IP-RAN

3.1.2.3 Arquitectura de los servicios LTE

Para los servicios que se encuentran operando se trabajará con el modelo descrito en la sección 1.3.4, ya que se tienen los servicios en L3VPN y con la reingeniería se seguirá operando basándose en el mismo modelo, pero adaptado a la nueva topología, es decir solo hay un cambio en el transporte de información de las VRFs tanto de servicio como de gestión, es decir que de igual manera la conectividad S1 y X2 siguen operando de la misma manera. De manera general entonces el servicio se lo puede observar aplicando lo que se expuso en la figura 1.26, ahora aplicado a la nueva arquitectura.

3.2 DISEÑO DE LA NUEVA SOLUCIÓN

El presente diseño consta de dos partes y un diseño general denominado de alto nivel, donde se describe de manera general la solución y funcionalidades, en la segunda parte se detalla la solución de manera didáctica, para comprender como se está aplicando este diseño general y como se procederá a configurar los equipos para cumplir con el plan de migración. El objetivo es presentar entonces el diseño físico y su funcionalidad y presentar un plan de migración de servicios a la nueva red con *Seamless MPLS*.

3.2.1 DISEÑO DE ALTO NIVEL

En general en la sección 1.3.4 se describió los tipos de tráfico que se tiene en LTE dentro de la IP-RAN, X2 y S1, la nueva solución cubre de igual forma los mismos requerimientos, esto se ha esquematizado nuevamente¹ en la figura 3.1, donde se puede apreciar los servicios para la red LTE, que este caso se presentan en la VRF serlte y geslte, que se tienen localmente en los *cell-sites* ya que se tiene MPLS directamente en estos equipos, gracias a la configuración de *Seamless MPLS*, el tráfico es transportado por una capa de pre-agregación, hasta llegar al EPC, que está conectado al core de la red MPLS, que a su vez se encuentra interconectado a core de servicios móviles MPC(*Mobile Packet Core*).

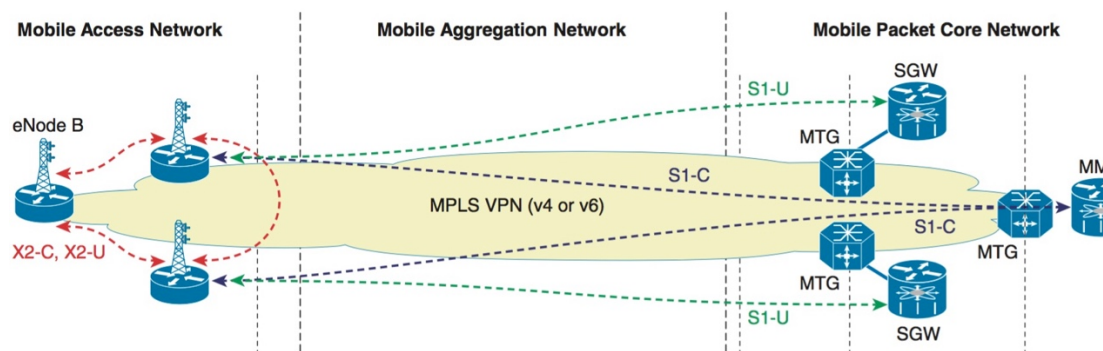


Figura 3.1 Diseño de alto nivel del Tráfico LTE con L3VPN ^[41]

En la figura 3.1 se puede observar los siguientes tipos de tráfico que es necesario aclarar:

¹ Esta sección tiene tres fuentes importantes en las que se basó el diseño, la referencia [41], referencia [27] y [43]

- S1-c, es el tráfico entre un eNodeB y el servicio de MME, por ello también le suelen llamar S1-MME.
- S1-u, es el tráfico entre el eNB y el servicio de datos SGW.
- X2-u, se refiere al tráfico entre dos eNB de la misma zona de agregación.
- X2-c, se trata el tráfico de *handover* entre dos celdas físicamente contiguas pero que pertenecen a dos zonas de agregación diferentes, entonces este tráfico pasa primero por el core y luego a su destino.

A través del modelo de transporte y el tipo de MPLS VPN en L3 seleccionados, se requiere filtrar apropiadamente la información de enrutamiento, de tal manera que se optimice los recursos en los nodos finales. Dentro de la MPLS VPN de servicios, mediante el atributo de *community* dentro del proceso MP-BGP (*Multiprotocol BGP*) se controla que en las zonas de agregación se aprendan las rutas necesarias, es decir las rutas de los servicios (tráfico S1) y de las zonas cercanas para el tráfico X2. En la figura 3.2 se representa, como los RT (route-targets¹) son empleados para lograr el objetivo deseado, entonces se tiene lo siguiente:

- Un único RT llamado **Common RT**, el cual concentra las rutas de todos los nodos, para uso del core móvil MPC (*Mobile packed core*) que se encuentran configurados en los routers que se denominan MTG (*Mobile Transport Gateway*).
- Un único RT llamado **MPC RT**, que contiene las rutas del core móvil que deben tener todos los nodos CSG en la VRF.

¹ En la sección 1.3.1.1 se tiene una descripción mas detallada de la operación de una L3VPN a cerca de los route-targets

- Cada red de acceso agregada tiene su propio RT, denotado RAN X, RAN Y y RAN Z, entonces las rutas de una zona de CSG se distribuyen a otra solo si tienen celdas contiguas que puedan tener un posible tráfico X2

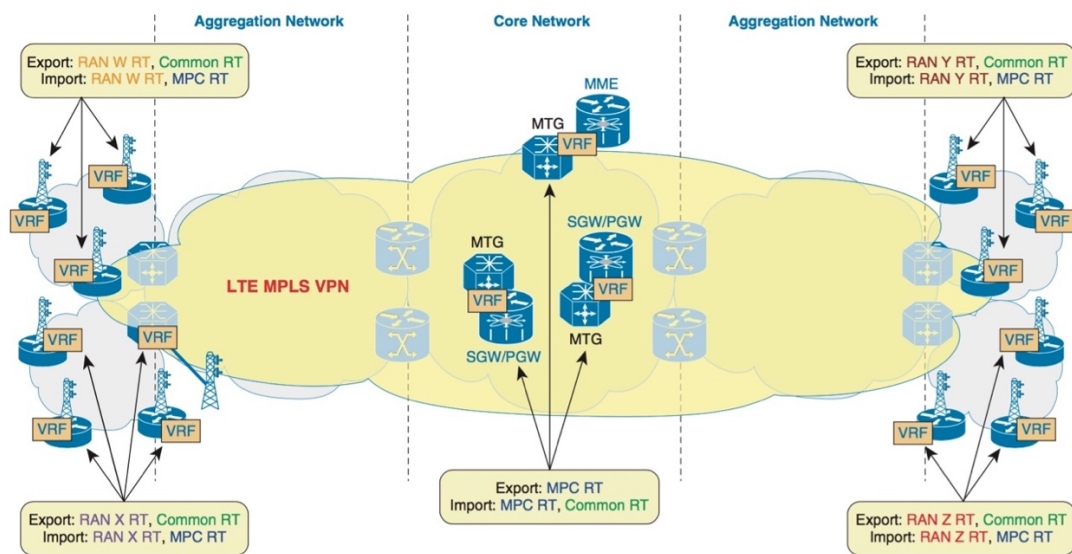


Figura 3.2 Diseño de Alto Nivel de la L3 VPN de servicios LTE ^[41]

La red de transporte planteada tiene las siguientes características:

- Un único Sistema autónomo en BGP.
- Se continua utilizando ISIS como IGP, pero en dominios separados, en el core se utiliza IS-IS L2, en la zona de Pre-Agregación ISIS L1, en la zona de Agregación de los CSG se tendrá ISIS L1, en un proceso y área distinto, el equipo que une las dos zonas que se denominará PAN (*Pre-Agregation Node*) tendrá IS-IS L1 pero en diferente proceso de enrutamiento para separar la capa de pre-agregación de la de agregación.
- Distribución jerárquica de las etiquetas gracias al RFC-3107, es decir la activación de lo que se conoce como BGP ipv4 unicast + label.

- En los nodos PAN se tiene redistribución de etiquetas con filtro de comunidades como se indicó en la figura 3.2.

Por otro lado es importante comprobar como se construye en MPLS el LSP de extremo a extremo, es por eso que se incluye esta última parte en el diseño de alto nivel. Previamente se aclara la nomenclatura utilizada para explicar las funciones de cada router dentro de la topología diseñada¹:

- **CSG**, como se había indicado anteriormente *Cell-Site Gateway*.
- **PAN**, para indicar el router ASBR entre la red de agregación y la de pre-agregación, viene de *Pre-Agregation Node*.
- **CN-ABR**, para diferencia al ASBR que se encuentra en el core.
- **CN-RR**, *route-reflector* de la capa de core.
- **Inline-RR**, *route-reflector*, que es una función de los equipos de core como *route-reflector* de la capa de agregación.

En la figura 3.3 se tiene el seguimiento de cómo las etiquetas son manejadas en cada una de las zonas en las que se ha dividido el protocolo de enrutamiento IGP y de esta manera ver como se consigue establecer la comunicación de los servicios en los e-Mbps. Para empezar los nodos CSG no tienen sesiones *BGP+label*², ya que se

¹ Esta nomenclatura viene de la referencia [41], ya que la solución aplicada es de Cisco, aunque el concepto es standar, a excepción claro esta de funcionalidades adicionales propias de este fabricante.

² Nos referimos con esta nomenclatura al RFC3107, que es como las etiquetas son transportadas por BGP. En la solución propuesta los CSGs no tienen sesiones *BGP+label* con los ASBRs de pre-agregación, pero si la red de acceso necesitara ampliar su accionar sobre servicios fijos masivos, entonces si es necesario también tener esto en lugar de una simple distribución.

redistribuye la información de las etiquetas desde el core hasta los nodos de pre-agregación de la siguiente manera:

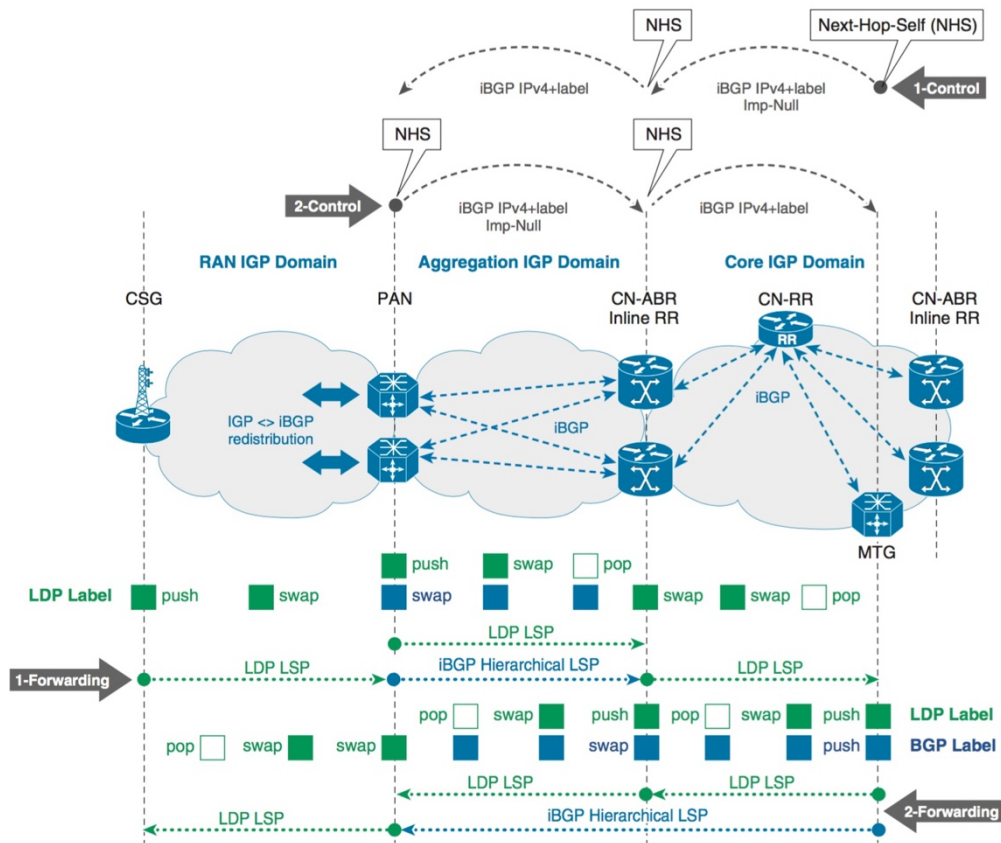


Figura 3.3 Diseño de Alto Nivel de la L3 VPN de servicios LTE ^[41]

- El punto de presencia en el core (CN-ABR) tiene sesiones *BGP+label* y actúan como *route-reflectors* o también denominado Inline-RR (*Route Reflector* en línea) con los equipos ASBR de pre-agregación PAN, los CN-ABRs entre si deben tener las mismas sesiones *BGP+label* entre si, en caso de crecimiento se puede pensar en un *route-reflector* para simplificar estas configuraciones en el core, esto depende del crecimiento de la red a futuro.
- Los MTG, son los PEs que concentran los servicios móviles(SGW,PGW y MME) y son parte del core. Los MTGs tienen sesiones *BGP+label*

directamente con los CN-ABR, publicando sus *loopbacks*¹ con una comunidad como se había indicado anteriormente.

- c) Todos los equipos PAN, de una zona de agregación requieren información de etiquetas para alcanzar a otros equipos PAN de otras zonas, por lo que se tienen sesiones *BGP+label* con los equipos CN-ABRs locales, los que son sus *route-reflectors*. Los equipos PAN publican sus *loopbacks* con una comunidad común a su zona y aprenden los prefijos de los CSGs y de los MTGs.
- d) Debido a que no se redistribuyen rutas entre el IGP del core (ISIS *level 2*) y la zona de agregación (ISIS *level1*), los CN-ABR tienen que reflejar los prefijos aprendidos de los MTGs hacia los routers PAN por *BGP+label* cambiando su origen y declarándose como *next-hop*², por lo que se coloca como parte del camino para alcanzar a la red solicitada y con esto conseguir unir el LSP a pesar de existir una separación en el IGP.
- e) Los MTGs deben tener la capacidad de manejar un escalable número de prefijos, ya que son los equipos que aprenderán todos los prefijos de todos los CSGs. Se tiene un control de los prefijos aprendidos en cada nodo PAN, por medio de un filtro basado en el atributo de comunidades³ para evitar que los nodos remotos aprendan prefijos innecesarios, los nodos PAN aprenden únicamente los prefijos marcados como comunes, los marcados como

¹ En las redes MPLS, las interfaces *loopback* de los router sirven para identificar la información principal en la que los procesos de enrutamiento y en este caso las rutas de los servicios de los MTGs vienen en el BGP con esta dirección de *loopback* como origen y vemos que en esta solución como esta información es distribuida.

² Se considera a *next-hop* como un atributo mandatorio en la información de enrutamiento del protocolo MGP, en el presente caso MP-BGP.

³ Las comunidades se las vio en la sección anterior

provenientes de los MTGs; los prefijos aprendidos son enseñados a todos los routers PAN de esta zona. En la región de acceso, los routers PAN manejan con un proceso de enrutamiento diferente el control de los prefijos que ingresan el IGP de acceso, llegando únicamente los prefijos correspondientes al servicio.

3.2.2 DISEÑO DE BAJO NIVEL

En el diseño de alto nivel, se han descrito las funcionalidades de la red, el objetivo de un diseño de bajo nivel es llegar a detallar esas funcionalidades para ver como se consigue aplicar la solución de *Seamless* MPLS al caso particular de la red IP-RAN de CNT E.P. en Quito, a continuación se detallan los aspectos técnicos del diseño aplicado.

3.2.2.1 Diseño de Conexiones físicas

Para el nuevo diseño físico de la red, se ha utilizado información de las referencias [13], [27], [30], [38], [39], [41], y [43]; también por otro lado información confidencial de las capacidades de la red de fibra óptica metropolitana de la CNT E.P en Quito, esto mediante consultas a la Jefatura de Red troncal de Fibra Óptica¹, con lo cual se han confirmado capacidades que permitirán implementar el diseño que se detalla en esta sección.

¹ Estas consultas no se han añadido por tratarse de información sensible para la CNT E.P.

Por escalabilidad de la red, apoyado en la demanda¹ se eligió el modelo de *Seamless* MPLS a aplicar que se revisó en la sección 1.3.5.3.6; para empezar el diseño se crea una capa de agregación con los routers que se denominaron PAN, en la sección anterior y basados en la referencia [38] y en la tabla 2-11 se tomó la decisión de zonificar al ciudad de Quito en tres áreas, una que cubre la parte norte, otra la parte centro-sur y la zona de los valles, tal como se muestra en la tabla 3-1; para la generación de la tabla se tomó en cuenta la cercanía geográfica de los nodos y la capacidad de la red de transmisión de la red de fibra. En la figura 3.1 se presenta el diseño físico de la red, presentándose las capas de core, agregación y pre-agregación.

Nodo ASG	No de CSG	No de CSG tránsito	Capacidad Máxima Gbps	Sector
CNC	1	0	10	VALLES
SRF	4	0	10	VALLES
NGN	1	0	1	VALLES
CBY	4	0	1	VALLES
TMB	7	0	10	VALLES
TIF	1	0	10	VALLES
LPZ	2	0	1	NORTE
INQ	10	18	10	NORTE
SIS	1	0	1	NORTE
LNQ	2	0	1	NORTE
MSR	6	0	10	NORTE
LCL	3	0	10	NORTE
CRD	6	0	10	NORTE
LFL	1	0	10	NORTE
CND	1	0	10	NORTE
CTC	8	0	10	NORTE
LLZ	1	0	10	NORTE
MBC	2	0	2	CENTRO-SUR
SLD	2	0	1	CENTRO-SUR
MTF	4	0	2	CENTRO-SUR
PTD	3	8	10	CENTRO-SUR
GMN	4	0	10	CENTRO-SUR
GJL	2	1	10	CENTRO-SUR
VLF	3	2	10	CENTRO-SUR
QCN	12	0	10	CENTRO-SUR
LCS	3	0	1	CENTRO-SUR
MSC	4	63	10	CENTRO-SUR

Tabla 3.1 Distribución de los nodos de pre-agregación en tres zonas de agregación

¹ Tomado de la referencia [32].

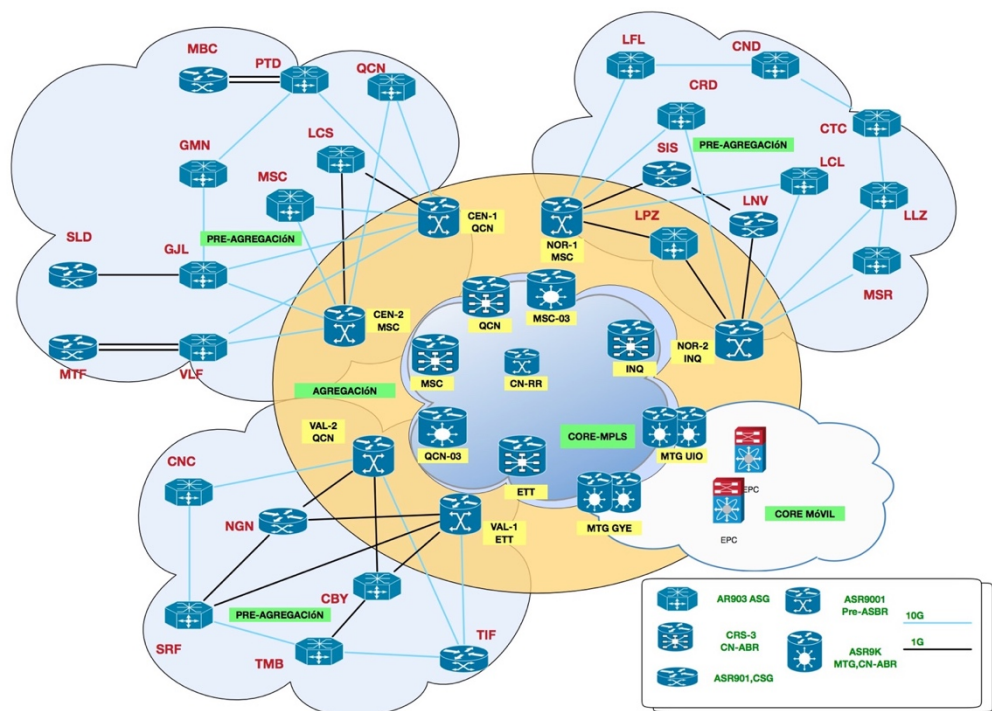


Figura 3.4 Diseño de bajo nivel de la IP-RAN de servicios LTE de Quito

3.2.2.2 Diseño del protocolo de enrutamiento IS-IS

El protocolo de enrutamiento seleccionado es IS-IS, que es el mismo ya existente en toda la red MPLS de la CNT, para la solución propuesta se ha segmentado en dominios diferentes para controlar de esta manera las tablas de enrutamiento correspondiente en cada dominio. En la figura 3.5 se puede apreciar que en la zona de core se encuentra configurado IS-IS L2, en la red de agregación L1 y para la pre-agregación se tiene nuevamente L1; los routers de Core operan del modo IS-IS-L1/L2, los router de servicio como los MTG se encuentran en IS-IS-L2, los equipos PAN están en IS-IS-L1 manteniendo un proceso de IS-IS en la red de agregación y otro hacia los *cell-site* que también se mantienen en IS-IS-L1

Para las direcciones de IS-IS de la red de Agregación incluidos los equipos PAN, se sigue el mismo formato que en la red de Core del MPLS, es decir basados en la tabla 2.3, es decir con el Área ID 0008 y la IPv4 *loopback*100 mapeada dentro de la dirección NET. En cambio para los equipos de la capa de acceso, el plan de direccionamiento IP de los equipos no cambia se mantiene el original incluso con el Área ID 0009.

3.2.2.3 Diseño del protocolo de enrutamiento BGP

Como se trata de una red MPLS, se tiene la distribución de etiquetas con el protocolo LDP, pero al encontrarse segmentado en dominios diferentes el protocolo de enrutamiento interno IS-IS ya únicamente se tiene información de etiquetas de forma local. Para transportar en diferentes dominios la información necesaria de las etiquetas se tiene el protocolo BGP, que para distribución de etiquetas se basa en la RFC 3107, de esta manera se tiene la información de las etiquetas para las interfaces loopbacks dentro de la red IP-RAN, tal como se describió en la sección 3.2.1.

Se tienen tres tipos de configuraciones de *BGP+label* a considerar de la siguiente manera:

- Los equipos MTG deben activar una sesión de *BGP+label* contra el route-reflector del core o con todos los equipos de core.
- Los equipos de core CN-ABR deben establecer sesiones *BGP+label* contra el route-reflector del core (con todos los equipos del core, a esto se lo conoce

como BGO *full-mesh*). Además los CN-ABR son route-reflector para los equipos PAN, por esto deben tener en la sesión BGP la característica de *next-hop-self*.

- Los equipos PAN deben establecer sesiones *BGP+label* contra el route-reflector, que en éste caso son los equipos CN-ABR y a su vez estos equipos son route-reflector para los equipos CSG.

En la zona de agregación se realiza una redistribución de BGP al IS-IS para permitir que las *loopbacks* de los MTG sean alcanzables, para el tráfico S1 y de las zonas cercanas que servirán para el tráfico X2, cuando el *handover* ocurre en los límites de dos zonas de acceso.

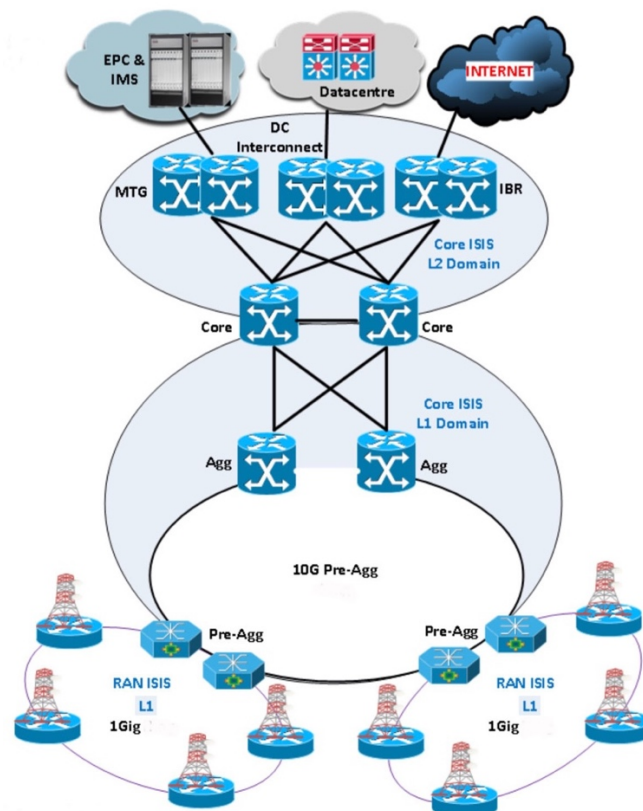


Figura 3.5 Segmentación de dominios IS-IS de la IP-RAN de LTE de Quito [27]

En la distribución de información *BGP+label* se encuentran filtros para que por comunidades discriminar la información de las IPs de las interfaces loopbacks, para lo cual se ha definido la tabla 3.2 que permite identificar el equipo y su roll dentro de la red, de la siguiente manera:

AS: Sistema autónomo de CNT

ZZ: Identificador de la zona, todas , 01 zona norte, 02 zona centro-sur y 03 zona valles

FF: Funcionalidad o roll del equipo, 01 Core-ASBR, 02 MTG o Core, 03 Core –RR, 04 – PAN, 05 ASG, 06 CSG.

CC: Identificador General común, 1 IP-RAN, 2 MTG, 3Core.

AS	:	ZZ	FF	CC
28006	:	01	01	1
Sistema Autónomo	.	Identificador de la zona de agregación	Identificador de la función del equipo	Identificador Común de la red

Tabla 3.2 Comunidades BGP para optimizar la publicación de las interfaces *loopback*

Ejemplos: 28006:00001 es la comunidad en común del acceso IP-RAN, 28006:00002 es la comunidad de los MTG, 28006:00003 es la comunidad de Core y de una zona IP-RAN es 28006:02XX1, es decir que los dígitos FF son informativos únicamente.

3.2.2.4 Diseño de la L3 VPN para los servicios LTE

Como se expuso en la sección 3.2.1, para los servicios del LTE se tiene una L3VPN para el tráfico S1 y el X2, dentro de la VRF de servicios se tienen las comunidades¹ para RT (*routed-target*) que en MPLS se utilizan para distribuir las rutas VPNV4 de la VRF de servicios para lo cual se ha designado los siguientes valores descritos en la tabla 3-3.

AS	:	RA	ZZ	III
28006	:	01	01	9
Sistema Autónomo	..	Identificador de zona de agregación	Identificador de la zona de pre-agregación	Identificador de RT de IP-RAN

Tabla 3.3 RT de la VRF de servicios

AS: Sistema autónomo de CNT

RA: Identifica la zona de agregación, R es la Región en que administrativamente se dividió CNT (son siete regiones, Pichincha pertenece a la regional 2) en este caso la 20 es Quito, esta nomenclatura puede ser utilizada en caso de extender este diseño a otras ciudades.

¹ En la sección 1.3.1.1 se describe en detalle que es un RT en MPLS

ZZ: Identificador de la zona, 00 no pertenece a zona de acceso alguna , 01 zona norte, 02 zona centro-sur y 03 zona valles

III: Para identificar dentro de los RT de la red que se trata del servicio IP-RAN de servicios 9 y 8 para lo que es gestión

Entonces se usará el RT 28006:20009, para exportar las rutas hacia la VRF de servicios LTE y 28006:20008 para exportar las rutas hacia la VRF de gestión, para identificar cada zona de servicios LTE con 28006:20019, 28006:20029, 28006:20039 y para identificar cada zona de gestión se tiene 28006:20018, 28006:20028, 28006:20038

La VRF de datos se encuentra creada en la red MPLS y se denomina movdcn en el equipo que se considera MTG de los servicios móviles, para este caso su RT es 28006:605001; por otro lado se tiene la VRF de gestión denominada movoyim_trust con su RT 28006:602003. En la figura 3.6 se encuentra esquematizado las dos VRFs que sirven tanto para la gestión, como para el servicio de la red LTE ; las VRFs se encuentran configuradas también en el MTG y los CSG de la red.

Las VRFs, tanto de servicio como la de gestión en MPLS dependen de las rutas VPNV4 como se había visto en la sección 1.3.1.1 las cuales dependen de dos aspectos fundamentales, el primero es las sesiones MP-BGP contra los *route-reflectors* en el address-family vpnv4 y recordemos que las rutas vienen con un origen, que son las loopbacks en ipv4 de los equipos que poseen estas rutas; es importante tomar en cuenta lo siguiente:

- Para los CSGs, la función de *route-reflector* es desempeñada por los router PAN, los cuales además tienen el atributo de BGP de *next-hop-self*, lo que hace que el tráfico pase por los mismos controlando el flujo tanto de tráfico como de inundación de prefijos aprendidos.
- Los equipos CSG sugeridos en el diseño además pueden presentar una funcionalidad denominada RTC (*Route Target Constraint*), la cual es un RFC 4684, que consiste en establecer una comunicación vía BGP con el route-reflector para limitar los prefijos no deseados de otras VRF del sistema, que no estén configuradas en los CSG y de esta manera precautelar la memoria y procesamiento en información innecesaria. Esta funcionalidad es más útil cuando se desea implementar algún servicio adicional en estos nodos.
- Para los router PAN estos los *route-reflector* son ahora los routers de core CN-ABR, publicando lo aprendido por los CSG y aprendiendo del core las rutas VPNV4 necesarias.
- Los CN-ABR, son los *route-reflector* principales y es en las sesiones BGP VPNV4 donde se realiza el control, según las comunidades definidas.
- Para que las rutas VPNV4 dependen de las *loopback* aprendidas en ipv4, es decir por el IGP y por las etiquetas que son transportadas por las sesión *BGP+label* definidas en la sección 3.2.2.3

Como se puede ver en la figura 3.7 se puede tener el caso en que un celda LTE que se encuentra agregada en un nodo CSG este cercana a otra celda que no pertenezca a la misma zona de agregación, razón por lo cual se hace necesario para estos casos identificar los RT dentro de la VRF movdcn, necesarios para distribuirlos entre las

zonas consideradas adyacentes, se puede incluso filtra la información de las loopbacks de los equipos necesarios, para evitar la inundación de información que solo consumirá recursos.

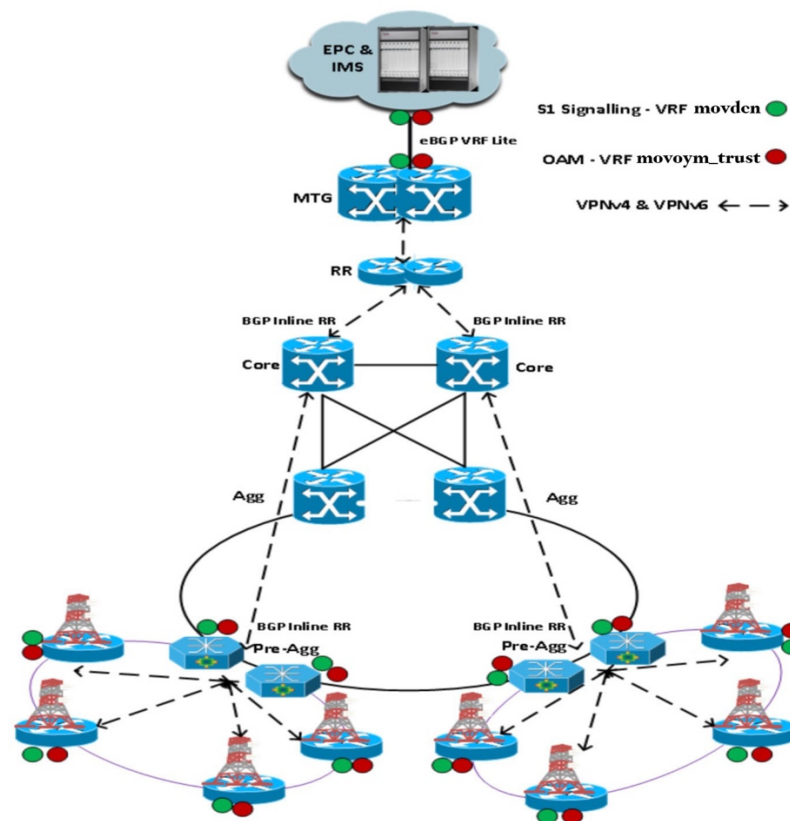


Figura 3.6 Modelo del diseño de las L3VPN de LTE [24]

3.2.2.5 Resiliencia en los servicios LTE

Debido a que el protocolo de enrutamiento interno IS-IS ha sido segmentado, para la distribución de etiquetas se depende no solo del protocolo LDP, son también del protocolo BGP, este último conocido por su alto tiempo de convergencia; en las red MPLS sin segmentación para la distribución de la información de las *loopbacks* necesarias para crear los LSPs de la red, se dependía únicamente del protocolo LDP y del IGP, mas ahora se requiere mejorar el tiempo de convergencia en caso de

eventos que interrumpen el tráfico en la red¹. A continuación se describen los mecanismos que permiten mejorar los tiempos de restablecimiento del servicio en caso de existir algún evento en la red y en donde se deben configurar.

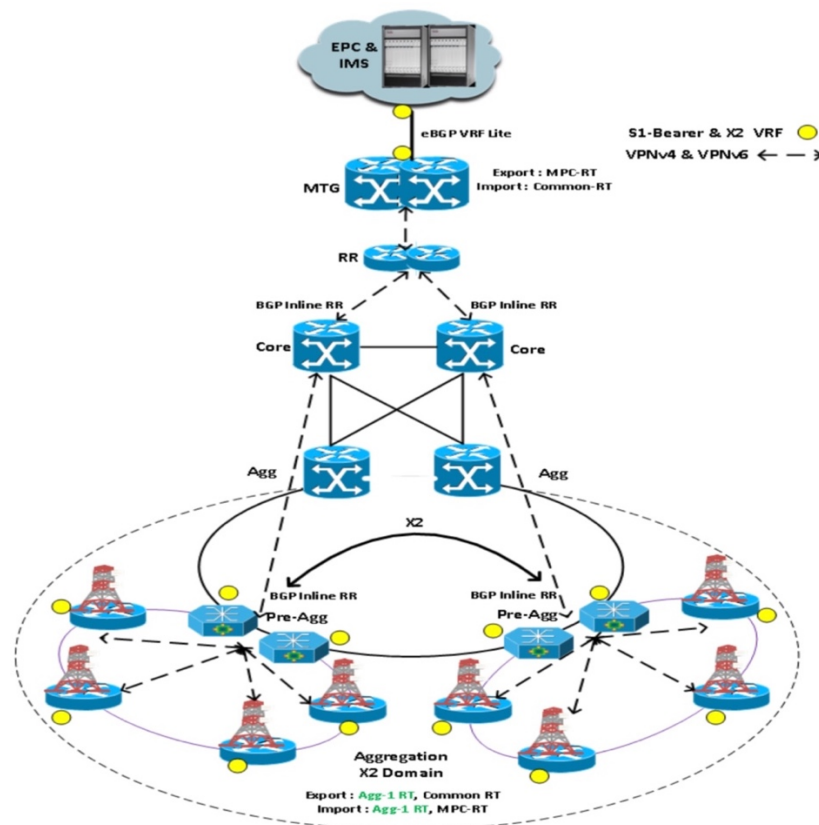


Figura 3.7 Tráfico X2 en caso de regiones contiguas L3VPN de LTE ^[27]

Los proveedores, en esta caso proponen los siguientes mecanismos, que se encuentran ligados a las capacidades de los equipos que intervienen en el diseño, entonces se tiene:

¹ Los siguientes mecanismos de mejora para lograr optimizar los tiempos de convergencia del sistema diseñado son estándar pero la aplicación es tomada de las referencias [43], [27], [30] y [41] del fabricante Cisco

- a) LFA FRR (*Loop Free Alternate Fast Reroute*), es una funcionalidad del sistema operativo del equipamiento que está basado en el RFC 5286, consiste en realizar un cálculo de un camino alternativo del protocolo interno de enrutamiento y tener esa información lista ante una posible falla, por lo general tradicionalmente esto se solía solventar en una red MPLS con un túnel de ingeniería de tráfico, pero la principal ventaja de esta funcionalidad es que solo hay que activarla y no planificar una configuración como es el caso de un túnel de ingeniería. Esta funcionalidad es muy útil en el caso de topologías en anillos, como se puede ver en la figura 3.8 cuando el tráfico de C1 a C2 presenta una falla y se desea enviar datos desde C2 hasta A1, C2 envía los datos a C3 y C3 retorna la información a C2 ya que todavía considera como mejor ruta por C2 para llegar a A1, con LFA C3 tiene ya lista una ruta de *backup* y dirige el paquete por C4.

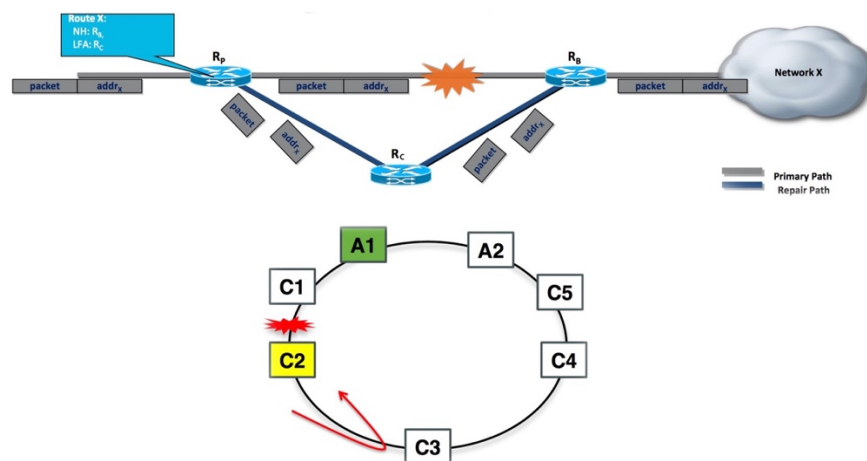


Figura 3.8 Funcionalidad de la RFC 5286 LFA FRR ^[43]

- b) Remote LFA FRR, esta funcionalidad, que viene del estándar desde el 2014 con la RFC 7490, extiende las funcionalidades de LFA para otras topologías. Un nodo en forma dinámica calcula el o los nodos remotos alternos libres de

laso (LFAs) y establece túneles LDP para intercambiar etiquetas con esos nodos, encaso de una falla como la de la figura 3.9 a) el tráfico MPLS continua por los caminos alternos. En la figura 3.9 b) se representa como las etiquetas trabajan en el caso de una falla, el tráfico de origen está en C2 y el destino de este tráfico es A1; C2 tiene una etiqueta 20 por C1 para llegar a A1, una etiqueta 99 por C3 para llegar a C5 y una etiqueta 21 por C5 para llegar a A1; para el caso de la falla entre C1 y C2, el tráfico es enviado por el LSP de C5 inmediatamente sin recalculo de LDP ya que las etiquetas necesarias ya fueron calculadas. Esta funcionalidad es parte del sistema operativo de los routers y solo necesita se activado no requiere de configuración adicional.

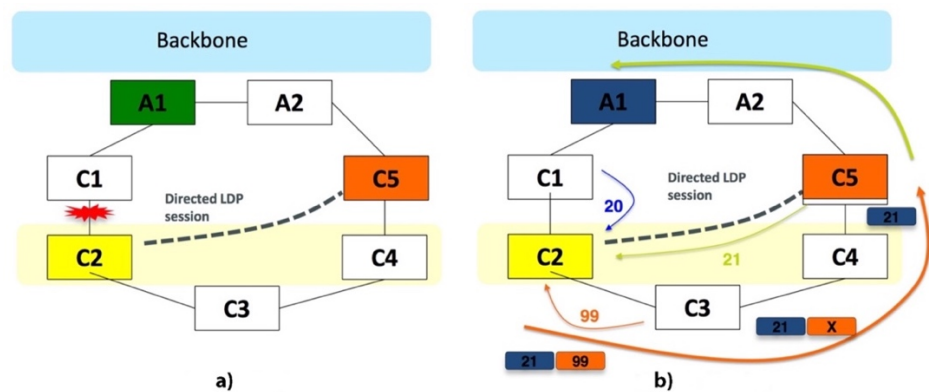


Figura 3.9 Funcionalidad de la RFC 7490 Remoto LFA FRR ^[43]

- c) BFD (*Bidirectional Forward Detection*) en el protocolo IS-IS, BFD ¹ es un protocolo que permite detectar caídas en el plano de datos e informar al plano de control para tomar decisiones de re-enrutamiento. BFD opera en modo

¹ BFD tomado de la referencia [42]

unicast punto a punto, utilizando UDP como medio de transporte con puerto de destino 3784 y puerto origen en el rango 49252-65535. Los siguientes protocolos del plano de control son típicos clientes de BFD que lo utilizan para una detección rápida de fallas en el plano de datos : BGP, IS-IS, OSPF, MPLS Traffic-Engineering, HSRP y Static routing.

Cuando BFD detecta que un vecino está inalcanzable, notifica a los protocolos de control registrados como clientes utilizan esta información para realizar sus labores de reprogramación del plano de control y de datos. En este sentido BFD realiza una función de detección de fallas similar al que realizan los protocolos de ruteo mediante paquetes HELLO, con la ventaja de una mejor escalabilidad debido a que:

- BFD es independiente del protocolo de control. De esta forma, varios protocolos de control pueden utilizar la misma sesión BFD para detectar fallas sobre enlaces de red.
- BFD entrega un mejor desempeño y granularidad en sus temporizadores, requisito fundamental para implementar una rápida convergencia.
- BFD está optimizado para detectar fallas en el plano de datos. En plataformas como los equipos de core (CRS o ASR9K) la implementación de BFD es distribuida, mientras que en otras es centralizada.

Los modos más comunes en los que opera BFD son Asynchronous y Echo Mode. El modo asíncrono es mandatorio y permite realizar detecciones del

plano de datos. El modo Echo es opcional y permite realizar detecciones de falla solo en el camino entre ambos nodos.

Para nuestro caso se utiliza BFD en todos los enlaces que utilicen la red de transporte y que requieran rápida detección y recuperación de fallas, utilizando como modo de operación el tipo Asíncrono para diferentes clientes, tales como: BGP, IS-IS y MPLS Traffic-Engineering. Los temporizadores que se utilizan son los indicados en la tabla 3.4

Tipo de Enlace	Timers BFD
Entre core - core , entre Pre-Agregación y core	Tx Interval: 100 ms Rx Interval: 100 ms Multiplier: 4
Acceso, Pre-Agregación	Tx Interval: 100 ms Rx Interval: 100 ms Multiplier: 3

Tabla 3.4 Valores de temporizadores para BFD

d) IS-IS SPF (*Support for Priority-Driven IP Prefix*) para prioridad de prefijos^[42]

Esta es una de las funcionalidades características de IS-IS; después de realizar el cálculo de SPF, IS-IS debe instalar las rutas actualizadas en la RIB (*Routed Information Base*) o tabla de enrutamiento del router, si el número de prefijos anunciados por IS-IS es grande, el tiempo que toma realizar este proceso puede ser significativo en términos de convergencia luego de un evento. La priorización de prefijos permite que un subconjunto de los prefijos anunciados por IS-IS sea designado como de alta prioridad; todas las actualizaciones relacionadas con dicho subconjunto serán instalados en la RIB antes que el resto de prefijos permitiendo de esta manera que el tiempo

de convergencia se reducido. En IS-IS se distinguen tres niveles de importancia: Prefijos de prioridad, alta, media y baja.

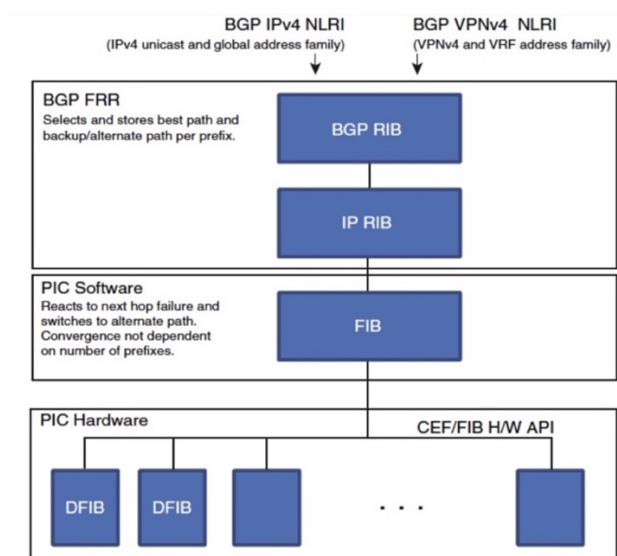


Figura 3.10 Funcionalidad de la RFC 7490 Remoto LFA FRR ^[43]

- e) En la configuración de la VPN, se recomienda configurar BGP FRR¹ (*Fast Reroute*) conocido en cisco como BGP PIC (*Prefix-Independent Protection*). Es una funcionalidad de los sistemas operativos de equipos de enrutamiento de algunos proveedores (por ejemplo Juniper o Cisco), que permite utilizar caminos alternos dentro del protocolo BGP y antes de una falla tenerlos activos, de tal manera que entren a operar lo antes posible. Este algoritmo de BGP utiliza un puntero para mover las rutas que se tienen activas por un *next-hop* al nuevo *next-hop* y no recalcular salto a salto como tradicionalmente trabaja el BGP; no depende de la cantidad de prefijos y no requiere

¹ Actualmente en la IETF es un draft, [draft-rtgwg-bgp-pic-02.txt](#)

configuraciones complejas, solo activar la funcionalidad tanto en los PEs como en los *router-reflector*. En la figura 3.10 se esquematiza esta funcionalidad identificando las tablas de enrutamiento involucradas en el algoritmo, se tiene la RIB tanto BGP como en el IGP, la FIB (*Forwarding Information Base*) , esta última representa ya a nivel de interface la correspondencia con el enrutamiento y el medio físico lo que hace que con esta base los paquetes tomen directamente la interface de salida y esta solución no dependa del número de prefijos.

Como se los puede ver en la figura 3.11, se muestra en los sitios que estas configuraciones de mejora de convergencia ayudarán a que se reduzca el tiempo de convergencia el momento de alguna falla.

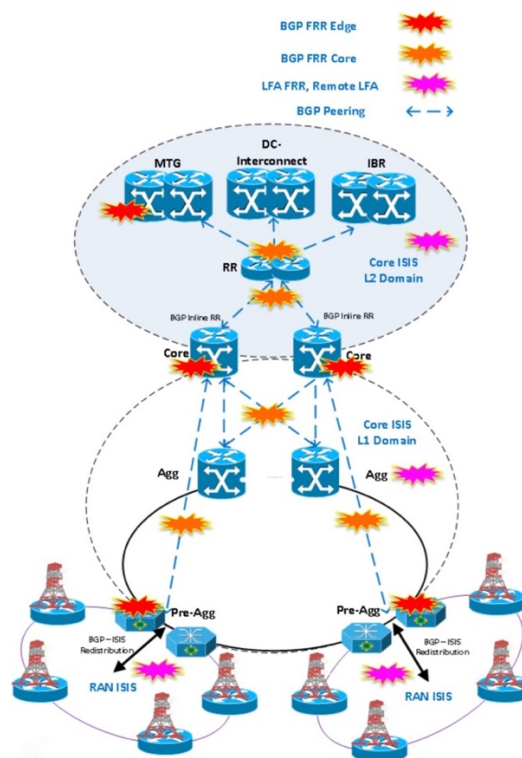


Figura 3.11 Esquema de mecanismo de resiliencia ^[27]

3.2.2.6 Calidad de servicio en los servicios LTE

En una red de transporte, la congestión puede ocurrir en cualquier segmento de la red, sin embargo los segmentos mas sensibles es donde se ha dimensionado una capacidad limitada, ya sea por razones económicas o por limitaciones de los proveedores; de todas maneras siempre que se diseña calidad de servicio en una red se recomienda tomarla en cuenta de extremo a extremo, el modelo a seguir esta basado en el RFC 2475 (*DiffServ*). Siempre que se trata de calidad de servicio, se sobreentiende que estará presente una congestión en los puntos mas vulnerables de la red, es entonces que el tráfico más sensible debe tener prioridad en cada segmento de red, definiéndose políticas en los equipos que demarcan el paso del tráfico de un segmento a otro. En el diseño de QoS que se vio en la sección 2.1.6, no se tomó en cuenta el tráfico de control, es decir el tráfico del protocolo de enrutamiento, el tráfico de sincronismo y tampoco el tráfico de gestión de la red, por esta razón se ha tomado el modelo de la referencia [27], que es un esquema que si considera lo antes mencionado y que se ilustra en la tabla 3-5, aquí se puede observar que se toma en cuenta el tráfico indicado, además se verifica la correspondencia del tipo del tráfico dependiendo del tipo de segmento de la red, se tiene entonces: la clasificación del tipo de tráfico (LTE QCI), el tipo de encolamiento seleccionado, la equivalencia en el *Traffic Class* o *Experimental bit* de la cabecera MPLS, para continuar con el QoS a través de la red MPLS y por último la equivalencia con los equipos de acceso en DSCP, consideradas como interfaces hacia los usuarios UNI(*User Network Interface*); en el modelo Diffserv se tiene la clasificación de calidad de servicio que se la denomina PHB (*Per Hop Behavior*) .

Traffic Class	LTE QCI	DiffServ PHB	Queue Type	Core, Aggregation & Access NNI MPLS/EXP	Mobile Access UNI DSCP
Network Management, LTE OAM	7	AF	CBWFQ	7	CS7
Network Routing/Control Protocols	6	AF	CBWFQ	6	CS6
Network Sync (1588v2 PTP) Real Time Voice, Video & Gaming LTE Signaling	1 2 3	EF	Priority	5	EF EF CS5
Reservado	4	AF	CBWFQ	4	AF41 / AF42 / AF43
Hosted Video	5	AF	CBWFQ	3	AF31 / AF32 / AF33
Reservado	8	AF	CBWFQ	1	AF11 / AF12 / AF13
Internet Best Effort	9	BE	CBWFQ	0	0

Tabla 3.5 Modelo de QoS para el servicio LTE ^[27]

En las figuras 3.12 y 3.13 se tiene, de manera descriptiva los mecanismos del modelo de calidad de servicio¹, tanto para el tráfico desde el usuario a la red de servicios “tráfico de subida”(o *Upstream* en inglés), como para el tráfico desde la red de servicios hasta el usuario conocido como tráfico de bajada (o de *Downstream* en inglés). Se debe tomar en cuenta también, que en una red móvil los accesos a los eNodB pueden ser inalámbricos, para este caso a las políticas definidas se las aplica dentro de una técnica conocida como parent shaping, la cual es una aplicación de los equipos de red que permite ajustar una política de calidad de servicio dentro de la limitación de la capacidad del canal, aprovechando los recursos de hardware de las interfaces físicas para almacenar los datos mientras sea posible hasta poder enviarlos . En las figuras indicadas se muestra cada segmento de red con una representación de la política aplicada, además se indica el tipo de tráfico si es de tiempo real RT, con aseguramiento de calidad AF(Assurance Forwardin) o sin calidad asegurada BE(Best Efford).

¹ En la referencia [44] se encuentra una amplia visión sobre calidad de servicio

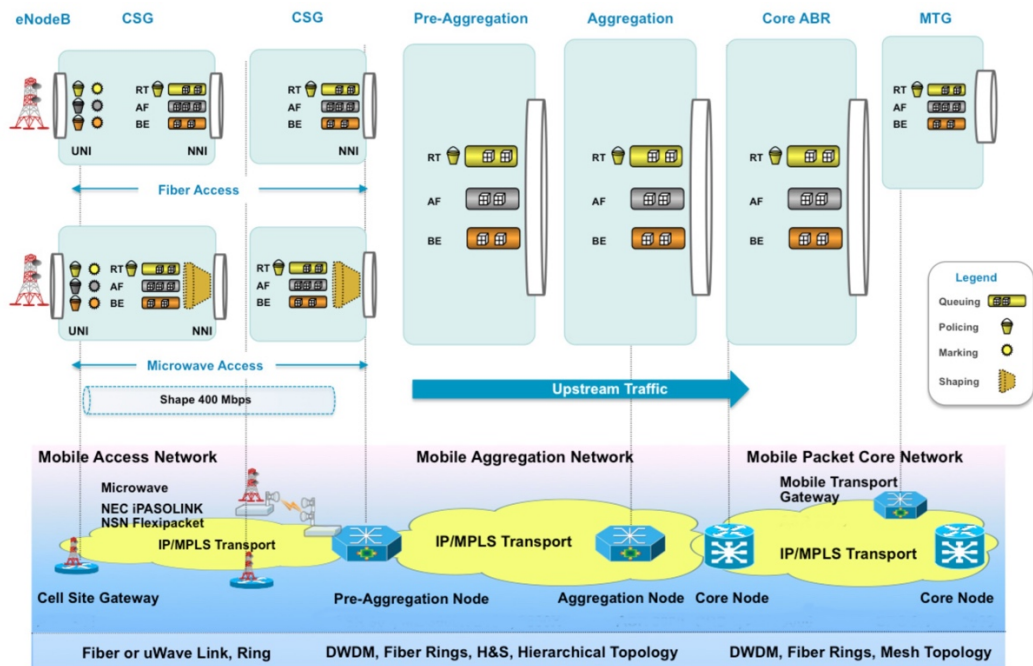


Figura 3.12 Modelo de Calidad de servicio del tráfico de subida [30]

Una vez clasificado el tráfico como en la tabla 3-5, las políticas de calidad de servicio se van configurando, en cada interface de entrada y a continuación en cada interface en el sentido del tráfico bajo las siguientes recomendaciones:

- a) Una vez clasificado el tráfico en las políticas se procede de dos maneras, si es tráfico en tiempo real, se garantiza un ancho de banda para este tráfico que además solo puede ser la Voz sobre IP. Para el resto de tráfico se garantiza un ancho de banda mínimo según la prioridad de cada tráfico. Por esto en calidad de servicio siempre será importante tener una muestra real de tráfico y seguir de esta manera ajustando esta garantía de ancho de banda, siempre habrá un momento en que la única solución sea ampliar el canal de transmisión.

- b) Los nodos de Core, Agregación y Acceso hacia la red o NNI (*Network to Network Interface*), al ser interfaces que están en redes MPLS trabajan con políticas en base a la clasificación según los bits TC o EXP.
- c) En las secciones de la red con acceso de microonda como es el caso de ciertos eNB, se trabaja con el mismo tipo de política pero dentro de una política de *shaping*.
- d) Los nodos de acceso de usuarios UNI (*User Network Interface*), las políticas aplicadas van a depender del equipamiento de acceso y del equipo del eNB, para confiar en la marcación o remarcar los paquetes.

3.2.2.7 Sincronismo en los servicios LTE

El equipamiento seleccionado tiene la capacidad de trabajar de dos maneras de distribución de sincronismo, en frecuencia y fase. Para el método de sincronismo en frecuencia existe un protocolo conocido como SyncE (*Synchronization using the Ethernet physical layer*) y para el método de fase se tiene el protocolo IEEE 1588-2008 PTP (*Precise Time Protocol*) o también conocido como PTPv2. Debido a que por el momento no se puede acceder a fuentes de sincronía por frecuencia, se ha diseñado la sincronía con PTPv2, pero queda la posibilidad de tener a futuro un esquema mixto que resultaría incluso redundante.

En la figura 3.14 se plantea la jerarquía del sincronismo, las fuentes principales ingresarían por los equipos MTG, los mismos que distribuyen la señal a los equipos

de Core, PAN y finalmente a los CSG, cabe mencionar que los CSG también poseen un sistema de sincronismo con GPS en caso de pérdida de la señal de sincronismo.

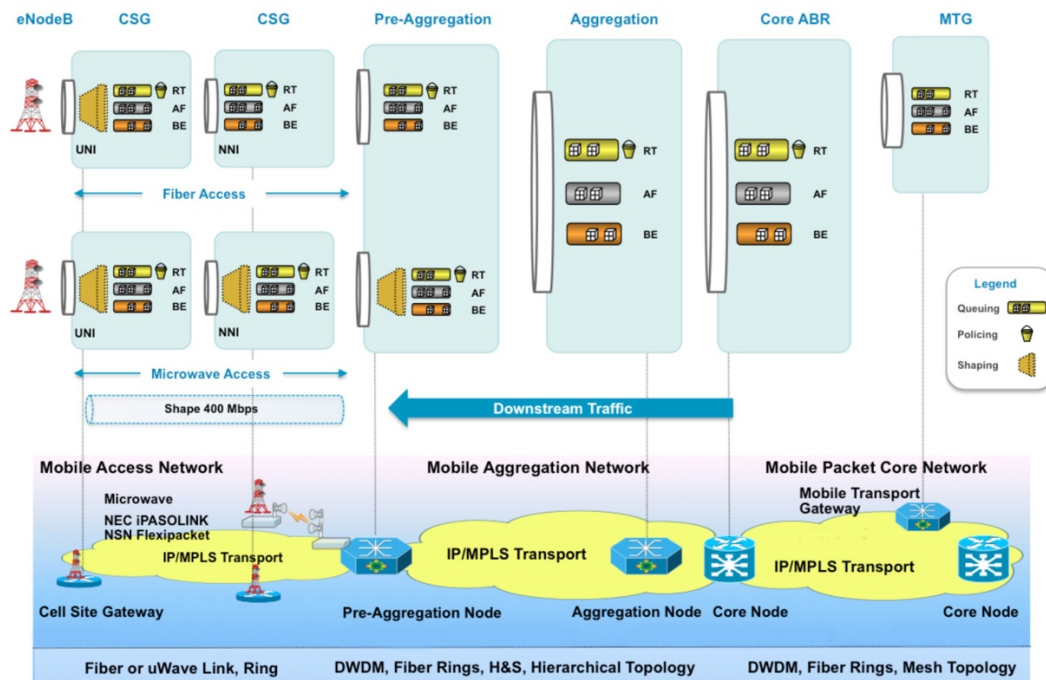


Figura 3.13 Modelo de Calidad de servicio del tráfico de bajada [30]

3.3 PLAN DE MIGRACIÓN DE FUNCIONALIDADES A LA NUEVA SOLUCIÓN^[45]

Una de las ventajas de la solución *Seamless MPLS*, es que no es una nueva topología que debe obligatoriamente implantarse en toda la red existente, sino que permite añadir nuevos segmentos de red, con las características de menor requerimiento en los equipos de acceso todo esto sin tener que realizar un cambio

radical en toda la red y en las ventanas de mantenimiento afectar masivamente al servicio de los clientes. Entonces el plan de migración puede ser progresivo y de no muy alto impacto en la interrupción de servicios, además existen trabajos que se pueden implementar en paralelo antes de la ventana de mantenimiento, en esta sección se ha puesto énfasis en las configuraciones de los equipos, que es la parte central del plan de migración, en cuanto a la logística y planificación no hacen parte relevante en este estudio, la información de las configuraciones presentadas fueron basadas en una guía confidencial por parte del proveedor Cisco Systems de la referencia [45] y adaptadas a la topología diseñada.

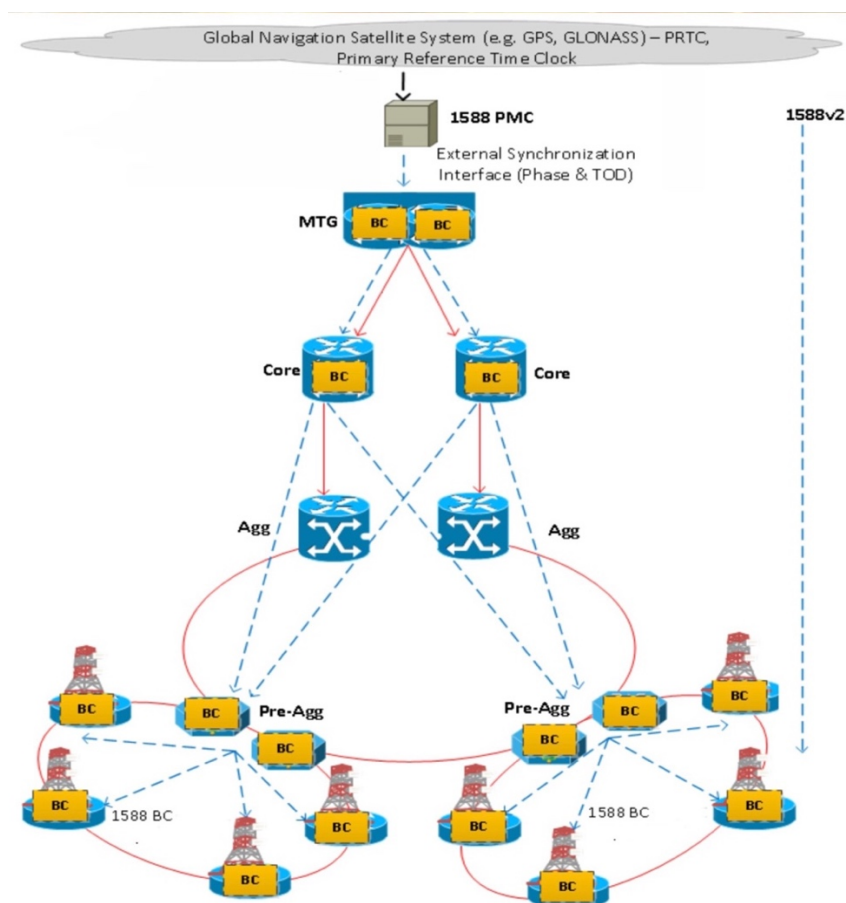


Figura 3.14 Esquema de sincronismo ^[27]

Tomando en cuenta lo antes mencionado el plan puede ser progresivo y de bajo impacto en la red, entonces a continuación se plantean las tareas a realizarse para alcanzar el objetivo de migrar la red actual IP-RAN LTE a la nueva solución con *Seamless MPLS*.

Al tratarse de una implementación que puede ser progresiva y es por esto que se ha determinado el siguiente plan, primero se prepara cada capa del nuevo diseño y posteriormente se irá migrando cada una de las tres zonas definidas (ver figura 3.4 o tabla 3-1), de esta manera la transición es gradual y el tiempo sin servicio es menor. Debido a que conceptualmente las tres zonas son similares entonces se tomará como modelo la zona del centro y en la figura 3.15 se encuentra representada, además con las direcciones IP de las interfaces *loopback*.

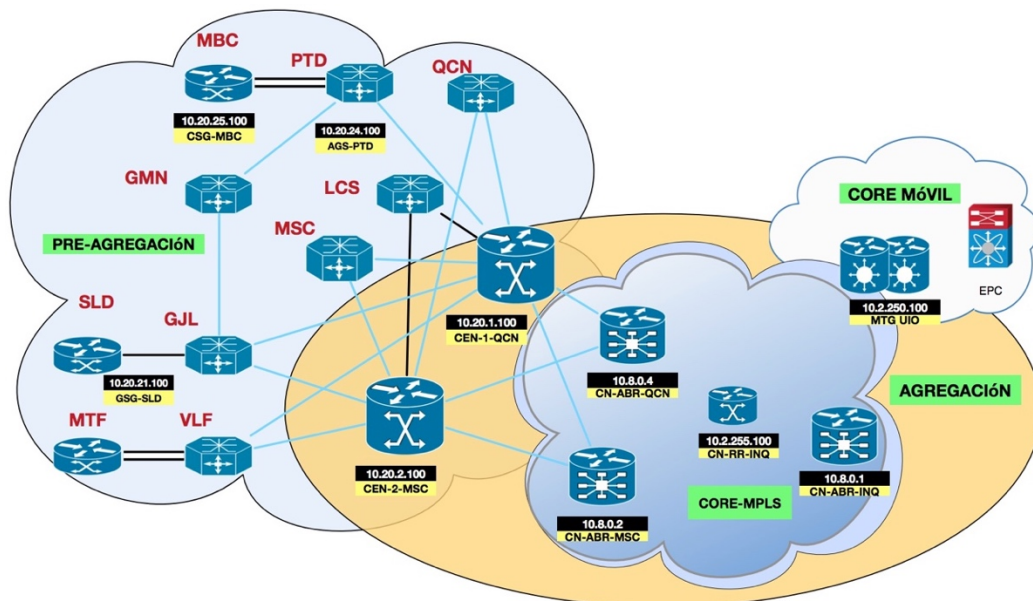


Figura 3.15 Modelo de Topología de Migración del proyecto de reingeniería de la IP-RAN

3.3.1 PREPARACIÓN DE LA CAPA DE CORE Y ROUTE-REFLECTOR

Luego de haber seleccionado los equipos de core en la red, tanto por ubicación topológica como por capacidades de Hardware, a continuación se tiene las configuraciones principales tanto del route-reflector como de los equipos de core. En cuanto al *route-reflector* se debe considerar las sesiones iBGP con los equipos de core y las sesiones iBGP con los equipos MTG, tal como lo representa la figura 3.16 y se deben considerar las siguientes configuraciones necesarias

a) Configuración del *Route-Reflector* CN-RR-INQ

Este equipo mantiene las sesiones iBGP hacia los CN-ABR dentro del core y con los MTG de los servicios, para que todos los equipos hagan una sola sesión BGP en lugar de tener varias sesiones, una sesión con cada equipo. Es importante tomar en cuenta que este equipo realizará el filtrado de las rutas de cada región para que no se vean entre sí, a menos que sean contiguas, en el caso de Quito no se presentó este caso.

Lo siguiente es la configuración de conectividad, es decir la manera en que este equipo se configura las interfaces, el protocolo de enrutamiento interno, continuamos con el protocolo de distribución de etiquetas, que en este caso es LDP y finalmente con la configuración de las sesiones iBGP hacia los CNT-ABR y los MTG

Configuración de Interfaces

```
interface Loopback100
  description Global Loopback
  ipv4 address 10.2.255.100 255.255.255.255
!
interface GigabitEthernet0/1/0/0    <<< Core interface
  description Hacia CN-ABR-QCN Gig0/2/0/9
  cdp
  ipv4 address 10.80.2.250 255.255.255.252
  negotiation auto
  load-interval 30
!
interface GigabitEthernet0/1/0/1    <<< Core interface

  description Hacia CN-ABR-MSD Gig0/3/0/9
  cdp
  ipv4 address 10.80.3.30 255.255.255.252
  negotiation auto
  load-interval 30
```

Configuración de IGP (IS-IS) y LDP

```
router isis core
  set-overload-bit on-startup 360
  net 49.0008.0100.0225.5100.00
  log adjacency changes
  lsp-gen-interval maximum-wait 5000 initial-wait 50 secondary-wait 200
  lsp-refresh-interval 65000
  max-lsp-lifetime 65535
  address-family ipv4 unicast
    metric-style wide
  ispf
  spf-interval maximum-wait 5000 initial-wait 50 secondary-wait 200
!
interface Loopback100
  passive
  address-family ipv4 unicast
!
!
interface GigabitEthernet0/1/0/0
  circuit-type level-2-only    <<< IS-IS en el core esta en L2
  bfd minimum-interval 15
  bfd multiplier 3              <<< BFD en la interface y IS-IS
  bfd fast-detect ipv4
  point-to-point
  address-family ipv4 unicast
!
!
interface GigabitEthernet0/1/0/1
  circuit-type level-2-only
  bfd minimum-interval 15
  bfd multiplier 3
  bfd fast-detect ipv4
  point-to-point
  address-family ipv4 unicast
!
mpls ldp                      <<< Configuración de LDP
  router-id 10.2.255.100
  graceful-restart
  log
neighbor
  graceful-restart
!
interface GigabitEthernet0/1/0/0
!
interface GigabitEthernet0/1/0/1
```

Configuración BGP

```
router bgp 28006
nsr
  bgp router-id 10.2.255.100
  address-family ipv4 unicast
    additional-paths receive
    additional-paths send
    additional-paths selection route-policy add-path-to-ibgp
  network 10.2.255.100/32
  allocate-label all
!
  address-family vpnv4 unicast
... se omiten lineas innecesarias
  session-group infra <<< session group para iBGP con CN-ABRs y
MTGs
  remote-as 28006
  password encrypted 082D4D4C
  update-source Loopback100
use session-group infra
  address-family ipv4 labeled-unicast
    route-reflector-client
    maximum-prefix 150000 85 warning-only
  address-family vpnv4 unicast
... se omiten lineas innecesarias
! !
  neighbor-group cn-abr <<< neighbor group para CN-ABRs in PoP-
1,2,3... etc.
  use session-group infra
  address-family ipv4 labeled-unicast
    route-reflector-client
    route-policy BGP_Egress_Transport_Filter out <<< Filtro de salida
que aisle las regiones
y que tiene todas loopbacks del IP-RAN de todas las regiones.
! !
  neighbor 10.8.0.1
  use neighbor-group cn-abr
!
  neighbor 10.8.0.2
  use neighbor-group cn-abr
!
  neighbor 10.8.0.3
  use neighbor-group cn-abr
!
  neighbor 10.8.0.21
  use neighbor-group cn-abr
!
  neighbor 10.8.0.110
  use neighbor-group cn-abr
!
  neighbor 10.2.250.100
  use neighbor-group mtg
!
  neighbor 10.2.1.100
  use neighbor-group mtg
!
!
route-policy BGP_Egress_Transport_Filter
  if community matches-any Deny_Transport_Community then
drop else
pass endif
end-policy
community-set Deny_Transport_Community
  28006:00001 <<< Esto se definió en la sección 3.2.2.3
end-set
```

b) Configuración del equipo de core CN-ABR-MSR.

A continuación se detallan las partes de la configuración del equipo de Core que se estiman mas relevantes para entender el funcionamiento y que se debe replicar

para cada equipo de Core, en este caso el equipo es el CN-ABR-MSC de la figura 3.15

Configuración de Interfaces del equipo CN-ABR-MSC

```
interface Loopback100
  ipv4 address 10.8.0.2 255.255.255.255
!
interface TenGigE0/7/0/1 <<< Interconexión con otro Core CN-ABR-QCN
  cdp
  service-policy output PMAP-NNI-E
  ipv4 address 10.80.0.13 255.255.255.252
  load-interval 30
  frequency synchronization
!
interface TenGigE0/0/0/0 << Interconexión dentro del Core
  cdp
  service-policy output PMAP-NNI-E
  ipv4 address 10.80.0.6 255.255.255.252
  load-interval 30
!
interface TenGigE0/4/0/5 <<< Interconexión Agregación CEN-1-QCN
  cdp
  service-policy output PMAP-NNI-E
  ipv4 address 10.80.2.45 255.255.255.252
  carrier-delay up 2000 down 0
  load-interval 30
  frequency synchronization
```

Configuración de IS-IS y protocolo LDP

```
router isis core-agg
  net 49.0008.0100.0800.0002.00
  nsf cisco
  log adjacency changes
  lsp-gen-interval maximum-wait 5000 initial-wait 50 secondary-wait 200
  lsp-refresh-interval 65000
  max-lsp-lifetime 65535
  address-family ipv4 unicast
    metric-style wide
    spf-interval maximum-wait 5000 initial-wait 50 secondary-wait 200
  propagate level 1 into level 2 route-policy drop-all <<< Disable default
  IS-IS L1 to L2 redistribution
!
interface Loopback100

passive
  point-to-point
  address-family ipv4 unicast
tag 1000 !
!
interface TenGigE0/7/0/1 <<< L1/L2 link
  bfd minimum-interval 15
  bfd multiplier 3
  bfd fast-detect ipv4
  point-to-point
  link-down fast-detect
  address-family ipv4 unicast
    mpls ldp sync
!
!
interface TenGigE0/0/0/0
  circuit-type level-2-only <<< L2 link
  bfd minimum-interval 15
  bfd multiplier 3
  bfd fast-detect ipv4
  point-to-point
  link-down fast-detect

  address-family ipv4 unicast
    mpls ldp sync
!
```

```

!
interface TenGigE0/4/0/5
  circuit-type level-1 <<< L1 link
  bfd minimum-interval 15
  bfd multiplier 3
  bfd fast-detect ipv4
  point-to-point
  link-down fast-detect
  address-family ipv4 unicast
  mpls ldp sync
!
!
route-policy drop-all <<< Route policy to disable IS-IS L1 to L2
redistribution
  drop
end-policy
mpls ldp
router-id 10.8.0.2
nsr
graceful-restart
session protection
!
interface TenGigE0/0/0/0
!
interface TenGigE0/4/0/5
!
interface TenGigE0/7/0/1

```

Configuración del protocolo BGP

```

router bgp 28006
  nsr
  bgp router-id 10.8.0.2
  bgp cluster-id 1001 <<< Redundant ABRs K1001 and K1002 have a cluster ID 1001 to
prevent loops.
  bgp graceful-restart
  ibgp policy out enforce-modifications
  address-family ipv4 unicast
  additional-paths receive <<< BGP add-path to receive multiple paths
from CN-RR
  additional-paths selection route-policy add-path-to-ibgp
  network 10.8.0.2/32 route-policy CN_ABR_Community <<< Color loopback
prefix in BGP with CN_ABR_Community
  allocate-label all
!
  address-family vpnv4 unicast
!
  session-group infra
  remote-as 100
  password encrypted 03085A09
  cluster-id 1001
  update-source Loopback100
  graceful-restart
!
  neighbor-group cn-rr <<< iBGP neighbor group para el route-reflector del
Core CN-RR
  use session-group infra
  address-family ipv4 labeled-unicast <<< Address family for RFC 3107 based
transport
  maximum-prefix 150000 85 warning-only
  next-hop-self <<< next-hop-self para insertar el Nuevo camino de
enr
!
  neighbor-group pan <<< iBGP neighbor group para los nodos de Pre-
Aggregation network(PAN)
  use session-group infra
  address-family ipv4 labeled-unicast <<< Address family para RFC 3107
bgp+label
  route-reflector-client <<< PANs son route-reflector Clients del Core
  next-hop-self <<< next-hop-self to insert into data path
!
  address-family vpnv4 unicast

```

```

!
neighbor 10.2.255.100    <<< iBGP con CN-RR
use neighbor-group cn-rr
!
neighbor 10.20.1.100    << iBGP con Pre agregación CEN-1-QCN
use neighbor-group agg
!
neighbor 10.20.2.100    << iBGP con Pre agregación CEN-2-MS
use neighbor-group agg
!
neighbor 10.20.3.100    << iBGP con Pre agregación VAL-1-ETT
use neighbor-group agg
!
neighbor 10.20.4.100    << iBGP con Pre agregación VAL-2-QCN
use neighbor-group agg
!
neighbor 10.20.5.100    << iBGP con Pre agregación NOR-1-INQ
use neighbor-group agg
!
neighbor 10.20.6.100    << iBGP con Pre agregación NOR-1-MS
use neighbor-group agg
!
route-policy add-path-to-ibgp
set path-selection backup 1 install
end-policy
!
route-policy CN_ABR_Community
set community CN_ABR_Community
end-policy
community-set CN_ABR_Community
28006:00003
end-set

```

c) Configuración de un equipo de MTG-UIO.

Se ha tomado en cuenta la configuración de este equipo, para confirmar como el servicio llega a los nodos finales. Los MTG deben estar interconectados al Core y además con una sesión iBGP hacia el route-reflector, como se muestra a continuación.

Configuración de las interfaces MTG-UIO

```

interface Loopback100
description Global Loopback
ipv4 address 10.2.250.1 255.255.255.255
!
interface TenGigE1/0/0/1
description Hacia CN-ABR-MS Ten0/4/0/7
cdp
service-policy output PMAP-NNI-E
ipv4 address 10.82.250.6 255.255.255.252
carrier-delay up 2000 down 0
load-interval 30
transceiver permit pid all
!
interface TenGigE0/0/0/1
description Hacia CN-ABR-INQ Ten0/0/0/2
cdp
service-policy output PMAP-NNI-E
ipv4 address 10.82.250.2 255.255.255.252
carrier-delay up 2000 down 0
load-interval 30
transceiver permit pid all

```

Configuración del protocolo interior IS-IS y protocolo LDP

```
router isis core
 set-overload-bit on-startup 400
 net 49.0008.0100.0225.0100.00
 nsf cisco
 log adjacency changes
 lsp-gen-interval maximum-wait 5000 initial-wait 50 secondary-wait 200
 lsp-refresh-interval 65000
 max-lsp-lifetime 600
 address-family ipv4 unicast
 metric-style wide

ispf
 spf-interval maximum-wait 5000 initial-wait 50 secondary-wait 200
!
interface Loopback100
 passive
 point-to-point
 address-family ipv4 unicast
!
!
interface TenGigE1/0/0/1
 circuit-type level-2-only
 bfd minimum-interval 15
 bfd multiplier 3
 bfd fast-detect ipv4
 point-to-point
 address-family ipv4 unicast
 fast-reroute per-prefix
 mpls ldp sync
!
!
interface TenGigE0/0/0/1
 circuit-type level-2-only
 bfd minimum-interval 15
 bfd multiplier 3
 bfd fast-detect ipv4
 point-to-point
 address-family ipv4 unicast
 mpls ldp sync
!
!
mpls ldp
 router-id 10.2.250.100
 discovery targeted-hello accept
 nsr
 graceful-restart
 session protection
 log
 neighbor
 graceful-restart
 session-protection
 nsr
!
interface TenGigE1/0/0/0
!
interface TenGigE0/0/0/1
```

Configuración de las sesiones BGP

```
router bgp 28006
 nsr
 bgp router-id 10.2.250.100
 bgp redistribute-internal
 bgp graceful-restart
 ibgp policy out enforce-modifications
 address-family ipv4 unicast
 additional-paths receive <<< Configuración BGP PIC o multiples caminos
recibidos por
 el route-Reflector del core.
 additional-paths selection route-policy add-path-to-ibgp
```

```
network 10.2.250.100/32 route-policy MTG_Community <<< Se impone la
comunidad común del
```

servicio que se definió en la sección

3.2.2.3

```
in BGP with MPC_Community
  allocate-label all
!
address-family vpnv4 unicast
!
!
session-group infra
  remote-as 100
  password encrypted 011F0706
  update-source Loopback100
!
neighbor-group cn-rr
  use session-group infra
  address-family ipv4 labeled-unicast
  maximum-prefix 150000 85 warning-only
  next-hop-self
  !
  address-family vpnv4 unicast
!!
neighbor 10.2.255.100 <<< Sesión Hacia el route reflector
  use neighbor-group cn-rr
!
route-policy MTG_Community
  set community MTG_Community
end-policy
!
community-set MTG_Community
  28006:00002
end-set !
route-policy add-path-to-ibgp
  set path-selection backup 1 install
end-policy
```

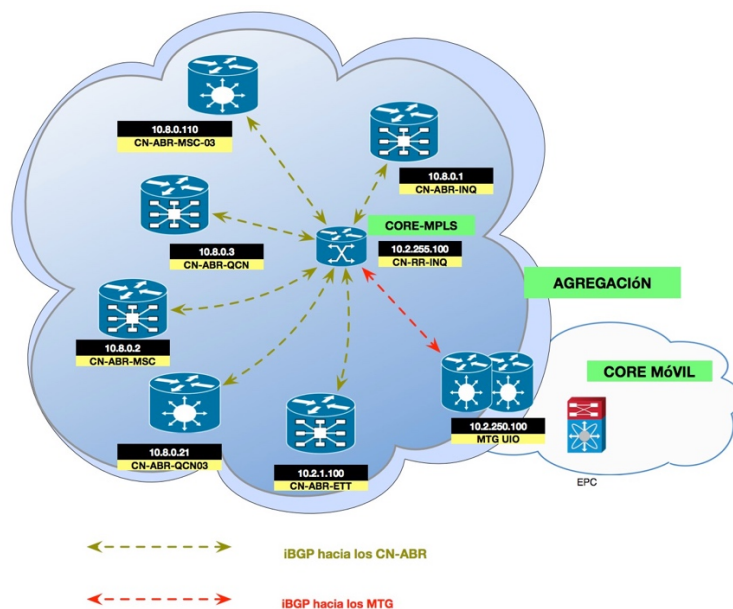


Figura 3.16 Sesiones iBGP del Route-Reflector del

3.3.2 PREPARACIÓN DE LA CAPA DE PRE-AGREGACIÓN

El equipo que hace el enlace entre el core y las zonas de pre-agregación son los equipos que se han denominado PAN (*Pre-Agregation Network*), son a su vez clientes de los route-reflectors CN-ABRs y al mismo tiempo son *route-reflectors* de cada zona de IP-RAN, es aquí también donde el IGP vuelve a dividirse entre la capa de agregación IS-IS L1 y IS-IS L1, pero en otro proceso de enrutamiento, a continuación se presenta la configuración que servirá para la implementación de la nueva red IP-RAN.

Configuración de Interfaces, IS-IS y LDP de CEN-1-QCN

```
interface Loopback100 ip                <<< Loopback para core-agg Proceso IS-IS
 ip address 10.20.1.100 255.255.255.255
!
interface Loopback101                  <<< Loopback para el proceso IS-IS de la IP-RAN de acceso
 ip address 10.20.1.99 255.255.255.255
!
mpls ldp router-id Loopback100 force
!
interface TenGigabitEthernet0/1        <<< Interface hacia la red de Core
 description Hacia CN-ABR-MS
 no switchport
 ip address 10.80.0.5 255.255.255.252
 ip router isis core-agg
 mpls ip
 synchronous mode
 bfd interval 50 min_rx 50 multiplier 3
 isis network point-to-point
!
interface TenGigabitEthernet0/2        <<< Interface hacia la RAN, proceso de IS-IS de acceso
 description Hacia AGS-PTD
 ip address 10.9.1.0 255.255.255.254
 ip router isis ran                    <<< Proceso IS-IS denominado ran, Segundo proceso luego
de core-agg
 mpls ip
 mpls ldp discovery transport-address 100.20.1.99 <<< LDP discovery cambia de una zona a
otra del
                                     IS-IS, ahora ala loopback101
 synchronous mode
 bfd interval 500 min_rx 500 multiplier 3
 isis circuit-type level-1
 isis network point-to-point
```

Configuración de IS-IS y LDP

```
router isis core-agg
 net 49.0008.0100.2001.0100.00
 is-type level-1
 ispf level-1
 metric-style wide
 fast-flood
 max-lsp-lifetime 65535
 lsp-refresh-interval 65000
 spf-interval 5 50 200
 prc-interval 5 50 200
 lsp-gen-interval 5 50 200
 no hello padding
 log-adjacency-changes
 passive-interface Loopback100
 bfd all-interfaces
```

```

!
router isis ran
 net 49.0009. 0100.2001.0100.00    <<< La red IP-RAN se encuentra en una area de IS-IS
diferente
 is-type level-1                    <<< IS-IS Level-1
 ispf level-1
 metric-style wide
 fast-flood
 max-lsp-lifetime 65535
 lsp-refresh-interval 65000
 spf-interval 5 50 200
 prc-interval 5 50 200
 lsp-gen-interval 5 50 200
 no hello padding
 log-adjacency-changes
 redistribute isis core-agg ip route-map Local_ABR_Loopback    <<< Redistribución loopback
ABR
 redistribute bgp 100 route-map BGP_TO_RAN_ACCESS level-1 <<< Redistribución MTG/EPC,por
comunidades
 passive-interface Loopback101    <<< interface loopback del proceso IS-IS de la RAN
 bfd all-interfaces
!
route-map BGP_TO_RAN_ACCESS permit 10    <<< Filtro basado en la comunidad BGP para IS-
IS,solo MTG
 match community EPC_Community
 set tag 1000    <<< Se marca este prefijo con un tag de 1000
 ip community-list standard EPC_Community permit 28006:0002 <<<Definido en 3.1.2.2.3 para
los servicios EPC

```

Configuración del protocolo BGP

```

router bgp 28006
 bgp router-id 10.20.1.2
 bgp cluster-id 110
 bgp log-neighbor-changes
 bgp graceful-restart restart-time 120
 bgp graceful-restart stalepath-time 360
 bgp graceful-restart
 no bgp default ipv4-unicast
 neighbor csg peer-group    <<< Peer group para CSGs y RAN local
 neighbor csg remote-as 28006
 neighbor csg password lab
 neighbor csg update-source Loopback101 <<< Loopback101 es usada como origen
 neighbor abr peer-group    <<< Peer group para conexión con el Core, con los CN-ABRs
 neighbor abr remote-as 28006
 neighbor abr password lab
 neighbor abr update-source Loopback0 <<< Loopback100 es usada como origen por el IS-IS
 neighbor 10.8.0.4 peer-group abr
 neighbor 10.8.0.2 peer-group abr
 neighbor 10.20.21.100 peer-group csg
 neighbor 10.20.24.100 peer-group csg
 neighbor 10.20.25.100 peer-group csg
!
 address-family ipv4    <<< BGP+label con el RFC 3107
  bgp redistribute-internal
  network 10.20.1.100 mask 255.255.255.255 route-map AGG_Community << Loopbac100 con la
comunidad de RAN
  redistribute isis ran level-1 route-map RAN_ACCESS_TO_BGP
  neighbor abr send-community
  neighbor abr next-hop-self    <<< Propiedad de Next-Hope-Self para unir los LSP
  neighbor abr send-label    <<< Envío de etiquetas con BGP
  neighbor 10.8.0.4 activate    <<< CN-ABR-QCN
  neighbor 10.8.0.2 activate    <<< CN-ABR-MSC

 exit-address-family
!
 address-family vpnv4
... Se omiten líneas intencionalmente
 address-family rtfiler unicast    <<< Activa la aplicación de aprender solo lo que se
utiliza
... Se omiten líneas intencionalmente
 route-map AGG_Community permit 10
  set community 28006:00001 28006:01041

 <<< 28006:00001 es la comunidad de agregación. 28006:01041 is identificado a este equipo
respecto a la red RAN

 route-map RAN_ACCESS_TO_BGP permit 10    <<< Acceso solo si tiene el tag 10
 match tag 10

```

3.3.3 PREPARACIÓN DE LA CAPA DE ACCESO

En la fase final de la preparación de equipos, ya que constituye el sitio donde se interconecta la red IP-RAN con el nodo celular, por eso el nombre de *Cell Site Gateway*, el equipo aquí recibe únicamente las rutas de interés y debido a la redistribución realizada en por los equipos PAN al IGP no se hace necesario levantar una sesión BGP+label entre el CSG y el equipo PAN. A continuación su configuración.

Configuración de Interfaces, IS-IS y LDP de CSG-MBC

```
interface Loopback100
ip address 10.20.25.100 255.255.255.255
isis tag 10
!
interface GigabitEthernet0/10 <<< NNI: Primera interface hacia el equipo PTD
description Hacia AGS-PTD
synchronous mode
service instance 10 ethernet
encapsulation untagged
bridge-domain 10
!
interface GigabitEthernet0/11 <<< NNI: Segunda interface hacia el equipo PTD
description Hacia AGS-PTD
synchronous mode
service instance 20 ethernet
encapsulation untagged
bridge-domain 20
!

interface Vlan10
ip address 10.9.1.3 255.255.255.254
ip router isis ran
mpls ip
bfd interval 50 min_rx 50 multiplier 3
isis network point-to-point
!
interface Vlan20
ip address 10.9.1.5 255.255.255.254
ip router isis ran
mpls ip
bfd interval 50 min_rx 50 multiplier 3
isis network point-to-point
!
router isis ran
net 49.0009.0100.2025.0100.00 <<< IS-IS Level-1
is-type level-1
ispsf level-1
metric-style wide
fast-flood
max-lsp-lifetime 65535
lsp-refresh-interval 65000
spf-interval 5 50 200
prc-interval 5 50 200
lsp-gen-interval 5 50 200
no hello padding
log-adjacency-changes
passive-interface Loopback100
bfd all-interfaces
```

3.3.4 PREPARACIÓN DE LA MPLS VPN LTE

En las secciones 3.2.1, 3.2.2.3 y 3.2.2.4 se establecieron los lineamientos para la creación de las MPLS VPNs, especialmente se debe considerar los identificativos que permiten manipular el enrutamiento de las dos VRF, en este caso la de servicios que se denominó movdcn y la de gestión que se denominó movoyim_trust, estas VRF están operando actualmente y son a las que se migrará el tráfico directamente. Se presentarán las configuraciones necesarias para que el servicio S1 y X2 llegue desde el EPC hasta el cliente final, estos trabajos se deben realizar tanto en el MTG como en los CSGs y en otra sección se representará como se debe configurar las sesiones VPNv4 para que este tráfico pueda operar, a esto se ha denominado plano de control.

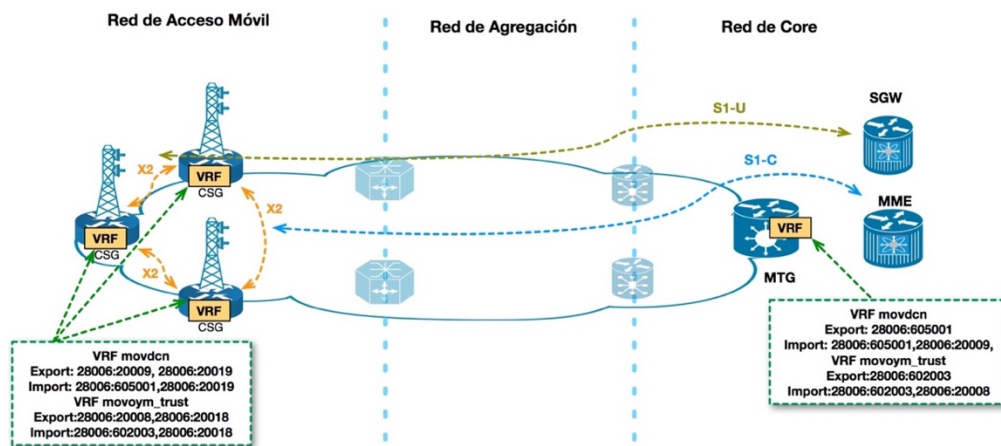


Figura 3.17 VRF movdcn de servicio y movoyim_trust de gestión con su

3.3.4.1 Preparación de la MPLS VPN LTE, servicios S1 y X2

En la figura 3.17 se ha representado la información MPLS, los route-targets que permiten el manejo del enrutamiento dentro de las dos VRF, debido a que las

configuraciones en los equipos son iguales para las dos VRFs, se desarrollará únicamente para la VRF movdcn, que es la que permite distribuir los servicios LTE a los usuarios finales. La configuración de la VRF es la siguiente sobre los CSGs y es la misma para todos, se ha añadido el soporte para BGP PIC.

Configuración para interfaces con los e-NB

```
interface GigabitEthernet0/1
description Hacia el eNB    <<< Conexión hacia el eNB
service-policy output PMAP-eNB-UNI-P-E
service instance 100 ethernet
encapsulation untagged
service-policy input PMAP-eNB-UNI-I
bridge-domain 100
!
!
interface Vlan100
vrf forwarding movdcn
ip address 113.26.23.1 255.255.255.0
!
Definición de VRF
vrf definition movdcn
rd 28006:20009
!
address-family ipv4
export map ADDITIVE
route-target export 28006:20019    <<< RT local en esta región, este caso zona 01
route-target import 28006:20019    <<< Habilita la comunicación X2 de la zona 01
route-target import 28006:605001    <<< RT del core des servicios móviles
route-target import 28006:20029    <<< Comunicación X2 con otra zona
exit-address-family
!
route-map ADDITIVE permit 10
set extcommunity rt 28006:20019 additive    <<< Filtro que permite importar rutas del CGS
hacia el core de                                servicios
```

Configuración de session VPNv4

```
router bgp 28006
!
address-family ipv4 vrf movdcn
bgp additional-paths install
bgp nexthop trigger delay 0
redistribute connected
exit-address-family
!
address-family vpnv4
bgp additional-paths install    <<< Configuración que habilita BGP PIC, tiene su similar
en el PAN.                                en el PAN.
bgp nexthop trigger delay 1
```

Configuración en de BGP en el equipo PAN

```
router bgp 28006
!
address-family ipv4
bgp additional-paths receive    <<< BGP PIC para MPLS VPN
bgp additional-paths install
bgp nexthop trigger delay 0
!
address-family vpnv4
bgp additional-paths install
bgp nexthop trigger delay 1
```

Se puede dar el caso de que el equipo PAN pueda tener eNBs directamente conectados, para ese caso se puede aplicar la misma configuración que en los CGS.

Configuración el equipo MTG(Mobile Transport Gateway)

Interface EPC

```
interface TenGigE0/0/0/2.1100
description Hacia el EPC Gateway.
vrf movdcn
ipv4 address 115.1.23.3 255.255.255.0
encapsulation dot1q 1100
```

Definición de la VRF

```
vrf movdcn
 address-family ipv4 unicast
import route-target
 28006:20009 <<< RT común de la RAN importado por MTG
 28006:605001 <<< RT propio del Servicio de Core.
!
export route-target
 28006:605001 <<< RT propio del Servicio de Core.
entire network.
```

Configuración de session VPNV4.

```
router bgp 28006
 bgp router-id 10.2.250.100
 bgp update-delay 120
vrf movdcn
 address-family ipv4 unicast
 redistribute connected
```

3.3.4.2 Preparación de la MPLS VPN LTE, plano de control

Se ha mostrado la configuración de la VRF movdcn en los equipos CSG, PAN y MTG, por tener una visión global de la relación de las sesiones BGP VPNV4 a lo largo de toda red tanto desde los MTGs como de los nodos de acceso, además es entonces de esta manera que se aclara indicando aspectos del plano de control de estas sesiones BGP VPNv4 extremo a extremo de la red implementada, así se puede comprender, de mejor manera el enrutamiento que lleva la información del servicio LTE, se representa de manera general en la figura 3.18 como se encuentran estas sesiones BGP VPv4 y luego se mostrarán las configuraciones que conforman estas sesiones entre los equipos.

Configuración de sesión VPNV4 en los CSG, ejemplo CSG-MBC.

```
router bgp 28006
  bgp router-id 10.20.25.100
  neighbor pan peer-group
  neighbor pan remote-as 28006
  neighbor 10.20.1.100 peer-group pan <<< Route-reflector equipo PAN CEN-1-QCN
  neighbor 10.20.2.100 peer-group pan <<< Route-reflector equipo PAN CEN-2-MS
  !
  address-family vpnv4
  neighbor pan send-community extended
    neighbor 10.20.1.100 activate
    neighbor 10.20.2.100 activate
  exit-address-family
  !
  address-family rtfilter unicast <<< Funcionalidad que restringe las rutas no necesarias
  de VPNV4
    neighbor pan send-community extended
    neighbor 10.20.1.100 activate
    neighbor 10.20.2.100 activate
  exit-address-family
```

Configuración de sesión VPNV4 en los PAN, ejemplo CEN-1-QCN

```
router bgp 28006
  bgp router-id 10.20.1.100
  neighbor csg peer-group <<< Peer group para CSGs y RAN local
  neighbor csg remote-as 28006
  neighbor csg password lab
  neighbor csg update-source Loopback101 <<< Loopback101 es usada como origen
  neighbor abr peer-group <<< Peer group para conexión con el Core, con los CN-ABRs
  neighbor abr remote-as 28006
  neighbor abr password lab
  neighbor abr update-source Loopback0 <<< Loopback100 es usada como origen por el IS-IS
  neighbor 10.8.0.4 peer-group abr <<< CN-ABR-QCN
  neighbor 10.8.0.2 peer-group abr <<< CN-ABR-MS
  neighbor 10.20.21.100 peer-group csg <<< CSG-SLD
  neighbor 10.20.24.100 peer-group csg <<< AGS-PTD
  neighbor 10.20.25.100 peer-group csg <<< CSG-MBC
  !
  address-family ipv4
    bgp nexthop trigger delay 2
    <SNIP>
  exit-address-family
  !
  address-family vpnv4
    bgp nexthop trigger delay 3
    neighbor csg send-community extended
    neighbor csg route-reflector-client
    neighbor abr send-community both
    neighbor 10.8.0.4 activate
    neighbor 10.8.0.2 activate
    neighbor 10.20.21.100 activate
    neighbor 10.20.24.100 activate
    neighbor 10.20.25.100 activate
  exit-address-family
  !
  address-family rtfilter unicast <<< Funcionalidad que restringe las rutas no necesarias
  de VPNV4
    neighbor csg send-community extended
    neighbor 10.20.21.100 activate
    neighbor 10.20.24.100 activate
    neighbor 10.20.25.100 activate
  exit-address-family
  !
```

Configuración de sesión VPNV4 en los CN-ABR, ejemplo CN-ABR-MS

```
router bgp 28006
  nsr
  bgp router-id 10.8.0.2
  address-family vpnv4 unicast
  !
  session-group infra
```

```

remote-as 100
password encrypted 03085A09
cluster-id 1001
update-source Loopback100
graceful-restart
!
neighbor-group cn-rr <<< iBGP neighbor group para el route-reflector del Core CN-RR
use session-group infra
address-family ipv4 labeled-unicast <<< Address family for RFC 3107 based transport
maximum-prefix 150000 85 warning-only
next-hop-self <<< next-hop-self para insertar el Nuevo camino de enrutamiento
!
neighbor-group pan <<< iBGP neighbor group para los nodos de Pre-Aggregation network(PAN)
use session-group infra
address-family ipv4 labeled-unicast <<< Address family para RFC 3107 bgp+label
route-reflector-client <<< PANs son route-reflector Clients del Core
next-hop-self <<< next-hop-self to insert into data path
!
address-family vpnv4 unicast
!
neighbor 10.2.255.100 <<< iBGP con CN-RR
use neighbor-group cn-rr
!
neighbor 10.20.1.100 << iBGP con Pre agregación CEN-1-QCN
use neighbor-group agg
!
neighbor 10.20.2.100 << iBGP con Pre agregación CEN-2-MS
use neighbor-group agg
!
neighbor 10.20.3.100 << iBGP con Pre agregación VAL-1-ETT
use neighbor-group agg
!
neighbor 10.20.4.100 << iBGP con Pre agregación VAL-2-QCN
use neighbor-group agg
!
neighbor 10.20.5.100 << iBGP con Pre agregación NOR-1-INO
use neighbor-group agg
!
neighbor 10.20.6.100 << iBGP con Pre agregación NOR-1-MS
use neighbor-group agg

```

Configuración de session VPNV4-CN-RR

```

router bgp 28006
nsr
bgp router-id 10.2.255.100
address-family ipv4 unicast
additional-paths receive
additional-paths send
additional-paths selection route-policy add-path-to-ibgp
network 10.2.255.100/32
allocate-label all
!
address-family vpnv4 unicast
... se omiten líneas innecesarias
session-group infra <<< session group para iBGP con CN-ABRs y MTGs
remote-as 28006
password encrypted 082D4D4C
update-source Loopback100
use session-group infra
address-family ipv4 labeled-unicast
route-reflector-client
maximum-prefix 150000 85 warning-only
address-family vpnv4 unicast
... se omiten líneas innecesarias
!!
neighbor-group cn-abr <<< neighbor group para CN-ABRs in PoP-1,2,3... etc.
use session-group infra
address-family ipv4 labeled-unicast
route-reflector-client

```

```

    route-policy BGP_Egress_Transport_Filter out <<< Filtro de salida que aíse las
regiones
    y que tiene todas loopbacks del IP-RAN de todas las regiones.
! !
neighbor 10.8.0.1
    use neighbor-group cn-abr
!
neighbor 10.8.0.2
    use neighbor-group cn-abr
!
neighbor 10.8.0.3
    use neighbor-group cn-abr
!
neighbor 10.8.0.21
    use neighbor-group cn-abr
!
neighbor 10.8.0.110
    use neighbor-group cn-abr
!
neighbor 10.2.250.100
    use neighbor-group mtg
!
neighbor 10.2.1.100
    use neighbor-group mtg
!
!
route-policy BGP_Egress_Transport_Filter
    if community matches-any Deny_Transport_Community then
drop else
pass endif
end-policy
community-set Deny_Transport_Community
    28006:00001 <<< Esto se definió en la sección 3.2.2.3
end-set

```

Configuración de session VPNV4-MTG

```

router bgp 28006
    nsr
    bgp router-id 10.2.250.100
    bgp redistribute-internal
    bgp graceful-restart
    ibgp policy out enforce-modifications
...
address-family vpnv4 unicast
!
!
session-group infra
    remote-as 100
    password encrypted 011F0706
    update-source Loopback100
!
neighbor-group cn-rr
    use session-group infra
    address-family ipv4 labeled-unicast
    maximum-prefix 150000 85 warning-only
    next-hop-self
!
    address-family vpnv4 unicast
! !
neighbor 10.2.255.100 <<< Sesión Hacia el route reflector
    use neighbor-group cn-rr
!
route-policy MTG_Community
    set community MTG_Community
end-policy
!
community-set MTG_Community
    28006:00002
end-set !
route-policy add-path-to-ibgp
    set path-selection backup 1 install
end-policy

```

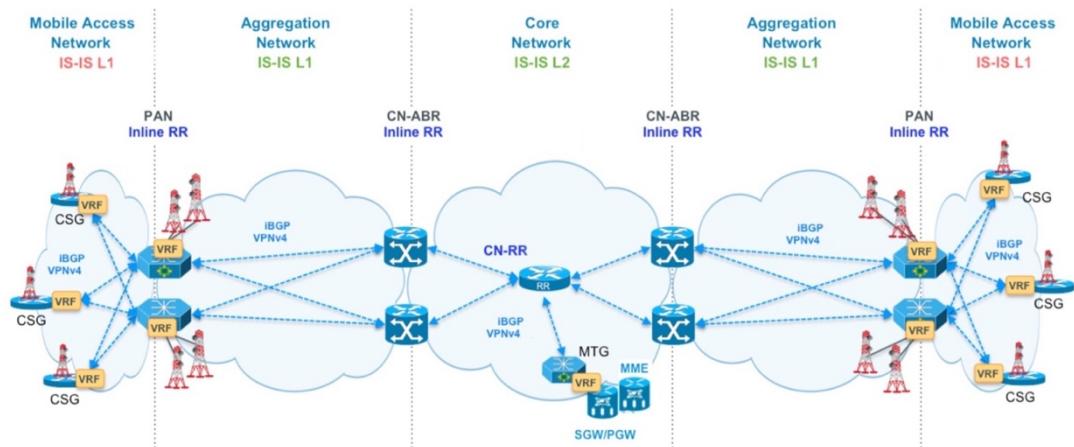


Figura 3.18 Sesiones BGP VPNv4, plano de control de la VRF movdcn

3.3.5 PREPARACIÓN DE LA RED PARA LA SINCRONÍA

En la presente sección se pone en práctica lo diseñado en la sección 3.2.2.7 y se procede a presentar la configuración requerida para el protocolo PTP (IEEE 1588) para el MTG, un CN-ABR, un PAN y un CSG.

Configuración el equipo MTG(Mobile Transport Gateway)

Interfaces hacia el reloj Maestro en MTG, con redundancia

!

```
interface GigabitEthernet0/0/1/0.300
service-policy output PMAP-NNI-E_PTP-Test
ipv4 address 200.10.1.8 255.255.255.254
load-interval 30
encapsulation dot1q 300
```

!

!

!

```
interface GigabitEthernet0/0/1/0.400
service-policy output PMAP-NNI-E_PTP-Test
ipv4 address 200.10.2.8 255.255.255.254
encapsulation dot1q 400
```

!

!

*** PTP profile ***

```
ptp
clock
domain 0
identity mac-address router
```

```
!
profile MTG-Slave      <<< Plantilla para reloj que ingresa al MTG
dscp 46
```

```

transport ipv4
port state slave-only
master ipv4 200.10.1.9
priority 120
!
master ipv4 200.10.2.9
priority 125
!
profile MTG-Master <<< Plantilla para cliente de reloj del MTG
dscp 46
transport ipv4
announce timeout 2
!

!
interface GigabitEthernet0/0/1/0.300 <<< Puerto que toma la señal de reloj "slave"
description Hacia reloj Master
ptp
profile MTG-Slave
interface GigabitEthernet0/0/1/0.400 <<< Puerto que toma la señal de reloj "slave"
description Hacia reloj Master
ptp
profile MTG-Slave
!
interface TenGigE1/0/0/1 <<< Puerto que distribuye reloj como Master
description Hacia CN-ABR-MSC Ten0/4/0/7 <<< Wan hacia el Core CN-ABR-MSC
ptp
profile MTG-Master
!
interface TenGigE0/0/0/1 <<< Puerto que distribuye reloj como Master
description Hacia CN-ABR-INQ Ten0/0/0/2 <<< Wan hacia el Core CN-ABR-INQ
ptp
profile MTG-Master
!

```

Configuración del equipo CN-ABR-MSC

```

ptp
clock
domain 0
identity mac-address router
!
profile MTG-Slave <<< Plantilla para reloj que ingresa al MTG
dscp 46
transport ipv4
port state slave-only
master ipv4 200.10.1.9
priority 120
!
master ipv4 200.10.2.9
priority 125
!
profile MTG-Master <<< Plantilla para cliente de reloj del MTG
dscp 46
transport ipv4
announce timeout 2
!

interface TenGigE0/0/0/0 << Interconexión dentro del Core

ptp
profile MTG-Slave
!

interface TenGigE0/4/0/5 <<< Interconexión Agregación CEN-1-QCN

ptp
profile MTG-Master

```

Configuración del equipo PAN, CEN-1-QCN

```

interface Loopback100 <<< Loopback para core-aggr Proceso IS-IS
ip address 10.20.1.100 255.255.255.255

```

```

!
interface Loopback101          <<< Loopback para el proceso IS-IS de la IP-RAN de acceso
 ip address 10.20.1.99 255.255.255.255

ptp clock boundary domain 0
 output lpps R0
 clock-port Core-slave slave
  transport ipv4 unicast interface Lo100 negotiation
  clock source 10.8.0.4
  clock source 10.8.0.2 1
 clock-port Acceso-Master master
  transport ipv4 unicast interface Lo101 negotiation
!
router isis agg-acc
 passive-interface Loopback1
 passive-interface Loopback2

```

Configuración de Interfaces, IS-IS y LDP de CSG-MBC

```

interface Loopback100
 ip address 10.20.25.100 255.255.255.255   <<< Hacia el reloj PAN
!
interface Loopback101
 ip address 10.20.25.99 255.255.255.255   <<< Hacia la red de transporte
!
ptp clock boundary domain 0
 lpps-out 1 4096 ns
 clock-port CSG_Slave_slave
  transport ipv4 unicast interface Lo100 negotiation
  clock source 100.20.1.100   <<< Reloj principal CEN-1-QCN
  clock source 100.20.2.100   <<< Reloj principal CEN-2-MS
 clock-port CSG_Master_master
  transport ipv4 unicast interface Lo101 negotiation
!
router isis agg-acc
 advertise passive-only
 passive-interface Loopback100   <<< Ingreso de las interfaces de sincronia al IS-IS
 passive-interface Loopback101
!
!
interface Vlan400
 description 1588 PTPv2 client
 ip address 10.21.25.100 255.255.255.0   <<< Gateway de la red de los eNBs locales
!
interface GigabitEthernet0/6
 description Hacia los eNB 1588 PTPv2 client
 service instance 400 ethernet
 encapsulation untagged
 bridge-domain 400

```

Configuración de un eNB conectado a CSG-MBC

```

interface Vlan400
 description UMMT3.0 1588v2 client
 ip address 10.21.25.1 255.255.255.0
 ptp announce interval 3
 ptp announce timeout 2
 ptp sync interval -4
 ptp slave unicast negotiation
 ptp clock-source 100.20.25.99
 ptp enable
!
interface GigabitEthernet0/0
 description UMMT3.0 1588v2 client
 switchport access vlan 400
!
network-clock-select hold-timeout infinite
network-clock-select mode nonrevert
network-clock-select 1 PACKET-TIMING
ptp mode ordinary
ptp priority1 128
ptp priority2 128
ptp domain 0
!
ip route 0.0.0.0 0.0.0.0 10.21.25.100

```

3.3.5.1 Preparación de la red para alta disponibilidad

Como se trato en la sección 3.2.2.5, al transportar las etiquetas de MPLS a través de varios segmentos del IS-IS utilizando BGP tiene limitantes en cuanto a la convergencias, con el propósito de mitigar este nuevo problema se han visto algunos mecanismos propuestos en este caso por el fabricante y que hoy en día son estándares. A continuación se presentan extractos de las configuraciones que permiten aplicar estos mecanismos.

3.3.5.2 LFA FRR con BFD

En resumen LFA y remote LFA son mecanismos que permiten pre calcular un camino alternativo y tener esta información lista para ser utilizada en caso de fallas, se suele aplicar estas configuraciones con la funcionalidad de BFD y lograr tiempos de convergencia alrededor de 50ms^[30]. A continuación se presentan ejemplos de configuración de los modelos utilizados en la solución propuesta.

Ejemplo, equipo cell-site ASR-901

```
interface Vlan10          <<<Interface hacia el equipo ASR 903
 ip router isis agg-acc
 mpls ldp igp sync delay 10
 bfd interval 50 min_rx 50 multiplier 3
 no bfd echo
 isis network point-to-point
!
interface Vlan11          <<<Interface hacia otro equipo ASR 901 de cell-site
 ip router isis agg-acc
 mpls ldp igp sync delay 10
 bfd interval 50 min_rx 50 multiplier 3
 isis network point-to-point
!
router isis agg-acc
 fast-reroute per-prefix level-1 all
 fast-reroute remote-lfa level-1 mpls-ldp
 bfd all-interfaces
```

```

!
remotelfa-frr enable
no l3-over-l2 flush buffers
!
mpls ldp discovery targeted-hello accept
!
!
nterface Vlan400
!

```

Ejemplo, equipo PAN ASR-903

```

cef table output-chain build favor memory-utilization
FRR
interface TenGigabitEthernet0/0/0    <<< Conexión Hacial el Core
 ip router isis agg-acc
 mpls ldp igp sync delay 10
 bfd interval 50 min_rx 50 multiplier 3
 no bfd echo
 isis network point-to-point
!
interface TenGigabitEthernet0/1/0    <<< Conexión Hacia el Core
 ip router isis agg-acc
 mpls ldp igp sync delay 10
 bfd interval 50 min_rx 50 multiplier 3
 no bfd echo
 isis circuit-type level-2-only
 isis network point-to-point
!
interface BDI10                      <<< Conexión hacia la capa de acceso y cell-sites como un 901
 ip router isis agg-acc
 mpls ldp igp sync delay 10
 bfd interval 50 min_rx 50 multiplier 3
 no bfd echo
 isis network point-to-point
!
!
router isis agg-acc
 fast-reroute per-prefix level-1 all
 fast-reroute remote-lfa level-1 mpls-ldp
 bfd all-interfaces
!
remotelfa-frr enable
no l3-over-l2 flush buffers
!
mpls ldp discovery targeted-hello accept
!

```

Ejemplo, equipo MTG

```

router isis core
 interface TenGigE0/0/0/0    <<<Interface hacia el Core
  circuit-type level-2-only
  bfd minimum-interval 15
  bfd multiplier 3
  bfd fast-detect ipv4
  point-to-point
  link-down fast-detect
  address-family ipv4 unicast
  fast-reroute per-prefix
  mpls ldp sync
!
!mpls ldp
router-id 10.2.250.100
discovery targeted-hello accept
!
interface TenGigE0/0/0/0
!

```

Ejemplo, equipo Core, CN-ABR-MS

```

explicit-path name TAKETHISPATH
 index 1 next-address strict ipv4 unicast 10.8.0.1    <<< Hacia CN-ABR-INO
 index 2 next-address strict ipv4 unicast 10.8.0.3    <<< Hacia CN-ABR-QCN
!
interface Loopback100
 ipv4 address 100.8.0.2 255.255.255.255
!

```

```

interface tunnel-te13
  ipv4 unnumbered Loopback100
  autoroute announce
  destination 10.8.0.3
  path-option 1 explicit name TAKETHISPATH
  logging events link-status
!
router isis core
  address-family ipv4 unicast
    mpls traffic-eng level-1-2
    mpls traffic-eng router-id Loopback100
!
  interface TenGigE0/0/0/0
    circuit-type level-2-only
    bfd minimum-interval 15
    bfd multiplier 3
    bfd fast-detect ipv4
    point-to-point
    link-down fast-detect
    address-family ipv4 unicast
      fast-reroute per-prefix
      mpls ldp sync
!
  interface TenGigE0/7/0/1
    circuit-type level-2-only
    bfd minimum-interval 15
    bfd multiplier 3
    bfd fast-detect ipv4
    point-to-point
    link-down fast-detect
    address-family ipv4 unicast
      fast-reroute per-prefix
      mpls ldp sync
!
  interface TenGigE0/4/0/5
    circuit-type level-2-only
    bfd minimum-interval 15
    bfd multiplier 3
    bfd fast-detect ipv4
    point-to-point
    link-down fast-detect
    address-family ipv4 unicast
      fast-reroute per-prefix level 2
      fast-reroute per-prefix exclude interface TenGigE0/0/0/0
      fast-reroute per-prefix exclude interface TenGigE0/7/0/1
      fast-reroute per-prefix lfa-candidate interface tunnel-te13
    mpls ldp sync
!
mpls traffic-eng
  interface TenGigE0/0/0/3
    bfd fast-detect
  !
  interface TenGigE0/4/0/5
    bfd fast-detect
  !
!mpls ldp
  router-id 10.8.0.2
  igp sync delay 10
  !
  interface tunnel-te13
  !
  interface TenGigE0/0/0/0
  !
  interface TenGigE0/4/0/5
    igp sync delay 5
  !
  interface TenGigE0/7/0/1

```

Ejemplo, equipo Core, CN-ABR-QCN

```

interface Loopback100
  ipv4 address 10.8.0.3 255.255.255.255
!
router isis core
  address-family ipv4 unicast
    mpls traffic-eng level-1-2
    mpls traffic-eng router-id Loopback100
!
  interface TenGigE0/0/0/0
    circuit-type level-2-only

```

```

bfd minimum-interval 15
bfd multiplier 3
bfd fast-detect ipv4
point-to-point
link-down fast-detect
address-family ipv4 unicast
mpls ldp sync
!
! !
mpls traffic-eng
interface TenGigE0/0/0/0
  bfd fast-detect
!
interface TenGigE0/0/0/5
  bfd fast-detect
!
!mpls ldp
router-id 10.8.0.3
discovery targeted-hello accept
!
interface TenGigE0/0/0/0
!
interface TenGigE0/0/0/5

```

3.3.5.3 BGP PIC o BGP FRR

En las primeras configuraciones de las secciones 3.1.3.2, 3.1.3.3, 3.1.3.4 y 3.1.3.5 ya se incluía esta configuración en el BGP, esta funcionalidad consiste en coordinar que la información de enrutamiento entre el *route-reflector* con los routers que tienen las sesiones iBGP, permitan tener rutas alternas listas en la tabla de enrutamiento, es decir en memoria para ser utilizadas en el caso de un evento en la ruta principal. A continuación, se resaltarán únicamente las líneas de configuración correspondiente a esta funcionalidad.

Ejemplo, equipo MTG

```

router isis cor

router bgp 1000
  address-family ipv4 unicast
    additional-paths receive
    additional-paths selection route-policy add-path-to-ibgp
  !
  address-family vpnv4 unicast
    additional-paths receive
    additional-paths send
    additional-paths selection route-policy add-path-to-ibgp
  !
  route-policy add-path-to-ibgp
    set path-selection backup 1 install
  end-policy
!

```

Ejemplo, equipo Core, CN-ABR-MS

```

router bgp 28006

```

```

address-family ipv4 unicast
  additional-paths receive
  additional-paths send
  additional-paths selection route-policy add-path-to-ibgp
!
route-policy add-path-to-ibgp
  set path-selection backup 1 advertise install
end-policy
!

```

Ejemplo, equipo Core RR, CN-RR-INQ

```

interface Loopback100
  description Global Loopback
  ipv4 address 10.2.255.100 255.255.255.255
router bgp 28006
  nsr
  bgp router-id 10.2.255.100
  address-family ipv4 unicast
    additional-paths receive
    additional-paths send
    additional-paths selection route-policy add-path-to-ibgp
  network 10.2.255.100/32
  allocate-label all
!

```

Ejemplo, equipo PAN CEN-1-QCN

```

cef table output-chain build favor convergence-speed <<< Este comando optimiza al equipo
para BGP-PIC
router bgp 28006
  address-family ipv4
  bgp redistribute-internal
  bgp additional-paths receive

```

Ejemplo, equipo CSG CSG-MBC

```

router bgp 28006
  address-family ipv4
  bgp additional-paths install

```

3.3.6 PREPARACIÓN DE LA RED PARA CALIDAD DE SERVICIO (QOS)

Como conclusión de lo que se trató en la sección 3.2.2.6, se presenta a continuación la clasificación del tráfico y las políticas que se deben configurar en cada punto de acceso del tráfico de acuerdo a los criterios mencionados en la sección 3.2.2.6 y a la referencia [45], adicionalmente para el tráfico de la sincronización se ha marcado la señal proveniente del reloj maestro con un DSCP de 46 y el valor de EF en la red de transporte. Las políticas configuradas en cada interface son de dos tipos, NNI

(Network to Network interface) para los nodos que interconectan puntos de acceso interno y UNI (User to Network interface) para los puntos de acceso con el usuario final, en este caso para la red LTE el usuario final es el eNB.

En la figura 3.19 se ha representado en general un camino extremo a extremo del servicio LTE, desde el eNB hasta el core de servicios y se ha nombrado los puntos de interconexión de acuerdo al detalle descrito en la lista a continuación:

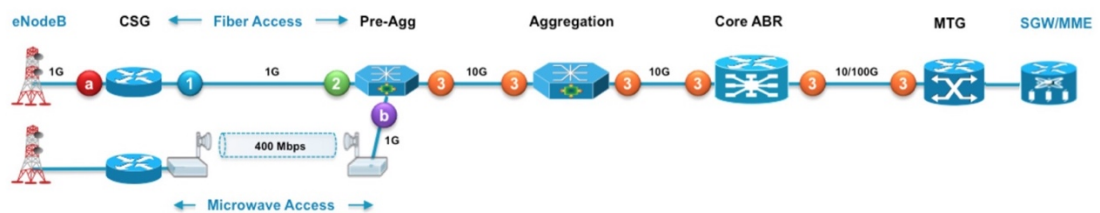


Figura 3.19 Esquema de calidad de servicio QoS ^[45]

- (a) Calidad de servicio en una interface de Usuario, conocido como UNI, esto del tráfico que viene de un eNB hacia un equipo *cell-site*, es decir un CSG.
- (b) Calidad de servicio entre un CSG y un equipo de pre-agregación, lo que se puede considerar como NNI, pero con la mención del caso especial de una última milla de microonda.
- (1) – (2), Calidad de servicio entre un CSG y equipo PAN, esto se considera como NNI común.
- (3) Calidad de servicio entre elementos de red NNI, este enlace está en la capa de Agregación unificando el core con la capa de pre-agregación..

Ejemplo, equipo CSG

Clasificación del tráfico, las interfaces de ingreso son del tipo UNI y su marcación basada en IP DSCP convencionalmente se lo hace en la estación base, por lo cual en el router CSG se tiene lo siguiente para el tráfico de salida (*upstream*) de la interface UNI:

```
!  
  
class-map match-any CMAP-NMgmt-DSCP << Tráfico de Gestión  
  match dscp cs7  
!  
class-map match-all CMAP-RT-DSCP << Tráfico de Voz en tiempo real  
  match dscp ef  
!  
class-map match-any CMAP-HVideo-DSCP << Tráfico de Video  
  match dscp cs3  
!  
La clasificación de la calidad de servicio en el sentido de entrada  
(downstream) de las interfaces UNI y esta basado en lo que se conoce  
como QoS groups:  
!  
class-map match-any CMAP-NMgmt-GRP << Tráfico de Gestión  
  match qos-group 7  
!  
class-map match-any CMAP-CTRL-GRP << Tráfico de enrutamiento, conocido como plano de control  
  match qos-group 6  
!  
class-map match-all CMAP-RT-GRP << Tráfico de Voz en tiempo real  
  match qos-group 5  
!  
class-map match-any CMAP-HVideo-GRP << Tráfico de Video  
  match qos-group 3  
!  
La clasificación de calidad de servicio, para las interfaces NNI en la  
dirección de entrada hacia la red MPLS viene dada por el campo TC  
(Traffic Class), mas conocido como EXP (Experimental bit)  
!  
class-map match-any CMAP-NMgmt-EXP << Tráfico de Gestión  
  match mpls experimental topmost 7  
!  
class-map match-any CMAP-CTRL-EXP << Tráfico de enrutamiento, conocido como plano de control  
  match mpls experimental topmost 6  
!  
class-map match-all CMAP-RT-EXP << Tráfico de Voz en tiempo real  
  match mpls experimental topmost 5  
!  
class-map match-any CMAP-HVideo-EXP << Tráfico de Video  
  match mpls experimental topmost 3
```

!
Luego de la clasificación se tiene entonces las políticas de calidad de servicio que se aplican a las interfaces, en este caso a la interface UNI del eNB, se tiene entonces, una política para el tráfico de entrada, conocido como tráfico de

Upstream y otra política en el tráfico de salida conocido como tráfico de *Downstream* de esta manera se tiene a manera de ejemplo lo siguiente:

```
!
!
interface GigabitEthernet0/2 << QoS en el punto (a). Interface hacia el eNB.
 service-policy output PMAP-eNB-UNI-P-E
 service-policy input PMAP-eNB-UNI-I
!
!
policy-map PMAP-eNB-UNI-I
 class CMAP-RT-DSCP
  police cir 20000000
  set qos-group 5
 class CMAP-NMgmt-DSCP
  police 5000000
  set qos-group 7
 class CMAP-HVideo-DSCP
  police 100000000
  set qos-group 3
 class class-default
  police 200000000
!
!
policy-map PMAP-eNB-UNI-P-E
 class class-default
  shape average 425000000
  service-policy PMAP-eNB-UNI-C-E
!
!
policy-map PMAP-eNB-UNI-C-E
 class CMAP-RT-GRP
  priority percent 5
 class CMAP-NMgmt-GRP
  bandwidth percent 1
 class CMAP-HVideo-GRP
  bandwidth percent 25
 class class-default
!
!
```

La calidad de servicio hacia las interfaces de red conocidas como NNI, en este caso se tiene la red a través de fibra óptica en el acceso, se tiene una correspondencia entre la clasificación del equipo CSG dentro las interfaces NNI gracias al modelo DiffServ, es decir con las marcas EF y AF. Nuevamente se tiene la política del tráfico de *Downstream* aplicado al ingreso de la interface en el punto (1) de la figura 3.19 y la política de *Upstream* en el tráfico de salida de la interface:

```
!
!
interface GigabitEthernet0/10 << Calidad de servicio en el punto (1). Enlace de fibra.
 Access Ring.
 service-policy input PMAP-NNI-I
 service-policy output PMAP-NNI-E
 hold-queue 1500 in
 hold-queue 1500 out
!
!
policy-map PMAP-NNI-I
 class CMAP-RT-EXP
  set qos-group 5
 class CMAP-NMgmt-EXP
  set qos-group 7
```

```

class CMAP-CTRL-EXP
  set qos-group 6
class CMAP-HVideo-EXP
  set qos-group 3
class class-default
!
!
policy-map PMAP-NNI-E
class CMAP-RT-GRP
  priority percent 20
class CMAP-NMgmt-GRP
  bandwidth percent 5
class CMAP-CTRL-GRP
  bandwidth percent 2
class CMAP-HVideo-GRP
  bandwidth percent 50
class class-default

!
!
policy-map PMAP-eNB-UNI-I
class CMAP-RT-DSCP
  police cir 20000000
  set qos-group 5
class CMAP-NMgmt-DSCP
  police 5000000
  set qos-group 7
class CMAP-HVideo-DSCP
  police 100000000
  set qos-group 3
class class-default
  police 200000000
!
!
policy-map PMAP-eNB-UNI-C-E
class CMAP-RT-GRP
  priority percent 5
class CMAP-NMgmt-GRP
  bandwidth percent 1
class CMAP-HVideo-GRP
  bandwidth percent 25
class class-default
!
!
policy-map PMAP-eNB-UNI-P-E
class class-default
  shape average 350000000
  service-policy PMAP-eNB-UNI-C-E
!
!
policy-map PMAP-NNI-I
class CMAP-RT-EXP
  set qos-group 5
class CMAP-NMgmt-EXP
  set qos-group 7
class CMAP-CTRL-EXP
  set qos-group 6
class CMAP-HVideo-EXP
  set qos-group 3
class class-default
!
!
policy-map PMAP-NNI-E
class CMAP-RT-GRP
  priority percent 20
class CMAP-NMgmt-GRP
  bandwidth percent 5
class CMAP-CTRL-GRP
  bandwidth percent 2
class CMAP-HVideo-GRP
  bandwidth percent 50
class class-default
!
!
!

```

Ejemplo, equipo Pre-agregación PAN

La clasificación las interfaces del tipo UNI en el caso de los nodos de pre-agregación PAN, están basados en IP-DSCP, esto para el caso de se tengan eNB conectados directamente en el equipo PAN.

```
!  
class-map match-any CMAP-NMgmt-DSCP << Tráfico de Gestión  
  match dscp cs7  
!  
class-map match-all CMAP-RT-DSCP << Tráfico de Voz en tiempo real  
  match dscp ef  
!  
class-map match-any CMAP-HVideo-DSCP << Tráfico de Video  
  match dscp cs3  
!
```

Para el caso de las interfaces NNI, se tiene la marcación en el campo TC o EXP

```
class-map match-any CMAP-NMgmt-EXP << Tráfico de Gestión  
  match mpls experimental topmost 7  
!  
class-map match-any CMAP-CTRL-EXP << Tráfico de enrutamiento, conocido como plano de control  
  match mpls experimental topmost 6  
!  
class-map match-all CMAP-RT-EXP << Tráfico de Voz en tiempo real  
  match mpls experimental topmost 5  
!  
class-map match-any CMAP-HVideo-EXP << Tráfico de Video  
  match mpls experimental topmost 3  
!
```

Para la política aplicada a la interface del punto (b) del enlace de microonda se aplica una técnica conocida como *traffic-shaping*, que consiste en limitar el ancho de banda de la interface almacenando lo más posible la información sin desecharla, a continuación se presenta ésta política para dos modelos de equipos que se suelen utilizar como pre-agregación

```
!  
ASR 903 Pre-Agregación  
interface GigabitEthernet0/22 << QoS punto (b). Enlace de microonda  
  service-policy output PMAP-uW-NNI-P-E  
  policy-map PMAP-NNI-uw-C-E  
    class CMAP-RT-EXP  
      priority  
    class CMAP-CTRL-EXP  
      bandwidth remaining ratio 1  
    class CMAP-NMgmt-EXP  
      bandwidth remaining ratio 2  
    class CMAP-HVideo-EXP  
      bandwidth remaining ratio 3  
    class class-default  
  !  
  policy-map PMAP-NNI-uw-P-E  
    class class-default  
      shape average 400M  
    service-policy PMAP-NNI-uw-C-E
```

ME3800X Pre-Agregación

```
interface GigabitEthernet0/22 << QoS punto (b). Enlace de microonda
service-policy output PMAP-uW-NNI-P-E
!
policy-map PMAP-NNI-uw-C-E
class CMAP-RT-EXP
priority
  police 80m
class CMAP-NMgmt-EXP
  bandwidth 10000
class CMAP-HVideo-EXP
  bandwidth 80000
class class-default
!
policy-map PMAP-NNI-uw-P-E
class class-default
  shape average 400M
  service-policy PMAP-NNI-uw-C-E
!
!
!
```

Para el enlace de fibra, en el punto (2) de la figura 3.19 se tiene entonces la siguiente política de salida o *Dowstream*

!

ME3800X-based Pre-Agregación Node

```
!
!
interface GigabitEthernet0/1 << QoS punto (2), interface conectada con fibra al nodo
de acceso
service-policy output PMAP-PreAG-NNI-E
end
!
policy-map PMAP-PreAG-NNI-E
class CMAP-RT-EXP
priority
  police cir 200000000
class CMAP-CTRL-EXP
  bandwidth 15000
class CMAP-NMgmt-EXP
  bandwidth 50000
class CMAP-HVideo-EXP
  bandwidth 200000
class class-default
!
!
```

ASR 903 Pre-Agregación

```
!
interface GigabitEthernet0/1 << QoS punto (2), interface conectada con fibra al nodo
de acceso
service-policy output PMAP-NNI-E
end
!
policy-map PMAP-NNI-E
class CMAP-RT-EXP
priority
class CMAP-CTRL-EXP
  bandwidth remaining ratio 1
class CMAP-NMgmt-EXP
  bandwidth remaining ratio 2
class CMAP-HVideo-EXP
  bandwidth remaining ratio 3
class class-default
!
!
```

EL tráfico en la interface de salida dentro de estos equipos PAN hacia los nodos de agregación, es para tráfico de Upstream y constituyen tráfico de transito únicamente

```
interface TenGigabitEthernet0/1 << QoS en el punto (3) de la figura 3.19. Interface
de agregación
service-policy output PMAP-NNI-E
end
policy-map PMAP-NNI-E
class CMAP-RT-EXP
  police cir 1000000000
  priority
class CMAP-CTRL-EXP
  bandwidth 150000
class CMAP-NMgmt-EXP
  bandwidth 500000
class CMAP-HVideo-EXP
  bandwidth 2000000
class class-default
```

Ejemplo, de Acceso, Core y MTG

La clasificación de la calidad de servicio en las interfaces NNIs esta basado en EXP

```
class-map match-any CMAP-RT-EXP
  match mpls experimental topmost 5
end-class-map
!
class-map match-any CMAP-CTRL-EXP
  match mpls experimental topmost 6
end-class-map
!
class-map match-any CMAP-NMgmt-EXP
  match mpls experimental topmost 7
end-class-map
!
class-map match-any CMAP-HVideo-EXP
  match mpls experimental topmost 3
end-class-map
```

La política NNI en estos tres equipos, corresponden al tratamiento en el punto (3)

```
policy-map PMAP-NNI-E
class CMAP-RT-EXP
  priority level 1
  police rate 1000000000 bps
!
!
class CMAP-CTRL-EXP
  bandwidth 150000 kbps
!
class CMAP-NMgmt-EXP
  bandwidth 500000 kbps
!
class CMAP-HVideo-EXP
  bandwidth 2000000 kbps
!
class class-default
!
end-policy-map
```

3.4 ANÁLISIS DEL DISEÑO PRESENTADO

En la figura 3.4 se representa el diseño propuesto y que se ha desarrollado a lo largo de este capítulo, se ha aplicado para su creación los conceptos de la técnica *Unified MPLS*, conocida en general con *Seamless MPLS* y presenta algunas características que se describen a continuación:

Continuidad extremo a extremo del servicio

El diseño actual ya no tiene el problema de discontinuidad que tenía el modelo original, todo se encuentra en un único sistema autónomo, mejorando de esta manera la administración en cuanto al aprovisionamiento ya que los nuevos eNB que se sigan añadiendo no implican ajustes de configuración en toda la red, sino que se agregan similarmente como se añade un nuevo nodo MPLS

Escalabilidad y Capacidad

La red MPLS, en el nuevo diseño se ha segmentado de manera jerárquica y ordenada, de tal manera que a futuro se asegura seguir creciendo en nuevos elementos de red, tanto en el acceso como en la pre-agregación o en el core; en consecuencia el aumento en los eNBs y nuevos servicios, no presentan problemas de capacidad del equipamiento. Este ordenamiento jerárquico tiene a su vez las siguientes ventajas:

Optimización de los equipos de red, ya que la segmentación libera recursos de memoria de los equipos de red, ya que solo deben conocer la información correspondiente a su zona y directamente solo los servicios que son de su correspondencia.

En los servicios se ha diseñado un sistema de filtrado de rutas VPNV4 por comunidades que permite que los equipos de la red optimicen aún mas sus recursos al no aprender por este medio rutas de otros servicios de la red en nodos donde no esté configurado ningún cliente de ese servicio, esto en una red plana MPLS no suele suceder, en una red plana todos los PEs conocen todas las rutas de todos los clientes para poderlos utilizar en cualquier momento.

Facilidad de seguir agregando más eNB y acceso de otros servicios a la red, sin que implique que

La jerarquía de la red permite también añadir nuevos servicios como por ejemplo plataformas de seguridad, tanto centralizados como distribuidos.

Este diseño de red también es compatible con IPV6 y Multicast.

La segmentación también facilita el manejo del análisis de problemas de enrutamiento y la operación de la red se simplifica, optimizando también de esta manera el talento humano requerido.

Se facilita con este diseño de la red, la configuración de calidad de servicio para cumplir con los SLAS que los clientes de la CNT E.P demanden.

El diseño presenta facilidades para implementar calidad de servicio y sincronismo, que son requerimientos fundamentales para el servicio LTE.

Resiliencia

Al segmentar la red, se dividió el camino extremo a extremo que existía en cuanto a la distribución de etiqueta, esta información se la obtiene por medio de las sesiones iBGP + label, por tal motivo se han implementado varias técnicas que permiten mejorar la disponibilidad del servicio y sus tiempos de convergencia. Además a

diferencia del diseño anterior se tienen distribuida la carga en varios equipos y no concentrado en los dos ASBRs de la solución de inter-AS.

El fabricante asegura con los mecanismos implementados y que se trataron en la sección 3.2.2.5 tiempos de convergencia menores a los 50ms, lo que asegura la estabilidad de servicios en tiempo real como la telefonía y la video conferencia.

CAPÍTULO 4

CONCLUSIONES Y RECOMENDACIONES

4.1 CONCLUSIONES

1. El crecimiento constante de la demanda de interconexión de datos, de internet y el desarrollo de equipos móviles como *tablets* y *smartphones*; han impulsado a la vez el crecimiento de una red móvil que satisfaga estos requerimientos; y han conducido al actual desarrollo implementado actualmente en las redes 4G LTE. El acceso de radio LTE requiere de transporte *IP* para acceder a los servicios del core de datos *EPC*; es por eso que esta red denominada IP-RAN adquiere importancia fundamental para la entrega de servicios al cliente final, por lo cual se busca que sea más eficiente, y la solución *Seamless* mejora esa eficiencia. En este trabajo se ha aplicado la solución de *Cisco Systems*, denominada *Unified MPLS*, pero se pueden encontrar soluciones similares en otros proveedores tales como: *Juniper*, *Alcatel-Lucent* y *Huawei*, entre otros.
2. Se utilizó de manera general el concepto de *Seamless MPLS*, pero es importante aclarar que dentro de la IETF (*Internet Engineering Task Force*), *Seamless* es un *draft* que lo elaboraron un conjunto de colaboradores, principalmente de las marcas Cisco y Juniper. En este draft^[22], se incluyen descripciones conceptuales y se ha iniciado un primer estudio; su última actualización fue en Diciembre del 2014, en este documento no están todos los aspectos involucrados en la solución, tal como se aplica en el presente trabajo. Los técnicos, tanto de *Cisco Systems* como de *Juniper* no han

seguido dando nuevos aportes al *draft*, sino mas bien siguieron desarrollando su propio producto, *Juniper* lo continuó llamando *Seamless* y Cisco lo denomina *Unified MPLS*; es sobre este desarrollo que se ha fundamentado el diseño de la solución y es de donde se han obtenido las guías de configuraciones de los equipos en la estrategia de migración.

En resumen para *Cisco Systems*, *Unified MPLS* ^[43] es el conjunto de los siguientes componentes, que se suman a la configuración clásica de una red MPLS.

- a) **Segmentación del IGP/LDP.** La aplicación de esta segmentación divide una red MPLS en zonas, que en un inicio es Plana de extremo a extremo y gracias a esta segmentación cada zona optimiza sus recursos de memoria y procesamiento, debido a que los equipos conocen únicamente rutas y etiquetas de su zona y lo que necesitan conocer de los servicios que utilizan los clientes que tienen conectados.
- b) **RFC 3107 (*bgp+label*).** Es el estándar que distribuye etiquetas de un dominio a otro, ya sea al interior del mismo sistema autónomo como entre dos sistemas diferentes; para el caso de *Unified MPLS* es fundamental su uso ya que sin él no podría existir la segmentación, debido a que no podrían unirse los LSPs de la red MPLS.
- c) **Control con filtros BGP (*BGP filtering*).** Con filtros BGP se puede controlar de una manera óptima todas las rutas de control y de servicio, dentro de la tabla de enrutamiento global de cada segmento, como también las L3VPN que se encuentren en el segmento. De esta manera los equipos de la zona pueden tener configurado MPLS y no

requerir de características muy exigentes en memoria y procesamiento, disminuyendo el costo de inversión en equipos de acceso y al mismo tiempo con ello disminuyendo en lo posible el tener una red en capa 2 con sus respectivos inconvenientes.

- d) Flexibilidad en el acceso.** Esta tecnología busca mejorar las prestaciones del equipamiento para que se adapten a todo tipo de nodo de acceso, pudiéndose tener algunas alternativas, tales como: redes en L2, extensión de red MPLS en el acceso, redistribución de etiquetas de manera estática al interior del IGP y MPLS-TP. Con estas posibilidades de conexión se busca integrar a la red MPLS a los acceso que tradicionalmente maneja un proveedor de servicio y optimizar la red MPLS para adaptarse a las redes de acceso existentes.
- e) LFA y R-LFA.** Debido a que las etiquetas son ahora distribuidas por BGP y que muchos segmentos de acceso, por la topología de transmisión con fibra óptica constituyen anillos de acceso, se requiere de mecanismos que ayuden a una mejor convergencia de la información de las etiquetas al momento de presentarse un evento en la transmisión. Comúnmente en topologías en anillo, ante una falla se tienen retardos en la convergencia y la pérdida de paquetes como consecuencia debido a que el tráfico se sigue enrutando hacia el nodo que se encuentra desconectado del anillo. Actualmente se tiene equipos MPLS que soportan los RFCs 5286 y 7490, los cuales aseguran un camino alternativo listo para ser utilizado y lograr que ese

tiempo de convergencia esté en el orden de los milisegundos disminuyendo la pérdida de información.

- f) **BGP PIC (*Prefix-Independent Protection*)**. Conocido también como BGP FRR (*Fast Reroute*) y algunos proveedores lo implementan en los equipos que tienen la funcionalidad de MPLS, como por ejemplo *Cisco* y *Juniper*. Esta funcionalidad de los equipos de enrutamiento permite una mejor convergencia en BGP en un ambiente en el que se tiene un *route-reflector*, donde por lo general se cuenta con una única ruta y que a pesar de tener caminos redundantes, se tiene que esperar la convergencia de: el IGP, LDP y de BGP para las rutas VPNV4 del servicio L3VPN para que esa redundancia entre en operación, esto toma un tiempo en el orden los segundos, mientras que con la funcionalidad de BGP PIC esto requiere de un tiempo de alrededor de 50 ms. En la topología *Seamless* se vuelve más útil esta funcionalidad ya que al separar al IGP en zonas las etiquetas son transportadas por BGP, lo que en un caso de falla agregaría más tiempo de convergencia.

3. La aplicación del *Unified* MPLS, por su características tiene los siguientes beneficios:

- a. **Escalabilidad**, gracias a la segmentación IGP/LDP se puede seguir añadiendo nodo de agregación y acceso. En la referencia [22] se tiene una demostración de la escalabilidad del *Seamless* MPLS
- b. **Orden Jerárquico**, la red *Seamless* es una red ordenada, en la cual se maneja de esa manera el tráfico y permite un mejor aprovechamiento de los recursos.

- c. **Seguridad**, al tener puntos de control de la distribución de prefijos de enrutamiento, se evita la inundación de información innecesaria que puede afectar a los equipos en su procesamiento y memoria.
 - d. **Multiservicios a lo largo de toda la red**, gracias a la flexibilidad que presenta la solución *Seamless* se puede conseguir distribuir los servicios, ya sea lo más cercano al core, agregación o acceso y a los clientes; dependiendo de su jerarquía puede estar ubicado en cualquier punto de esta clasificación.
 - e. **Simplicidad**, Una vez implementada la solución el ir agregando clientes y servicios se vuelve una tarea similar a la que se lo hacía con una red MPLS común, pero además el poder tener en el acceso MPLS vuelve a la operación más simple.
 - f. **Disponibilidad**, con los mecanismos como LFA y RemoteLFA y otras configuraciones el *Seamless* asegura también alta disponibilidad de la red, lo que redunda en el cumplimiento de los SLAs con los clientes
4. Como resultado de la aplicación de *Seamless* MPLS se tiene una red ordenada con una jerarquía que hace posible separar a la red en zonas de acceso consiguiendo una optimización de equipamiento y de transporte; además de que en esta nueva topología se hace posible distribuir nuevos servicios al tener el dominio MPLS más cerca del usuario final.
5. Utilizar similares tecnologías y homologar el equipamiento utilizado en redes MPLS es muy importante para simplificar la operación y el mantenimiento de la red, lo que facilita una disminución del costo de operación. Esto siempre está incluido en las propuestas de los proveedores de equipamiento de telecomunicaciones. En este caso el *Seamless* MPLS aplicado simplifica la

operación y mantenimiento de la red IP-RAN ya que no se trata de una red diferente al backbone MPLS, simplemente un acceso más.

6. En la proyección de LTE está el servicio de IMS para el tráfico de voz, el cual puede ser centralizado o distribuido, la flexibilidad del *Seamless* MPLS permitirá estar listos para cualquiera de las dos posibilidades. Aunque en la actualidad se ha tomado como una ventaja el tener el servicio de voz en 3G y únicamente datos en LTE.
7. La solución propuesta, al estar basada en la topología *Seamless* o *Unified* MPLS de Cisco es compatible con ipv6, en cuanto al servicio. Para esto se debe indicar que en una red MPLS el transporte de ipv6 es un servicio más; en la referencia [30], se especifica que este servicio se lo puede configurar utilizando el estándar RFC 4659 denominado 6VPE, bajo el cual un backbone MPLS IPv4 puede transportar IPv6. Los eNB y EPC entonces pueden ser IPv6 o pueden trabajar en *dual stack*, es decir en IPv4 y en IPV6 al mismo tiempo.
8. Por facilidades en la transmisión, la actual red MPLS de la ciudad de Quito, se diseñó con los PEs de acceso directamente conectados al CORE; en ese entonces no existían aplicaciones que tornen necesario una capa de agregación antes de llegar al CORE. En la aplicación del modelo *Unified MPLS* se presenta la necesidad de crear esta capa de agregación.
9. Debido a la optimización de recursos, tanto por la segmentación del IGP, como por el control de los prefijos VPNV4 distribuidos hacia la IP-RAN, se permite tener en esta zona equipos que tengan las funcionalidades de MPLS y no demanden recursos de memoria y procesamiento muy elevados, lo cual implica tener un equipo no muy costoso y funcional a la vez en el acceso.

10. En la segmentación que se consiguió para la IP-RAN de la ciudad de Quito con la aplicación de *Seamless* o *Unified* MPLS, se pueden identificar tres zonas principales, dentro de las cuales se simplifica el control de los enlaces de transmisión, se pueden entonces medir parámetros como: ancho de banda, jitter, delay, entre otros; y, de esta manera se puede ir balanceando carga para mejorar el aprovechamiento del transporte. Este balanceo de tráfico actualmente se lo hace en forma manual, pero en el mercado ya existen técnicas para automatizar esta tarea y el haber ya organizado jerárquicamente la red con *Unified* MPLS facilita la red para implementar nuevas tecnologías como por ejemplo SDN y optimizar la red de forma automática.
11. Para la implementación de *Unified* MPLS en una red existente, es imprescindible el conocimiento de la red, en cuanto a: topología, servicios y su transmisión, de tal manera que se pueda conseguir una segmentación equilibrada. Siempre para la segmentación se debe conocer si existen los recursos de transmisión y espacio físico en los nodos que se requiera añadir nuevo equipamiento.
12. Es necesario verificar que el equipamiento disponible para el *Unified* MPLS cumpla con los estándares y funcionalidades requeridos. En caso de cumplimiento parcial, se puede considerar un reemplazo o tener en cuenta que en ese segmento de red esa funcionalidad no estará disponible.
13. En la sección 1.3.5.3 se tienen varios mecanismos de transporte con *Unified* MPLS, la elección del modelo a implementarse depende mucho de la proyección de crecimiento, ya que cada modelo tiene su grado de dificultad.

4.2 RECOMENDACIONES

1. Esta tesis permite un estudio posterior sobre el la expansión de la red a nivel Nacional, debido a que el modelo que se ha aplicado permite añadir más ciudades dentro del core Nacional, de igual manera se debe crear una red de agregación que permita añadir las demás ciudades y sus respectivas redes de acceso.
2. Para agregar servicios fijos en un nodo de acceso de la topología *Seamless*, se recomienda que los equipos dentro de la zona de agregación soporten LDP DoD (RFC 7032), para pasar servicios fijos además de los móviles, optimizando los recursos de la red.
3. El modelo de transporte de la solución implementada puede ser modificado para adaptarse no solamente a soluciones móviles como en el presente caso, sino que se puede tener en los nodos de acceso servicios fijos masivos. Para esto se debe realizar un estudio de capacidades requeridas, tanto en equipamiento como en ancho de banda.
4. Para realizar implementaciones de *Seamless* y adaptarse a la operación de la empresa proveedora de servicios de telecomunicaciones, se debe elegir el modelo adecuado al tamaño de la red y su proyección
5. El diseño encontrado de *Seamless* MPLS no se limita únicamente a redes de acceso del tipo móvil, se puede aplicar a otros servicios como *Cloud computing* centralizado o servicios de internet masivos y corporativos, para lo cual se puede utilizar BNG (*Broadband Network Gateway*) para la agregación de clientes de FTTH o GPON.

6. La migración de servicios de una red plana MPLS a un modelo *Seamless* MPLS puede ser gradual o implementarse únicamente donde es necesario hacerlo, es decir no es mandatorio cambiar toda la red, por ello se recomienda realizar un estudio previo.
7. En el diseño realizado de la solución puede añadirse configuraciones de seguridad en la red en la L3VPN de servicio y añadir por ejemplo encriptación. Esto se sugiere para aplicaciones de datos que requieran una atención especial a la seguridad.
8. La topología diseñada para esta solución permite agregar configuraciones adicionales y trabajar con *multicast* dentro del servicio LTE para entregar servicios de video bajo demanda y otras aplicaciones similares.
9. Se puede también trabajar en un esquema de simulación de la solución implementada, para ello se puede plantear dos etapas de prueba, una confirmando el LSP extremo a extremo y otra mediante la configuración de L3VPN de servicios.
10. La estrategia de migración es también un conjunto de procedimientos que permitirán nuevas implementaciones a nivel nacional.
11. En la implementación de estas tecnologías, se hace necesario siempre un trabajo en conjunto entre el personal de ingeniería de la empresa de telecomunicaciones y del proveedor para afinar las configuraciones necesarias.

REFERENCIAS BIBLIOGRÁFICAS

- [1] Cisco Visual Networking Index: Global Mobile Data Traffic Forecast Update, 2013–2018, White paper from Cisco Systems.
- [2] Cisco Visual Networking Index: The Zettabyte Era: Trends and Analysis, 2013-2018, white paper from Cisco Systems.
- [3] Allegiance for Mobil *Backhaul* Unification, Cisco Systems and NEC. 2012.
- [4] IP RAN, Sanog 15 presentation 2006, Kasu Venkat Reddy, Solution Architect & Santanu Dasgupta, Consulting Systems Engineer.
- [5] Architectural Considerations for *Backhaul* of 2G/3G and Long Term Evolution Networks, Cisco Systems, 2010 white paper
- [6] Afif Osseiran, Jose F. Monserrat, Werner Mohr, “MOBILE AND WIRELESS COMMUNICATIONS FOR IMT-ADVANCED AND BEYOND”, Editorial WILEY 2011 John Wiley & Sons Ltd.
- [7] Christopher Cox, “AN INTRODUCTION TO LTE,L, LTE-ADVANCED, SAE AND 4G MOBILE COMMUNICATIONS”, Editorial WILEY 2012
- [8] S. Akhtar, "2G-5G Networks: Evolution of Technologies, Standars, and Deployment", College of Information Technology, UAE University, 2008.
- [9] 3GPP TS 27.001 V12.0.0 (2014-10), “General on Terminal Adaptation Functions (TAF) for Mobile Stations (MS) (Release 12)”
- [10] 3GPP TS 36.401 V12.2.0 “Evolved Universal Terrestrial Radio Access Network (E-UTRAN); Architecture description (Release 12)”
- [11] TS 23.002 V13.1.0 “Network architecture (Release 13)”
- [12] Esa Metsälä y Juha Salmelin, “MOBILE *BACKHAUL*”, Editorial WILEY 2012
- [13] Jariz Aizprúa B. “LLD- Low Level Design CNT LTE – IP RAN”, 2013. **Documento confidencial**

- [14] Ina Minei & Julian Lucek, "MPLS-Enabled Applications", Editorial WILEY , third edition 2011.
- [15] Cisco Systems, white paper "Benefits to Using Layer 3 Access for IP Radio Access Networks".
- [16] Dave Wisely, Philip Eardley y Louise Burness, "IP for 3G-Networking Technologies for Mobile Communications", Editorial WILEY 2002.
- [17] 3GPP TS 23.203 V10.10.0 "Technical Specification Group Services and System Aspects; Policy and charging control architecture"
- [18] Santanu Dasgupta, Sanog (*South Asian Network Operators' Group*), 2008, "Introduction to MPLS Technologies"
- [19] Mukhtiar A. Shaikh, Yousuf Hasan, Mossadaq Turabi, "MPLS Tutorial", SANOG VIII, 2006
- [20] Luc De Ghein, "MPLS Fundamentals", Cisco Press 2007
- [21] Cisco Systems "Implementing Cisco Service Provider Next-Generation Edge Network Services" Student Guide v1.2 2014
- [22] Internet-Draft, IETF, "draft-ietf-mpls-seamless-mpls-07", Diciembre 30, 2014
- [23] Huawei, "AAA System Overview", 2013
- [24] Cisco Systems, BRKSPG-2022 "Evolved Packet Core for LTE Networks", Febrero 2, 2012
- [25] Internet RFC 4364, "BGP/MPLS IP Virtual Private Networks (VPNs)", Febrero 2006
- [26] Vinod Joseph, Srinivas Mulugu, "Network Convergence", Editorial Elsevier 2014.
- [27] Cisco Systems, BRKSPM-2008 "Unified MPLS Design and Deployment Case Study for Mobile Service Provider", 2014
- [28] Juniper Networks, "Building Multi-Generation Scalable Networks with End-to-End MPLS", 2012
- [29] Juniper Networks, "Universal Access and Agregation Mobile *Backhaul* Desing Guide", 2013.
- [30] Cisco Systems, "UMMT 3.0 Design Guide Technical Paper", Julio 12 del 2012.

- [31] Juniper Networks, "Scaling MPLS - Seamlessly", Septiembre 2012.
- [32] CNT E.P, "PROYECCIÓN INTERNET FIJO V6 FINAL" Junio 2014, **Documento confidencial.**
- [33] HUAWEI TECHNOLOGIES CO., LTD., "ATN 905&910&910I&910B&950B Multi-Service Access Equipment, V200R005C00 Hardware Description", Marzo 2015.
- [34] HUAWEI TECHNOLOGIES CO., LTD., "ATN 905&910&910I&910B&950B Multi-Service Access Equipment, V200R005C00 Hardware Description", Marzo 2015.
- [35] HUAWEI TECHNOLOGIES CO., LTD., "HUAWEI CX600 Metro Services Platform, V600R003C00, Product Description", Junio 2012
- [36] Cisco Systems, "Cisco ASR 900 Route *Switch* Processor", Julio del 2015
- [37] Cisco Systems, "End-of-Sale and End-of-Life Announcement for the Cisco ASR 903 Router - IOS XE Release 3.9S"
- [38] Carlos E. Gallegos "LLD- Low Level Design CNT LTE – IP RAN", 2015, Actualización.
Documento confidencial
- [39] CNT E.P, Ingeniería de Trasmisiones "Diseño de Red LTE" **Documento confidencial**
- [40] Juniper Networks, "*Seamless* MPLS Architecture", March 2015.
- [41] Cisco Systems, "Fixed Model Convergence 2.0", Septiembre del 2013.
- [42] CNT E.P, "AMPLIACIÓN DE LA COBERTURA DEL BACKBONE NACIONAL IP/MPLS E INTERNET DE LA CORPORACION NACIONAL DE TELECOMUNICACIONES CNT E.P." Diciembre del 2011, Documento Confidencial
- [43] Cisco Systems, "Mobile Transport Evolution with Unified MPLS " BRKSPG-3684, 2013
- [44] Cisco Systems, "Implementing Cisco Quality of Service" Student Guide, 2006.
- [45] Cisco Systems, "UMMT 3.0 Implementation Guide - Large Network, Multi-Area IGP Design with IP/MPLS Access", Sep 11. 2012 **CISCO CONFIDENTIAL**