

IP sobre ATM - clave en la convergencia de las comunicaciones

Homero Ortega Boada, Carolina Villabona R., Wilder E. Castellanos H.

Grupo I+D en Teleinformática – GTI

Industrial de Santander

Universidad

Abstract—Este artículo busca demostrar la necesidad de una solución combinada de IP sobre ATM para las redes de próxima generación, es decir para las redes que responden al reto de la convergencia. El artículo presenta un método original para llegar a comprender el verdadero sentido de MPLS, la razón de ser del IPoATM y su necesidad en las redes de próxima generación. También se esclarecen nuevos protocolos que permiten la Calidad de Servicio (QoS) como ARP, NHRP. Se espera que el lector tenga conocimientos de IP y ATM. Las figuras y tablas a que se hace referencia en el artículo están incluidas en la presentación anexa en forma de animaciones, en las memorias “II Congreso de Electrónica y Tecnologías Avanzadas”, Universidad de Pamplona, Marzo del 2002.

Palabras claves: ATM, Calidad del servicio (QoS).

I. INTRODUCCIÓN

ESTAMOS marcando la pauta de una nueva era en la tecnología de comunicaciones. La desregulación de las telecomunicaciones y la explosión del tráfico de Internet han forzado a los operadores a revisar los esquemas tradicionales para proveer servicios de voz y datos. Ahora los usuarios tienen una gran necesidad para acceder a la información de una manera eficiente y rápida; es por esto que las comunicaciones de voz, datos, inalámbricas y de cable se están integrando para facilitar más las comunicaciones y los negocios.

El ambiente actual de las telecomunicaciones está conformado por una amplia variedad de redes, por ejemplo PSTN (Public Switched Telephony Network), Internet. El crecimiento de Internet ha estimulado mucho el crecimiento de la transmisión de datos sobre PSTN, creando la necesidad de reacomodar

esta red de telefonía tradicional y además, en muchos casos, ha impulsado el surgimiento de redes de datos paralelas para prestar cada uno de los nuevos servicios requeridos por los usuarios (televisión por cable, internet, servicios móviles, etc.). La mayoría de estas redes son altamente especializadas y diseñadas para proveer un servicio específico. Podemos referirnos a éstas como “redes integradas verticalmente”. Esta situación de integración vertical es el resultado de un proceso de evolución a través de la historia que dificulta la creación de mecanismos que reduzcan costos de operación, portabilidad del servicio, etc.

Todos somos testigos del gran impacto que ha provocado el Internet en la sociedad actual. La era de la información es la era de Internet. El mundo está cambiando a medida que cambia la forma en que nos comunicamos. Internet puede verse como una red de redes que permite la intercomunicación entre computadoras en lugares remotos, conectados a diferentes tecnologías (como X.25, Frame Relay, Ethernet, Token Ring PSTN, ISDN, etc.), ver Fig. 1. Hoy es normal no solo buscar información en la Internet, sino también establecer llamadas telefónicas, comunicaciones de videoconferencias y multimedia, es en este momento cuando comienza a hablarse de calidad de servicio (QoS) de Internet.

Así que cuál es la solución para esta situación de múltiples redes? La solución es la Convergencia de las distintas redes en una nueva arquitectura de red, conocida como Red de Próxima Generación (NGN Next Generation Network), en donde las líneas divisoras entre la transmisión de voz, datos y multimedia desaparecerán..

Algunos requisitos de las redes de próxima generación (NGN) serán:

- Las redes tendrán que ser capaces de manejar una variedad de tráfico, desde transferencia de archivos sencillos hasta servicios de multimedia.
- Soportar acceso por cable e inalámbrico con mucho más ancho de banda del disponible actualmente.
- Proporcionar altos niveles de QoS (generalmente llamado “carrier-class”) donde se requiera.
- Entregar voz, datos y los servicios de multimedia en tiempo real a gran número de usuarios residenciales y corporativos.

II. LA CALIDAD DE SERVICIO (QoS)

Este artículo ha sido recibido el 4 de Marzo del 2002. Es el resultado de investigaciones realizadas en el grupo de I+D de Teleinformática-GTI de la Universidad Industrial de Santander.

Homero Ortega B. Ph.D. of Sciences. Profesor investigador, miembro de GTI y del capítulo de comunicaciones de IEEE. Escuela de Ingeniería Eléctrica Electrónica y de Telecomunicaciones. Universidad Industrial de Santander. También trabajó como investigador en la empresa ERICSSON en temas relacionados.

A.A. 678 Bucaramanga, Colombia, hortegab@uis.edu.co

Carolina Villabona R., Wilder Eduardo Castellanos H. son estudiantes en tesis de la Escuela de Ingeniería Eléctrica Electrónica y de Telecomunicaciones de la Universidad Industrial de Santander (UIS), actualmente realizan investigaciones en el ámbito de NGN mediante un proyecto entre la empresa Telebucaramanga y la UIS para el desarrollo servicios, e-mail: carolina_villabona@ieee.org y wher@ieee.org respectivamente.

Típicamente Internet presta servicios de comunicación basado en el mejor esfuerzo, pues fue ideada para la transmisión de datos por ráfagas. Sin embargo, se ha venido adaptando a una amplia variedad de servicios, cada uno de ellos con requerimientos y parámetros diferentes. Por ejemplo con la posibilidad de cursar servicios de tiempo real como la comunicación de voz, video sobre la Internet surgen interrogantes de calidad como retardos, sincronismo, acceso en el momento. Precisamente el problema de QoS se acentúa. La QoS, desde la perspectiva del usuario, se percibe como la diferencia entre lo que él solicitó del servicio y lo que verdaderamente está recibiendo.

En conclusión, estamos observando la tendencia a crear una serie de servicios sobre redes que no fueron diseñadas para ellos: Las redes de conmutación de circuitos, las cuales garantizan una buena QoS para la voz pretenden ser usadas para manejar comunicaciones de datos con ancho de banda por demanda. Las soluciones son módems conectados a las líneas telefónicas que jamás sobrepasan los 64 kbps o las líneas ISDN (Integrated Services Digital Network) que resultan exageradamente costosas. Algo similar ocurre al pretender usar redes de datos, el Internet, para ofrecer comunicaciones en tiempo real. El problema consiste en cómo combinar las ventajas de el uso óptimo de los recursos de red que ofrecen la conmutación de paquetes con la QoS para la voz que ofrece la conmutación de circuitos.

III. ATM

A mediados de los años 80 se concibe ATM (Asynchronous Transfer Mode) como la tecnología que garantiza la QoS para todo tipo de comunicación. En 1984 la CCITT (ahora UIT) ya había adoptado recomendaciones para ISDN. Sin embargo, esta solución se centra exclusivamente en la interfaz entre el usuario y la red, razón por la cual los recursos de red seguían siendo usados en forma no óptima. ATM viene a ser un desarrollo posterior también conocido como ISDN de banda ancha (BISDN) con un enfoque radicalmente diferente, con las siguientes principales particularidades:

- Se trata de una solución tanto a nivel de red como a nivel de acceso (nivel de usuario) con la posibilidad de transmitir flujos de datos constantes y a ráfagas.
- Al contrario de lo que ocurre con las redes de medio compartido, ATM no está limitada por la velocidad o la distancia.
- Sus estrategias para garantizar QoS para voz, datos y video la hace atractiva para aplicaciones comerciales.
- Sus grandes ventajas se originan con el uso de pequeñas celdas de igual tamaño en lugar de paquetes de datos de longitud variable.

ATM tiene la cualidad de poder establecer *Contratos de Tráfico* entre la red y el usuario para garantizar QoS en una comunicación. La función que se utiliza para esto se conoce como *Control de admisión para una conexión* (CAC) y es la que determina si una nueva petición de conexión debe aceptarse o rechazarse. En caso de aceptarse se realiza el

contrato que incluye los requisitos de QoS, descriptores de tráfico y la categoría de servicio ATM:

A. *Parámetros de QoS de red usados en ATM.* Determinan el compromiso de la red frente a una conexión determinada:

- *CER (Cell Error Ratio).* Proporción de celdas con uno o más bits erróneos, pero excluyendo los bloques de celdas severamente erróneos.
- *CMR (Cell Minsinsertion Ratio).* Rata de celdas/Segundo que se envían hacia una conexión de destino equivocada. También se excluyen los bloques de celdas severamente erróneos.
- *SECBR (Severely-Errored Cell Block Ratio).* Se entiende por bloque de celdas severamente erróneo cuando se presentan M celdas perdidas por error o se encuentran erróneamente en un bloque de N celdas recibidas. M y N son definidos por el operador de red. SECBR es la proporción de tales bloques en el total de bloques transmitidos en una conexión.

B. *Parámetros de QoS negociables entre el usuario y la red durante el establecimiento de la conexión:*

- *CLR (Cell Loss ratio).* Rata de celdas perdidas con respecto al número total de celdas transmitidas. Está en el rango 10^{-1} a 10^{-15}
- *CTD (Cell Transfer Delay).* Es el retardo total de la celda desde que sale del origen hasta que llega al destino. Debido a que generalmente celdas de una misma conexión experimentan retardos diferentes, el CTD se define mediante una función de densidad de probabilidad. Los estándares permiten la negociación del máximo CTD.
- *CDV (Cell Delay Variation).* Mide la variación del retardo total de las celdas de una conexión.

C. *Descriptores de tráfico del origen.* La red debe estar en capacidad de determinar si el usuario está cumpliendo con el contrato establecido para determinada conexión. Para este fin se usan los siguientes parámetros:

- *PCR (Peak Cell Rate).* Velocidad de celdas/segundo que el origen nunca podrá exceder. De aquí se determina el intervalo mínimo de celdas como $T=1/PCR$
- *SCR (Sustainable Cell Rate).* Velocidad media de celdas/segundo, a la que se compromete transmitir el origen durante un periodo de tiempo.
- *MBS (Maximum Bursts Size).* Corresponde a la máxima cantidad de celdas consecutivas durante la velocidad máxima – PCR.
- *MCR (Minimum Cell Rate).* Velocidad mínima en celdas/Segundo a que se compromete transmitir el origen.
- *CDVT (Cell Delay Variation Tolerance).* Especifica el nivel de variación del retardo de celdas que debe ser tolerado por la conexión.

D. *Categorías de servicio ATM.* En 1996 el foro ATM [2] definió 5 categorías de servicios a saber:

- *CBR (Constant Bit Rate).* Está designada principalmente a emular la conmutación de circuitos, permitiendo una buena calidad a las comunicaciones de voz y ciertos tipos de video, que requieren una velocidad constante durante toda la transmisión.

- *rt-VBR (Real-Time Variable Bit Rate)*. Designada a comunicaciones de tiempo real, pero con tasa de bits variable. Este es el caso de ciertos tipos de video.
- *nrt-VBR (non-real time Variable-Bit-Rate)*. Designada a aplicaciones que producen tráfico en forma de ráfagas. Es el caso de transferencias bancarias.
- *ABR (Available Bit Rate)*. Designada para fuentes de datos capaces de adaptar dinámicamente su velocidad según la disponibilidad que ofrezca la red. De esta manera se explota al máximo el ancho de banda que pueda aparecer disponible en la red en un momento dado.
- *UBR (Unspecified Bit Rate)*. En esta categoría no se proporcionan ninguna garantía de QoS. Es apropiado para aplicaciones que pueden tolerar o adaptarse fácilmente a la pérdida de celdas.

Todos los parámetros mencionados anteriormente resultan enlazados como se muestra en la tabla I.

E. Control de la congestión en ATM

ATM permite la administración del ancho de banda total disponible de manera que se garantizan recursos para las comunicaciones de tiempo real y se evita la posible congestión que puedan causar las comunicaciones por ráfagas. El problema principal radica en el control de las ráfagas de manera que no vayan a causar congestión en los nodos (cuando las colas de celdas entrantes superen la memoria de almacenamiento disponible). El procedimiento que siguen los nodos involucrado es como se muestra en la Fig. 2:

Primero que todo los nodos a lado y lado de la conexión proceden a reservar recursos de ancho de banda correspondientes a la CBR, luego se reservan recursos para las fuentes con VBR. El ancho de banda restante (ABR) es el que se usará para las comunicaciones de datos por ráfagas.

Para el control de congestión de ABR, los nodos que toman parte de la conexión deben intercambiar información sobre su estado de congestión a través de las celdas de gestión de recursos (RM-Resource management). Se tienen 3 métodos principales, los cuales se ilustran en la Fig. 3:

- *Método de realimentación binaria*. También conocido como método de realimentación positiva. Si una celda pasa por un nodo que presenta congestión, la celda es marcada mediante un bit llamado EFCI=1 (Explicit Forward Congestion Indication). El nodo final detectará este bit y pondrá un bit CI=1 (Congestion Indication) a las celdas RM que transmite hacia el origen. El origen procederá a disminuir, en forma exponencial, la tasa de bits transmitida. De otra suerte, el origen aumentará esa tasa de bits en forma lineal. Este método tiene la desventaja de presentar variaciones abruptas en las tasas de bits transmitidos.
- *Método de realimentación explícita de la velocidad*. Para cubrir la desventaja del método anterior, cada nodo en el camino va calculando el ancho de banda que puede ofrecer en el momento para la conexión (se usa generalmente el método de idoneidad max-min). Este ancho de banda se compara con el valor recibido del nodo anterior, en caso de ser menor, la celda RM es modificada con el nuevo valor y

se transmite al siguiente nodo. De esta manera el destino tiene la información sobre el ancho de banda que ofrece el nodo de mayor congestión y la envía al origen a través de una celda RM. Con esta información, el origen procede a adaptar su velocidad de transmisión.

- *EPRCA (Enhanced Proportional Rate Control Algorithm)*. Combina los dos métodos anteriores en redes compuestas de nodos que implementan uno u otro método.

IV. IP SOBRE ATM

Conociendo las enormes ventajas que ofrece IP para la interconexión, así como las ventajas de ATM para la prestación de servicios con QoS garantizada, podríamos considerar que una red óptima debería ser la combinación de estas dos tecnologías. Para llegar a comprender los últimos avances en este sentido es necesario revisar varios modelos que se vienen usando hasta el momento.

A. *El modelo LANE (LAN Emulation)*. LANE es una especificación del Foro ATM con el fin de acelerar la implementación de ATM en las redes empresariales. Estas redes generalmente implementan IP sobre una red LAN estandarizada (como Ethernet o Token Ring). LANE es en términos sencillos una capa que se monta sobre ATM con el fin de ofrecer a IP una interfaz idéntica a la de las LAN legales. De esta manera, cualquier software que se ejecuta sobre una red LAN legal, lo hará sobre la capa LANE sin necesidad de modificación alguna. Como resultado tenemos una red LAN emulada - ELAN (Emulated LAN) la cual está compuesta de clientes LEC (LAN Emulation Clients), un servidor de emulación - LES (LAN Emulation Server), un servidor de difusión y desconocidos (BUS) y un servidor de configuración LECS (LAN Emulation Configuration Server), ver la Fig.4. Puede decirse que LANE es un solucionador entre direcciones MAC y direcciones ATM. Cada elemento de la red tiene una dirección ATM única. Cuando un LEC de origen necesita comunicarse con una dirección MAC perteneciente a un LEC de destino, el LEC de origen realiza una solicitud de resolución de direcciones al servidor LES, el cual responde con la dirección ATM del destino, finalmente se establece la VCC (Virtual Channel Connection) entre origen y destino. Mientras transcurre todo este proceso, el tráfico de datos en cuestión es atendido por BUS. LECS permite la asignación de nuevos LECs a la ELAN y su asociación con el LES. LANE resulta una solución óptima para conexiones UBR y ABR, sin embargo, al mantener ocultos los detalles de la red ATM, impide que los atributos de QoS de ATM estén disponibles a los protocolos de la capa de red.

B. *El modelo IP sobre ATM clásico (CLIP)*. Es una especificación de la IETF (RFC 2255), según la cual IP trata a ATM como otra subred en la que se tienen computadoras (host), así como enrutadores. Múltiples subredes IP, en adelante LIS (Logical IP Subnetwork), componen la red ATM. Una LIS se caracteriza por que todos sus host usan el mismo prefijo de red IP (los mismos números de red y los mismos números de subred). Se tiene un servidor ATM ARP

(ATM Resolution Protocol) para resolver direcciones IP a direcciones ATM dentro de la LIS, ver Fig. 5. En caso de una nueva conexión el host de origen conoce la dirección IP del destino, solicita al servidor ATM ARP resolver esta dirección a la correspondiente dirección ATM para proceder finalmente a establecer una VCC directa a través de ATM.

El protocolo ARP surge como un sistema de diálogo entre clientes y servidor ATM ARP. En el caso de una comunicación entre un host de origen y un host de destino que se encuentran en diferentes LIS, el ATM ARP del LIS de origen ordenará una VCC hacia un enrutador. Este último puede desconocer el destino y por lo tanto ordenará una VCC hacia otro enrutador y así sucesivamente hasta llegar a la LIS de destino. El enrutador de destino, una vez conoce la dirección ATM (gracias al servidor ATM ARP en el LIS de destino), ordenará una VCC con el host de destino.

La principal desventaja de este método consiste en que una VCC entre usuarios de diferentes LIS debe pasar obligatoriamente por cada enrutador, en el cual se hace un nuevo análisis de la dirección IP. Esto limita las posibilidades de QoS ofrecidas por ATM.

A pesar de esta desventaja, este modelo nos permite ya observar una cooperación entre la capa IP y ATM. En cierta manera este modelo nos recuerda el establecimiento de una llamada de voz en una red PSTN. El número IP es comparable con el número telefónico marcado. Este número solo se utiliza para que el procesador pueda reservar recursos de conmutación en la red, es decir una ruta a través de muchos nodos (comparable a un VPI (Virtual Path Identifier)) y un Time-Slot entre cada nodo (comparable a un VCI (Virtual Channel Identifier)). En otras palabras este modelo de IPoATM bien podría reemplazar una red telefónica PSTN en la cual la dirección IP sirve para identificar al usuario final, mientras que ATM permite la conmutación en sí.

C. El modelo NHRP (Next-Hop Resolution Protocol). En este modelo los host se conocen como NHC (Next Hop Client), también se tienen servidores NHS (Next-Hop server) montados sobre los mismos enrutadores usados en el modelo CLIP, véase la Fig. 6. NHRP busca cubrir la desventaja del modelo CLIP, permitiendo encontrar la dirección ATM del NHC-destino, de manera que se pueda establecer una VCC directa con el NHC-origen. Durante el establecimiento de una comunicación, el NHC-origen comienza solicitando a su NHS, mediante **el protocolo NHRP**, la resolución de la dirección IP del NHC-destino a una dirección ATM. El NHS, en caso de no tener a cargo el LIS con el NHC-destino, reiniciará el paquete de solicitud de resolución NHRP al siguiente NHS a lo largo del camino de destino. El NHS final resuelve la dirección ATM del NHC-destino y la devuelve por el mismo camino mediante el protocolo NHRP. Finalmente el NHC-origen podrá establecer una VCC directa con el NHC-destino.

D. El modelo MPoA (Multiprotocol over ATM). Este modelo busca usar las ventajas de NHRP para la comunicación entre LISs (capa3 modelo OSI) y las de LANE para la comunicación dentro de un mismo LIS (capa 2 del modelo OSI). Los elementos de una red MPoA son los siguientes:

los clientes de la red se conocen como MPC (MPoA Client); los enrutadores tienen funciones de servidores y conocen como MPS (MPoA Server). Estos atienden las solicitudes de resolución de direcciones que llegan en el protocolo NHRP. Los MPCs se comunican con el MPS dentro de una misma LIS a través una ELAN. Es muy normal encontrar una red de datos conectada a la red MPoA para lo cual se usa un dispositivo de borde ED (Edge Device), el cual juega el papel de MPC, véase la Fig. 7. El proceso de establecimiento de una conexión es muy similar al visto en el modelo NHRP. Si una computadora 1 desea establecer una comunicación con una computadora 2, conociendo la dirección IP de esta, procede a comunicarse con el ED de origen, el cual es a su vez un MPC. Si el MPC-origen detecta que se trata de un flujo con larga vida intentará establecer una VCC con el MPC-destino para lo cual solicitará la resolución de la dirección IP a una dirección ATM usando el protocolo NHRP que viajará por una cadena de servidores MPS hasta llegar al MPS de destino. Este último devolverá al MPC-origen la dirección ATM del MPC-destino, el cual es a su vez un ED para la red LAN que contiene a la computadora 2. En caso de una comunicación con paquetes de corta duración, no se establecerá una VCC directa, sino una conexión similar al modelo CLIP.

V. MPLS

Los diferentes métodos de superposición vistos hasta el momento nos ayudarán a comprender la esencia de MPLS (Multiprotocol Label Switching), una tecnología que ha dado un enorme impulso a la conmutación de paquetes con QoS [3]. MPLS no es un modelo de superposición, sino de paridad, de manera que IPoATM se convierte prácticamente en una sola red que combina las ventajas de IP y ATM. De esta manera desaparece la necesidad de gestionar dos infraestructuras de red, cada una con sus propios esquemas de direccionamiento, enrutamiento y problemas de gestión. Puede decirse que la conmutación por etiqueta integra la conmutación de capa 2 el enrutamiento de la capa 3 en un solo elemento de red, el LSR (Label-Switching Router). En MPLS el proceso de envío de paquetes se simplifica utilizando una etiqueta corta de longitud fija en lugar de la dirección IP de destino. Es en verdad una realización práctica de la *conmutación orientada a conexión* con la especificidad de que la comunicación es IP y los punteros (que ahora se llamarán etiquetas) definidos entre los nodos de conmutación (que ahora se llamarán LSR) se transmiten en los campos VCI y VPI de celdas ATM. Una etiqueta identifica un circuito virtual en entre dos LSR vecinos y tiene validez solo entre ellos. El camino que se establece en la red para una conexión se conoce como LSP (Label-Switched Path), este es comparable a un canal virtual en ATM, excepto que un LSP es unidireccional. Se dice que un grupo de paquetes que se envían de la misma forma pertenecen a la misma FEC (Forward Equivalence Class), los LSR que sirven de entrada o salida del dominio MPLS se conocen como LER (Label Edge Router). Una etiqueta MPLS contiene 32 bits de los cuales 20 bits son el campo de etiqueta en sí, 3 bits para especificar la clase de

servicio, un bit de pila jerárquica y 8 bits indican el tiempo de vida. Véase la Fig. 8.

El protocolo de distribución de etiquetas, LDP (Label Distribution protocol), es el protocolo usado para la asociación de etiquetas. La Fig. 9 muestra un ejemplo para el proceso de distribución de etiquetas. Cuando un LSR1 detecta que un LSR2 es el siguiente salto para FEC con dirección IP 10.5.16, procede a enviar un mensaje LDP a LSR2 solicitando etiqueta, este último responde con un mensaje LDP que asocia la etiqueta 8 a la FEC=10.5.16.

Se tienen dos casos para la asignación de etiquetas:

- La asignación de etiquetas normalmente ocurre en forma dinámica según la situación presente en el “tráfico de datos”, lo que se conoce como **asignación de etiquetas derivada del tráfico**. La asignación se inicia con una nueva comunicación que demande el uso de etiquetas, el LSP se establece bajo demanda solamente cuando hay tráfico para enviar.
- El establecimiento de etiquetas es ordenado, mediante señalización (es decir, se deriva del tráfico de control) entre nodos de la red. Por ejemplo, un LSR puede en cierto momento decidir actualizar sus tablas de análisis y cambiar las etiquetas que venía manejando para sus entradas. Esto puede suceder ante situaciones anormales que pueden presentarse en la red o en un nodo. El LSR procederá a informar las nuevas etiquetas a su nodos vecinos, esto es lo que se conoce como **asignación de etiquetas derivada de la topología**. Otro caso de señalización es cuando la capa IP usa protocolos de QoS propios de IP como RSVP, el cual reserva recursos en la red para determinado servicio, lo cual a su vez puede conllevar a cambios de etiquetas en las tablas análisis del LSR, esto es lo que se conoce como **asignación de etiquetas derivada por solicitud**.

Cuando un paquete entra al dominio MPLS, el LER de entrada crea una nueva etiqueta que será insertada en las celdas que llevarán este paquete. Los LSR siguientes en el dominio MPLS solo intercambiarán la etiqueta entrante por la etiqueta saliente. Cuando el paquete abandona el dominio MPLS, el LER de salida realiza la extracción de etiqueta. Véanse los ejemplos presentados en la Fig. 10.

En las redes de próxima generación es normal encontrar un dominio MPLS sobre otro dominio MPLS, por ejemplo para separar la señalización que se cruzan entre los MG (Media Gateways) y la información que viaja en el núcleo de la red. Por esta razón es normal encontrar celdas con varios niveles de etiquetas, lo que se conoce como “**pila de etiquetas**”. La etiqueta de más alta pila es la única que determina la decisión de envío. Véase la Fig. 11.

Una técnica que optimiza el uso de etiquetas es el “**El intercalado de VC**”. Gran parte de las conexiones de datos son de tipo multipunto-punto, las cuales se podrían ver como múltiples conexiones punto-punto, Fig. 12. La idea es traducir a una misma etiqueta todos los VC de entrada en un LSR que van a un mismo destino. Esto requiere de algunas estrategias para que el LSR final (o receptor) pueda reconocer los diferentes paquetes, pues estos le llegan con una misma etiqueta. Una de las ideas es ensamblar cada paquete de manera que las celdas de un mismo paquete no resulten

intercaladas con celdas de otros paquetes, para esto se usa el bit “fin de paquete” de AAL5, Fig. 13.

MPLS proporciona grandes ventajas a la **ingeniería de tráfico y para el establecimiento de redes virtuales, como el caso de VPN (Virtual Private Network)**. Es importante analizar como con MPLS no es necesario transportar en la cabecera de cada paquete información sobre la secuencia de LSR que se va a seguir en un LSP, como se haría en las redes de datagrama convencionales. Estas últimas determinarían la ruta salto a salto de acuerdo con el objetivo del camino más corto. Los LSR entran a colaborar para ofrecer nuevos LSP como rutas alternativas para un destino, en caso de observar congestión en la ruta actual, Fig. 14.

VI. LA CONVERGENCIA DE LAS COMUNICACIONES

Los operadores de telecomunicaciones actuales, con frecuencia deben gestionar una red por cada tipo de aplicación (datos, voz, televisión por cable, sistemas móviles, etc.). El esfuerzo por reducir los costos operacionales de estas redes conlleva a la demanda de una red sencilla capaz de transportar todos los servicios que actualmente se ofrecen en redes diferentes.

La Fig. 15 muestra la forma en que se implementan las soluciones de servicio actualmente en donde cada tipo de servicio se presta en una red separada, la Fig. 16 muestra el modelo de una red multiservicio en donde se tendrá una conectividad común a los diferentes tipos de servicios.

Consideraremos ahora el Modelo de Convergencia propuesto por el Consorcio TINA (Telecommunication Information Networking Architecture) [4] para NGN, Fig. 17. Aquí observamos un modelo por capas en el cual la separación del acceso, la conectividad y la provisión del servicio es fundamental. Por ejemplo, tanto las conexiones inalámbricas como las de cable requerirán los mismos métodos de acceso.

En la Fig. 17 encontramos en primer lugar la *Capa de Contenido*, donde se definen los diferentes servicios que la red soporta

En la *Capa de Control y Aplicación* encontramos diferentes servidores que controlan y gestionan los diferentes servicios y accesos que permite la red. Por ejemplo, Servidores de Telefonía cuya función es controlar el establecimiento de llamadas en la red, Servidores SIP (Session Initiation Protocol), Servidores H.323 para servicios multimedia, servidores de acceso móvil, Servidor central de conmutación Móvil (Controla las llamadas móviles y los servicios de IN (Intelligent Network) inalámbricos), Servidor de servicios de Datos Móvil (Controla los servicios de datos IP para dispositivos móviles), Servidores de Contenido (por ejemplo, e-commerce usando API's abiertas como Parlay, JAIN (Java Application for Integrated Networks), CORBA (Common Object Request Broker Architecture)).

Los Conmutadores del Backbone son el corazón de la *Capa de Conectividad* de la red. Hay cierta incertidumbre en la elección de la tecnología de conmutación de paquetes de NGN. La batalla entre IP y ATM ha ardido desde hace muchos años, ambas con poderosos argumentos. La opción

predominante es que, a largo plazo, IP será la tecnología de elección para la conectividad en las redes. Pero hasta el momento, se considera que ATM es la tecnología apropiada para transportar servicios de tiempo real como la voz, con una calidad equivalente a la que se proporciona actualmente en PSTN. A los usuarios empresariales se les deberá proveer una conectividad fácil que les permita un acceso sencillo de alta capacidad para el transporte de voz y datos. La red de conectividad o transporte deberá ofrecer retardo mínimo y máxima flexibilidad, pero sobre todo diversos niveles de QoS. La Capa de Acceso provee los distintos accesos a la red, que pueden ser acceso inalámbrico de tercera generación (3G), acceso desde redes LAN o desde la telefonía (es decir, acceso por conmutación de circuitos o por conmutación de paquetes). A los usuarios residenciales, tecnologías de acceso como ADSL (Asymmetric Digital Subscriber Line) o las recientes soluciones de datos propuestas en GSM (sistema de telefonía móvil digital de Europa), mejor conocida como GPRS (General Packet Radio Services) y que pueden encontrar aplicación en los PCS (Personal Communication services) permitirán utilizar la infraestructura de NGN para agregar servicios de datos de banda ancha.

VII. CONCLUSIONES

Para que la convergencia sea una realidad, algunos de los requerimientos y atributos que las NGN y los servicios que tendrán que soportar son:

- La red de transporte fundamental en NGN estará indiscutiblemente basada en paquetes. Esta infraestructura de transporte por paquetes debe soportar una gran variedad de niveles de calidad de servicio (QoS) haciendo posible servicios que involucren una combinación arbitraria de voz, video, y datos.
- La necesidad de soportar servicios multimedia de banda ancha conlleva a que la infraestructura de control deba ser optimizada para prestar servicios de paquetes con QoS. No obstante, tendrá que exhibir un alto grado de confiabilidad y robustez, como la señalización en el PSTN.
- La infraestructura de control y la arquitectura de servicio de NGN tiene que estar abierta de tal manera que permita a otros proveedores de servicio integrar fácilmente sus servicios con ésta.
- Dada la naturaleza de los servicios soportados por NGN, el CPE (Customer Premises Equipment) encargado de la entrega del servicio debe ser más sofisticado que el teléfono en PSTN. Las NGN, necesitarán soportar un amplio rango de dispositivos conectados por cable e inalámbricos y aparatos que van desde teléfonos de pantalla a sofisticados puestos de trabajo multimedia y centrales de multimedia.
- Finalmente, la división entre control y gestión desaparecerá completamente en NGN. Las funciones de gestión y servicio, para una red de gran magnitud, están empezando a requerir desempeño en tiempo real debido al crecimiento de las expectativas provocadas por Internet. Las funciones de gestión y control en NGN ya no serán

distinguibles y necesitarán ser soportadas por la misma arquitectura.

Indudablemente IP sobre ATM combinado con MPLS serán el principal soporte a las nuevas infraestructuras de red.

BIBLIOGRAFIA

- [1] León Alberto, Widjaja Indra. Redes de datos. Mc Graw Hill. 2002
- [2] Foro ATM. <http://www.atmforum.com/>
- [3] Davie, B.; P. Doolan, y Rekhter: Switching in IP Networks, Morgan Kaufmann, San Francisco 1998.
- [4] <http://www.tinac.com/>