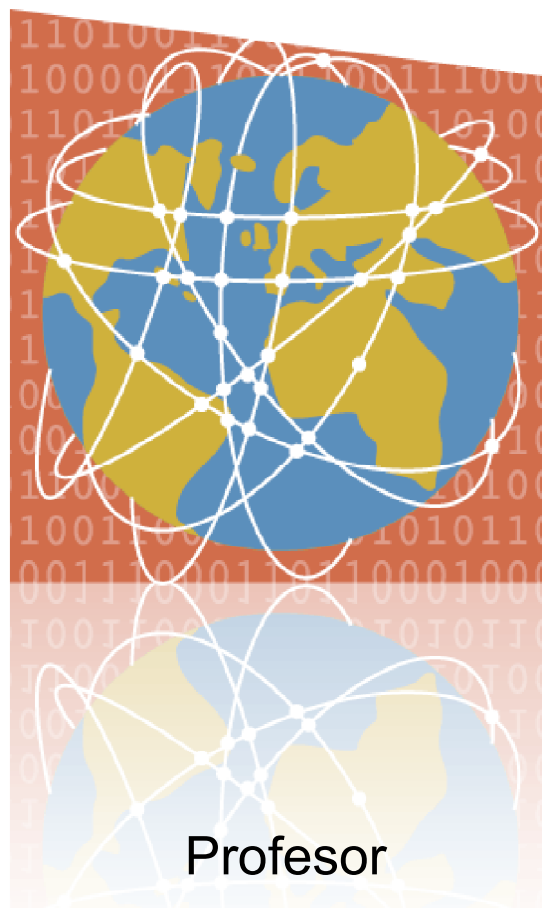


Protocolos de Interconexión de Redes

Tema I

Arquitectura TCP/IP



Profesor

Alberto E. García Gutiérrez

Este Tema se publica bajo licencia:

[Creative Commons 3.0 BY-NC-SA](https://creativecommons.org/licenses/by-nc-sa/3.0/)

Tema I

Arquitectura TCP/IP

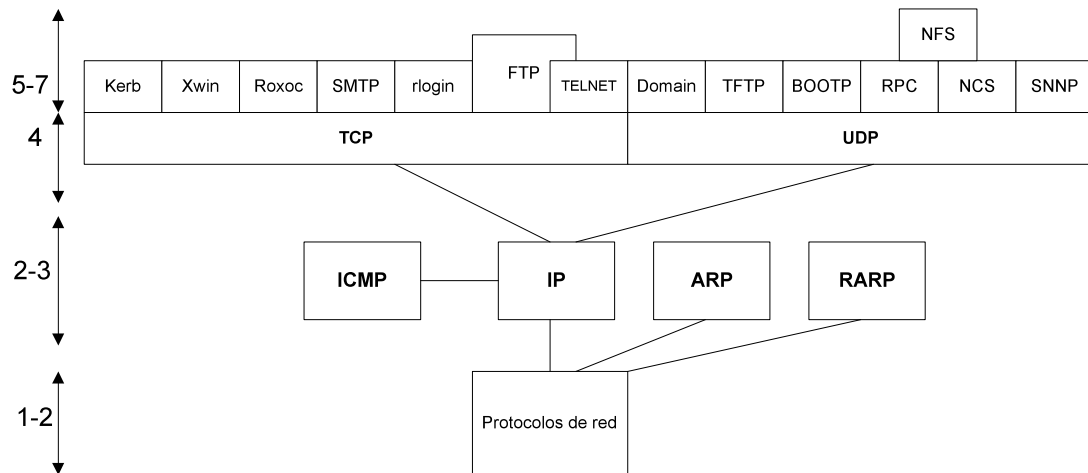
1. EL MODELO INTERNET

La arquitectura *TCP/IP*, y en definitiva el marco de *INTERNET*, está definida por el *IETF* basándose en la red *ARPAnet* y el modelo *DoD* (*Department of Defense*). Dicha arquitectura es anterior al modelo OSI, por lo que no es posible establecer correspondencias exactas entre ambos modelos. El modelo Internet se basa en cuatro capas:

Aplicación	Proceso
Presentación	/
Sesión	Aplicación
Transporte	Host a Host
Red	Internet
Enlace	Acceso a la
Físico	Red

Modelo OSI Modelo Internet

- **Capa de Acceso a la Red:** Realiza el intercambio de datos entre un host y la red, y entre los dispositivos de la misma red, haciendo uso de las direcciones físicas de los equipos implicados. Su adaptación a las normas de red existentes incluye la *conmutación de circuitos (X.21)*, *conmutación de paquetes (X.25)*, *Ethernet*, *IEEE 802.x*, *ATM* y *Frame Relay*.
- **Capa de Interred:** similar a la capa de red de OSI, realizando el encaminamiento de mensajes a través de las diferentes redes. Hace uso de encaminadores denominados “*gateways*”, e implementa un sistema de direcciones lógicas denominadas direcciones IP, relacionadas con las direcciones físicas mediante el *protocolo de resolución de direcciones (ARP)*.
- **Capa de Host a Host:** Es similar a la *capa de Transporte OSI*, manteniendo la integridad de los datos punto a punto. Hace uso de los protocolos de *control de transmisión (TCP)* y de *datagramas de usuario (UDP)*.
- **Capa de proceso/aplicación:** Abarca las capas de sesión, presentación y aplicación, por lo que lo componen gran número de protocolos diferentes: *FTP*, *Telnet*, *SMTP*, *SNMP*, *NFS*, etc.



2. CAPA DE ACCESO A LA RED

La arquitectura *TCP/IP* se ha adaptado a prácticamente la totalidad de las redes existentes hoy día:

2.1. ETHERNET II

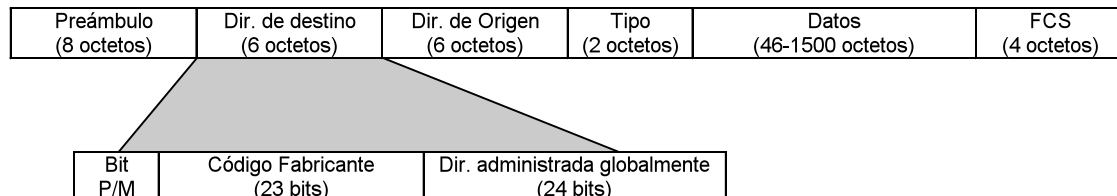
Basada en la primera norma *Ethernet*, denominada *DIX 1.0* (*Digital, Intel y Xerox*, 1980), es la *LAN* dominante entre las redes *TCP/IP*.

Normalmente la topología bus caracteriza a estos sistemas, tanto en sus variantes de cable coaxial fino, como de *cable coaxial grueso* (*10Base2* y *10Base5*, respectivamente), aunque en los últimos años el denominado “*cableado estructurado*” ha impuesto el uso de topologías mixtas, como la bus-estrella y árbol (basadas en el cable trenzado *10BaseT* la utilización de concentradores o *Hubs*).

Dichas topologías suponen la necesidad de realizar una compartición coherente del medio, para lo cual se utiliza uno de los métodos más eficaces, la detección de portadora, fundamental para la reducción del número de colisiones en medios compartidos. Concretamente, *Ethernet* hace uso del *CSMA* (*Carrier Sensing Multiple Access*), en su variante *CD* (*Collision Detection*), según el cual aquel equipo que desee transmitir, debe primero escuchar si el medio está libre, y en cualquier caso, seguir escuchando mientras transmita para

detectar si se produce alguna colisión, caso en el que deberá esperar un tiempo hasta realizar la retransmisión del mensaje.

La estructura de datos utilizada para transmitir y recibir se denomina *marco*, y presenta el siguiente formato:



- *Preámbulo*: Son 7+1 octetos, de patrón fijo *10101010* y *10101011* respectivamente. Señalan el principio del marco, y por tanto no son computados en la longitud del mismo.
- *Dirección de destino/origen*: única para cada nodo de la red, consta de 48 bits, repartidos en tres campos:
 - ✓ *P/M (Physical/Multicast)*: indica si la dirección representa a un equipo individual o a un grupo predefinido de nodos (*multidifusión*).
 - ✓ *Código de fabricante*: 23 bits que identifican al fabricante del dispositivo *Ethernet*.
- *Dirección física o Dirección administrada globalmente*: tres bytes que asigna el fabricante, y que asegura su carácter único.

Existen tres categorías de direcciones *Ethernet*:

- ✓ *Direcciones físicas*: con bit *P/M* a 0.
 - ✓ *Direcciones multidifusión*: con bit *P/M* a 1. La red define el grupo de dispositivos receptores (*RFC 1700*).
 - ✓ *Direcciones de difusión*: todos los bits a 1 (*FF:FF:FF:FF:FF:FF*).
-
- *Tipo (EtherType)*: son 16 bytes que indican el tipo de dato enviado, permitiendo la *demultiplexación* y la selección de la pila de protocolo adecuada.

0800 Internet Ipv4

0805 X.25 Level 3

0806 ARP

809B AppleTalk

80D5 IBM SNA

8137-8138 Novell

- *Datos*: contiene la unidad de datos de las capas superiores, conteniendo entre 46 y 1500 octetos
- *FCS (secuencia de verificación de marco)*: son 32 bits basados en un CRC.

La longitud mínima del marco es por tanto $6+6+2+46+4 = 64$ octetos, más el preámbulo, asegura los 576 bits requeridos por la detección de colisiones.

2.2. IEEE 802

Las normas *IEEE 802* contemplan dos capas que en conjunto suplen los niveles *OSI* del físico y enlace:

- *Control de enlace lógico (LLC)*: 802.2

Aplicación		
Presentación		
Sesión		
Transporte		
Red		
Enlace	LLC IEEE 802.2	
	MAC + IEEE 802.3	MAC + IEEE 802.5
Físico	CSMA/CD	Token Ring

- *Control de Acceso al Medio (MAC)*, tanto *Ethernet 802.3* y *Token Ring 802.5*

2.2.1. DIRECCIÓN MAC

Las direcciones físicas en estos casos quedan definidas al nivel del protocolo *MAC*, mediante formatos de 16 y 48 bits, siendo este último el más utilizado, compatible con *IEEE 802*, *ISO 8802* y similar a *Ethernet*:

		Identificador único de organización (22 bits)	Dirección administrada por la organización (24 bits)
--	--	---	--

Bit U/L (0 = dir. Administrada universalmente / 1 = dir. Administrada localmente)

Bit I/G (0 = dir. Individual / 1 = dir. De grupo)

- I/G: '0' significa que la dirección es individual, '1' que es *multidifusión*.
- U/L: '0' indica que la dirección tiene el formato universal, es decir 22 bits de "*identificador único de organización*", a los que se suman 24 bits de "*dirección administrada por la organización*" correspondiente. '1' significa que los 46 bits son administrados localmente por el software del dispositivo de red.

2.2.2. IEEE 802.2

La capa de control de enlace lógico realiza, como misión más importante la *multiplexación* y posterior *demultiplexación* de los datos procedentes de los diferentes protocolos superiores, con una sustancial variación respecto a *Ethernet II*. Los diferentes protocolos de capa superior acceden a la capa *LLC* a través de un *LSAP* (*Link Service Access Point*), que no es sino una dirección lógica que identifica al protocolo del que proceden o se envían los datos.

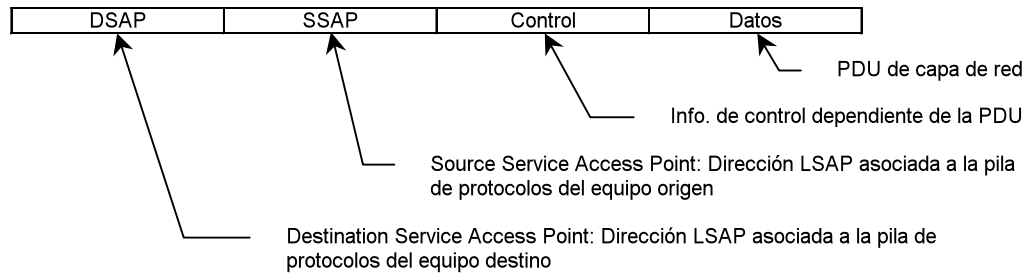
Además, *LLC* proporciona varios servicios de entrega que determinan el nivel de integridad de la comunicación establecida, ya que los dispositivos disponen de memorias intermedias limitadas que hacen que en determinados casos la probabilidad de pérdida de los marcos aumenta. Para ello el control de flujo hace uso de diferentes métodos, desde el más simple de Parada y espera propio de servicios de *datagramas* sin conexión, hasta los más complejos de “ventana deslizante” de los servicios orientados a la conexión.

De igual manera, aunque la capa *MAC* realiza la detección de errores, la recuperación de los mismos compete a la capa *LLC*, haciendo uso de las *ARQ* (solicitudes de repetición automática), ya sean estas “*de parada y espera*” (acuse de recibo por cada marco en servicios sin conexión), o “*de retroceso*” (retransmisión selectiva de marcos en servicios de conexión).

Así, *LLC* implementa tres tipos básicos de servicio:

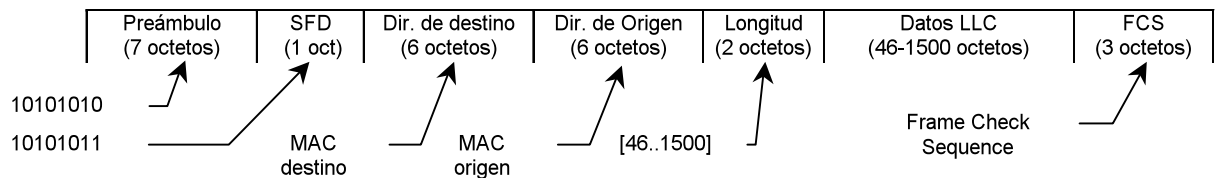
- *Servicio de datagramas sin acuse de recibo* (Tipo 1): Punto a punto, multipunto y difusión sin conexión ni detección o recuperación de errores ni control de flujo.
- *Servicio de circuito virtual* (Tipo 2): Modo conexión con secuenciamiento, control de flujo, detección y recuperación de errores.
- *Servicio de datagramas con acuse de recibo* (Tipo 3): Datagramas punto a punto, a medias entre los de tipo 1 y 2.

En cualquier caso, la *unidad de datos de protocolo* (*PDU*) resultante sería:

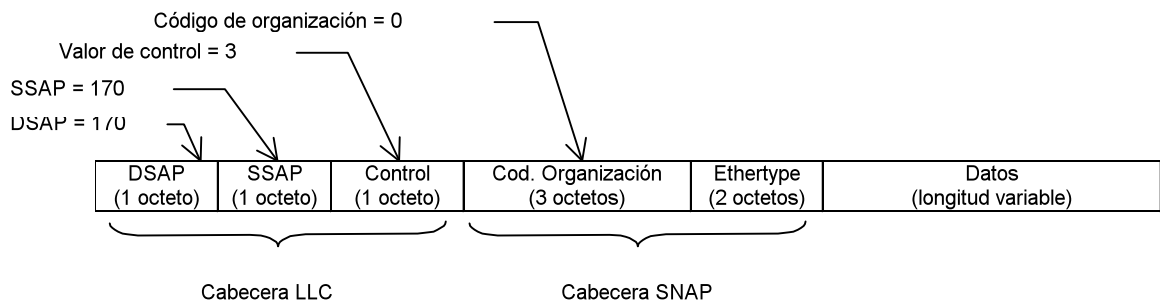


2.2.3. IEEE 802.3

El formato del marco *IEEE 802.3* es similar al marco *Ethernet II*:



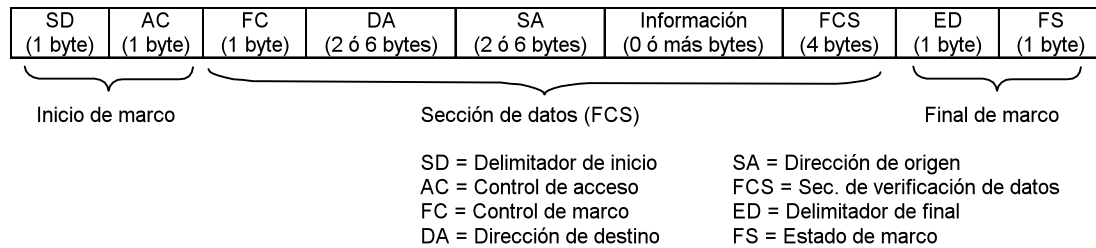
De hecho, la única diferencia importante radica en que el marco *802.3* utiliza el campo de longitud de 2 octetos, en vez del *EtherType* de *Ethernet II*. De esta forma, cuando se ejecuta *TCP/IP* sobre una capa física *IEEE*, la información *EtherType* se codifica en el marco utilizando el *protocolo de acceso a la subred (SNAP)*, como extensión de la cabecera *LLC* (ver *RFC 1042*).



Con esta modificación, un marco *IEEE 802.3* con *encapsulación SNAP* dispone de 1492 octetos para datos, frente a los 1500 de *Ethernet II*. Sin embargo, *SNAP* es compatible con todas las *LAN IEEE 802.x*, incluido *token ring*.

2.2.4. IEEE 802.5

El formato del marco *token ring* varía considerablemente respecto de los vistos hasta el momento:



Consta de tres secciones:

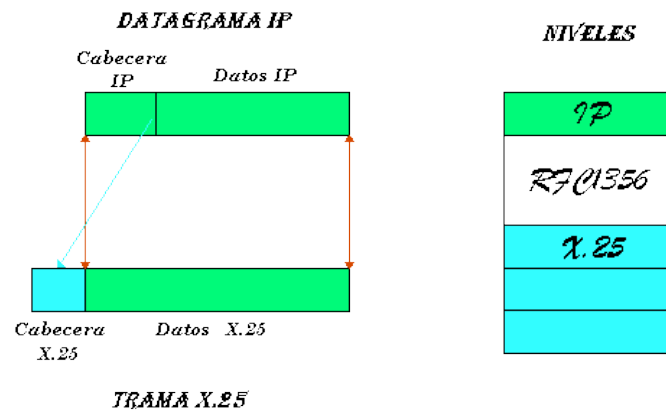
- *SFS (secuencia de inicio de marco).*
- *Sección de datos.*
- *EFS (secuencia de final de marco).*

Las cuales a su vez, incluyen los siguientes campos:

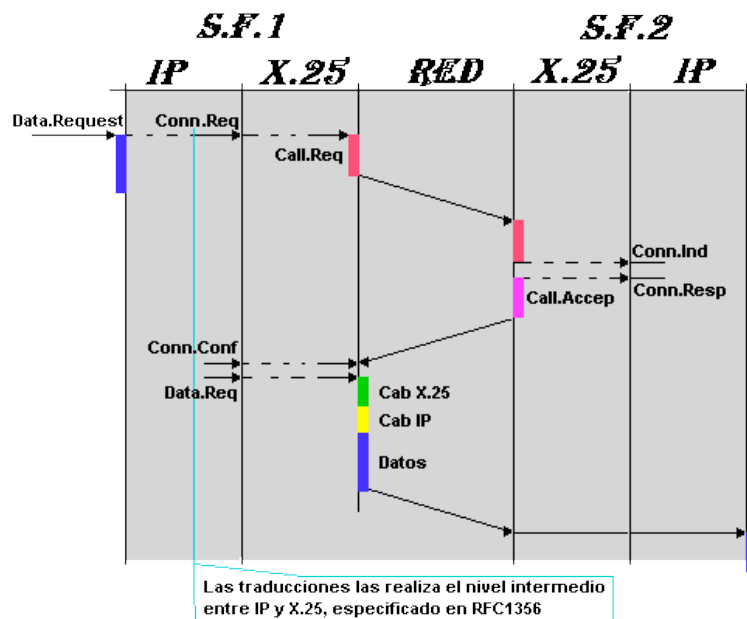
- *SD (Delimitador de inicio):* a base de utilizar señales eléctricas de distinto formato al normal, viola las reglas de codificación de datos del marco.
- *AC (Control de acceso):* incluye bits de prioridad y reserva para establecimiento de prioridades, así como un bit de control para administración de red, y un bit de señal para indicar si el marco es de señal o de datos.
- *FC (Control de marco):* indica si el marco contiene datos LLC, o si es un marco de control MAC.
- *DA/SA (Direcciones de Destino/Origen):* admite direcciones de 16 y 48 bits.
- *FCS (Secuencia de verificación de marco):* CRC sobre los campos FC, DA, SA e información.
- *DE (Delimitador de final):* Similar al SD, este campo incluye dos bits de control, uno para indicar que el marco es intermedio o final, y otro de error para casos de un FCS erróneo.
- *FS (Estado de marco):* Mediante bits de control indica los reconocimientos

2.3. X.25

Aunque *X.25* y *TCP/IP* han convivido durante mucho tiempo, su combinación no resulta sencilla, ya que *X.25* es una norma orientada a conexión, que no permite la difusión o multidifusión de mensajes. Sin embargo, la mayoría de los servicios de *TCP/IP* utilizan difusiones, ya que el *protocolo de resolución de direcciones (ARP)* utiliza mensajes de difusión para encontrar las direcciones físicas de los nodos de destino. Por ello, este tipo de protocolos requieren de un esfuerzo adicional para funcionar en una infraestructura de red orientada a conexión.



Cuando las computadoras remotas se comunican a través de una red pública como X.25, la red opera con independencia de los protocolos de cada dispositivo. Si los nodos finales ejecutan *TCP/IP*, los *datagramas IP* son encapsulados en marcos X.25 y se transmiten a través de la *WAN*. La *desencapsulación* del dispositivo de destino permite recuperar el *datagrama IP* original. Esta técnica recibe el nombre de “*tunnelling*”, y solo tiene sentido cuando ambos extremos de la conexión X.25 se ponen de acuerdo para intercambiar paquetes *TCP/IP*.



El esquema de direccionamiento usado por redes X.25 se recoge en el estándar X.121. Las direcciones físicas de X.121 constan de 14 dígitos, con 10 dígitos asignados al vendedor que proporciona servicios de X.25. De esta forma es difícil establecer una asignación de direcciones Internet a direcciones X.25. La solución es tener una tabla donde se mapeen las direcciones X.25 con las direcciones Internet.

2.4. Frame Relay

Frame Relay es una norma para redes de área extensa de conmutación de paquetes y banda ancha. Es una actualización de la norma *X.25*, y la *ITU* es el organismo encargado de su normalización. Aunque *Frame relay* abarca tanto el nivel físico como el de enlace, no se encarga de la corrección de errores ni del control de flujo.

El campo de datos de un marco *frame relay* se denomina *payload*. La implementación de la red define el tamaño de dicho campo, por lo que las redes *frame relay* son configurables. Por su parte, el soporte del protocolo *TCP/IP* utiliza la encapsulación *SNAP*.

2.5. ATM

Los expertos en redes piensan que *ATM* es la tecnología con mayores posibilidades de éxito. Un posible inconveniente reside en que *ATM* es una red orientada a conexiones que requiere una conexión entre cada par de nodos. Este tipo de redes no admiten directamente las transmisiones de difusión (de uno a varios) que *TCP/IP* y otros protocolos utilizan ampliamente. La industria *ATM* ha desarrollado una emulación *LAN* para superar este problema. Evita que *TCP/IP* tenga que modificarse para funcionar sin mensajes de difusión dotando a *ATM* de la apariencia de una *LAN* convencional como *Ethernet* o *token ring*. El coste de este método consiste en una pérdida significativa de ancho de banda consumido por el proceso de emulación. Hablaremos más adelante de este método con más detenimiento.

3. CAPA DE INTERRED

La capa de red es la que se encarga de la distribución de los datos a través de la red. *IP* es el protocolo de Internet, incluido en las *RFC 791*, *919* y *950*, aunque hace uso de un conjunto de protocolos adicionales, especialmente el *ICMP* (*Internet Control Messaging Protocol*).

IP va a realizar tres funciones primordialmente: el direccionamiento, la fragmentación / reensamblaje de datagramas y la entrega de datagramas inter-redes.

El tamaño de un datagrama *IP* se especifica en un campo de dos bytes, por lo que su valor máximo es de 65535 bytes, pero muy pocas redes admiten este valor. Normalmente el nivel de enlace no fragmenta, por lo que el nivel de red adapta el tamaño de cada paquete para que viaje en una trama; con lo que en la práctica el tamaño máximo de paquete viene determinado por el tamaño máximo de trama característico de la red utilizada. Este tamaño máximo de paquete se conoce como *MTU* (*Maximum Transfer Unit*); a continuación aparecen algunos ejemplos de valores de *MTU* característicos de las redes más habituales:

Protocolo a nivel de enlace	MTU
<i>PPP (valor por defecto)</i>	<i>1500</i>
<i>PPP (bajo retardo)</i>	<i>296</i>
<i>SLIP</i>	<i>1006 (límite original)</i>
<i>X.25</i>	<i>1600 (varía según las redes)</i>
<i>Frame relay</i>	<i>al menos 1600 normalmente</i>
<i>SMDS</i>	<i>9235</i>
<i>Ethernet version 2</i>	<i>1500</i>
<i>IEEE 802.3/802.2</i>	<i>1492</i>
<i>IEEE 802.4/802.2</i>	<i>8166</i>
<i>Token Ring IBM 16 Mb/s</i>	<i>17914 máximo</i>
<i>IEEE 802.5/802.2 4 Mb/s</i>	<i>4464 máximo</i>
<i>FDDI</i>	<i>4352</i>
<i>Hyperchannel</i>	<i>65535</i>
<i>ATM</i>	<i>9180</i>

Valor de MTU para los protocolos más comunes a nivel de enlace.

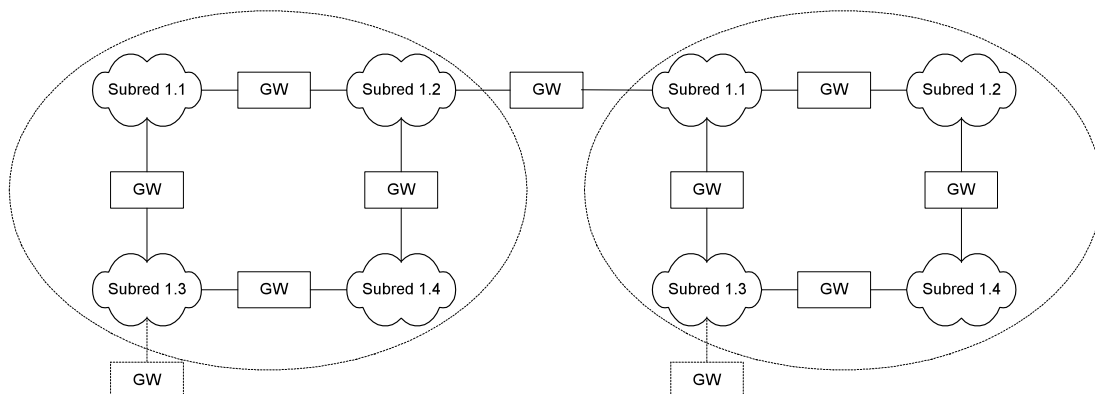
3.1. DIRECCIONAMIENTO IP

Los protocolos de capas superiores no utilizan directamente las direcciones del hardware de la red. En su lugar hacen uso de un sistema de direcciones físicas para identificar a los hosts, e identificaciones lógicas denominadas direcciones *IP*, de tal forma que incluso el encaminamiento está codificado en la misma.

Las direcciones *IP* tienen una longitud de 32 bits, divididos en dos campos:

- *Identificador de red (netid)* a la que está conectado el host.
- *Identificador de host (hostid)*, único para cada host en una red dada.

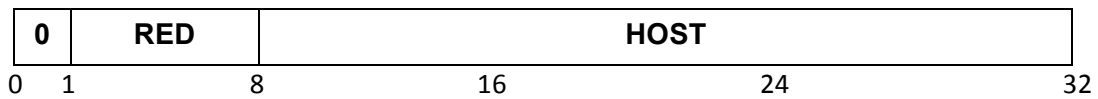
Todos los host que pertenecen a una misma red deben tener el mismo identificador de red. Los hosts con distintos identificadores de red deben comunicarse a través de un encaminador. Así una interred *TCP/IP* es una red de redes interconectadas a través de encaminadores.



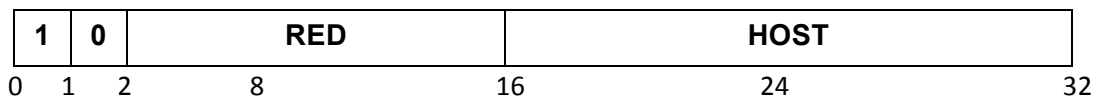
Bajo esta configuración, se consideraron tres posibles escenarios:

- Pocas redes con muchos hosts.
- Una cantidad media de redes con número intermedio de hosts.
- Gran cantidad de redes con pocos hosts.

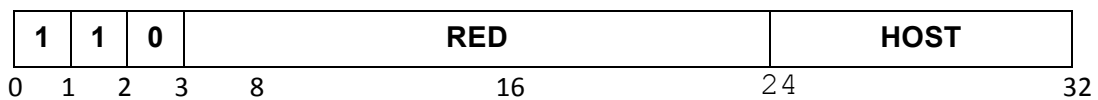
Ante dichas situaciones, se establecieron cinco clases de direcciones *IP* con distinto número de bits en los identificadores de red:



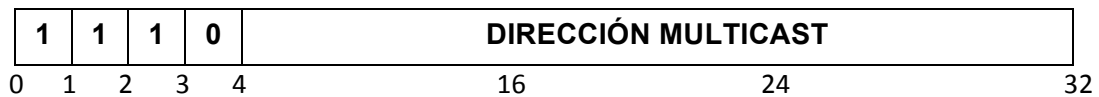
Clase A: Redes con más de 2^{16} hosts.



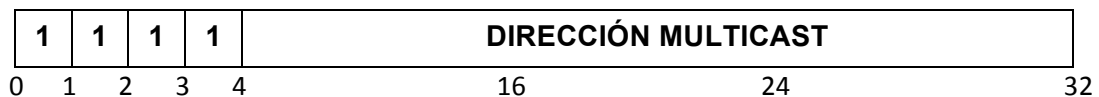
Clase B: Redes entre 2^8 y 2^{16} hosts.



Clase D: Para redes que soportan la entrega multidifusión.



Clase E: Son redes de este tipo tienen un uso experimental.



Los 32 bits son representados mediante la notación decimal de los 4 octetos, así, las direcciones *IP* disponibles serían:

Clase	Rango	Id. de red	Id. de Host
A	1 - 126	126	16.777.214
B	128 - 191	16.384	65.534
C	192 - 223	2.097.152	254
Multicast	224 - 255	-	-

Hay que tener en cuenta que:

- Los identificadores de red y de host con valor 0 significan “red”. Ejm.: la dirección IP 155.123.0.0 indica la red 155.123, mientras que la dirección 0.0.0.35 identifica al host 35 de la red local.
- El identificador de red 127 es una dirección de retorno, de tal forma que los mensajes que van dirigidos a dicho identificador se reflejan.
- Los identificadores de host con valor 255 quedan restringidos a las difusiones. Un mensaje dirigido a la dirección 255.255.255.255 es enviado a todos los hosts de la red. Un mensaje dirigido a la dirección 130.20.255.255 se envía a todos los hosts de la red 130.20.
- El último octeto de una dirección IP no puede tener ni los valores 0 ni 255.
- La dirección 0.0.0.0 identifica al host actual.
- Las redes 127.0.0.0, 128.0.0.0, 191.255.0.0, 192.0.0.0 y el rango 240.0.0.0 en adelante están reservados.
- La dirección 127.0.0.1 se utiliza para pruebas de loopback: todas las implementaciones de IP devuelven a la dirección de origen los datagramas enviados a esta dirección sin intentar enviarlos a ninguna parte.
- Las redes 10.0.0.0, 172.16.0.0 a 172.31.0.0, y 192.168.0.0 a 192.168.255.0 están reservadas para redes privadas ('intranets') por el RFC 1918; estos números no se asignan a ninguna dirección válida en Internet y por tanto pueden utilizarse para construir redes, por ejemplo detrás de un cortafuego, sin riesgo de entrar en conflicto de acceso a redes válidas de la Internet.

Cuando las redes no van a estar conectadas a Internet, los administradores pueden utilizar cualquiera de las direcciones IP existentes. Sin embargo, cuando la red es conectada a Internet, debe hacer uso de su identificación, por lo que el número de direcciones puede resultar a veces demasiado restringido. Para ello se desarrolló un procedimiento de subred (RFC 950), por el que se permite distribuir identificadores de host de una red dada en varias subredes. El mecanismo utiliza algunos bits del identificador de host para identificar la subred:

- Si no se utiliza subred $Dir. IP = id. de red + id. de host$
- Si se utiliza subredes $Dir. IP = id. de red + id. de subred + id. de host$

Para poder identificar cuál es la parte reservada para el identificador de subred, se utiliza una “máscara de subred”, que consiste en otro número de 32 bits, en el que un bit a 0 indica que en su posición hay un bit del identificador de host.

Dir.IP	1	0	1	0	0	0	0	1	0	1	1	0	1	0	1	1	0	1	1	1	1	0	1	1	1	1
Máscara	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1

Las máscaras de subred predeterminadas son:

- Clase A: 255.0.0.0
- Clase B: 255.255.0.0
- Clase C: 255.255.255.0

Bits de subred	Número de subredes	Nº subredes (subred cero)	Bits de host	Número de hosts	Máscara
0	0	0	16	65534	255.255.0.0
1	0	2	15	32766	255.255.128.0
2	2	4	14	16382	255.255.192.0
3	6	8	13	8190	255.255.224.0
4	14	16	12	4094	255.255.240.0
5	30	32	11	2046	255.255.248.0
6	62	64	10	1022	255.255.252.0
7	126	128	9	510	255.255.254.0
8	254	256	8	254	255.255.255.0
9	510	512	7	126	255.255.255.128
10	1022	1024	6	62	255.255.255.192
11	2046	2048	5	30	255.255.255.224
12	4094	4096	4	14	255.255.255.240
13	8190	8192	3	6	255.255.255.248
14	16382	16384	2	2	255.255.255.252

15	32766	32768	1	0	255.255.255.254
16	65532	65536	0	0	255.255.255.255

Subredes y máscaras que pueden definirse en una red clase B

Bits de subred	Número de subredes	Nº subredes (subred cero)	Bits de host	Número de hosts	Máscara
0	0	0	8	254	255.255.255.0
1	0	2	7	126	255.255.255.128
2	2	4	6	62	255.255.255.192
3	6	8	5	30	255.255.255.224
4	14	16	4	14	255.255.255.240
5	30	32	3	6	255.255.255.248
6	62	64	2	2	255.255.255.252
7	126	128	1	0	255.255.255.254
8	254	256	0	0	255.255.255.255

Subredes y máscaras que pueden definirse en una red clase C

3.1.1. Superredes: Routing classless (CIDR)

El rápido crecimiento de la Internet está creando varios problemas, el más importante de los cuales es el rápido agotamiento de las direcciones. Las redes clase A ya prácticamente no se asignan, y resulta muy difícil obtener una clase B. Muchas empresas no tienen bastante con una clase C pero una B les resulta excesiva. Entre las medidas que se están adoptando para paliar el problema de escasez de direcciones está la de asignar un conjunto de redes clase C donde antes se asignaba una clase B. De esta forma se puede ajustar mejor el rango de direcciones asignado a las necesidades reales previstas de cada organización

Esto ha resuelto un problema pero ha creado otro: el crecimiento desmedido de las tablas en los routers. Muchos routers de la Internet funcionan con unas tablas de encaminamiento que solo recogen una pequeña parte de todas las redes existentes, y disponen de una ruta por defecto por la cual se sale hacia el resto de la Internet. Sin embargo, a medida que nos aproximamos al centro o 'backbone' de la Internet las tablas tienen que ser necesariamente más exhaustivas; en las tablas de los denominados 'core routers' se mantienen entradas para la mayoría de las redes existentes en la Internet.

Al asignar a una organización varias clases C donde normalmente hubiera bastado con una clase B hace falta definir varias entradas en las tablas de routing (una por cada clase C asignada) cuando antes habría bastado una para toda la clase B. Actualmente hay mas de 100.000 redes registradas en la Internet. Además del costo en memoria RAM que supone el mantener tablas extremadamente grandes en los routers los algoritmos de búsqueda se complican y no funcionan adecuadamente ya que fueron diseñados pensando en muchas menos entradas. El crecimiento de la Internet se está produciendo a un ritmo que duplica el número de redes conectadas cada 9 meses, mientras que la tecnología sólo permite duplicar la capacidad y potencia de los routers cada 18 meses. En esta situación el problema de la explosión de las tablas de routing se convirtió en un problema aún más grave que la escasez de direcciones. Según cálculos hechos por la IETF en 1993 de seguir produciéndose el crecimiento normal en el número de redes y rutas la Internet se colapsaría hacia 1998.

Para solucionar este problema se adoptó en 1993 un sistema denominado **CIDR** (*Classless InterDomain Routing*) descrito en el RFC 1519. Se trata de dos medidas complementarias:

La primera consiste en establecer una jerarquía en la asignación de direcciones. Antes de CIDR la asignación de números de red se hacía por orden puramente cronológico, independientemente de la ubicación geográfica, lo cual equivalía en la práctica a una asignación aleatoria del número de red. Con CIDR se han asignado rangos por continentes:

- 194.0.0.0 a 195.255.0.0 para Europa
- 198.0.0.0 a 199.255.0.0 para Norteamérica
- 200.0.0.0 a 201.255.0.0 para Centro y Sudamérica
- 202.0.0.0 a 203.255.0.0 para Asia y la zona del Pacífico

A su vez dentro de cada uno de estos rangos se ha dado una parte a cada país, y dentro de éste un rango a cada proveedor de servicios Internet. Con esta distribución regional de los números y los cambios pertinentes en el software las entradas en las tablas de routing pueden agruparse, con lo que las tablas se simplifican; por ejemplo un router en Japón puede poner una sola entrada en sus tablas indicando que todos los paquetes dirigidos a las redes 194.0.0.0 hasta 195.255.0.0 vayan a la interfaz que da acceso a Europa. Se ha establecido pues un criterio geográfico jerárquico en las direcciones IP.

Una consecuencia curiosa de la asignación de rangos de direcciones por proveedor es que si una empresa cambia de proveedor normalmente tendrá que 'devolver' a este sus direcciones, y solicitar direcciones nuevas al nuevo proveedor; por supuesto tendrá que modificar las direcciones IP de todas sus máquinas.

La segunda medida adoptada por CIDR es en realidad es un caso particular de la anterior. Consiste en dar a cada organización (bien directamente o a través de su proveedor correspondiente) un conjunto de redes clase C ajustado a lo que son sus necesidades previstas, dándole siempre un rango contiguo y un número de redes que sea potencia entera de 2 (es decir 1, 2, 4, 8 redes, etc.) elegidas de modo que tengan una máscara común en la parte de red; por ejemplo un grupo de 8 redes clase C puede hacerse variando únicamente los 3 últimos bits (22 a 24) de la parte de red, por lo que el grupo deberá tener comunes los primeros 21 bits de la parte de red, y deberá empezar necesariamente por un valor múltiplo de ocho y abarcar los siete valores siguientes.

Por ejemplo, supongamos que la Universidad de Cantabria solicita a su proveedor direcciones de red y justifica la previsión de tener 1200 hosts en un plazo razonable; como 4 redes clase C no serían suficientes se le asignan 8 redes, que permiten llegar a 2032 hosts (254 x 8). Supongamos que somos el proveedor de la Universidad de Cantabria y que tenemos disponible para nuestro uso el rango 195.100.0.0 a 195.100.255.0 (obsérvese que en este caso el número de redes clase C válidas es de 256, no 254, pues las redes 195.100.0.0 y 195.100.255.0 son perfectamente utilizables); supongamos que hemos ido asignando redes clase C a nuestros clientes por orden cronológico, y que nos queda libre a partir del 12, es decir de la red 195.100.12.0; si le asignamos de 195.100.12.0 a la 195.100.19.0 el rango es contiguo pero 12 no es múltiplo de 8, por lo que la máscara no es común (es decir, los primeros 21 bits no son iguales en las 8 redes). El primer rango de 8 redes con máscara común sería 195.100.16.0 a 195.100.23.0, que es el que le asignaríamos a nuestro cliente. Las redes 195.100.12.0 a 195.100.15.0 quedarían libres para otros clientes, por ejemplo podrían formar un grupo contiguo de 4 redes (22 bits iguales) para alguno que necesitara conectar entre 509 y 1016 hosts.

Con esta asignación de direcciones el aspecto normal de una tabla de routing para enviar paquetes a la Universidad de Cantabria sería:

- *ip route 195.100.16.0 255.255.255.0 interfaz*
- *ip route 195.100.17.0 255.255.255.0 interfaz*
- *ip route 195.100.18.0 255.255.255.0 interfaz*
- *ip route 195.100.19.0 255.255.255.0 interfaz*
- *ip route 195.100.20.0 255.255.255.0 interfaz*
- *ip route 195.100.21.0 255.255.255.0 interfaz*
- *ip route 195.100.22.0 255.255.255.0 interfaz*
- *ip route 195.100.23.0 255.255.255.0 interfaz*

donde 'interfaz' especificaría la interfaz por la que se accedería a la Universidad de Cantabria.

El uso de CIDR permite sintetizar estas ocho entradas en una sola que sería como sigue:

- *ip route 195.100.16.0 255.255.248.0 interfaz*

Obsérvese que CIDR es en realidad el mismo mecanismo que las subredes, pero aplicado en sentido inverso. Las subredes permiten dividir una red, ampliando la parte red a costa de la parte host de la dirección. El CIDR funde diferentes redes en una, reduciendo la parte red y ampliando la parte host. Por este motivo el CIDR también se conoce como supernet addressing.

Cuando se utiliza CIDR como en el ejemplo anterior los hosts finales que se encuentran en distintas redes clase C no pueden hablar directamente entre ellos, han de hacerlo a través de un router, a menos que soporten CIDR que no es lo normal en hosts.

Un grupo CIDR de clases C siempre funciona como subnet-zero, es decir:

- No existe una dirección que haga referencia al grupo; en nuestro ejemplo la dirección *195.100.16.0* haría referencia a la primera red clase C únicamente.
- No existe una dirección broadcast del grupo; en nuestro ejemplo la dirección *195.100.23.255* es la dirección broadcast de la red *195.100.23.0* únicamente.

El espacio de redes clase A, que suponen la mitad del espacio total, está asignado actualmente sólo en un 50% aproximadamente. Se está estudiando la posibilidad de dividir la parte no utilizada de este rango de direcciones en redes de menor tamaño para su asignación mediante CIDR, igual que se hace actualmente con el rango *194.0.0.0-203.255.0.0*.

3.2. PROTOCOLO IP

La unidad de transferencia del protocolo *IP* es el datagrama, con un tamaño máximo de 64 kb.

0	4	8	16	31
Versión	HLEN	TOS	Longitud Total	
Identificación			Flags	Offset
TTL		Protocolo	Checksum	
Dirección IP Origen				
Dirección IP destino				
Opciones				Padding
DATOS				

- *Versión*: Estos 4 bits identifican la versión del protocolo IP utilizada, debiendo coincidir en ambos extremos.
- *HLEN*: Mediante 4 bits se da la longitud de la cabecera del datagrama, medida en palabras de 32 bits, siendo de 20 para todos aquellos datagramas sin campo de *Opciones* ni de *Padding*.
- *Longitud Total*: indica la longitud del datagrama medida en octetos, incluyendo los datos (valor máximo 64 k).
- *TOS (Type of Service)*: Son ocho bits codificados:
 - Los tres primeros bits indican *PRECEDENCIA*, especificando la importancia del datagrama. La prioridad actúa alterando el orden de los paquetes en cola en los routers, pero no modifica la ruta de estos.
 - Bits *D*, *T*, *R* y *C*: especifican el tipo de transporte deseado para el datagrama, indicando respectivamente “*bajo retardo*”, “*rendimiento alto*”, “*fiabilidad alta*” y “*bajo costo*”.

- El último bit no tiene un uso específico.

Valor	Descripción
0000	Valor por defecto
0001	Mínimo costo
0010	Máxima fiabilidad
0100	Máximo rendimiento
1000	Mínimo retardo
1111	Máxima seguridad

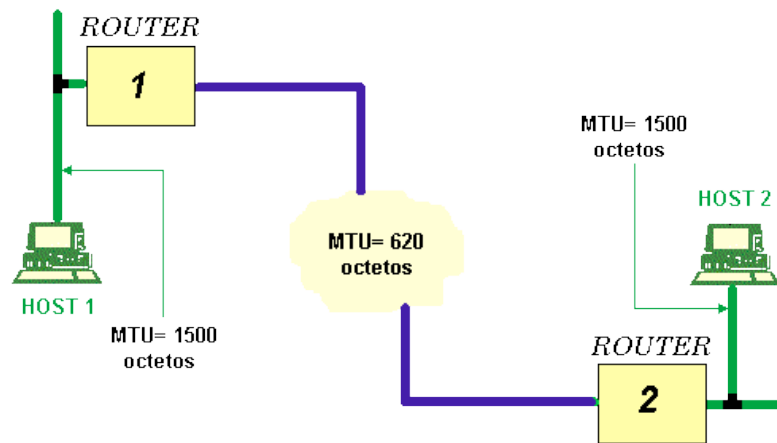
Combinaciones válidas de los bits DTRC

Para cada aplicación existe un valor de TOS recomendado. Por ejemplo, para Telnet se recomienda 1000 (mínimo retardo), para FTP 0100 (máximo rendimiento) y para TNP (news) 0001 (mínimo costo). Además, algunos routers utilizan el subcampo TOS para encaminar los paquetes por la ruta óptima en función del valor especificado (podrían tener una ruta diferente según se desee mínimo retardo o mínimo costo, por ejemplo); también pueden utilizar el valor del campo TOS para tomar decisiones sobre que paquetes descartar en situaciones de congestión (por ejemplo descartar antes un paquete con mínimo costo que uno con máxima fiabilidad). Algunos routers simplemente ignoran este subcampo.

La especificación del protocolo IP establece que todo host debe ser capaz de aceptar y reensamblar datagramas de al menos 576 octetos, la misma dimensión que todo encaminador debe ser capaz de gestionar. IP se encarga de fragmentar los datagramas largos en datagramas compatibles, de tal forma que en la cabecera de cada datagrama se incluya información que permite al host receptor identificar la posición del fragmento y reensamblar el datagrama original. Tal información la constituyen los siguientes tres campos:

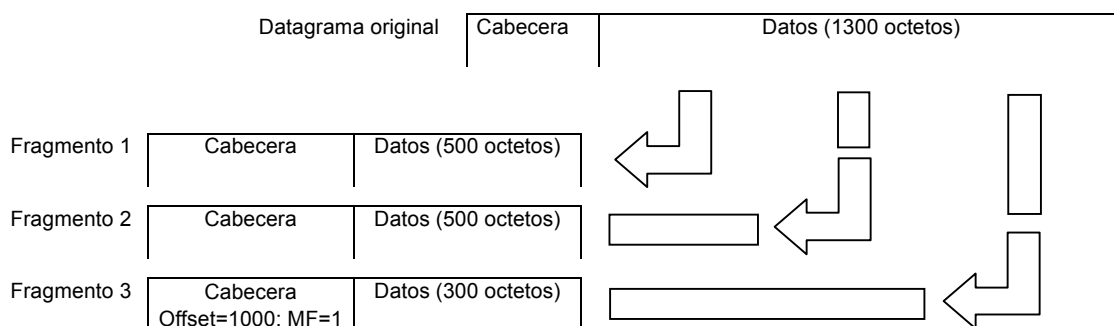
- *Identificación*: mediante 16 bits se indica a qué datagrama pertenecen los fragmentos. Cuando un datagrama atraviesa un router que conecta con una red cuyo MTU es inferior, este debe ser fragmentado, para adecuarlo precisamente a dicha limitación.
- *Flags*: Tres bits, de los cuales los dos primeros significan:
 - *Flag de no fragmentación (DF)*: obliga a que los datagramas conserven su tamaño original. Por ejemplo, si un ordenador arranca su sistema operativo a través de la red solicitará que el ejecutable correspondiente se le envíe desde algún servidor a través de la red como un único datagrama (ya que en ese estado él aun no está capacitado para reensamblar datagramas). Si un datagrama con el bit DF puesto no puede pasar por una red el router lo rechazará con un mensaje de error al emisor. Existe una técnica para averiguar el MTU de una ruta (denominada 'path MTU discovery') que consiste en enviar un datagrama

grande con el bit DF puesto al destino deseado; si se recibe un mensaje de error se envía otro más pequeño, hasta que el emisor averigua a base de tanteos cual es el valor de MTU de la ruta correspondiente, y a partir de ahí puede utilizarla para todos los datagramas sin riesgo de que sean fragmentados en el camino (siempre y cuando la ruta no cambie sobre la marcha).



- *More Fragment Bit (MF)*: indica cuál es el último fragmento dentro de un datagrama.
- *Offset*: Mediante 13 bits se indica la posición relativa del fragmento en el datagrama original.

Cuando un datagrama excede el tamaño máximo y se fragmenta, la cabecera del fragmento incluye el parámetro de offset, que especifica la posición del primer octeto del fragmento en el datagrama global.



- *TTL (tiempo de vida)*: indica el máximo tiempo de vida de un datagrama como el máximo tiempo que puede estar ese paquete circulando en la red, siendo actualizado en cada transición a través de una pasarela.
- *Protocolo*: Indica el protocolo que se encuentra por encima del protocolo IP. El campo protocolo especifica a que protocolo del nivel de transporte corresponde el datagrama. La tabla de protocolos válidos y sus correspondientes números son controlados por el IANA (Internet Assigned Number Authority) y se especifican (junto con muchas otras tablas de números) en un RFC muy especial, denominado 'Assigned Numbers', que se actualiza regularmente; el vigente actualmente es el RFC 1700. Algunos de los posibles valores del campo protocolo son los siguientes:

Valor	Protocolo	Descripción
0		<i>Reservado</i>
1	<i>ICMP</i>	<i>Internet Control Message Protocol</i>
2	<i>IGMP</i>	<i>Internet Group Management Protocol</i>
3	<i>GGP</i>	<i>Gateway-to-Gateway Protocol</i>
4	<i>IP</i>	<i>IP en IP (encapsulado)</i>
5	<i>ST</i>	<i>Stream</i>
6	<i>TCP</i>	<i>Transmission Control Protocol</i>
8	<i>EGP</i>	<i>Exterior Gateway Protocol</i>
17	<i>UDP</i>	<i>User Datagram Protocol</i>
29	<i>ISO-TP4</i>	<i>ISO Transport Protocol Clase 4</i>
38	<i>IDRP-CMTP</i>	<i>IDRP Control Message Transport Protocol</i>
80	<i>ISO-IP</i>	<i>ISO Internet Protocol (CLNP)</i>
88	<i>IGRP</i>	<i>Internet Gateway Routing Protocol (Cisco)</i>
89	<i>OSPF</i>	<i>Open Shortest Path First</i>
255		<i>Reservado</i>

Ejemplo de valores y significados del campo protocolo en un datagrama

- *Checksum*: Asegura la integridad de la cabecera de los datagramas, realizando el cálculo del checksum exclusivamente de la cabecera, y no de los datos. No es un CRC, sino el complemento a uno en 16 bits de la suma complemento a uno de toda la cabecera, tomada en campos de 16 bits.
- *Opciones*: Su tamaño varía desde 0 hasta 11 palabras de 32 bits. Tiene un formato predefinido, siendo el primer byte obligatorio para todas las opciones posibles.

0 8 16

Las principales opciones son:

- *Grabación de ruta* (código 7): consiste en grabar la dirección de cada una de las pasarelas que el datagrama atraviese. Para ello, además del código, se incluye un campo de longitud con el número de palabras de 32 bits reservadas para la información de ruta, y un puntero que indica cuáles son los primeros 4 octetos libres donde poder realizar la grabación.

0 8 16 24

Codigo (7)	Longitud	Puntero	Padding
1ª dirección IP			

- *Encaminamiento fuente* (código 3 ó 9): consiste en forzar la ruta a seguir, ya sea de forma estricta (especificando el camino exacto), o débil (permite rutas alternativas entre dos pasarelas).

0 8 16 24

Codigo (3 ó 9)	Longitud	Puntero	Padding
dirección IP del primer salto			

- *Timestamp* (código 4): Consiste en grabar tanto las direcciones de las pasarelas atravesadas, como la hora en la que fue procesado el datagrama.

3.3. PROTOCOLOS DE UTILIDAD A IP

El protocolo IP hace en mayor o menor medida uso de los siguientes protocolos:

3.3.1. PROTOCOLO DE RESOLUCIÓN DE DIRECCIONES (ARP)

El protocolo IP utiliza su propio esquema de direcciones lógicas, por lo que para poder entregar un datagrama en la red local, se debe proporcionar a la capa *MAC* la dirección física del host de destino. En concreto, *IP* llama a *ARP* utilizando la dirección *IP* del Host de destino, y *ARP* devuelve la dirección física correspondiente.

Para ello el protocolo *ARP* del host origen envía un marco de solicitud *ARP* mediante la dirección de difusión *FF:FF:FF:FF:FF:FF*, incluyendo la dirección *IP* y *MAC* del emisor, y la dirección *IP* del destinatario. El marco es recibido por todos los host, que comparan su propia dirección con la de destino. El host destinatario devuelve un marco de respuesta *ARP* conteniendo su dirección *IP* hacia la dirección *IP* de su par. Una vez recibida la respuesta se transfiere la información a *IP*.

Las direcciones resueltas vía *ARP* se mantienen un cierto tiempo en caché. En realidad cuando se manda una trama *ARP* todas las máquinas de la red, no sólo la destinataria, aprovechan este mensaje para 'fichar' al emisor, anotando en su caché la asociación dirección *IP*- dirección LAN. De esta forma si mas tarde necesitan contactar con dicha máquina podrán hacerlo directamente, sin necesidad de enviar antes una trama *ARP* broadcast.

Las entradas *ARP* caché expiran pasados unos minutos, para permitir que cambie la correspondencia dirección *MAC*-dirección *IP*, por ejemplo por avería de una interfaz LAN o cambio de la dirección *IP* de un host.

En caso de que más de un ordenador en una red tengan la misma dirección *IP* todos ellos responderán al mensaje *ARP* broadcast, lo cual puede causar serios problemas al emisor, quedando normalmente inaccesibles una o todas las máquinas con dirección duplicada (en algunos casos se han reportado incluso caídas del sistema operativo en grandes equipos por este motivo). Un caso típico de este problema se da cuando dos usuarios comparten un disquete con el software de red, en el cual están incluidos además del software los ficheros de configuración que contienen la dirección *IP*. En una red es fundamental establecer mecanismos que permitan minimizar el riesgo de tener direcciones de red duplicadas.

El mecanismo *ARP* se utiliza en todas las redes broadcast para descubrir el destinatario de los paquetes. En las redes basadas en conmutadores, como éstos son simplemente puentes

multipuerta, los paquetes broadcast llegan a todas las redes conectadas al conmutador, por lo que el ARP funciona exactamente igual.

3.3.2. PROTOCOLO DE RESOLUCIÓN INVERSA DE DIRECCIONES (RARP)

ARP utiliza como complemento un protocolo que permite que un host determine la dirección *IP* correspondiente a una dirección hardware, denominado *RARP*. Este protocolo es de gran utilidad en estaciones de trabajo sin disco que arrancan desde la red. En estos casos, la estación sin disco solo conoce su dirección *IP* a través de la red, para lo cual envía mediante difusión un paquete *RARP* de solicitud de dirección *IP* asociada a su dirección hardware. Dicho paquete es recibido por el servidor que responde con un paquete *RARP* de respuesta con la dirección *IP* correspondiente.

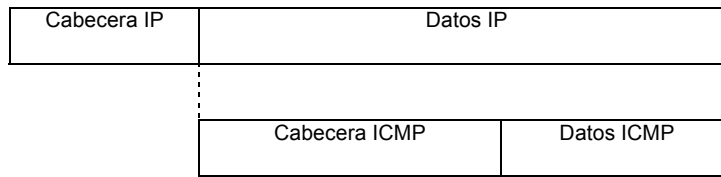
Como la estación solicitante emite una trama broadcast el servidor *RARP* puede estar en la misma LAN o en otras que estén conectadas con esta por puentes o conmutadores LAN, pero no puede haber routers entre ella y el servidor *RARP* (los routers filtran los paquetes broadcast de las LANs). Por otro lado, el protocolo *RARP* solo permite al servidor enviar la dirección *IP* del cliente en el paquete, cuando sería interesante aprovechar para enviar en el mismo datagrama una serie de parámetros de configuración a la estación que arranca. Para ofrecer estas facilidades adicionales se inventó *BOOTP* (*BOOT*strap Protocol), que tiene fundamentalmente las siguientes ventajas sobre *RARP*:

- El mensaje inicial se envía utilizando *UDP*; al ser este un protocolo de transporte reconocido por *IP* puede ser enviado a través de routers; esto permite mayor flexibilidad en la ubicación del servidor; en particular éste puede estar en una ubicación remota respecto al cliente; en la LAN del cliente debe haber al menos un retransmisor (relay) *Bootp* que se ocupe de redirigir el paquete *UDP* al servidor remoto (en este caso dicho reenvío se hace en modo unicast).
- El formato de un mensaje *BOOTP* permite enviar muchos parámetros *IP* al cliente, no únicamente la dirección *IP*. Entre estos se encuentran por ejemplo la máscara de subred, el *MTU*, rutas estáticas, el valor por defecto del parámetro *TTL*, etc.

3.3.3. PROTOCOLO DE MENSAJES DE CONTROL DE INTERNET (ICMP)

ICMP (*RFC 792*) es un protocolo estándar que dota a *IP* de capacidad de mensajería, y aunque se describe por separado, constituye una parte integral de *IP*, de tal forma que los mensajes *ICMP* son transportados en datagramas *IP*.

El paquete del protocolo *ICMP* se encapsula en el campo de datos del protocolo *IP*:



Los tipos de mensajes *ICMP* más comunes son:

- *Destino inalcanzable*: informan cuándo un host, una red, un puerto o un protocolo son inalcanzables. Dependiendo de su significado real podemos encontrar la siguiente lista de códigos:

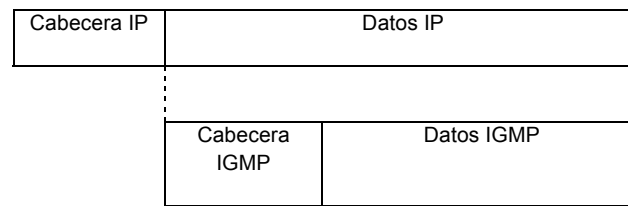
0	Red inalcanzable
1	Host inalcanzable
2	Protocolo inalcanzable
3	Puerto inalcanzable
4	Fragmentación necesaria y DF activo*
5	Fallo encaminamiento fuente

* Significa que el usuario activó el Flag de No Fragmentación y el mensaje llega a una pasarela que requiere de fragmentación.

- *Exceso de tiempo*: indican la imposibilidad de entregar un datagrama al haber expirado su tiempo de vida.
- *Problema de parámetro*: informan de problemas en octetos concretos del mensaje.
- *Eliminación de origen*: informan de situaciones de desbordamiento en los routers y hosts de destino.
- *Redirección*: Cuando un mensaje es enviado a un host por una ruta poco óptima, el encaminador puede informar al host de origen acerca del encaminador más adecuado para enviar el datagrama.
- *Mensajes de solicitud y de respuesta de eco*: son mensajes accionados por medio del PING.
- *Solicitud y respuesta de timestamp*: son mensajes para estimación de problemas y rendimientos, enviándose los tiempos de emisión, de recepción y de transmisión.
- *Control de flujo y congestión (Source Quench)*: Ante situaciones de congestión, el encaminador manda paquetes de este tipo hasta detener incluso al emisor.
- *Detección de bucles*: cuando se detectan paquetes cuyo tiempo de vida expira se avisa a los host de la posibilidad de que los datagramas estén en un bucle sin fin, en una red congestionada, o que su tiempo de vida sea muy pequeño.
- *Máscara de subred*: mensaje para solicitar la máscara de la subred *ICMP*.

3.3.4. PROTOCOLO DE GESTIÓN DE GRUPOS INTERNET (IGMP)

IP permite el multicast a través de múltiples redes, para lo cual el host debe informar a las pasarelas multicast locales. Las pasarelas locales contactan con otras pasarelas, pasando información de los miembros del grupo y estableciendo rutas para que cada miembro del grupo multicast reciba una copia de cada datagrama enviado al grupo. Los host comunican su pertenencia al grupo multicast utilizando *IGMP*, que se encapsula en el campo de datos del protocolo IP.



El protocolo *IGMP* tiene dos fases:

1. Cuando un host se une a un grupo multicast envía un mensaje *IGMP* a todos los hosts miembros. De igual forma, las pasarelas multicast envían información de los miembros del grupo al resto de pasarelas de internet. Para ello hacen uso del protocolo *DVRMP* (*Distance Vector Multicast Routing Protocol*), que es encapsulado dentro de la zona de datos de *IGMP*.
2. Las pasarelas locales preguntan periódicamente a los hosts de la red local acerca de su pertenencia al grupo multicast.

IGMP utiliza los recursos de la red de forma eficiente, ya que en la mayoría de los casos el único tráfico que incorpora es el mensaje periódico de la pasarela multicast al host, así como la respuesta de este acerca de la pertenencia al grupo.