

ESPECIALIDAD FORMATIVA

Operación en sistemas de comunicaciones de voz y datos

IFCM0110

UF1865: Operación y supervisión de los equipos y servicios.

El siguiente documento está creado con fines únicamente docentes y corresponde al registro diario de cada una de las jornadas de los cursos de formación impartidos por Luis Orlando Lázaro Medrano, y por lo tanto sólo se autoriza la lectura del mismo a los alumnos dados de alta en la plataforma denominada Portal del Alumno, cuyo acceso está restringido con nombre de usuario y contraseña. Y en ningún caso se autoriza la reproducción o difusión de este documento a terceros sin la aprobación expresa y por escrito de Luis Orlando Lázaro Medrano. El objetivo de este documento es únicamente ilustrar la actividad educativa en el aula, sin ninguna finalidad comercial, y siempre que sea posible, y la jornada educativa lo permita, se incluirá el nombre del autor y la fuente, adecuándose a los artículos 32.1 y 32.2. de la Ley de propiedad intelectual vigente en España.

Capturas de pantalla y textos electrónicos de varias web únicamente para ilustrar la actividad educativa

Contenido

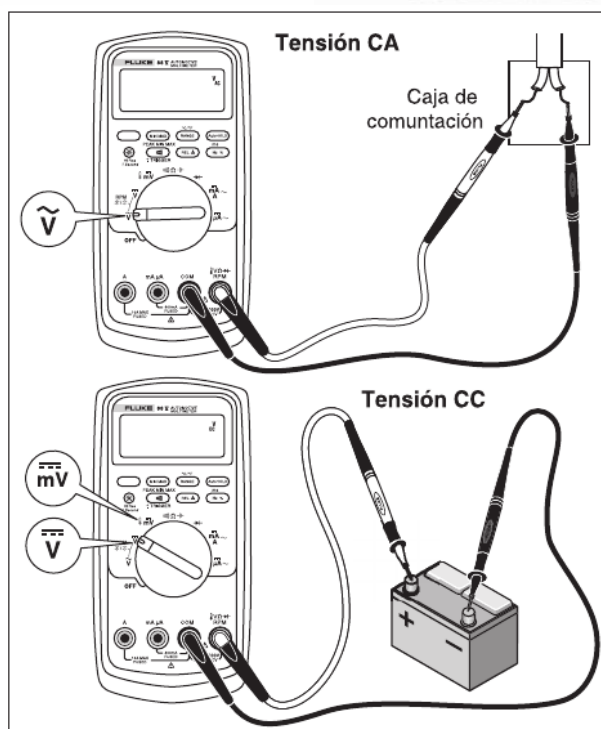
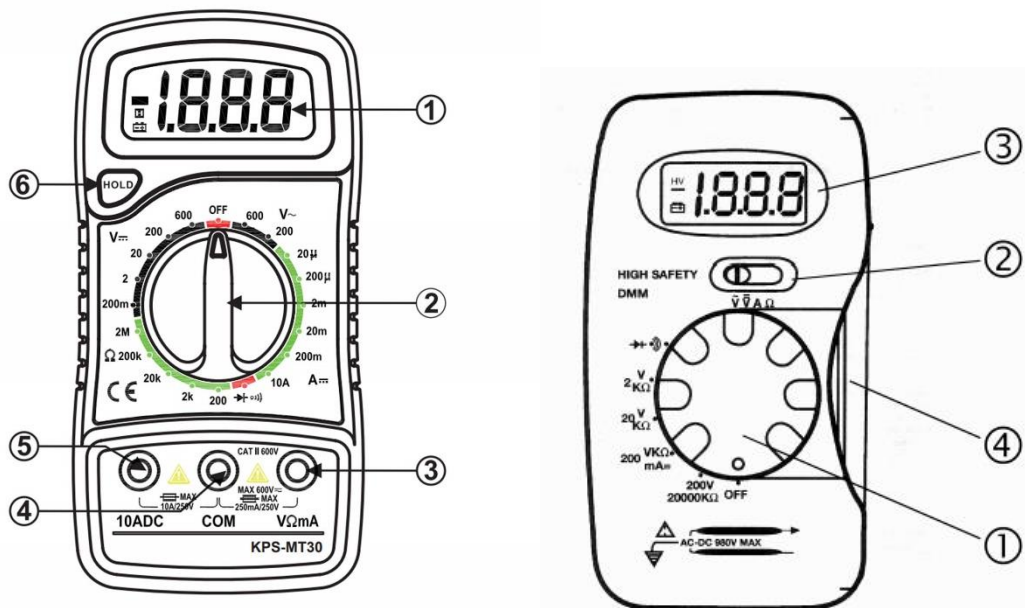
Multímetro	1
Medir continuidad	1
Medir Voltaje	2
Redes de comunicaciones de voz y datos	4
Elementos de una red de comunicaciones.....	4
Medios de transmisión	4
Tarjetas o Adaptadores de red	4
Equipos de interconexión.....	4
Niveles funcionales de una red de comunicaciones.....	4
Vídeos: Curso de Fibra óptica	5
Resumen de los Vídeos:.....	6
Video 5:.....	11
Empalmes:	11
Video 6:.....	12
Video 10:.....	16
Video 11:.....	17
Video 12:.....	21
Video 14:.....	24
Video 15:.....	25
Video 16:.....	26
Red de acceso	28
Red troncal de transporte	28
Red de distribución.....	28
Multiplexación.....	29
Multiplexores.....	29
Técnicas de multiplexación.....	30
Funciones de conmutación, transporte y señalización	31
Arquitectura de redes de voz y datos.....	34
Redes de acceso: guiadas y no guiadas	34
Redes de acceso vía cobre: xDSL	34
Redes de acceso vía radio: WLL, MMDS, LMDS	37
Redes de acceso vía fibra óptica: HFC, PON y CWDM	40
Redes troncales	45
MTA (Modo de transferencia asíncrono – ATM).....	46
JDP (Jerarquía Digital Plesiócrona – PDH)	47
JDS (Jerarquía Digital Síncrona – SDH).....	49
Mecanismos de codificación y cifrado de la información	50

Sistemas de seguridad en el transporte de datos	61
1. Incidencias en dispositivo de acceso a redes públicas	67
1.1 Incidencias habituales	67
1.1.1. Incidencias internas	69
1.1.2. Incidencias Externas (atribuibles al proveedor de servicios)	112
1.2. Gestión de incidencias en equipos de acceso a redes públicas	114
1.2.1. Sistemas de gestión/monitorización de equipos	115
1.2.2. Herramientas de gestión de incidencias	129
1.3. Herramientas de monitorización de equipos para la localización y notificación de incidencias	132
RESUMEN.....	136

Luis Orlando Lázaro Medrano

Luis Orlando Lázaro Medrano

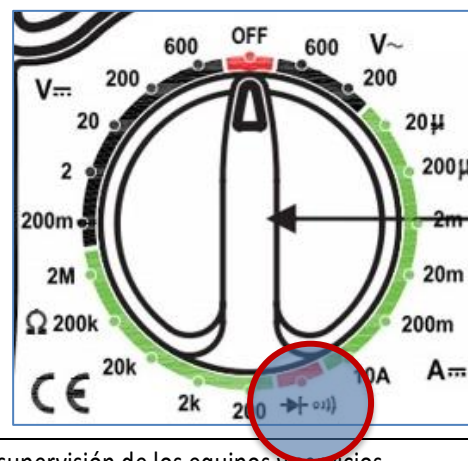
Multímetro



Medir continuidad

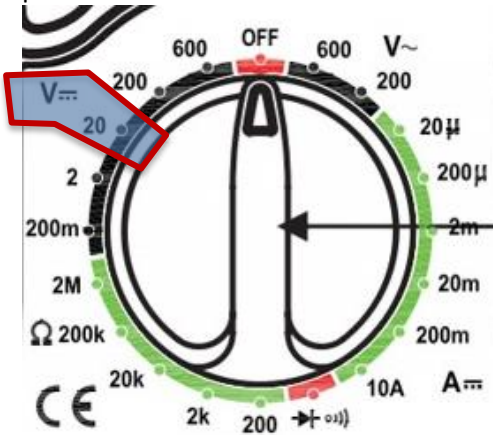
Para medir continuidad de cualquier cable: par trenzado, coaxial, alimentación....

Lo ponemos en el símbolo del diodo (componente electrónico que permite la circulación de la corriente eléctrica a través de él en un solo sentido) o uno similar a una wifi dentro del rango de medida en ohmios (Ω) (oposición al flujo de corriente en un circuito eléctrico)



Medir Voltaje

Si miramos el cargador (transformador de corriente alterna en continua) nos dice que tendrá entre 18 y 20V. Con lo cual ponemos la ruleta del multímetro en la V con la raya continua (V=) en la posición 20. Y tenemos que tocar con el electrodo negro la parte exterior del conector del transformador y con el rojo la parte interior.



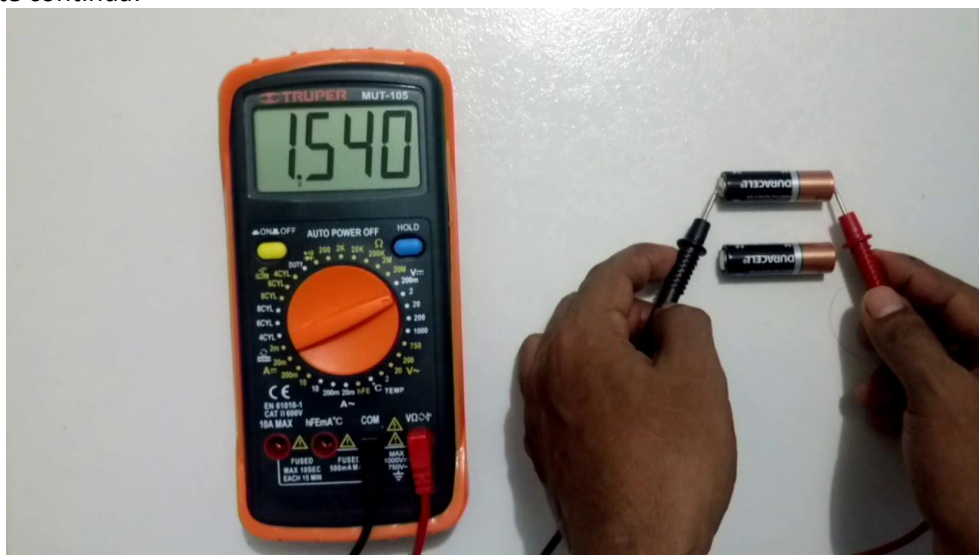
El De los móviles es de 5 V, así que lo dejamos en corriente continua (V=) en la posición 20, porque 2 es muy poco. Y tenemos que tocar los extremos del terminal (mini/micro USB,USB-C) del cargador con los electrodos, seguramente tendremos que ayudarnos de cualquier cable pelado para hacer buen contacto.



Para la batería de portátil, tendremos que localizar los polos positivo y negativo y mantendremos el Voltaje en Corriente continua y la escala en 20.



En cambio, para las pilas pequeñas: AA/AAA cuyo voltaje es de 1.5 V, tendremos que cambiar la escala a 2 de V corriente continua:



Y medir Voltaje en las regletas de Corriente Alterna cambiaremos a la V con la virgulilla ~ e insertaremos los electrodos en los 2 agujeros del enchufe (con cuidado) y lo pondremos en 600 ya que la corriente alterna en España es de 220V a una frecuencia de 50Hz

Redes de comunicaciones de voz y datos

Elementos de una red de comunicaciones.

Una red de comunicaciones está formada por un conjunto de elementos que permite la transmisión de información y datos entre diferentes equipos.

En este sentido, una red de comunicaciones fundamentalmente la forman los siguientes **elementos**:

Elementos	Funcionalidad
Equipos de interconexión	Son los elementos que permiten interconectar diferentes redes.
Medios de transmisión	Son los elementos que permiten conectar equipos de interconexión y equipos finales a una red.
Equipos finales	Son los elementos que generan o reciben los datos. Deberán estar equipados con los adaptadores de red necesarios.

Medios de transmisión

Visto en UF1863 (Página X)

Medios de transmisión guiados y no guiados

El par trenzado: UTP, STP, FTP

El cable coaxial: delgado y grueso, banda base y banda ancha

La fibra óptica: monomodo y multimodo

Medios no guiados: Transmisión por radiofrecuencia

Tarjetas o Adaptadores de red.

Visto en UF1863 (Página X)

Tarjetas de red para redes cableadas: RJ-45, BNC

Tarjetas de red para redes inalámbricas

Equipos de interconexión.

Visto en UF1863 (Página X)

Repetidores o Hub

Puente o switch

Punto de acceso

Encaminador o Router

Pasarelas o Gateways

Niveles funcionales de una red de comunicaciones

Para que los servicios de telecomunicaciones lleguen a los usuarios y equipos finales es necesaria una red que distribuya dichos servicios desde donde se generan hasta la ubicación donde están estos usuarios y equipos finales.

Así, las redes pueden ser más o menos grandes ya que implica su distribución por urbanizaciones, calles, edificios, hogares, etc.

Es por ello que la red presenta tres niveles funcionales.

Nivel de acceso

Es la parte de la red que hace que dichos servicios de telecomunicaciones lleguen hasta los hogares o sedes de empresas. Da el 'acceso' a estos usuarios.

Nivel troncal de transporte

Es la parte de la red que da servicio al nivel de acceso y constituye la parte de red que tiene la mayor capacidad de tráfico que luego se capilariza en los niveles de acceso. También es denominado habitualmente como **backbone**. Se encarga en definitiva de que los servicios puedan llegar a cualquier situación geográfica.

Nivel de distribución

Es la parte de la red que se encarga de la conmutación y multiplexación de la información que procede de los proveedores de servicio y adaptarla a las características de la red troncal de transporte.

Vídeos: Curso de Fibra óptica

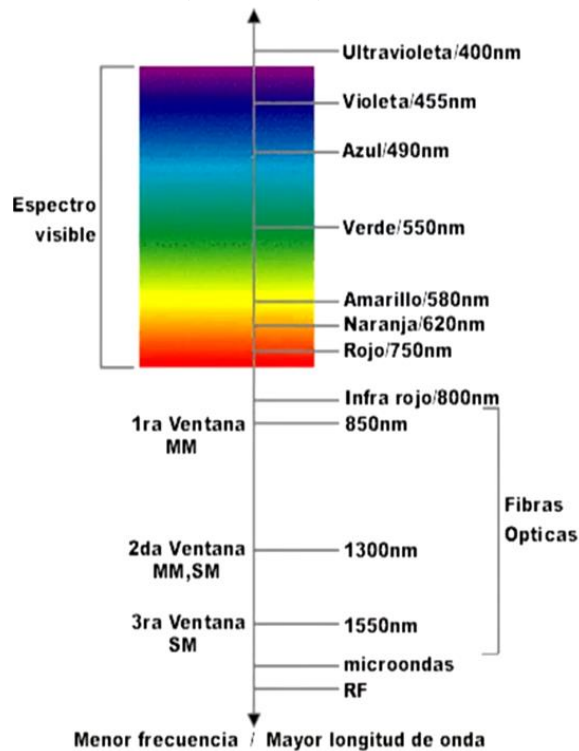
Vídeos muy interesantes sobre el funcionamiento de la fibra óptica y como llega hasta el usuario final. Estos vídeos nos sirven de introducción para entender los conceptos que tenemos que estudiar en los capítulos siguientes:

1. https://www.youtube.com/watch?v=M-PahKowgVM	Introducción
2. https://www.youtube.com/watch?v=WgHL3CWQRpc	En este vídeo sentaremos las bases para entender los siguientes capítulos del curso. Nos introduciremos en la estructura de red FTTH, hablaremos de sus elementos Splitter, ONT y OLT, y del lugar dónde se instalan.
3. https://www.youtube.com/watch?v=2nHnJI84EXU	Veremos una breve introducción de los protocolos que utilizamos al enviar y recibir la información en redes FTTH, así como las distintas longitudes de onda y su utilización.
4. https://www.youtube.com/watch?v=iIF-K8coKdc	En este vídeo, vamos a poder ver y visitar los diferentes elementos que tiene la operadora en central..OLT, ROM, galería de cables, etc.
5. https://www.youtube.com/watch?v=wO4KR6hU-oM	En este vídeo veremos cómo se llevan los cables de F.O. desde la sede de la operadora hasta el edificio del cliente, pasando por los diferentes elementos que nos podremos encontrar en el despliegue: cables, cajas de empalme, cámaras de registro, etc.
6. https://www.youtube.com/watch?v=kz_fGKZhpwl	Hoy veremos las cajas ópticas que se instalan en los edificios de cliente, en concreto, en fachadas y postes. Estas cajas son preconectorizadas y podrás ver la forma de conectar la acometida hasta el cliente.
7. https://www.youtube.com/watch?v=oPpSBwlxXDg	Cuando introducimos la acometida de fibra en cliente, podemos realizar una transición de cable, se denomina Roseta en paso. Presta atención a la colocación de la fibra en las Rosetas ya que es la parte más delicada de la instalación.
8. https://www.youtube.com/watch?v=-q-BmCwVLqE	Las cajas ópticas de exterior se instalan en distintos elementos de la Planta exterior, según la infraestructura del edificio del cliente. Podrás ver los entornos más comunes de instalación.
9. https://www.youtube.com/watch?v=I60GJQfdBdY	Actualmente, el proceso de fusiónado de fibra óptica es sencillo y rápido gracias a la evolución en las máquinas de empalme. Para una buena finalización, es imprescindible tener cuidado en la limpieza de la máquina de empalme y la fibra óptica. Te lo mostramos en detalle.
10. https://www.youtube.com/watch?v=EcG7mC7f83w	En este vídeo podrás ver una iniciación a los ajustes básicos de una máquina fusionadora y a los mantenimientos mínimos que puede requerir. No olvides leer los manuales del fabricante de la que utilices, te ahorrará mucho tiempo y alargará el tiempo útil de la herramienta.
11. https://www.youtube.com/watch?v=Dv2ePLjzpw	En este vídeo podremos ver, como es el despliegue dentro de un edificio y la forma de actuar en cajas preconectorizadas para llegar al domicilio del cliente.
12. https://www.youtube.com/watch?v=2fOiOY6088w	En este vídeo, seguimos en el edificio de cliente, pero vamos a ver un tipo de despliegue diferente. Son cajas sin preconectorizar, donde las fibras tendremos que fusionarlas, en todos los puntos del despliegue de cliente. Concretamente vamos a ver el manejo de las fibras del divisor y del "Riser", en la Caja Terminal Óptica.
13. https://www.youtube.com/watch?v=RZHdp18jlgw	En este vídeo terminaremos la instalación empezada en el vídeo anterior. Concretamente veremos la caja de derivación, que se instala en las diferentes plantas del edificio, para poder dar el servicio de fibra de una manera más sencilla para el instalador, acercando la fibra, al domicilio de cliente.
14. https://www.youtube.com/watch?v=BBkPLBtNTTI	En este vídeo, podremos ver la instalación de equipos en el domicilio de cliente y las diferentes opciones que nos pueden ofrecer las operadoras.
15. https://www.youtube.com/watch?v=WLiLwAw6Cpk	En este vídeo podremos ver la importancia de la limpieza en los conectores de fibra óptica y las diferencias con las redes fusionadas.
16. https://www.youtube.com/watch?v=kspS_XyDqws	Para la detección de incidencias en la fibra óptica, tenemos tres herramientas que nos facilitan tanto la instalación como su posterior mantenimiento. Veremos una breve introducción y su funcionamiento.

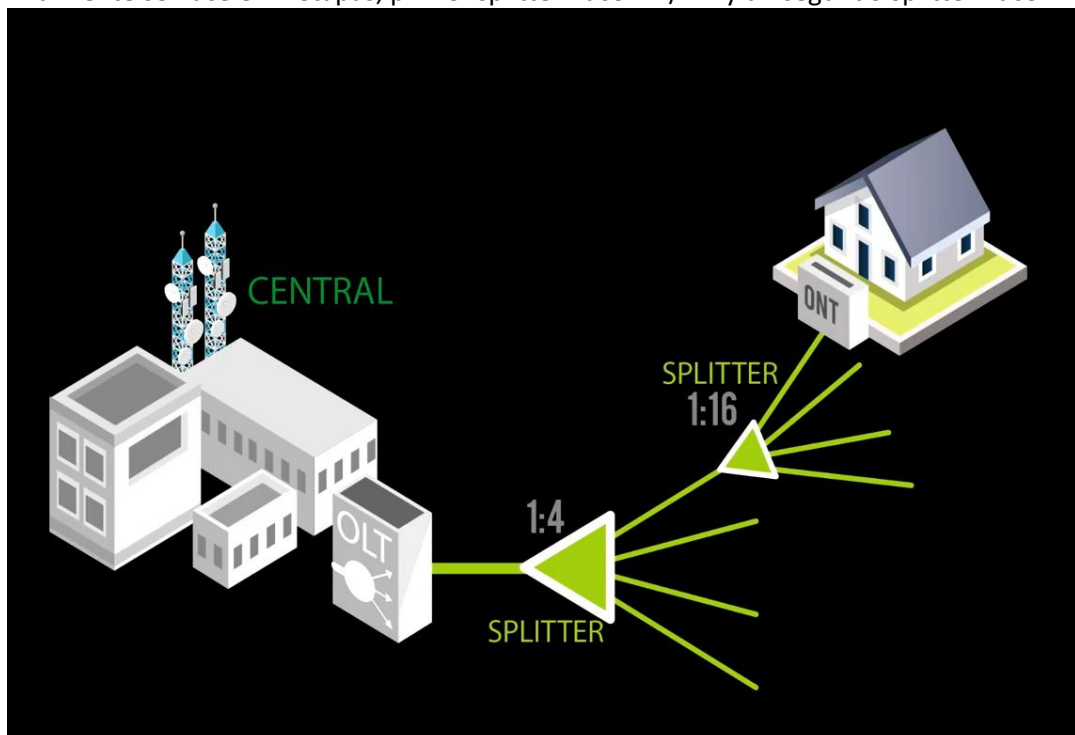
Resumen de los Vídeos:

Queremos ver como esta montada la estructura de la red, desde la distribuidora hasta los hogares:

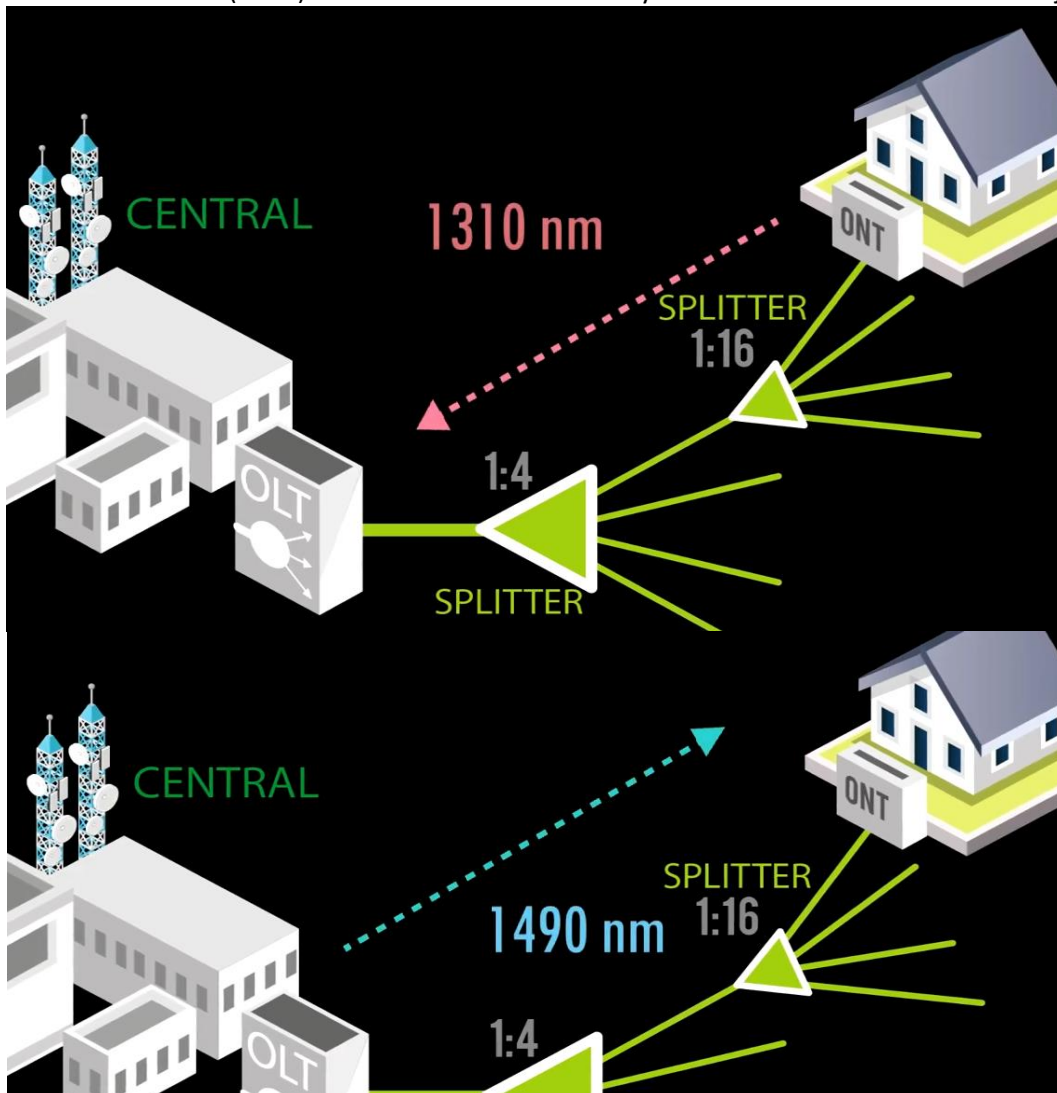
- ⇒ Red FTTH (Fiber To The Home)
- ⇒ 3 ventanas de longitudes de honda (no visibles): 850nm, 1300nm, 1550nm



- ⇒ GPON: Gigabit Passive Optical Network
- ⇒ Punto/Multipunto: De una fibra que sale de central llega a varias fibras de clientes
- ⇒ Fibra bidireccional: con una sola fibra hace transmisión y recepción
- ⇒ ONT: Convierte la señal de óptico a eléctrico en el domicilio del cliente (activo)
- ⇒ OLT: Dispositivo de oficina central
- ⇒ Splitter: Divisor Optico (Pasivo) El laser llega por una de las entradas y sale por varias salidas.
- ⇒ Normalmente se hace en 2 etapas, primer splitter hace 1:2/1:4 y un segundo splitter hacer 1:8/1:16

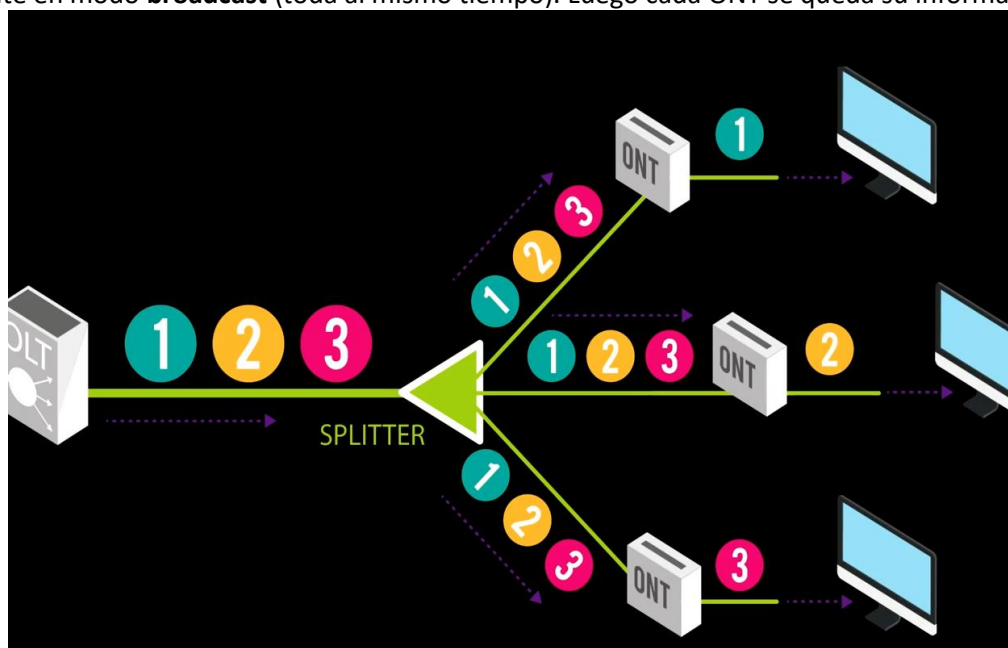


Señal **bidireccional**: onda (onda) de 1490 de la OLT a la ONT y 1310nm de la ONT a la OLT de bajada:

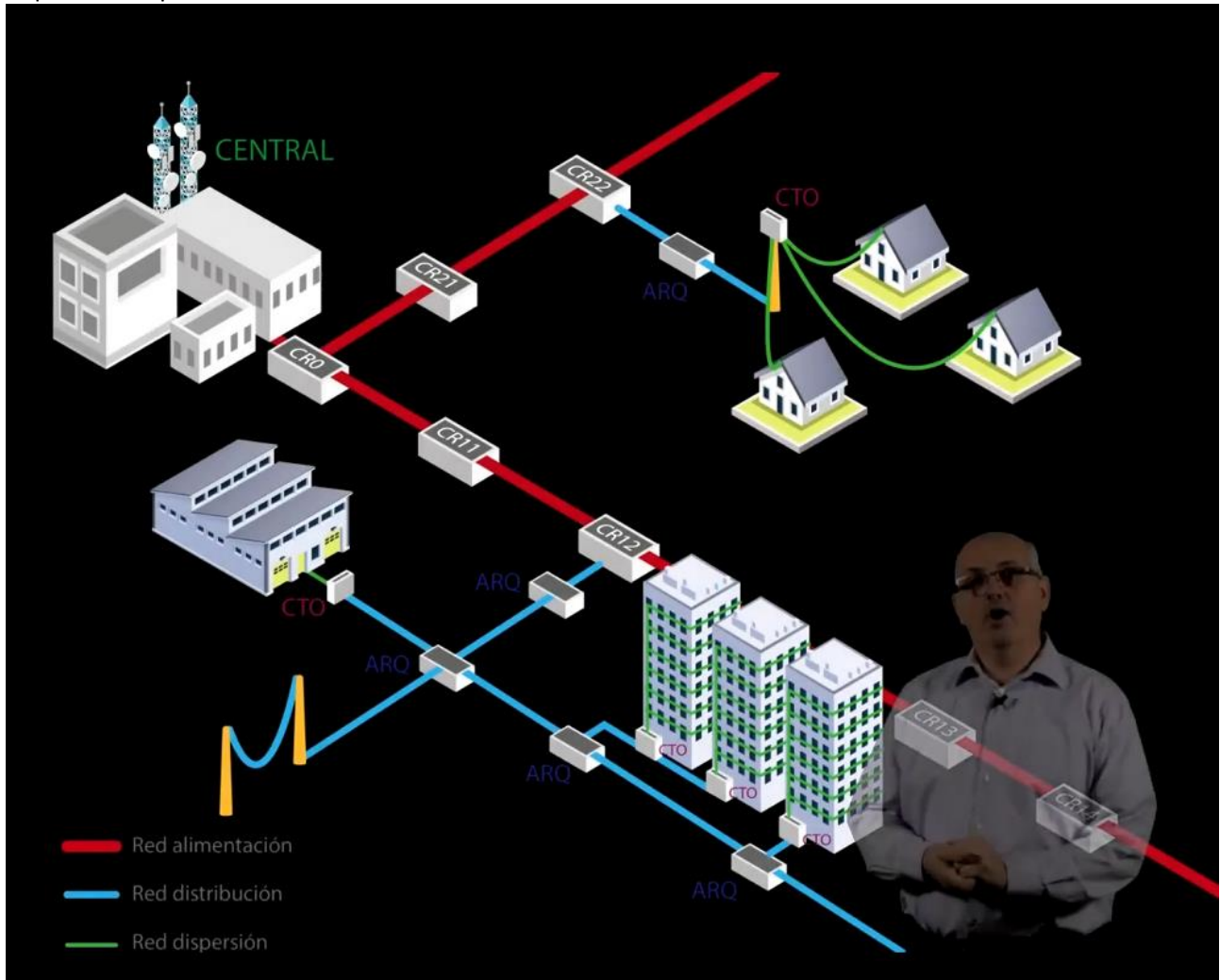


Mediante login y password, número de serie u otro sistema de autenticación se autentica y se le asigna un SLOT temporal. Que se le asigna por TDMA (Time Division Multiple Access).

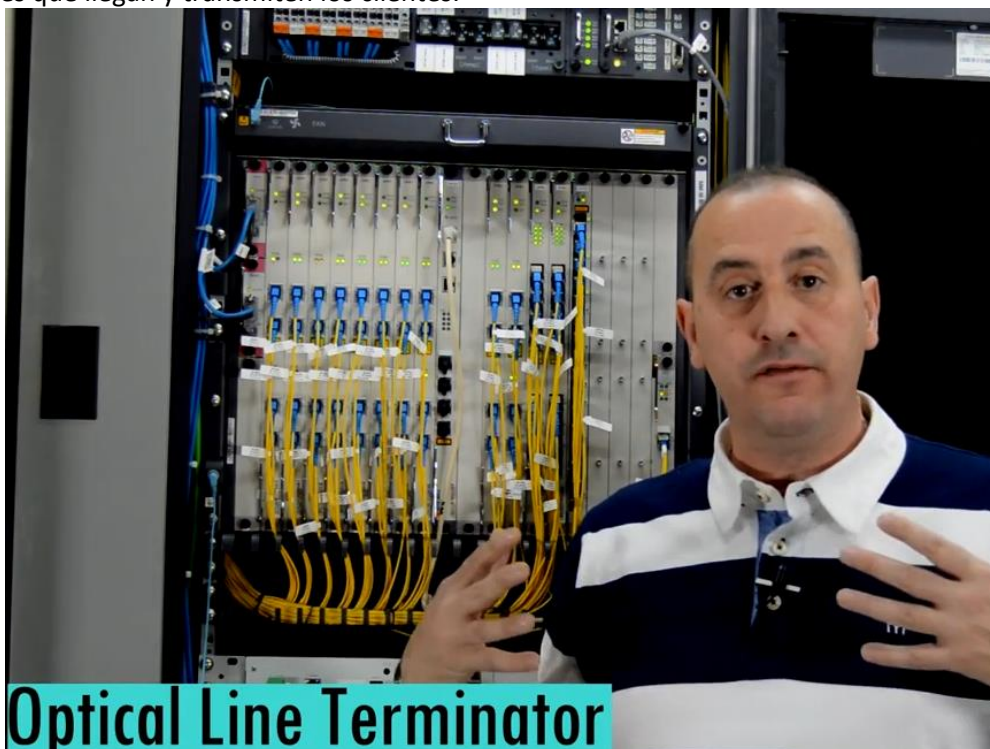
Y transmite en modo **broadcast** (toda al mismo tiempo). Luego cada ONT se queda su información:



Esquema completo:



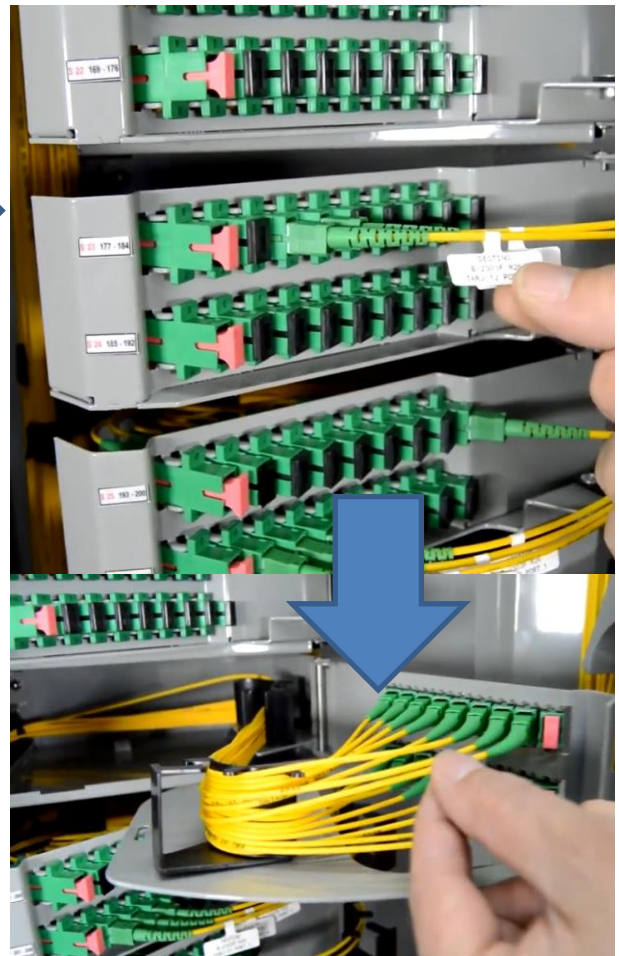
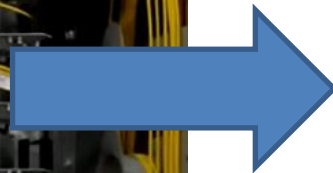
En la central nos encontramos con la OLT (Optical Line Terminator), concentrador que llegan todas las informaciones que llegan y transmiten los clientes.

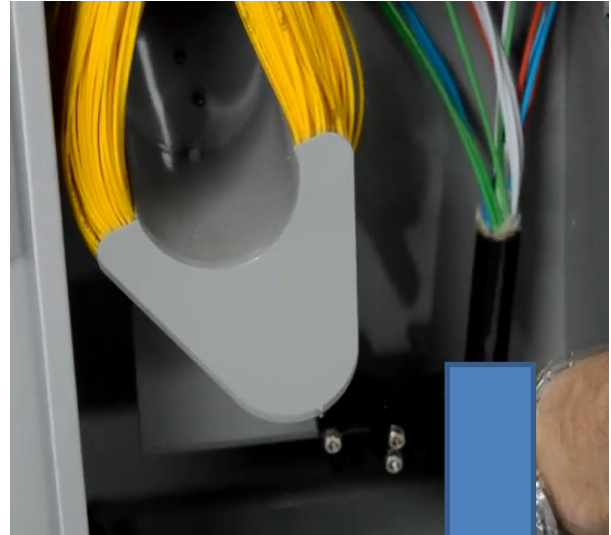


A cada OLT llega una fibra (bidireccional) y luego de todas las OLT juntas salen 2 fibras una de transmisión y otra de recepción)



ROM: Repartidor Optico Modular. Llegan las fibras de la OLT o de otros equipos de transmisión con las fibras de alimentación (las que reparten por la ciudad)

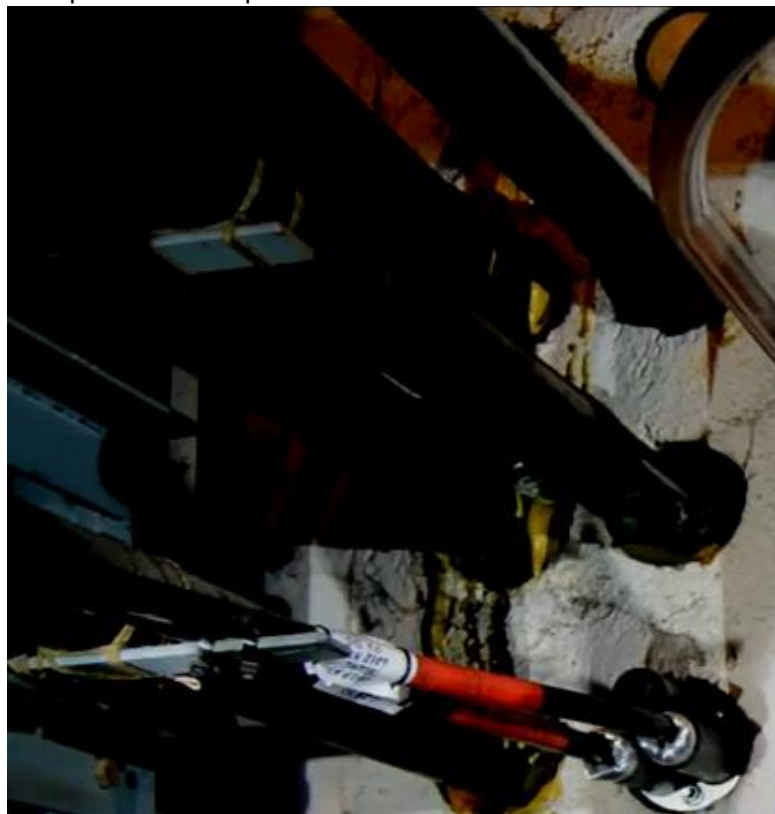




Del ROM baja a la sala de cables:

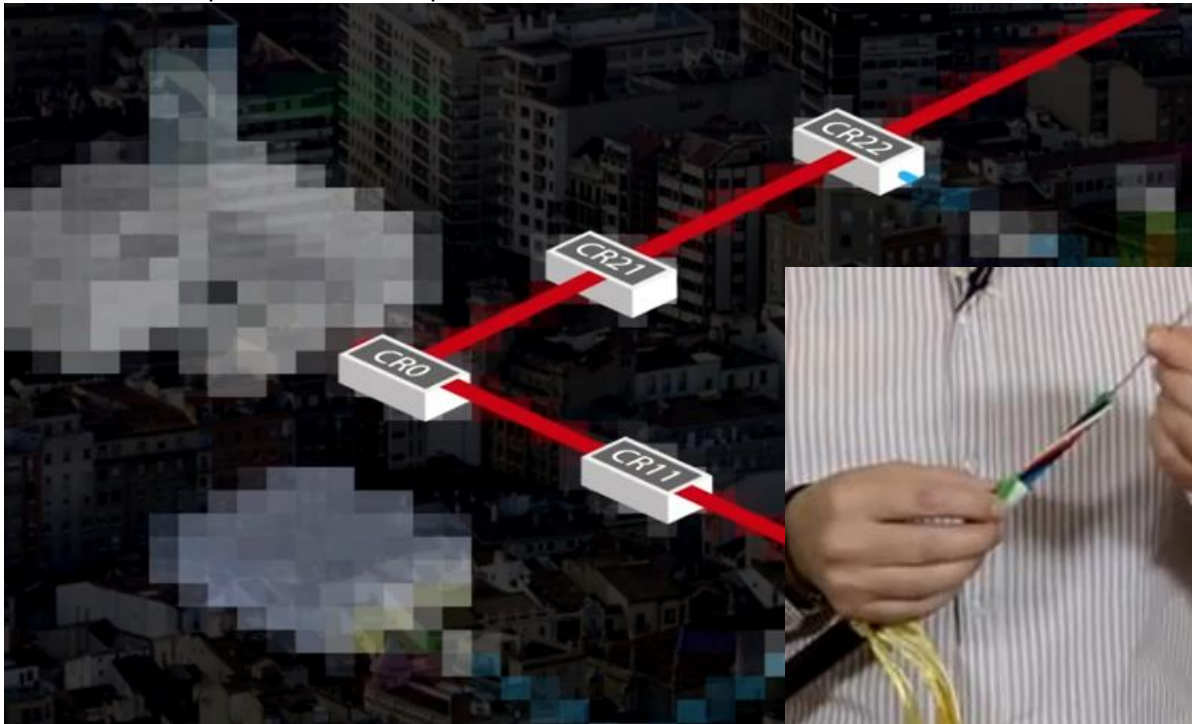


Los cables van de la OLT al ROM y luego bajan a la sala de cables van a los conductos que son donde salen los cables hacia el exterior para distribuir por la ciudad:

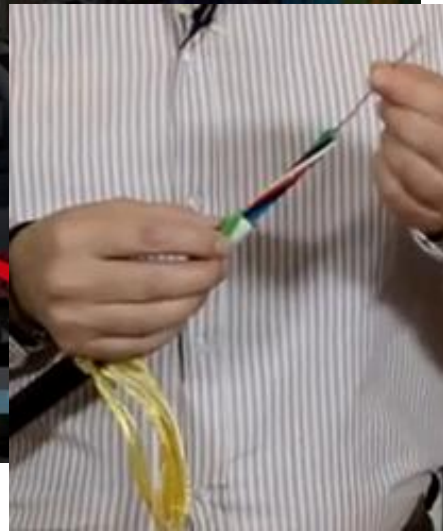


Video 5:

Salimos de las dependencias de la operadora a la calle:

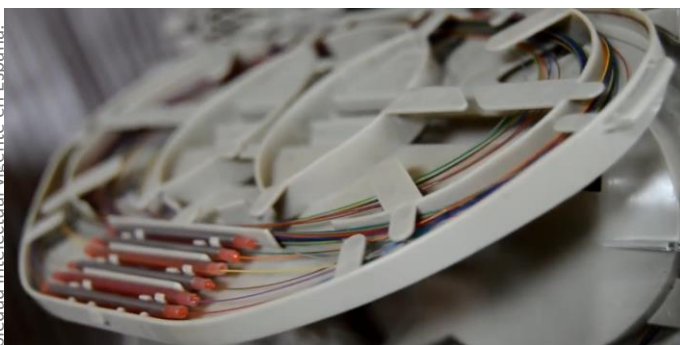


Con cables que pueden llevar varias fibras: 256...



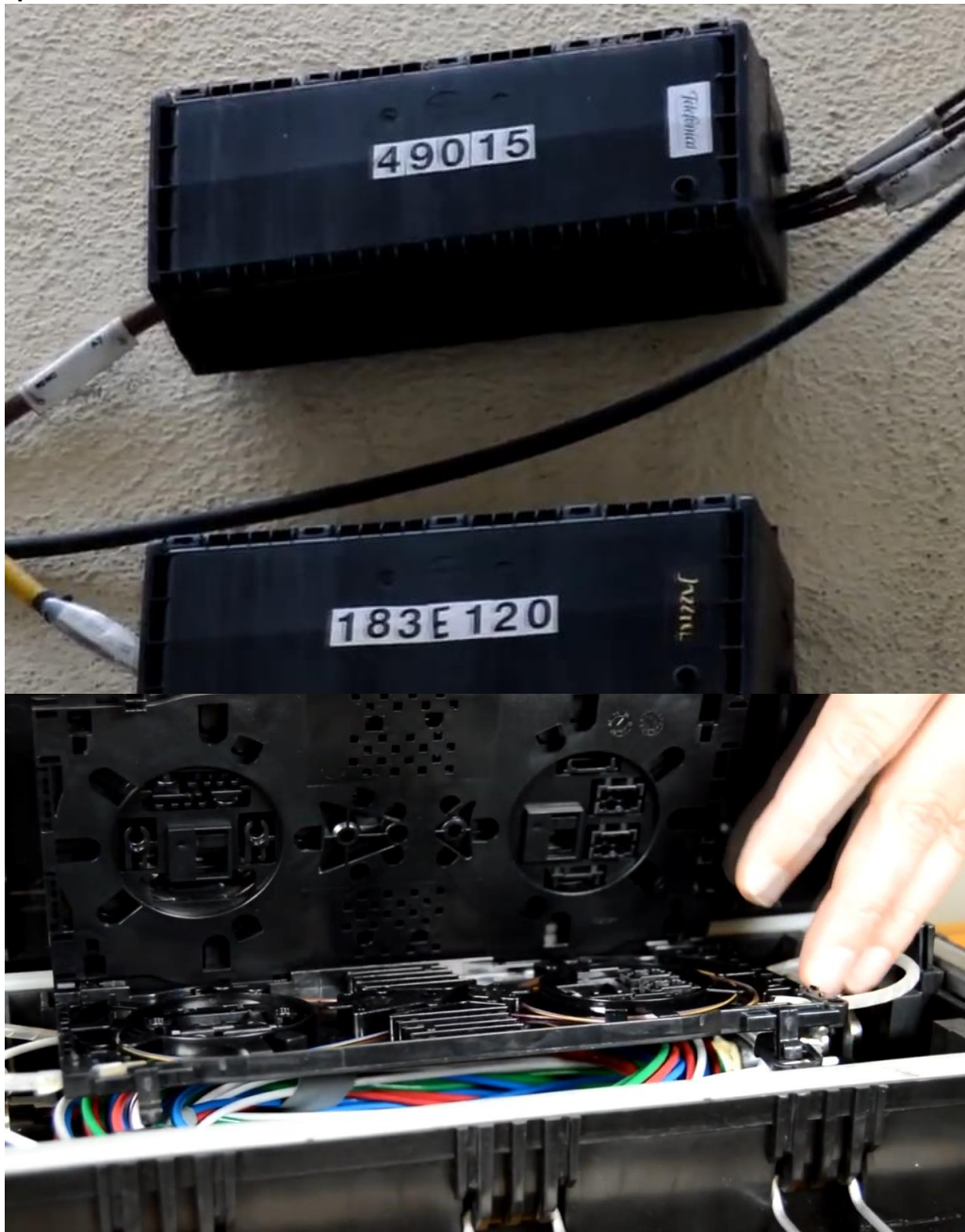
Empalmes:

Se pueden hacer en el **Subsuelo** usando Cajas de empalme ubicadas en las Cajas de Registro (CR0...):



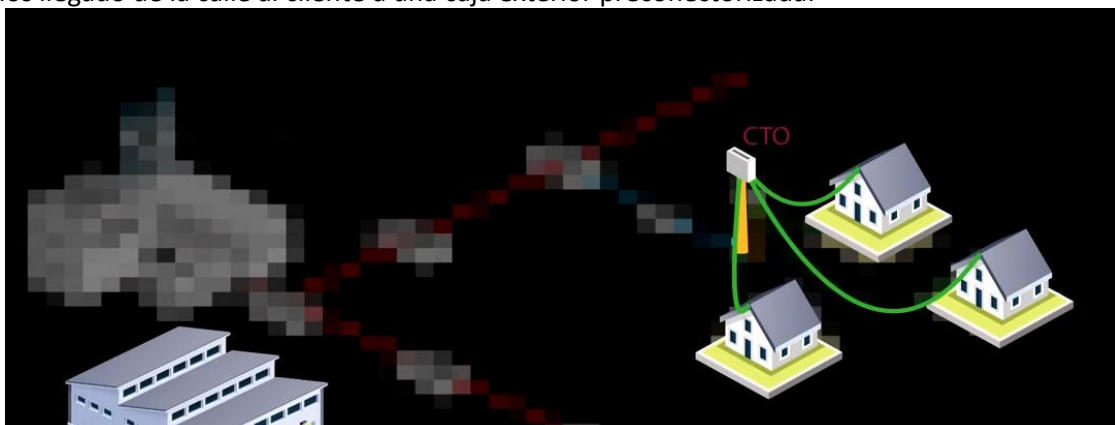
El siguiente documento está creado con fines únicamente docentes y corresponde al registro diario de cada una de las jornadas de los cursos de formación impartidos por Luis Orlando Lázaro Medrano, y por lo tanto sólo se autoriza la lectura del mismo a los alumnos dados de alta en las plataformas de formación, cuyo acceso está restringido con nombre de usuario y contraseña. Y en ningún caso se autoriza la reproducción o difusión de este documento a terceros sin la aprobación expresa y por escrito de Luis Orlando Lázaro Medrano. El objetivo de este documento es únicamente ilustrar la actividad educativa en el aula, sin ninguna finalidad comercial, y siempre que sea posible, y la jornada educativa lo permita, se incluirá el nombre del autor y la fuente, adecuándose a los artículos 32.1 y 32.2. de la Ley de propiedad intelectual vigente en España.

De fachada-poste



Video 6:

Ya hemos llegado de la calle al cliente a una caja exterior preconectorizada:



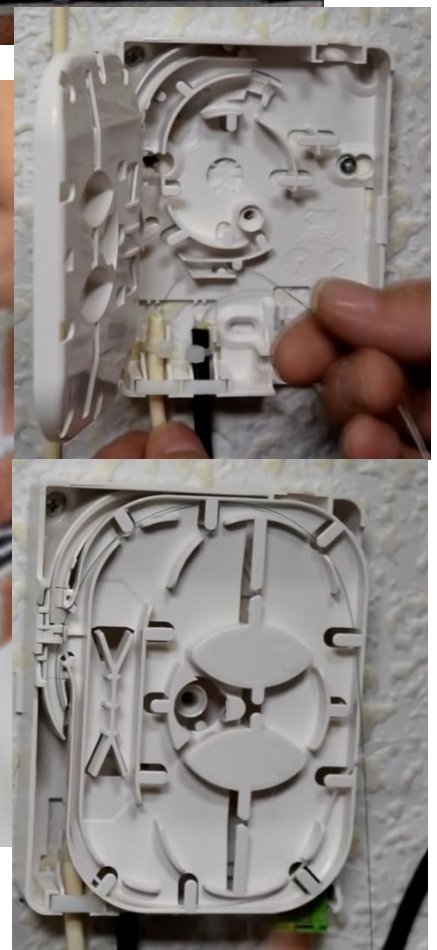
CTO: Caja Terminal Óptica (La caja negra de la imagen)



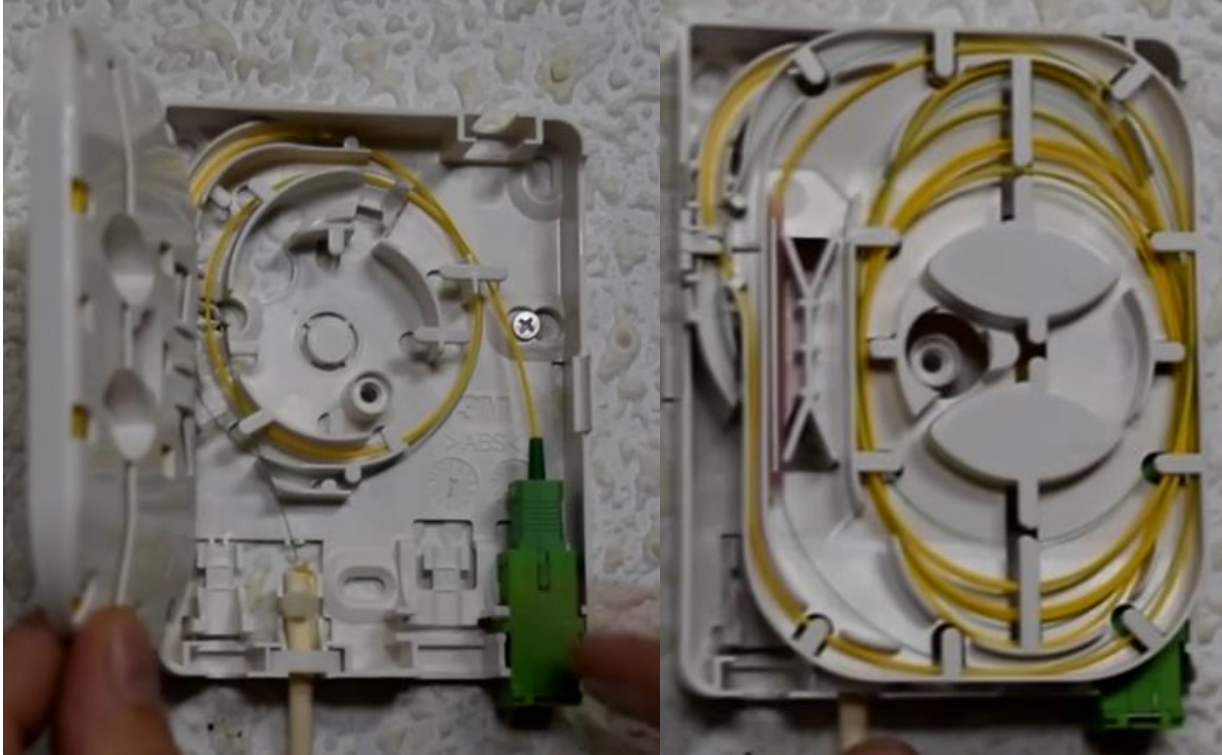
Llegamos hasta la Roseta en paso



Se Fusionan la fibra negra y la blanca



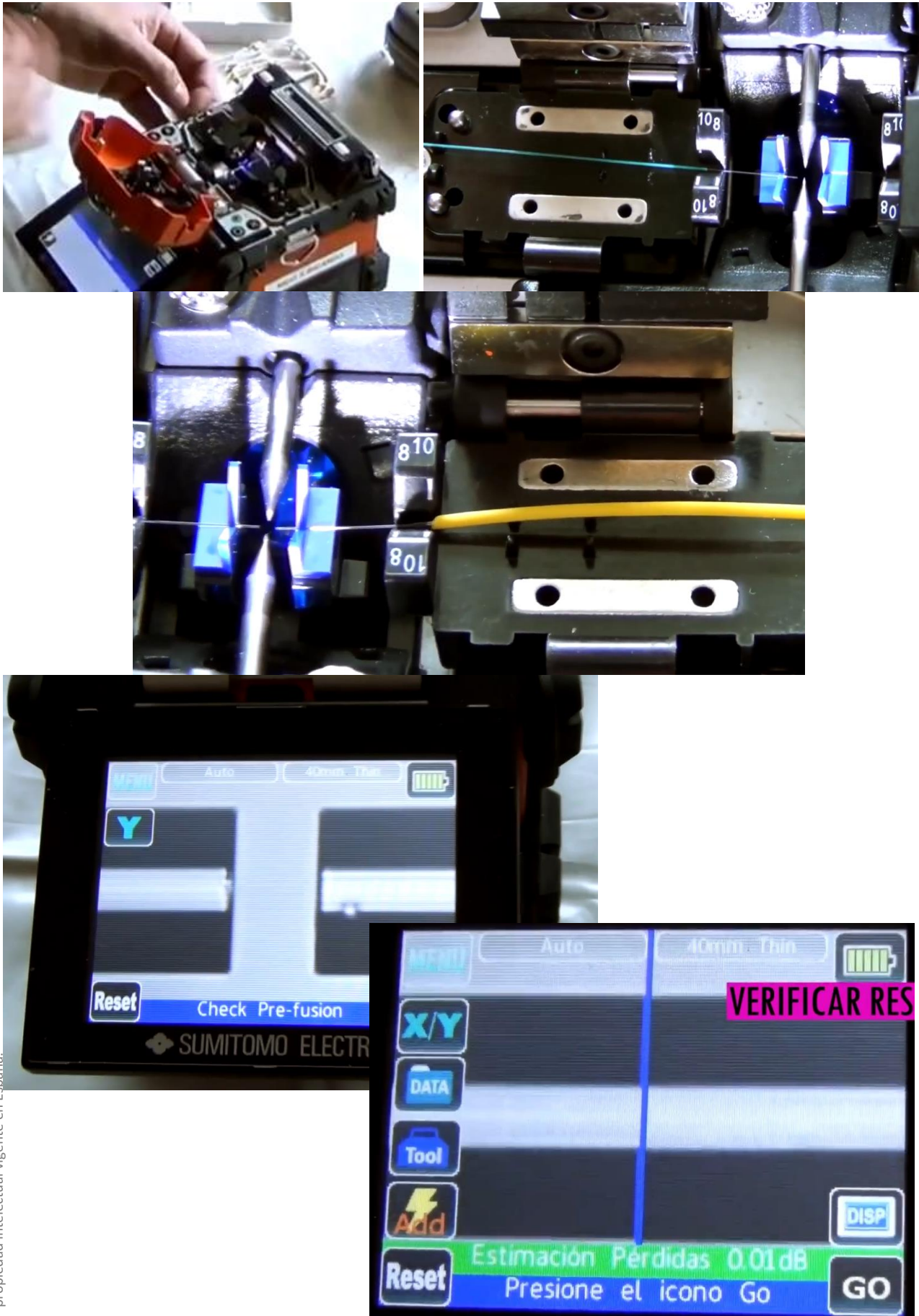
Roseta final: fusión de la fibra blanca al latiguillo (pig-tail)



En las viviendas unifamiliares también se puede llegar con una Caja de Empalme en pedestal:



Video 9: fusión en el domicilio del cliente



El siguiente documento está creado con fines únicamente docentes y corresponde al registro diario de cada una de las jornadas de los cursos de formación impartidos por Luis Orlando Lázaro Medrano, y por lo tanto sólo se autoriza la lectura del mismo a los alumnos dados de alta en las plataformas de formación, cuyo acceso está restringido con nombre de usuario y contraseña. Y en ningún caso se autoriza la reproducción o difusión de este documento a terceros sin la aprobación expresa y por escrito de Luis Orlando Lázaro Medrano. El objetivo de este documento es únicamente ilustrar la actividad educativa en el aula, sin ninguna finalidad comercial, y siempre que sea posible, y la jornada educativa lo permita, se incluirá el nombre del autor y la fuente, adecuándose a los artículos 32.1 y 32.2. de la Ley de propiedad intelectual vigente en España.

Luego le ponemos la funda protectora, la cual metemos al horno:



Video 10:
Configuración de la fusionadora: Tipo de Fibra:



Tipo de funda protectora-tipo de calefactor:



Video 11:

Instalación en edificios modernos -> RITI: Armario de telecomunicaciones del edificio



Cable de la operadora: indica las fibras en este caso 16



El cable de cliente es el mas pequeño que se ve en la imagen inferior.
Abrimos el módulo del operador:



Y abrimos el módulo del cliente:



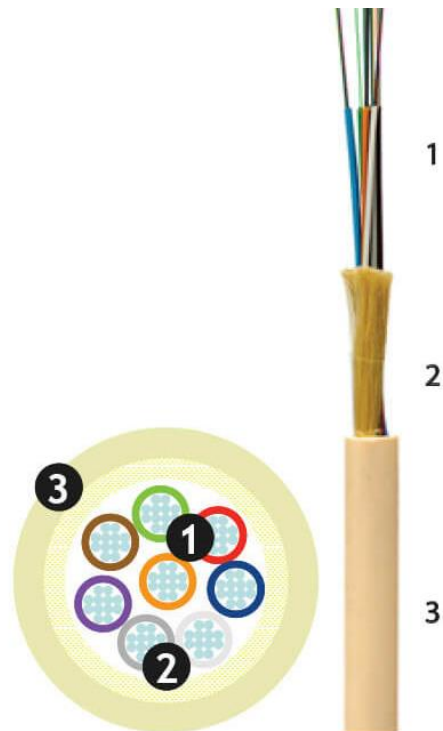
El cable que se usa se llama Cable RISER:

Cable RISER

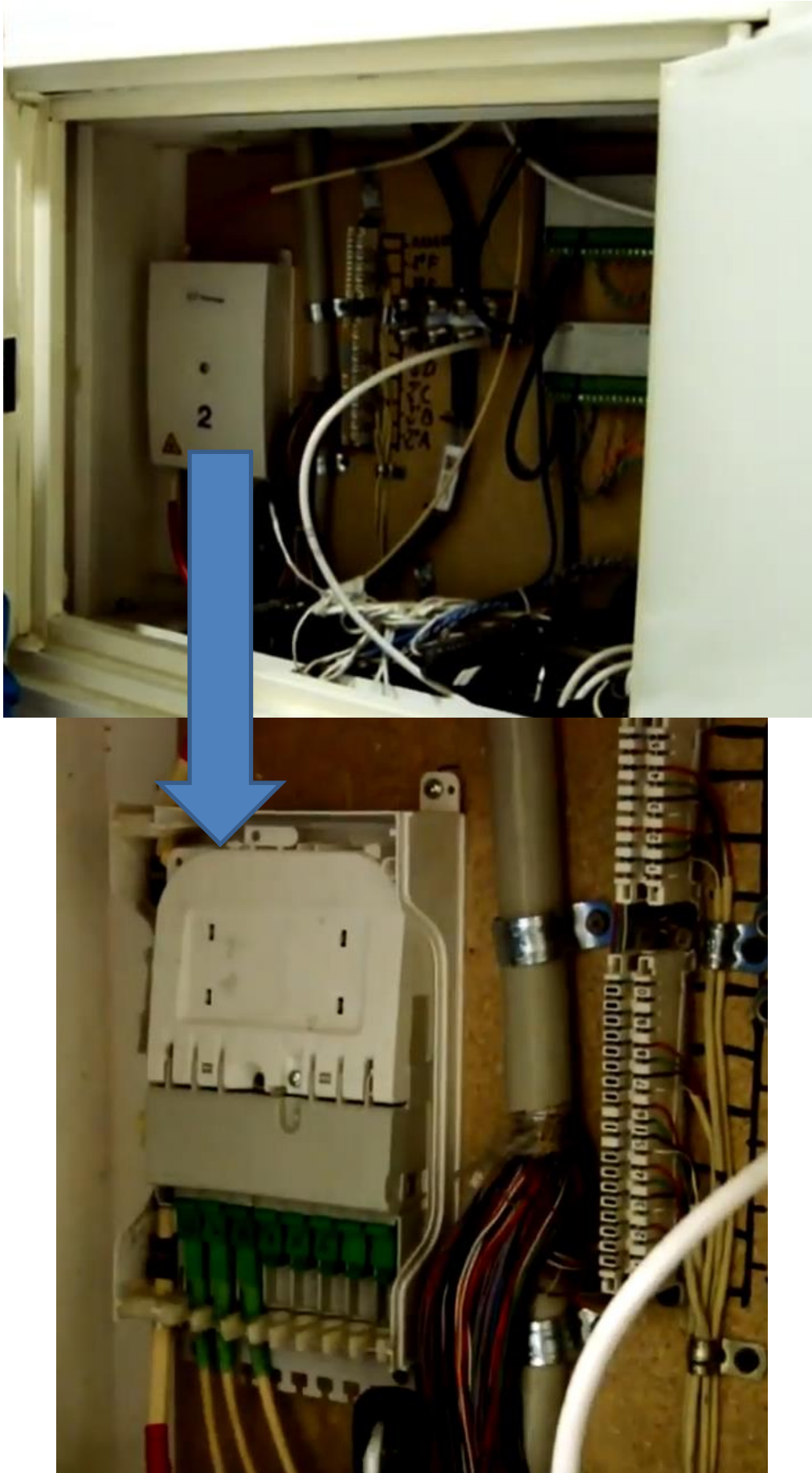
16 FIBRAS (multitubo 4x4)

Detalle de construcción

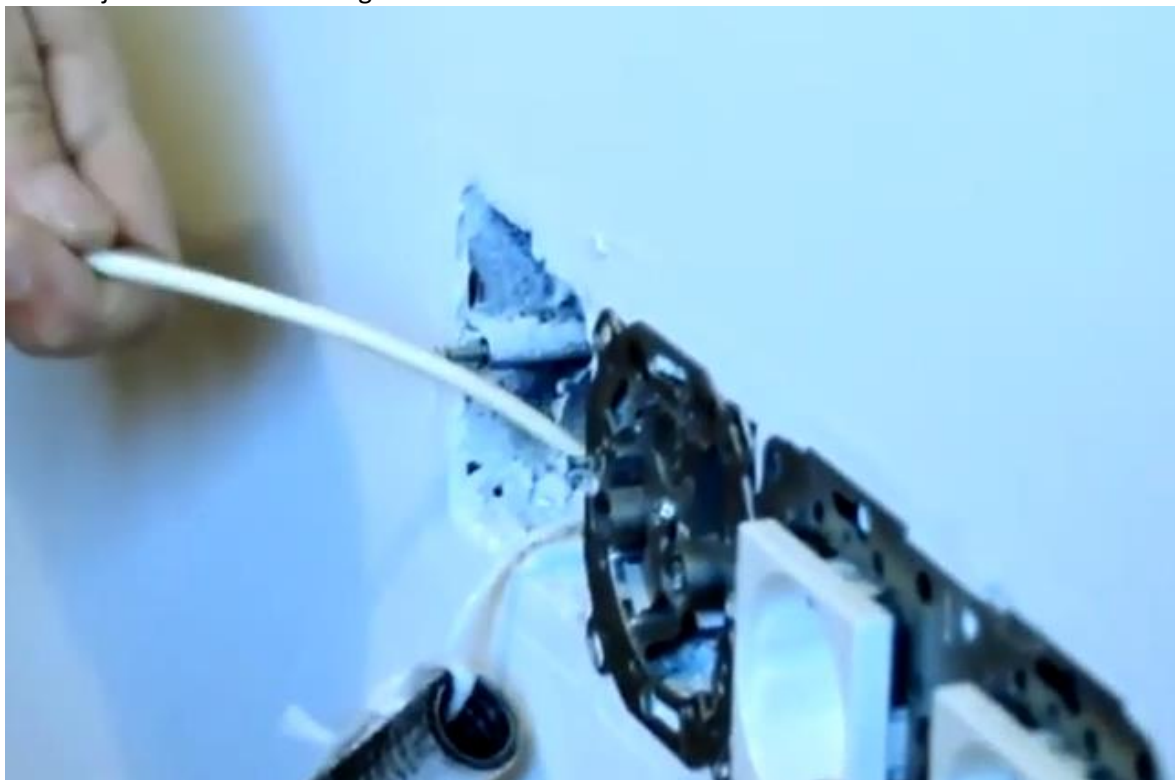
1. Micromódulos que contienen fibras ópticas
2. Cabos de aramida como elemento de refuerzo a la tracción
3. Cubierta de termoplástico retardante de llama, de baja emisión de humos y cero halógenos (LSZH)



Y desde el RITI llegamos a la caja de Derivación en la planta del cliente



Y desde la caja de derivación se llega a los enchufes de la casa del cliente:

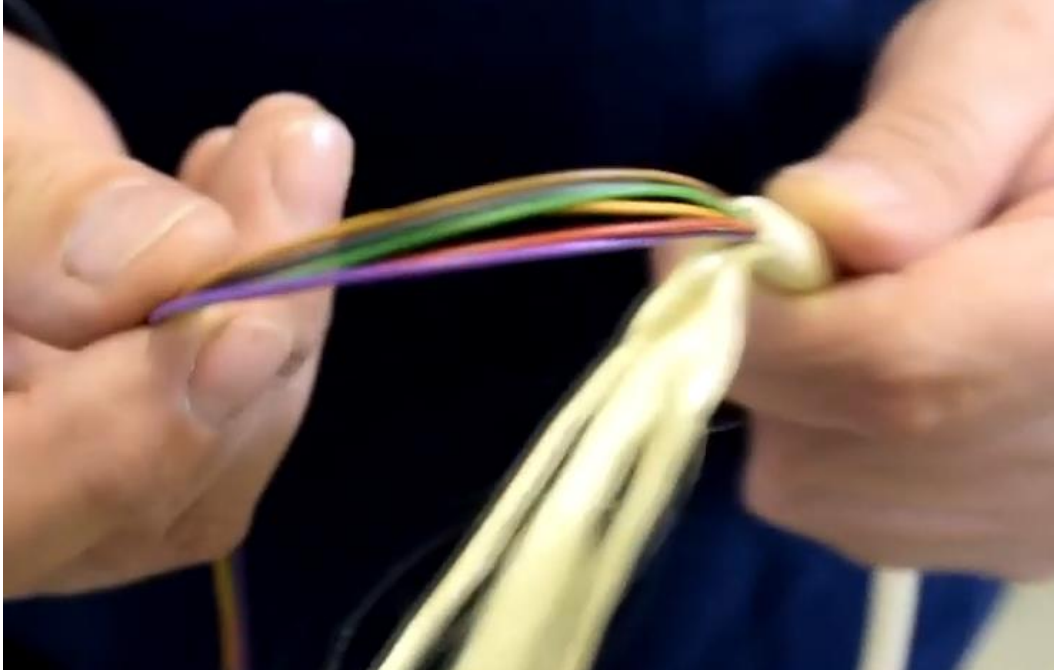


Video 12:

Cajas de interior fusionadas, dónde no hay cuarto para el RITI (por ejemplo en garajes)



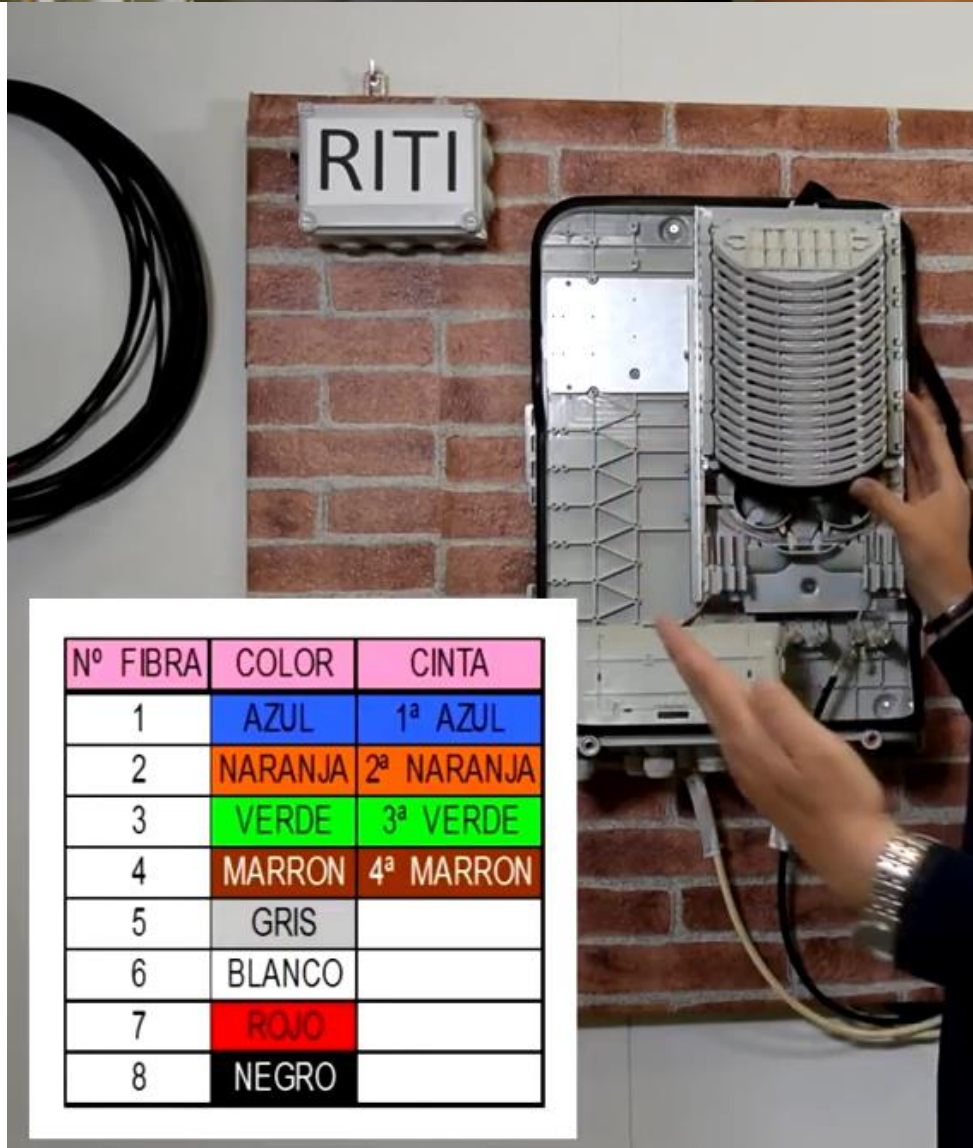
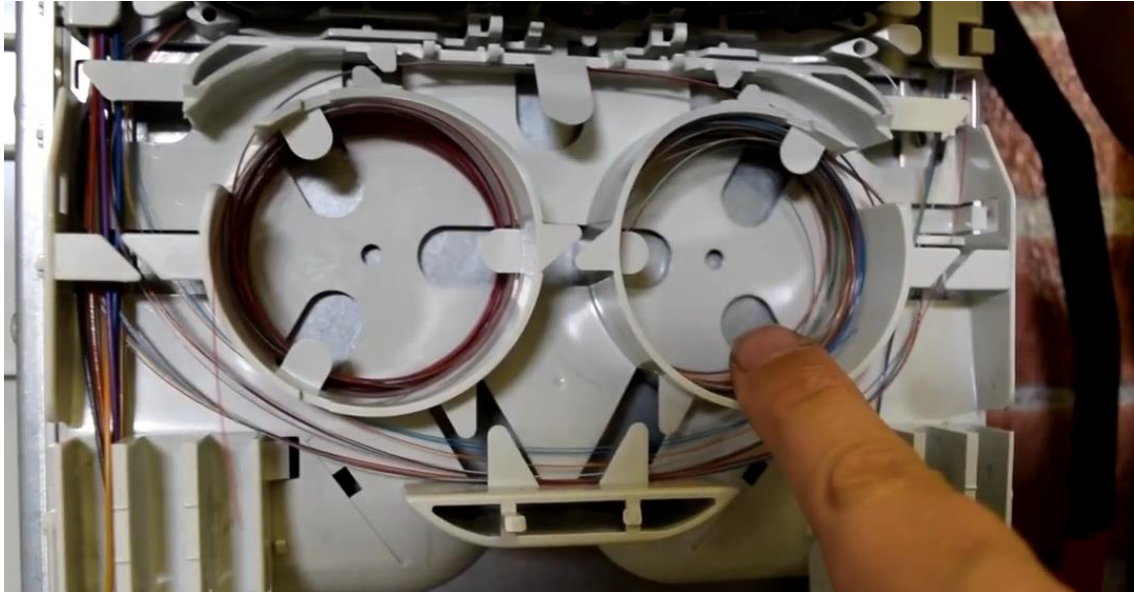
El Riser que llega al módulo del operador:



Usando el separador (negro) dividimos el RISER en:

Cinta de color azul del 1-8 y la cinta de color naranja del 9-16

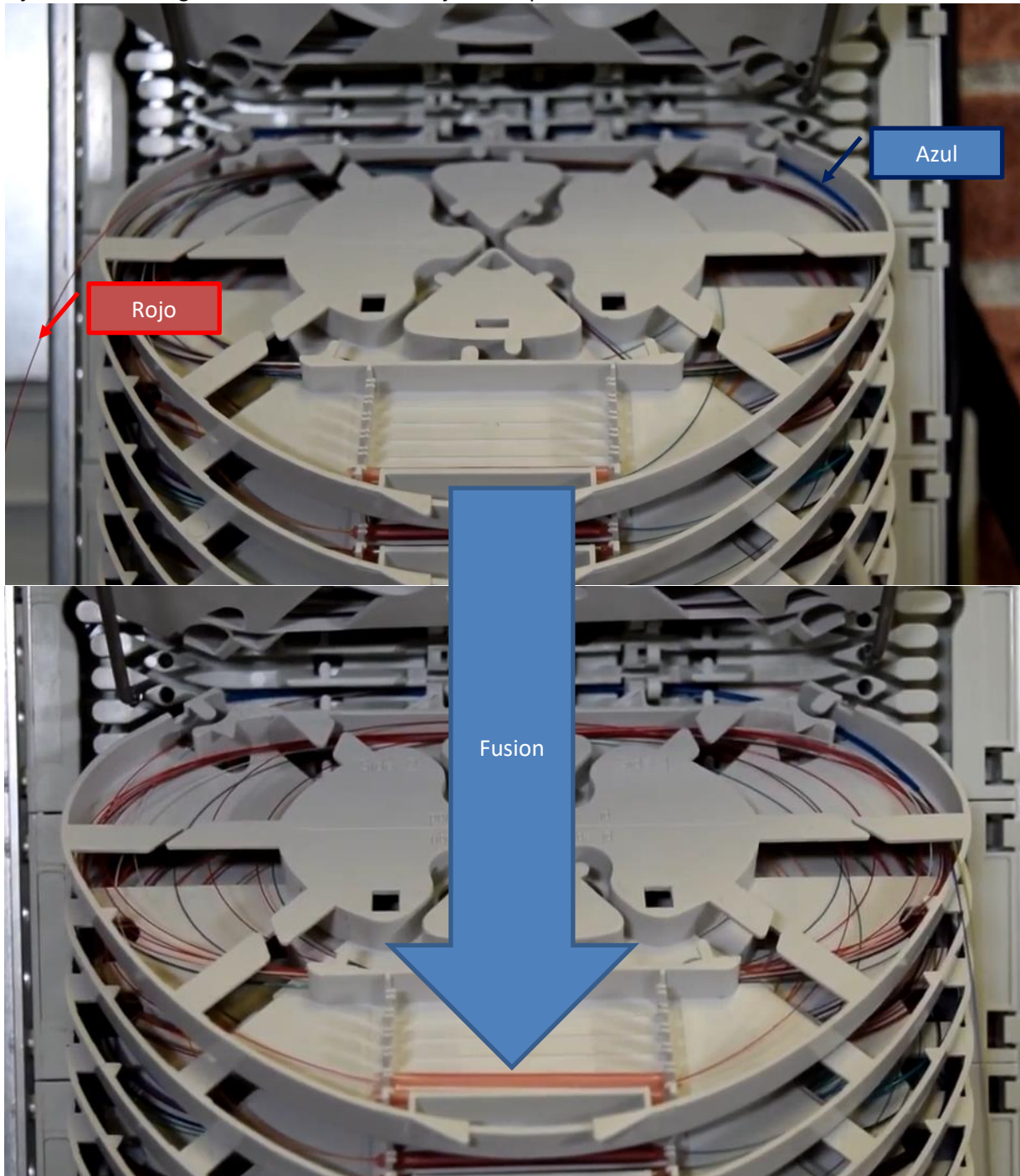
Y luego conectamos en la bandeja que llega hasta el cliente:



Nº FIBRA	COLOR	CINTA
1	AZUL	1ª AZUL
2	NARANJA	2ª NARANJA
3	VERDE	3ª VERDE
4	MARRON	4ª MARRON
5	GRIS	
6	BLANCO	
7	ROJO	
8	NEGRO	

Seleccionamos el color en el divisor (bandeja negra) por ejemplo el Rojo y luego en la toma del cliente, por ejemplo en el piso 4º miramos el color: en este caso azul. Buscamos ese color en las bandejas de la caja de la operadora y solo faltaría fusionarlas.

La roja del divisor negro con el azul de la bandeja correspondiente:



Vídeo 14:

Instalación de equipos en la casa del cliente: ONT, ONT+Router y Equipo integrado. Ejemplos:

- ⇒ Solo ONT: (Optical Node Terminal) Conversor óptico – eléctrico y entonces el router suele estar virtualizado



⇒ ONT+Router:



⇒ Equipo integrado:

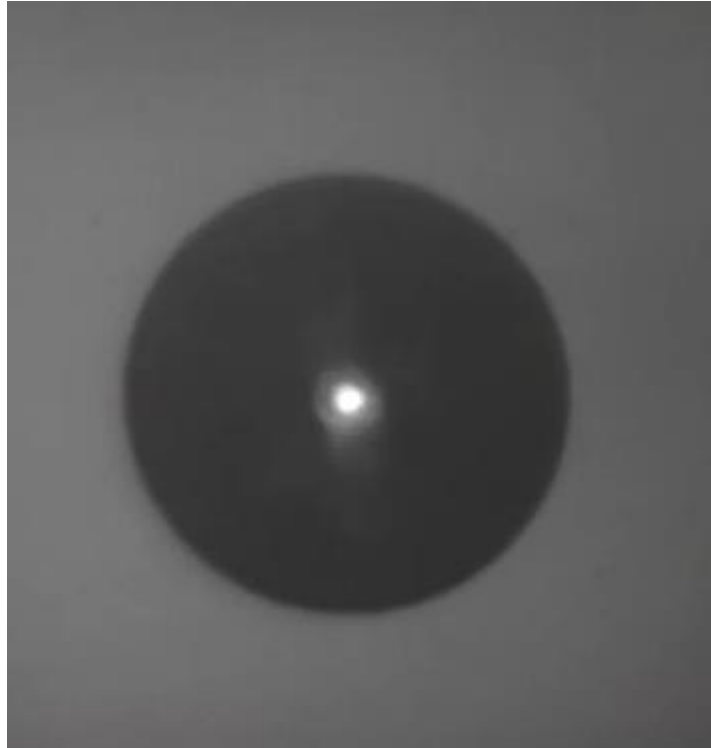


Video 15:

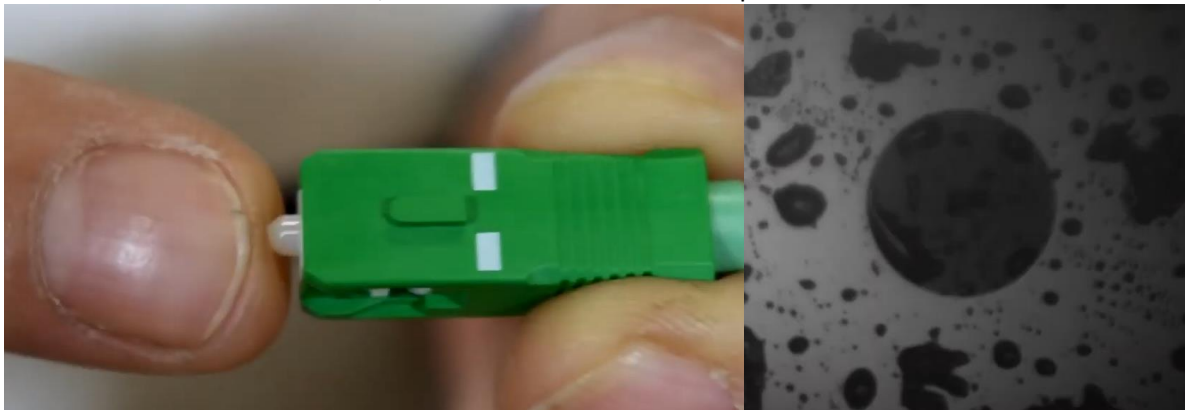
Fusión o Conectores:



Vista al microscopio:



Si toca con las manos el conector SC, mirar como se ve al microscopio:



Video 16:

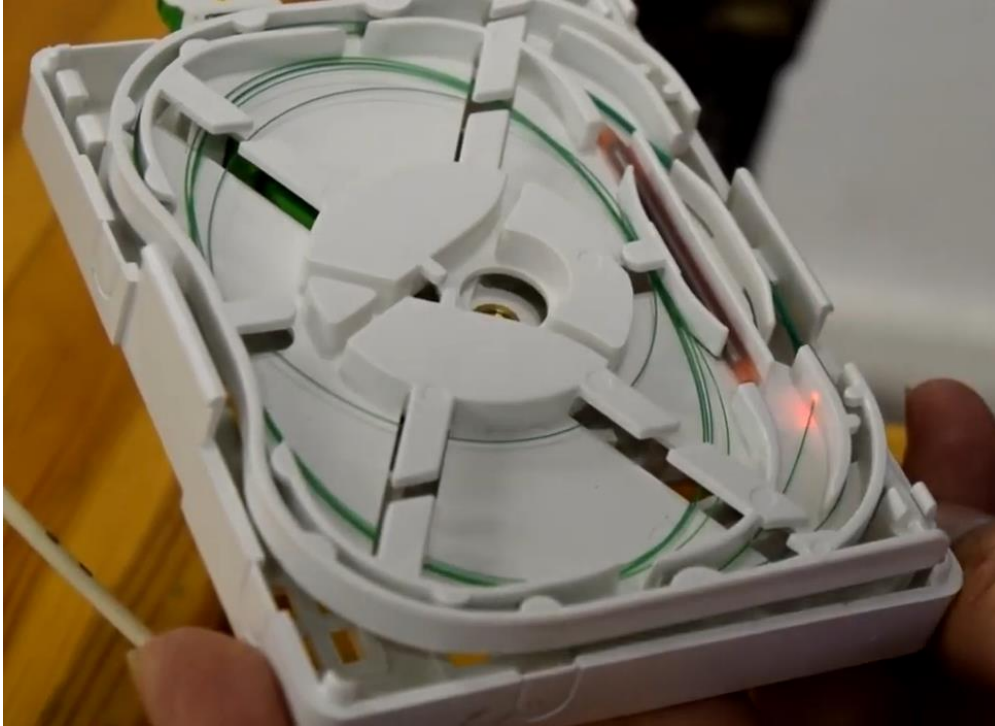
Test de fibra (Medidor de potencia GPON): nos marca la Landa (nm) y la potencia óptica (dbm) sabiendo que la ONT sincroniza a una valor entre -18 y -23 dbm



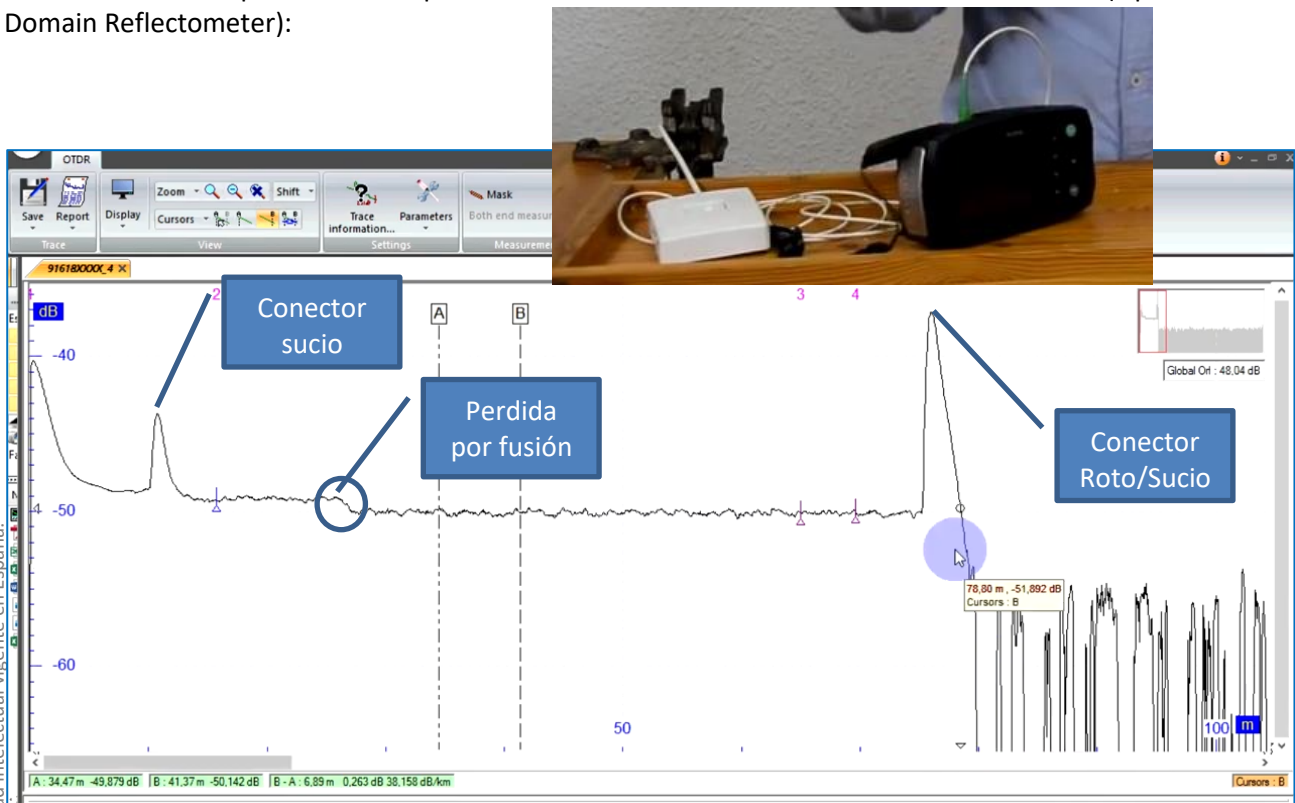
Localizador Visual de fallos:



Lo encendemos y vemos que nos parpadea porque hay una rotura de fibra:



Para hacer una comprobación más precisa usamos la Gráfica Reflecto-métrica usando OTDR (Optical Time Domain Reflectometer):



Web sobre Cómo es el despliegue de la red de fibra FTTH GPON de Movistar:

<https://bandaancha.eu/articulos/como-despliegue-red-fibra-ftth-gpon-9852>

Red de acceso

La red de acceso, como ya se ha descrito anteriormente, constituye el nivel funcional de la red que [hace llegar los servicios de telecomunicaciones a los usuarios finales](#), así como atender a las peticiones que estos solicitan.

Estas redes de acceso se pueden implementar de varias formas, dando lugar a tres tipos de redes de acceso:

Redes de acceso vía cobre

Son las que usan como medio de transmisión el cobre. Destaca entre ellas las tecnologías [xDSL](#) (base de las tecnologías ADSL, vDSL...).

Redes de acceso vía fibra óptica

Son las que usan la fibra óptica como medio de transmisión, dando lugar a las redes [PON](#) (Red Óptica Pasiva), [CWDM](#) (Coarse wavelength Division Multiplexing - Multiplexación por división aproximada de longitud de onda) y redes [HFC](#) (redes híbridas coaxial-fibra).

Redes de acceso inalámbricas

Son aquellas que usan la transmisión por radiofrecuencia para la transmisión de la información. Destacan principalmente las tecnologías [WLL](#) (Wireless Local Loop" o bucle de abonado sin hilos) y [LMDS](#) (Sistema de Distribución Local Multipunto - Local Multipoint Distribution Service - despliegue de servicios fijos de voz, acceso a Internet, comunicaciones de datos en redes privadas, y video bajo demanda).

Red troncal de transporte

La red de transporte es el nivel funcional de una red de comunicaciones que [hace llegar los servicios de telecomunicaciones a la red de acceso](#) y que agrupa la mayor parte del tráfico de un conjunto de usuarios. Se denomina habitualmente el [backbone](#) de la red.

[La mayor parte de esta red se implementa sobre fibra óptica debido a la gran cantidad de tráfico que debe soportar.](#)

Los motivos que llevan a usar la fibra óptica para este nivel de red son los siguientes:

- ⇒ Baja atenuación.
- ⇒ Gran ancho de banda.
- ⇒ Fácil instalación.
- ⇒ Inmunidad ante interferencias electromagnéticas.
- ⇒ Alta seguridad.
- ⇒ Integración con cualquier tipo de red.

El tipo de [fibra óptica habitualmente](#) utilizado es la fibra [monomodo](#), ya que es la que cumple los anteriores requisitos.

En este nivel de red han aparecido las siguientes tecnologías:

- ⇒ Tecnología PDH (jerarquía digital plesiócrona - Plesiochronous Digital Hierarchy) permite enviar varios canales telefónicos sobre un mismo medio usando técnicas de multiplexación por división de tiempo y equipos digitales
- ⇒ Tecnología SDH (jerarquía digital síncrona - Synchronous Digital Hierarchy) dispositivo digital que trabaja realizando multiplexación por división de tiempo
- ⇒ Tecnología DWDM (Dense Wavelength Division Multiplexing) multiplexado denso por división en longitudes de onda

Red de distribución

La red de distribución es el nivel funcional de una red de comunicaciones que [se encarga de realizar las tareas de conmutación y multiplexación de la información procedentes de los proveedores de servicio y adaptarlas y entregarlas al siguiente nivel de la red de telecomunicaciones](#), es decir, [a la red troncal](#).

[Es misión de esta red de distribución realizar el establecimiento y liberación de las conexiones con los diferentes bucles de abonado](#), además de [gestionar el ancho de banda de la información transmitida entre los proveedores de servicios y los usuarios](#).

En este nivel de red es donde [se instalan y configuran los siguientes elementos](#):

Conmutadores

Son dispositivos que **permiten realizar conexiones de varios usuarios e interconectarlos empleando un mismo enlace físico**, optimizando así los recursos de la red.

Multiplexores

Son dispositivos que **permiten transmitir por el mismo canal diferentes fuentes de información de diferentes usuarios (multiplexación)**, haciendo que cada una de ellas se transmita sin interferir con el resto. Para ello se emplea la **multiplexación en frecuencia, en el tiempo y la multiplexación estadística**.

Multiplexación

En una red de comunicaciones se transmite mucha información de diferentes usuarios y a diferentes puntos.

Es por ello que se basa en una arquitectura e infraestructuras compartidas que son usadas por muchos usuarios y equipos.

Para que sobre un mismo soporte pueda transmitirse información de diferentes usuarios es preciso emplear técnicas como la multiplexación, que **permite que cada información llegue a su destino correcto y no se solape con el resto de información de otros usuarios**.

A continuación, veremos estas técnicas y en especial la multiplexación como la técnica más usada en las modernas redes de comunicaciones.

En un enlace de comunicaciones entre dos o más equipos siempre se busca maximizar la capacidad del canal, es decir, **se busca el máximo rendimiento de las comunicaciones**.

Ello se consigue utilizando el mismo canal para transmitir varias comunicaciones independientes a la vez. Una de las técnicas más empleadas para esto es la citada multiplexación.

Se define la multiplexación como la técnica que permite la transmisión por el mismo canal de diferentes comunicaciones independientes entre sí y de diferentes dispositivos o equipos asegurando que la transmisión de cada una de ellas sea fiable, segura y que no interfieran entre ellas consiguiendo con ello la máxima eficiencia de transmisión del canal.

La multiplexación es una técnica empleada en el tratamiento de señales que permite por el mismo canal enviar varias comunicaciones a la vez sin que interfieran unas con otras. Con ello se mejora la eficiencia de los canales de transmisión.

Multiplexores

Los multiplexores son dispositivos electrónicos encargados de realizar la multiplexación de los canales de información en las redes de comunicaciones.

Su **ubicación** se realiza, como ya se ha descrito anteriormente, **en la red de distribución junto con otros dispositivos como los conmutadores**.

Existen diferentes tipos de multiplexores:

- Multiplexores terminales.
- Multiplexores de inserción y extracción.
- Multiplexores de distribución o DXC.

Son estos últimos los más empleados en la red de distribución.

El multiplexor dispone de varias entradas de datos y una única salida por donde se transmite el conjunto de canales de entrada ya multiplexados.



Los multiplexores **basan su funcionamiento en circuitos combinatoriales capaces de seleccionar una entrada entre varias posibles y cortocircuitar entrada con salida**, de forma que la información de esa entrada vaya a la salida.

Con una conmutación temporal de varias entradas podemos ir enviando a la salida la información de varias entradas y con ello realizar la multiplexación de los canales de entrada.

Uno de los parámetros a definir en un multiplexor es el número de entradas de los que dispone.

Existen diferentes tipos de multiplexores en función del tipo de multiplexación a realizar y sobre el tipo de medio de transmisión que conmuta.

Así existen multiplexores para:

- ⇒ Redes de pares trenzados.
- ⇒ Redes de fibra óptica.
- ⇒ Redes de coaxiales.
- ⇒ Redes inalámbricas.

Existen además numerosos fabricantes en el mercado que suministran este tipo de equipos con diferentes modelos. En su hoja de especificaciones o datasheet figuran las características de cada uno de ellos.

Técnicas de multiplexación

Existen tres tipos de técnicas de multiplexación claramente diferenciadas:

- ⇒ Multiplexación por división en frecuencia (FDM).
- ⇒ Multiplexación por división en el tiempo (TDM).
- ⇒ Multiplexación por división en la longitud de onda (WDM).

Cada una de ellas presenta sus ventajas e inconvenientes, que veremos a continuación.

Dependiendo del tipo de multiplexación a realizar se deberá emplear un equipo u otro de multiplexación.

No obstante, existen determinados modelos de multiplexados que pueden aplicar varias técnicas de multiplexación.

La multiplexación se puede aplicar para cualquier tipo de red tanto cableada (par trenzado, fibra óptica) como inalámbrica.

Multiplexación por división en frecuencias (FDM)

La multiplexación por división en frecuencias o FDM (Frequency Division Multiplexing) tiene sus orígenes en las transmisiones analógicas.

Se basa en transmitir cada comunicación en diferentes ventanas de frecuencias no solapadas entre sí y con ventanas de guardas.

De este modo, aprovechamos al máximo todo el ancho de banda de canal dividiendo dicho ancho de banda en canales. En cada uno de ellos transmitimos un canal de comunicación independiente del otro.

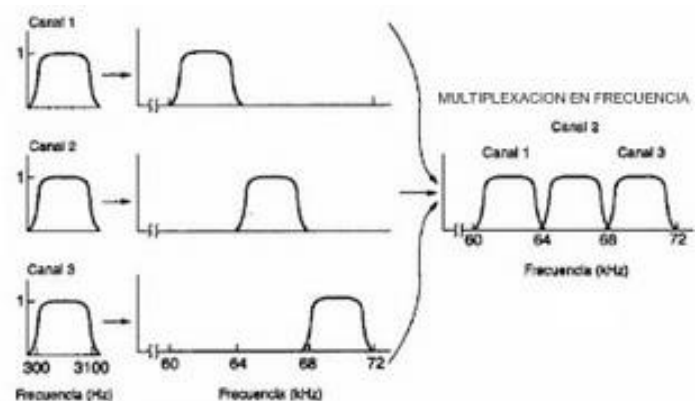
La capacidad del canal está limitada por el propio ancho de banda del canal.

La multiplexación por división en frecuencias presenta las siguientes ventajas:

- Bajo coste, es una tecnología ya muy madura.
- Posibilidad de conectarse en cascada entre varios equipos.

Las desventajas que ofrece son las siguientes:

- Presenta un número limitado de canales a transmitir en función del ancho de banda del canal.
- Presenta baja eficiencia al precisar bandas de guarda.
- Necesidad de mantener el sincronismo en las frecuencias de funcionamiento.
- Necesidad de utilizar filtros en los equipos.



Multiplexación por división en el tiempo (TDM)

La multiplexación por división en el tiempo o TDM (Time Division Multiplexing) surge a partir de las transmisiones digitales, cuando se emplea la transmisión de datos binarios (bits).

Se basa en dividir el tiempo en intervalos. En cada intervalo se envía una trama de bits de una comunicación.

De este modo, aprovechamos al máximo todo el ancho de banda de canal dividiendo dicho ancho de banda en canales donde transmitimos un canal de comunicación independiente del otro.

La capacidad del canal está limitada por el propio ancho de banda del canal.

Multiplexado estadístico o asíncrono

La multiplexación estadística consiste en asignar intervalos de tiempos a cada canal en función del tráfico existente y demandado por el canal y no de forma fija como en la multiplexación por división en el tiempo (TDM).

De esta forma a aquellos canales con tráfico crítico o con mucho tráfico se le asignan más intervalos y a aquellos con menos tráfico o tráfico no crítico se le asignan menos canales.

Por tráfico crítico se entienden servicios de telecomunicaciones con QoS (quality of service) como VoIP, TvIP, streaming de vídeo, etc.

Este multiplexado (aunque más complejo) es el más eficiente, ya que reparte los recursos de la red en función de cada tipo de tráfico y su carga, con lo que se consigue optimar el canal portador.

Multiplexación por División de Longitud de Ondas o WDM

Existe otro tipo de multiplexación también muy empleada que es la multiplexación por división de longitud de ondas o WDM (WaveLength Division Multiplexing).

Es un tipo de multiplexación similar a la multiplexación por división de frecuencias, pero en vez de hablar de frecuencias hablamos de longitud de ondas (λ).

Como ya sabemos, frecuencia (f) y longitud de onda (λ) están relacionadas mediante la siguiente expresión:

$$\lambda = c/f$$

Este tipo de multiplexación se emplea en transmisiones por infrarrojos, en las que el canal de transmisión se divide en ventanas de longitud de onda y cada canal de comunicación se envía en una determinada ventana de transmisión.

De este modo somos capaces de enviar n comunicaciones simultáneamente, cada una de ellas en una ventana o longitud de onda diferente.

El número de ventanas o canales que se pueden transmitir dependerá de la longitud de la ventana de cada canal y de la longitud de la ventana total del medio de transmisión, es decir:

$$N^{\circ} \text{ canales} = (\text{Longitud de la ventana del medio de transmisión}) / (\text{Longitud de la ventana de un canal})$$

La multiplexación por división de longitud de onda cada vez se emplea más en las transmisiones, ya que suelen ofrecer un gran ancho de banda para los nuevos servicios de telecomunicaciones.

En el emisor requiere de un multiplexor óptico y en el receptor de un demultiplexor óptico.

Funciones de conmutación, transporte y señalización

La conmutación es una técnica ampliamente utilizada en comunicaciones consistente en poner en contacto un equipo con otro empleando una infraestructura común de comunicaciones para la transmisión de los datos.

Con ello se pretende dar eficiencia al sistema, ya que varios equipos pueden emplear la misma infraestructura para enviar datos y no crear redes y recursos individuales para cada transmisión, lo que encarecería enormemente la infraestructura.

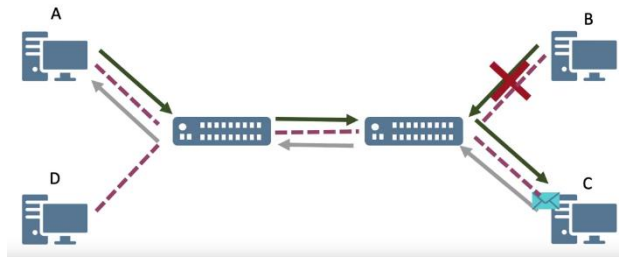
Esta técnica, por tanto, permite que en un momento dado el equipo emisor y el equipo receptor estén conectados para la transferencia de la información. Cuando termina la transferencia, se liberan los recursos para que puedan ser usados para otra transmisión de otros equipos.

Existen dos tipos o técnicas de conmutación:

Conmutación de circuitos

Es una técnica de conmutación basada en el establecimiento de una conexión física entre los dos extremos (emisor y receptor), empleando para ello y conectando todos los elementos y nodos intermedios para que durante la transferencia exista ese camino físico para el intercambio de la información.

Cuando se termina la transferencia, se liberan todas las conexiones intermedias y quedan a disposición de la red para otra comunicación del mismo o de diferentes equipos.



Esta conmutación de circuitos permite ser implementada de dos formas:

- ⇒ Conmutación de circuitos **espacial**:

Es aquella en la que durante la transferencia **el circuito establecido está permanente y en exclusiva para el emisor y el receptor y sólo se libera cuando haya finalizado la transferencia.**

- ⇒ Conmutación de circuitos **temporal**:

En esta **se crean espacios temporales de transmisión**, de forma que cada comunicación emplea una serie de intervalos de tiempo para transmitir, pero todas las comunicaciones emplean el mismo circuito o enlace físico.

En la conmutación de circuitos se establece una conexión física a través de varios enlaces y nodos intermedios para unir emisor y receptor.

En la conmutación de circuitos temporal, en cambio, se establecen slots temporales.

Cada slot temporal es utilizado por un canal de comunicaciones, pero todos ellos usan el mismo enlace o circuito de conexión.

Este tipo de conmutación es el empleado por el servicio telefónico RTC (red telefónica conmutada).

Conmutación de paquetes

La conmutación de paquetes es una técnica que **'trocea' la información en paquetes de longitud fija y envía cada paquete desde el emisor al receptor.**

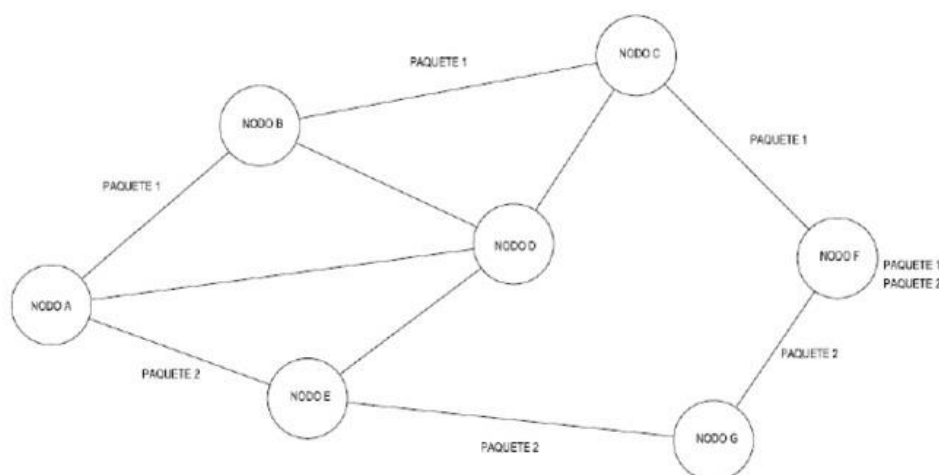
Cada paquete puede utilizar una ruta diferente para llegar al destino y para poder llegar a ese destino cada paquete incluye una cabecera (además de la información) que contiene la dirección de origen y destino del paquete.

Cada nodo de la red analiza la cabecera del paquete y decide si es para él o si debe enrutarlo a otro nodo para hacerlo llegar a su destino.

Este tipo de conmutación es más eficiente que la conmutación de circuito, ya que aprovecha mejor los recursos del sistema.

Internet utiliza para la transmisión de los datos la conmutación de paquetes.

Uno de los problemas de este tipo de conmutación es el retardo sufrido por los paquetes (al tener que pasar por muchos nodos), las pérdidas de paquetes y la llegada desordenada de paquetes del mismo mensaje al receptor (ya que llegan por rutas diferentes). Estos dos últimos problemas implican que el sistema debe incluir técnicas de detección y corrección de paquetes.



La conmutación de paquetes admite dos implementaciones:

Modo circuito virtual u orientado a conexión

En este tipo de conmutación de paquetes, **todos los paquetes pertenecientes al mismo mensaje siguen la misma ruta**, por lo que previamente se debe establecer un circuito virtual entre emisor y receptor.

Así, cada paquete además de la dirección origen y destino incluye el número de conexión por el que va dirigido.

La ventaja es que evita la llegada desordenada de paquetes, aunque incluye más retardo debido al tiempo de establecimiento del circuito virtual.

Modo datagrama u orientado a no conexión

En este tipo de conmutación de paquetes, **cada paquete puede ir por rutas diferentes dependiendo del estado de la red y de cada uno de los enlaces**.

Esto implica que los paquetes puedan llegar desordenados y serán aplicaciones de niveles superiores a nivel de red quienes deban aplicar técnicas de detección y corrección de paquetes.

La ventaja es que hay menos retardo que en el modo circuito virtual u orientado a conexión.

Señalización

La **señalización** es un **procedimiento de gestión y control que incluyen las redes de comunicaciones y que se emplea como medio para establecer, mantener y finalizar las diferentes conexiones y/o llamadas**.

Gestiona por tanto **los conmutadores para que puedan realizar sus funciones adecuadamente en función de las necesidades del tráfico en la red**.

Esta señalización cobra especial relevancia en las redes de conmutación telefónica (RTC).

En este caso, **la señalización realiza las siguientes funciones**:

- ⇒ Establecer la comunicación con el abonado incluyendo los tonos de marcado, tono de llamada, señal de ocupado, etc.
- ⇒ intermedios de la red.
- ⇒ Transmisión de la información por la red.
- ⇒ Conexión con el abonado destino.
- ⇒ Realizar las funciones de tarificación.
- ⇒ Liberar la conexión una vez finalizada.
- ⇒ Control y gestión del sistema.

Existen **dos procedimientos para realizar la señalización**:

- **Señalización por canal común.**
Esta señalización se caracteriza por **enviar los mensajes de señalización por el mismo canal donde se transmite la información**.
- **Señalización por canal asociado.**
Esta señalización se caracteriza por **enviar los mensajes de señalización por un canal dedicado especialmente a este fin e independiente del canal donde se envía la información**.

La función de transporte de la información en las redes de comunicaciones se encarga de hacer llegar al destinatario la información enviada por el emisor de forma correcta, ordenada y sin errores.

Además, deberá asegurar que dicha información sea entregada en los plazos establecidos en la QoS de cada servicio, algunos muy críticos como pueden ser:

- ⇒ Servicio de VoIP.
- ⇒ Servicio de TvIP.
- ⇒ Servicio de streaming de vídeo.

Arquitectura de redes de voz y datos

Redes de acceso: guiadas y no guiadas

Para que los servicios de telecomunicaciones lleguen a los usuarios y equipos finales es necesaria una red que distribuya dichos servicios desde donde se generan hasta la ubicación donde están estos usuarios y equipos finales.

Es por ello que las redes de comunicaciones están estructuradas en diferentes niveles, que son: nivel de acceso, nivel troncal o de transporte y nivel de distribución.

El nivel o red de acceso es la parte de la red que permite que los servicios de telecomunicaciones lleguen hasta los hogares, usuarios y empresas. Representa el último tramo de la red y proporciona por eso el 'acceso' de los servicios.

Esta parte de la red, debido a que debe hacer llegar los servicios a todos los usuarios y ubicaciones, es la que presenta mayor capilaridad y por ello la más costosa y en numerosas ocasiones las más difícil de implementar.

En numerosas bibliografías a este tramo de la red se le conoce como 'la última milla'.

Las redes de acceso pueden implementarse de dos formas:

Redes de acceso guiadas

Son aquellas que usan un cable (medio guiado) como medio de transmisión.

En este caso existen dos variantes fundamentalmente:

- ⇒ Redes de acceso basado en el **cobre**: son las que usan como medio de transmisión el cobre, destacando entre ellas las tecnologías **xDSL**.
- ⇒ Redes de acceso basado en la **fibra óptica**: son las que usan la fibra óptica como medio de transmisión, dando lugar a las redes **PON** (Red Óptica Pasiva), **CWDM** (Coarse wavelength Division Multiplexing - Multiplexación por división aproximada de longitud de onda) y redes **HFC** (redes híbridas coaxial-fibra).

Redes de acceso no guiadas

Son las redes inalámbricas que emplean la radiofrecuencia para la transmisión de la información.

Destacan principalmente las tecnologías **WLL** (Wireless Local Loop" o bucle de abonado sin hilos) y **LMDS** (Sistema de Distribución Local Multipunto - Local Multipoint Distribution Service - despliegue de servicios fijos de voz, acceso a Internet, comunicaciones de datos en redes privadas, y video bajo demanda)

Las redes de acceso (también denominadas la última milla) constituyen la última parte de la red que hace llegar los servicios de telecomunicación a los usuarios.

Las redes basadas en el cable, como las tecnologías xDSL y la fibra óptica, presentan una serie de ventajas como:

- ⇒ Son altamente fiables y seguras.
- ⇒ Ofrecen en general un gran ancho de banda.
- ⇒ Permiten adaptarse a las nuevas tecnologías.

Por el contrario, presentan el gran inconveniente del alto coste de instalación, ya que su alta capilaridad exige habitualmente apertura de zanjas, permisos, tendidos de cables, servidumbres, etc.

En cambio, las redes inalámbricas en sus diferentes implementaciones presentan las siguientes ventajas:

- ⇒ Su despliegue es mucho más rápido que las redes cableadas.
- ⇒ Son muy aptas para zonas u orografías accidentadas o complicadas.
- ⇒ Su coste es sensiblemente menor que las redes cableadas.

Por el contrario, presentan ciertos inconvenientes como:

- ⇒ En general, su ancho de banda es menor que las redes cableadas.
- ⇒ Su nivel de seguridad es menor que las redes cableadas (cada vez menos).

Redes de acceso vía cobre: xDSL

Ya hemos visto que las redes de acceso se pueden implementar mediante redes cableadas y guiadas, siendo las más utilizadas las redes con tecnologías xDSL, que son redes basadas en el cable de cobre.

Las tecnologías xDSL surgieron al aprovechar el hilo telefónico ya instalado en la mayoría de los hogares y edificios (su sustitución por otro tipo de cable era muy costoso) para ofrecer servicios de telecomunicaciones de banda ancha y de alta velocidad.

xDSL representa una tecnología que admite diversas variantes o **modalidades** como:

- ⇒ ADSL.
- ⇒ HDSL.
- ⇒ SDSL.
- ⇒ RADSL.
- ⇒ VDSL.
- ⇒ Etc.

Cada una de ellas tiene unas características y prestaciones diferentes que la hacen adecuada para un determinado servicio u otro.

La más conocida, por ser la más ampliamente utilizada, es la tecnología **ADSL** (Asymmetric Digital Subscriber Line). Se emplea habitualmente en el bucle de abonado.

Proporciona un gran ancho de banda, con una alta velocidad y conexión permanente. Además, se realiza individualmente sobre cada línea, sin tener que afectar al resto de líneas de otros usuarios.

Esto último da gran facilidad y operatividad al **operador**, que puede activar este tipo de tecnología individualmente a cada cliente que lo solicite. Además, no requiere un acondicionamiento especial en la central telefónica que sostiene a la línea en cuestión.

El ADSL basa su funcionamiento en crear dos canales de datos de ancho de banda asimétricos (uno de subida y otro de bajada). El canal de bajada (de central al usuario) es mucho mayor que el de subida (de usuario a central), que se adapta perfectamente al servicio de datos de Internet, ya que habitualmente el flujo de información de descarga es mucho mayor que el de subida.

Es por ello que el ADSL se emplea para el acceso a datos de Internet, dado que ofrece gran ancho de banda gracias a esta asimetría de los canales.

Además de los canales de datos, el ADSL incorpora un tercer canal destinado al servicio de voz, es decir, al servicio telefónico básico.

Esta estructura de tres canales obliga a poner dispositivos denominados **splitter** (que son básicamente filtros) que permitan separar en el usuario cada canal para que pueda emplearse simultáneamente el servicio de voz telefónico habitual con el servicio de datos en el mismo PTR del usuario.

Un inconveniente de esta tecnología es su fuerte dependencia con la distancia entre el usuario y la central. A mayor distancia, la velocidad y ancho de banda operativa se reduce debido a las interferencias y a la atenuación que aumenta en los cables. Es por ello que la distancia máxima permitida para ofrecer ADSL de buen ancho de banda se ha establecido en 6 km.

No obstante, esta tecnología ha evolucionado a versiones como **ADSL2** y **ADSL2+** que permiten alcanzar mayores distancias y mayores velocidades que el ADSL original.

En la siguiente tabla puede verse cómo ha evolucionado esta tecnología y con ello ha aumentado sus prestaciones en cuanto a ancho de banda y velocidades de transmisión.

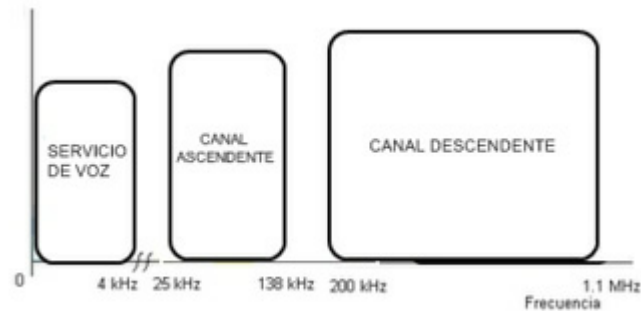
Nombre Comercial	Velocidad de bajada máxima	Velocidad de subida máxima
ADSL	8 Mbit/s	1,0 Mbit/s
ADSL (G.DMT)	12 Mbit/s	1,3 Mbit/s
ADSL2	12 Mbit/s	1,0 Mbit/s
ADSL2+	24 Mbit/s	1,0 Mbit/s
ADSL2+M	24 Mbit/s	3,5 Mbit/s

El funcionamiento del ADSL consiste en utilizar parte del espectro de frecuencia del hilo telefónico y que no empleaba el servicio telefónico básico de voz, cuya banda cubría desde los 300 Hz hasta los 3.400 Hz.

El servicio de datos que incluye el ADSL usa la banda de frecuencias desde los 24 KHz hasta los 1,1 Mhz, y mediante multiplexación y demultiplexación permite al usuario tener ambos servicios en el mismo PTR.

Esta multiplexación realizada en la central se consigue con unos dispositivos (habitualmente modems) que luego son concentrados en los llamados DSLAM, mientras que en el usuario se instala otro módem o set-top-box que debe incluir unos filtros (splitters) para separar ambas bandas y dar por un lado el servicio de voz (que se conecta al terminal telefónico) y el servicio de datos (conectado habitualmente a un router).

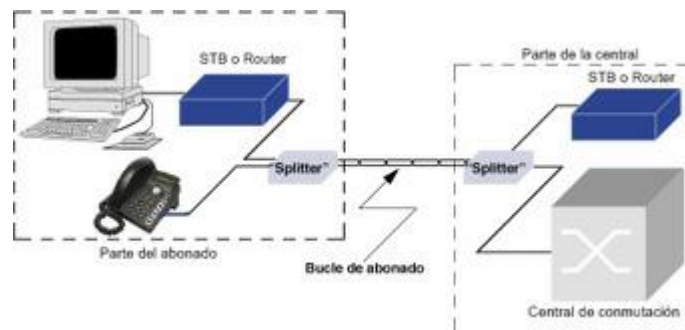
En la siguiente figura podemos ver el espectro de frecuencias asignados a cada banda del ADSL:



Los **splitters** o **filtros** anteriormente mencionados **realizan dos funciones**:

- ⇒ **Separar** o **combinar bandas de trabajos de diferentes servicios**.
- ⇒ **Eliminar interferencias entre ambas bandas**, es decir, que las señales del servicio telefónico de voz no se acoplen en la banda del servicio de datos y viceversa.

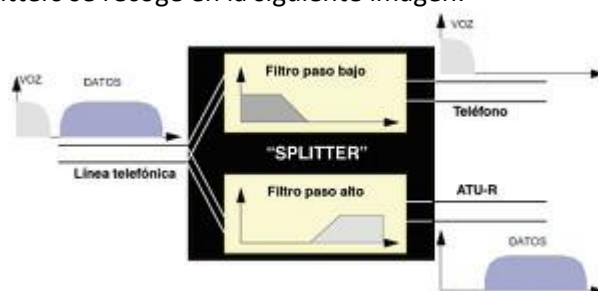
En la siguiente figura podemos ver cómo es el esquema y los elementos que componen esta estructura del ADSL.



En la figura anterior puede comprobarse la presencia del splitter como un conjunto de filtros paso bajo y paso alto que separa la banda del servicio telefónico básico y entrega su salida a un terminal telefónico y la banda del servicio de datos que entrega su salida a un router o un set-top-box.

En la parte de la central ocurre lo mismo ya que el servicio telefónico se conecta a la matriz de conmutación del operador de telefonía (RTC) y el servicio de datos se conecta a un proveedor de servicios de Internet (ISP).

El funcionamiento de los splitters se recoge en la siguiente imagen:



Como se aprecia en la anterior imagen, el **splitter es un conjunto de dos filtros**:

- ⇒ **Filtro paso bajo**: que deja pasar solo las **frecuencias bajas**, es decir, el servicio **telefónico** básico.
- ⇒ **Filtro paso alto**: que deja pasar solo las **frecuencias altas**, es decir, el servicio de **datos**.

Las frecuencias de corte de ambos filtros están ajustados al ancho de banda de cada una de las bandas que se quiere dejar pasar.

Además, **este filtrado evita las interferencias de unas bandas con otras, con lo que la señal es más 'limpia' e inmune y eso se traduce en una mayor velocidad de transmisión**.

Aunque el ADSL es una tecnología muy utilizada, existen otras tecnologías que forman parte de este grupo de sistemas xDSL.

Entre ellas destacan las variantes siguientes:

HDSL (High Data Rate Digital Subscriber Line)

Se trata de una tecnología con anchos de banda y caudales simétricos en el enlace ascendente y descendente pero de alta velocidad.

No se permiten distancias superiores a 3,6 Km desde el usuario a la central.

SDSL (Symmetric Digital Subscriber Line)

Se trata de una versión comercial y mejorada del HDSL con ancho de banda y caudales también simétricos en el enlace ascendente y descendente, pero con mayores tasas de velocidad.

IDSL (ISDN Digital Subscriber Line)

Se trata de una tecnología xDSL empleada habitualmente sobre redes RDSI. Cada vez está más en desuso.

RADSL (Rate Adaptative Digital Subscriber Line)

Es una tecnología similar al ADSL pero que permite el ajuste de los valores de transmisión en función de la longitud y otros parámetros de transmisión sobre el par de cobre.

Esto permite adaptar diferentes velocidades de transmisión en función del estado y propiedades del cobre instalado.

VDSL y VDSL2 (Very High Speed Digital Subscriber Line)

La tecnología VDSL es la que ofrece mayores tasas de transferencia con caudales que oscilan entre 13 y 52 Mbps en el canal descendente y de 1,5 a 2,3 Mbps en el canal ascendente.

Su principal inconveniente es que limita la distancia entre el usuario y la central a 1,5 km.

VDSL2 es una mejora de la VDSL que permite mayores tasas de transferencia en función de la distancia usuario-central.

Esta tecnología es capaz de soportar servicios de telecomunicaciones de altas requerimientos de transmisión como Vídeo bajo demanda (VoD), streaming de vídeo, etc

Redes de acceso vía radio: WLL, MMDS, LMDS

Como ya se ha descrito anteriormente la red de acceso también puede implementarse mediante tecnologías inalámbricas, es decir, por radiofrecuencias.

Entre ellas destacan principalmente las tres tecnologías más utilizadas:

Tecnología WLL

WLL son las siglas de Wireless Local Loop, es decir, bucle local inalámbrico. Se trata de una tecnología vía radio para ofrecer servicios de banda ancha (fundamentalmente telefonía e Internet) a los usuarios en la última parte de la red, es decir, la conocida como la última milla.

Trabajan en las frecuencias licenciadas, por lo que está asignado habitualmente a operadores de telecomunicaciones privadas que proporcionan los servicios de telefonía y de Internet a sus clientes.

La estructura de esta tecnología se basa en la instalación de un conjunto de estaciones base interconectadas entre sí por radiofrecuencia que concentran todo el tráfico de datos y que transmiten dicha información a los terminales de los usuarios mediante radioenlace.

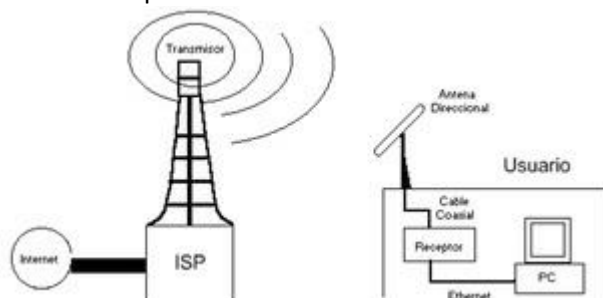
Como toda tecnología por radioenlace, precisa de visibilidad directa entre las estaciones base, por lo que su ubicación debe realizarse en lugares altos como edificios, torretas, montes, etc.

La gran ventaja de esta tecnología es su rápida implementación (no requiere cable), un ancho de banda alto para los usuarios (del orden de Mbps) y fácilmente escalable.

Como desventaja tiene que es más sensible a la orografía del terreno y a las interferencias con otras señales de radiofrecuencia.

La transmisión se puede realizar siguiendo una estructura punto-punto o punto-multipunto.

En la siguiente figura podemos ver el esquema de funcionamiento de la tecnología WLL:



En la figura anterior se puede observar que el sistema se basa en una antena transmisora que da servicio por radio a los abonados que requiere de una antena receptora en sus tejados, la cual con la electrónica necesaria se conecta al PTR (Punto de terminación de red) de la vivienda.

Esta antena transmisora puede dar **cobertura** a una zona de **varios kilómetros** y a una **horquilla de entre 50 y 200 usuarios**, dependiendo de la orografía de la zona de cobertura.

Esta antena se conecta a la red de telefonía conmutada (RTC) o a los proveedores de servicios de Internet (ISP) generalmente por cable, ya sea por par trenzado o fibra óptica.

Es imprescindible en esta arquitectura la visibilidad entre antena transmisora y receptora.

Esta tecnología se ha empleado con mucha frecuencia para dar servicio telefónico y de Internet a zonas rurales o de difícil acceso.

Tecnología MMDS

MMDS son las siglas de Microwave Multipoint Distribution System, es decir, **Sistema de distribución multipunto por microondas**.

Se trata de una tecnología que se creó para dar servicio de vídeo y televisión por radio a zonas con orografía complicada y zonas alejadas (entornos rurales, polígonos industriales, etc.).

Las velocidades de transmisión rondan los 2-5 Mbps para alcances de hasta 20 km, y aunque se empleó originariamente para vídeo luego se amplió para el servicio de datos.

La banda de MMDS utiliza la frecuencia de microondas, es decir, de los 2 Ghz a 3 Ghz en banda L. Es por ello que requiere antenas cuya banda de trabajo esté en estas frecuencias, así como un decodificador para la recepción en el abonado.

En la siguiente imagen puede verse el ejemplo de una antena receptora para esta tecnología de MMDS.



Esta tecnología ha sido sustituida por su sucesor, el LMDS, cuyo funcionamiento es similar pero ofrecía mayores velocidades de transmisión.

Tecnología LMDS

LMDS son las siglas de Local Multipoint Distribution System, es decir, Sistema de distribución multipunto local.

Se trata de una tecnología inalámbrica para la distribución de servicios de datos y/o telefónico en la frecuencia de los 26 Ghz.

Requiere, al igual que la tecnología MMDS, de un conjunto de estaciones base interconectadas entre sí (bien por radioenlace o de forma cableada) y donde cada una de ellas da cobertura inalámbrica a un conjunto de usuarios que precisará una antena receptora junto con su decodificador.

La distancia entre antenas no debe superar los 3 km (menor que MMDS) pero se alcanzan mayores velocidades de transmisión que oscilan entre los 50 a 622 Mbps.

Esta tecnología ha sido ampliamente utilizada en entornos rurales y en zonas con orografía complicada para sustituir los antiguos TRAC (servicio telefónico básico por radio), ofreciendo con ello este servicio de voz además de servicio de Internet.

La tecnología LMDS es una tecnología inalámbrica muy utilizada para dotar de servicios de banda ancha a zonas aisladas (rurales) y de difícil orografía donde los costes de dotarlo con redes cableadas los hacen inviables.

WIMAX

WiMAX (Worldwide Interoperability for Microwave Access - interoperabilidad mundial para acceso por microondas), es una norma de transmisión de datos que utiliza las ondas de radio en las frecuencias de 2,5 a 5,8 GHz y puede tener una cobertura hasta de 70 km.

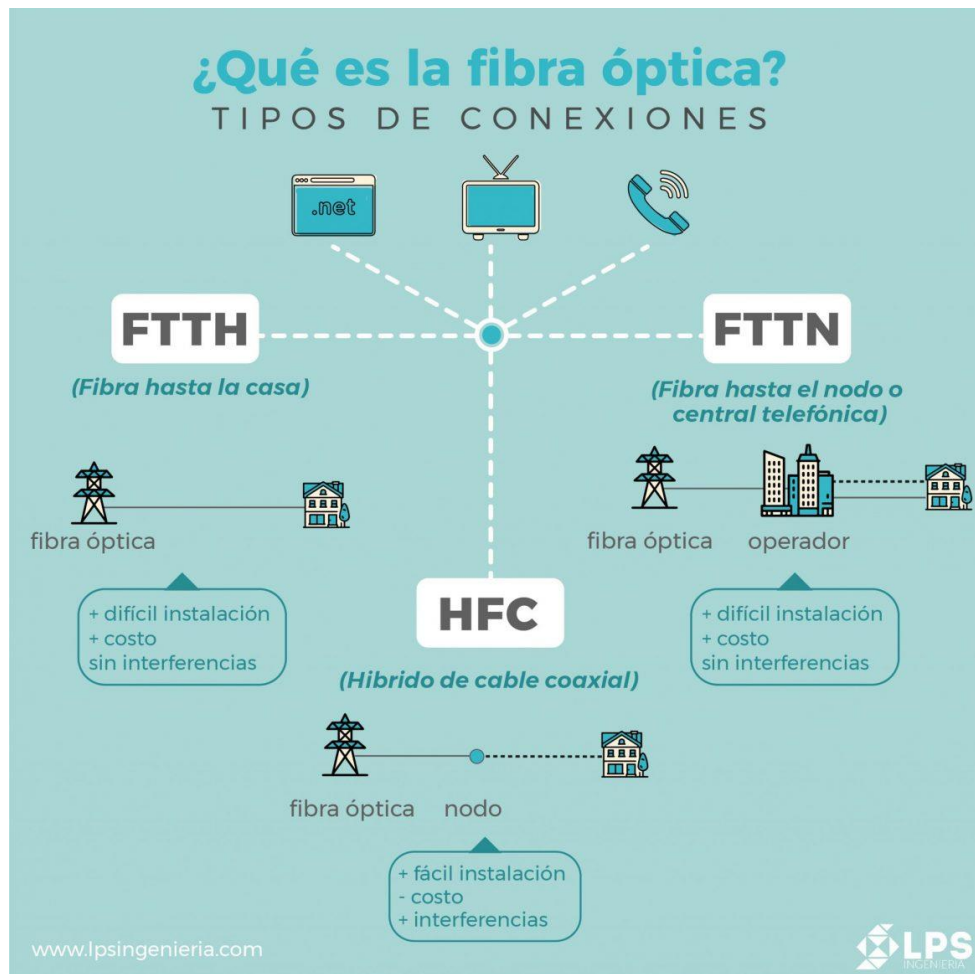
Es una tecnología dentro de las conocidas como tecnologías de última milla, también conocidas como bucle local que permite la recepción de datos por microondas y retransmisión por ondas de radio



Simulador WIMAX: <https://airlink.ui.com/#/ptp>

Redes de acceso vía fibra óptica: HFC, PON y CWDM

Las redes de acceso también se pueden implementar con redes de fibra óptica como medio de transmisión cableada.



En este caso, dan lugar a diferentes tecnologías, como son:

- ⇒ Tecnología HFC.
- ⇒ Tecnología PON.
- ⇒ Tecnología CWDM.

A continuación, veremos con más detalles cada una de ellas.

Tecnología HFC

HFC son las siglas de Hybrid Fiber Coaxial, es decir, redes híbridas de fibra y coaxial.

Se trata de una red de acceso que **aprovecha parte del coaxial ya instalado (la parte más capilarizada, es decir, el último tramo del abonado) y lo conecta a redes de fibra óptica (la parte troncal de la red)** para con ello ofrecer mayores velocidades de transmisión al usuario.

Esto se realiza porque **la velocidad de transmisión tiene una fuerte dependencia con la distancia desde la central al nodo de conmutación**. Es por ello que **se intenta minimizar el tramo de coaxial** (que supone el cuello de botella de la red) empleando el otro tramo con fibra óptica, cuya atenuación y ancho de banda es mucho más elevado.

De esta forma se consigue aumentar sensiblemente las prestaciones de red sin suponer altos costes de instalación.

Estas redes HFC presentan **dos niveles jerárquicos**:

- ⇒ **Un primer nivel con topología en estrella que parte del nodo o central y que llega a nodos secundarios distribuidos en una ciudad** con armarios tipo muxfin (MULTipleXor Flexible de Interfaces Normalizadas) o similar y **que realiza la conversión óptico-eléctrica y su distribución por coaxial**.



⇒ Un segundo nivel con topología en bus que partiendo desde los nodos secundarios llega a todos los abonados mediante coaxial.

Con las redes HFC se han podido distribuir servicios de voz, vídeo y datos a alta velocidad bajo la denominación **Triple play**.

Cada nivel de la red es la responsable de unas determinadas funciones con una gestión propia, es decir:

Red de transporte

Corresponde a la parte de **fibra** y suele estar formada por **un anillo de fibra primaria** al que se conectan **diversos anillos de fibra secundarios** también en fibra.

Concentra el tráfico troncal y sus velocidades de transferencia son muy altas.

Cada nodo secundario puede alimentar a unos 500 abonados y un nodo primario a unos 10.000 abonados.

Red de distribución

Corresponde a **la parte de coaxial** que, **partiendo del nodo secundario óptico y siguiendo una topología en bus, llega hasta los abonados.**

En el abonado **es necesario un elemento adaptador que se ubica generalmente en fachada y que realiza además la conversión a cobre.**

En esta tecnología HFC han aparecido dos estándares que definen sus especificaciones:

- **DOCSIS:** es un estándar que surgió en 1997 por los principales operadores de cable de Estados Unidos y que ha tenido gran implementación.
- **DAVIC:** es un estándar definido por la DVB (Digital Video Broadcast) pero que no ha tenido tanta implantación como la anterior.

Tecnología PON

PON son las siglas de Passive Optical Network. Representa una tecnología de red de acceso basada en una **red óptica pasiva**, es decir, **no incluyendo elementos activos (amplificadores) en la red de distribución desde la central hasta el abonado.**

Este tipo de tecnología es ampliamente **utilizada en las redes FTTH.**

Su alta expansión se debe a las continuas demandas de ancho de banda y prestaciones que requieren los nuevos servicios de telecomunicaciones y en especial los servicios Triple play.

Las redes de cable y de ADSL (aunque cubría gran parte de estas redes) tienen el problema de que están limitadas en prestaciones, ya que a lo sumo pueden ofrecer 100 Mbps en el canal descendente y 50 en el canal ascendente, además de su fuerte dependencia con la distancia entre la central y el abonado (máximo 6 km).

Es por ello que la fibra óptica es la solución alternativa que permite superar estos obstáculos, ya que sus **ventajas** son:

- ⇒ **Gran ancho de banda.**
- ⇒ **Descenso** continuo del **precio** de sus elementos y en especial de la fibra y de los láseres ópticos.
- ⇒ Aumento de la **distancia hasta 20 Km** entre la central y los abonados.
- ⇒ **Inmunidad** ante las **interferencias electromagnéticas**.
- ⇒ Posibilidad de incorporar **nuevos servicios de telecomunicación con QoS**.

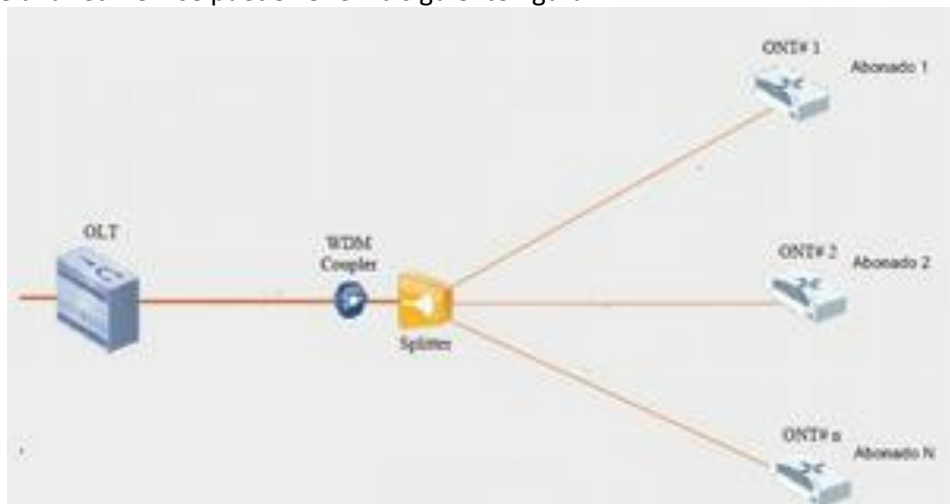
Una red PON la forman los siguientes **elementos**:

- ⇒ Un módulo **OLT** (Optical Line Terminal), que es un dispositivo óptico ubicado en la central.
- ⇒ Un divisor óptico denominado **splitter**.
- ⇒ Numerosos **ONTs** (Optical Network Terminal), que son dispositivos ópticos situados en el domicilio del abonado.

La red PON se comporta de dos maneras:

- ⇒ Como red **Punto-Multipunto** en la dirección **Central-Abonado**, es decir, en la dirección **OLT-ONT** (**canal descendente**).
- ⇒ Como red **Punto-Punto** en la dirección **Abonado-Central**, es decir, en la dirección **ONT-OLT** (**canal ascendente**).

El esquema de una red PON se puede ver en la siguiente figura:



En ella se aprecia cómo el módulo OLT se conecta a un splitter (divisor óptico) que da acceso a las unidades ONT ubicadas en cada abonado.

Todos los tramos se realizan con fibra óptica. En el tramo **OLT-Splitter** se emplea habitualmente **fibra multimodo** y en el tramo **Splitter-ONT** es más habitual usar **fibra monomodo**.

Las redes **PON** emplean la **multiplexación por longitud de onda (WDM)** para los canales **ascendente y descendente**.

Bajo esta tecnología y arquitectura PON se han definido diversas variantes o **estándares** entre las que destacan las siguientes:

Redes APON (ATM PON)

Se trata de una red **PON** que emplea la **tecnología ATM (modo de transferencia asíncrono)** para el canal descendente y que consigue con ello velocidades de transferencia de 622 Mbps de salida del OLT (debe repartirse entre todas las ONTs conectadas).

Redes BPON (Broadband PON)

Se trata de una red **APON mejorada** para ofrecer servicios de gran ancho de banda. En ella **se pueden configurar los canales ascendente y descendente a diferentes velocidades**.

Así, podemos conseguir las siguientes velocidades en cada canal:

- ⇒ Canal ascendente con velocidades de 155 Mbps y 622 Mbps.
- ⇒ Canal descendente con velocidades de 155 Mbps, 622 Mbps y 1,25 Gbps.

Sus especificaciones quedan definidas en la normativa ITU-T G.983.x.

Redes GPON (Gigabit PON)

Son redes PON con velocidades de transferencia de [Gigabit Ethernet](#).

Así, podemos obtener las siguientes velocidades en cada canal:

- ⇒ Canal ascendente con velocidades de 155 Mbps, 622 Mbps, 1,25 Gbps y 2,5 Gbps.
- ⇒ Canal descendente con velocidades de 1,25 Gbps y 2,5 Gbps.

Sus especificaciones quedan definidas en la normativa ITU-T G.984.x.

Redes EPON (Ethernet PON)

Son redes PON que emplean la fibra óptica para crear redes Ethernet.

Por ello siguen las especificaciones del 802.3 y su gran ventaja es que trabajan con tramas Ethernet, por lo que reducen los costes en su implementación.

Con esta tecnología se consiguen las siguientes velocidades de transferencia en cada canal:

- ⇒ Canal ascendente con velocidad de 1,25 Gbps.
- ⇒ Canal descendente con velocidad de 1,25 Gbps.

Redes 10GPON (10 Gigabit PON)

Son redes PON que consiguen velocidades de transferencia del orden de los [10Gbps](#).

Son las redes PON más avanzadas que existen en estos momentos y son capaces de admitir grandes flujos de datos.

Con esta tecnología se consiguen las siguientes velocidades de transferencia en cada canal:

- ⇒ Canal ascendente con velocidad de 1,25 Gbps o 10 Gbps.
- ⇒ Canal descendente con velocidad de 10 Gbps.

En la siguiente tabla recogemos una comparativa de todos estos estándares.

Tecnología	APON	BPON	GPON	EPON	10GPON
Estándar	ITU-T.G983x	ITU-T.G983x	ITU-T.G984x	802.3ah	802.3ah
Velocidades de transmisión (Mbps)	Canal ascendente: 155, 622 Canal descendente: 155 y 622	Canal ascendente: 155, 622 Canal descendente: 155, 622 y 1244	Canal ascendente: 155, 622, 1244, 2488 Canal descendente: 1244 y 2488	Canal ascendente: 1244 Canal descendente: 1244	Canal ascendente: 1244 y 10000 Canal descendente: 10000
Tipo de fibra empleada	Monomodo	Monomodo	Monomodo	Monomodo	Monomodo
Máxima distancia entre OLT y ONU	20 km	20 km	10-20 km	10 km	20 km
Arquitectura de transmisión	Asimétrica o simétrica	Asimétrica o simétrica	Asimétrica o simétrica	Simétrica	Asimétrica o simétrica

A continuación describiremos con más detalle cada uno de los elementos que forman parte de una red PON.

OLT

El OLT es un dispositivo óptico activo [ubicado en la central del proveedor y que se conecta al splitter o divisor óptico mediante fibra](#).

Su función principal es la de [hacer de router para gestionar el tráfico demandado por los abonados](#). Es capaz de gestionar cientos de abonados, además de permitir la conectividad con otras redes externas.

En resumen, las funciones del OLT son:

- ⇒ [Gestionar y enrutar el tráfico de las ONTs](#).
- ⇒ [Conectar la red PON con otras redes](#).

El OLT actúa como concentrador de redes externas que ofrecen los servicios de voz, datos y vídeo. Es por ello que su cabecera [permite la conectividad con](#):

- ⇒ La red [RTC](#) para los servicios de voz [telefónica básica](#).
- ⇒ Proveedores [ISP](#) para los [servicios de datos a través de un Gateway](#).
- ⇒ Proveedores de [Vídeo \(VoD\)](#) [a través de un Gateway](#).

El OLT como dispositivo hardware está formado por varios [módulos](#):

- ⇒ **Módulo P-OLT (Provider OLT)**

Es el [módulo](#) encargado de [recoger las tramas de las redes de voz y datos procedentes del exterior \(red RTC e ISP\) y de inyectarlas a la red PON mediante multiplexación TDM](#). Emplea la ventana de los 1490 nm.

Además, recoge las tramas de datos y voz procedentes de la red PON y las dirige a la red RTC o ISP donde corresponda en la ventana de 1310 nm.

⇒ Módulo V-OLT (Video OLT)

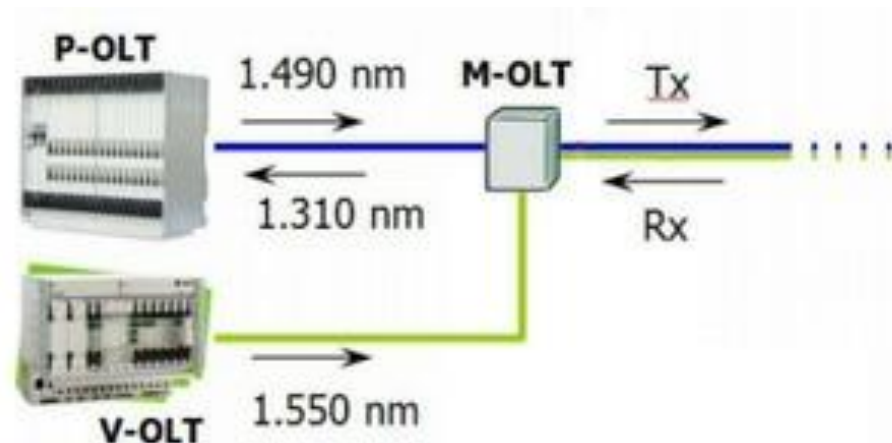
Es el módulo encargado de recoger las tramas de la red de vídeo exterior (VoD) y de inyectarlas en la red PON, empleando para ello la ventana de los 1550 nm.

También recoge las tramas de vídeo de la red PON y las dirige a la red exterior de vídeo también en la misma ventana de los 1550 nm.

⇒ Módulo M-OLT (Multiplexer OLT)

Es el encargado de realizar la multiplexación de ambos servicios suministrados por el P-OLT y V-OLT bajo una multiplexación WDM.

En la siguiente figura podemos ver los módulos que componen el OLT y las ventanas de trabajo que emplea para cada canal.



ONT

El ONT es un dispositivo óptico ubicado en el domicilio del cliente capaz de recibir y filtrar la información recibida por el OLT para entregársela al abonado. A su vez, recoge la información y petición del abonado para encapsularla y entregársela al OLT para que pueda procesarla.

Existen dos tipos de ONT:

- ⇒ H-ONT (Home ONT): es un dispositivo que se coloca en el hogar. Es el empleado en las redes FTTH.
- ⇒ B-ONT (Building ONT): es un dispositivo que se coloca en el RITI o RITU del edificio. Es el empleado para redes FTTB.

El ONT en definitiva realiza un filtrado de la información separando los servicios de voz, datos y vídeo.

Para ello incorpora dos filtros ópticos:

- ⇒ Filtro OAF (Optical Analogic Filter): es quien se encarga de obtener la señal de vídeo a la longitud de onda de los 1550 nm y la entrega al fotodiodo APD (Fotodiodo De Avalancha) para realizar la conversión de frecuencia.
- ⇒ Filtro ODF (Optical Digital Filter): es quien se encarga de obtener la señal de voz y datos a la longitud de onda de los 1490 nm y la entrega al fotodiodo digital DPD (Dinámicos).

Splitter Óptico

Se trata de un dispositivo pasivo óptico que actúa como divisor de señales, ya que recibe las señales procedentes del OLT y las entrega a las diferentes ONT conectados a él.

Sus funciones son las de:

- ⇒ Multiplexar las señales recibidas.
- ⇒ Dividir la señal de entrada entre un número de salida.
- ⇒ Combinar las señales recibidas y entregarlas por una única salida.

Además, al tratarse de un elemento pasivo no requiere alimentación.

El inconveniente principal del splitter es la pérdida de potencia óptica debido a la división de señal en múltiples salidas. Esta atenuación de la señal de salida viene dada por la siguiente expresión.

$$\text{Atenuación}_{\text{salida}} = 10 \times \log (1 / N)$$

Siendo N el número de salidas del divisor.

Existen divisores ópticos o splitters de 2, 4, 8, 16, 32 y 64 salidas.

Las redes PON son redes pasivas de fibra óptica que permiten altas tasas de transferencia para los usuarios y con ello ofrecer los nuevos servicios de Triple Play (voz, datos y televisión).

Tecnología CWDM

CWDM son las siglas de Coarse WDM, es decir, una tecnología basada en la [transmisión óptica empleando multiplexación por división en longitud de onda](#).

CWDM emplea la fibra óptica multimodo para la transmisión de datos, y se emplea para transmisiones de corta distancia a diferencia de la tecnología DWDM para largas distancias.

CWDM tuvo sus orígenes en la transmisión de vídeo (CATV). Su característica principal era que sus componentes ópticos (sobre todo láseres) eran de menor complejidad, ya que se permitía obtener un ancho de banda menor que en la tecnología DWDM y con ello su coste era menor.

Es por ello que CWDM tuvo tanta implantación para redes de corta distancia, pues su [coste](#) era [menor](#) a pesar de tener un ancho de banda alto (aunque inferior al DWDM).

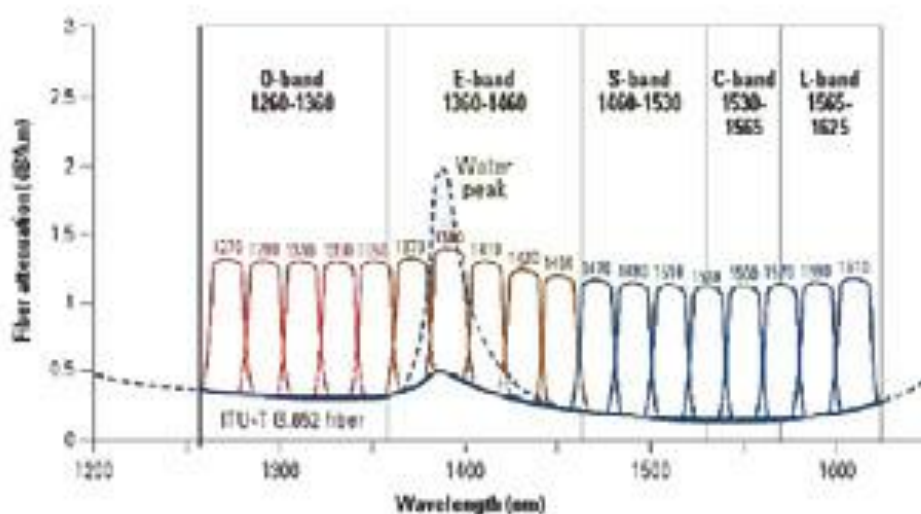
CWDM se transmite en 18 canales o longitudes de onda, con un ancho de banda de 20 nm en el intervalo de 1270 a 1610 nm en las bandas O, E, S, C y L.

Es por ello que [dispone de gran ancho de banda para distancias máximas de 120 km](#).

Esta tecnología [emplea Laser DBF](#), es decir, láseres de realimentación distribuidos, cuyos [costes de fabricación](#) son sensiblemente [menores](#) que los empleados en las redes DWDM.

CWDM se emplea en tecnología PON, FTTH, CATV y ATM, entre otros.

En la siguiente figura podemos ver los canales o longitudes de onda en los que trabaja la tecnología CWDM.



En la anterior figura se pueden observar los 18 canales de 20 nm que emplea la tecnología CWDM.

A destacar el 'pico de agua' en la longitud de los 1390 nm que provoca una alta atenuación en este canal y que por ello a menudo este canal no se emplea para las largas distancias usando esta tecnología.

Redes troncales

Las redes troncales o de transporte [constituyen el segundo nivel](#) y constituyen la [parte de la red que tiene la mayor capacidad de tráfico](#) que luego se capilariza en los niveles de acceso.

También es denominado habitualmente 'backbone'. [Se encarga en definitiva de que los servicios puedan llegar a cualquier situación geográfica](#).

La misión última de la red troncal es concentrar y distribuir todo el tráfico generado entre los diferentes usuarios que proceden de las diferentes redes de acceso.

La mayor parte de esta red se implementa sobre fibra óptica debido a la gran cantidad de tráfico que debe soportar este nivel de red.

Los motivos que llevan a usar la fibra óptica para este nivel de red son los siguientes:

- ⇒ Baja atenuación.
- ⇒ Gran ancho de banda.
- ⇒ Fácil instalación.
- ⇒ Inmunidad ante interferencias electromagnéticas.
- ⇒ Alta seguridad.
- ⇒ Integración con cualquier tipo de red.

El tipo de fibra óptica habitualmente utilizado es la fibra monomodo, ya que es la que cumple los anteriores requisitos.

En este nivel de red han aparecido las siguientes tecnologías:

- ⇒ Tecnología PDH.
- ⇒ Tecnología SDH.
- ⇒ Tecnología ATM.

Veremos a continuación y con más detalle cada una de estas tecnologías.

Las redes troncales constituyen la parte de la red del operador que concentra todo el tráfico de las diferentes redes de acceso. Constituye el backbone del operador con un gran ancho de banda y altas tasas de transferencia. Suele estar implementado con fibra óptica monomodo.

MTA (Modo de transferencia asíncrono – ATM)

ATM son las siglas de **Asynchronous Transfer Mode**. Representa una tecnología de multiplexación y conmutación de celdas o pequeños paquetes de longitud fija aprovechando las ventajas de la conmutación de circuitos y la conmutación de paquetes.

Proporciona un ancho de banda variable y ajustable según las necesidades que van desde los 2 Mbps hasta los 10 Gbps.

Las ventajas del ATM son las siguientes:

- ⇒ Capacidad garantizada y retardo de transmisión constante.
Esto lo obtiene de la conmutación de circuitos.
- ⇒ Flexibilidad y eficiencia para el tráfico ráfagueado.
Esto lo obtiene de la conmutación de paquetes.
- ⇒ Permite integrar diferentes servicios con QoS diferentes.

ATM se caracteriza por ser asíncrona y es por ello que lo hace más eficiente que otras tecnologías como TDM, que son síncronas.

Esto caracteriza a ATM en el sentido de que:

- ⇒ No hay control del flujo y corrección de errores nodo a nodo.
- ⇒ Trabajo en modo orientado a conexión.
- ⇒ Tiene una cabecera eficiente.
- ⇒ El campo de información es relativamente pequeño, lo que favorece una conmutación rápida.
- ⇒ Utiliza paquetes o slots de longitud fija.

Una red ATM está formada principalmente por los siguientes elementos:

⇒ Conmutadores o switches ATM

Son dispositivos encargados del tránsito de celdas ATM dentro de la red.

Leen la cabecera de la celda ATM y en función de esta la conmuta a uno u otro nodo.

⇒ Puntos finales ATM

Actúan como interfaz de entrada o salida a la red ATM.

Leen el contenido de la celda ATM (además de su cabecera) y en función de ello lo entrega a un destinatario final que puede ser un servidor, conmutador LAN, router, etc.

Dentro de los conmutadores ATM existen dos tipos:

- Conmutador UNI (User Network Interface): conecta destinatarios finales como router, servidores, etc., a un conmutador ATM.
- Conmutador NNI (Network to Network Interface): conecta dos conmutadores ATM.

Estos conmutadores, tanto UNI como NNI pueden ser privados o públicos. En el primer caso su propiedad y ubicación están en el cliente final (es privado) y en el segundo caso su propiedad está ubicada en el operador ATM.

La celda ATM es una trama de formato fijo de 53 octetos o bytes de donde los primeros 5 bytes se corresponden a la cabecera (cabecera pequeña para rápida conmutación) y los 48 bytes restantes es información del usuario.

En la siguiente figura podemos ver el formato de trama de la celda ATM.

Cabecera	Información
5 bytes	48 bytes

El formato de cabecera de la celda ATM puede admitir dos formatos:

- ⇒ Formato UNI: que es el empleado por los conmutadores UNI.
- ⇒ Formato NNI: que es el empleado por los conmutadores NNI.

El formato de cabecera de celda ATM UNI es el que se muestra en la siguiente figura:

----- 8 bits -----		
Campo de Control de flujo (GFC)	Identificador de Camino virtual (VPI)	
Identificador de Camino virtual (VPI)	Identificador de Canal virtual (VCI)	
Identificador de Canal virtual (VCI)		
Identificador de Canal virtual (VCI)	Identificador de tipo de carga (PTI)	Prioridad de pérdida de celda (CLP)
Campo de control de errores (HEC)		

En ATM se define lo que se denomina canal virtual (VC) o camino virtual (VP).

Un canal virtual (VC) es la multiplexación de varios flujos de información entre dos nodos ATM.

Un camino virtual (VP) es la multiplexación de varios canales virtuales en la red ATM.

Por tanto, un camino virtual (VP) es un conjunto de canales virtuales (VC).

Es por ello que en la cabecera ATM existen los campos de identificación del canal virtual (VCI) y el identificador de camino virtual (VPI) para que el conmutador pueda saber a qué nodo debe dirigir la celda ATM.

El campo de control de flujo (GFC) es un campo encargado del control y gestión del tráfico, aunque en la práctica apenas se emplea.

El campo de identificador de tipo de carga (PTI) indica si la celda contiene datos de usuarios y datos de control.

El campo de prioridad de pérdida de celda indica si dicha celda ATM puede ser descartada en caso de congestión de la red.

El campo de control de errores (HEC) es un campo checksum.

En cambio el formato de la cabecera de celda ATM NNI tiene un formato ligeramente diferente al anterior.

----- 8 bits -----		
Identificador de Camino virtual (VPI)		
Identificador de Camino virtual (VPI)	Identificador de Canal virtual (VCI)	
Identificador de Canal virtual (VCI)		
Identificador de Canal virtual (VCI)	Identificador de tipo de carga (PTI)	Prioridad de pérdida de celda (CLP)
Campo de control de errores (HEC)		

Es decir, en este caso solo varía los primeros bits de la cabecera, eliminando el campo de control de flujo (GFC) y ampliando el del identificador de camino virtual (VPI).

PDH (Jerarquía Digital Plesiócrona – PDH)

PDH son las siglas de Plesichronous Digital Hierarchy y representa [unas de las primeras tecnologías troncales que aparecieron en la década de los 70.](#)

Esta tecnología se empleaba en aquella época para la transmisión de las señales telefónicas de voz que se digitalizaban según la técnica de muestreo PCM (modulación por impulsos codificados - Pulse Code Modulation).

PDH se caracterizaba por emplear transmisiones digitales que permitían tasas de transferencia hasta 565 Mbps fraccionados en múltiplos de 2 Mbps.

Esta tecnología creaba canales múltiplos de 2 Mbps porque un canal de 2 Mbps era lo que cabía por transmitir 30 canales de voz telefónico muestreado con PCM.

Esta tecnología se basaba en un entrelazado de bit que obligaba a sincronizar las distintas fuentes de información a transmitir.

Como las diferentes fuentes no tenían por qué transmitir a las mismas tasas de transferencia sobre este sistema se crearon diferentes mecanismos de control para absorber esas asincronías de las fuentes (de ahí el término de plesiócrona que indica casi síncrona).

Este fue uno de los motivos por los que esta tecnología fue sustituida por otras como la SDH, que no presentaba este problema y en la que además las redes cada vez más eran más digitales.

Veamos un ejemplo.

Justifique mediante cálculos el número de canales de voz digitalizadas (30) que permitía transmitir un canal de 2 Mbps empleando la tecnología PDH.

Solución:

El servicio telefónico básico tiene un ancho de banda que va desde los 300 Hz hasta los 3.400 Hz. Es por ello que su ancho de banda es:

$$\text{Ancho de banda} = 3.400 - 300 = 3.100 \text{ Hz.}$$

Si digitalizamos la señal y con objeto de no perder información, aplicamos el teorema de muestreo que nos dice que:

$$F_{\text{muestreo}} \geq 2 \times \text{ancho de banda.}$$

Es decir, en nuestro caso particular quedaría como:

$$F_{\text{muestreo}} \geq 2 \times 3.100 \text{ Hz} = 6.200 \text{ Hz.}$$

Redondeamos y aplicamos una frecuencia de muestro de 8 Khz.

Para la cuantificación usaremos 8 bits, es decir, 256 valores. Con ello, la tasa binaria para un canal de voz es:

$$\text{Tasa binaria} = F_{\text{muestreo}} \times \text{Nº bits cuantificación} = 8 \text{ Khz} \times 8 = 64 \text{ Kbps.}$$

Si multiplicamos esa tasa binaria por 30 canales de voz, nos sale que:

$$\text{Caudal de 30 canales de voz digitalizada} = 30 \times 64 \text{ Kbps} = 1920 \text{ Mbps} \approx 2 \text{ Mbps.}$$

Aquí está la justificación de cómo en un canal de 2 Mbps podemos transmitir 30 canales de voz de señal de voz telefónico digitalizada.

Veamos otro ejemplo.

En Estado Unidos y Japón también se empleó la tecnología PDH para la transmisión de caudales de señal de voz telefónica digitalizada, pero con canales de 1,5 Mbps.

Indique y justifique cuántos canales de voz digitalizada podemos transmitir por ese canal de 1,5 Mbps.

Solución:

Partimos del mismo punto de que el servicio telefónico tiene el mismo ancho de banda y se digitaliza con el mismo número de bits de cuantificación, es decir:

$$\text{Ancho de banda} = 3.400 - 300 = 3.100 \text{ Hz.}$$

Muestreamos aplicando el teorema de muestreo.

$$F_{\text{muestreo}} \geq 2 \times 3.100 \text{ Hz} = 6.200 \text{ Hz} \approx 8 \text{ Khz.}$$

Para la cuantificación usaremos 8 bits, es decir, 256 valores, luego con ello, la tasa binaria para un canal de voz es:

$$\text{Tasa binaria} = F_{\text{muestreo}} \times \text{Nº bits cuantificación} = 8 \text{ Khz} \times 8 = 64 \text{ Kbps.}$$

Si disponemos de un caudal de 1,5 Mbps, el número de canales de voz telefónica digitalizada es de:

$$\text{Número de canales de voz} = (1,5 \text{ Mbps}) / (64 \text{ kbps}) = (1536 \text{ Kbps}) / (64 \text{ kbps}) = 24 \text{ canales}$$

Luego para el caso de EEUU y Japón para un canal de 1,5 Mbps sólo podemos transmitir 24 canales de señal de voz telefónica digitalizada.

Veamos un ejemplo más.

Sobre el ejemplo anterior, en comunicaciones militares es admisible pérdida de calidad de la información telefónica transmitida (sobre todo timbre) a cambio de obtener un mayor número de canales a transmitir por el caudal siempre y cuando la comunicación sea audible e interpretable.

Es por ello que se admite un menor nivel de cuantificación en la digitalización de la señal.

Si sobre el canal de transmisión de 1,5 Mbps, aplicamos un nivel de cuantificación de 4 niveles, indique y calcule cuantos canales de voz telefónica digitalizada podemos transmitir.

Solución:

Partimos del mismo punto de que el servicio telefónico tiene el mismo ancho de banda y donde se le aplica el teorema de muestreo.

$$\text{Ancho de banda} = 3.400 - 300 = 3.100 \text{ Hz.}$$

$$F_{\text{muestreo}} \geq 2 \times 3.100 \text{ Hz} = 6.200 \text{ Hz} \approx 8 \text{ KHz.}$$

Dado que cuantificamos con 4 niveles, eso indica que necesitamos sólo 2 bits para cuantificar cada muestra de la señal, es decir:

$$2^x = 4 \text{ niveles} \Rightarrow x = 2 \text{ bits para cuantificar.}$$

Luego la tasa binaria de cada canal de voz digitalizada queda en:

$$\text{Tasa binaria} = F_{\text{muestreo}} \times \text{Nº bits cuantificación} = 8 \text{ KHz} \times 2 = 16 \text{ Kbps.}$$

Si disponemos de un caudal de 1,5 Mbps, el número de canales de voz telefónica digitalizada es de:

$$\text{Número de canales de voz} = (1,5 \text{ Mbps}) / (16 \text{ kbps}) = (1536 \text{ Kbps}) / (16 \text{ kbps}) = 96 \text{ canales}$$

Es decir, hemos multiplicado por 4 el número de señales de voz telefónica digitalizada que podemos transmitir por un canal de 1,5 Mbps.

Eso sí, se ha perdido parte de calidad de cada señal

JDS (Jerarquía Digital Síncrona – SDH)

SDH son las siglas de Synchronous Digital Hierachy. Representa a una tecnología troncal de transporte que surgió de la evolución de su antecesora, la tecnología PDH, corrigiendo los problemas que presentaba este último.

Esta tecnología tuvo gran expansión a raíz de la digitalización de todas las comunicaciones de voz y de datos ya que alcanzaban tasas de transferencia desde los 155 Mbps hasta los 40 Gbps con canales de 2 Mbps.

La gran ventaja de SDH es su compatibilidad con diferentes tecnologías incluidas PDH y ATM, lo que permitía una gran interoperabilidad de equipos y con ello soportaba una amplia variedad de servicios desde voz, datos, redes alquiladas, redes de televisión, etc.

Estructura de la trama

SDH basa su funcionamiento en una estructura fija denominada trama STM-1, que es un slot de tiempo con una duración de 125 µseg y se corresponde con una matriz de 9 filas y 270 columnas cuyos elementos son octetos de 8 bits. La velocidad de esta trama es de 155 Kbps.

La estructura de una trama STM-1 tiene el siguiente aspecto:



Puede observarse cómo la trama se dispone en 9 filas por 270 columnas.

Esta trama presenta las siguientes partes:

– Los primeros 9 bytes de cada fila se denominan tara de cabecera u Overhead.

– Los restantes bytes de cada fila se denominan carga útil.

La tara de cabecera se divide en tres partes:

– El R-SOH que contiene información de gestión y control entre los repetidores de la tecnología SDH y cuyas

funciones son, entre otras, las de:

– Señal de alineamiento.

– Monitorización de errores.

– Señalización de conmutación automática.

– Estado de sincronización.

– Puntero de unidad administrativa que contiene información de gestión de la red.

–El M-SOH que contiene información de la monitorización de la calidad de transmisión. Se compone del Contenedor Virtual (VC) que es la carga útil que viaja sin cambios en la red y de algunos bytes que se añaden en los puntos de terminación del servicio.

La transmisión de la trama se realiza fila a fila hasta completar las 9 filas que componen la trama con una orientación de izquierda a derecha y de arriba abajo.

Velocidades

PDH está basado en una trama básica denominada SMT-1, que es la descrita anteriormente.

Pero admite diferentes modos de transmisión de esta trama dando lugar a diferentes modalidades de SMT, que vienen recogidas en la siguiente tabla:

Tipo de SMT	Tasa de transferencia
SMT-1	155 Mbps
SMT-4	622 Mbps
SMT-16	2,5 Gbps
SMT-64	10 Gbps
SMT-256	40 Gbps

El valor que se adjunta junto al SMT indica el factor de múltiplo sobre la trama SMT que se aplica y que da lugar a las diferentes tasas de transferencia.

Ventajas y desventajas respecto a PDH

La SDH surgió como evolución de la PDH que corregía los problemas que este último presentaba.

En este sentido las ventajas de SDH frente a PDH son las siguientes:

- La multiplexación en SDH es mucho más rápida y eficiente que en PDH.
- Gracias a lo anterior, se obtienen tasas de transferencia mucho más altas.
- Permite transmitir tráfico de numerosos formatos, caudal y tipos, por lo que se hace flexible a cualquier tipo de servicio de telecomunicación.
- Es perfectamente compatible e interoperable con otras tecnologías como ATM o PDH.
- Permite su integración con tecnologías ópticas.

Pero SDH presenta algunas desventajas como:

- Precisa de un exacto sincronismo con todos los nodos SDH. Esto encarece enormemente la tecnología al precisar de un reloj de sincronismo altamente estable.
- La cabecera de la trama es mayor que en PDH y esto implica una mayor ineficiencia en la tasa binaria de datos a transmitir.

Mecanismos de codificación y cifrado de la información

La codificación es una técnica usada en comunicaciones que consiste en la conversión de un sistema de datos en otro sistema de datos.

Cuando queremos transmitir datos por un medio de transmisión no podemos transmitirlo tal cual, sino que debemos convertirlo en otro formato adaptado al medio de transmisión para que pueda ser transmitido. Esto es codificar los datos.

La codificación la realizan unos dispositivos denominados CODEC (Codificación- Decodificación) que se encargan (antes del envío del datos) de codificarlos (en la parte del emisor) y de decodificarlos (en la parte del receptor).

Evidentemente emisor y receptor deben emplear el mismo código de codificación para codificar y decodificar la información transmitida.

Dado que existen dos tipos de transmisión: analógica y digital, se emplean técnicas de codificación diferentes para cada una de ellas.

Así encontramos:

- Técnicas de codificación de datos digitales.

Para una transmisión digital.

- Técnicas de codificación de datos analógicos.

Para una transmisión analógica.

Veremos a continuación con más detalle cada una de estas técnicas.

La codificación es una técnica de tratamiento de señales que consiste en convertir la señal de un formato a otro con objeto de adaptarla al canal, mejorar su transmisión, aumentar la eficiencia de las infraestructuras, etc.

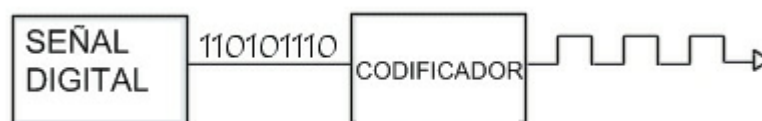
En una transmisión digital lo que se transmite son 0 y 1.

Pero estos 0 y 1 no se pueden transmitir tal cual sino que deben ser codificados (traducidos) a estados de una señal, como por ejemplo:

- Ausencia o presencia de señal.
- Cambio del voltaje de una señal.
- Diferencias de voltajes con respecto a una señal de referencia.

Esta traducción de 0 y 1 a estados de una señal es lo que realiza un CODEC. Emplea para ello diferentes técnicas de codificación de datos digitales.

Lo vemos en la siguiente figura:



En la anterior figura podemos ver cómo el flujo de bits tras pasar por el códec se ha transformado en flujo de estados de una señal, que es lo que realmente se transmite.

El receptor detectará estos cambios de señal y, empleando la misma técnica de codificación, será capaz de traducirlo en señales binarias de 0 y 1 tal y como se enviaron por parte del emisor.

Existen numerosas técnicas de codificación de señales digitales, destacando las siguientes:

- Codificación NRZ.
- Codificación NRZI.
- Codificación Manchester.
- Codificación de Miller.
- Codificación bipolar.

A continuación veremos con más detalle el funcionamiento de cada una de ellas.

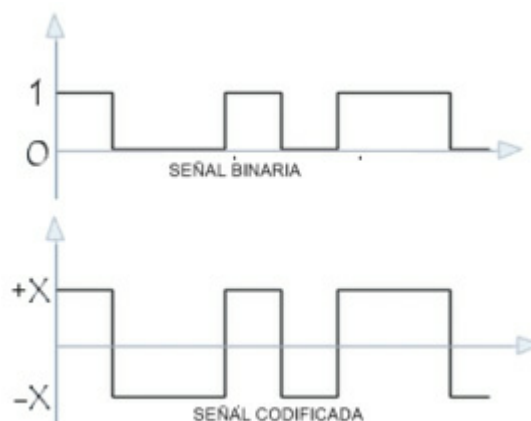
Codificación NRZ

Se trata de una de las técnicas de codificación de señales digitales más sencillas y simples, ya que fue de las primeras que aparecieron.

Se trata de una codificación de no retorno a cero (Not Return zero) y consiste en transmitir un valor de amplitud de señal X cuando queremos transmitir un 1 y un valor de amplitud de señal negativo $-X$ cuando queremos transmitir un 0.

Nunca hay un valor 0 voltaje de señal en el canal (por eso lo de no retorno a cero).

En la siguiente figura puede verse cómo funciona esta codificación:



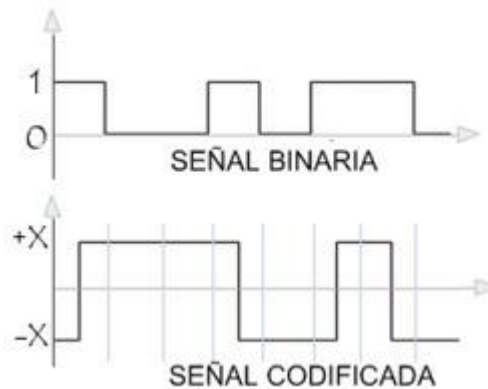
La figura de arriba representa el flujo de bits a transmitir y cuando queremos transmitir un 1 enviamos una señal con voltaje X y cuando queremos transmitir un 0 enviamos la misma señal y con el mismo voltaje pero en negativo.

Así el receptor en función del voltaje recibido (X o $-X$) decodificará si se ha enviado un 1 o un 0.

Codificación NRZI

Se trata de una codificación de datos digitales en la cual la señal que se transmite por el canal cambia de estado cuando queremos transmitir un 1 y permanece en el mismo estado cuando transmitimos un 0.

Lo vemos en la siguiente figura:



La figura de arriba representa el flujo de bits a transmitir. Cuando queremos transmitir un 1 provocamos un cambio de estado en la señal. Si en cambio queremos enviar un 0 no cambiamos el estado de la señal transmitida.

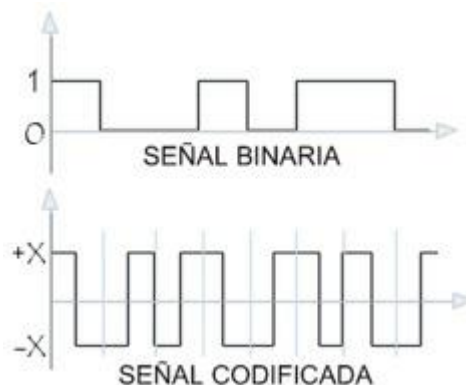
Así, el receptor, en función de los cambios de estado de la señal recibida, va decodificando si el emisor ha enviado un 0 o un 1.

Codificación de Manchester

Se trata de una codificación de datos digitales en la cual se introduce una transición del estado cada vez que transmitimos un bit. Es por ello que también se le denomina codificación en dos fases.

Esta codificación podría equivaler a realizar una OR exclusiva (XOR) con la señal de reloj de la transmisión digital.

En la siguiente figura podemos ver cómo funciona esta codificación:



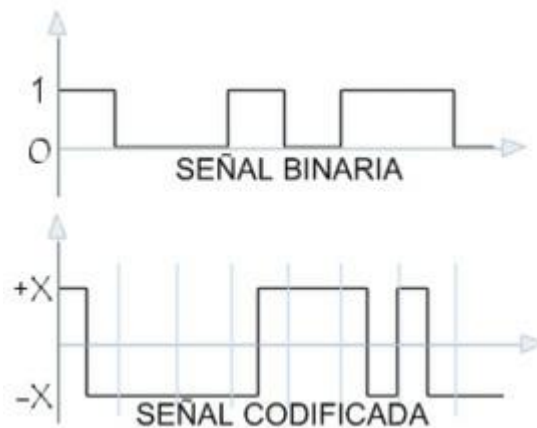
La figura de arriba representa el flujo de bits a transmitir y cuando queremos transmitir un 1 enviamos un pulso con una transmisión de X a $-X$. En cambio, cuando queremos enviar un 0 enviamos un pulso con una transición de $-X$ a X .

Así, el receptor, en función de estos pulsos y sus transiciones, va decodificando si el emisor ha enviado un 0 o un 1.

Codificación de Miller

Se trata de una codificación de datos digitales muy similar al código Manchester, pero en este caso la transición en medio del intervalo solo se produce cuando se transmite un 1.

Lo vemos en la siguiente figura:



La figura de arriba representa el flujo de bits a transmitir. Cuando queremos transmitir un 1 la señal incluye una transición en medio del intervalo. En cambio, cuando queremos enviar un 0 no provocamos ningún cambio de la señal.

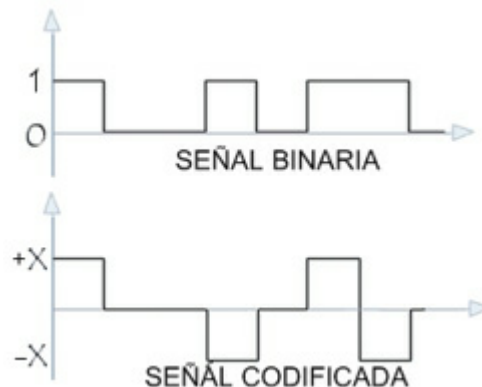
Esta técnica permite un mayor índice de datos a transmitir.

Codificación Bipolar

Se trata de una codificación de datos digitales que, a diferencia de los anteriores, presenta tres estados (los anteriores solo 2).

En esta codificación, cuando se quiere transmitir un 0, se envía una señal con valor 0 y cuando se quiere enviar un 1 se envía la misma señal con valor X y $-X$ alternativamente.

En la siguiente figura se puede ver cómo funciona esta técnica de codificación:



La figura de arriba representa el flujo de bits a transmitir. Cuando queremos transmitir un 1 la señal incluye una transición en medio del intervalo. En cambio, cuando queremos enviar un 0 no provocamos ningún cambio de la señal.

Esta técnica permite un mayor índice de datos a transmitir.

Ejemplo:

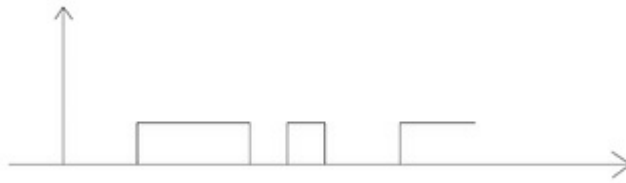
Queremos transmitir la siguiente secuencia de bits por un medio digital en la cual se va a emplear la codificación NRZ.



Dibuje la secuencia de bits codificada que se envía al canal.

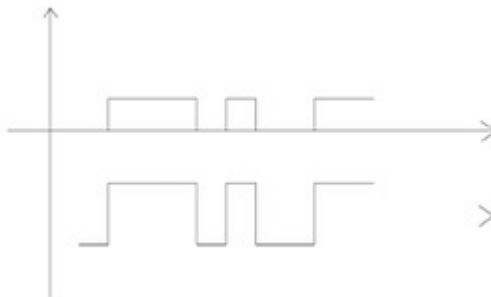
Solución:

Dibujemos en primer lugar en un diagrama de tiempos los bits originales que se quieren enviar, es decir:



La codificación NRZ consiste en transmitir un pulso con amplitud $-X$ cuando se quiere transmitir un 0 y un pulso con amplitud $+X$ cuando se quiere transmitir un 1.

Superponemos en el mismo diagrama de tiempo los bits originales a enviar y la señal codificada según codificación NRZ.



El hecho de que siempre se envíe una señal (negativa o positiva) favorece que el canal tenga siempre señal y sea más inmune al ruido.

Por otro lado, su espectro es mayor, por lo que requiere más ancho de banda.

Veamos otro **ejemplo**.

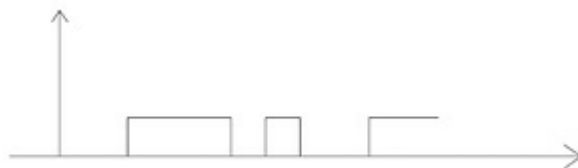
Queremos transmitir la misma secuencia de bits del ejemplo anterior pero ahora empleando la codificación bipolar.

0	1	1	1	0	1	0	0	1	1
---	---	---	---	---	---	---	---	---	---

Dibuje la secuencia de bits codificada que se envía al canal.

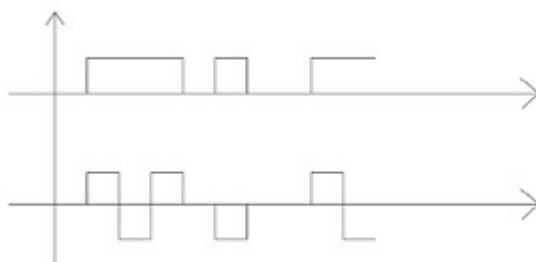
Solución:

Al igual que en el ejemplo anterior, debemos dibujar en un diagrama temporal los bits originales que se quieren enviar:



En la codificación bipolar, cuando se quiere enviar un 0 no se transmite señal y cuando se quiere enviar un 1 se envía alternativamente una señal con amplitud X y $-X$.

Superponemos en el mismo diagrama de tiempo los bits originales a enviar y la señal codificada según codificación bipolar.



Esta es una de las codificaciones que se caracteriza por tener tres estados: ausencia de señal, señal positiva y señal negativa.

Los datos analógicos se caracterizan porque pueden adoptar un número infinito de valores.

Los datos analógicos pueden ser transmitidos por medios analógicos o digitales, ya que el formato que tenga el dato es independiente del medio de transmisión por el que se quiere transmitir.

Para ello, se emplean técnicas de codificación que permiten traducir estos datos analógicos a estados de una señal adaptada al medio de transmisión (analógica o digital) y así poder transmitir los datos analógicos. Existen diferentes técnicas de codificación para datos analógicos. Depende de si lo transmitimos por un canal digital o por un canal analógico.

Si lo queremos transmitir por un canal digital, lo primero que se realiza es una digitalización de la señal (pasamos el dato de analógico a digital) y obtenemos un flujo de bits sobre el cual podemos emplear cualquiera de las técnicas de codificación de datos digitales descritos anteriormente.

También podemos transmitir dichos datos analógicos empleando un medio analógico (transmisión analógica) y para ello existen, entre otras, las siguientes técnicas de codificación:

- Modulación en amplitud.
- Modulación en frecuencia.
- Modulación en fase.

Técnicas de modulación existen muchas, aunque las tres anteriores son las más sencillas y básicas, ya que el resto suelen ser combinaciones de las tres anteriores.

Veremos a continuación cómo funciona cada una de estas técnicas de codificación de datos analógicos.

Hemos comentado anteriormente que los datos analógicos los podemos transmitir por un canal digital empleando la digitalización de la señal.

Veamos en qué consiste esta técnica.

La digitalización consiste en ir tomando muestras de la señal analógica en intervalos regulares y cuantificando la amplitud en cada una de las muestras tomadas.

El intervalo con que se toman las muestras sigue el Teorema de Muestreo o Teorema de Nyquist, que establece que muestreando a una frecuencia al menos el doble que la máxima frecuencia de la señal original se puede obtener una señal digitalizada sin pérdida de la información.

$$\text{Teorema del Muestreo Frecuencia}_{\text{muestreo}} \geq 2 \times \text{frecuencia}_{\text{máxima señal}}$$

Si cuantificamos (tomamos valores discretos de la amplitud de cada muestra), obtenemos así un flujo de bits que representa la digitalización en cada tiempo de la señal analógica y con ello ya podemos enviarlo por un canal digital empleando cualquiera de las técnicas de codificación digital.

En la siguiente figura podemos ver cómo funciona la digitalización de una señal analógica:

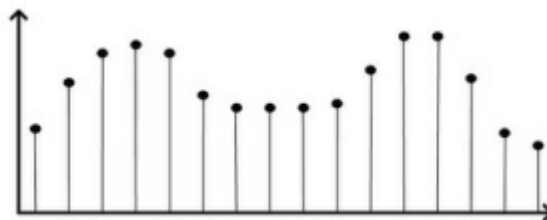


En la figura de arriba podemos ver el ejemplo de una señal analógica, la cual puede adoptar cualquier tipo de valor de amplitud.

Con la digitalización tomamos valores a intervalos regulares a frecuencias f_1 , $2 \times f_1$, $3 \times f_1$, $4 \times f_1$, etc. y en cada uno de ellos tomamos el valor que le corresponde en la señal analógica.

Si luego dibujamos la señal a partir de sus muestras podemos obtener una representación digitalizada de la señal analógica.

Lo vemos en la siguiente figura:



Pero la digitalización implica también una cuantificación de los valores de la señal muestreada, ya que el valor tomado puede ser infinito.

Para ello, cuantificamos la señal, es decir, tomamos el valor máximo y mínimo de la señal analógica y podemos cuantificarlo en un número discreto de valores, por ejemplo, 2, 4, 8, 16, etc.

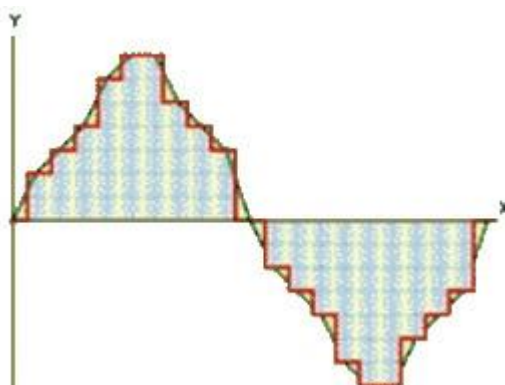
Si, por ejemplo, tomamos un número discreto de 4 valores, eso indica que la señal digitalizada solo podrá tener uno de esos posibles valores y el valor tomado es el valor más cercano al valor de la señal analógica. En cada una de las muestras se toman los valores de la señal analógica, pero la cuantificación implica que debemos ajustarnos a uno de los 4 valores discretos posibles y esos valores son los que se transmiten.

Esos cuatro valores posibles se pueden representar mediante 2 bits ($2^2=4$ valores).

A partir de estos cuatro valores a intervalos regulares podemos obtener la señal digitalizada.

Si en vez de cuatro valores cuantificamos 8 valores ello implica que podemos obtener una señal digitalizada más fiel a la señal analógica original pero ello implica transmitir más bits. Ocho valores implica tres bits ($2^3=8$ valores).

Lo vemos en la siguiente figura.

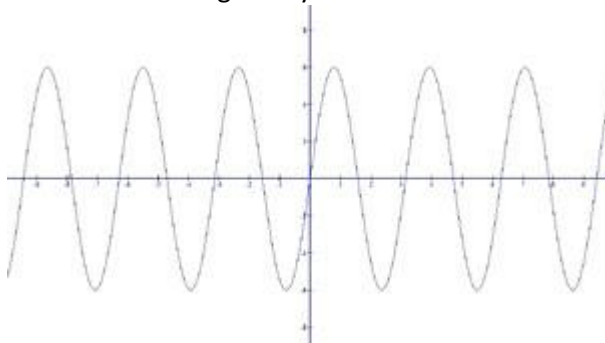


Lo mismo ocurre si muestreamos a intervalos de mayor frecuencia. Obtenemos una señal digitalizada más fiel a la señal analógica original, pero ello implica un mayor número de bits a transmitir.

El teorema del Muestreo nos indica que con una frecuencia de muestreo de al menos el doble de la frecuencia máxima de la señal analógica es suficiente para no perder información.

Vemos un **ejemplo**.

Si queremos digitalizar la siguiente señal analógica cuya frecuencia máxima es de 1 Mhz.



Calcula:

- Frecuencia mínima de muestreo para no perder información.
- Número de bits de cuantificación necesarios si tomamos un valor de 16 posibles.

—Cuántos bits se deben transmitir si tomamos un total de 10 muestras con el nivel de cuantificación anterior.

Solución:

Resolvemos cada uno de los apartados:

—Dado que la frecuencia máxima de la señal analógica es de 1 Mhz, por el Teorema del muestreo la frecuencia del muestreo debe ser:

$$\text{Frecuencia}_{\text{muestreo}} \geq 2 \times \text{frecuencia}_{\text{máxima señal}}$$

Por lo que si tomamos el mínimo, será de:

$$\text{Frecuencia}_{\text{muestreo}} = 2 \times \text{frecuencia}_{\text{máxima señal}} = 2 \text{ Mhz}$$

—Si la señal cuantificada solo puede adoptar uno de los 16 valores posibles, el número de bits necesarios será de:

$$2^x = 16 \text{ bits } x = 4$$

—Si se envían 10 muestras de la señal donde cada una precisa de 4 bits por la cuantificación, el flujo de bits que se transmite será de:

$$\text{Flujo de bits} = 10 \text{ muestras} \times 4 \text{ bits} = 40 \text{ bits.}$$

Es decir, en total se envían 40 bits. En realidad se suelen enviar más bits, ya que en una transmisión se suelen incluir bits de redundancia y control de errores.

Vemos otro ejemplo.

Dada la anterior señal analógica la cual se quiere digitalizar empleando una frecuencia de muestreo de 5 Mhz y 3 bits de cuantificación, se pide lo siguiente:

—Dibuje la señal digitalizada.

—¿Se pierde información por la digitalización?

—¿Cuántos valores son posibles en la cuantificación?

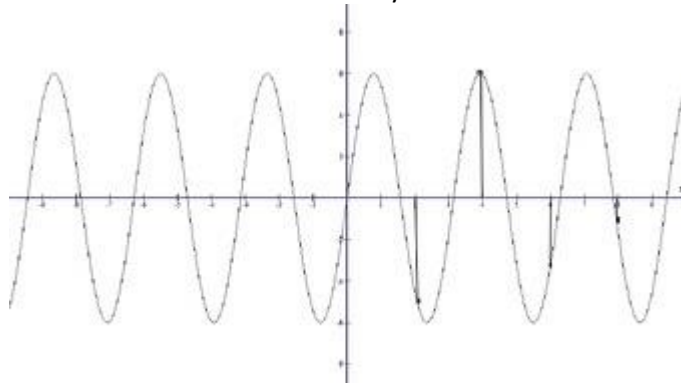
Solución:

Resolvemos cada uno de los apartados:

—La digitalización consiste en formar muestras a intervalos de 2 Mhz, 4 Mhz, 6 Mhz, etc. y donde con tres bits de cuantificación nos dan $2^3 = 8$ valores posibles.

Dado que la máxima y mínima amplitud de la señal es de 16 V, los valores posibles serán de -16, -12, -8, -4, 4, 8, 12 y 16.

Dibujamos la señal digitalizada con esos valores discretos y muestreada a esas frecuencias.



—El que se pierda información o no depende de si cumple el Teorema del Muestreo.

$$\text{Frecuencia}_{\text{muestreo}} \geq 2 \times \text{frecuencia}_{\text{máxima señal}}$$

Dado que la frecuencia de la señal es de 1 Mhz y muestreamos a 5 Mhz, cumple el teorema, luego no se pierde información en la digitalización.

—El número de posibles valores en la cuantificación es:

$$2^3 = 8 \text{ valores}$$

Y cuyos valores posibles son los indicados anteriormente en función del valor máximo y mínimo de pico de la señal analógica: -16, -12, -8, -4, 4, 8, 12 y 16.

Si ahora queremos transmitir un dato o señal analógica por un canal analógico, empleamos otras técnicas de codificación.

En este caso se emplea la técnica de modulación, que consiste en transformar una señal de mayor frecuencia (señal portadora) en función de la señal analógica a transmitir (señal moduladora).

El resultado es una señal modulada de mayor frecuencia con la que se consiguen dos cosas:

- Adaptar la señal transmitida al canal.
- Permitir la multiplexación de canales en el medio de transmisión.

Con los dos logros anteriores se obtiene una mayor eficiencia en la transmisión.

La modulación puede ser de tres tipos:

- Modulación en amplitud.
- Modulación en frecuencia.
- Modulación en fase.

Vemos cada uno de ellas con más detalle a continuación.

Son muchos los tipos de modulación existentes en comunicaciones. Todas son variantes y combinaciones de las tres modulaciones básicas: modulación en amplitud, en frecuencia y en fase.

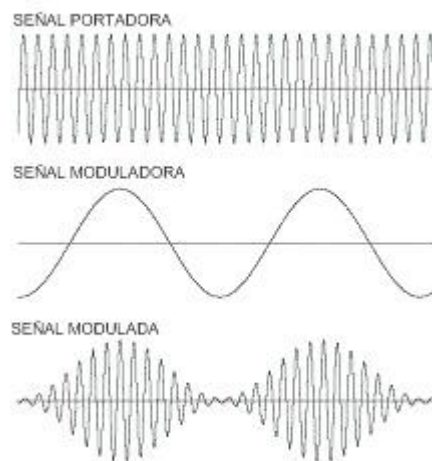
Modulación en amplitud

Es una técnica de codificación de señales analógicas basada en modular (transformar) la amplitud de la señal portadora en relación a la señal a transmitir (moduladora).

El resultado es una señal modulada en amplitud que mantiene la misma forma que la señal moduladora pero es de mayor frecuencia.

La frecuencia de la señal portadora es aquella que está adaptada al medio de transmisión (generalmente una frecuencia mucho mayor que la señal moduladora).

Lo vemos en la siguiente figura:



La operación que permite la modulación en amplitud es una operación de multiplicación de señal moduladora y señal portadora.

Las expresiones matemáticas que representan la modulación en amplitud se basan en la multiplicación de señales sinusoidales.

Así si, por ejemplo, tenemos una señal moduladora expresada por la siguiente ecuación:

$$S_{\text{moduladora}}(t) = A_{\text{moduladora}} \times \cos(w_{\text{moduladora}} \times t)$$

Siendo $A_{\text{moduladora}}$ la amplitud de la señal moduladora y $w_{\text{moduladora}}$ la frecuencia de la señal moduladora (téngase en cuenta que $S_{\text{moduladora}}(t)$ es la señal analógica a transmitir).

Y la señal portadora expresada por la siguiente expresión:

$$S_{\text{portadora}}(t) = A_{\text{portadora}} \times \cos(w_{\text{portadora}} \times t)$$

La señal moduladora viene dada por la multiplicación de ambas señales, es decir:

$$S_{\text{modulada}}(t) = S_{\text{moduladora}}(t) \times S_{\text{portadora}}(t) = A_{\text{moduladora}} \times \cos(w_{\text{moduladora}} \times t) \times A_{\text{portadora}} \times \cos(w_{\text{portadora}} \times t)$$

La modulación en amplitud es una de las modulaciones más sencillas y básicas de las existentes.

Su principal problema es que es sensible al ruido, ya que ésta se acopla a la amplitud de la señal modulada y puede provocar errores en la recepción.

La radio AM (Amplitud Modulada) es un ejemplo del uso de esta técnica de modulación.

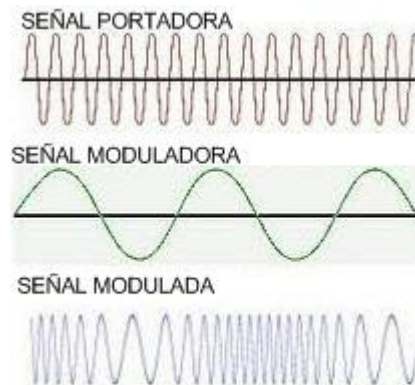
Modulación en frecuencia

Es una técnica de codificación de señales analógicas basada en modular (transformar) la frecuencia de la señal portadora en relación a la señal a transmitir (moduladora).

El resultado es una señal modulada en frecuencia en función de la amplitud de la señal moduladora.

Así, a mayor amplitud de la señal moduladora, más frecuencia en la señal modulada y a menor amplitud de la señal moduladora, menos frecuencia en la señal modulada.

Lo vemos en la siguiente figura:



A diferencia de la modulación en amplitud, la señal modulada mantiene la misma amplitud.

Al igual que la modulación en frecuencia, se basa en la multiplicación de una señal portadora con una señal moduladora, es decir:

$$S_{\text{moduladora}}(t) = A_{\text{moduladora}} \times \cos(w_{\text{moduladora}} \times t)$$

$$S_{\text{portadora}}(t) = A_{\text{portadora}} \times \cos(w_{\text{portadora}} \times t)$$

$$S_{\text{modulada}}(t) = S_{\text{moduladora}}(t) \times S_{\text{portadora}}(t) = A_{\text{moduladora}} \times \cos(w_{\text{moduladora}} \times t) \times A_{\text{portadora}} \times \cos(w_{\text{portadora}} \times t)$$

Y donde aplicando operaciones de cosenos podemos obtener una señal modulada con una frecuencia que varía en función de la amplitud de la señal moduladora.

Un ejemplo de aplicación de este tipo de modulación es la radio FM (Frequency Modulation).

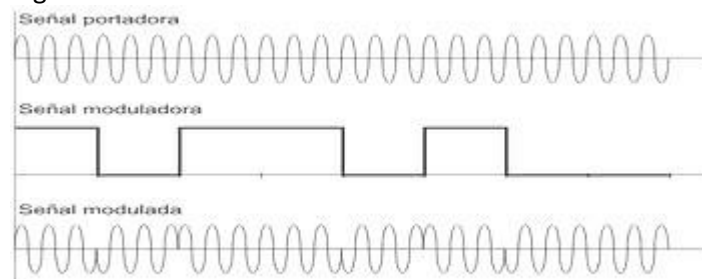
Modulación en fase

Es una técnica de codificación de señales analógicas basada en modular (transformar) la fase de la señal portadora en relación a la señal a transmitir (moduladora).

El resultado es una señal modulada en fase en función de la amplitud de la señal moduladora.

Así, cada vez que hay un cambio en la amplitud de la señal moduladora hay un cambio de fase en la señal portadora.

Lo vemos en la siguiente figura:



La amplitud y frecuencia de la señal modulada, a diferencia de las modulaciones anteriores, no varía.

Al igual que las modulaciones anteriores, se basa en la multiplicación de una señal portadora con una señal moduladora, es decir:

$$S_{\text{moduladora}}(t) = A_{\text{moduladora}} \times \cos(w_{\text{moduladora}} \times t + \phi_{\text{moduladora}})$$

$$S_{\text{portadora}}(t) = A_{\text{portadora}} \times \cos(w_{\text{portadora}} \times t + \phi_{\text{portadora}})$$

$$S_{\text{modulada}}(t) = S_{\text{moduladora}}(t) \times S_{\text{portadora}}(t) = A_{\text{moduladora}} \times \cos(w_{\text{moduladora}} \times t + \phi_{\text{moduladora}}) \times A_{\text{portadora}} \times \cos(w_{\text{portadora}} \times t + \phi_{\text{portadora}})$$

Y donde aplicando operaciones de cosenos podemos obtener una señal modulada con una fase que varía en función de la amplitud de la señal moduladora.

Existen numerosas técnicas de modulación, pero la gran mayoría suelen ser combinaciones de las anteriormente descritas.

Veamos un **ejemplo**.

Dadas las siguientes señales moduladas recibidas en un receptor, indique qué tipo de modulación se ha empleado en cada una de ellas.

Solución:

En el análisis de las señales moduladas podemos observar que en la figura 1 la señal mantiene su amplitud y fase constante, variando únicamente su frecuencia, por lo que podemos concluir que en la Figura 1 se ha empleado modulación en frecuencia.

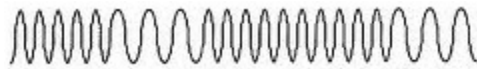


Figura 1 = Modulación en frecuencia.

En la figura 2 podemos observar, en cambio, que la amplitud de la señal va variando constantemente manteniendo invariable su frecuencia y su fase, por lo que podemos deducir que la modulación empleada en el emisor ha sido de una modulación en amplitud.

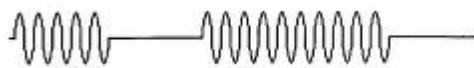


Figura 2 = Modulación en amplitud.

Por último, en la figura 3 podemos observar que la fase de la señal va dando 'saltos' en el tiempo manteniendo constante su frecuencia y amplitud, por lo que la modulación empleada para esta señal es una modulación en fase.



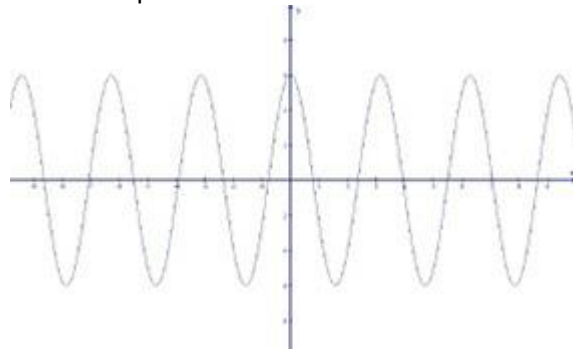
Figura 3 = Modulación en fase.

Veamos otro ejemplo.

Dada la siguiente señal analógica que se quiere transmitir con una amplitud de A voltios y $-A$ voltios y frecuencia 1 KHz , se emplea una modulación en fase de forma que la señal portadora de 1 Mhz , emplea un cambio de fase 180° en cada cambio de ciclo de la señal moduladora.

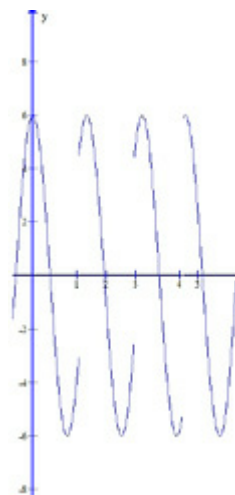
Dibuje en base a lo anterior la señal modulada correspondiente.

A continuación, razone si el ruido blanco puede afectar a la señal modulada transmitida.

**Solución:**

La modulación en fase se caracteriza por modular la fase de la señal portadora en función de la amplitud de la señal moduladora.

Dado que en este caso la señal moduladora solo presenta dos estados de amplitud (A voltios y $-A$ voltios), luego solo habrá dos estados de fase y donde cada vez que hay un cambio de ciclo en la señal moduladora (hay un cambio en la amplitud) la fase de la portadora 'pega un salto' de 180° .



Por último, el ruido blanco (siempre presente en todo medio de transmisión) afecta a cualquier tipo de modulación pero afecta mucho más a una modulación en amplitud (por interferencias en suma de señales) que a una modulación en fase.

El ruido impulsivo, en cambio, sería más perturbador a este tipo de modulación en fase (puede crear un salto de fase) que el ruido blanco

Sistemas de seguridad en el transporte de datos

La seguridad en las redes es un concepto que cada vez más se tiene en cuenta en este mundo hiperconectado.

Hasta hace poco las redes solo eran utilizadas por determinadas multinacionales y organizaciones por lo que su seguridad también estaba acotada a ellos.

Hoy día, todo el mundo (empresas, organizaciones, particulares e incluso objetos) están conectados a la red (generalmente a Internet) y es por ello que cada vez surge más el concepto de seguridad en las redes como aquel que permite que su conectividad y acceso se realicen de forma segura.

Todos queremos estar interconectados con todos. Están conectados a las redes (Internet) las personas, los equipos, las aplicaciones, etc.

Gracias a las redes podemos comunicarnos con todos, en cualquier parte del mundo y a cualquier hora.

Pero esto hace también que seamos vulnerables a intrusiones y ataques no deseados de personas (hackers o crackers) o aplicaciones (virus, rootkit, códigos maliciosos, etc.).

En base a esto surgió la seguridad informática como conjunto de técnicas o mecanismos que intentan proteger el almacenamiento, procesamiento y transmisión de la información que circula por las redes y en particular por Internet.

Con la seguridad informática conseguimos que nuestras conversaciones telefónicas y nuestros mensajes sean privados y confidenciales y que, por ejemplo, nuestros mensajes lleguen al destinatario deseado y no a otros y que además lleguen sin alteración ni manipulación.

En una red constantemente están transmitiéndose datos de un equipo a otro. Es por ello que debemos asegurar que dichas transmisiones sean seguras y de ello se encarga la seguridad informática. Es decir, debemos proteger la información que se transmita con objeto de que llegue a su destino de forma segura, no sea alterada ni manipulada, ni sea leída por terceros no deseados.

Para conseguir una comunicación segura se emplea como técnica de seguridad en redes el cifrado de las conexiones, es decir: el mensaje no se transmite tal cual se ha generado, sino que se codifica (según una clave) en un mensaje cifrado que es el que se transmite. El receptor (que conoce la clave) es capaz de descifrar este mensaje cifrado para poder visualizar el mensaje original.

Evidentemente, la clave debe ser conocida sólo por emisor y receptor para que la comunicación sea realmente segura.

Como se ha descrito anteriormente, para conseguir una comunicación segura ciframos nuestros mensajes y documentos antes de enviarlos.

De esto se encarga la criptografía, que es una técnica o ciencia capaz de crear mensajes ocultos.

La criptografía consiste en aplicar un algoritmo matemático sobre un documento original y legible con objeto de obtener otro documento no legible (cifrado) que es el que se transmite.

Este algoritmo matemático emplea una clave denominada clave del algoritmo, que solo es conocida por emisor y receptor y que debe aplicarla sobre el algoritmo para cifrar y descifrar el documento.

El algoritmo es un conjunto de reglas ampliamente conocido, pero la que es privada es la clave del algoritmo que es realmente la fortaleza del sistema.

Los algoritmos criptográficos emplean claves que son necesarias para el cifrado de los mensajes y/o documentos.

Pero estas claves pueden ser descifradas por hackers o crackers con la técnica denominada ataque de fuerza bruta, es decir, probar todas las combinaciones de símbolos posibles hasta dar con la clave buscada.

La constante ciberdelincuencia ha hecho que las empresas, organizaciones y organismos gubernamentales dediquen cada vez más esfuerzos, recursos y personal a la seguridad de los datos con objeto de protegerlo de los hackers, crackers y códigos maliciosos.

Para evitar el descifrado de la clave del algoritmo y con ello dar fortaleza al sistema de cifrado se deben tomar ciertas medidas en la elección de la clave del algoritmo. Entre ellas están:

–La clave de cifrado debe ser de gran longitud: por ejemplo, emplear claves de 512, 1024 o incluso 2048 bytes de forma que el atacante necesite de muchos recursos hardware y software para conseguirlo. Con ello se genera la desmotivación del atacante.

–Cambiar la clave regularmente: con ello se consigue que, en el caso de descifrarla, solo la tiene disponible un corto espacio de tiempo.

–Emplear todo tipo de símbolos disponibles: el uso de caracteres especiales (% , & , # , etc.), junto con valores numéricos y alfanuméricos hace más difícil su descifrado.

–No emplear palabras conocidas o identificables: es decir, fechas de nacimiento, película favorita, etc., que puede el atacante asociar con la persona atacada.

–Detectar intentos fallidos continuos en un intervalo corto de tiempo: Los algoritmos criptográficos existentes pueden ser de dos tipos:

•Emplean una criptografía simétrica: es decir, emplean la misma clave para cifrar (para el envío del mensaje o documento) que para descifrar (en la recepción del mensaje o documento).

•Emplean una criptografía asimétrica: es decir, emplean la misma clave para cifrar (para el envío del mensaje o documento) pero emplean otra diferente para descifrar (en la recepción del mensaje o documento).

Veremos a continuación las diferencias existentes entre uno y otro.

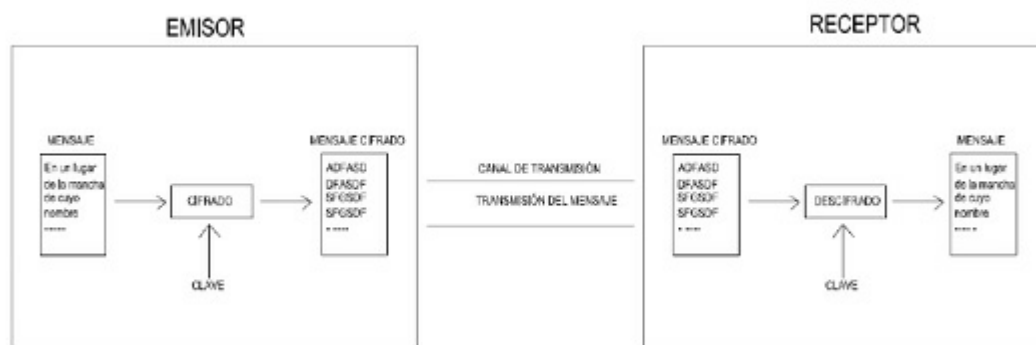
Criptografía simétrica

En la criptografía simétrica se emplea la misma clave del algoritmo para cifrar que para descifrar.

Es el tipo de criptografía más sencilla y se ha empleado desde la antigüedad.

En ella el emisor emplea una clave de cifrado antes de enviar el mensaje. El documento cifrado es el que se envía y una vez recibido por el receptor emplea la misma clave para descifrar.

Evidentemente, la clave debe ser la misma para ambos para obtener el documento cifrado y luego obtener el documento descifrado.



La criptografía simétrica presenta varios problemas o inconvenientes:

–La clave del algoritmo debe ser conocida por emisor y receptor y en algún momento debe ser enviada al 'otro' para poder emplear el algoritmo. Evidentemente, no podemos usar el mismo canal 'inseguro' para enviar la clave.

–Además, debemos tener una clave diferente por cada pareja emisor-receptor, ya que no se puede emplear la misma clave para varios receptores. Por ello, para cada usuario receptor (que puede ser un trabajador de una empresa) debemos disponer de una clave diferente. Esto hace que tengamos que tener almacenadas una gran cantidad de claves, una por cada receptor al que queramos enviar el mensaje.

En base a lo anterior, en los años 70 surgió la criptografía asimétrica, que solucionaba los problemas de la criptografía simétrica y es el método criptográfico más empleado actualmente.

Criptografía asimétrica

La criptografía asimétrica soluciona los problemas de la criptografía simétrica al emplear una clave de cifrado distinta para cifrar y otra para descifrar.

El emisor, cuando va a enviar un documento o mensaje, emplea una clave (denominada clave pública) con la que cifrar el documento. Esta clave puede ser conocida por cualquiera, ya que se emplea para cifrar.

Pero el documento solo puede ser cifrado por otra clave (clave privada) que solo conoce el receptor al que va dirigido el mensaje o documento que emplea dicha clave privada para descifrar el documento.

No existe relación matemática entre la clave pública o privada, por lo que conociendo la clave pública con que se cifró el mensaje no se puede descifrar. Solo puede descifrarse con la clave privada que solo conoce el receptor al que va dirigido el mensaje.

Con la criptografía asimétrica hemos resuelto los siguientes problemas:

–No es preciso transmitir la clave de cifrado.

Es una clave que puede ser conocida por cualquiera (por eso se le denomina clave pública), ya que sólo se emplea para cifrar el documento o mensaje.

–No hay problema de almacenamiento de clave.

Por cada documento que se envía a diferentes usuarios solo se emplea una clave de cifrado. Si aumenta la clave privada para el descifrado, pero esa clave privada la tiene almacenada cada receptor, porque el emisor solo guarda la clave de cifrado (clave pública).

No obstante, la criptografía asimétrica presenta ciertos problemas o vulnerabilidades, que son las siguientes:

–Se debe proteger la clave privada.

La clave privada empleada por cada receptor debe ser protegida para que nadie pueda emplearla para descifrar el mensaje cifrado recibido.

–La clave privada debe ser transportada.

La clave privada, en algún momento, debe ser transportada y para ello se emplea un mecanismo como el llavero de claves.

–Son poco eficientes.

Las claves de cifrado y descifrado suelen ser largas. Se tarda tiempo en cifrar y descifrar los documentos, lo que lo hace ineficiente en este sentido.

En seguridad informática, la autenticación se define como la confirmación de que un usuario, equipo o aplicación es quien dice ser y no otro.

Es decir, evitamos la suplantación y al impostor.

Esta técnica constituye uno de los pilares básicos en toda comunicación segura. Para conseguir esta autenticación generalmente en la transmisión de un documento o mensaje se debe loguear el usuario o máquina, es decir, le exigimos que introduzca su usuario y contraseña con lo cual el sistema confirma que la persona que envía el documento o mensaje es quien dice ser.

Esta técnica de autenticación está plenamente integrada en nuestras vidas cotidianas: nuestro DNI representa la autenticación de que esa persona es quien dice ser; el PIN de nuestro móvil indica al sistema que quien tiene acceso a usar la red móvil es el abonado que dice ser; el código de nuestra tarjeta de crédito indica que la persona que intenta sacar el dinero del cajero es realmente el titular de la cuenta, etc.

En seguridad informática, la integridad se define como la seguridad de que los datos almacenados son realmente los datos que se espera almacenar, es decir, que no han sido alterados ni manipulados.

Por tanto, cuando se intentan recuperar dichos datos son los mismos que se almacenaron, ya que no han sufrido alteración alguna en su forma o contenido.

Esto es otro de los pilares básicos en la seguridad de los datos.

Esta técnica, al igual que la autenticación, está plenamente integrada en nuestra vida cotidiana.

Por ejemplo, nuestro DNI lo forman ocho números seguidos de una letra. Existe un algoritmo que comprueba que dicha letra corresponde a esa numeración, ya que la letra se obtiene a través de una combinación aritmética de los números. Cuando se introduce un DNI en un sistema informático, lo primero que se hace es comprobar si dicho DNI es válido, y para ello aplica la combinación aritmética a los números. Si la letra que obtiene es la misma que la introducida en el DNI, el DNI es válido. Con esto se verifica la integridad del DNI introducido y el usuario puede operar con el sistema.

Hemos visto que los algoritmos criptográficos tanto simétricos como asimétricos necesitan una clave privada que ambos, emisor y receptor, deben compartir y que en algún momento debe ser transportada por un canal seguro.

Para transportar esta clave de cifrado (distribuir la clave) existen mecanismos de transporte de claves siendo el más habitual la tarjeta inteligente.

Esta tarjeta es un dispositivo generalmente de plástico y provisto de un chip electrónico en el que se almacena nuestra clave. Antes de usarla nos pedirá un PIN de acceso a la clave.

La implementación de esta tarjeta inteligente puede ser de dos tipos:

–Mediante una tarjeta de memoria flash.

–Mediante una tarjeta preprocesadora.

La primera es más insegura ya que cuando se introduce la tarjeta en el equipo para utilizar la clave se realiza una copia temporal de dicha clave en el equipo. Aquí es donde está la vulnerabilidad.

En cambio, en la segunda, la clave nunca sale de la tarjeta preprocesadora, ya que el proceso matemático de cifrado y descifrado que emplea la clave lo realiza el propio chip de la tarjeta preprocesadora.

Otra clasificación de las tarjetas inteligentes se puede realizar en función del interfaz de comunicación que emplea. Así encontramos:

–Tarjetas de contacto:

Son aquellas en las que debe existir un contacto (generalmente metálico) entre el equipo y la tarjeta inteligente para poder operar. Es el caso más habitual y más seguro.

–Tarjeta sin contacto:

Es aquella en que no existe contacto físico entre la tarjeta inteligente y el equipo. Generalmente la transmisión se realiza por radiofrecuencia. Es más inseguro (utiliza un canal como el aire, que es inseguro) pero es más rápido y por ello se emplea en lugares donde se busca gran rapidez, como por ejemplo en estaciones de tren, aeropuertos, etc.

Las aplicaciones en red son realmente los puntos críticos en la seguridad de las redes.

Son numerosas las aplicaciones que presentan vulnerabilidades y el hecho de que se empleen en red hace más fácil que códigos maliciosos puedan propagarse de un equipo a otro de manera muy rápida.

Es por ello que en el transporte de datos existen mecanismos que deben estar centrados en aquellas aplicaciones que más se utilizan en las redes para detectar sus vulnerabilidades y crear parches a dichos programas para asegurar su integridad.

A continuación, veremos aquellas aplicaciones y sus protocolos asociados que más se emplean en red, donde se describirá su funcionamiento y cómo la seguridad informática actúa sobre ellos para conseguir unas transmisiones fiables y seguras.

SSL representa un protocolo que trabaja en el nivel de transporte del modelo OSI o TCP/IP y que asegura una transmisión segura de la información entre los equipos, generalmente entre un equipo cliente y un servidor.

Que trabaje a nivel de transporte lo hace transparente de la arquitectura de los elementos de interconexión o encaminadores que trabajan a nivel 3 o nivel de red.

SSL surgió para proteger las conexiones existentes entre clientes y servidores web con el protocolo http que se empleaban para el comercio electrónico. Esta protección debía asegurar al cliente que se había conectado al servidor auténtico, y enviarle en consecuencia datos confidenciales, como por ejemplo los datos de su número de tarjeta de crédito.

El protocolo se encarga de encapsular el trabajo de los elementos de la capa superior, construyendo un canal de comunicaciones entre los dos extremos objeto de la comunicación. Esto lo realiza empleando la técnica de Handshake, encargada de intercambiar la clave que se utilizará para crear un canal seguro mediante un algoritmo eficiente de cifrado simétrico.

SSH representa un protocolo en el que se definen una serie de reglas que permiten una transmisión segura de un equipo a otro al cifrar la información que transmite de extremo a extremo.

Es muy empleado para configuración remota de equipos (sustituyendo al TELNET), para transferencia de archivos (sustituyendo al FTP), crear canales seguros de comunicaciones, etc.

En la actualidad existen dos versiones de este sistema SSH: SSH1 y SSH2, siendo esta última la más usada al tener mejoras sobre la primera.

Una de las aplicaciones más usadas que emplea dicho protocolo es el OpenSSH.

SSH constituye un protocolo seguro de transferencia de información de un equipo a otro. Esta seguridad la proporciona porque la información intercambiada va cifrada de un equipo a otro. Muchos programas TELNET, como Putty, integran este protocolo para gestionar equipos en modo remoto y con comunicaciones seguras.

IPsec es un protocolo que define una serie de reglas que, trabajando a nivel de red (a diferencia de SSL que trabaja a nivel de transporte), permiten una comunicación segura y fiable entre extremos o dicho de otro modo entre equipos.

Para ello este protocolo usa técnicas criptográficas además de la autenticación para la transmisión segura de los datos.

NIVEL DE APLICACIÓN
NIVEL DE PRESENTACION
NIVEL DE SESION
NIVEL DE TRANSPORTE
NIVEL DE RED IPSEC
NIVEL DE ENLACE
NIVEL FÍSICO

Un cortafuegos o firewall es una aplicación software especializada que se intercala entre las aplicaciones y la tarjeta de red para realizar un filtrado de los paquetes.

El objetivo no es más que controlar los paquetes que entran y salen del equipo por la red con objeto de actuar correctamente ante paquetes sospechosos.

Es una herramienta básica para la seguridad de las redes y todos los equipos siempre deberían tenerla instalada y activada.

El firewall se puede instalar tanto en equipos clientes como en equipos servidores actuando de forma diferente aunque empleando las mismas técnicas.

Veremos a continuación cómo trabaja según sea un equipo cliente o un equipo servidor.

Cortafuegos en un equipo cliente

En un equipo cliente el cortafuegos actúa aplicando un filtrado de paquetes.

En el tráfico saliente, es decir, de paquetes que salen del equipo hacia la red, el cortafuegos o firewall analiza la cabecera de cada paquete y, en función de las reglas que tenga definidas en el cortafuegos, realiza una acción u otra. Por ejemplo, si detecta que la máquina cliente hace spam, bloquea el puerto 25 (puerto de correo electrónico).

En el tráfico entrante, es decir, de paquetes que entran al equipo cliente, analiza la cabecera de cada paquete y en función de las reglas de configuración actúa de una forma u otra.

Cortafuegos en un equipo servidor

En un equipo servidor el cortafuegos también actúa aplicando un filtrado de paquetes.

En el tráfico saliente, es decir, de paquetes que salen del equipo hacia la red, el cortafuegos o firewall analiza la cabecera de cada paquete y en función de las reglas que tenga definidas en el cortafuegos realiza una acción u otra. En este tráfico es donde generalmente actúa el cortafuegos puesto que en un servidor la mayor parte del tráfico es saliente.

En el tráfico entrante, en cambio, el firewall analiza qué paquetes quieren acceder a qué puertos del servidor para bloquear en el caso de un ataque al servidor si así lo tiene en las reglas de configuración.

1. Incidencias en dispositivo de acceso a redes públicas

Tipos de Incidencias en una red:

Tipo	Generadas	Síntomas	Soluciones
Físicas	Por componentes de tipo hardware	Perdidas de rendimiento en un servicio, corte de conexión ...	Reemplazar el dispositivo o componente dañado
Red	Asociadas al funcionamiento de la red y no se incluídas en los fallos físicos	Perdidas de la señal, rendimientos o bajada velocidad de transmisión...	Cambios en la configuración y actualizaciones del software
Usuarios	Configuraciones, Averías, Peticiones de ayuda o Mal uso de los servicios	Se suelen derivar de otro tipo de incidencia	Se suelen derivar de otro tipo de incidencia
Servicios	Asociados al mal funcionamiento de los servicios que proporciona la red	Cortes de servicios. Perdidas de rendimiento. Baja velocidad del servicio	Soluciones muy diversas, dependiendo del servicio, pero podemos implementar redundancia en los servicios.
Ambientales	Asocia a factores ambientales	Inundaciones, Fuego, Altas temperaturas.	Evitar altas temperaturas, ventiladores, ventilación de aire...
Seguridad	Fallos de seguridad	Accesos no autorizados Programas. Intrusiones externas. Acceso físico no autorizado.	cortafuegos, antivirus... Permisos de usuarios. Política de contraseña

1.1 Incidencias habituales

Los problemas de conectividad en redes pueden tener múltiples causas. La solución al problema una vez identificado el fallo, no debe de ser muy difícil.

A continuación se presentan algunos problemas más comunes:

- ⇒ No se puede conectar con un servidor de acceso remoto.
- ⇒ Al intentar conectar, aparece un mensaje que indica que el servidor de acceso remoto no responde.
- ⇒ Se corta continuamente la sesión con el servidor
- ⇒ Las conexiones se desconectan de forma anormal.
- ⇒ Al intentar conectar, se recibe un mensaje de error del hardware.
- ⇒ Al intentar conectar mediante una ISDN recibimos el mensaje «Sin respuesta».
- ⇒ Las conexiones configuradas usando TCP/IP generan errores.
- ⇒ Las conexiones establecidas mediante conexión compartida a Internet (ICS) generan errores.
- ⇒ No hay respuesta en una conexión de red de área local.
- ⇒ Los clientes de conexiones entrantes no pueden ver ningún recurso excepto el equipo de conexión entrante.
- ⇒ La tarjeta inteligente no funciona.
- ⇒ Al intentar establecer una conexión de acceso remoto mediante una tarjeta inteligente, recibe un mensaje de error.
- ⇒ Hay conflictos de direcciones IP.
- ⇒ Hay problemas en la conexión de fibra óptica.
- ⇒ El proveedor de servicios de internet no funciona correctamente.
- ⇒ Hay conflictos en los canales Wi-Fi.
- ⇒ Los cables trenzados no están correctamente grimpados.
- ⇒ Los servidores DNS no funcionan correctamente.
- ⇒ Problemas de conectividad achacables a los proveedores de servicios de internet (ISP).
- ⇒ Problemas de enrutamiento.
- ⇒ Problemas con las subredes.
- ⇒ Problemas de alimentación eléctrica.
- ⇒ Sobrecalentamientos, etc.

Hay problemas con Internet ¿qué hago?

Si se pierde la conexión a Internet debemos comprobar si el Router está encendido y en funcionamiento. En caso afirmativo, puede que se haya perdido la conexión con el servidor. En este caso deberíamos reiniciar el Router. Debemos esperar dos o tres minutos a que se establezca correctamente la sincronización. Si sigue sin funcionar, el problema no está en el Router y debemos de buscar a buscar otras causas.

Para empezar, identificaremos si el problema se limita a un equipo, al navegador o recae en toda la conexión.

En primer lugar deberíamos comprobar si se produce la misma incidencia en más de un navegador, si es así lo mejor sería reiniciar el equipo. También el problema puede estar en el cortafuegos de nuestro software de seguridad.

También deberíamos de comprobar si la incidencia afecta más de un equipo de la red.

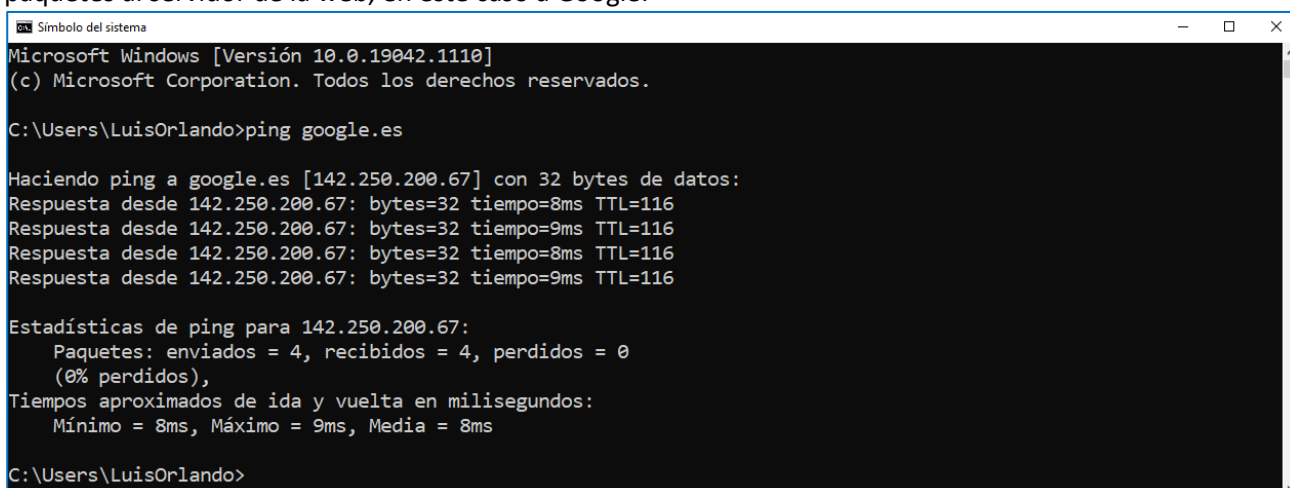
Si el problema persiste, podemos acceder al intérprete de comando del Windows pulsando la combinación de teclas de Windows más la tecla R.

Ahí escribiremos **cmd** y pincharemos en Aceptar.

Una vez en el intérprete de comandos, tecleamos el comando ping de la siguiente forma:

ping www.google.es

Si todo funciona correctamente se nos informará del tiempo en milisegundos que tardan en llegar los paquetes al servidor de la web, en este caso a Google:



```

Microsoft Windows [Versión 10.0.19042.1110]
(c) Microsoft Corporation. Todos los derechos reservados.

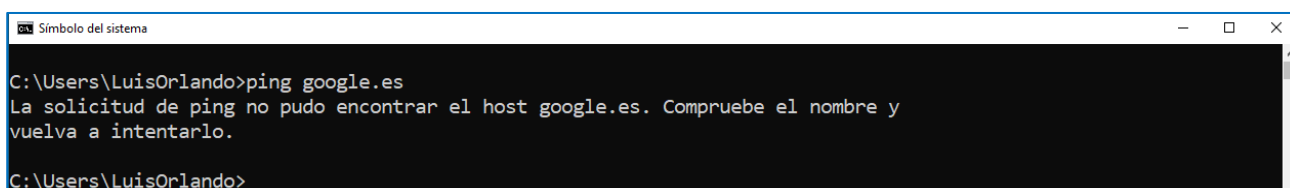
C:\Users\LuisOrlando>ping google.es

Haciendo ping a google.es [142.250.200.67] con 32 bytes de datos:
Respuesta desde 142.250.200.67: bytes=32 tiempo=8ms TTL=116
Respuesta desde 142.250.200.67: bytes=32 tiempo=9ms TTL=116
Respuesta desde 142.250.200.67: bytes=32 tiempo=8ms TTL=116
Respuesta desde 142.250.200.67: bytes=32 tiempo=9ms TTL=116

Estadísticas de ping para 142.250.200.67:
    Paquetes: enviados = 4, recibidos = 4, perdidos = 0
            (0% perdidos),
    Tiempos aproximados de ida y vuelta en milisegundos:
        Mínimo = 8ms, Máximo = 9ms, Media = 8ms

C:\Users\LuisOrlando>
  
```

Si no hay conexión, se nos indicará que no se puede acceder al host de destino:



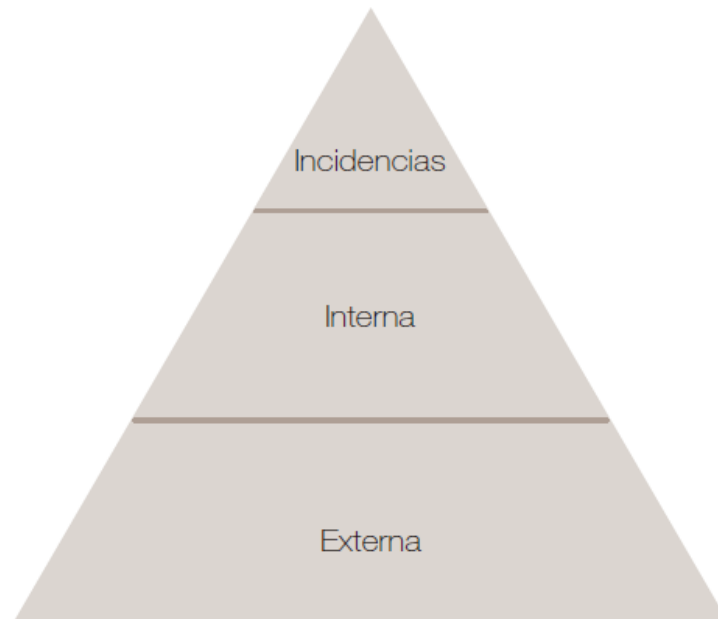
```

C:\Users\LuisOrlando>ping google.es
La solicitud de ping no pudo encontrar el host google.es. Compruebe el nombre y
vuelva a intentarlo.

C:\Users\LuisOrlando>
  
```

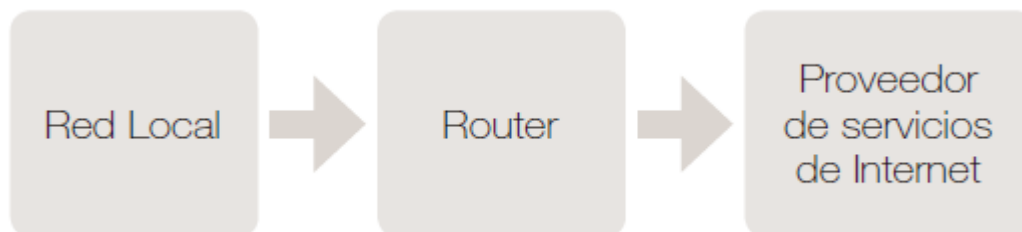
Si no tenemos conexión, probablemente haya algún conflicto en la línea de Fibra/ADSL, algún problema con la conexión interna de la red hasta el Router, o incluso errores en las DNS o con la configuración del equipo, etc.

Las incidencias pueden ser muy variadas.



1.1.1. Incidencias internas

Las incidencias internas, son aquellas que afectan desde el Router hacia adentro de la red.



Problemas de Fibra/ADSL y Router

Muchos de los errores de acceso Internet se deben al cable telefónico que trae la señal ADSL desde el servidor. También el propio Router suele generar muchos problemas.

Si el Router no puede sincronizar con la línea, la luz indicadora de Conexión parpadea lentamente o está apagada.

Si después de haber reiniciado el Router, no podemos sincronizar con el servidor, posiblemente el fallo esté en el cableado que trae la señal.

Si tenemos sincronización con el servidor, pero Internet no funciona, podemos restaurar la configuración completa del Router con los parámetros de originales de fábrica.

Si tras restaurar el Router a sus valores originales, las cosas siguen igual, podría ser debido a un problema de enrutado en la infraestructura del operador, algo sobre lo que no podemos actuar nosotros mismos.

Primeras Verificaciones:

Para averiguar dónde está el problema, en casi todos los casos tenemos que verificar si el problema es:

- ⇒ Interno, de la configuración de nuestro equipo o del propio interfaz de red
- ⇒ De nuestra red Local, puede ser de la configuración de la red o de los elementos hardware internos de la misma: cable, switch, router...
- ⇒ De nuestro proveedor de servicios, puede ser de una mala configuración de nuestro equipo o de la configuración del propio hardware del ISP: router...

Lo primero que debemos comprobar es si funciona nuestra tarjeta de red, luego nuestra red Local, una vez comprobado todo esto, comprobamos si funciona nuestra red WAN.

Usando los comandos:

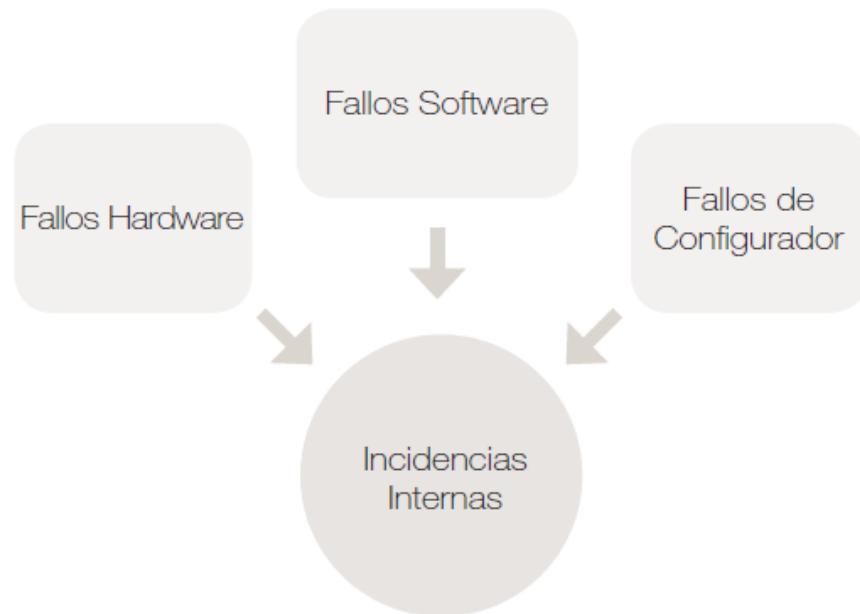
- Usando ipconfig, vemos nuestra configuración de red
- Hacemos ping a nuestra IP, por ejemplo: ping 192.168.1.44

- Hacemos ping desde la red local a un equipo de la red, por ejemplo: ping 192.168.1.70
- Hacemos ping al router: ping 192.168.1
- Hacemos ping a una URL conocida y fiable: ping google.es
- Y si no funciona hacemos un ping a una ip publica conocida: ping 8.8.8.8

Si falla en alguno de estos pasos ya tendremos una pista de dónde esta el problema para solucionarlo.

Las incidencias Internas podemos dividir las en tres categorías:

- ⇒ Fallos HW.
- ⇒ Fallos SW.
- ⇒ Fallos de configuración interna o de los interfaces de interconexión.



Errores en la conexión Wi-Fi

Si el equipo ha dejado de conectarse a la red Wi-Fi y no hay forma de que recupere la conexión, puede haber diversas causas.



En primer lugar reiniciamos el equipo y el Router Wi-Fi. Esto suele resolver el problema en el 80% de los casos, ya que una vez hecho esto, se suele solucionar y todo vuelve a funcionar correctamente otra vez. En caso de que siga habiendo problemas, deberíamos de asegurarnos de que esté activada la señal Wi-Fi en el Router (cada el fabricante es diferente pero lo normal es mediante un interruptor, una combinación de teclas o un software instalado en el equipo).

Después hemos de pinchar sobre el icono de la barra de tareas en forma de escala de señal, para ver la lista con las redes Wi-Fi detectadas.

Ahora debemos comprobar que nuestra red aparece en la lista de redes Wi-Fi disponibles. En caso de que aparezca, intentaremos conectarnos a ella. Si tampoco hay conexión, tendremos que asegurarnos de que la clave de seguridad que hemos introducido es correcta (muchos errores de conectividad se derivan de errores en las claves).

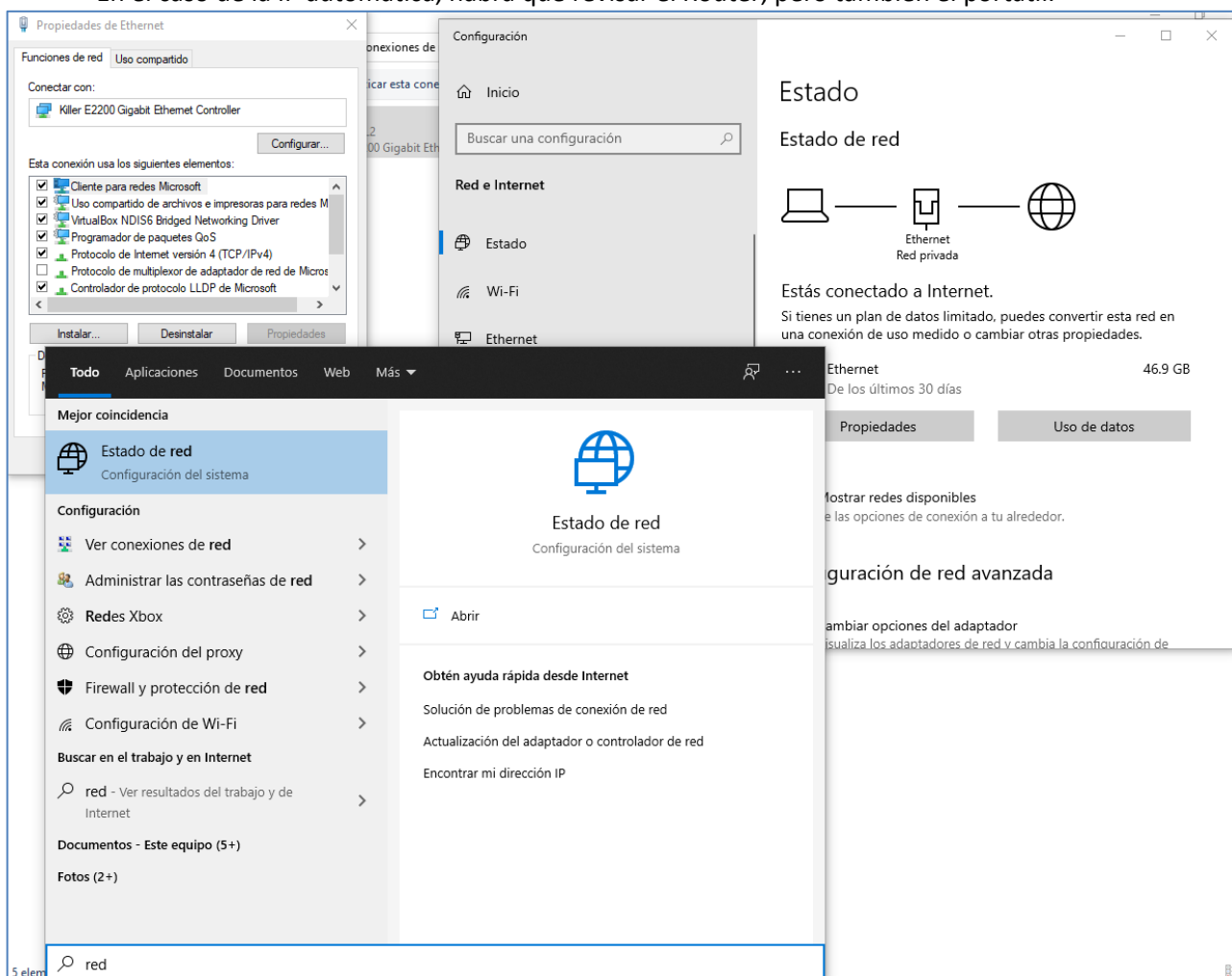
Si la red no aparece entre las redes disponibles o si aparece, no podemos conectar, habrá que revisar el punto de acceso.

También puede pasar que a pesar de tener conectado el equipo al Router Wi-Fi, no hay de tráfico entre el ordenador y la red. Esto se puede deber a muchas razones. Como hemos dicho con anterioridad, lo primero va a ser siempre reiniciar el equipo y el Router Wi-Fi.

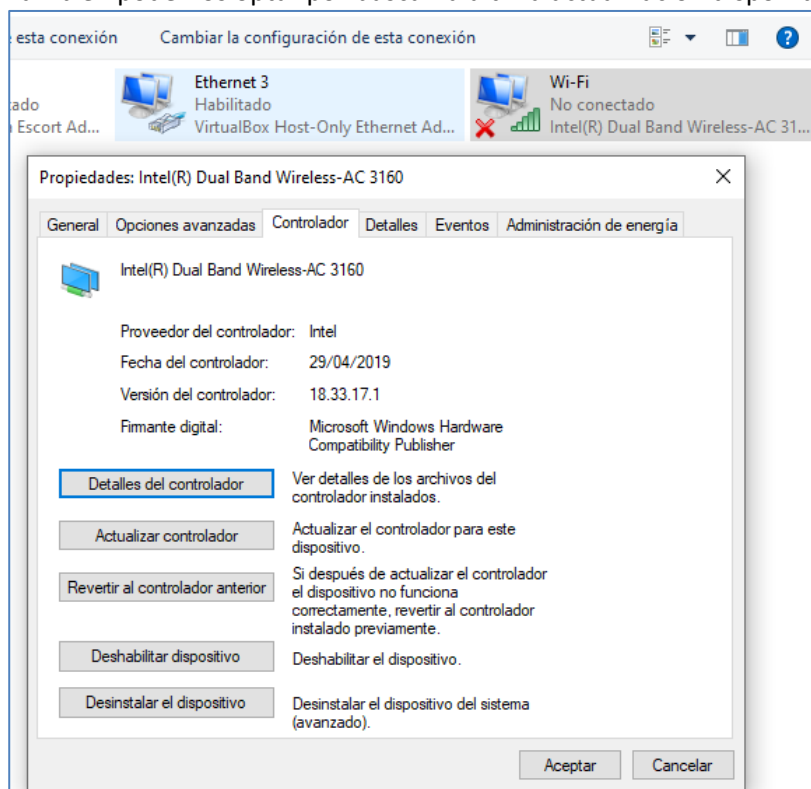
Si tras haber reiniciado, seguimos con problemas, habrá que comprobar que el Router está asignando la dirección IP correctamente.

Para ello:

- ⇒ Buscaremos el texto:
 - Red (y nos saldría **estado de red**) configuración Windows 10
 - Si buscamos conexiones (Ver conexiones de red) configuración Windows 7
- ⇒ Y luego iremos a Centro de redes y recursos compartidos
- ⇒ Conexión de red inalámbrica, lo que nos mostrará los detalles del adaptador.
- ⇒ Detalles podremos ver si se nos está asignando una IP automática por DHCP o es automática.
- ⇒ En el caso de la IP automática, habrá que revisar el Router, pero también el portátil.

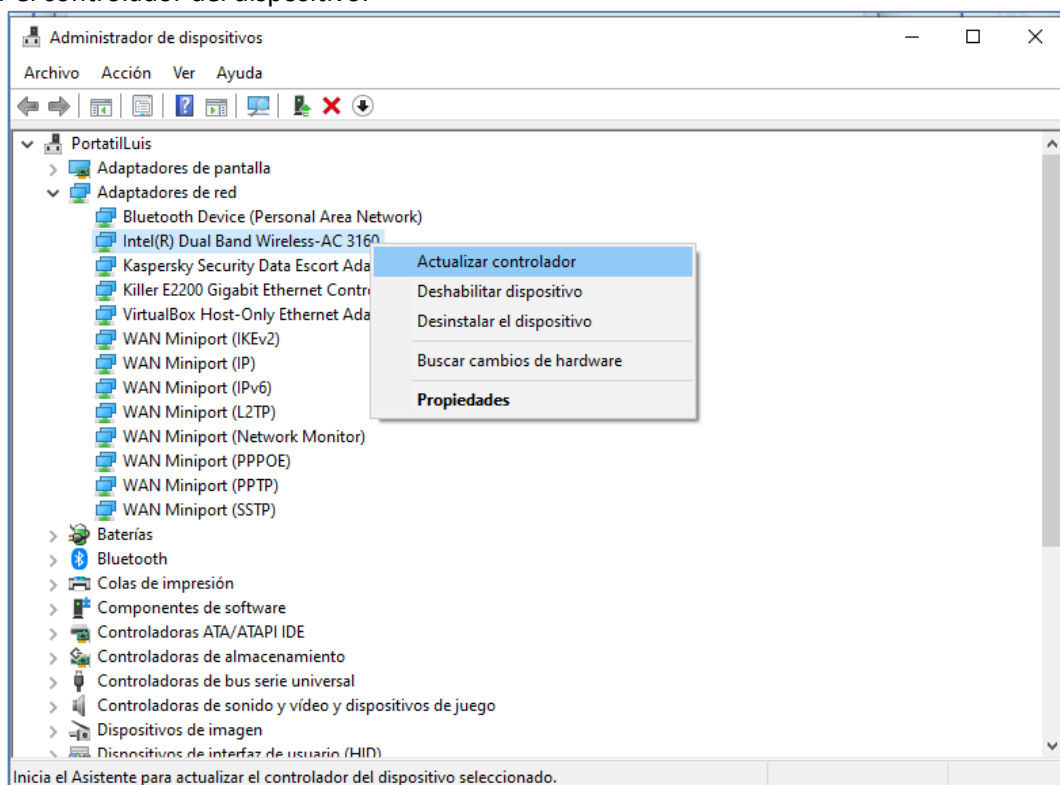


También podemos optar por buscar la última actualización disponible para el controlador de la tarjeta Wi-Fi.

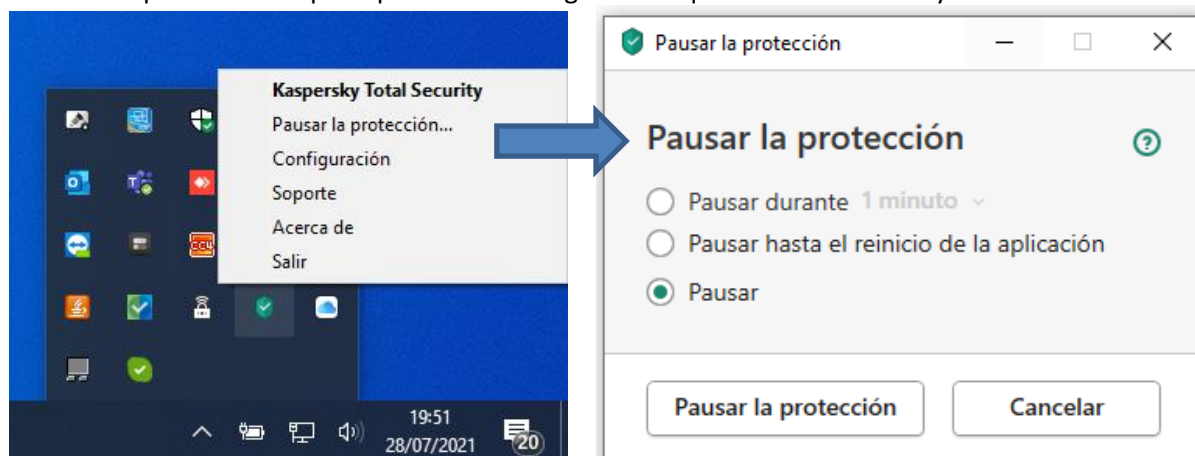


Para poder ver los elementos hardware del equipo, incluido el adaptador Wi-Fi, hemos de acceder al administrador de dispositivos de Windows. Para ello, debemos pulsar la combinación de la Tecla de Windows y la Tecla Pausa. Ello nos dará acceso a la opción propiedades del sistema. Dentro de esta ventana, debemos pinchar sobre la opción Administrador de dispositivos.

Aquí aparecerán todos los elementos hardware del equipo. En la categoría adaptadores de Red, debe aparecer el adaptador Wi-Fi. Deber doble clic sobre él, se debe mostrar una ventana en la que aparecerán se mostrarán todos los detalles del dispositivo. Accedemos a la pestaña controlador y ahí podremos actualizar el controlador del dispositivo.



Si tenemos instalado algún software de seguridad instalado (antivirus, cortafuegos, etc.), sería bueno desactivarlo temporalmente para averiguar si causante del problema puede ser este. Hay que tener en cuenta que la función principal del cortafuegos es bloquear accesos desde y hacia la red.



Router Wi-Fi

En lo referente a los Router Wi-Fi, hay muchos tipos de ellos atendiendo a factores como:

- ⇒ Estándares Wi-Fi: 802.11b/g/a/n
- ⇒ Seguridad: WEP, WPA o WPA2

Frecuencias de trabajo: entre 2,4 y 5 GHz

A esto hay que añadir las propias tecnologías de algunos fabricantes que cada uno ha desarrollado con la idea de acelerar y mejorar la transmisión de datos.

Una de las primeras cuestiones de debemos hacer si el Router Wi-Fi falla es actualizar su firmware a la última versión disponible en la web del fabricante.

Una vez hecho eso, lo más prudente sería, reiniciar el Router Wi-Fi y a continuación restablecer la configuración de fábrica del Router. Ahora podríamos ir modificando ciertos parámetros empezando por lo más sencillo y después configurando lo más complejo.

Podríamos desactivar la seguridad, ajustamos la banda en 2,4 GHz y seleccionamos el estándar más antiguo (802.11b u 802.11g). Esta sería la configuración más básica, menos problemática, pero también la menos eficiente.

Si todo funciona, ahora podríamos empezar aumentando la seguridad a WEP y después a WPA, también iremos cambiando la cambiando la banda de emisión y después el estándar de transmisión.

Interferencias inalámbricas

Los problema también pueden deberse a interferencias que dificulten la comunicación inalámbrica entre el equipo y el Router. Hay teléfonos inalámbricos DECT que funcionan en la banda de 2,4 GHz, en el área de cobertura demasiadas redes Wi-Fi, que operen a en el mismo canal o incluso puede haber en zonas próximas a edificios sensibles inhibidores de señal.

Para solucionar esto, podemos cambiar el canal de emisión. También podemos cambiar el Router Wi-Fi e instalar otro que funcione a 5 GHz. Esta es una banda muy poco utilizada aún y seguro que apenas tendremos interferencias.

Los servidores DNS, muy importantes

Algunos problemas como la velocidad lenta, los derivados de la propia dificultad de acceder a Internet o el no poder acceder a determinadas direcciones, pueden estar causados por la configuración de los servidores DNS que tenemos en la configuración de red.

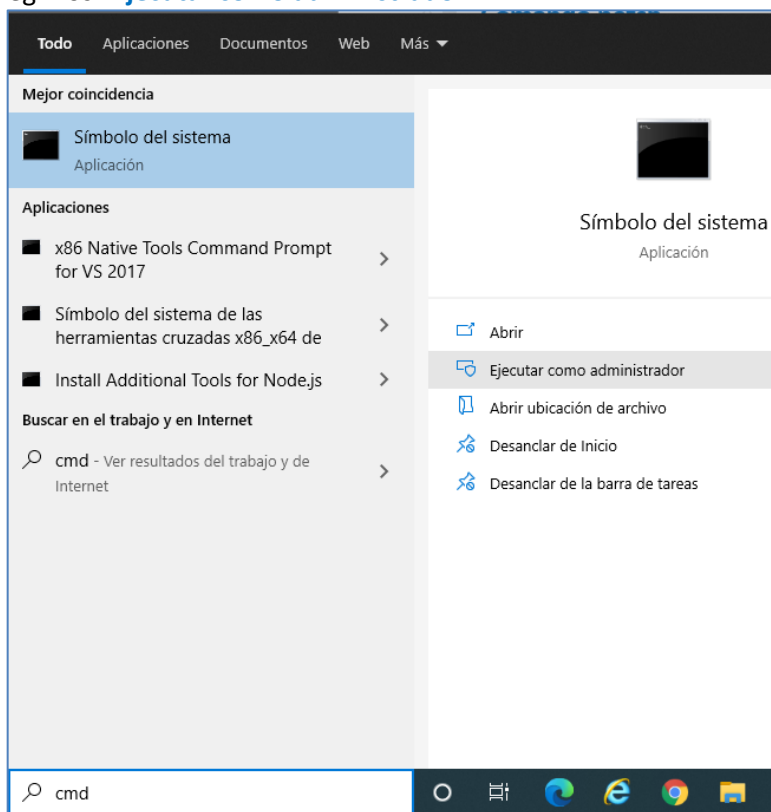
Los operadores de Internet cuentan con sus propios servidores DNS que, en muchos casos, pueden estar sobrecargados, caídos o no ser de mucha calidad. Esto provoca malos funcionamientos.

A veces es necesario cambiar la configuración de los servidores DNS que asigna por defecto el Router. Esto puede mejorar la conexión y sobre todo la velocidad a la hora de acceder a muchas páginas.

Más adelante trataremos el problema de los servidores DNS en profundidad. Decir ahora que servidores DNS como el de Google, OpenDNS o incluso los de Norton, además de gran disponibilidad, aportan seguridad a la conexión al bloquear el acceso a páginas malware o phishing.

Comando netsh

Otro comando que podemos usar es netsh, el cual permite cambiar la IP de una interfaz de red... Para hacerlo tenemos que ejecutar la línea de comandos como administradores, así escribimos **cmd** en buscar programas y elegimos **"Ejecutar como administrador"**



Desde la línea de comandos lo primero que vamos a verificar es las conexiones ipv4 que tenemos y sus nombres, para lo que escribimos: **netsh interface ipv4 show config**

```
Administrador: Símbolo del sistema
C:\WINDOWS\system32>netsh interface ipv4 show config

Configuración para la interfaz "Ethernet"
DHCP habilitado:                Sí
Dirección IP:                   192.168.1.44
Prefijo de subred:              192.168.1.0/24 (máscara 255.255.255.0)
Puerta de enlace predeterminada: 192.168.1.1
Métrica de puerta de enlace:   0
Métrica de interfaz:            25
Servidores DNS configurados a través de DHCP: 192.168.1.1
Registrar con el sufijo:       Solo el principal
Servidores WINS configurados a través de DHCP: ninguno

Configuración para la interfaz "Ethernet 3"
```

Vemos que tenemos varios interfaces de red, y el que nos interesa es el primero, el interfaz de red local llamando **Ethernet** y vemos que tiene activado DHCP.

Vamos a ver como podemos cambiar la ip a estática (static), mascara de subred y puerta de enlace, para hacerlo escribimos en la línea de comandos:

netsh interface ipv4 set address name="Ethernet" static 192.168.1.51 255.255.255.0 192.168.1.1

```
Administrador: Símbolo del sistema
C:\WINDOWS\system32>netsh interface ipv4 set address name="Ethernet" static 192.168.1.51 255.255.255.0 192.168.1.1
C:\WINDOWS\system32>
```

Si escribimos de nuevo **netsh interface ipv4 show config** veremos cómo ha cambiado la dirección IP pero no tenemos ninguna dirección DNS con lo cual no vamos a poder navegar por internet con normalidad:

```
Administrador: Símbolo del sistema

C:\WINDOWS\system32>netsh interface ipv4 set address name="Ethernet" static 192.168.1.51 255.255.255.0 192.168.1.1

C:\WINDOWS\system32>netsh interface ipv4 show config

Configuración para la interfaz "Ethernet"
DHCP habilitado: No
Dirección IP: 192.168.1.51
Prefijo de subred: 192.168.1.0/24 (máscara 255.255.255.0)
Puerta de enlace predeterminada: 192.168.1.1
Métrica de puerta de enlace: 1
Métrica de interfaz: 25
Servidores DNS configurados estáticamente: ninguno
Registrar con el sufijo: Solo el principal
Servidores WINS configurados estáticamente: ninguno
```

Por lo tanto tenemos que indicarle también la configuración DNS, en este caso le podemos decir que use la IP de Google **netsh interface ipv4 set dnsservers "Ethernet" source=static 8.8.8.8**

```
Administrador: Símbolo del sistema

C:\WINDOWS\system32>netsh interface ipv4 set dns name="Ethernet" static 8.8.8.8

C:\WINDOWS\system32>netsh interface ipv4 show config

Configuración para la interfaz "Ethernet"
DHCP habilitado: No
Dirección IP: 192.168.1.51
Prefijo de subred: 192.168.1.0/24 (máscara 255.255.255.0)
Puerta de enlace predeterminada: 192.168.1.1
Métrica de puerta de enlace: 1
Métrica de interfaz: 25
Servidores DNS configurados estáticamente: 8.8.8.8
Registrar con el sufijo: Solo el principal
Servidores WINS configurados estáticamente: ninguno
```

Configuración de IP

Asignación de IP:	Manual
Dirección IPv4:	192.168.1.51
Longitud del prefijo de subred IPv4	24
Puerta de enlace de IPv4:	192.168.1.1
Servidores DNS IPv4:	8.8.8.8

Editar

También podemos añadir el servidor secundario escribiendo **netsh interface ipv4 add dnsservers "Ethernet" 192.168.1.1 index=2**

Y para cambiar la configuración Ip a automática con DHCP escribimos:

netsh interface ipv4 set address name="Ethernet" source=dhcp

```
Administrador: Símbolo del sistema

C:\WINDOWS\system32>netsh interface ipv4 set address name="Ethernet" source=dhcp

Configuración de IP

Asignación de IP: Automático (DHCP)

Editar
```

1.1.1.1. Fallos Hardware

Los fallos hardware, los podemos clasificar conforme al siguiente esquema:

Fallos Hardware:

- Organización
 - Rack
 - Paneles de distribución
- Medios de transmisión
 - Cableado estructurado
 - Medios inalámbrico
- Equipos eléctricos
 - Regletas
 - S.A.I.
- Equipos de conexión
 - Hub
 - Switch
 - Modem
 - Router

Los equipos de distribución

Los equipos de distribución se encargan de controlar y repartir el flujo de datos en la red local. Los más utilizados son los concentradores (Switch) y los encaminadores (Routers).

Para conectar los equipos de distribución y los equipos (ordenadores, portátiles, TPV, etc.) entre sí se utiliza lo que llamamos el cableado estructurado. Este se compone de elementos pasivos como rosetas, los conectores, el cable cables de par trenzado coaxial o de fibra optica, los armarios Rack, los paneles de parcheo, etc.

Electrónica de red

Para establecer la conexión entre los distintos equipos de la red es necesario instalar unos dispositivos electrónicos:

- Switch: Centralizan las comunicaciones y distribuyen la información.
- Router: proporcionan acceso a internet
- Router Wi-Fi: proporcionan acceso a internet para redes inalámbricas.

Rack

Un armario rack se utiliza para alojar en su interior todos los elementos de distribución de la electrónica de la red y centralizar en su interior todas conexiones de la red.

Dentro de un rack podemos incluir:

- Los Hub's.
- Los Switch.
- El Router.
- Los ordenadores que hacen las veces de servidores de red.
- Los sistemas físicos de copias de seguridad automáticas
- etc.

La principal función del rack es la de la organización de los distintos elementos necesarios para la conexión de la red. Dentro de ellos se colocan los paneles (meros intermediarios entre los cables de par trenzado que vienen de los ordenadores y los latiguillos que llevarán la conexión hasta los puertos de los Hub's o Switch's). También, como dijimos antes se, colocan los elementos de la electrónica de la red (Switch, Hub y por supuesto el Router).

Hay armarios rack de tamaños diferentes. Hay algunos pequeños para alojar un solo panel y un solo Switch o Hub. En el otro extremo los hay de tamaño bastante grande que incluyen prestaciones adicionales como puede ser la climatización, fuente de alimentación, etc.

En el interior del rack se concentran las señales de voz y datos y desde ahí se reparten por toda la red de área local.

Los racks se componen de un armazón metálico con unas medidas normalizadas (el ancho suele ser de 19 pulgadas). La altura y el fondo varían según el fabricante y las necesidades que tengamos.

Tienen unos bastidores con unas guías horizontales. En estos bastidores se atornillan los elementos que se van a fijar en su interior. También tienen otras guías verticales para fijarlos en la pared.

Podemos clasificar los armarios rack de varias formas:

- **Accesibilidad:** Algunos modelos tienen en los paneles laterales algunas aberturas que facilitan las labores de mantenimiento y de conexión de nuevos elementos.
- **Tamaño.** Hay rack de tipo columna diseñados para grandes redes. Los hay pensados para redes más pequeñas, estos son de fijación mural. El ancho normalizado es de 19 pulgadas pero también los hay de 10 pulgadas.
- **Unidades de rack («U»).** Es una unidad de medida que define la capacidad de almacenamiento del rack. Un rack («U») equivale a 1,75 pulgadas (44,45 mm) de alto.

Paneles de distribución de cableado

Los patch panel, paneles de distribución de cableados, se utilizan para conectar los distintos cables de red. Las líneas de entrada y salida de los equipos se conectan a estos elementos. Forman parte de estas conexiones ordenadores, impresoras, servidores...

Los patch suelen estar colocados en el rack, armarios de distribución.

En la red LAN un patch panel conecta entre sí los equipos de la red, además proporciona líneas salientes que permitirán a la LAN conectarse a Internet o a otra red de tipo WAN.

Las conexiones efectivas de los equipos se realizan con los llamados cables de parcheo, patch cord.

Este sistema de conexión, patch panel y patch cord, permite realizar cambios de forma rápida y sencilla permitiendo la detección y corrección de errores de una manera intuitiva. Normalmente estos cables estarán en la parte frontal del panel mientras que en la parte posterior estarán las conexiones permanentes, las que no son susceptibles de cambios.

Hay diferentes modelos de patch panel, en función del tipo de cableado: fibra, coaxial, par trenzado y sin trenzar que utilizaremos según los requerimientos de los elementos a interconectar.

Estarán destinados a soportar no sólo datos y teléfonos, sino aplicaciones de video y audio.

Además, son necesarios elementos adicionales como unidades guía cables y regletas de red eléctrica.

Son estructuras metálicas con placas de circuitos que permiten interconexión entre los equipos. Un Patch-Panel posee unos puertos RJ-45, cada uno se asocia a una placa de circuito, y ésta a su vez se propaga en pequeños conectores RJ-45 hembra. Es aquí donde se pinchan los conectores RJ-45 machos de los cables provenientes de los cajetines u otros Patch-Panels.

Lo definiremos como una matriz de conectores hembra RJ 45 que se utiliza para realizar conexiones cruzadas entre los equipos activos a través del cableado horizontal. Permite un gran manejo y administración de los servicios de red, ya que cada punto de conexión del Patch-Panel maneja el servicio de una salida de telecomunicaciones.

Permite la conexión entre equipos por tanto deben ser de primera calidad debido a que por sus puntos transitan señales de alta velocidad.

La idea del Patch-Panel además de seguir estándares de redes, es la de estructurar y ordenar los cables que conectan equipos en la red.

Características

El Patch-Panel puede venir o no montado de fábrica. Preparado para pinchar el cable, o sólo viene el troquel para instalar cada uno de los elementos. Los hay de 12, 24, 48, 96 puertos (siempre en potencias de 2).

Los patch-panels modulares aceptan las mismas rosetas que se ubican en los puestos de trabajo en cualquier orden, tipo y color.

Opcionalmente de pueden montar en pared, o se puede utilizar directamente para instalar en rack con frente estándar de 19". Hay modelos compactos permitiendo ahorrar espacio en el rack.

Se utiliza un ordenador de cableado y etiquetando cada puerto con su correspondiente puesto de trabajo, se asegura una perfecta administración de la red una vez concluida la instalación.

Cuando la conexión se realiza con fibra óptica, existen patch panels especiales que permiten acomodar en 1 HU 12 ports ST ó SC y en 2 HU 24 ports ST ó SC.

El cableado Estructurado:

Denominamos cableado estructurado al tendido de cables en el interior de un edificio con el objeto de implantar una red de área local. Suele emplearse cable de par trenzado de cobre, aunque hay redes funcionando con fibra óptica o cable coaxial.

Este sistema de cableado planificado es de obligado cumplimiento en las contrataciones públicas en el entorno de la UE ya que presenta las siguientes ventajas:

- ⇒ Está pensado para hacer frente a modificaciones y al posible crecimiento de la instalación.
- ⇒ Permite la transmisión de datos, vídeo, voz, sistemas de alarma, etc.
- ⇒ Facilita la detección, administración y solución de las posibles averías del cableado. Se usa una topología jerarquizada en forma de estrella (árbol).
- ⇒ Está sujeto a la Norma estándar EN-50173, que incluye la normativa EIA/TIA-568. En nuestro país, AENOR, desde enero de 2004 ha ratificado la norma EN-50173.

Subsistema de cableado horizontal:

Se extiende desde el rack hasta las áreas de trabajo. El cableado horizontal se inicia en el área de trabajo y termina en los paneles de parcheo conectándose a los conmutadores mediante latiguillos, lo que permite modificar cualquier configuración de la conexión de forma rápida y segura.

La longitud máxima para un cable horizontal es de 90 metros.

El número de puntos de conexión en una instalación se determina en función de la superficie útil o de los metros lineales de fachada. Se aplica siguiente normativa: habrá punto de acceso (toma de red) por cada 8 ó 10 metros cuadrados útiles o por cada 35 metros lineales de fachada.

Las opciones de encaminamiento hacia el área de trabajo, más habituales son:

- ⇒ Canaletas por la pared.
- ⇒ Falso techo. El cable transita a través del falso techo.
- ⇒ Falso suelo. El cable se encamina mediante conductos por debajo del suelo.

Subsistema de cableado vertical

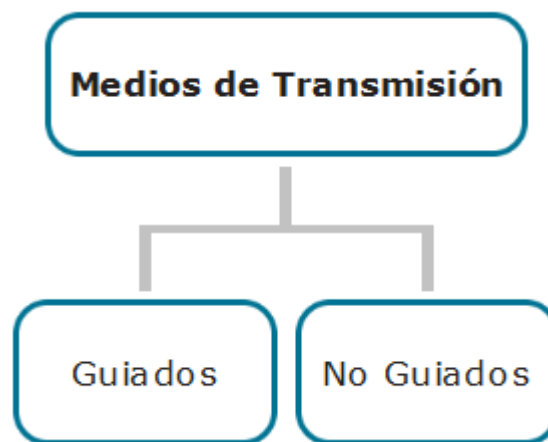
Posteriormente, hay que interconectar los armarios de distribución de cada planta mediante otro conjunto de cables, que deben atravesar verticalmente el edificio de planta a planta.

Esto se suele hacer a través de las canalizaciones existentes en el edificio, pero si esto no es posible, es necesario habilitar nuevas canalizaciones o aprovechar algunas ya existentes (huecos de ascensor o escaleras).

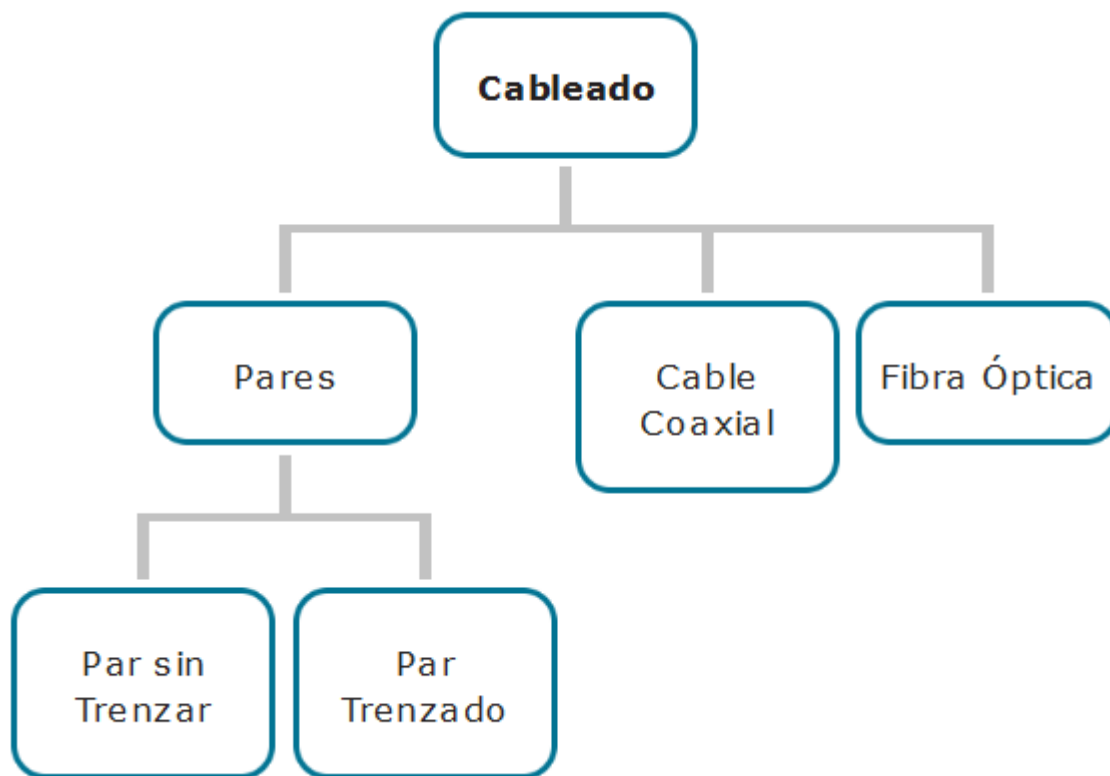
Obsérvese que el cableado para el subsistema vertical debe soportar el tráfico de datos de todas las plantas.

Por tanto, es necesario utilizar un cable que permita transportar mayor número de datos, es decir, mayor ancho de banda. Por ejemplo, cable de fibra óptica o cable trenzado multipar (de 25, 50 ó 100 pares) categoría 5 ó 5e.

Los medios de Transmisión de datos:



Medios Guiados: cableado (Ya visto en los módulos anteriores)



Cuidados al Ejecutar la Conexión

Es muy imprescindible realizar un buen montaje de los conectores, ya que estos son los puntos más propensos a que el ruido se introduzca a través de ellos.

Algunas recomendaciones generales a la hora de realizar correctamente el montaje de los conectores son:

- ⇒ Mantener el trenzado de los pares lo más cerca posible del punto donde encajan con el conector.
- ⇒ Eliminar solamente el fragmento de pantalla protectora del cable que sea necesario para realizar la conexión. No eliminar más pantalla protectora de la estrictamente necesaria.
- ⇒ No se deben realizar empalmes de los cables. Es más aconsejable sustituir los cables cortos por otro más largo. Siempre se podrán utilizar estos trozos pequeños para los latiguillos del armario Rack.
- ⇒ Dejar una longitud suficiente para que el cable no esté forzado ni tirante. Es importante tratar de que el cable nunca tire del conector.
- ⇒ Un cable UTP nunca debe almacenarse torcido. Deben deshacerse todas las torceduras que haya.

- ⇒ Algo que el estándar destaca es que la torsión del cable no debe superar cierta medida: para cables UTP ésta es de 4 veces el diámetro del cable.
- ⇒ No se debe usar cable UTP para conexión Ethernet en exteriores, ya que este no está recubierto. Esto lo hace extremadamente flexible, pero en contra tenemos que aumenta considerablemente la influencia de la impedancia y la temperatura externas e incluso de la humedad ambiental. En su lugar debe usarse cable STP.
- ⇒ Tratar de no usar cable UTP para cableado en el suelo ya que es muy sensible a la torsión por alargamiento además del aplanamiento y posible corte de las fibras.
- ⇒ Aunque los pares Par A y Par D no se utilizan, no deben nunca dejarse desconectados, para así evitar posibles cortocircuitos.
- ⇒ La parte más difícil del cableado es aplanar y enrasar los pares para engastarlos a los conectores RJ11 o RJ45 según sea el caso.
- ⇒ La forma más adecuada de hacerlo es destrenzar cada par y aplanarlo por separado, y luego utilizar un peso para alinear las puntas.
- ⇒ Es posible combinar dos redes en un solo cable UTP, usando los pares T/R1 y T/R4 como los pares de la segunda red; sin embargo para hacerlo es necesario destrenzar y volver a trenzar correctamente entre sí los dos pares para que la impedancia resultante de una red no afecte la otra. En cualquier caso es preferible utilizar cables distintos para redes distintas.

Montaje de los conectores RJ-45 macho

Los pasos que deben seguirse para montar un conector RJ-45 macho son los siguientes:

1. Introducir la funda protectora del conector por el extremo del cable. Si lo que se va a montar es un cable cruzado para conexión directa, es conveniente utilizar fundas de un color diferente para así poder diferenciarlas.
2. Eliminar un trozo no muy grande de la funda protectora del extremo del cable, utilizando alguna herramienta de corte. Con eliminar unos 3 cm. es suficiente.
3. Retirar el plástico protector y si tiene pantalla protectora de cobre, enrollarla sobre el plástico protector.
4. Destrenzar los pares y estirarlos. Ordenar los cables según los colores establecidos por las normas TIA586A o TIA586B, según sea para conexión a Switch o conexión directa.
5. Cortar todos los extremos de los hilos para que queden todos enrasados a la misma altura. Es muy importante que todos queden a la misma altura o de lo contrario puede que no se conecten correctamente algunos de ellos.
6. Introducir todos los hilos en el conector, haciendo fuerza para asegurar que lleguen hasta el fondo. Así mismo, asegurar que la pantalla se introduce ligeramente en el conector para que haga contacto con el chasis metálico. Si no hiciera contacto, habría que cortar un poco más los cables.
7. Después de haber comprobado que la pantalla hace correctamente contacto con el conector y que todos los hilos llegan hasta el fondo, fijar el conector al cable, utilizando la crimpadora para ello.
8. Por último, colocar la funda protectora sobre el conector.

Montaje de los conectores RJ-45 hembra

Los conectores RJ-45 hembra se utilizan en enchufes de pared y paneles de distribución.

El montaje de estos conectores se realiza siguiendo estos pasos:

1. Eliminar un trozo no muy grande de la funda protectora del extremo del cable, utilizando alguna herramienta de corte. Con eliminar unos 3 cm. es suficiente.
2. Retirar el plástico protector y si tiene pantalla protectora de cobre, enrollarla sobre el plástico protector.
3. Doblar el terminal de tierra hasta situarlo en el interior del conector.
4. Destrenzar los pares y estirarlos. Ordenar los cables según los colores establecidos por las normas TIA586A o TIA586B, según sea para conexión a Switch o conexión directa.
5. Introducir cada hilo en su lugar correspondiente y engastarlos utilizando una herramienta tipo crimpadora. Hay algunos tipos de conectores que no necesitan herramientas de engaste, sino que su montaje se realiza ensamblando varias piezas del componente.
6. Colocar la carcasa protectora del conector y montarlo sobre el enchufe de pared.

Los parámetros más importantes a medir

Cuando se realiza la certificación de una instalación de cableado en una red hay que medir una serie de parámetros.

Estos parámetros se pueden medir en cualquiera de los extremos del cable.

Los parámetros más importantes a medir son los siguientes:

⇒ **Continuidad:**

Nos indica si hay alguna rotura en algún punto del cable. Este parámetro de continuidad se puede medir con un simple comprobador de continuidad. Si en vez de esto, utilizamos un comprobador tipo TDR, además de saber si hay alguna rotura, también podremos saber en qué punto exacto del cable está.

⇒ **Mapeado de hilos:**

Se usa para comprobar si los cables están correctamente montados o si hay algún cortocircuito en algún punto de este.

Este es un error muy grave ya que impide que la información llegue de un extremo al otro. Este parámetro se puede medir con un comprobador de continuidad.

⇒ **Resistencia:**

Se trata de medir la resistencia al paso de la corriente eléctrica a través del cable. Un alto valor de la resistencia indica que puede haber una atenuación alta de la señal. Esto puede ser consecuencia de algún defecto en el cableado o en el montaje de los conectores.

⇒ **Longitud:**

Es la distancia que existe entre los dos extremos del cable. La longitud de las conexiones no debe exceder la establecida en las recomendaciones de los fabricantes y de los estándares.

Si la longitud es mayor puede haber retardos excesivos en las señales y los protocolos de comunicaciones pueden fallar.

⇒ **Atenuación por inserción:**

Esto se produce por la pérdida de energía de la señal cuando ésta atraviesa un medio que tiene una determinada resistencia a su paso.

Este valor se calcula dividiendo la energía de la señal de entrada entre la energía de la señal de salida y se mide en decibelios.

Un valor elevado puede ser debido a:

- La longitud excesiva del cableado del enlace.
- Los conectores no están correctamente montado.
- La temperatura es elevada (puede influir en la atenuación de algún tipo de cableado, especialmente los que están compuestos de PVC).

⇒ **Diafonía:**

La diafonía es un fenómeno que se produce debido a que la corriente eléctrica que circula por un cable, crea un campo electromagnético a su alrededor. Este campo puede interferir en la señal eléctrica que puede circular por otro cable que se encuentre cercano a este. Como consecuencia, pueden aparecer distorsiones en esta segunda señal.

Es el típico caso de escuchar otra conversación telefónica de fondo, mientras hacemos del teléfono fijo.

Cuanto más trenzado esté el cable, menor será el efecto de diafonía en él.

⇒ **Diafonía del extremo cercano:**

Es la diafonía proporcionada por la diferencia entre la cantidad de señal de un cable y la cantidad de señal que se acopla en otro cable y que vuelve en el sentido contrario al de circulación de la señal original.

Este valor se mide en decibelios. Este valor debería ser alto, lo que indica que hay poca diafonía. Para controlar este valor hay que mantener los cables trenzados hasta lo más cerca posible de los conectores.

⇒ **Diafonía del extremo lejano:**

Es un parámetro muy parecido al anterior. Mide la diferencia de la señal con la señal acoplada en otro cable que va en el mismo sentido.

⇒ **Igualdad de nivel de diafonía del extremo lejano:**

Para calcular este valor se transmite una señal por un determinado par y se mide la diafonía en otro de los pares. Hay que tener en cuenta que en un cable de cuatro pares existen doce medidas posibles. Un valor bajo de este parámetro indica una buena instalación del cableado.

⇒ **Ratio de atenuación a diafonía:**

Es un valor calculado que mide la diferencia entre el valor de diafonía del extremo cercano y el valor de atenuación en la línea. Este valor nos indica si la potencia de la señal mayor que la potencia del ruido de fondo. Este parámetro debe ser lo más alto posible.

⇒ **Pérdida por retomo**

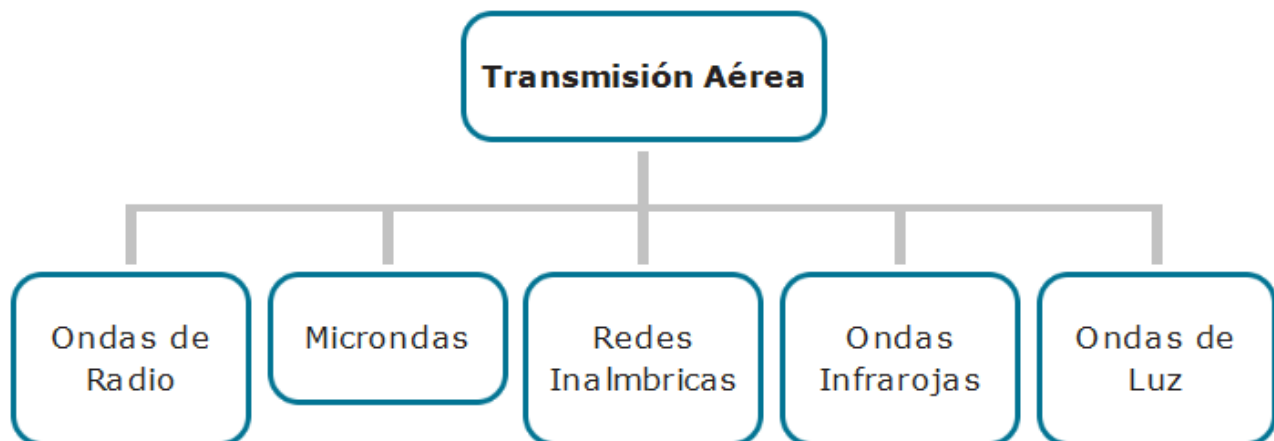
Este fenómeno se produce debido a la diferencia de resistencia al paso de la corriente eléctrica que existe en una línea (impedancia). Estas diferencias suelen aparecer en varios puntos del cableado y lo que se produce es un efecto de rebote de parte de la señal. Esta que vuelve al extremo de origen.

Lo que se mide es la cantidad de señal que vuelve del cable, expresado en decibelios, y su nivel de aceptación depende del tipo de medio.

Cuando existe un valor excesivo de pérdida por retorno suele ser porque las longitudes de destrenzado de los cables en los conectores son muy altas o no se ha mantenido hasta el extremo la pantalla protectora.

Medios Inalámbricos (Vistos en módulos anteriores)

Dependiendo de la frecuencia radioeléctrica de la señal, hay diferentes tipos de conexiones inalámbricas. Cada una de ella tiene propiedades y usos diferentes.



Ondas de radio

Las ondas de radio se generan fácilmente, recorren distancias largas, penetran fácilmente en los edificios y pueden viajar en todas direcciones desde la fuente emisora.

Cuando estas redes a base de ondas de radio cubren distancias largas, las autoridades en telecomunicaciones deben realizar un estricto control para que las diferentes transmisiones no se interfieran entre sí.

En aquellas redes que usan señales de radio y que cubren distancias más cortas, no es necesario solicitar permisos especiales.

El gran problema de las ondas de radio es el elegir bien las frecuencias de emisión, ya que existe un riesgo importante de solapamiento de las comunicaciones.

No están muy recomendadas para la transmisión digital, por no ser muy fiables.

Ondas infrarrojas

La transmisión mediante ondas infrarrojas se utiliza principalmente para las comunicaciones de corto alcance como por ejemplo en mandos a distancia de televisores, equipos de audio, etc. Los ordenadores portátiles también suelen tener un puerto de comunicación infrarroja.

Estos dispositivos infrarrojos a distancia son baratos, pero tienen un inconveniente importante: no atraviesan los obstáculos que se pueden encontrar en su camino.

No es necesario tener ningún tipo de licencia administrativa para poder trabajar con un sistema de transmisión mediante ondas infrarrojas.

Estos medios de comunicación infrarroja no pueden utilizarse en el exterior de locales cerrados, ya que el Sol emite una cantidad considerable de radiaciones infrarrojas que interfieren en las señales enviadas.

Ondas de luz

Las ondas de luz también permiten la comunicación entre dos dispositivos, siempre y cuando haya una visión directa entre ellos, ya que la transmisión es en línea recta y además no atraviesa los objetos.

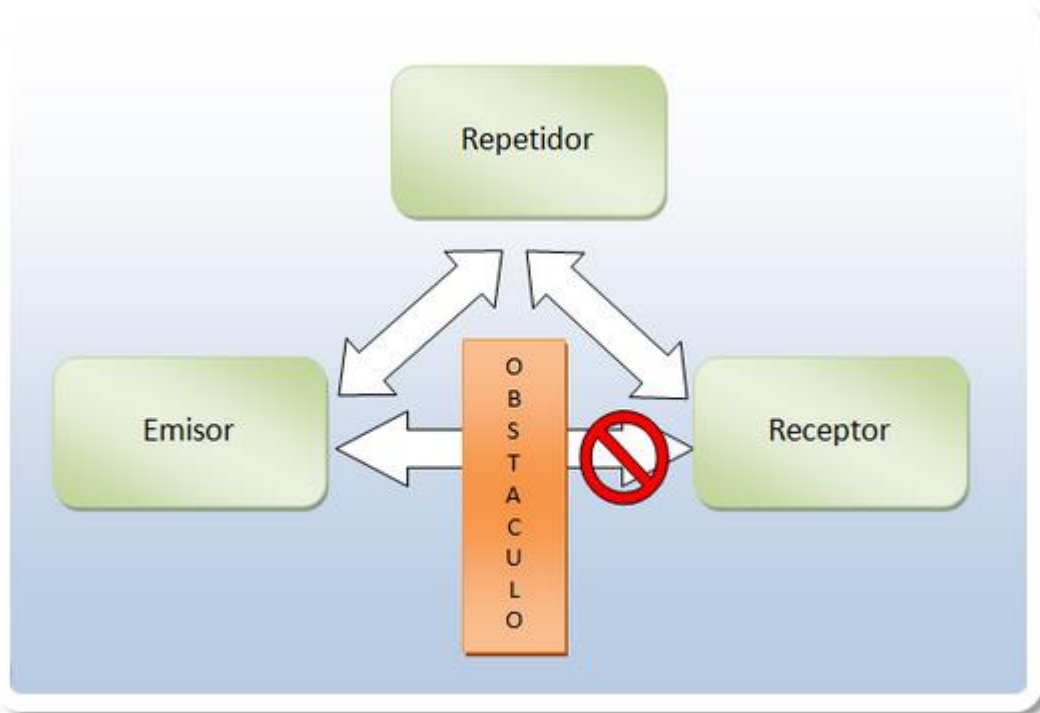
Una de las señales lumínicas más utilizadas es el láser. No hay dispersión alguna en el haz de luz, que se mantiene enfocado en un punto muy estrecho a lo largo de todo el recorrido.

Para la comunicación laser solo se necesita un emisor del haz laser y un receptor. Esto tiene un coste bastante bajo y una velocidad de transmisión alta.

En su contra, tenemos que la instalación de los emisores y receptores es algo complicada, ya que han de estar perfectamente alineados. Además se debe tener en cuenta que el rayo láser es muy sensible a interferencias debidas a la humedad, la niebla densa o las corrientes de aire de convección.

Microondas

Las transmisiones con microondas pueden ser tanto terrestres como vía satélite. La frecuencia en la transmisión oscila entre 1 y 10 GHz y la velocidad es del orden de 10 Mbps. A frecuencias superiores a 1.000 Hz, las microondas viajan en línea recta y, por tanto, se pueden enfocar en un haz de bastante estrecho.



El concentrar toda la energía de la emisión de microondas en un pequeño haz, va a proporcionar una relación señal/ruido muy alta. Esto proporciona una amplitud muy pequeña del ruido.

En este tipo de transmisiones, las antenas del emisor y el receptor deben estar perfectamente alineadas entre sí.

Las ondas de radio, atraviesan bien los obstáculos. Las microondas, las ondas infrarrojas y las ondas de Luz, no. Si queremos realizar comunicaciones a distancias largas con microondas, debemos usar repetidores de señal.

Las redes locales inalámbricas (con tecnología Wi-Fi) utilizan microondas para transmitir la información.

En la actualidad, la mayoría de redes inalámbricas emiten a una frecuencia de 2,4 GHz, si bien cada vez más se hacen transmisiones en la banda de frecuencias de 5 GHz. Estas bandas están menos saturadas y además ofrecen transmisiones con cada vez menos ruido.

Debido a que la capacidad de transmisión de una red inalámbrica está directamente afectada por la potencia de la señal, ésta puede adaptar la velocidad de transmisión a la "cobertura" disponible.

Cuando la señal es más débil a medida que aumenta la distancia, los protocolos establecen una tasa adaptativa, por lo que la velocidad de transmisión disponible puede reducirse.

Una red inalámbrica transmite la información en la banda de frecuencia de 2,4 GHz. En esta frecuencia, se han establecido diferentes rangos, denominados canales, que se pueden utilizar para realizar las comunicaciones.

Canales de redes inalámbricas:

Canal	Frecuencia (GHz)
1	2.412
2	2.417
3	2.422
4	2.427
5	2.432
6	2.437
7	2.442
8	2.447
9	2.452
10	2.457
11	2.462
12	2.467
13	2.472

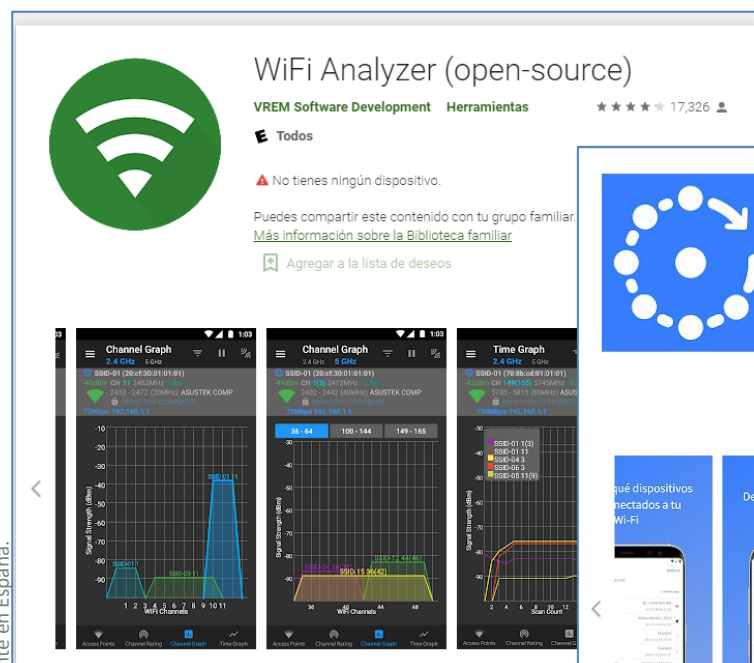
En la tabla anterior se enumeran los diferentes canales que se pueden utilizar en las redes inalámbricas instaladas en la actualidad. Se puede observar que los canales están separados en un rango de 5 MHz.

Al instalar una red inalámbrica hay que tener mucho cuidado a la hora seleccionar el canal de comunicación sobre el que va a funcionar. Dos redes cercanas no pueden transmitir en el mismo canal, porque se crean interferencias entre ellas. Esto suele dar muchos problemas de conectividad.

Tampoco es conveniente transmitir en canales contiguos, ya que cada canal tiene un rango de transmisión de 22 MHz y sólo una separación de 5 MHz con los canales anterior y siguiente.

Esta situación también hay que evitarla dentro de una misma red cuando existen varios puntos de acceso cercanos.

Existen muchos programadas para verificar, canales, dispositivos conectados...: Fing, Wifi Analyzer...



Fing ha ayudado a 40 millones de usuarios en todo el mundo a entender:

- ¿Quién está usando mi WiFi?
- ¿Alguien está robando mi WiFi y mi banda ancha?
- ¿Me han hackeado? ¿Es segura mi red?
- ¿Hay cámaras ocultas en el B&B en el que me estoy quedando?
- ¿Por qué Netflix ha empezado a cargar?

Instalaciones eléctricas dedicadas (I.E.D.):

Son instalaciones de uso exclusivo para el equipamiento del sistema de cableado estructurado y los equipos de la red. Dispondrán del sistema de protección de toda la infraestructura y toma de tierra para protección frente a sobrecargas.

Sus principales características son:

- No se comparte del suministro con el resto de los circuitos de las instalaciones, como por ejemplo alumbrado.
- La I.E.D. básica suministra energía a la electrónica de red y a los servidores.
- La I.E.D. ampliada suministra además energía a los distintos puestos de trabajo.
- Se divide en dos circuitos: de corriente SAI o ininterrumpida y de corriente no-SAI.

S.A.I.:

SAI son las siglas en español de Sistema de Alimentación Ininterrumpida.

Un SAI es un dispositivo basado en baterías que suministra energía a los distintos dispositivos electrónicos en ausencia de suministro eléctrico normal.

A todos nos ha pasado alguna vez que mientras estamos trabajando con el ordenador hemos sufrido un corte de electricidad y nos ha sorprendido sin poder guardar el trabajo que estábamos haciendo. Situación muy desagradable.

Un SAI sirve precisamente para evitar eso: pérdidas accidentales de información debida a la interrupción del suministro eléctrico. Hay algunos modelos que también protegen de algunos otros problemas que tiene el suministro de corriente eléctrica como pueden ser los picos o caídas de la tensión.

El suministro de corriente eléctrica se basa en la corriente alterna (AC), ya que ésta corriente se transporta mucho mejor que la corriente continua (DC).

En España, el suministro de corriente proporciona una tensión de 230 Voltios (V), y oscila a 50Hz (es decir, cambia de +230V a -230V, 50 veces por segundo).

Es frecuente que el suministro de energía eléctrica sufra fluctuaciones. Estas pueden afectar y mucho al funcionamiento de los aparatos electrónicos (que necesitan de un suministro constante). Esto puede afectar, no solo a su funcionamiento, sino que también puede afectar a integridad.

Los SAI tienen en su interior una batería que acumula energía eléctrica. El SAI se conecta a la red y con esa energía de la red, va recargando dicha batería. Todos aparatos eléctricos que se quieran proteger, se deben conectar a la SAI.

Si hay una caída del suministro de electricidad en la red, el SAI proporcionará la energía necesaria a los dispositivos que tenga conectados, sacándola de la batería, y evitando de esta forma que se apague.

Funcionamiento de un S.A.I.:



La funcionalidad de un SAI es ofrecer el suministro eléctrico cuando se produzca corte de este por parte de las compañías eléctricas. A lo largo del tiempo han evolucionado y se han perfeccionado considerablemente, ofreciendo además protección ante otros problemas que presenta el suministro de la corriente eléctrica.

La SAI se conecta al suministro de la red eléctrica y los diversos equipos de la red se conectan a él. Veamos a continuación los problemas más frecuentes que presenta en el suministro eléctrico.

Problemas derivados del suministro eléctrico

- La ausencia de suministro: no hay suministro, se pierde completamente la tensión eléctrica. Como consecuencia, se apagan todos los equipos de la red y, lo que es más grave, hay pérdida de información.
- Microcorte: se produce ante la ausencia de suministro eléctrico durante un periodo muy breve de tiempo. Si se produce un microcorte, puede que los equipos sigan funcionando, pero lo normal es que se apaguen y se vuelvan a reiniciar. Lo más grave es que si vienen acompañados de subidas de tensión, los circuitos electrónicos internos pueden sufrir daños definitivos.
- Bajadas de tensión: se produce cuando baja la tensión eléctrica. Si el voltaje no es suficiente durante un periodo de tiempo, los circuitos electrónicos de los distintos equipos pueden funcionar incorrectamente y por lo tanto se pueden producir cuelgues, reinicios o malfuncionamiento. Incluso se pueden apagar los equipos o se pueden reiniciar.
- Subidas de tensión: es lo contrario al punto anterior. El suministro eléctrico llega con más voltaje de lo esperado. Se pueden dañar los circuitos, provocando daños permanentes en ellos.
- Cuando se producen cortes en el suministro, microcortes o bajadas de tensión, peligra la integridad de la información, aunque no suele haber problemas físicos. En este caso es el hardware el que puede sufrir daños.
- Picos de tensión: son fuertes subidas de tensión eléctrica. Sube bastante el voltaje, pero lo hace durante un periodo de tiempo muy corto. Es una cuestión bastante problemática, pues puede provocar daños irreversibles en el hardware.
- Ruido eléctrico o interferencias: se produce cuando la forma de la onda de la corriente alterna sufre deformaciones. La corriente alterna tiene una forma de onda sinusoidal pasando de los -240 V a los +240 V suave y constantemente. El ruido consiste en modificaciones puntuales de la frecuencia de cambio. Esto se puede deber a que el tendido eléctrico pase cerca de algún motor eléctrico. Éstos producen una inducción electromagnética que provoca ese ruido en la transmisión de la electricidad. No es especialmente grave para los ordenadores.
- Aumento o descenso de la frecuencia: no es muy frecuente, Se produce cuando la corriente alterna no cambia a la frecuencia que debe de hacerlo (50 Hz en España, 50 veces por segundo).

La batería, el cargador y el inversor:

Las baterías de los SAI son muy similares a las baterías de los coches. Normalmente están fabricadas de plomo y ácido. Tienen una vida limitada, de unos pocos años (en teoría), después se degradan y acaban por no funcionar.

No se suelen las baterías de Níquel, las Hidruro Metálico o las de iones de litio, Estas presentan el problema del tiempo de respuesta ante el corte del fluido eléctrico, además son bastante más caras.

A través de la red eléctrica viene corriente alterna (AC). La SAI almacena corriente continua (DC). Por lo tanto lleva en su interior un circuito llamado cargador, que convierte la AC en DC y de esta forma puede ir almacenándola.

De igual forma, si la SAI debe suministrar AC a los equipos, pero la batería sólo puede proporcionar DC, se necesita un circuito llamado inversor, que recoge la DC de la batería, y la convierte en AC, cambiando también la tensión de salida (desde los 12 o 24V que suele tener una batería, a los 220V necesarios para un equipo).

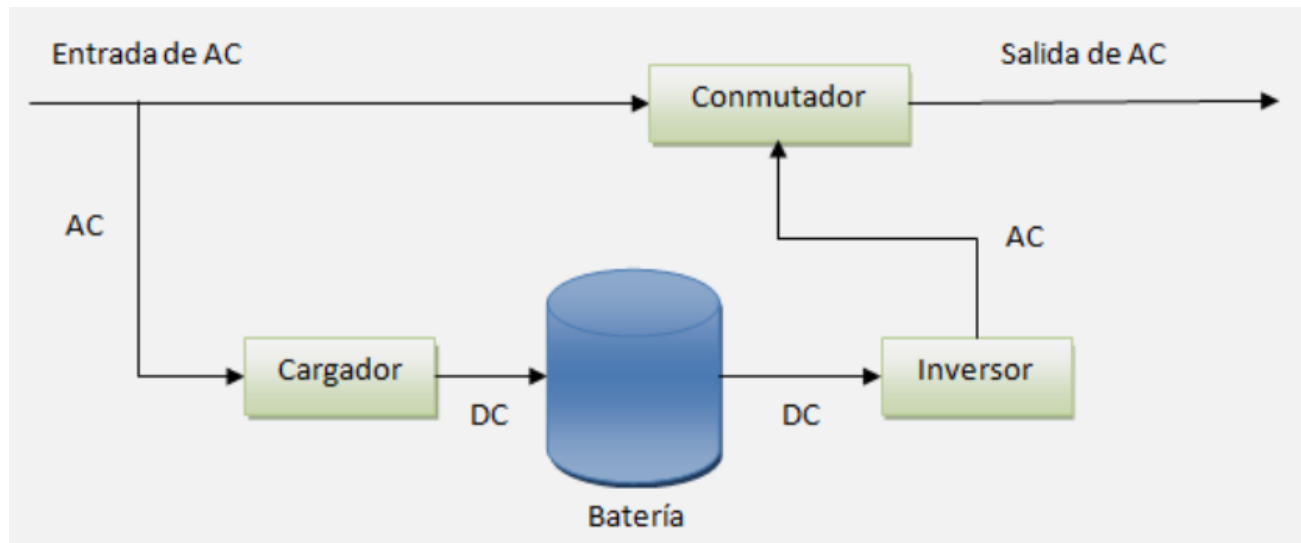
Tipos de SAI:

Hay varios tipos de SAI, y cuentan con todo tipo de circuitos auxiliares. Hay dos grandes grupos: Los off-line y los on-line.

Off-Line pasivos, también llamados "Stand-by"

Los SAI off-line son los más comunes, los más sencillos y los más económicos. Tienen de un cargador que suministra carga a la batería a partir de la corriente alterna de la red eléctrica, y un inversor que proporciona corriente continua a los dispositivos conectados a esa batería.

El SAI Off-line sólo proporciona energía de su batería cuando hay una ausencia de suministro, un microcorte o una bajada de tensión importante. Para ello, hay un circuito conmutador que se encarga de detectar cuando se producen estos problemas, iniciando así el suministro desde la batería.

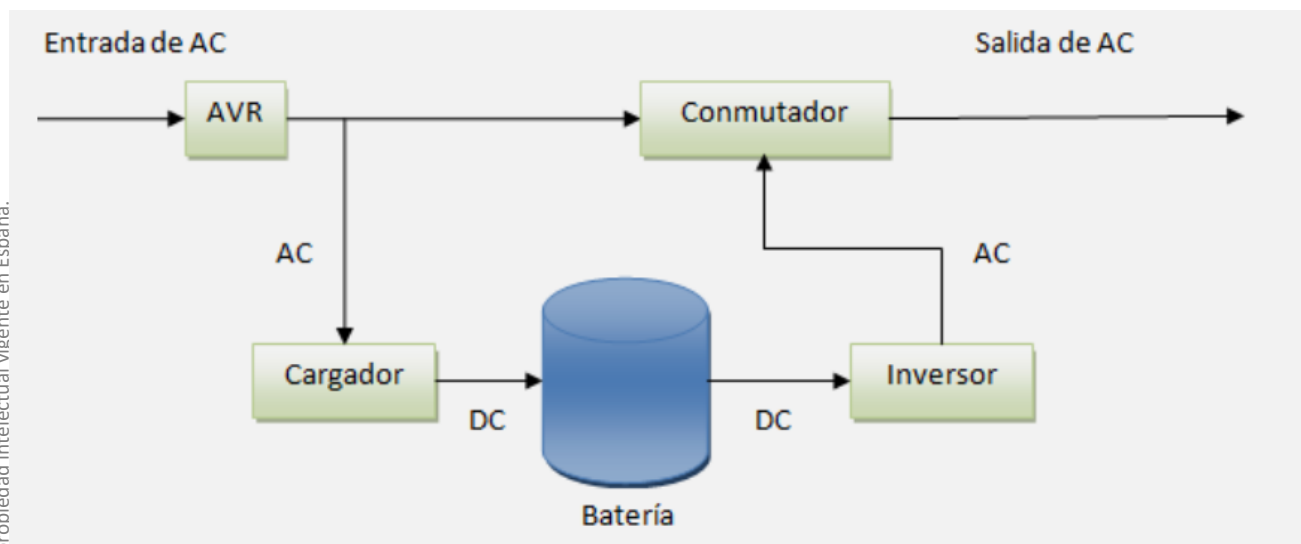


El conmutador tarda un tiempo en detectar el problema y comenzar el suministro desde la batería. Normalmente este tiempo es corto pero durante el, no hay suministro de energía. A este tiempo, se le llama tiempo de conmutación o de transferencia. Es este el gran problema que presenta este tipo de SAI.

Off-Line interactivos, o con estabilizador o AVR

El estabilizador (AVR) es un circuito que suministra a la salida una tensión eléctrica constante (220V), independientemente de las fluctuaciones que tenga la entrada.

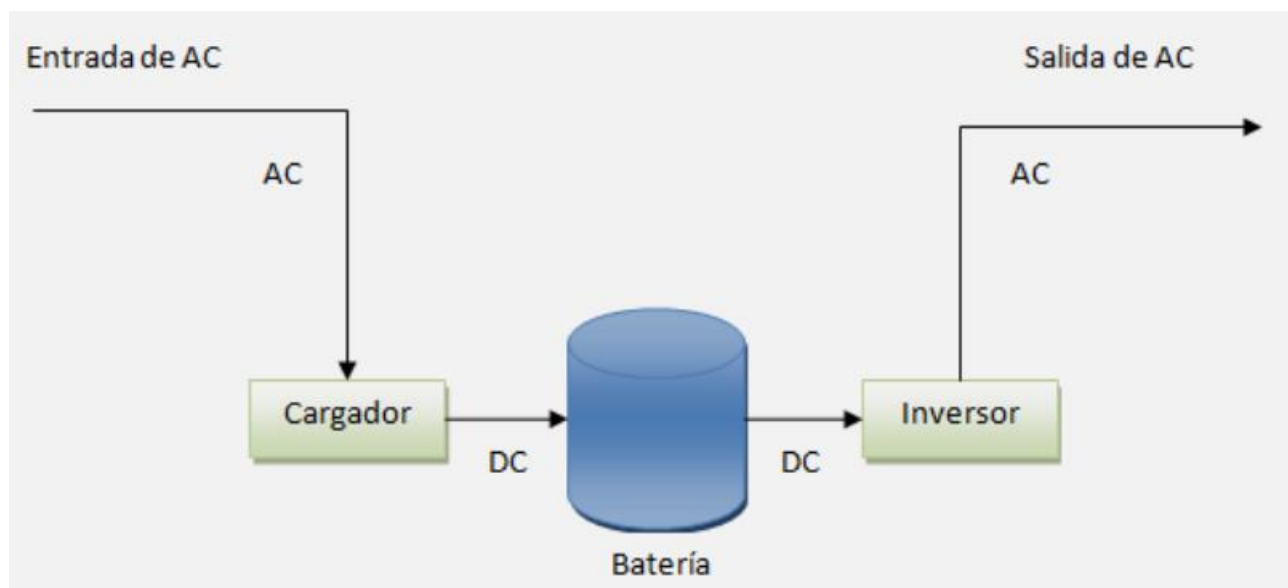
Un SAI interactivo tiene de un estabilizador en la entrada de corriente alterna, que proporciona la energía estabilizada sin fluctuaciones al cargador de la batería y por lo tanto, la salida de AC tiene también una tensión constante, sin fluctuaciones.



On-line, en línea o de doble conversión

Son los SAI más avanzados y también de más coste. No hay conmutador, y en la salida la corriente nunca se obtiene directamente de la entrada.

La batería está cargándose continuamente y la salida se obtiene siempre de la batería, nunca de la entrada.

**Sustitución de la batería:**

El procedimiento de sustitución de la batería se realiza con el SAI apagado, y el equipo que protege desconectado. Esto muchas veces puede ser inviable por diversas razones. Para evitar el parón en el servicio de los equipos, se usa lo que llamamos el **bypass activo**.

Con el bypass activo, conseguimos que la corriente pase directamente al equipo protegido desde la entrada de la red. Mientras se desconecta la batería, es una especie de puente. Su objetivo es que se pueda sustituir la batería sin necesidad de apagar el SAI ni el equipo que protege.

El bypass debe estar activado el tiempo estrictamente necesario para la sustitución de la batería, puesto que supone un riesgo para la seguridad (imaginemos un corte de suministro durante el proceso).

La sustitución de la batería es un procedimiento bastante frecuente en las labores de mantenimiento de cualquier SAI. Los fabricantes de estos proporcionan baterías de recambio, dado la vida limitada de las baterías de plomo y ácido (aquí también entra en juego la llamada obsolescencia programada, pero eso es otra cuestión ajena a este certificado de profesionalidad).

Los pasos a seguir para la sustitución de la batería son los siguientes:

- Apagar el equipo y el SAI. Si la SAI dispone de bypass, lo mejor es activarlo. Localizar el compartimento de la batería y abrirlo.
- Extraer la batería, y desconectar el polo negativo. El polo negativo es el que tiene el cable negro, y el signo "-". Es conveniente desconectar primero éste porque la carcasa del SAI también va conectada ahí, y de esa forma se evitan descargas accidentales al tocarla.
- A continuación, se desconecta el polo positivo. Normalmente tiene el cable de color rojo y el signo "+".
- Se conectan los cables a la nueva batería, primero el positivo y luego el negativo.
- Se coloca la batería en su compartimento y se cierra.
- Se desactiva el by-pass (si se hubiese activado), o se pone el SAI y el equipo en marcha, si los hubiéramos apagado.

Capacidad y tiempo de suministro:

Los SAI no están pensados para proporcionar energía a los distintos equipos durante un tiempo ilimitado. La idea es mantener la corriente el tiempo necesario para que se puedan guardar los datos. Es decir, estamos hablando de minutos.

Este proceso de guardar datos se puede hacer manualmente o bien automáticamente.

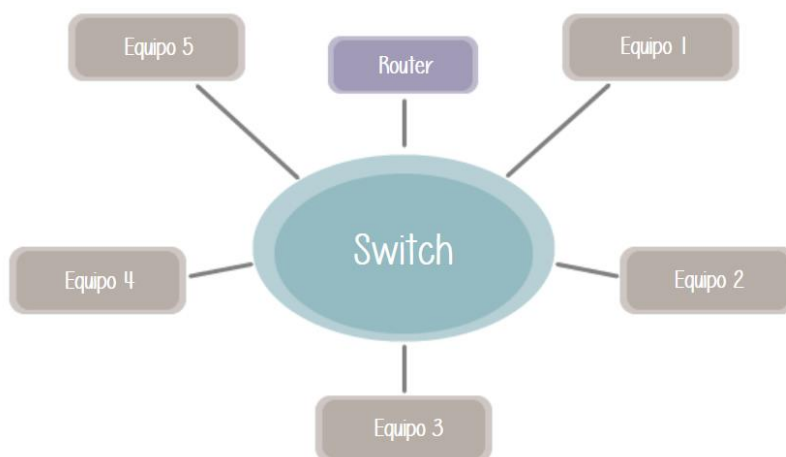
Para poder hacerlo automáticamente, casi todos los SAI se conectan al equipo mediante un cable USB. En el equipo se instala un software de control que es el que se encarga realizar las tareas de guardado de datos cuando el SAI dispara una alarma indicando el corte de suministro eléctrico.

Mediante ese cable también se suelen obtener otros datos del SAI, como la cantidad de carga de las baterías, o el tiempo que podrá suministrar energía el SAI si se va la corriente. De esta forma se puede monitorizar el SAI y controlar su funcionamiento desde el ordenador.

Los fabricantes de SAI proporcionan la capacidad de estas, utilizando una unidad de medida poco conocida: el Voltio-Amperio (VA). Esa unidad está relacionada con la potencia que es capaz de proporcionar el SAI.

Otros elementos Hardware a revisar: (Ya visto en módulos anteriores)**Concentradores – Hub**

Switch: Topología en estrella de las redes locales:

**Alimentación Eléctrica por Ethernet (PoE):**

La Alimentación Eléctrica por Ethernet (Power Over Ethernet, PoE), es una tecnología que permite el envío de alimentación eléctrica junto con los datos en el cableado de una red Ethernet. Es algo así como los dispositivos móviles que a través de un cable USB pueden comunicarse con el ordenador a la vez que pueden cargar su batería.

En año 2003 se publicó el estándar IEEE 802.3af que definía la primera versión de esta tecnología. En el año 2009 se publicó una revisión y ampliación en el estándar IEEE 802.3at.

La tecnología PoE permite suministrar alimentación eléctrica a dispositivos conectados a una red Ethernet, simplificando en gran medida la infraestructura de cableado necesaria para su funcionamiento.

De esta forma, un dispositivo que soporte la tecnología PoE obtendrá tanto los datos como la alimentación eléctrica por el cable de red Ethernet. Recordemos que de los 4 pares del cable RJ-45, dos de ellos están libres y se pueden utilizar para la alimentación eléctrica.

Entre los dispositivos que utilizan esta tecnología cabe destacar: puntos de acceso inalámbricos Wi-Fi, cámaras de video IP, teléfonos de VoIP, Switch remotos y en general cualquier dispositivo que esté conectado a una red Ethernet, que no tenga un consumo eléctrico muy elevado.

En el mercado se pueden encontrar multitud de modelos de Switch que incluyen puertos Power Over Ethernet.

Router: (Ya vistas las funciones básicas)

Otras funciones de los Routers:

Si bien la función principal de un Router, es la del encaminamiento del tráfico de red, suelen proporcionar otras funcionalidades:

- Adaptación de los datos entre diferentes tecnologías de transmisión:
Pensemos en los Routers domésticos que suministran los proveedores de internet. Estos Routers realizan el intercambio de datos entre la red interna del usuario que utiliza tecnologías como pueden ser Ethernet y Wi-Fi y la red de acceso de los operadores de internet, que utilizarán tecnologías de última generación como ADSL, cable (HFC) o fibra óptica (FTTH).
- Proporcionar los parámetros de configuración de red:
Esta función se lleva a cabo mediante el servicio DHCP. Simplifica mucho la conexión de un dispositivo a la red ya que todos los parámetros de red se configuran de forma automática.
- Filtrado de datos:
Esta función se debe a cuestiones de seguridad. Este filtrado de datos se lleva a cabo mediante un elemento conocido como cortafuegos.
- Traducción de direcciones de red:
Debido a la escasez de direcciones IP prácticamente todas las redes utilizan un mecanismo de traducción de direcciones de red conocido como NAT que permite el uso de direcciones privadas en redes conectadas a Internet.
Esta función la suele realizar el Router, especialmente en los Router domésticos.
- Servidor proxy:
Un servidor proxy es un equipo que actúa de intermediario entre un navegador web e Internet. Estos servidores proxy mejoran el rendimiento de Internet puesto que lo que hacen es almacenar una copia de las páginas web más utilizadas en una caché.
Cuando un navegador solicita una página web almacenada en la caché del servidor proxy, el servidor proxy la proporciona, sin necesidad de acceder a la red. Los servidores proxy también ayudan a mejorar la seguridad, ya que filtran algunos contenidos web y software malintencionado.
- Punto de acceso inalámbrico (Wi-Fi).
- Redirección de puertos.
- Balanceo de carga/tráfico.
- Gestión de conexiones VPN.

Veamos algunas ideas a tener en cuenta a la hora situar un Router inalámbrico para ofrecer una buena cobertura:

- Cuanto más lejos se quiera llegar, más alto se deberá colocar el Router inalámbrico.
- Se debe situar el Router inalámbrico en una posición central de la casa, ya que la cobertura que ofrecen los router es circular.
- Una ventana, un patio de luces, un hueco de escaleras, pueden ayudar a distribuir la señal.

También es importante tener en cuenta los obstáculos que pueden haber. Hay determinados objetos que pueden absorber o reflejar la señal llegando a degradar e incluso bloquear la misma.

Algunos de los posibles obstáculos son:

- Azulejos: Pueden atenuar bastante la intensidad de la señal, más si tienen componentes metálicos. Especial cuidado hay que tener con los azulejos porcelánicos.
- Paredes: Son, quizás, uno de los obstáculos principales. Se debería evitar en lo posible el número de paredes a cruzar, aunque esto no es algo tan sencillo. No solo es importante el número de paredes a atravesar, también lo es la composición de estas. Yeso, hormigón, escayola, madera se comportan de forma distinta con la señal inalámbrica.
- Armarios: Al igual que las paredes, los armarios de madera o escayola también atenúan la señal. Los armarios metálicos, pueden eliminar completamente la señal.

- **Cristal revestido:** El cristal transparente no suele afectar a la señal. Si el cristal está cubierto con una película metalizada o tiene una estructura de alambre en su interior afectan y mucho a la calidad de la señal. Hay que tener especial cuidado con los espejos.
- **Techos:** Si se dispone de falsos techos o altillos, no estaría de más el valorar situar el inalámbrico allí. En este caso hay que tener en cuenta:
Los materiales del techo y del altillo.
Hay que aumentar las medidas de seguridad de acceso, pues se está ofreciendo cobertura al piso superior.
- **Elementos naturales:** Elementos como el agua, árboles y arbustos, que pueden encontrarse en el exterior en un jardín, patio o balcón, también degradan la señal. Especial cuidado hay que tener con los acuarios.
- **Reflexión:** Algunos objetos no absorben la señal sino que la reflejan de la misma forma que un espejo refleja la luz, es el caso de algunas paredes. Esta reflexión se puede utilizar como aliado a la hora de emitir la señal Wi-Fi.
- **Interferencias o ruido:** El ruido provocado algunos dispositivos electrónicos como teléfonos inalámbricos, hornos microondas, etc., pueden interferir en el Router inalámbrico.
Las fuentes de alimentación de los distintos aparatos eléctricos también interfieren bastante en la señal. Frigoríficos, monitores, televisiones, batidoras, etc., atenúan la señal.
- **Encerrar el Router inalámbrico:** No se debe colocar el Router inalámbrico en el interior de un mueble. Mucho menos si este es metálico.
Tampoco se debe situar el equipo sobre una mesa o armario metálico.

1.1.1.2. Fallos Software

En el software de red se incluye todo el software relacionado con la interconexión de los distintos equipos.

Todas aquellas aplicaciones necesarias para que las redes funcionen correctamente.

El software de red hace posible la comunicación entre los ordenadores, permiten compartir recursos (software y hardware) y ayudan a controlar la seguridad en la red.

Dentro de él, incluimos programas informáticos que se encargan de las comunicaciones entre los usuarios de la red y que permiten compartir información (bases de datos, documentos, gráficos, vídeos, etc.) y recursos (impresoras, unidades de disco, sistemas de copia de seguridad, etc.).

De la misma forma que un equipo no puede trabajar sin un sistema operativo, una red de ordenadores no puede funcionar sin un sistema operativo de red. Si este no existe, los equipos no podrán compartir recursos y como consecuencia los usuarios de la red podrán acceder a dichos recursos. En otras palabras: el sistema operativo de red es imprescindible.

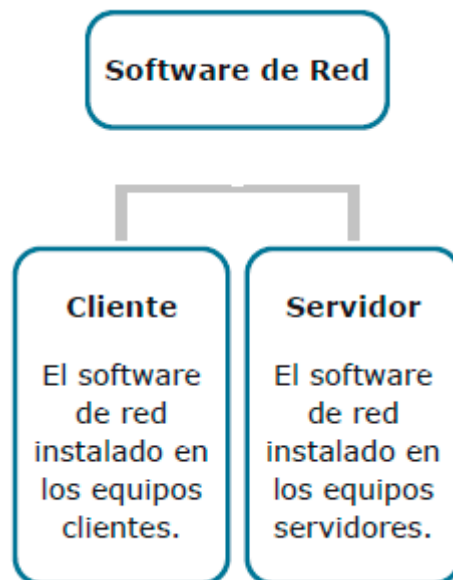
Lo más común es que el sistema operativo de red esté integrado en el propio sistema operativo del equipo. Es el caso de Windows 2000 server/professional, Windows NT server/workstation, Windows 7/8/9, etc.



Un sistema operativo de red:

- Conecta todos los equipos y periféricos de la red.
- Coordina las funciones de todos los periféricos y equipos.
- Proporciona seguridad controlando el acceso a los datos y periféricos.

Las dos componentes principales del software de red son:



El sistema operativo en red controla los recursos, así como la forma de compartirlos y acceder a ellos.

Cuando se planifica una red, la tarea de elegir el sistema operativo de red viene determinada en función de la propia arquitectura de red, bien sea cliente/servidor o de trabajo en grupo. Esto es lo primero que se debería de determinar.

Esta decisión se basa normalmente en cuestiones de seguridad. Las redes basadas en arquitectura cliente/servidor, permiten incluir más medidas de seguridad que las disponibles para una red de trabajo en grupo. Si el tema de la seguridad no se considera relevante, puede resultar más apropiado un entorno de red trabajo en grupo.

Aportaciones de Software de Red:

- Capacidad de compartir información de forma rápida y económica.
- Puesta en marcha inmediata.
- Gracias a las interfaces gráficas, este software resulta realmente fácil de usar.
- Hacen posible que varias personas compartan simultáneamente datos y periféricos.
- Pueden reducir la necesidad de comunicación por escrito, incrementar la eficiencia y hacer que prácticamente cualquier tipo de dato esté disponible simultáneamente para cualquier usuario que lo necesite.

El cortafuegos o firewall:

La seguridad de un ordenador es quizás de lo más importante. Lógicamente en él puede haber almacenada mucha información importante y sensible. Aunque Internet ofrece muchas oportunidades, también esconde muchos peligros.

En general, la principal medida de seguridad es el sentido común. Es de sobra conocido que no debemos aceptar archivos de desconocidos, instalar programas de origen dudoso, acceder a conexiones Wi-Fi abiertas, etc.

Pero no solo esto. Necesitamos programas de seguridad, como los antivirus o los cortafuegos.

Algunos Router proporcionan la protección de un firewall, pero con esto tampoco es suficiente. Debemos de instalar en los equipos de la red cortafuegos y así asegurarnos de que no hay conexiones entrantes al equipo sin la autorización oportuna.

Abrir la red de área local a Internet, que es completamente público, puede ser peligroso puesto que pueden producirse accesos indebidos desde el exterior. Desde accesos de curiosos que solo van a ver que hay dentro del ordenador hasta otros procedentes del espionaje de la competencia.

El problema inverso también se puede dar. Una organización puede restringir los accesos al exterior o, al menos, a un determinado tipo de documentos. No es extraño que se programe el cortafuegos para impedir el acceso a determinadas páginas (redes sociales, páginas pornográficas, etc.).

Para resolver todas estas situaciones, se instala en el acceso a la red un nodo denominado cortafuegos o firewalls, que se encarga de limitar los accesos en ambas direcciones, haciendo invisible la red de área local desde el exterior y restringiendo los accesos desde adentro hacia afuera o viceversa.

Los distintos cortafuegos, operan en los distintos niveles de la arquitectura OSI. Así, un cortafuegos que opere en niveles bajos, será más fácilmente configurable y mucho menos flexible. Otros firewalls que operan en las capas superiores de la arquitectura, investigan el contenido de cada paquete de datos. Esto los hace bastante lentos pero extraordinariamente fiables.

Un firewall es un software que se usa para impedir el acceso no autorizado a Internet (u otra red) por parte de programas maliciosos que están instalados en nuestro equipo, o de programas o intrusos que intentan atacar nuestros equipos desde el exterior.

Los primeros elementos que ofrecían funcionalidad de cortafuegos fueron los Routers IP. Estos podían filtrar los paquetes de datos en función de las características de la red y de la configuración que establecía el administrador de red, que es el máximo responsable en materia de seguridad.

Después han ido apareciendo cortafuegos que actúan en las capas superiores, con lo que pueden incorporar nuevas funcionalidades que mejoran la seguridad.

Algunos de estas funcionalidades son las siguientes:

Traducción de direcciones: Consiste en que las direcciones IP utilizadas por los servidores de la Intranet sólo tienen validez dentro de la propia red de área local. El cortafuegos sustituye cada dirección IP de la Intranet en los paquetes entrantes y saliente por otras direcciones IP virtuales, protegiendo de esta forma contra accesos indeseados a través de direcciones Internet que realmente no existen en la Intranet.

Protección frente a virus: Al operar en las capas altas, los cortafuegos analizan la información que fluye hacia la Intranet, pudiendo detectar anomalías en los datos y programas.

Auditoría: El firewall puede auditar recursos concretos de la Intranet y avisar a través de un sistema de mensajería electrónica del intento de violación de algún recurso o de accesos indebidos.

Gestión de actividad: A través de agentes SNMP o DMI, propios de gestión de red, se puede monitorizar el firewall con el fin de realizar informes sobre la actividad de la red.

La función principal del firewall es, controlar las conexiones entrantes y salientes a través de la red.

Imaginemos un equipo que se infecta con un programa malicioso. Probablemente intentará acceder al servidor de la red bien para transmitir datos o para descargar otros programas más peligrosos aún. Como ya hemos dicho anteriormente, lo mejor es contar con firewall que vigile y proteja la red.

No debemos caer en las falsas ideas de:

- Atacar a mi ordenador, ¿para qué?
- ¿Qué me pueden sacar?
- No tengo nada importante en mi equipo.
- No, el mío no lo van a atacar.

Este es el error más grave, ya que no siempre los curiosos están buscando obtener información de otros equipos, sino simplemente experimentar, o en algunos casos, causar daños por el simple hecho de que saben cómo hacerlo.

El modo de funcionamiento de un cortafuegos ofrece protección en varios casos fundamentales:

- Impide a usuarios no autenticados enviar paquetes potencialmente peligrosos al interior de la red, normalmente destinados a puertos específicos de los equipos.
- Prohíbe a los usuarios de la red local recibir contenidos no autorizados o programas destructivos.
- No permite que usuarios del exterior puedan obtener información del interior de la red local.

Veamos ahora un ejemplo sencillo de configuración de un cortafuegos para eliminar tráfico no deseado o potencialmente peligroso.

Supongamos que tenemos en la red varios servidores FTP y HTTP a los que deseamos permitir la conexión con usuarios remotos. Además, también tenemos servidores Windows y Linux que necesitamos proteger. Una posible tabla de filtrado que podría configurarse en estas condiciones puede ser la siguiente:

Protocolo	Dirección origen	Dirección destino	Puerto destino	Filtrado
TCP	*	10.0.1.1	21	No
TCP	*	10.0.1.2	80	No
TCP	*	*	*	Si
UDP	*	*	*	Si

- La primera línea especifica que el cortafuegos permite el paso de los paquetes destinados al servidor FTP cuya dirección es 10.0.1.1.
- La segunda línea indica que el cortafuegos también permite el paso de los paquetes destinados al servidor Web 10.0.1.2.
- Las dos últimas líneas especifican que, para los demás servicios y puertos, el cortafuegos deberá descartar los paquetes entrantes.

Además de los paquetes entrantes, el cortafuegos también filtran el tráfico saliente y la conexiones no deseadas. De esta forma, se puede impedir que un usuario de la red interna pueda conectarse a servidores con contenidos poco éticos o potencialmente peligrosos.

Existen muchos productos en el mercado contruidos tanto como dispositivos independientes (que podrían confundirse a primera vista con un puente o un encaminador), como productos software que se instalan sobre el sistema operativo.

Los sistemas Microsoft Windows incluyen una herramienta denominada Centro de Seguridad a la que se accede desde el panel de control. Este es el cortafuegos de Windows.

En el caso de Linux, algunas distribuciones incluyen programas cortafuegos bastante avanzados.

Las distribuciones SuSE Linux también incluyen varios cortafuegos que se pueden instalar en el equipo: Personal-Firewall, un cortafuegos encargado de impedir la recepción de mensajes del exterior, y SuSEfirewall, una utilidad más profesional que ofrece gran cantidad de parámetros de configuración.

¿Cómo funciona un firewall? De fuera a dentro:

En primer lugar, los firewall crean una base de datos con los programas que necesitan acceso a internet en nuestro ordenador. Estos son los llamados programas seguros o permitidos. A estos programas se les permite el acceso. A otros no.

Es muy importante que a la hora de configurar el cortafuegos, el equipo esté limpio de todo tipo de virus, spyware u otros malware, para así evitar que estos, permitan su propio acceso a internet.

La mejor forma de garantizar todo lo anterior, es que el ordenador esté recién formateado. Aunque no siempre es posible, resulta muy útil para poder generar una base de datos limpia para nuestro firewall. Este es un punto a tener muy en cuenta.

Una vez formateado nuestro equipo, comenzaremos a instalar todos los programas que necesitemos, el antivirus que nosotros creamos más conveniente y por último el firewall.

Si instalamos un firewall cuando ya existen programas malignos en el sistema, hay que tener en cuenta que estos podrían alterar el funcionamiento del firewall e incluso no permitir su correcta instalación, para evitar así ser detectados.

Es por esto que siempre es conveniente tener un antivirus y un firewall después de una instalación limpia del sistema operativo.

Cuando los distintos programas empiecen a conectarse a internet, el firewall irá alertando de ello, daremos o no permiso de acceso. Si es un programa sospechoso, no debería permitirle conectarse.

De esta forma el firewall sabrá que queremos que ocurra con el acceso a internet de cada programa.

¿Cómo funciona un firewall? De dentro a fuera:

Los firewalls también alertan y protegen de ataques que recibimos desde fuera de la red. Pero no sólo eso, este software monitoriza la entrada de datos desde Internet u otra red, detectando los posibles intentos de accesos no autorizados o de ataques externos.

Por otra parte, crean registros o logs, guardando toda la información relativa a los ataques y su procedencia, guardando la dirección IP. Con la dirección esta IP, podremos saber la ubicación del equipo atacante.

Configuración:

Hay 2 tipos de configuraciones: básica y avanzada.

- ⇒ Configuración básica: Suele ser bastante fácil de realizar, pero algunos Firewall incluyen más parámetros de configuración. También hay diferencias en tipo de parámetros que se controlan en este tipo de configuración, pero por lo general va a ser siempre de una configuración bastante sencilla.
- ⇒ Configuración avanzada: No es que sea fácil ni difícil, lo que ocurre es que precisa de unos conocimientos algo más avanzados por parte de la persona que vaya a realizar la configuración, ya que debe conocer, no solo el elemento sobre el que se va a establecer una determinada regla, sino que también hay que saber el orden exacto que deben seguir estas reglas para que funcionen correctamente y el efecto que van a tener en nuestras comunicaciones.

1.1.1.3. Fallos de configuración interna/interfaces de interconexión

Direccionamiento a nivel de enlace:

Las redes de área extensa (WAN) utilizan enlaces punto a punto para realizar la comunicación de cada uno de sus nodos. Así forman topologías de red normalmente irregulares. Este tipo de enlaces comunican tan solo dos equipos (uno en cada extremo). Las Redes de área local (LAN), por simplicidad y reducción de costes, utilizan un medio compartido por el que transmitir la información.

Esta importante diferencia hace que los protocolos de comunicación de bajo nivel en LAN y WAN sean bastante diferentes.

Imaginemos un ordenador que envía un mensaje a través de una red de difusión. Este mensaje permanece en el medio compartido en espera de que el destinatario puede pasar a recogerlo.

Anteriormente hemos dicho que todos los dispositivos están conectados al mismo medio de transmisión, es por ello que todos ellos pueden "ver" ese mensaje enviado (al menos a nivel físico). Pero solo el nivel de enlace de dispositivo destinatario del envío podrá cogerlo para él.

Todo esto nos lleva a la conclusión de que es necesario algún mecanismo que identifique a los dispositivos y los diferencien unos de otros.

Las direcciones a nivel de enlace, están formadas por una secuencia más o menos larga de dígitos binarios que identifican a los distintos dispositivos. Dependiendo del protocolo utilizado, estas direcciones pueden tener un mayor o menor número de dígitos. Así en el estándar de Ethernet, las direcciones MAC son números binarios de 48 dígitos que se suelen expresar como grupos de ocho bits representados en hexadecimal y separados por puntos.

Una dirección MAC podría ser la siguiente:

23.4F.05.34.FA.05

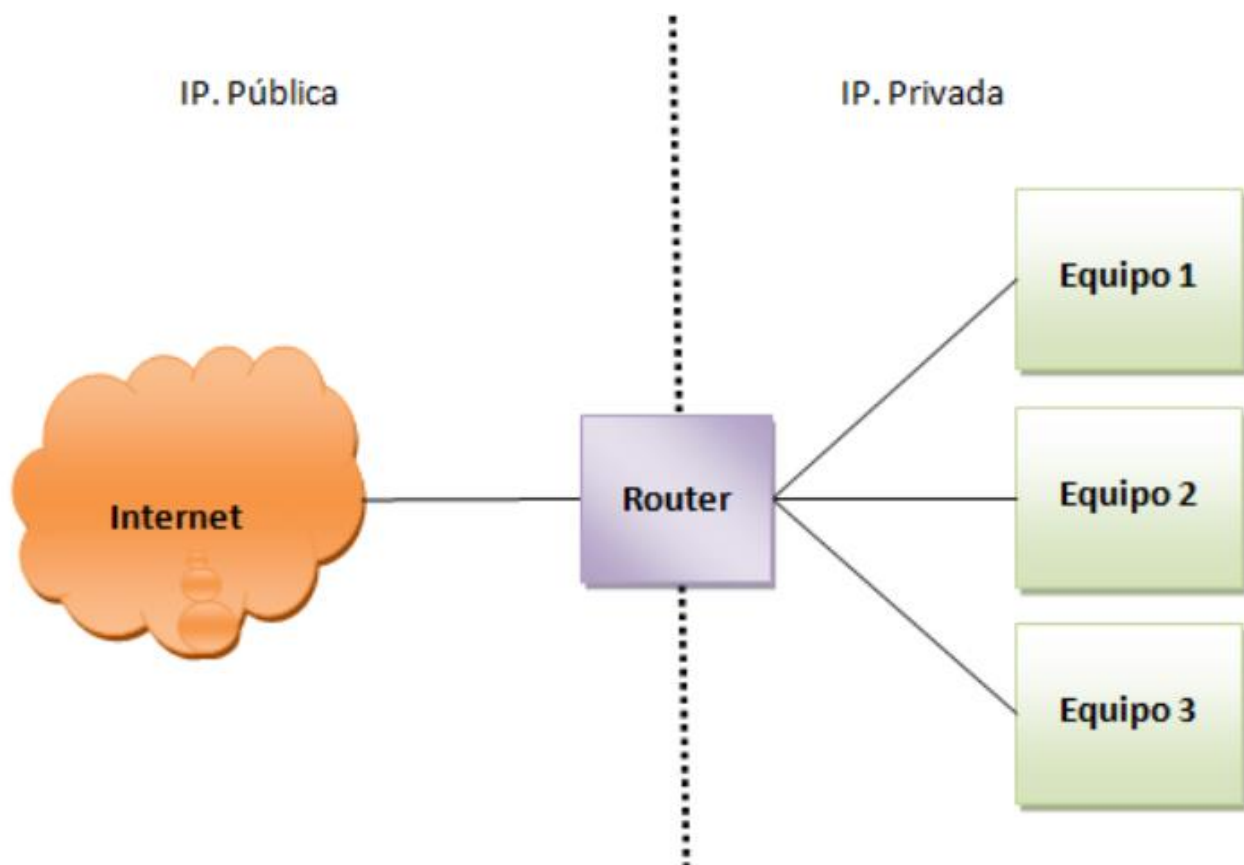
Las direcciones MAC están grabadas de fábrica en los adaptadores de red. Así no se pueden modificar a posteriori, garantizando así la correcta identificación de los equipos. Para impedir que en una red puedan existir dos estaciones con las mismas direcciones MAC, los fabricantes suelen coordinarse para asignar las MAC a sus equipos y a sus dispositivos.

En las direcciones Ethernet y Token Ring, los 24 primeros bits, identifican al fabricante, mientras que los dígitos restantes identifican de forma única cada adaptador, en una asignación realizada por el propio fabricante.

Cuando un conmutador recibe un mensaje con esta dirección de destino, este lo reenvía a través de todos sus puertos, para que llegue a todos los equipos que tienen conectados.

Direccionamiento a nivel de red: IPs, segmentación de redes, IP Pública, Privada...

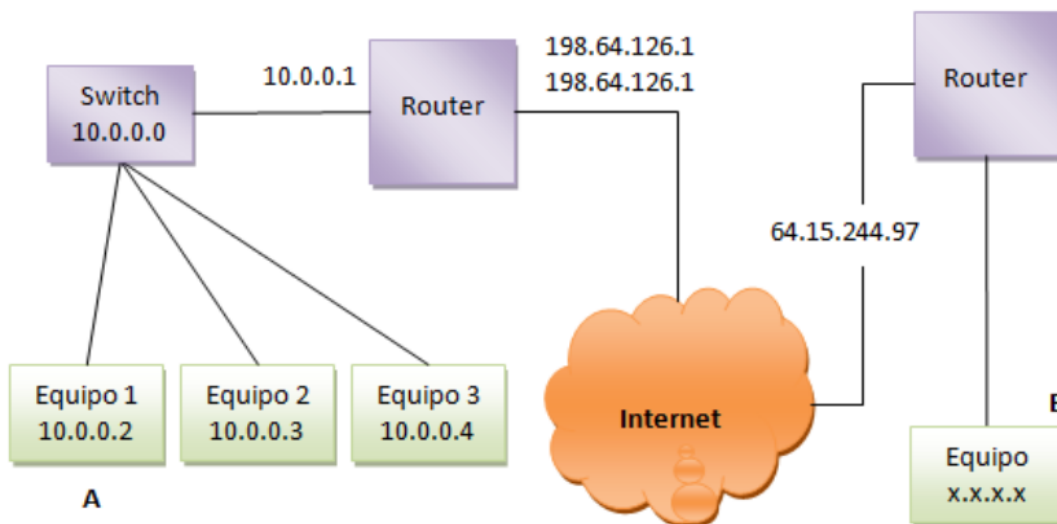
Ya visto en módulos anteriores



Traducción de las direcciones de red(NAT)

Ya visto en módulos anteriores

Supongamos que tenemos una red local conectada a Internet cuya estructura es la siguiente:



Imaginemos que la estación A cuya dirección IP es 10.0.0.2 va a enviar un paquete a la estación B (cuya dirección IP es 68.15.224.97), la cual contestará con otro paquete de respuesta. Los pasos que se siguen en la traducción NAT para este caso son:

1. La estación A envía el paquete a la red. La dirección origen es 10.0.0.2 y la dirección destino es 68.15.224.97
2. El enrutador detecta ese paquete y comprueba que la dirección de destino está en una red externa. Antes de enviar el paquete, sobrescribe la dirección de origen (con una dirección pública disponible) y guarda esta comunicación en tabla NAT.

El paquete llevará ahora su dirección de origen, 198.64.126.1 y la de destino, 68.15.224.97. La tabla NAT tendrá esta forma:

IP privada	Ip pública
10.0.0.2	68.15.224.97

3. La estación B recibe el paquete y contesta con otro cuyas direcciones de origen y destino son: 68.15.224.97 y 198.64.126.1.
4. El enrutador recibe el paquete, comprobando que existe en su tabla local una comunicación establecida entre 10.0.0.1 y 68.15.224.97. Por lo tanto, antes de meter este paquete en la red local, modifica la dirección de destino. Ahora el paquete llevará como dirección de origen 68.15.221.97 y 10.0.0.1 como dirección de destino.
5. Finalmente, la estación A comprueba que el paquete que circula por la red lleva su dirección de destino; por lo tanto, podrá tomarlo como propio.

Si se produjeran varias comunicaciones simultáneas, entonces a cada equipo de la red privada se le asignará una dirección IP pública del rango disponible.

Por ejemplo, si otro equipo de la red privada se conectara con un equipo de la red pública al mismo tiempo que lo hace A, entonces el enrutador puede tener las siguientes entradas en su tabla NAT:

IP privada	Ip pública
10.0.0.2	68.15.224.97
10.0.0.3	80.59.249.7

Y las direcciones IP públicas asignadas pueden ser 198.64.126.1 y 198.64.126.2.

Asignación dinámica de direcciones:

Otro mecanismo que consigue ahorrar en el número de direcciones IPv4 disponibles consiste en la asignación automática de direcciones solo a aquellos equipos que las necesiten. Esta asignación también facilita la configuración de los equipos, de forma que no es necesario realizar una configuración manual de los parámetros de la red.

Los mecanismos de asignación dinámica de direcciones aprovechan mejor el espacio IPv4 debido a que sólo se asignan a los equipos que necesitan conectarse en un momento dado. Una vez que se desconectan, la dirección se libera para que pueda ser utilizada por otro equipo.

Este mecanismo de asignación de direcciones funciona muy bien cuando se trabaja con conexiones no permanentes de acceso a Internet donde los equipos que no están conectados no tienen asignada ninguna dirección.

Existen dos protocolos que se pueden utilizar para que un equipo obtenga su configuración de red de forma automática:

BOOTP

El protocolo BOOTP (Boot Strap Protocol) se utiliza para que los equipos que no tienen instalado un sistema operativo puedan obtener una configuración de red de forma automática y así poder obtenerlo y cargarlo desde la red.

Para que un equipo pueda obtener su configuración automática y, posteriormente, cargar el sistema operativo desde otro equipo, necesita ejecutar un programa, que puede almacenarse en un dispositivo de almacenamiento externo como puede ser un Pendrive o una tarjeta de memoria o en la propia memoria del adaptador de red.

En la red debe existir un servidor que reciba las peticiones de los clientes para asignarles su configuración de red.

Al contrario que ocurre en otros protocolos similares, como DHCP, un servidor BOOTP envía la dirección IP al cliente después de haber consultado su información local.

Esta información se almacena en una tabla de correspondencias entre direcciones MAC y direcciones IP (que son las que se asignan). Esta tabla se debe haber definido previamente de forma manual y asigna siempre la misma dirección IP al mismo equipo.

BOOTP se ha utilizado mucho en entornos con sistemas operativos Unix y Linux, pero actualmente ha quedado en desuso, sobre todo por la utilización de protocolos más avanzados y flexibles, como DHCP.

y DHCP (este es el más utilizado).

El protocolo DHCP (Dynamic Host Configuration Protocol), permite asignar a los equipos de la red, una dirección IP automáticamente sólo cuando la necesiten.

Una vez dejan de necesitarla, se liberan para el uso por otros equipos.

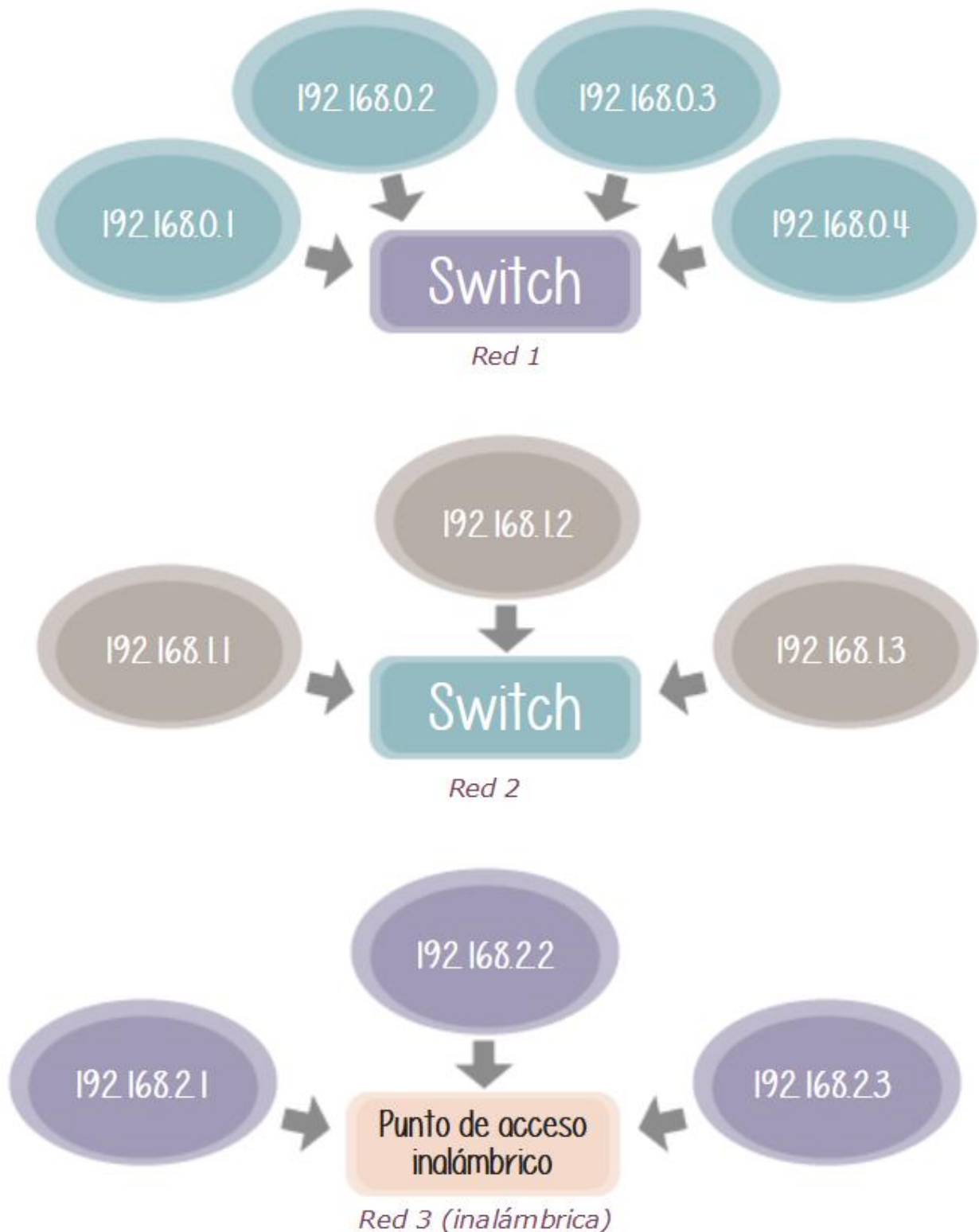
DHCP apareció en la década de los 90 con el objetivo de mejorar el protocolo BOOTP que se utilizaba en máquinas Unix para asignar direcciones a ciertos dispositivos que trabajaban en pequeñas redes locales.

El protocolo BOOTP permite la asignación automática de direcciones a equipos pero, esta asignación es estática, es decir, a cada equipo se le asigna siempre la misma dirección, que está almacenada como una tabla del servidor BOOTP. En el caso de DHCP, esta asignación es dinámica.

El Router para unir redes:

Supongamos que tenemos tres redes locales independientes entre sí. Las dos primeras son redes que están construidas con cable par trenzado, cableados y conectados todos los dispositivos a su respectivo Switch. La tercera red es inalámbrica con un punto de acceso inalámbrico.

Cada dispositivo tiene, como es lógico, una dirección IP. La máscara de subred para cada uno de los dispositivos es 255.255.255.0. Por lo tanto, cada red física tiene un rango de asignaciones de direcciones IP diferentes, es decir, son redes lógicas diferentes.



El rango de direccionamiento (red lógica) de cada red, utilizando la nomenclatura CIDR sería:

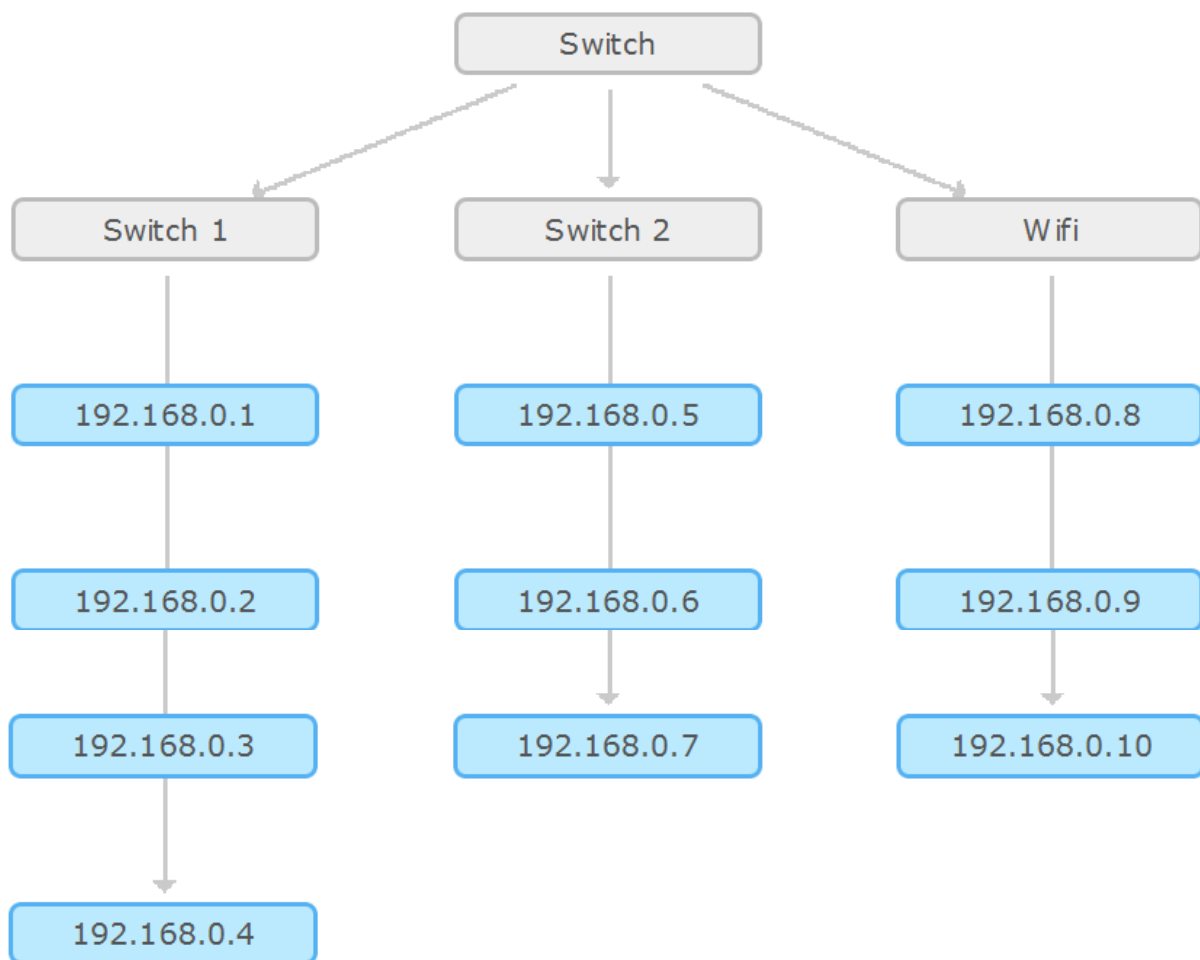
- ⇒ Red 1: 192.168.0.0 / 24
- ⇒ Red 2: 192.168.1.0 / 24
- ⇒ Red 3: 192.168.2.0 / 24

Si queremos interconectar todos los dispositivos de esas tres redes entre sí, tenemos dos soluciones posibles: mediante una red única o mediante el uso de subredes.

Red única:

La primera forma es crear una única red que conecte todos los dispositivos de las tres redes anteriores.

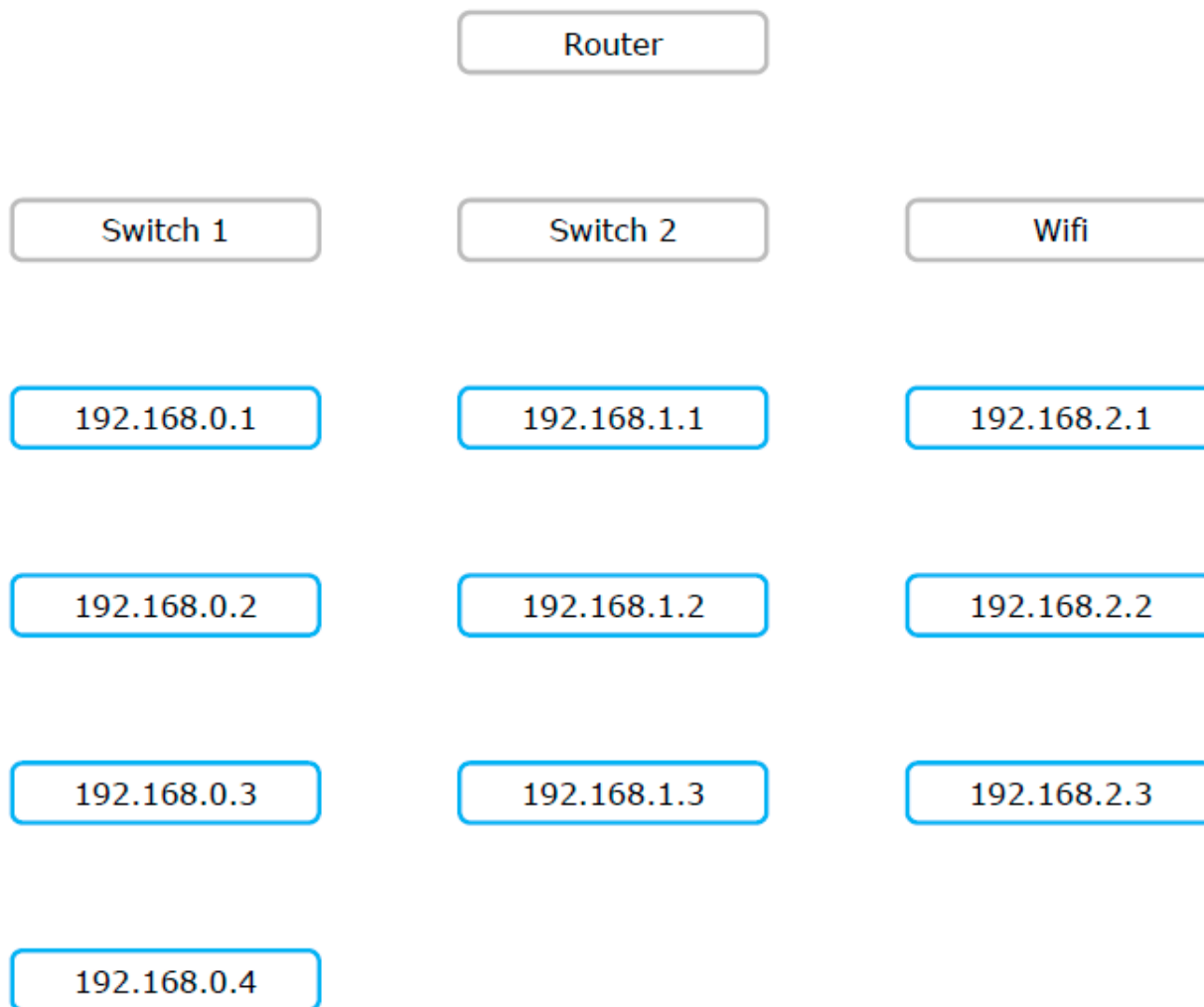
En este caso, vamos a usar un nuevo Switch, al que conectamos los dos Switch y el punto de acceso inalámbrico anteriores. Además, habría que cambiar las direcciones IP de dos de las redes para evitar conflictos y además conseguir que todos los dispositivos pertenezcan a la misma red lógica.



No hemos unido las tres redes anteriores, las hemos fusionado en una sola red más grande. Además, tenemos un solo rango de direcciones, el 192.168.0.0 / 24, que es el que se utilizaba originalmente en la primera red.

Subredes – Opción 1:

La segunda opción es unir las tres redes, pero manteniendo la independencia de cada una de ellas. En este caso, el dispositivo de interconexión necesario no es un Switch, es un Router.



De esta forma, los dispositivos de cada red pertenecen a su propia red lógica, independiente para cada red. El Router va a redirigir el flujo de datos entre las diferentes redes pero cada red mantiene su identidad y su propio rango de direcciones.

Los Routers no solo sólo se usan para conectar redes separadas físicamente bien sea en el mismo edificio, en edificios, ciudades o incluso países diferentes. Dentro de la red de una misma empresa se pueden tener diferentes redes lógicas, de forma que dos equipos conectados a la misma red física pueden pertenecer a redes lógicas diferentes.

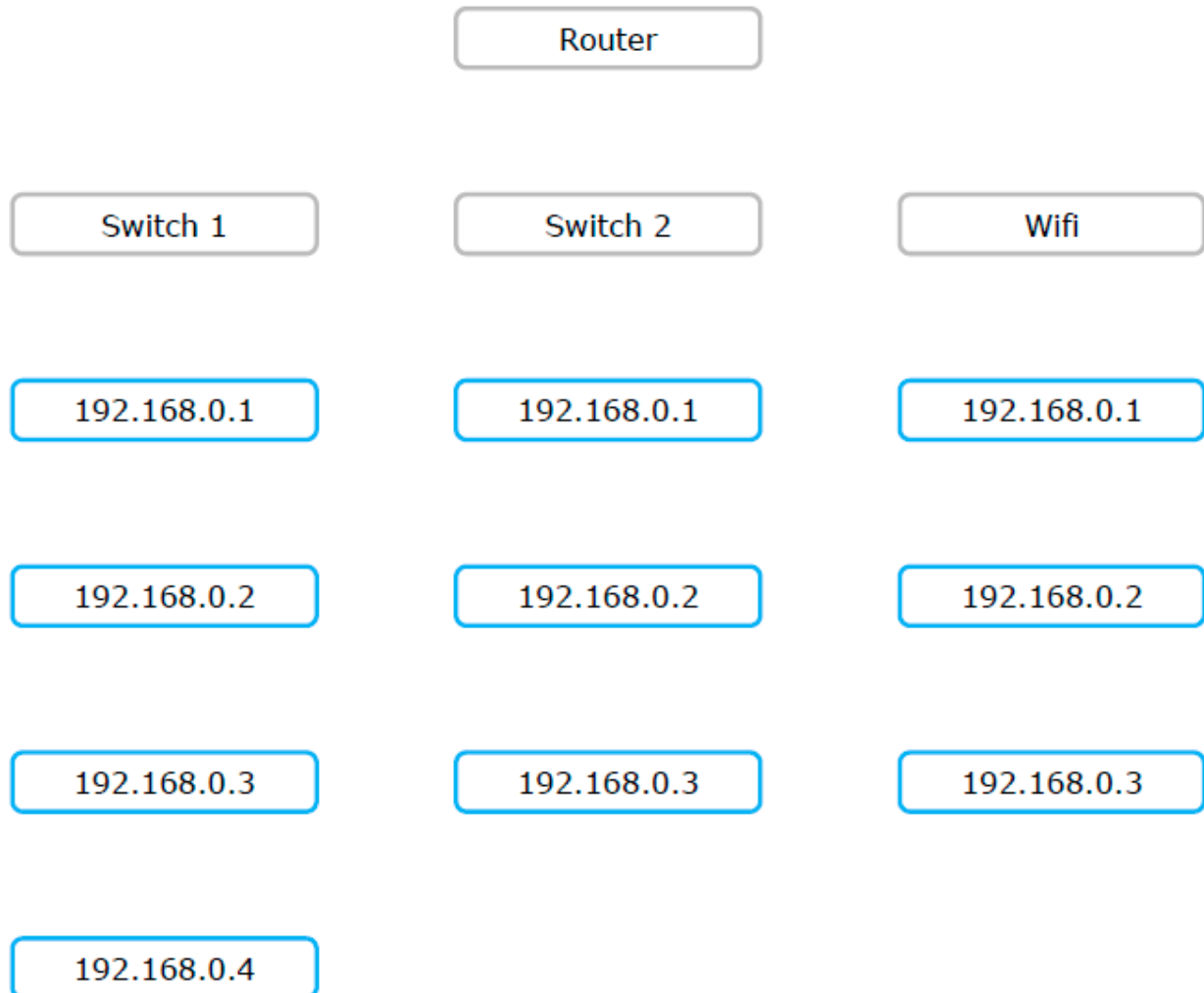
Como podemos ver, el Router une las tres redes anteriores, pero estas no pierden su identidad, es decir, su rango de direccionamiento.

Subredes – Opción 2:

La tercera solución es la más utilizada puesto que no necesita cambiar las IP privadas originales.

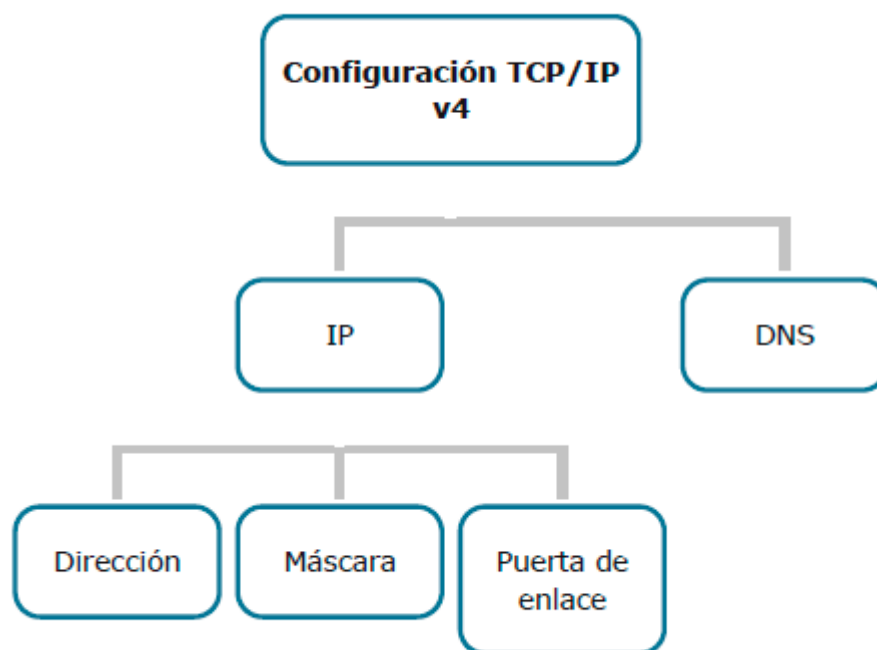
La idea es mantener las tres redes de partida de forma que cada equipo conserve su IP. Pero ¿puede haber dos equipos con la misma IP dentro de la misma Red?

La respuesta es clara: **NO**. Esto ocasionaría un problema grave de identificación en la red.



Para evitar el grave problema de duplicidad de direcciones IP en la red, la solución pasa por crear subredes con identificación exclusiva de cada una. Usando por ejemplo un switch de capa 3 que nos permita generar VLANs independientes.

La configuración IP: (Ya visto en módulos anteriores)



Coordinación entre el nivel de transporte y el nivel de aplicación:

El sistema de servidores DNS almacena y relaciona la información relativa a los dominios de los sitios web. Esta información se almacena en lo que llamamos zonas, que está formada por un conjunto de dominios y/o subdominios. La zona tiene supremacía sobre el dominio, de forma que es la zona la que guarda la información sobre cómo está organizado el dominio.

Normalmente, las zonas se almacenan en una base de datos en los servidores DNS, desde donde se gestionan.

Si se quiere definir una zona para una determinada red, guarde la información de configuración de ésta se deberá almacenar en al menos servidor DNS. Si no hay definidas zonas, los servidores DNS funcionarán como una caché de nombres. Es decir, el servidor no define ninguna zona y se limita a recibir peticiones y buscar en su caché a ver si puede resolverlas. En caso contrario, reenvía todas esas solicitudes a otros servidores DNS conocidos.

El uso de servidores para el caché de nombres acelera bastante el acceso a redes de área extensa, ya que la resolución de nombres es mucho más rápida.

Se entiende por caché a una tabla local que almacena correspondencias ya resueltas que se van actualizando conforme van quedando obsoletas.

La creación de las zonas de los dominios se basa en cuestiones relativas al tamaño de la organización, los subdominios definidos, la velocidad de los enlaces que comunican los servidores DNS, tasa de tráfico, etc.

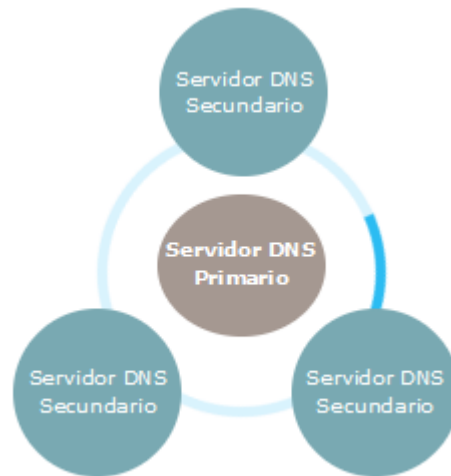
En determinadas situaciones como puede ser un dominio grande, se crea un subdominio en una zona distinta a la que se encuentra el dominio padre. Será la zona que contiene el dominio padre la que autorice y delegue el control a la nueva zona para que gestione el subdominio de forma autónoma.

La base de datos de una zona determinada se almacena en un servidor DNS primario, de tal forma que los dominios que ésta contiene se gestionan desde el servidor DNS primario.

No es conveniente tener un único servidor primario para la zona. Esto se debe a que en caso de cualquier tipo de fallo no se podrán resolver esas direcciones. Por lo tanto, se deben configurar al menos un servidor DNS secundario, que se encargan de mantener copias actualizadas de la información de zona.

Esta actualización de la información, se hace volcando total o parcialmente la información de zona. Este proceso se llama transferencia de zona.

Duplicidad de las zonas en servidores secundarios:



La información de zona se estructura en forma de Registros de Recursos (RR), cuyos tipos aparecen expuestos en la siguiente tabla:

Tipo	Descripción
A	Asocia un nombre de dominio de equipo con su dirección IPv4. Es el registro más utilizado.
AAAA	Asocia un nombre de dominio de equipo con su dirección IPv6 completa de 128 bits. Su nombre proviene de los registros A para IPv4 pero las direcciones IPv6 son cuatro veces más grandes.
A6	Otro registro que asocia un nombre de equipo con su dirección IPv6, pero que puede especificar direcciones completas como AAAA o direcciones fragmentadas.
ATMA	Asocia direcciones de dominio con direcciones ATM.
CNAME	Se utiliza para asignar otro nombre (alias) de un nombre de dominio equipo.
HINFO	Guarda información adicional de un equipo, como el tipo de CPU que lleva o el sistema operativo instalado. No se recomienda utilizar este tipo de registro por cuestiones de seguridad en la red.
MX	Registra un servidor de correo.
NS	Es una referencia a los nombres de dominios DNS de servidores que tienen autoridad para una zona, es decir, guardan la información de los dominios de esa zona.
PTR	Asigna una dirección IP a un nombre de dominio para llevar a cabo correspondencias inversas de verás más adelante en este apartado.
RDSI	Asigna un nombre de dominio a un número de teléfono de RDSI.
SOA	Es el primer registro que se crea cuando se agrega una zona nueva. Se usa para informar que el equipo especificado es el que guarda la información del dominio.
SRV	Se usa para localizar servidores que proporcionen tipos de servicios específicos (POP, LADR, etc.).
TXT	Se trata de un registro definido para guardar información de tipo texto que cualquier usuario puede leer.
X25	Asocia un nombre de dominio a una dirección X.25.

El protocolo DNS se define en los estándares RFC 1034 y 1035. Estos definen que se utiliza una base de datos distribuida por la red en ordenadores llamados servidores DNS, que almacenan tablas de correspondencias entre direcciones de nombres de dominio y direcciones IP.

Este servidor consulta primero en sus registros de recursos de zona la dirección solicitada (en caso de ser un servidor primario o secundario) y devuelve la dirección IP si la encuentra. Si no la encuentra ahí, consultará entonces la caché donde están almacenadas temporalmente consultas anteriores. En caso de que tampoco la encuentre ahí, puede consultar otros servidores DNS, operación que se denomina consulta recursiva. Finalmente, devuelve la dirección IP solicitada a la estación.

En función del sistema operativo que tengamos instalado, la tabla local puede estar almacenada en distintos directorios del disco duro:

En Microsoft Windows ésta se encuentra en el archivo HOSTS.SAM de la carpeta WINDOWS.

En Linux, esta tabla local se guarda en el archivo `/etc/hosts`.

En todos los casos se pueden incluir entradas manualmente con direcciones de equipos bien conocidos.

En la mayoría de estas tablas suele aparecer la entrada referente al equipo local cuando tenemos instalado un servidor web local, de la siguiente forma:

```
127.0.0.1 <-----> localhost
```

En las primeras versiones del protocolo DNS, los servidores DNS resolvían correspondencias directas, es decir, tras recibir una dirección de nombre, estos devuelven la dirección IP asociada. En las siguientes versiones del protocolo, también resuelven correspondencias inversas, devolviendo una dirección de nombre cuando reciben una solicitud con una dirección IP.

Estas resoluciones inversas se utilizan por ejemplo cuando un servidor DNS ha solicitado una resolución directa a otro y por seguridad quiere comprobar que esa correspondencia es cierta. Este proceso de comprobación es bastante frecuente.

El problema que plantean las resoluciones inversas es que, conforme se había organizado el espacio de nombres de los dominios y éstos en las zonas, estas consultas deben hacer una búsqueda en todos los dominios de la red. Es por ello que las consultas inversas tardan demasiado tiempo en resolverse.

Para solucionar este problema, en el estándar de DNS se ha reservado el dominio especial `in-addr.arpa` (funciona con IPv4) o `ip6.arpa` (funciona sobre IPv6), que debe definirse en una zona tantas veces como dominios existan en ella.

A diferencia de los dominios del DNS, el dominio `in-addr.arpa` (o `ip6.arpa`) está definido por todos los subdominios en orden inverso, junto con su dirección IP asociada (también aparecen los octetos en orden inverso, ya que el nombre de un dominio se lee al revés a una dirección IP).

Los registros que guardan estas correspondencias inversas son punteros de registros de recursos (PTR). Es posible que la herramienta de administración del servidor DNS cree automáticamente las resoluciones inversas, aunque a veces hay que establecerlas manualmente.

Con objeto de facilitar la administración de zonas en los servidores DNS, algunos sistemas operativos para equipos cliente, permiten una actualización dinámica. Consiste en que el equipo cliente notifica de forma automática al servidor DNS cualquier cambio en su dirección IP o de su nombre de dominio.

Así, el administrador de la zona no tendrá que configurar ese nuevo registro en el servidor, ya que éste lo hará automáticamente cuando reciba la notificación desde el equipo cliente.

En el estándar RFC 1032 se especifica cómo registrar un nombre de dominio en Internet e incluye los formularios necesarios.

Así mismo, la configuración de un servidor DNS se encuentra definida en el RFC 1033. Hay que tener en cuenta que, aunque las implementaciones de los servicios DNS se basan en las especificaciones estándares incluidas en los documentos RFC, es posible que surjan ciertas incompatibilidades entre los servidores DNS de Windows y los servidores DNS Linux, sobre todo en versiones antiguas.

Servidores DNS:

¿Qué es el Domain Name System?

Es una base de datos distribuida, con información que se usa para traducir los nombres de dominio, fáciles de recordar y usar por las personas, en números de protocolo de Internet (IP).

El servicio de DNS permite que una web o una dirección de correo electrónico sean localizadas desde cualquier lugar del mundo mediante un nombre de dominio.

El servicio de DNS (Domain Name System) se encarga de traducir los nombres de los dominios de las diferentes peticiones de acceso a sus correspondientes direcciones IP, el cual se utiliza para identificar a cualquier equipo conectado a la red. Es decir, para poder localizar físicamente y consultar los recursos.

Este procedimiento tarda unos pocos milisegundos. A lo largo de todo un día de funcionamiento y con las miles de peticiones recibidas, esto puede significar cantidades importantes de tiempo. En estos casos o si hay problemas concretos, puede resultar beneficioso utilizar el servicio de servidores DNS diferentes a los provistos por el ISP (Proveedor de servicios de Internet).

El Domain Name System (DNS), o Sistema de Nombres de Dominio, comprende personas, instituciones reguladoras, archivos, máquinas y software trabajando conjuntamente. Una pieza fundamental en este sistema es el servidor de DNS.

Aunque muchas veces no reparamos en ello, los servidores DNS son una pieza clave para que podamos navegar por Internet con cierta agilidad. La traducción de un nombre de dominio a una dirección IP es fundamental para que podamos acceder a una web o a un servidor de correo de una manera sencilla y sin tener que andar memorizando direcciones

IP que, además, podrían ser susceptibles de cambiarse.

Estos servicios están normalmente vinculados a los ISP, es decir, a nuestros proveedores de acceso a Internet. Estos no suelen caracterizarse por ofrecer el mejor de los servicios y suelen presentar incidencias que ralentizan la navegación.

Estas DNS en muchos casos, pueden estar sobrecargados o caídas, lo que provoca problemas de conexión.

Por ello, a veces un simple cambio de los servidores DNS que nos asigna por defecto el Router puede mejorar mucho la conexión y resolver los problemas.



A continuación se muestran las DNS de los principales servidores de Internet (ISP).

Servicio de DNS de los ISP:

ISP	IP	Nombre
Acens	217.116.0.177	ns2.acens.net
Acens	217.116.0.176	ns1.acens.net
Adam	212.36.64.16	dns1.adam.es
Adam	212.36.64.17	dns2.adam.es
Arrakis	195.5.64.6	ns2.landsraad.net
Arrakis	195.5.64.2	ns1.landsraad.net
Arsys	217.76.129.4	prometeo.servidoresdns.net
Arsys	217.76.128.4	atlante.servidoresdns.net
Arsys	217.76.129.131	dns8.servidoresdns.net
Arsys	217.76.128.131	dns7.servidoresdns.net
AunaCable	62.81.31.250	dns.iddeo.es
AunaCable	62.81.0.1	mayor.red.retevision.es
AunaCable	62.81.16.131	131-16-81-62.libre.auna.net
AunaCable	62.81.16.132	132-16-81-62.libre.auna.net
AunaCable	62.81.0.35	35-0-81-62.libre.auna.net
AunaCable	62.81.0.36	36-0-81-62.libre.auna.net
Comunitel	212.145.4.98	ns2.comunitel.net
Comunitel	212.145.4.97	ns1.comunitel.net
ConexiónFutura	213.139.0.52	ns1.es.easynet.net
ConexiónFutura	193.43.232.4	dns1.conexionfutura.com
ConexiónFutura	213.139.0.51	ns0.es.easynet.net
Euskaltel	212.55.8.133	dns2.euskaltel.es
Euskaltel	212.55.8.132	dns.euskaltel.es
Iberbanda	217.11.96.234	dns.mad.iberbanda.es

Iberbanda	217.11.108.234	dns.bcn.iberbanda.es
Iberdrola	213.172.33.35	ns2.neo.es
Iberdrola	213.172.33.34	ns1.neo.es
Iberdrola	193.127.102.104	ns2.presenzia.net
Iberdrola	193.127.102.45	ns1.presenzia.net
Jazztel	62.14.4.65	dnscache2.jazzvisp.com
Jazztel	62.14.4.64	dnscache1.jazzvisp.com
Jazztel	62.14.63.145	Desconocido
Jazztel	62.14.2.1	dns1.inversas.jazztel.es
ONO	62.42.230.135	dns01.ono.com
ONO	62.42.230.136	newsf01.ono.com
ONO	62.42.230.24	resolv.ono.com
ONO	62.42.63.51	dns03.ono.com
ONO	62.42.63.52	Desconocido
OpenForYou	213.195.64.129	dns1.ibercom.com
OpenForYou	213.195.79.129	dns2.ibercom.com
Retena	212.21.224.5	sn1.retena.es
Retena	212.21.224.6	sn2.retena.es
Reterioja	212.21.224.4	sn2.reterioja.es
Reterioja	212.21.224.3	sn1.reterioja.es
SuperBanda	212.4.96.22	ns1.grupalia.com
SuperBanda	212.4.96.21	ns2.grupalia.com
Tele2	130.244.127.161	dns1.swip.net
Tele2	130.244.127.169	dns2.swip.net
TeleCable	212.89.0.31	dns.telecable.es
Telefónica	80.58.0.33	33.Red-80-58-0.pooles.rima-tde.net
Telefónica	194.179.1.100	minerva.ttd.net
Telefónica	194.179.1.101	artemis.ttd.net
Telefónica	80.58.0.97	Red-80-58-0.pooles.rima-tde.net
Telefónica	213.0.184.68	minerva.ttd.net
Telefónica	194.224.52.37	ns2.telefonica-data.com
Telefónica	80.58.32.97	97.Red-80-58-32.pooles.rima-tde.net
Telefónica	80.58.32.33	33 Red-80-58-32.pooles.rima-tde.net
Telefónica	213.0.184.69	artemis.ttd.net
Telefónica	194.224.52.4	esifw1.tsai.es
Telefónica	194.224.52.36	ns1.telefonica-data.com
Telefónica	194.224.52.6	esifw2.tsai.es
Terra	213.4.141.1	dns2.terra.es
Terra	195.235.96.90	tpdns2.terra.es
Terra	213.4.132.1	dns1.terra.es
Terra	195.235.113.3	dns.terra.es
Tiscali	212.166.64.2	dns2.tiscali.es
Tiscali	212.166.64.1	dns1.tiscali.es
Wanadoo	62.37.236.200	dns2.wanadoo.es
Wanadoo	62.37.225.58	pdns03pub.uni2.es
Wanadoo	62.37.237.140	dns1.wanadoo.es
Wanadoo	62.37.225.56	dns.comtenidos.com
Wanadoo	62.37.225.57	dns2.comtenidos.com
Wanadoo	62.37.228.22	m2cachedns2.uni2.es
Wanadoo	62.37.228.22	m2cachedns2.uni2.es
Wanadoo	62.37.228.20	m2cachedns.uni2.es

Ya.com	62.151.20.6	ns2.bs-ya.com
Ya.com	62.151.20.7	ns.bs-ya.com
Ya.com	62.151.2.65	dns.yaonline.es
Ya.com	62.151.2.8	dns.ya.com
Ya.com	62.151.8.100	dns2.ya.com

DNS Públicas:

Como ya hemos dicho anteriormente, cuando contratamos una conexión ADSL el proveedor de servicios de internet (ISP), nos facilita unos DNS propios de la compañía. Estos DNS suelen funcionar bastante bien, pero puede ocurrir que por alguna circunstancia estos servidores DNS no funcione todo lo bien que sería deseable, provocando lentitud al cargar páginas o directamente que no funciones internet.

En otros casos y por cuestiones más políticas y administrativas que informáticas los servidores DNS de los ISP bloquean el acceso a determinadas páginas webs. Por ejemplo nos referimos a ciertos países que intentan bloquear el acceso de sus ciudadanos a la red.

Para evitar esas dos cuestiones, podemos utilizar una serie de DNS públicos que hay en la red. A continuación listamos los más conocidos y utilizados.

Google Public DNS:

Son los servidores DNS libres y gratuitos más conocidos. Tiene una disponibilidad de servicio de entorno al 100% del tiempo. En poco tiempo el servidor DNS de Google tiene en torno a 70.000 millones de peticiones diarias, un dato que sigue en ascenso y hace que este servidor sea uno de los más utilizados.

Sus principales características son el reducido tiempo de respuesta y una disponibilidad del 100%.

Los DNS de Google son los siguientes:



DNS primario	DNS secundario
8.8.8.8	8.8.4.4

Open DNS:

Otro servidor DNS con una disponibilidad también muy alta. Además, Open DNS protección ante sitios web fraudulentos, bloqueando el acceso a dichas webs. También tiene opciones interesantes referentes al control parental para menores.



Los DNS de Open DNS son los siguientes:

OpenDNS

DNS primario DNS secundario

208.67.222.222 208.67.220.220

OpenNIC:

OpenNIC es un proyecto de DNS y de registro de nombres de dominio libre del ICANN (Internet Corporation for

Assigned Names and Numbers).

Este proyecto se apoya en la propia comunidad de usuarios de la red que colabora en la implantación de una red de

servidores DNS libre y descentralizada.

Una cuestión a tener en cuenta es que estos servidores que no guardan registros de las consultas que realizan los

usuarios (o se borran a las 24 horas).

Los DNS de OpenNIC son los siguientes:



DNS primario DNS secundario

64.0.55.201 184.154.13.11

202.83.95.227 216.87.84.211

DNS de OpenNick

Level 3 Communications:

Level 3 Communications es una empresa de telecomunicaciones que posee una importante infraestructura de red de transporte de datos. También ofrece un conjunto de servidores DNS que están entre los mejor valorados entre los usuarios por su buen tiempo de respuesta.



Los DNS de Level 3 Communications son los siguientes:

Label3

DNS primario DNS secundario

4.2.2.1 4.2.2.2

4.2.2.3 4.2.2.4

4.2.2.5 4.2.2.6

DNS de Label3

1.1.1.3. Fallos de configuración interna / interfaces de interconexión (continuación).

ScrubIT:

El servicio se caracteriza por sus potentes filtros de sitios con contenido pornográfico, phishing u otros contenidos maliciosos.

Es muy apreciado en entornos empresariales y sobre todo familiar por el filtrado que hacen.

Los DNS de ScrubIT son los siguientes:



DNS primario DNS secundario

67.138.54.100 207.225.209.66

DNS de ScrubIT

Norton ConnectSafe:

La muy conocida empresa de seguridad Symantec ofrece en Norton ConnectSafe un conjunto muy completo de servidores DNS. Posee servidores con tres niveles de seguridad:

- ⇒ Nivel 1. Con el par de servidores DNS 198.153.192.40 y 198.153.194.40. Proporciona un servicio gratuito de resolución de nombres que bloquea sitios malware, phishing y sitios web fraudulentos.
- ⇒ Nivel 2. Con el par de servidores DNS 198.153.192.50 y 198.153.194.50. Además de las funcionalidades del nivel anterior, añade bloqueo de sitios web pornográficos. Muy usados en el entorno del control parental.
- ⇒ Nivel 3. Con el par de servidores DNS 198.153.192.60 y 198.153.194.60. Además de lo aportado por los niveles anteriores, añade otro nivel más de filtrado de contenidos ya que se bloquea las peticiones a páginas web catalogadas como no aptas para ver en familia.



Los DNS de Norton ConnectSafe son los siguientes:

Norton DNS

DNS primario DNS secundario

198.153.192.40 198.153.194.40

198.153.192.50 198.153.194.50

198.153.192.60 198.153.194.60

DNS de Norton DNS

1.1.1.3. Fallos de configuración interna / interfaces de interconexión (continuación).

Comodo Secure DNS:

Comodo Secure DNS ofrece alta disponibilidad (los servidores están localizados 15 lugares diferentes).

También

protección frente a sitios web con malware, sitios de phishing. Otro valor añadido que posee es el bloqueo de páginas

de publicidad constante.

Los DNS de Comodo Secure DNS son los siguientes:



DNS primario DNS secundario

8.26.56.26 156.154.70.22

DNS de Comodo Secure DNS

Dyn Internet Guide:

Dyn es un ISP que ofrece también servidores DNS gratuitos y, al igual que ocurre con Level 3, los resultados de su

funcionamiento son muy buenos y suele estar entre los mejores tiempos de respuesta.

Ofrece filtrados y protección frente a páginas potencialmente peligrosas.

Los DNS de Dyn Internet Guide son los siguientes:



DNS primario DNS secundario

216.146.35.35 216.146.36.36

DNS de Dyn Internet Guide

DNS Advantage :

DNS Advantage proporciona una velocidad de respuesta de las más altas del espectro de servidores DNS.

Los DNS de DNS Advantage son los siguientes:



DNS primario DNS secundario

156.154.70.1 156.154.71.1

Errores en la conexión WiFi:

Un problema que puede surgir en la configuración de la red WiFi, es la seguridad de las claves de acceso, el evitar que alguien ajeno a la organización puede acceder a la red. Esto hay que tenerlo muy en cuenta. Las redes Wi-Fi son muy útiles, sencillas de utilizar y cada vez más habituales. Pero deben protegerse mediante claves de acceso. Estas claves están a su vez protegidas mediante WEP o WPA, que cifran la información de la red inalámbrica.

Pero eso dos sistemas de cifrado no son iguales. Son absolutamente diferentes.

⇒ WEP (Wired Equivalent Privacy):

El cifrado WEP fue el primer estándar de seguridad para redes Wi-Fi. Hoy está superado.

No se debe usar WEP para proteger la red inalámbrica si hay otras alternativas. Su protección es demasiado débil. Se puede crackear un cifrado WEP en pocos minutos usando las herramientas adecuadas. Estas están cada vez mas disponibles en internet a la mano de casi cualquiera.

⇒ WPA (Wi-Fi Protected Access)

Este cifrado aparece para corregir las deficiencias del cifrado WEP.

Introdujo mejoras de seguridad como el TKIP (Temporal Key Integrity Protocol), que varía por sí solo la contraseña Wi-Fi cada cierto tiempo.

Tiene dos variantes:

- WPA-Personal. Usa el sistema PSK. En él, todos los usuarios de la red inalámbrica tienen una misma contraseña Wi-Fi, que el propio usuario define.
- WPA-Enterprise. Ofrece seguridad adicional al obligar al usuario a identificarse con un nombre y contraseña.

⇒ WPA2

Es el estándar más moderno para proteger redes inalámbricas y el que recomienda la Wi-Fi Alliance.

Hay también dos versiones:

(WPA2-Personal y WPA2-Enterprise).

WPA2 es compatible con WPA, lo que significa que en la red Wi-Fi se pueden usar PCs o dispositivos (Router, adaptadores de red...) que admitan uno u otro sistema.

WPA2 no es compatible, sin embargo, con sistemas WEP. En una misma red Wi-Fi no podrán coexistir dispositivos que sólo admitan WEP con otros válidos para WPA2.

Veamos algunas consideraciones importantes a tener en cuenta a la hora de proteger la red Wi-Fi:

- ⇒ Debemos asegurarnos de que el Router, los dispositivos inalámbricos y los adaptadores de red sean compatibles con WPA2.
- ⇒ Al configurar la red Wi-Fi debemos de usar WPA2-Personal con cifrado AES.
- ⇒ Hay que definir una contraseña Wi-Fi fuerte. Esto significa que:
 - Tenga al menos 15 caracteres.
 - Combine letras mayúsculas y minúsculas, números y caracteres especiales (\$, #, @, etc.).
 - No tenga NINGUNA información personal, como nombres, fechas de cumpleaños o aniversarios, etc.
 - No contenga palabras completas en ningún idioma.
 - No sea obvia. Claves como "qwerty", "1234" o "contraseña" no son nada seguras.
- ⇒ Es recomendable cambiar la contraseña Wi-Fi cada cierto tiempo siguiendo las mismas recomendaciones.

1.1.2. Incidencias Externas (atribuibles al proveedor de servicios)

Los proveedores de servicios de Internet (ISP) comenzaron a surgir a finales de la década de los 80 y principios de los 90 del pasado siglo. Son las empresas y/o organizaciones que proporcionan a sus abonados acceso a Internet así como servicios relacionados, principalmente, hospedaje de páginas web, servicios de correo electrónico a través de la línea telefónica o dedicada.

Estos proveedores conectan sus clientes a los clientes de otros proveedores de servicio por medio de las distintas redes que hay repartidas a lo largo y ancho de la nube.

En la mayoría de los casos, los proveedores de servicios de Internet (también llamados Proveedores de acceso a Internet) son las propias empresas que proporcionan servicios de telecomunicaciones, como pueden ser las propias empresas de telefonía. La mayoría de las empresas telefónicas ahora funcionan como proveedores de acceso a Internet también.

Tipos de ISP:

Hay diferentes tipos de Proveedores de servicio de Internet disponibles en la actualidad. Veamos algunos de ellos.

- ⇒ ISP de Acceso: Emplean una amplia gama de tecnologías para proporcionar la conexión de sus clientes a sus redes. Estas tecnologías van desde la banda ancha a la más antigua conexión por línea telefónica conmutada. Muchos de estos proveedores de acceso también proporcionan email y servicios de servidores.
- ⇒ ISP de Buzón de correo: Ofrecen servicios de servidor de buzón de email y servidores de email para enviar, recibir y almacenar email. Muchos ISP de buzón de correo son también proveedores de acceso.
- ⇒ ISP de Servidores: Ofrecen email, Protocolo de Transferencia de Archivos (FTP), servicios de servidores web, máquinas virtuales, servidores de cloud y físicos.

- ⇒ ISP de Tránsito: Proporcionan grandes cantidades de ancho de banda necesarias para conectar los ISP de servidores y los ISP de acceso juntos.
- ⇒ ISP Virtuales (VISP): Compran servicios de otros ISP para permitir a los clientes el acceso a Internet.
- ⇒ ISP gratuitos: Proporcionan el servicio de forma gratuita y a menudo muestran anuncios mientras los usuarios están conectados. Sobre todo los hay para el alojamiento de páginas web.

1.1.2.1. Caídas de servicios por parte del proveedor de servicios

Una de las peores situaciones con las que se puede encantar cualquier persona que esté navegando en la web es encontrarse

con un corte o caída en el servicio de Internet, sobre todo, cuando su trabajo depende de estar conectado la mayor parte de su tiempo en el ciberespacio. Seguro nos resulta familiar la situación de ir a un banco y encontrarnos

con que no funciona internet.

Esta problemática puede deberse a distintos motivos y tienen una gran relevancia en la economía global. Pensemos

en la cantidad de procesos administrativos y comerciales que muchas empresas están realizando en el mundo digital.

Cada vez que hay un problema importante de conectividad a internet, las pérdidas son muy importantes.

Pero ¿por

qué ocurre esto? La respuesta no es sencilla, pero lo que si está claro es que puede deberse a multitud de factores.

Veamos algunos de ellos:

Caída de Servicios

- ⇒ Un fallo generalizado en el Servidor web.
 - Daño eléctrico generalizado.
 - Desconexión temporal o permanente de los grandes servidores.
 - Problemas estructurales de algunos nodos.
 - Problemas directos de conexión.
 - Ataques directos de Hacker. Esta es la causa de caída del servicio de Internet más común.
- ⇒ Robo de cables de transmisión.

Los robos de cables de alimentación, de transmisión o los propios conectores son el ataque más común a las centrales de abastecimiento. Las cajas de conexiones, el cableado estructurado de par o de fibra óptica, los cables telefónicos, son objetos de robos constantes debido al alto precio del cobre. Esto genera pérdidas millonarias en cableado, y por supuesto limitan las conexiones.
- ⇒ Deshabilitación del sistema básico de conexión.
- ⇒ Una conectividad de poca potencia o de estructuración baja, que influye muy negativamente en la
- ⇒ navegación.
- ⇒ Microcortes.

Los microcortes:

A menudo la línea experimenta en muchas ocasiones microcortes, o incluso bajadas de velocidad esporádicas, que mejoran al reiniciar el Router al día siguiente.

Los microcortes y caídas esporádicas de velocidad no es una cuestión exclusiva de configuraciones o particularidades de cada uno de los operadores. Afecta de una forma o de otra a todos los operadores del servicio de acceso a internet.

Línea telefónica convencional, la fibra óptica, los datos móviles, 3G, 4G, comunicaciones vía satélite. Es indiferente, en todos los casos nos encontramos con la misma problemática: Hay caídas del servicio de internet.

Para entender qué son esos microcortes, cómo se producen y qué es lo que normalmente los ocasiona debemos entender claramente qué es lo que ocurre cuando se conecta un Router.

Durante el proceso de arranque del Router, se carga y ejecuta un programa que está almacenado en el interior del mismo. Este funciona a modo de sistema operativo y permite la gestión de comunicación con el servidor.

Una vez que se ha cargado este software en la memoria del dispositivo, se empieza a sondear la línea telefónica de forma que el Router se comunica con el módem ADSL que hay en el nodo al que se haya conectado.

En esta comunicación se produce una negociación entre los protocolos de transporte y enlace de ambos extremos (equipo del nodo y equipo del cliente). Se tienen que poner de acuerdo ambos en cómo va a ser esa transmisión, sus parámetros.

Ahora hay que establecer el canal de comunicación. Para ello se envían una serie de señales a unas determinadas frecuencias. Se comprueba cómo se comporta la transmisión a diferentes potencias y a diferentes frecuencias.

Para ello, se analiza con qué potencia llegan los tonos que envía el otro extremo y así calculan el ruido que existe en el canal en cada frecuencia.

Hay un parámetro denominado margen de SNR (Relación Señal – Ruido) que indica cuánto mayor tiene que ser la señal por encima del ruido para que pueda activar el canal ADSL con una determinada calidad.

Una vez que el Router ha sincronizado con la central, ya hay una velocidad inicial de funcionamiento, sin embargo, si las condiciones iniciales que se negociaron cambian (ruido de fondo, interferencias, etc.) notaremos extraños comportamientos en la conexión.

Estos son los micro cortes. La comunicación no puede mantenerse en los mismos niveles que se negoció al comenzar la sesión. Por eso muchas veces reiniciar el router es una solución muy optima a estos problemas, ya que se vuelven a negociar los parámetros de comunicación, contemplando la nueva situación de ruido y, por consiguiente, cambiando de frecuencia y/o canal.

1.2. Gestión de incidencias en equipos de acceso a redes públicas

Las empresas que usan continuamente los de servicios de redes de comunicaciones no pueden permitirse las situaciones de averías o de mal funcionamiento, ya que esto puede suponer pérdidas económicas o desconfianza por parte de los clientes.

Una característica que define a las redes de comunicaciones es su gran complejidad. Esta complejidad, se debe a la gran cantidad de servicios que soportan y al elevado número de usuarios a los que dan servicio.

Por ello, las tareas de administración y localización de incidencias pueden resultar bastante complejas y deberían de recaer en equipos de administradores bien cualificados.

Las incidencias en una red de comunicación pueden aparecer de muchas formas. Pueden ser problemas que afecten a una gran cantidad de usuarios, la caída de un servidor, o problemas más puntuales, el caso de un usuario no puede acceder a la red.

Una vez que se han identificado las causas que producen un problema, hay que proceder a solucionarlo.

Para ello puede que sea necesario involucrar varios departamentos de la empresa e incluso las acciones a emprender pueden suponer la parada momentánea de la red de comunicación.

Debido a la complejidad en la resolución de algunos los problemas que se pueden presentar, es más eficiente utilizar un protocolo bien definido en vez de ponerse directamente a teclear comandos o conectar el comprobador de red en los enlaces, sin tener muy claro que es lo que pasa.

Un método estructurado permite aprovechar mejor el tiempo y evita que los problemas puedan complicarse más de lo que ya lo están. Hemos de evitar la máxima que decía: “problema llama a problema”. Hay que tener en cuenta que cada sistema de comunicaciones tiene sus peculiaridades, y por lo tanto estos métodos estructurados han de ser específicos para cada red. Proponemos aquí un método que puede servir de base para la realización de uno específico para cada caso.

Los pasos básicos que se deberán seguir son los siguientes:

1. Establecer de una forma lo más clara posible cuál es el problema y los síntomas del mismo.
2. A partir de estos datos, intentar averiguar todas las posibles causas que pueden haberlo producido.
3. Recopilar toda la información posible sobre la problemática que se ha produciendo.

Deberíamos de hablar con todas las personas que trabajan en el sistema, tanto usuarios como administradores.

Deberíamos utilizar herramientas que permitan obtener información de cómo está funcionando el sistema, como analizadores de red, comprobadores, utilidades de verificación, snnifers, sistemas de monitorización, etc.

4. De toda la información recopilada en el paso anterior debemos deducir una lista con los problemas reales que pueden producirse, eliminando todo aquello que no guarda relación con la situación actual. Así, nos podemos centrar en los problemas que realmente tienen relevancia.
5. Ahora hay que establecer un plan de actuación para afrontar los problemas reales que se han identificado. Hay que comenzar tratando el problema que tenga más posibilidades de ser la causa del mal funcionamiento.
6. Conforme vayamos actuando, hay que ir haciendo comprobaciones para ver si se van solucionando los problemas.
7. En caso de que los problemas persistan, habría que volver al paso 4 para elaborar otro plan de acción basado en el siguiente problema encontrado.

Como medida preventiva ante posibles contingencias, se recomienda:

- ⇒ Mapa de Red: Mantener un mapa actualizado de la red.
- ⇒ Protocolos: Disponer de una lista de los protocolos que funcionan en la red.
- ⇒ Direcciones y Puertos: Mantener una lista de todas las direcciones y puertos que son accesibles desde el exterior.
- ⇒ Servidores de Red: Tener información sobre la configuración de los servidores de la red, los servicios instalados y los recursos compartidos.
- ⇒ Incidencias Anteriores: Tener información sobre los problemas que han surgido en la red y las soluciones que se les dieron en su momento. La experiencia es un grado y hay que aprender de los propios errores.

1.2.1. Sistemas de gestión/monitorización de equipos

El administrador de un sitio web o de una red basa su trabajo en software de monitorización. Con este puede mantener el control del sistema y puede detectar componentes lentos o defectuosos, siempre en tiempo real.

Estas herramientas de supervisión envían de forma automática actualizaciones de estado, o pueden activar copias de seguridad en caso de fallos o sobrecargas del servidor, conexiones de red y otros factores.

Las herramientas para la monitorización de una red de comunicación pueden ser muy variadas. Las hay que analizan la señal que circula por un cable, programas que monitorizan todo el tráfico de los enlaces, software que intercepta los paquetes de datos, etc.

Entre las distintas herramientas de monitorización de red podemos destacar:

- ⇒ Comprobadores de red
- ⇒ Monitores de red:
 - Muestran un mapa de la actividad de la red en un intervalo de tiempo determinado.
 - Capturan los mensajes que circulan por ella. No decodifican el contenido de los mensajes, sino que se limitan a vigilar los que circulan, su tamaño, los que han llegado con error y el número de ellos que se envían y se reciben por estación.
- ⇒ Analizadores de red

Las soluciones de monitorización y análisis de tráfico tienen una gran importancia para la gestión de las redes de telecomunicaciones.

Estas herramientas son muy útiles. No solo se pueden utilizar para solucionar problemas en poco tiempo, sino que también se pueden usar para prevenir fallos, detectar amenazas, y tomar decisiones correctas en relación a la planificación de la de la red.

Muchos fabricantes de los propios dispositivos de interconexión de redes, ofrecen también aplicaciones software de monitorización de red que distribuyen con sus productos. Lógicamente cada una es diferente pero todas tienen mismo fin: mostrar al administrador el estado de la red, estadísticas de actividad, posibilidad de configuración remota de dispositivos, fallos, sobrecargas, etc.

A modo de ejemplo, presentamos algunas de las aplicaciones de monitorización de equipos que podemos encontrar en el mercado:

- ⇒ **Nagios** es de las herramientas más populares, por no decir que es el sistema más popular de monitorización de red de código abierto disponible. Fue diseñado originalmente para ejecutarse en Linux. Proporciona supervisión de los servicios de red (SMTP, POP3, HTTP, NNTP, ICMP, SNMP, FTP, SSH) y de los recursos de los servidores (carga del procesador, uso de disco, los registros del sistema). Es implementa sobre un Servidor Web, preferiblemente apache. Más adelante la veremos en profundidad.
- ⇒ **Zabbix** es un software gratuito de vigilancia de red muy completo y con soporte comercial. Esto es muy importante. Puede monitorizar los diferentes servicios de la red, servidores y casi cualquier hardware de red. El usuario puede definir las opciones de visualización y de mapa de red. Posee un método de comunicación muy sencillo que permite una configuración rápida de las notificaciones de alerta de eventos predefinidos. Cuenta con tres módulos principales: el servidor, los agentes, y el usuario. Para almacenar los datos del seguimiento, puede utilizar distintos sistemas de gestión de base de datos, tales como MySQL, PostgreSQL, Oracle o SQLite. Sin necesidad de instalar ningún software especial, Zabbix permite comprobar la disponibilidad y capacidad de respuesta de los distintos servicios, como SMTP, FTP o HTTP. Para supervisar las estadísticas, tales como carga de la CPU, utilización de la red y espacio en disco, el agente de Zabbix debe estar instalado en la máquina host. Incluye también soporte para el monitoreo a través de SNMP, TCP y controles ICMP, IPMI y parámetros personalizados como una opción para instalar un agente en los hosts.
- ⇒ Cacti
- ⇒ NetSupervisor
- ⇒ Zenoss
- ⇒ Munin
- ⇒ EliteView
- ⇒ Axence NetTools: es un potente software gratuito para Windows 7 que ofrece a los administradores de la red 12 herramientas para gestionar la comunicación

1.2.1.1. Descripción general. Principios de funcionamiento. Alarmas

Hay varios tipos de herramientas que se encargan del monitoreo y análisis de la red. En particular, los denominados **sniffers** son de gran utilidad. Sniffer es una marca registrada de Network Associates, Inc. Sniffer es una denominación aceptada para aquellas herramientas cuya función principal es monitorizar y analizar tráfico, o sea, examinar paquetes, protocolos y tramas enviadas a través de la red.

La idea es la captura y visualización de las tramas de datos, así como el análisis de tráfico de estas tramas: sus orígenes y destinos.

Los administradores de la red, una vez tienen los paquetes de datos y la información del flujo de tráfico, pueden ver el comportamiento de la red, las aplicaciones y servicios disponibles, la carga de cada uno de estos, la utilización del ancho de banda, anomalías en materia de seguridad, etc.

Los sniffers existen desde que existen las propias redes de ordenadores y se han usado fundamentalmente con dos objetivos:

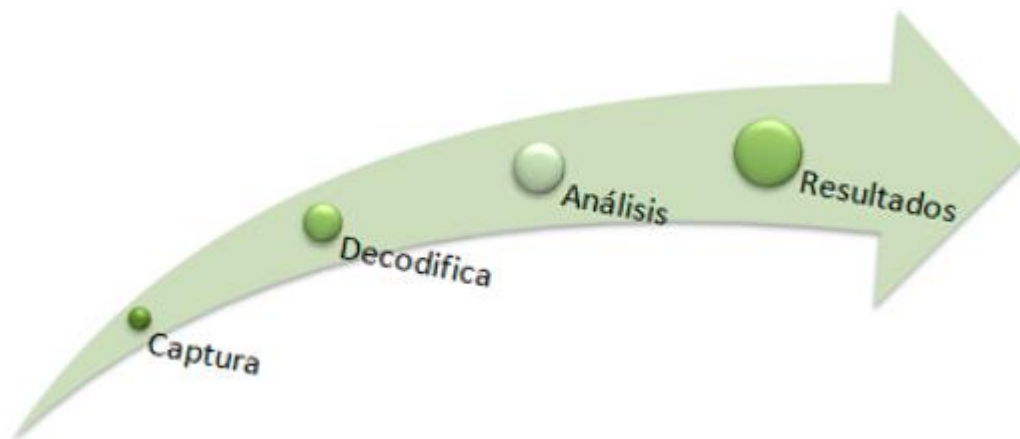
- ⇒ Apoyar a los administradores en el correcto funcionamiento y mantenimiento de la red.
- ⇒ Facilitar a aquellos individuos malintencionados el acceso a ordenadores, servidores y dispositivos como Routers y Switch sin la autorización pertinente.



Este tipo de programas son muy utilizados por hackers y atacantes informáticos, como una herramienta de ataque a redes. Pero también, es utilizado por administradores de redes para mantener la seguridad de sus redes (identificando las vulnerabilidades que puedan presentar) e, incluso, para gestionar de manera eficiente la red, ya que puede identificar la estabilidad de dicha red con una facilidad tremenda y realizar auditorías de redes LAN en muy poco tiempo.

Son las dos caras de una misma moneda.

Un sniffer puede estar basado en hardware y/o software, pero absolutamente todos ellos, interceptan y recolectan el tráfico de datos a través de la red. Una vez se ha capturado el tráfico, el sniffer puede decodificarlo y analizar el contenido de los paquetes para después, poder ofrecer la posibilidad de monitorizarlo en pantalla.



Un sniffer puede capturar los paquetes que están circulando por la red, una vez capturados, puede acceder a su contenido.

Sin embargo, si tenemos varias redes unidas, para poder capturar tráfico de todas ellas, hay que habilitar algunas extensiones o modificar la infraestructura de la red (esto último no suele ser muy factible).

Hay un procedimiento llamado **técnica de puerto espejo**, para conseguir que los Switch que gestionan las distintas redes, reenvíen todos los paquetes de datos hacia un puerto donde puede capturarlos el sniffer.

Los sniffers tienen en su contra que las tarjetas de red Ethernet están construidas de tal forma que, en su modo normal de operación, sólo capturan las tramas que van dirigidas hacia ellas.

Por lo tanto, la tarjeta no procesa todo el tráfico de datos que llega a ella y habría que activar un modo especial de funcionamiento de la tarjeta. Este modo de funcionamiento, permite a la tarjeta de red procesar todo el tráfico que le llega, siendo éste el modo de trabajo que un sniffer necesita para llevar a cabo su misión.

En el caso de las de redes inalámbricas, pretende ver los datos transferidos entre un punto de acceso a otro o entre dos nodos. Al igual que en las redes cableadas, se deben chequear algunos aspectos: la tarjeta de red inalámbrica debe permitir ser puesta en modo monitor, se necesita escoger el sistema operativo donde se trabajará y en dependencia, se deberán instalar los programas y el driver que permita que la tarjeta trabaje ese modo.

En la actualidad los sniffers de paquetes son muy populares en el mundo de las redes de comunicaciones.

Utilidades de los sniffers:

Para la gestión de la red, los sniffer ofrecen una gran cantidad de prestaciones:

- ⇒ Vigía de tráfico en redes LAN y WLAN.
- ⇒ Captura de tráfico a través de las diferentes interfaces de red.
- ⇒ Capacidad de examinar, salvar, importar y exportar capturas de paquetes en diferentes formatos de captura.
- ⇒ Comprensión de protocolos de las diferentes capas de la arquitectura de comunicaciones.
- ⇒ Aplicación de filtros para limitar el número de paquetes que se capturan o se visualizan.
- ⇒ Cálculo de estadísticas y gráficas detalladas con indicadores como paquetes transmitidos y perdidos, velocidad promedio de transmisión, gráficos de flujo de datos, entre otras.
- ⇒ Detección de los nodos que se encuentran en la red, ofreciendo información como sistema operativo, fabricante de la interface, entre otras.
- ⇒ Reconstrucción de sesiones TCP.
- ⇒ Análisis y recuperación de tráfico VoIP.
- ⇒ Generan informes de tráfico en tiempo real y permiten configurar alarmas que notifiquen al usuario ante eventos significativos como paquetes sospechosos, gran utilización del ancho de banda o direcciones desconocidas.

Aplicaciones de los sniffers:

Los usos típicos de un sniffer, bien sea por administradores de red o intrusos, incluyen los siguientes:

- ⇒ Conversión del tráfico de red en un formato entendible por los humanos.
- ⇒ Visualización de información relevante como un listado de paquetes y conexiones de red, estadísticas detalladas de las conexiones IP, entre otras.
- ⇒ Captura automática de contraseñas y nombres de usuario de la red.
- ⇒ Análisis de fallos para descubrir problemas en la red.
- ⇒ Monitorización del tráfico, mediante el cual es posible descubrir cuellos de botella.
- ⇒ Recuperación de archivos y mensajes intercambiados
- ⇒ Detección de puntos de acceso no autorizados.
- ⇒ Detección de intrusos.

Limitaciones de los sniffers:

Como se puede apreciar sniffers son de gran utilidad y presentan muchas aplicaciones, pero a su vez, poseen una serie de limitaciones:

- ⇒ Por lo general, una misma herramienta no tiene todas las utilidades necesarias para la monitorización y análisis de tráfico.
- ⇒ Las interfaces inalámbricas, tiene bastantes limitaciones para poder trabajar en modo monitor.
- ⇒ Algunas solo trabajan en modo consola.
- ⇒ Hay alguna problemática en la gestión de la información capturada.
- ⇒ No realizan tareas activas para el reconocimiento de las redes, los nodos y otros elementos. Otra gran limitación, para la monitorización de redes, es el uso de encriptación y de los cifrados de claves inalámbricas

- ⇒ La potencia de señal recibida va en función de la localización física del sniffer inalámbrico, pues esta no será la misma en función de que esté más alejado o más cercano a los puntos de acceso.
- ⇒ Un sniffer inalámbrico es capaz de capturar solamente el tráfico del área local donde está instalado. Si hay otras redes gestionadas con Switch, no puede monitorizarlas
- ⇒ Se agudizan las dificultades con los driver de interfaces inalámbricas que permiten trabajar en modo monitor.

La complejidad de las redes informáticas de las empresas va en aumento día a día. Las redes, cada vez más, soportan aplicaciones y servicios estratégicos de las organizaciones empresariales. Por todo lo anterior, el análisis y monitorización de redes se ha convertido en una labor cada vez más importante para evitar problemas.

Con el término Monitorización nos referimos a un sistema que constantemente rastrea una red de ordenadores para detectar sistemas lentos, mal funcionamiento, etc. y que notifica al administrador de la red los posibles fallos bien por correo electrónico, por sms u otros medios.

Hay un software llamado Nagios que se utiliza para controlar y observar el funcionamiento de una red, con el objetivo de prevenir errores en un sistema.

Nagios es un sistema de monitorización de servidores, aplicaciones y redes, alertando en caso de problemas como caídas o restauración de servicios. Es quizás el más utilizado en la actualidad.

Originalmente su nombre era NetSaint y se transformó posteriormente en el proyecto Nagios. Se trata de una herramienta que supervisa hasta el mínimo detalle del estado de los equipos, informando a los administradores en caso de detectarse cualquier tipo de problema.

1.2.1.2. Bloques funcionales. Procedimientos de análisis e identificación de fallos

El núcleo de Nagios, que es el que forma la lógica de control de la propia aplicación, contiene el software necesario para la monitorización de los servicios y de los equipos de la red que se hayan definido.

Este núcleo de Nagios usa diversos componentes que provee la aplicación. Aquí cabe reseñar que también se pueden usar otros componentes realizados por terceros.

Nagios no es un sistema de monitorización y gestión basado en SNMP. Realiza toda su labor basándose en una muchos módulos software que realizan chequeos de parte de la red que corresponda en cada caso. Independientemente de esto, también permite la captura de paquetes SNMP para notificar sucesos.

Por otro lado, Nagios, muestra los resultados de la monitorización y del uso de los diversos componentes en una interfaz web a través de un conjunto de páginas HTML que vienen incorporadas en la propia distribución base. Esto permiten al administrador de la red una completa visión de qué está ocurriendo, cómo, dónde y el por qué.

Además, Nagios guardará los históricos en una base de datos para que al detener y reanudar el servicio de monitorización, todos los datos sigan sin cambios en el mismo estado en que se encontraban.

1.2.1.3. Procedimientos de recuperación de fallos. Ejemplos y casos prácticos

Caso Práctico 1: Problemas de conexión a la red de un equipo

Los problemas de acceso a la red de un equipo son de los más comunes a los que se enfrenta el administrador de red. Llamadas de los usuarios con argumentos como "no puedo acceder a Internet", "no puedo entrar a mi carpeta en el servidor" o "no puedo ver las carpetas compartidas en la red" son las que más horas consumen a los administradores.

Ante una situación de este tipo, lo primero que debe hacer el administrador es comprobar si el fallo es debido a una instalación incorrecta del adaptador, a que no hay conexión con la red o a que existe algún protocolo que no funciona correctamente. A partir de ahí, se puede centrar en cuestiones más específicas hasta que llegue a la causa o causas que han generado el problema.

En general, los pasos que se deben seguir para resolver este tipo de problemas son:

1. Consultar la configuración de red del equipo y compararla con la que figura en nuestra base de datos. Los parámetros que hay que consultar son los correspondientes a TCP/IP y a la red Microsoft (si es que el equipo lleva Windows instalado). En caso de que no coincidan, establecer los valores que correspondan.
2. En caso de que el equipo obtenga su configuración de red a través del protocolo DHCP, comprobar que esta asignación se realiza correctamente. En caso contrario, revisar la configuración del servidor DHCP.
3. Utilizar el comando PING para comprobar de una forma rápida si en equipo puede comunicarse con su encaminador o puerta de enlace por defecto (o servidor proxy). En caso de que la respuesta sea positiva, entonces descartaremos un problema de configuración del equipo o de acceso a la red. De lo contrario, nuestro trabajo deberá centrarse primero en el equipo con problemas.

Si el equipo no recibe respuesta a los mensajes PING enviados, seguiremos los siguientes pasos:

1. Comprobar si el adaptador está correctamente instalado o si existe algún conflicto con otro dispositivo.
2. Comprobar si el cable de conexión está en perfectas condiciones. Para ello, seguir las indicaciones establecidas en el apartado siguiente.
3. En último caso, la falta de conexión con la red puede ser debida a una cantidad de ruido excesivo en el cable.

En caso de que el equipo tenga acceso a la red, deberemos seguir otras comprobaciones, que son:

1. Comprobar la configuración de encaminamiento en el encaminador por donde debe circular la información.
2. Realizar consultas a los servidores DNS para comprobar que la resolución de nombres funciona correctamente.

Caso Práctico 2: Problemas de conexión de los distintos dispositivos de red

Los fallos de conectividad producidos en los dispositivos de interconexión de red (ya sean ordenadores, concentradores, conmutadores, puentes, etc.) tienen una repercusión muy importante ya que producen la incomunicación de equipos individuales o de partes o segmentos de la red. Estos fallos pueden ser debidos a problemas físicos de los dispositivos o el cable, a cuestiones de configuración o a las características del tráfico.

Como regla general, un cable tiene conexión cuando ha sido enchufado a dos dispositivos activos en sus dos extremos.

En ese momento, se activan los indicadores luminosos de los puertos de esos dispositivos (normalmente en color verde o naranja), lo que indica que hay comunicación por ese enlace.

Cuando se detectan problemas de comunicación entre dos dispositivos, lo primero que hay que hacer es comprobar si el enlace es la causa. Para ello, se deberán seguir los siguientes pasos, siempre aplicados en los dos extremos del enlace involucrado:

1. Comprobar el indicador luminoso de estado del puerto. En caso de que no se encuentre activo, entonces el problema puede ser debido a ese puerto, al puerto del otro extremo o al cable que los conecta. En algunos casos, también pueden verse involucrados otros dispositivos intermedios, como enchufes de pared o adaptadores para distintos tipos de cable.
2. Verificar los conectores para asegurarse de que se ajustan correctamente a los puertos y no están sueltos. En algunos casos, cuando se manipula la instalación, puede ocurrir que el conectar se haya soltado ligeramente. También puede ocurrir que se haya deteriorado la pequeña pinza de plástico que llevan los conectores, por lo que se recomienda cambiar el cable ya que no es capaz de ajustar al puerto y puede soltarse en cualquier momento.

3. Realizar una inspección de los pines del conector y del puerto, para comprobar si se ha deteriorado alguno o si están sucios. También se debe comprobar que el conector se ha insertado en el puerto correcto.
4. Comprobar que el puerto del dispositivo no ha sido desactivado desde su programa de configuración. Hay que tener en cuenta que dispositivos como tarjetas de red o conmutadores permiten opciones para desactivar sus puertos de forma individual. También puede haber sido desactivado automáticamente debido a errores de configuración (en el caso de los adaptadores de red, consultar el apartado anterior). Cuando se desactiva un puerto conectado a un extremo de un cable entonces el puerto del otro extremo también pasa al estado de no conectado y el indicador luminoso se apaga.
5. Cuando los indicadores luminosos de los puertos están activos quiere decir que hay conexión pero no se asegura que el enlace funciona correctamente. Puede ocurrir que el cable esté en mal estado, por lo que se recomienda utilizar otro cable, asegurándonos que es del tipo correcto y está en perfectas condiciones (utilizando un comprobador de cableado, por ejemplo). Si no se dispone de comprobador, entonces se podrá probar el cable en puertos distintos o en otros dispositivos.
6. Comprobar que la longitud del cable no excede los límites máximos establecidos en los estándares.
7. Comprobar si los parámetros de negociación de los puertos son correctos. En muchos casos, sobre todo cuando se conectan dispositivos de diferentes fabricantes, la negociación automática que realizan puede no llegar a buen término. Los parámetros que se suelen negociar en un enlace son la velocidad de transmisión, el control de flujo o la transmisión simple/dúplex.
8. Si se trata de un cable de cobre, comprobar que el enlace es cruzado o normal para ese caso concreto.
9. Si se trata de un cable de fibra óptica, comprobar que se está utilizando el tipo correcto (monomodo o multimodo) para los tipos de puertos a conectar y la distancia a cubrir. Hay que comprobar también que el puerto de transmisión de un extremo está conectado con el puerto de recepción del otro extremo.

1.2.1.4. Escalados. Eventuales planes de contingencia/business continuity

¿Por qué se necesita un Plan de Contingencia?

A medida que las empresas son más dependientes de los ordenadores y de las redes para manejar sus actividades, la disponibilidad de los sistemas informáticos es de crucial importancia. Actualmente, la mayoría de las empresas necesitan un nivel alto y continuo de disponibilidad, ya que les resultaría extremadamente difícil funcionar sin los recursos informáticos.

Los procedimientos manuales, si es que existen, sólo serían prácticos por un corto periodo de tiempo.

En caso de un desastre informático, la interrupción prolongada de los servicios de gestión de la información, puede llevar pérdidas financieras significativas.

Lo más grave es que puede afectar a la credibilidad de la empresa con la consiguiente disminución de público o clientes y, como consecuencia, la empresa puede terminar en un desastre total.

Cabe preguntarse "¿Por qué se necesita un plan de contingencia para desastres si existe una póliza de seguro para esta eventualidad?" La respuesta es que si bien, el seguro puede cubrir los costos materiales de los activos de una organización en caso de un siniestro inesperado, no servirá para recuperar el negocio. No ayudará a conservar a los clientes y, en la mayoría de los casos, no proporcionará fondos suficientes para mantener funcionando el negocio hasta que se haya recuperado la normalidad.

Por lo tanto, la capacidad para recuperarse con éxito de los efectos de un desastre, dentro de un periodo predeterminado, debe ser un elemento primordial en el plan estratégico de seguridad para la vuelta a la normalidad de una empresa.

Imagínenos que se interrumpa la actividad de los ordenadores durante un periodo de tiempo prolongado; imaginemos la pérdida de todos los datos de la empresa o la destrucción de equipos vitales del sistema

¿Cómo manejar una situación tan catastrófica?

Si se llega a este punto, la única forma de afrontar este problema es tener una solución completa, efectiva y totalmente probada. Un plan de contingencia robusto, que abarque todos los posibles escenarios adversos.

¿Qué es un desastre?

Consideraremos como una situación adversa grave la interrupción más o menos prolongada de los recursos informáticos y de comunicación de una organización, que no pueda solventarse dentro de un periodo de tiempo aceptable y que además necesita un equipo de apoyo para su recuperación y vuelta a la normalidad.

Ejemplos sencillos son los grandes incendios, las inundaciones, los terremotos, las explosiones, los actos de sabotaje, etcétera.

Pero, desde nuestro punto de vista también lo es la caída de tensión prolongada, la interrupción, a lo largo de un periodo más o menos largo, del servicio de internet. Incluso podemos considerar como desastre, la rotura o fallo de un simple cable de red, de un conector (de los cientos de ellos que puede haber en la empresa). O si un virus deja no operativa la intranet, la extranet, eso también es un desastre.

En definitiva, todo lo que pueda paralizar el correcto funcionamiento de las redes, va a ser un desastre.

Plan de contingencia:

La reanudación de las actividades habituales después de un desastre es una de las situaciones más difíciles a las que una organización debe enfrentarse. Es probable que no se pueda regresar al lugar de trabajo o que no se disponga de ninguna de los recursos acostumbrados. Incluso, es posible que no esté disponible todo el personal. La preparación es la clave del éxito para enfrentar los problemas.

Desde un punto de vista económico, no se puede asegurar una protección total contra todo tipo de riesgos, particularmente amenazas naturales a gran escala, que pueden arrasarse zonas extensas. Como consecuencia, siempre se tiene que tolerar algún riesgo residual.

Los directivos de una gran empresa deberán decidir el alcance de sus actuaciones ante un posible desastre, es decir el nivel al que tendrán que estar preparados.

Por ejemplo, la mayor parte de las empresas implementan una estrategia que proteja contra desastres locales, pero pocas cubren desastres a nivel nacional o incluso internacional. Normalmente, las organizaciones que cuentan con varias localizaciones, pueden tener una estrategia de recuperación que funcione en caso de que un sitio sea destruido o dañado, pero si varios sitios son destruidos al mismo tiempo el plan de ataque es más complicado o simplemente no existe.

Un plan de contingencia es el proceso que determina qué hacer si una catástrofe afecta a una empresa y es necesario recuperar la red y los sistemas.

Normalmente, es más fácil pensar en un plan de contingencia que llevarlo a cabo. Con la cantidad de trabajo que la mayoría de los gerentes tienen, tienden a ir posponiéndolo.

Es complicado determinar el principio, decidir por dónde empezar.

Metodología para el plan de contingencia:

El diseñar e implementar un plan de contingencia para recuperación de desastres no es una tarea fácil; se pueden necesitar esfuerzos y gastos importantes, sobre todo si se está partiendo de cero.

Una solución adecuada necesita los siguientes requerimientos

- Debe ser diseñada y elaborada de acuerdo con las necesidades de la empresa.
- Se necesitará la construcción o adaptación de un sitio para los equipos computacionales.
- Requerirá del desarrollo y prueba de muchos procedimientos nuevos, y éstos deben ser compatibles con las operaciones existentes.
- Se necesitará la colaboración del personal de muchos departamentos diferentes, deberán trabajar conjuntamente cuando se desarrolle e implemente la solución.
- Implicará un compromiso entre costo, velocidad de recuperación, medida de la recuperación y alcance de los desastres cubiertos.

- Es **imprescindible, un método estructurado** que asegure que se tienen en cuenta todos estos factores y que se les trata adecuadamente.

Las etapas principales para la planificación e implementación de un plan de contingencia son:

1. Identificación de los riesgos.
2. Evaluación de dichos riesgos.
3. Asignación de prioridades a las aplicaciones.
4. Establecimiento de los requerimientos de recuperación.
5. Elaboración de la documentación.
6. Verificación e implementación del plan.
7. Distribución y mantenimiento del plan.

Veamos detalladamente cada uno de ellas.

Identificación de los riesgos:

La primera fase del plan de contingencia, está relacionada con tres simples preguntas: ¿qué está bajo riesgo?, ¿qué puede ir mal? y ¿cuál es la probabilidad de que suceda?

¿Qué está bajo riesgo?

Necesitamos **determinar todos los componentes del sistema susceptibles de ser dañados**, dando lugar a la pérdida de conectividad, **tanto equipos como datos**. **Un diagrama de la arquitectura de todos los componentes del sistema nos facilitará la realización de un inventario** de los elementos que pueden necesitar ser sustituidos tras un desastre. No hay que olvidar que **el software también necesita ser reemplazado**, y que todos los productos software importantes han de ser identificados.

Una omisión en el inventario puede dar lugar a una recuperación errónea después del desastre. El sistema de aplicación no funcionará correctamente si alguno de sus componentes no está disponible; en tal caso, es aconsejable **estar constantemente a la expectativa de los nuevos elementos que pueden haberse olvidado**.

Uno de los aspectos menos agradables a tener en cuenta, y que a menudo se pasa por alto, **es que las personas esenciales se ven afectadas por el desastre**, puede haber pérdidas humanas, y es necesario reemplazarlas para realizar sus labores.

Una buena formación en el manejo y funcionamiento de los sistemas por parte de todos los miembros del equipo técnico puede ayudar y mucho a reducir el impacto que supone la pérdida de uno de los colaboradores.

Al menos, los manuales de las aplicaciones más importantes para la empresa deberían encontrarse disponibles en un sitio externo.

¿Qué puede ir mal?

Lo más difícil en el plan de contingencia es responder a esa pregunta. La respuesta varía desde lo evidente hasta lo casi increíble.

La ley de Murphy es clara al respecto: **¿Qué puede ir mal?: TODO**. No podemos olvidar nada. Desde una explosión de un equipo hasta la simple rotura de un pequeño cable.

¿Cuál es la probabilidad de que suceda?

Si se tuviera una cantidad ilimitada de recursos y fuera posible protegerse contra todas las calamidades, esta pregunta carecería de interés. Sin embargo, no se dispone de recursos infinitos; de hecho, los recursos son bastante escasos. Por lo tanto, **se deben seleccionar los tipos de desastres contra los que uno intentará protegerse**.

Obviamente, estos preciados recursos se emplearán en aquellos desastres que tengan la mayor probabilidad de afectar a la organización.

Por ejemplo, se podría intentar proteger los sistemas de la caída sobre el edificio de un meteorito procedente del espacio exterior. Sería Magnífico prepararse para tal situación, pero siendo razonables, esto no va a ocurrir. Sería mejor prever una posible caída de la tensión eléctrica.

Responder a la pregunta: **¿cuál es la probabilidad de que suceda?** también **requiere de ciertas consideraciones presupuestarias**. Ello puede ayudar a asumir distintos presupuestos para comprender cuáles son los costes para diferentes niveles de protección y preparación.

Finalmente, se puede estar expuesto a ciertas amenazas cuya protección no está al alcance del presupuesto, pero, al menos, se es consciente de su existencia y, por lo tanto, es posible mejorar el plan en un futuro.

Evaluación de los riesgos:

Es el proceso que determina el coste, para la empresa, de sufrir un desastre que afecte a su actividad. Si una inundación impidiera la actividad comercial durante cinco días, la compañía perdería cinco días de ventas, además del deterioro físico de los edificios e inventario. En el caso de los sistemas informáticos, la preocupación principal es comprender la pérdida económica que supone la interrupción de los servicios, incluyendo los que se basan en las redes.

Por ejemplo, si la empresa realiza negocios en Internet, ¿cuál es el costo de tener el servidor web no operativo?

¿Cuál es el impacto sobre la productividad de la empresa?

Los costes de un desastre pueden clasificarse en las siguientes categorías:

Costes reales de reemplazar el sistema informático

El coste real de los equipos y el software es fácil de calcular, y es imprescindible disponer de un buen inventario de todos los componentes de la red.

Costes por falta de producción

Estos pueden determinarse midiendo la producción generada asociada a la red. La pérdida de producción, debida a la interrupción de la red, se puede ser calcular utilizando esta información.

Costes por pérdida de negocio

Los costes por pérdida de negocio son los ingresos perdidos por las organizaciones cuando la red no está disponible.

Si el sistema de solicitud de pedidos no funciona y la empresa sólo es capaz de procesar el 15% del volumen diario habitual de ventas, entonces se ha perdido el 85% de ese volumen de ventas.

Costes de reputación

Los costes de reputación son más difíciles de valorar pero es conveniente incluirlos en la evaluación. Estos costes se producen cuando los clientes pierden la confianza en la empresa. Los costes de reputación crecen cuando los retrasos en el servicio a los clientes son más frecuentes.

Asignación de prioridades en las aplicaciones:

Cuando ocurre un desastre y se inicie la recuperación de los sistemas, debe conocerse qué aplicaciones recuperar en primer lugar. No hay que perder el tiempo restaurando los datos y sistemas equivocados, sino que **hay que restaurar primero sus aplicaciones esenciales, las que la actividad empresarial necesita.**

Es necesario determinar por anticipado cuáles son las aplicaciones fundamentales del negocio. Habrá aplicaciones "muy importantes" dependiendo de a quién se le pregunte. El departamento de recursos humanos afirmará que el sistema de nóminas es el más importante, el departamento de ventas dirá que es su sistema de entrada de pedidos, el departamento de producción insistirá en su control de inventario y para el departamento de compras el más importante es su sistema de facturación. **Desgraciadamente, no todos estos sistemas pueden ser el más importante; por lo tanto, es fundamental decidir cual lo es realmente.**

Ventas

Web

Personal

Compras

Servicios

Administración

CPR

Fiscal

El plan de contingencia debería incluir la lista de los sistemas y la prioridad de estos.

Tenemos que tener claro todo lo que se va a restaurar y todo lo que necesitamos para poder hacerlo.

Un sistema de red de trabajo en grupo está compuesto por los servidores, donde las aplicaciones almacenan sus datos, las estaciones de trabajo que los procesan, las impresoras, la red que interconecta todo, y el software de las aplicaciones.

Las aplicaciones cliente/servidor o distribuidas aumentan la complejidad ya que necesitan que distintas partes de la aplicación residan en máquinas separadas.

Podemos pensar en construir una infraestructura superior a la necesaria para las aplicaciones de mayor prioridad.

Imaginemos que si actualmente la red tiene 50 estaciones de trabajo, se puede comenzar a trabajar inmediatamente en la reconstrucción de las 50 estaciones de trabajo. Sin embargo, si las aplicaciones prioritarias sólo necesitan cinco estaciones de trabajo, nos centraremos en la reconstrucción de esas cinco estaciones de trabajo y concentraremos los esfuerzos en lograr que las aplicaciones funcionen correctamente.

Es mucho mejor intentar que funcione un sistema pequeño, que no uno más grande, ahorrando así mucho tiempo en el proceso.

Una de las ventajas del enfoque basado en el sistema de aplicaciones es la cantidad de tiempo necesaria para recuperar una aplicación comparada con la cantidad de tiempo requerida para restaurar un servidor en su totalidad. Si la aplicación tiene sólo 600 MB de datos y el servidor 6 GB, es obvio que se ahorra una gran cantidad de tiempo recuperando únicamente la aplicación.

Establecimiento de requisitos de recuperación:

Es imprescindible en esta fase del proceso de elaboración del plan de contingencias definir un periodo de tiempo aceptable y viable para lograr que la red esté de nuevo activa. La principal preocupación debe ser disponer de las aplicaciones más importantes en primer lugar.

Hay que hacer una valoración correcta sobre el tiempo necesario y no realizar estimaciones poco realistas.

El tiempo tiene que ser suficiente para recuperar las copias de seguridad de los sistemas de almacenamiento, actualizar el sistema de copias de seguridad e incluso para adquirir los sistemas necesarios.

Un sistema de cinta que recupera datos a 2 MB por segundo realizará la labor mucho más rápido que uno que lo ejecute a 500 KB por segundo.

Hay que ser precavido y no suponer que se pueden hacer muchas cosas al mismo tiempo; uno se puede encontrar cometiendo desafortunados errores que frenan la labor si no se presta atención al trabajo que se tiene entre manos.

Elaboración de la documentación:

Elaborar un documento que sirva de referencia para toda la organización es quizás lo más difícil del plan de contingencia.

Implicará un sobreesfuerzo para las personas que lo realicen, pero ayudará a conocer el sistema y puede que algún día salve la empresa.

La empresa debe aportar los recursos necesarios para escribir y mantener este documento de forma que su realización sea un éxito.

Uno de los problemas con el que nos podemos encontrar en la elaboración del plan de contingencias es que la tecnología de redes cambia muy rápidamente, por lo que supone un esfuerzo extra mantenerlo al día.

La aparición de nuevos dispositivos, nuevos sistemas de aplicación introducen su propio nivel de complejidad.

Ejemplo:

Suponiendo que se ha de realizar la recuperación de un sistema, y este sistema es una gran base de datos relacional. Se necesitará no solo la instalación de la base de datos, que puede realizarla un administrador de redes, sino un administrador de bases de datos pues son necesarios unos conocimientos más complejos.

La actualización del plan de contingencia debe realizarse periódicamente, por ejemplo una vez al año.

Aunque al principio supone una gran cantidad de trabajo, una vez que se dispone del plan de contingencias redactado las actualizaciones son relativamente fáciles.

Este punto es de vital importancia.

Contenido del plan de contingencia:

El plan de contingencia debe intentar definir las cinco áreas siguientes:

Listas de notificación, números de teléfono, mapas y direcciones

Cuando ocurre un desastre debe estar preestablecido a quien se avisa en primer lugar.

Por ejemplo, si hay un incendio, llamar primero a los bomberos y luego al director general. Si hay un fallo eléctrico, primero al electricista y después al responsable de hardware.

Los mapas con las ubicaciones del centro de operaciones temporal y la instalación externa pueden ahorrar mucho tiempo.

Prioridades, responsabilidades, relaciones y procedimientos

Se empezará por recuperar inmediatamente las aplicaciones de mayor prioridad.

Deberán estar identificadas las instrucciones y responsabilidades precisas.

La relación entre tareas debe estar documentada de manera que pueda identificarse cualquier problema que pudiera surgir.

Hay que especificar, de manera detallada, las operaciones y tareas relacionadas con labores de instalación y recuperación necesarias, debiendo ser fáciles de leer y seguir.

También habría que incluir aquí los números de teléfono de las empresas de asistencia que pudieran necesitarse (servicios técnicos, proveedores de servicios de internet, etc.).

Información sobre adquisiciones y compras

Hay que tener claro como y a quien se va a comprar todo lo necesario para recuperar a la organización de un desastre. Es aconsejable disponer de copias de las facturas, recibos para mostrarlos como prueba de compra.

Es bueno disponer de lista de los números de serie de los equipos hardware.

Actualmente, los productos para las redes de área local y para las comunicaciones se venden a través de grandes sistemas de distribución.

Diagramas de las instalaciones

Los diagramas de red simplifican en gran medida la labor de construir una red.

Un diagrama detallado de la red, necesaria para las primeras aplicaciones, facilita y agiliza la reanudación de las actividades.

La asignación de [etiquetas a los cables](#) y su almacenamiento en un lugar adecuado, no llevará mucho tiempo y evitará muchas confusiones en el futuro.

Otra de las ventajas de utilizar un diagrama de conexiones detallado es la posibilidad de emplear contratistas para realizar las instalaciones.

Sistemas, configuraciones y copias de seguridad en cinta

Para ahorrar gran cantidad de tiempo en el proceso de recuperación existe la posibilidad de almacenar algunos sistemas de repuesto con la capacidad de gestionar tareas diferentes. Se puede planificar la instalación de una configuración genérica que permita ejecutar las aplicaciones de mayor prioridad sin problemas.

Si se desconoce los productos que la gente tiene en sus PC, un producto para inventario de LAN puede ayudar en la recopilación de esta información. Después de que la red alternativa se encuentre funcionando, con cierta garantía, se podrá restaurar los PC con sus configuraciones anteriores utilizando la información de configuración extraída de los informes de inventario.

Asegure la disponibilidad de un sistema de copias de seguridad en funcionamiento.

Debe mantenerse un sistema de reserva, incluyendo adaptadores SCSI, cables y software de unidades de dispositivo, en un sitio alternativo, siempre que sea posible.

Verificación e implementación del plan:

Una vez redactado el plan, hay que probarlo. Hay que verificar que el plan funciona correctamente. Para ello, se debe ser escéptico sobre el propio trabajo, de manera que pueda uno probarse a sí mismo que funciona.

Esto no es fácil porque se habrá invertido una gran cantidad de tiempo y energía personal en el proceso de desarrollo del plan, aunque lo mejor sería, situarse de forma imparcial ante la confiabilidad del plan. Lo más práctico es realizar las pruebas para encontrar problemas, no para verificar que el plan funciona. Si existen errores hay que tomar nota de ellos y corregir el plan.

Comprobación del plan por partes:

No se puede provocar la caída del sistema para comprobar si la recuperación se hace correctamente. Hay formas mejores de verificar un plan de contingencia sin provocar grandes interrupciones en el trabajo de la organización.

Algunas de las cosas en las que habitualmente no se piensa a la hora de comprobar pueden ahorrar mucho tiempo posteriormente. Por ejemplo, [llamar a los números telefónicos](#) de los colaboradores incluidos en las listas telefónicas del plan [para confirmar si son correctos](#); llamar a los vendedores y comprobar si disponen de existencias de productos, ya que puede que hayan modificado su política de inventario.

También [hay que verificar los procedimientos que se emplearán para recuperar los datos](#). O sea, comprobar el software para la realización de las copias de seguridad, hay que asegurar la recuperación de las aplicaciones de mayor prioridad de la manera deseada.

En este punto, es necesario actualizar el plan para incluir información sobre el establecimiento de cuentas de usuario.

Hay que comprobar cada una de las operaciones del plan individualmente y como consecuencia ver si se obtiene un sistema de red en funcionamiento.

También sería conveniente verificar el plan con otras personas de la organización que se encuentren familiarizadas con los procedimientos empleados.

Cada día hay que revisar la parte del plan relacionada con las operaciones de copias de seguridad verificando la finalización correcta de las mismas. Supervisar esta acción es imprescindible, asegurando que algunas personas de la organización saben realizar copias de seguridad adecuadamente, y comprobar su finalización.

Distribución y mantenimiento del plan:

Por último, cuando se disponga de una versión del plan definitiva y sobre todo verificada, es necesario distribuirlo a las personas interesadas.

Hay que controlar las versiones del plan, de forma que no haya confusiones.

Es necesario asegurar la disponibilidad de copias extra del plan para su depósito en una instalación exterior o en cualquier otro lugar además del lugar de trabajo. Deberíamos mantener un listado con todas las personas y ubicaciones que poseen una copia del plan. Cuando se actualice el plan, habrá que sustituir todas las copias anteriores y eliminar las versiones antiguas.

El mantenimiento del plan es un proceso sencillo.

Se comienza con una revisión del plan existente, se examina en su totalidad y se llevan a cabo los cambios en cualquier información que pueda haber variado.

A continuación, se vuelven a evaluar los sistemas de aplicación y se determina cuáles son los más importantes para la organización.

Las modificaciones a esta parte del plan causarán modificaciones en los procedimientos de recuperación. Sin embargo, esto no debería verse como un problema ya que la sección de procedimientos tendrá que actualizarse debido a otros cambios.

Si se han realizado modificaciones al sistema de copias de seguridad, hay que asegurarse de incluir la información sobre el funcionamiento del nuevo sistema.

Este proceso llevará su tiempo, pero también tiene otros beneficios importantes, que se percibirán aunque nunca tengan que utilizarse. Más personas conocerán la red, lo que asegurará a la empresa una base técnica más amplia para mantener correctamente la red.

Otra ventaja es la adquisición de una perspectiva global sobre la red, para el núcleo de administradores de sistemas de información y puede ayudar a identificar las futuras áreas conflictivas.

Uno de los aspectos más difíciles en la gestión y administración de LAN, es dar a conocer la situación actual. El mantenimiento y verificación de un plan de migración ayudará a que se produzca y mejore la comunicación entre el personal de la organización.



1.2.2. Herramientas de gestión de incidencias

La gestión de incidencias y peticiones de servicios (el "ticketing") es área que tiene una gran cantidad de software Open Source en cualquiera de sus modalidades (Software Abierto, Software Libre, servicios a través de Cloud, etc.).

El hacer uso de este tipo de software es una buena solución. No sólo por cuestiones económicas, muchas veces el gasto no va más allá del subyacente a los gastos de un servidor web sencillo, sino por la funcionalidad que ofrecen.

Enumeramos a continuación algunas de las herramientas más utilizadas.

RT (Request Tracker)

Es el sistema de Gestión de Incidencias Open Source más utilizado en la actualidad. Tiene licencia GNU y está desarrollado sobre Perl. Para funcionar sólo necesita de un servidor web Apache, bien sea sobre Linux, MAC OS o Windows, el intérprete de PHP, Las extensiones Perl y una base de datos como MySQL, PostgreSQL, Oracle o SQLite (preferiblemente MySQL).

Por lo tanto la instalación y use de RT tiene coste cero ya que podemos instalar todo lo necesario sin coste alguno (Linux+Apache+MySQL, por ejemplo).

Como todas estas herramientas, ofrece un interfaz Web tanto para los usuarios como para los administradores. Unos de los puntos fuertes que tiene esta herramienta es el uso que hace del eMail. En realidad, en muchos casos la interacción con los usuarios se realiza casi exclusivamente a través de este medio. Los correos electrónicos se convierten automáticamente en tickets y las respuestas llegan a los clientes de igual manera. Sólo los administradores trabajan sobre la interfaz Web.

La última versión ha tenido importantes mejoras en la interfaz de usuario, el aspecto visual de las notificaciones y la capacidad para agrupar tickets en diferentes niveles y extraer de ellos información del funcionamiento general del sistema.

RT no dispone de una App para dispositivos móviles, pero su diseño "responsive", hace posible su uso en dispositivos móviles.

OTRS: Open Ticket Request System

OTRS es otro software que distribuye bajo licencia GNU, por lo que, puede instalarse de forma totalmente gratuita.

Es una aplicación web y por lo tanto requiere, al igual que RT (Request Tracker), Apache y una Base de Datos MySQL, PostgreSQL, DB2, Oracle o MS Sql Server.

Tiene una interfaz Web y buena integración con el correo electrónico.

También tiene una FAQ de preguntas frecuentes integrada.

Esto suele ser muy interesante.

En este caso sí que hay una App para iPhone, aunque no para android.

También ofrece facilidades para Time Tracking, Calendarios, Workflows, gestión de SLAs, gestión de problemas y Catálogo de Servicios y dispone de un buen sistema de reporting y un pequeño módulo para la gestión de encuestas a clientes.

Aunque OTRS nació en la comunidad Open Source, desde hace algún tiempo está gestionado por una empresa. Como consecuencia no hay en la actualidad de una comunidad de desarrolladores que den soporte a la aplicación. Este soporte, se puede contratar con la empresa, previo pago de su importe, claro está.

osTicket

Esta es otra herramienta Open Source totalmente gratuita.

Vuelve a ser una aplicación web y por lo tanto requiere de Apache, del intérprete de PHP y de una Base de Datos MySQL, PostgreSQL, DB2, Oracle o MS Sql Server.

Permite el registro de tickets bien desde una página Web o a través del correo electrónico. Se pueden configurar visualmente los formularios de soporte, definiendo los campos que consideremos necesarios en cada caso.

Una funcionalidad interesante y que adolecen otras aplicaciones de este tipo, es el crear filtros para enviar, de forma automática, los tickets que cumplan las condiciones establecidas, a los distintos administradores.

Aunque algo básicos, también dispone de sistemas de reporting, control de SLAs y generación de cuadros de mando.

La herramienta es gratuita aunque es posible contratar servicios adicionales como la instalación, formación de operadores, alojamiento en servidores y sobre todo para la personalización de la herramienta.

1.2.2.1. Descripción general. Funcionalidades

Como hemos dicho con anterioridad, los gestores de incidencias, son herramientas que permiten hacer el seguimiento de las incidencias según un control de flujo definido.

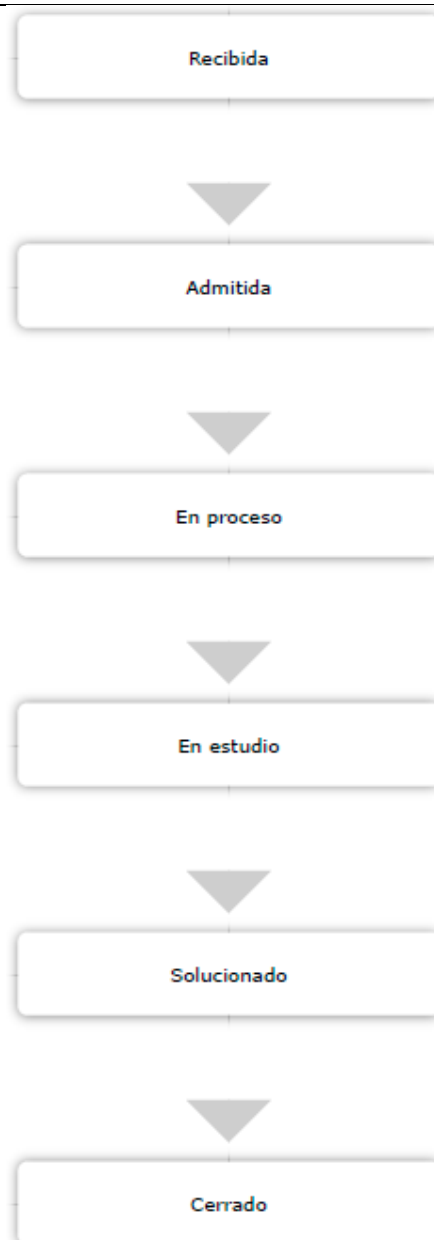
Este mecanismo de control de flujo se activa en el momento en el que se hay una petición de asistencia ante una incidencia determinada.

El usuario puede abrir la incidencia directamente a través de la web o del correo electrónico, o un administrador puede abrirla por su cuenta.

Cada usuario puede en cualquier momento ver el estado de las incidencias que tiene abiertas. También tiene acceso a un histórico de los ticket que ha ido abriendo y que en su momento se cerraron.

Los administradores, o grupos de ellos, pueden ver y tratar las peticiones que tienen asignadas y pueden ir cambiando el estado de la incidencia según vayan solucionándose los temas.

Cambios de estado en el Ticketing:



En cada cambio de estado, el sistema va registrando quién y cuándo se produce el cambio. De esta forma, puede analizarse el tiempo empleado para dar respuesta a cada petición. Esta es una información importante, pues muestra la calidad del servicio de resolución de incidencias.

Este flujo de tratamiento de las incidencias es totalmente personalizable para cada caso concreto. Pueden programarse alertas automáticas por correo electrónico a los usuarios o a los administradores a medida que se van produciendo la petición avanza por el flujo, y con la misma plataforma puede disponerse de disturbios de estado.

1.2.2.2. Procedimientos de gestión/documentación de incidencias

Trac es un sistema de gestión documental, lo que conocemos con el nombre de wiki. Está escrita en Python. Por lo tanto y para su instalación, necesitamos no un servidor web, sino un servidor de aplicaciones.

Es software libre y de código abierto, desarrollado y mantenido por la empresa Edgewall Software. Hasta mediados de 2005 estaba disponible bajo la Licencia Pública General de GNU, pero desde su versión 0.9 se distribuye bajo la Licencia BSD modificada.

El funcionamiento básico pasa por enlazar la información entre una base de datos de errores del software, un sistema de control de versiones y el contenido de una wiki. La mezcla de toda esa información es bastante importante y útil.

Se puede en la interfaz web, cualquier sistema de control de versiones, como pueden ser Subversion, Git, Mercurial, Bazaar o Darcs.

En la actualidad, es uno de los sistemas más utilizados. La versión original está diseñada para integrarse con repositorios de control de versiones de Subversion, pero hay varios plugins que permiten utilizar otros. Especialmente funciona bien con Git.

Veamos cómo podemos instalar y poner en marcha el sistema Trac con un repositorio Git.

1.2.2.3. Notificaciones y escalados (internos y/o a proveedor de servicios)

Una incidencia/incidente es una interrupción no planificada o una reducción de la calidad de un servicio de Telecomunicaciones. Como ya hemos visto, el proceso de gestión de incidencias, cubre cualquier tipo de ellas, ya sean fallos, errores, consultas, etc.

Conceptos fundamentales:

Límites de tiempo

Modelos de incidencias

Incidencias graves

de la resolución de una determinada incidencia, no sea capaz de resolverla en primera instancia.

En este caso, hemos de tener claro que no podemos dejar la incidencia sin resolver, ni tampoco podemos dejarla dando vueltas eternamente. ¿Qué hacer entonces? Quizás debamos derivar la incidencia a otro especialista o a algún superior que pueda tomar decisiones de cómo actuar. A este proceso se le denomina escalado.

Básicamente hay dos tipos diferentes de escalado:

Escalado Funcional

Se requiere el apoyo de algún especialista de más alto nivel para resolver el problema.

Transfiere un incidente o un problema a un equipo técnico con un nivel de experiencia mayor para asistir en la resolución del problema.

Escalado Jerárquico

Debemos acudir a un responsable de mayor autoridad para tomar decisiones que se escapan de las atribuciones definidas para ese nivel. Informa o involucra a los niveles superiores de gestión para la resolución del problema.

Es el caso de la necesidad de asignar más recursos para la solución de la incidencia.

Este escalado no busca tanto contar con un soporte técnico más específico conforme al registro que se está gestionando, y busca más contar con la autorización oportuna para el avance de los trabajos, o bien notificar de una situación anómala o potencialmente peligrosa.

Incidencias graves

Es importante establecer reglas concretas acerca de cuándo deben realizarse (o incluso automatizar en la medida de lo posible) estos escalados jerárquicos, no dejándolos a la decisión y parecer de los técnicos.

1.2.2.4. Procedimiento de escalado (y seguimiento) de problemas no resueltos. Documentación y seguimiento

La normativa SO/IEC 20000-1 especifica los requisitos para el proveedor de servicios para planificar, establecer, implementar, operar, supervisar, revisar, mantener y mejorar un sistema. Los requisitos incluyen el diseño, la transición, la entrega y la mejora de los servicios para cumplir con los requisitos de servicio acordados.

Especificaciones ISO 20000-1:

- ⇒ Se deben registrar todas las incidencias.
- ⇒ Se deben adoptar procedimientos para gestionar el impacto de las incidencias.
- ⇒ Los procedimientos deben definir el registro, la priorización, el impacto en el negocio, la clasificación, la actualización, el escalado, la resolución y el cierre formal de todas las incidencias.
- ⇒ Se debe mantener informado al cliente del progreso de la incidencia sobre la que haya informado o de su petición de servicio, y se le debe advertir por adelantado sobre si sus niveles de servicio no se pueden conseguir, acordando con él las acciones a tomar.
- ⇒ Todo el personal implicado en la gestión de incidencias debe tener acceso a información relevante como, por ejemplo errores conocidos, resoluciones de problemas y la base de datos de gestión de la configuración.
- ⇒ Los incidentes graves se deben clasificar y gestionar de acuerdo con un proceso.
- ⇒ La normativa SO/IEC 20000-2 profundiza la norma anterior, incluyendo un código de buenas prácticas.

Buenas prácticas ISO 20000-2:

El proceso de gestión de incidencias puede ser proporcionado por un servicio de atención al cliente, que actúe como punto de contacto diario con los usuarios.

La gestión de incidencias debería ser un proceso tanto proactivo como reactivo, que responda a los incidentes que afecten, o que eventualmente pudieran afectar al servicio.

La gestión del servicio debería ser un proceso centrado en la restauración del servicio a los clientes y no en la determinación de la causa de las incidencias.

El proceso de gestión de incidencias debería incluir lo siguiente:

1. La recepción, el registro, la asignación de prioridad y la clasificación de las llamadas.
2. La resolución de primer nivel o la derivación.
3. La consideración de cuestiones de seguridad.
4. El seguimiento y la gestión de ciclo de vida de las incidencias.
5. La verificación y el cierre de las incidencias.
6. El contacto de primer nivel con los clientes.
7. El escalado.

Para evitar la aparición de las incidencias no resueltas en el sistema, debemos seguir ambas recomendaciones.

Los pasos que componen el proceso son:

1. Aceptación y registro de la incidencia.
2. Clasificación y soporte inicial.
3. Comparación con registro de errores conocidos.
4. Investigación y diagnóstico.
5. Resolución y recuperación.
6. Cierre.
7. Seguimiento y monitorización.

1.3. Herramientas de monitorización de equipos para la localización y notificación de incidencias

Hoy en día tenemos tal dependencia de las redes, que en muchas ocasiones, no podemos permitir esta o algunos de sus servicios deje de funcionar, aunque sea por poco tiempo.

Poder controlar el funcionamiento en tiempo real, tener alertas que avisen siempre que surja una incidencia, puede ser de muy importante.

IpMonitor es una herramienta específica para estas cuestiones. Entre sus funciones más importantes está la de poder controlar el funcionamiento de la red vía web, y el emitir alertas de todos aquellos aspectos consideremos que pueden ser un malfuncionamiento de la red.

IpMonitor es una aplicación que permite monitorear si un equipo se encuentra activo en la red, es decir, si le ha sido asignada una dirección IP para poder compartir y disponer de recursos.

Controla el estado de un equipo mediante avisos, ya sea mensajes a un correo electrónico o alertas que emita, permite saber si una IP ha dejado de funcionar, además de registrar el estado completo del equipo.

También informa sobre las características de red de cualquier equipo y que tenga una dirección IP Dinámica. Hace seguimiento continuo a los equipos dentro de la red, si por ejemplo el equipo al que se quiere acceder para utilizar algunos recursos o archivos de ella, ha cambiado su dirección IP, este software nos permite saber cuál es ahora su IP y poder acceder a través de la nueva dirección IP.

Instalación:

Para la instalación se siguen los siguientes pasos:

1. Se hace clic al instalador (setup), aparecerá una ventana que indica que el sistema se está preparando para la instalación.
2. Luego aparecerá de bienvenida a la instalación.
3. Se verá después la ventana de aceptación de términos licencia.
4. Aparecerá entonces una ventana que pide los datos de usuario y organización.
5. A continuación se abrirá una ventana que pide la ubicación (carpeta) donde quiere se instale el programa.
6. Inmediatamente se abrirá una ventana que le dará la opción de comenzar a instalar, de volver al comienzo de la instalación para escoger una opción distinta o simplemente cancelar la instalación.
7. Ahora comenzará la instalación.
8. Tras lo anterior, aparecerá una ventana que indica que la instalación ha finalizado.
9. Una vez finalizada la instalación se abrirá una ventana para configurar las opciones y herramientas de IP Monitor.

Como podemos ver es la típica instalación de software, que no tiene mayor dificultad en su ejecución.

Monitoreo:

ipMonitor incluye un conjunto de monitores que se utilizan para ver los usos, el equipo de infraestructura y los servicios esenciales :

Los monitores garantizan la calidad para realizar el análisis de resultados de pruebas para los usos críticos tales como servidores SQL, soluciones comerciales y usos dinámicos servidores.

El SNMP y los monitores proporcionan los métodos estándares de la industria para supervisar los dispositivos tales como rebajadoras, interruptores y balanceadores de la carga.

Supervisa los aspectos claves de prueba de los sistemas operativos de Windows.

Los monitores Uptime prueban disponibilidad de los protocolos basados en la capa TCP/IP tales como HTTP, HTTPS, SNMP, etc.

Un monitor es un proceso de fondo que comprueba continuamente un recurso en intervalos sincronizados.

Los métodos de prueba dependen de las capacidades del monitor, y los “parámetros de la prueba” que son incorporados durante la configuración del monitor. Los parámetros flexibles de la sincronización proporcionan la capacidad de intensificar o de disminuir la prueba durante cada uno de los estados operacionales de un monitor.

Cada vez que una prueba del monitor falla, se incrementa la cuenta secuencial del fallo y se comprueba número configurado de los fallos permitidos antes de producir una alerta. Una prueba sin fallos en cualquier momento reinicia la cuenta secuencial de fallos a cero.

Cuando un monitor alcanza su número máximo de los fallos de prueba, se accionará una alarma de la siguiente manera:

Se examina cada perfil del usuario para considerar si se le debe o no alertar en función de la incidencia producida.

En caso afirmativo, se comprueban si hay alarmas dentro del perfil del usuario.

Se ejecuta cualquier alarma activa.

Estado de los monitores:

Los monitores tienen cuatro estados operacionales, así como un estado deshabilitado.

Estado del Monitor	Color	Comentario
UP + LISTENING	Verde	El servidor/ dispositivo está respondiendo como lo esperado o el ipMonitor está esperando a escuchando la entrada del SNMP
WARN	Amarillo	Indica un resultado inesperado. La prueba está en marcha, pero no se ha accionado ninguna alarma.
WARN	Rojo Brillante	Se están enviando las alarmas. Un monitor progresará de un estado FAIL a un estado LOST cuando el número máximo de alarmas se ha procesado.
LOST	Rojo oscuro	El recurso supervisado continúa estando en un estado de error. Se han enviado todas las alarmas configuradas.
SUSPENDED+ MAINTENANCE	Gris	Un monitor es deshabilitado, no inicializado o en modo de mantenimiento.

Para ver el estado de un monitor:

Acceder a la lista de monitores dentro del interfaz de la configuración haciendo clic en la opción “Monitors” del menú.

Acceder a los informes en tiempo real dentro de la interfaz de reportes. Se clasifican éstos en códigos de color que determinan el estado del monitor.

Características de los monitores:

Nombres del DNS

Como regla general, el uso de los nombres del DNS se recomienda para los monitores basados TCP/IP. La sensibilidad del porcentaje de disponibilidad puede medir el tiempo más exactamente quitando las operaciones de búsqueda del DNS de la ecuación. La respuesta del monitor se debe terminar dentro de un número especificado de segundos. Esto incluye el tiempo tomado para realizar operaciones de búsqueda del DNS.

Si la red utiliza un servidor de DHCP (protocolo dinámico de la configuración del host) para asignar dinámicamente direcciones IP, incorporara una dirección IP solamente si es “reservado”, sino incorpora un Nombre de Dominio. Si un monitor se configura para utilizar una dirección IP y esa IP es asignada dinámicamente a otro recurso, el monitor podría no más supervisar con éxito el recurso objetivo.

Si se asigna un nombre DNS a un monitor, los parámetros que miden el tiempo se pueden afectar por el tiempo adicional requerido para realizar las operaciones de búsqueda de la DNS.

Aunque los parámetros de la sincronización por defecto para probar deben asumir el tiempo que se necesita para realizar las operaciones de búsqueda del DNS, el medir el tiempo es una variable necesaria a considerar cuando se utilizan los tiempos de prueba.

Dependencias

En el caso de que los recursos supervisados dependan de unos o más recursos críticos para funcionar, los grupos pueden ser creados y asignados por dependencias del monitor. Éstos evitan que las alarmas redundantes sean enviadas de cada monitor en un grupo cuando solamente una sola alarma de un monitor de la dependencia sería suficiente.

Horario de mantenimiento

Los horarios de mantenimiento permiten que se suspenda temporalmente la supervisión para los monitores o los grupos individuales de monitores durante períodos previstos de mantenimiento.

Exploración de la red:

ipMonitor, tiene una característica de la exploración de la red que incluye una lista de recursos potenciales. Los recursos son agrupados por la dirección IP o Nombre de Dominio.

ipMonitor sugiere los monitores basados en el tipo de recurso. Los resultados de exploración enumeran los recursos disponibles en la red, permitiéndonos elegir entre monitores individuales, o agregan grupos enteros de monitores a tu instalación del ipMonitor.

Las características incluyen:

El control sobre el ipMonitor de los métodos de descubrimiento utiliza realizar la exploración de la red.

Control sobre la gama de las direcciones del IP que serán exploradas.

La capacidad de agregar puertos no estándar a la lista de los puertos que el ipMonitor sondeará.

La capacidad de volver a los resultados de exploraciones depositadas y de agregar monitores y a grupos durante un tiempo.

Software de monitorización de switches usando SNMP. A media página sale un esquema de funcionamiento de SNMP en 4 capas: <https://www.manageengine.com/es/network-monitoring/what-is-snmp.html>

RESUMEN

Niveles Funcionales: Acceso, Troncal de Transporte y Nivel de distribución

Videos: Splitter, ONT y OLT, ROM, Cajas de Registro (contenedor destinado a ocultar instalaciones eléctricas o empalmes de cables, segmentar redes), Cajas de Empalme.

Conmutadores: Conmutación de Circuitos y de Paquetes: Orientado a la Conexión y orientado a la no-conexión

Multiplexación: División de Frecuencia (FDM), de Tiempo (TDM) de longitud de onda (WDM)

Redes guiadas:

- Cobre xDSL:
 - o ADSL (Asymmetric Digital Subscriber Line)
 - o Dos canales de datos de ancho de banda asimétricos (uno de subida y otro de bajada)
 - o Distancia máxima de 6Km
 - o Splitter: Separar bandas de trabajo y Eliminar interferencias. Con el Splitter separamos los canales de voz (300 Hz hasta los 3.400 Hz) y de datos (24 KHz hasta los 1,1 Mhz)
 - Redes inalámbricas:
- Problemas:
- Necesitas Visibilidad directa
 - Interferencias por otras señales, por tormentas... muy sensible a la orografía del terreno

Redes de Fibra:

- ⇒ HFC: Fibra en la operador y coaxial en el hogar
 - o Primer Nivel: topología estrella del Nodo central a los nodos secundarios conversión óptico-coaxial
 - o Segundo Nivel: topología bus del nodo secundario al abonado con coaxial
 - Nivel Transporte: backbone – anillo fibra primaria a anillos fibra secundaria
 - Nivel distribución: viene del nodo secundario óptico y en bus hasta el abonado
 - ⇒ Tecnología PON: Red Óptica Pasiva, sin elementos activos (amplificadores)
 - o OLT: ubicado en la central.
- Funciones:
- Gestionar y enrutar el tráfico de las ONTs.
 - Conectar la red PON con otras redes
- Permite conectividad con:
- La red telefónica básica. (voz)
 - Proveedores ISP para datos a través de un Gateway.
 - Proveedores de Vídeo a través de un Gateway
- o Splitter: entre el abonado y la central
 - o ONT: ubicado en el domicilio del abonado
 - o Utiliza multiplexación por longitud de onda
 - o Estándares: APON, BPON, **GPON**, EPON, 10GPOM

Práctica: Error de conexión Internet/Wifi

Parámetros importantes que tenemos que medir cuando realizamos cualquier tipo de cableado:

- Continuidad, Mapeado de los Hilos, Resistencia, Longitud, Atenuación por inserción-conectores, Diafonía

Instalaciones eléctricas Dedicadas: SAIs: Tipos: Off-line y On-Line

Alimentación Eléctrica por Ethernet (PoE) Usa 2 de hilos (de los 8) que no transportan datos

Otras Funciones que tiene el Router:

- Pasar datos d: Fibra ⇔ ADSL ⇔ HFC
- Parámetros de configuración de red (usando el servicio DHCP)
- Cortafuegos
- Traduciendo direcciones de red con NAT
- Proxy: Intermediario entre un navegador web e Internet
- Wi-Fi

- Redireccionando puertos
- Balanceo de carga/tráfico.
- Gestión de conexiones VPN

Funciones de un Firewall:

- Traducción de direcciones
- Protección frente a virus
- Auditoría
- Gestión de actividad

Asignación de direcciones **IP**: Automática con BOOTP o DHCP (el mas usado) y Manual (escribiéndola)

Configuración IP: Dirección IP, Mascara de Subred, Puerta de Enlace y los servidores DNS

ISP (Internet Service Provider) Proveedor de Servicios de Internet

Sniffer es una denominación aceptada para aquellas herramientas cuya función principal es monitorizar y analizar tráfico, o sea, examinar paquetes, protocolos y tramas enviadas a través de la red.