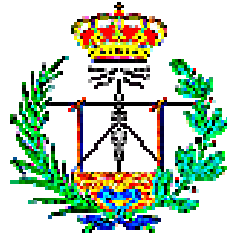


UNIVERSIDAD DE SEVILLA
ESCUELA SUPERIOR DE INGENIEROS
INGENIEROS DE TELECOMUNICACIÓN



DEPARTAMENTO DE INGENIERÍA ELECTRÓNICA
ÁREA DE TEORÍA DE LA SEÑAL Y COMUNICACIONES

PROYECTO FIN DE CARRERA:

**ESTUDIO DE TÉCNICAS DE
IGUALACIÓN PARA VDSL**

AUTOR: MANUEL GARCÍA GIL
DIRECTOR: FRANCISCO J. SIMOIS TIRADO

JUNIO DE 2001

INDICE

CAPÍTULO 1: INTRODUCCIÓN	5
1.1 – OBJETIVOS DEL PROYECTO	5
1.2 – ESTRUCTURA DE LA MEMORIA.....	6
CAPÍTULO 2: INTRODUCCIÓN A LA TECNOLOGÍA VDSL	9
2.1 – GENERALIDADES ACERCA DE LAS TECNOLOGÍAS SOBRE COBRE	9
2.2 – LA ARQUITECTURA VDSL.....	20
CAPÍTULO 3: ESTUDIO DEL MEDIO DE TRANSMISIÓN: EL BUCLE DE ABONADO.....	25
3.1 – INTRODUCCIÓN.....	25
3.2 – PRIMER ACERCAMIENTO AL PAR DE COBRE.....	27
3.3 – MODELO EN ESCALERA DE UN CABLE UTP	33
3.4 – CROSSTALK	36
3.5 – OTRAS FUENTES DE RUIDO EN EL BUCLE	42
3.6 – OBTENCIÓN DE UN MODELO DE CANAL	43
CAPÍTULO 4: INTRODUCCIÓN A LA MODULACIÓN MULTIPORTADORA. EL SISTEMA DMT	45
4.1 – INTRODUCCIÓN.....	45
4.2 – DESCRIPCIÓN DE UN SISTEMA BÁSICO	47
4.3 – INDEPENDENCIA ENTRE SUBCANALES.....	49
4.4 – MEDIDA DE PRESTACIONES DE UN SISTEMA DE MODULACIÓN MULTICANAL	51
4.5 – LA CARGA DE LOS SUBCANALES	59
4.6 – EL SISTEMA DMT	68
CAPÍTULO 5: IGUALACIÓN EN DMT APLICADA A VDSL. EL IGUALADOR EN EL DOMINIO DEL TIEMPO (TEQ) Y EL IGUALADOR EN EL DOMINIO DE LA FRECUENCIA (FEQ).....	87
5.1 – INTRODUCCIÓN A LAS TÉCNICAS DE IGUALACIÓN.....	87
5.2 – IGUALACIÓN EN EL DOMINIO DEL TIEMPO	89
5.3 – INICIALIZACIÓN DEL IGUALADOR TEQ: EL MMSE-TEQ.....	91
5.4 – IGUALACIÓN EN EL DOMINIO DE LA FRECUENCIA	94
5.5 – EL IGUALADOR FEQ	95
5.6 – INICIALIZACIÓN DEL FEQ.....	103
5.7 – COMENTARIOS ACERCA DE LOS IGUALADORES	105

CAPÍTULO 6: TÉCNICAS ESTADÍSTICAS PARA LA ESTIMACIÓN DE LA PROBABILIDAD DE ERROR EN UN SISTEMA DIGITAL.....	109
6.1 – INTRODUCCIÓN.....	109
6.2 – EL MÉTODO DE MONTE CARLO.....	110
6.3 – EL MÉTODO DE EXTRAPOLACIÓN DE COLA (“TAIL EXTRAPOLATION”).....	112
CAPÍTULO 7: SIMULACIONES.....	117
7.1 – INTRODUCCIÓN.....	117
7.2 – CONFIGURACIÓN DE LA TRANSMISIÓN.....	118
7.3 – FASE DE ENTRENAMIENTO. ASIGNACIÓN DE BITS A SUBCANALES	121
7.4 – GENERACIÓN DE SÍMBOLOS DMT Y TRANSMISIÓN	123
7.5 – EFECTO DEL CANAL SOBRE LA SEÑAL	127
7.6 – RECEPCIÓN DE SÍMBOLOS DMT Y DETECCIÓN.....	130
7.7 – UTILIZACIÓN DE UN IGUALADOR TEQ	136
7.8 – UTILIZACIÓN DE UN IGUALADOR FEQ.....	145
7.9 – SIMULACIÓN DE UN SISTEMA CON PROBABILIDAD DE ERROR REAL	150
CAPÍTULO 8: CONCLUSIONES Y LÍNEAS FUTURAS DE INVESTIGACIÓN..	155
8.1 – CONCLUSIONES	155
8.2 – LÍNEAS FUTURAS DE INVESTIGACIÓN	158
GLOSARIO	159
REFERENCIAS.....	161
ANEXO: PROGRAMAS PRINCIPALES.....	163

1.1 – OBJETIVOS DEL PROYECTO

En el presente proyecto fin de carrera se pretenden analizar las prestaciones de ciertos esquemas de igualación empleados en sistemas VDSL. Para ello, se van a ver todas las fases y componentes de una transmisión VDSL, haciendo hincapié en las principales características que la diferencian del resto de la familia. Se estudiarán las técnicas involucradas en la familia xDSL, en especial la modulación multiportadora (utilizaremos el sistema de transmisión DMT).

Será necesario hacer un estudio del bucle de abonado, que es el medio de transmisión utilizado por esta tecnología.

Una vez realizados estos estudios previos, se procederá al diseño de unas rutinas de MATLAB que simulen el comportamiento de los sistemas estudiados, particularizando los parámetros para el caso de VDSL. Aunque VDSL utiliza, como otros miembros de la familia xDSL, el mismo medio físico de transmisión (el par trenzado de cobre) y la misma técnica de modulación (modulación multiportadora), nos encontraremos con unas peculiaridades que harán necesaria la utilización de técnicas propias.

En la búsqueda de soluciones al problema de maximizar la tasa de transmisión, llegaremos a la implementación de dos tipos de igualadores: uno funciona en el dominio del tiempo y el otro en el dominio de la frecuencia.

Implementado ya el sistema de comunicaciones, el siguiente paso será medir las prestaciones que el mismo nos ofrece, con y sin las posibilidades de mejora que nos ofrecen los igualadores. Ya que se trata de un enlace digital de datos, nuestra medida de la calidad aquí va a ser la probabilidad de error de símbolo.

Por otro lado, y dada la situación y desarrollo de las comunicaciones hoy en día, los protocolos y servicios actuales (sobre todo los de banda ancha) requieren enlaces de datos con probabilidades de error cada vez menores. Esto, que en principio supone un obstáculo, en tanto en cuanto nos alarga en exceso los tiempos de simulación, se resuelve mediante el empleo de técnicas de estadística matemática que permiten detectar las mismas probabilidades de error pero con tiempos de simulación considerablemente menores.

1.2 - ESTRUCTURA DE LA MEMORIA

La memoria del proyecto está estructurada en capítulos, el primero de los cuales es esta introducción. A continuación comentaremos brevemente el contenido de los siguientes.

En el capítulo 2 se hace una introducción a VDSL, partiendo de los conceptos necesarios para la transmisión digital de datos a alta velocidad sobre líneas de cobre (DSL).

En el capítulo 3 se hace un estudio del bucle de abonado. Este estudio es de capital importancia, puesto que las características del medio de transmisión condicionarán el tipo de modulación a utilizar en las transmisiones. Se propone un modelo, tanto en lo que se refiere a la atenuación que provoca como al ruido que añade.

El siguiente capítulo está dedicado a la modulación multiportadora. Frente al uso alternativo de una modulación tradicional de portadora única, se justifica el uso de la modulación multiportadora basándonos en las características adversas del canal de transmisión en cuanto a que la distorsión que introduce no es constante con la frecuencia. Nos centraremos especialmente en el sistema de modulación DMT, que es el utilizado en los estándares VDSL, y, por tanto, el que emplearemos en nuestras simulaciones. Una vez que se haya estudiado DMT, se entenderá la conveniencia de emplear

igualación, con la finalidad de aumentar el rendimiento del enlace sin empeorar su calidad.

En el capítulo 5 se analizan los sistemas de igualación para DMT. En concreto, se estudian dos esquemas de igualación, uno que funciona en el dominio del tiempo y otro en el dominio de la frecuencia. Se estudia la problemática de su inclusión en un escenario VDSL, ya que su uso aumenta la complejidad computacional del proceso: por un lado, hay que considerar la fase de inicialización del igualador, y por otro, hay que considerar el procesamiento de la señal por este nuevo bloque.

El capítulo 7 se dedica al análisis de las simulaciones realizadas y a la obtención de conclusiones de las mismas. Pero antes, se incluye un pequeño capítulo (capítulo 6) dedicado a las técnicas estadísticas para la simulación de sistemas de comunicación digitales. Se estudian aquí el método de Monte Carlo y la técnica de “extrapolación de cola”. Como veremos, el método de Monte Carlo es muy costoso de implementar (en cuanto al tiempo de simulación necesario) cuando se quieren detectar probabilidades de error muy pequeñas. En estos casos, es conveniente utilizar la otra técnica anteriormente mencionada.

El último capítulo, está dedicado a las conclusiones y líneas futuras de investigación. Se hace un análisis global del estudio llevado a cabo en el proyecto.

2.1 - GENERALIDADES ACERCA DE LAS TECNOLOGÍAS SOBRE COBRE

Utilizar la tecnología VDSL (Very high rate Digital Subscriber Line) es una de las posibilidades a la hora de proporcionar servicios de banda ancha al usuario final. El objetivo de este capítulo es describir el escenario en el que se implanta dicha tecnología.: la red de acceso de un operador de telefonía. Se describen las diferentes formas que históricamente ha habido para transmitir datos por la red telefónica, haciendo hincapié en las modernas tecnologías xDSL (cuyo objetivo es aumentar las velocidades de transmisión), y en particular, VDSL.

2.1.1 – Descripción de la situación y objetivos buscados

Estamos bastante acostumbrados al uso de los modems de banda vocal y a sus limitaciones. Actualmente, estos modems (analógicos) pueden transmitir hasta aproximadamente 56Kbps sobre una línea telefónica estándar. Sin embargo, hace 20 años el límite práctico era de tan sólo 1.2Kbps. Nadie piensa que se pueda mejorar mucho más la velocidad alcanzada hoy en día; al fin y al cabo, el ancho de banda utilizado por este tipo de dispositivos no supera los 3.3 KHz [10]. Y hoy es fácil conseguir uno de estos dispositivos a un precio asequible. Son el resultado de los avances en algoritmos, tratamiento digital de señales y tecnología de semiconductores.

Los modems de banda vocal se instalan en los locales del abonado, en el extremo final de líneas de transmisión para voz, y transmiten señales que se propagan por la red conmutada sin apenas degradación. Esto ocurre porque la red trata a estas señales como si realmente fueran señales vocales. Esta es hoy en día su gran ventaja: un modem puede conectarse rápidamente al extremo de cualquier línea telefónica existente; y en el mundo existen unos 600 millones de tales líneas [10]. Sin embargo, las velocidades de transmisión que pueden alcanzar son bastante lentas comparadas con las velocidades de los equipos actuales.

Las limitaciones en cuanto a velocidad de los modems de banda vocal no provienen de la línea telefónica en sí, sino de cómo los equipos de la red telefónica tratan a las señales. Existen filtros que limitan el ancho de banda de la señal producida por los modems a menos de 4 KHz (en la central de conmutación, la señal vocal es digitalizada mediante un muestreo a 8 KHz, por lo que previamente debemos truncar su espectro en una frecuencia inferior a los 4 KHz). Si eliminamos dichos filtros, nos encontraremos con que podemos transmitir frecuencias en la región de los MHz, aunque con una atenuación bastante considerable. De hecho, es esta atenuación (que aumenta con la longitud de la línea y con la frecuencia) la que va a limitar las restricciones en cuanto a velocidad. En la figura 2.1 mostramos cómo se conecta a la línea un módem tradicional.

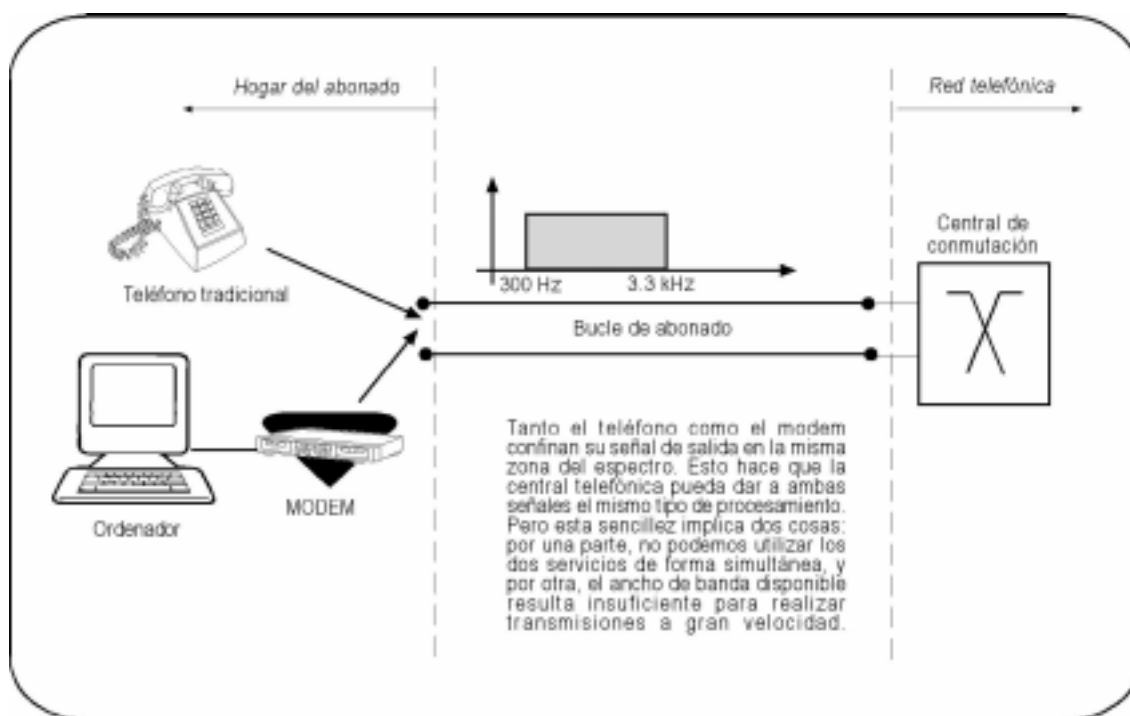


Figura 2.1 : Funcionamiento de un módem tradicional

Las configuraciones de los despliegues de las líneas de abonado pueden ser muy diferentes en distintos lugares del mundo. En algunos países la longitud media del bucle de abonado es de unos 5.5 Km. En otros, como Estados Unidos, la media es superior [14]. Si la una línea es muy larga, se hace necesario el uso de bobinas de carga y se inutiliza la línea para el uso de los accesos DSL (en el capítulo 3 se estudiarán las bobinas de carga, y se

entenderá el porqué de esta afirmación). En la figura 2.2 mostramos la estructura de una red telefónica, indicando cómo se integran en ella los bucles de abonado.

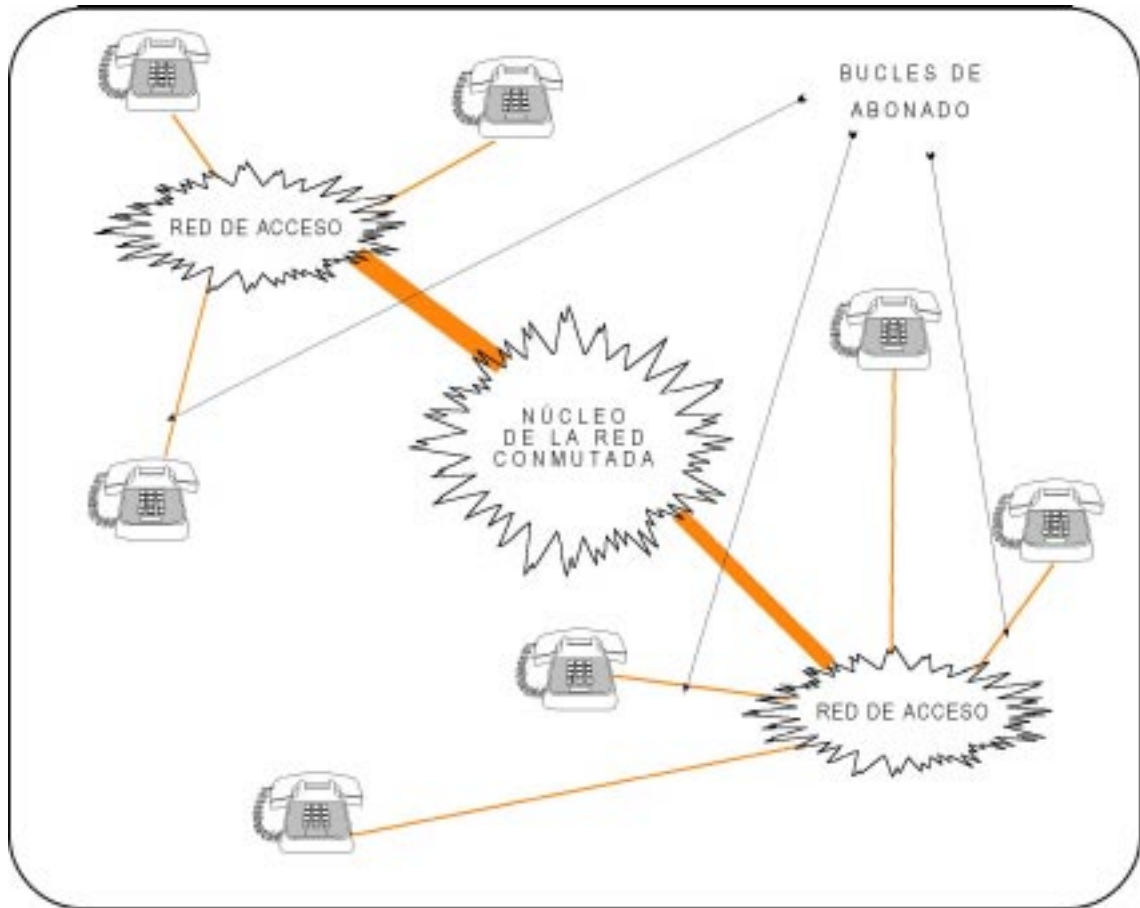


Figura 2.2 : Estructura de una red telefónica.

Una condición importante para que la planta de cableado actual de las compañías telefónicas pueda aprovecharse para implementar sistemas DSL es que la longitud del bucle de abonado sea lo más corta posible. La mayor parte de las compañías telefónicas del mundo han utilizado programas para reducir al máximo la longitud del bucle de abonado, pero maximizando la cobertura de las centrales existentes. Una de las técnicas típicas es la instalación de nodos remotos de acceso, los cuales son alimentados (conectados con la central) por líneas T1/E1 (hoy en día pueden usarse para este menester enlaces HDSL, capaces de transportar señales E1/T1).

Por supuesto, para que sea posible utilizar la planta de cableado como soporte para los servicios DSL, hay que diseñar modems específicos para

este tipo de aplicaciones, y colocarlos tanto en las centrales como en los locales del abonado: ahora la red no va a tratar sus señales como señales de voz. Evidentemente, las investigaciones necesarias para posibilitar la transmisión de datos a alta velocidad sobre líneas telefónicas de cobre [10] comenzaron hace muchos años.

2.1.2 – La línea digital de abonado (DSL)

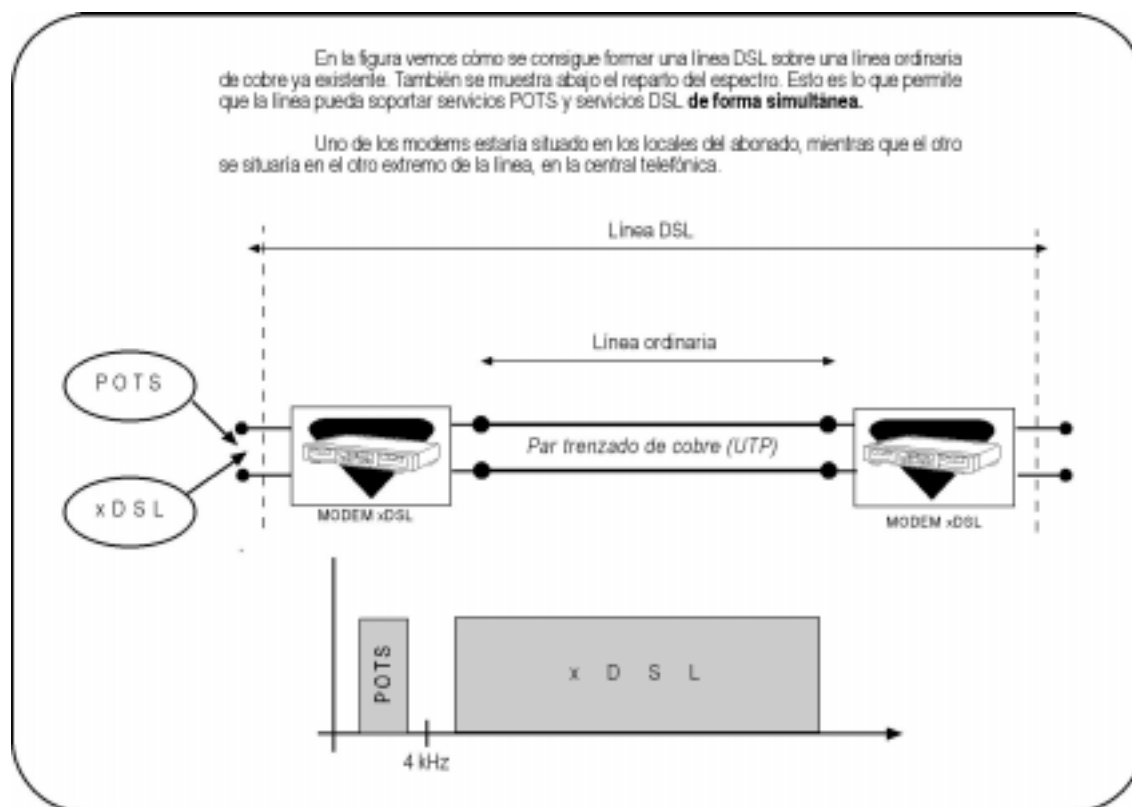


Figura 2.3 : Formación de una línea DSL sobre una línea ya existente, mediante colocación de modems DSL

El acrónimo DSL (Digital Subscriber Line, o línea digital de abonado) tiene su origen en la compañía Bellcore, así que son ellos los responsables de la confusión existente entre una línea y sus modems. En general, cuando nos referimos a DSL, nos referimos a un modem DSL (o en todo caso, a un par de modems), y no a una línea*. Es cierto que si conectamos dos modems DSL a una línea (cada uno en un extremo) conseguimos una línea digital de abonado (DSL), pero cuando la compañía telefónica dice que instala líneas DSL (ADSL,

* No se instala una nueva línea, sino que se aprovecha la ya existente (par trenzado de cobre). La instalación de los modems DSL hace que la línea pueda soportar una velocidad mayor.

HDSL, VDSL, etc) lo que hace en realidad es instalar modems DSL, no nuevas líneas. Véase la figura 1.3 para clarificar este hecho.

Uno de los primeros servicios que ha utilizado DSL es el acceso básico RDSI. Un enlace de este tipo funciona en modo duplex transmitiendo datos a 160 Kb/s sobre líneas de cobre de hasta 5.5 Km. Esta tasa de datos se multiplexa/desmultiplexa formando dos canales B (cada uno de 64 Kb/s) y un canal D (16 Kb/s); el resto de capacidad se utiliza para transportar taras, insertadas por los equipos, para el mantenimiento del enlace.

Otra de las posibilidades que se nos presenta consiste en utilizar la línea de cobre existente, y de forma simultánea, para los nuevos servicios DSL y para el POTS (Plain Old Telephone System, o telefonía básica analógica tradicional). También podemos conseguir, utilizando modems DSL, dos líneas POTS a partir de una sola línea POTS, obviando la instalación de una segunda línea física. En definitiva, el desarrollo de tecnologías que permitan utilizar servicios DSL nos va a abrir un montón de posibilidades de mejora en la calidad de los servicios sin el coste que supone cambiar la instalación física de la red de acceso.

2.1.3 – T1/E1

Al principio de los 60s, los ingenieros de la Bell Labs crearon un sistema multiplex para voz que digitalizaba la señal vocal analógica para formar un flujo binario de 64 Kb/s (muestreando a 8 KHz y expresando cada muestra con una palabra de 8 bits). Luego organizaron 24 de estos flujos para formar una trama de 193 bits, con convenciones sobre cómo y donde mapear cada uno de los 24 flujos en ella. Esta señal, cuya tasa binaria resulta ser de 1.544 Mb/s, fue bautizada con el nombre de DS1; pero también se la conoce con el sinónimo más coloquial de T1.

En Europa, y en el CCITT (hoy ITU-T), se definió una trama parecida, pero agrupando 30 canales vocales en lugar de 24. El resultado fue un flujo binario de 2.048 Mb/s dividido en tramas de 256 bits. A esta señal se le dio el nombre de E1.

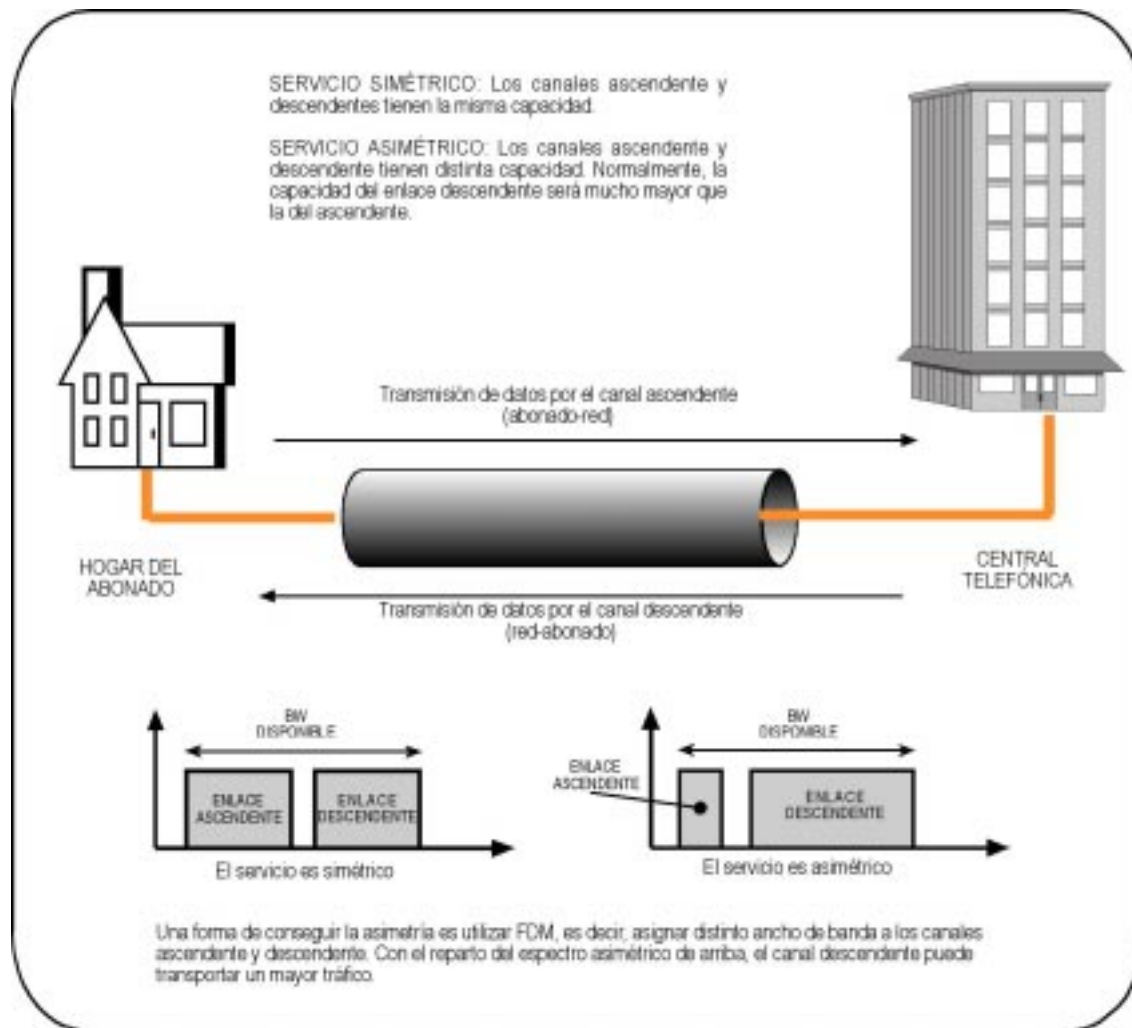


Figura 2.4 : Diferencia entre servicios simétricos y asimétricos

Hasta hace poco, los sistemas E1 o T1 se implementaban sobre líneas de cobre, usando transceptores ordinarios y un código de línea AMI (Alternate Mark Inversion, inversión de marcas alternadas), el cual tiene la propiedad de transportar muy bien el reloj (información de temporización). Este código AMI requiere repetidores de señal cada cierta distancia, y consume un ancho de banda de 1.5 MHz, con un pico de potencia en 750 kHz (en el estandar americano). Esto puede parecer un método poco elegante, pero ha funcionado durante muchos años, y hoy en día existen cientos de miles de estas líneas en el mundo.

Originalmente, las compañías telefónicas utilizaban los enlaces E1/T1 para la transmisión de datos entre sus centrales en el núcleo de la red conmutada (no en la red de acceso formada por los bucles de abonado).

Posteriormente, las compañías comenzaron a ofrecer (tarificando) este servicio, y se utilizó en redes privadas, conectando PBXs (centralitas) y multiplexores T1/E1 juntos sobre la WAN (red de área extensa). Hoy, los circuitos T1/E1 se usan para muchas aplicaciones más, tales como la conexión de routers de Internet, transporte de tráfico hasta antenas celulares, o conectando servidores multimedia a una central.

Una aplicación de estos circuitos que está cobrando cada vez más popularidad es la llamada "planta de alimentación", la cual podemos definir como la sección de una red telefónica que se encarga de distribuir las señales desde las centrales principales hasta los nodos de acceso, desde los cuales parten ya los bucles de abonado (líneas de cobre individuales).

Sin embargo, el sistema E1/T1 no es muy adecuado para realizar la conexión directamente a los usuarios particulares. El código AMI consume demasiado ancho de banda y provoca muchas interferencias (crosstalk): Esto hace que las compañías sólo utilicen un circuito E1/T1 en cada grupo de 50. Así que una oferta de este tipo para los usuarios pasa necesariamente por cambiar la mayor parte de la planta de cableado actual. Por otro lado, la mayoría de las aplicaciones multimedia que funcionan hoy en día demandan un flujo de datos asimétrico (los enlaces ascendente y descendente funcionan con distintas tasas binarias, véase la figura 2.4), y por lo general, demandan más velocidad que la que puede ofrecerse con un enlace E1/T1. El transporte hasta los hogares de las señales de estos servicios de banda ancha será llevado a cabo por enlaces ADSL o VDSL.

2.1.4 – HDSL – Línea digital de abonado de alta velocidad

HDSL (High data rate Digital Subscriber Line) es una nueva forma de conseguir que un enlace E1/T1 sea transmitido sobre pares trenzados de cobre. Usa menos ancho de banda que el código AMI y no requiere el uso de repetidores. HDSL soporta flujos de 1.544 Mb/s o de 2.048 Mb/s ocupando anchos de banda desde los 80 KHz a los 240 KHz (depende de la técnica específica que se emplee). Sin embargo, requiere el uso de dos pares de cobre para transportar un T1 y tres para un E1.

La mayor parte de los circuitos HDSL se situarán en la "planta de alimentación", por lo cual podemos decir que se utilizan para conectar a los abonados en cierta forma, pero no en el sentido de ofrecer un servicio individual con enlaces HDSL a cada uno de ellos (a cada abonado le llega tan sólo un par de cobre, por lo que no podemos utilizar esta técnica).

Las aplicaciones típicas incluyen conexiones de PBXs en redes, entre estaciones base de telefonía celular, y en redes privadas de datos. Como HDSL es la tecnología más madura entre todas las xDSL con tasas del orden del megabit, ha sido muy utilizada durante la primera etapa para aplicaciones de internet y acceso remoto a LANs, pero está siendo sustituida por ADSL y SDSL de forma paulatina.

2.1.5 – SDSL – Línea digital de abonado con línea única

SDSL (Single line Digital Subscriber Line) es una versión de HDSL pero utilizando un solo par de cobre. Se trata de transmitir una señal E1/T1 sobre un cable de tipo UTP, en la mayor parte de los casos, de forma simultánea con el servicio POTS. Como vemos, SDSL presenta una ventaja importante sobre HDSL, dado que, al utilizar un único par de cobre, ya podemos ofrecer de forma sencilla servicios individuales para el abonado, que cuenta con una línea telefónica individual hasta su domicilio.

SDSL es útil en las aplicaciones que se ajusten mejor a un modo de funcionamiento simétrico, tales como conexión de servidores y acceso remoto a redes LAN; es, por tanto, una opción que complementa a la solución ADSL (funcionamiento asimétrico).

2.1.6 – ADSL – Línea digital de abonado con funcionamiento asimétrico

ADSL siguió los pasos de HDSL, pero es un servicio totalmente enfocado a su uso en el bucle de abonado, el último tramo en un circuito de conexión telefónico. Como su nombre indica, se trata de transmitir un flujo de

datos asimétrico, asignando mucha más capacidad al enlace descendente (sentido red-abonado) que al ascendente (sentido abonado-red). La razón por la que se utiliza este modo asimétrico está más relacionado con la configuración de la planta de cableado que con las tecnologías de transmisión. Los cables de telefonía se envuelven (agrupando unos 50 pares) juntos en cables más gruesos. A veces, estos grupos se agrupan entre ellos, con lo que se forman grupos que pueden ser de miles de pares, especialmente si nos fijamos en los cables según salen de las centrales de conmutación. Un par de abonado está formado por multitud de secciones empalmadas (según la compañía Bellcore, en Estados Unidos cada bucle de abonado tiene una media de 22 empalmes). La idea de trenzar sobre sí mismos a los pares de cobre se le ocurrió a Alexander Bell, con objeto de minimizar el efecto indeseable de las interferencias producidas por los acoplamientos inductivos o capacitivos que tienen lugar cuando los pares se agrupan en paquetes de 50 o más para envolverlos todos en un cable mayor. Sin embargo, el hecho de trenzar los pares de cobre no soluciona totalmente el problema. Las señales se acoplan, y el acoplamiento es tanto más intenso al incrementa la frecuencia o la longitud de la línea. Y además, resulta que si se intentan transmitir señales simétricas por muchos pares de un mismo grupo, se limita mucho la tasa máxima de datos alcanzable, debido a los acoplamientos. Observemos la figura 2.5 para ver con claridad el agrupamiento de los pares de abonado.

Pero afortunadamente, la mayoría de los servicios susceptibles de implementarse con ADSL demandan un uso asimétrico del canal. VoD (video on demand), tele-compras, acceso a internet, acceso a LANs remotas, accesos multimedia, servicios especiales de conexión entre PCs, etc....Todas estas aplicaciones se caracterizan por precisar una gran capacidad de canal en el sentido red-usuario y una baja capacidad en el sentido usuario-red. Por ejemplo, una transmisión de video tipo MPEG con controles VCR simulados requiere unos 1.5 o 3 Mb/s en el enlace descendente, pero la aplicación funcionará correctamente con tan solo 64 Kb/s (o incluso con 16 Kb/s) en el ascendente. Los protocolos IP (para internet) o los de acceso a las LANs quizás requieran una mayor capacidad para el enlace ascendente, pero una relación de 10:1 en las velocidades de transmisión de información entre los dos sentidos es a menudo suficiente y no degrada la calidad del servicio.

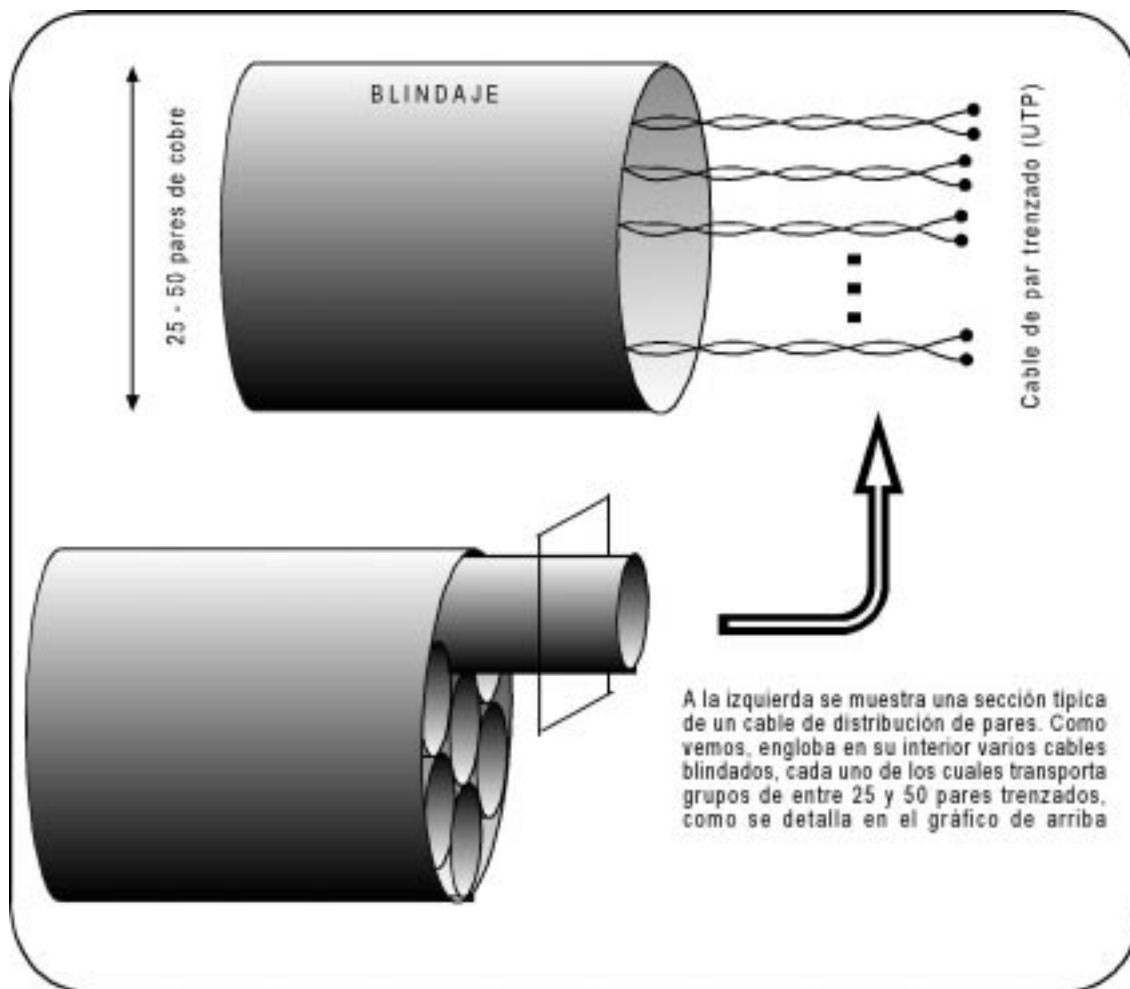


Figura 2.5 : Detalle del agrupamiento de los pares de abonado

En cuanto a las velocidades del enlace ascendente, en ADSL suelen situarse entre 16 Kb/s y 640 Kb/s. Los equipos fabricados para trabajar con ADSL suelen venir preconfigurados para funcionar en un rango variado de velocidades; pero siempre operan en una banda de frecuencia por encima de la banda POTS, y sin alterar para nada este último servicio, incluso en los casos en los que falle el módem ADSL.

Como ADSL puede transmitir video digital comprimido, entre otras cosas, incluye mecanismos de corrección de errores, con idea de reducir el efecto del ruido impulsivo en las señales de video. Estos mecanismos introducen en el enlace un retardo de 20 ms [10], lo cual es demasiado para aplicaciones LAN y IP. Por lo tanto, para un funcionamiento óptimo del

sistema, conviene que el módem ADSL conozca el tipo de señal que va a procesar, para saber si tiene que aplicar control de errores o no.*

Y más aún, va a ser utilizado para conmutación de circuitos (su uso actual), conmutación de paquetes (por ejemplo, en su uso en un router IP), y eventualmente, para conmutación de células ATM. ADSL va a conectar ordenadores personales y televisores al mismo tiempo. Todas estas características van a provocar que los protocolos utilizados para llevar a cabo todas los servicios requeridos se compliquen mucho, llevando al módem ADSL a trabajar mucho más allá del concepto de la transmisión y recepción ordinaria de datos.

2.1.7 – VDSL. Línea digital de abonado con tasa de bits muy alta

En los orígenes de VDSL (Very high data rate Digital Subscriber Line), su denominación era la de VADSL, ya que en sus primeras manifestaciones un sistema VDSL era equivalente a un ADSL pero operando sobre bucles mucho más cortos, lo cual provocaba que la tasa binaria se incrementase considerablemente.

El objetivo marcado es conseguir una velocidad de hasta unos 52 Mb/s en el enlace descendente y unos 1.6-2.3 Mb/s en el ascendente. Y en muchos aspectos, VDSL es más simple que ADSL. La longitud más corta de las líneas relajan más las condiciones de transmisión, así que la tecnología básica del transceptor VDSL no tiene que ser tan compleja como la del ADSL, aún teniendo en cuenta que funciona a una velocidad diez veces superior. Por otra parte, VDSL está muy orientado a soportar arquitecturas ATM, por lo que no es necesario que resuelva problemas de manejo de paquetes y canalización, los cuales sí tenían que ser resueltos por los transceptores ADSL.

* En realidad, esto no es un problema de ADSL, los mecanismos de control de errores, y el retardo introducido en el enlace consecuencia de ellos, son dos características en cualquiera de las tecnologías existentes, sobre cualquier línea de transmisión, ya sea cable coaxial o UTP.

No obstante, y dado que el sistema objeto de estudio en este proyecto es VDSL, dedicaremos el siguiente apartado a realizar un estudio un poco más profundo de su arquitectura.

2.2- LA ARQUITECTURA VDSL

VDSL (Very high-speed digital subscriber line) es una tecnología de acceso de banda ancha que permite el servicio de reparto de datos hasta una velocidad de 52 Mb/s en el último tramo de la red, aprovechando las líneas estándares de telefonía que ya se encuentran desplegadas.

Soporta servicios tanto simétricos como asimétricos, y también tiene una gran capacidad para soportar los servicios demandados tanto por el público en general (que desea disfrutar de lo último en aplicaciones multimedia de banda ancha) como por empresas (que desean accesos de alta velocidad para engancharse a redes ATM, STM, IP y otros tipos de redes de transporte).

En la mayoría de las redes, el medio de transmisión utilizado para el transporte de datos es la fibra óptica. Sin embargo, el despliegue de la fibra no llega hasta los locales del abonado (FTTH, fiber to the home), sino que llega hasta la central (FTTEx, fiber to the exchange) o hasta unas unidades de reparto, más próximas a los locales del abonado (FTTCab, fiber to the cabinet; FTTC, fiber to the curb). En todos estos casos, el despliegue de fibra termina en un elemento llamado ONU (unidad de red óptica), localizado en la central de conmutación local o en el nodo de reparto. En el modo FTTEx, los locales del abonado deben estar próximos a una central de conmutación, la cual puede ofrecerles el servicio directamente; en esta situación suelen estar las empresas. Por el contrario, los clientes particulares se alojan normalmente en zonas residenciales alejadas de las centrales de conmutación, por lo que se suele desplegar fibra óptica hasta un nodo intermedio o armario ubicado cerca de cada zona residencial. Con estos dos modos de despliegue, tanto las

empresas como los clientes particulares tienen garantizada su cobertura VDSL, sin restricciones en cuanto a la distancia a la central principal.

En la figura 2.6 se muestra la arquitectura de red para un despliegue de servicios VDSL [4]. El núcleo de la red soporta varios tipos de servicios (voz, video, datos, etc.) y protocolos (ATM, STM, IP, etc.), y está conectado a la red de acceso a través de un nodo de acceso (AN) situado en la oficina central. Los enlaces de fibra conectan las OLT (optical line termination, situadas en la central principal) con las ONU (optical network unit, situadas en las centrales de conmutación o en los nodos intermedios). La unidad de terminación de red NT (network termination) proporciona la adaptación de los protocolos en los locales del abonado.

El elemento más relevante de todo el sistema (al menos, para nuestro estudio), es la VTU (VDSL transceiver unit), dado que es la unidad responsable de implementar la tecnología VDSL. En una transmisión VDSL funcionan dos VTUs:

- Una está situada en la ONU (unidad de red óptica, situada o bien en la central local o bien en el nodo intermedio o armario). Adapta los datos desde el enlace VDSL hasta la red de fibra.
- La otra está en la NT, en los locales del abonado.

Entre estas dos VTUs, el medio de transmisión que tenemos es un par trenzado de cobre, el mismo par de cobre utilizado para ofrecer los servicios de telefonía básica. Este servicio es lo que llamaremos POTS (Plain Old Telephone Service). Para compaginar el servicio POTS con los actuales servicios de bucle digital, debemos instalar tanto en la ONU como en la NT unos "splitters" (repartidores) que separen la señal correspondiente al POTS de la señal correspondiente a los nuevos servicios (en este caso, la señal VDSL).

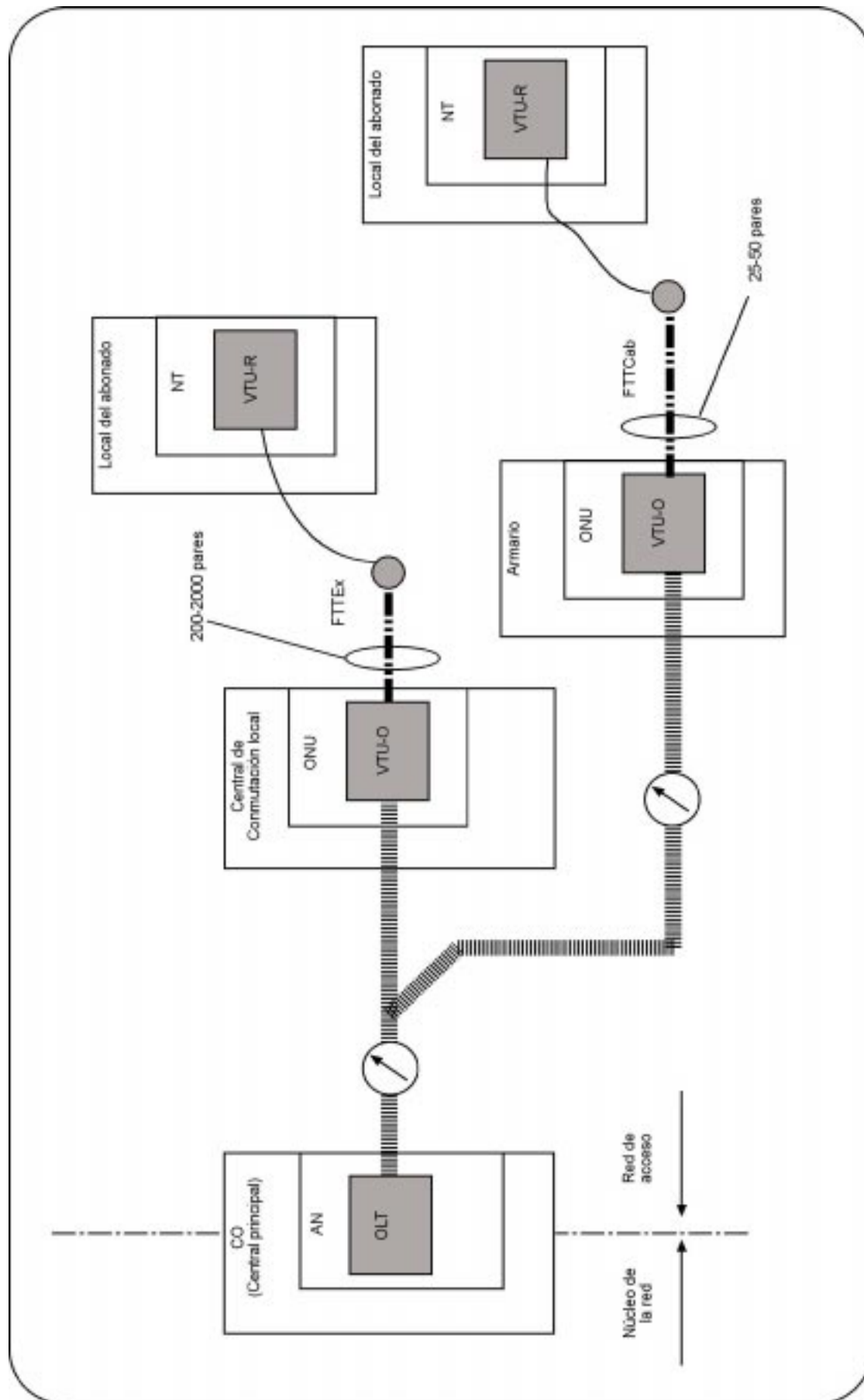


Figura 2.6 – Escenario típico en el despliegue de una red que ofrece servicios VDSL

Hemos afirmado que las dos VTUs se encuentran conectadas mediante un cable de cobre. Son necesarias algunas explicaciones relativas a este hecho.

- Desde las ONU no salen cables individuales, sino que en una ciudad se realiza el despliegue agrupando entre 25 y 50 pares en un único tubo para posteriormente, y desde armarios situados en la calle para tal fin, separarlos y distribuir cada uno hasta el hogar del abonado. Como consecuencia de esto, se producirá otro efecto que comentamos en el párrafo siguiente.
- El par de cobre puede no ser el mismo a lo largo del trayecto entre las VTUs. Es decir, podemos tener varios cables empalmados que formen uno sólo, combinando cables de diferente calibre (diámetro), cables puenteados, etc.

Estos dos efectos hay que tenerlos en cuenta a la hora de diseñar los codificadores de línea para VDSL, dado que el segundo efecto será el que determine la forma de la respuesta impulsiva del canal (par trenzado de cobre) y el primer efecto será la causa de que aparezca la limitación más importante en un enlace VDSL: las interferencias (crosstalk) procedentes de otros cables que viajan en el mismo tubo.

3.1 - INTRODUCCIÓN

Dentro del estudio sobre VDSL que se está llevando a cabo en este proyecto, en este capítulo vamos a adentrarnos un poco en el estudio del canal de comunicaciones utilizado por este tipo de sistemas de transmisión. Un conocimiento completo de las características específicas de este medio es lo que va a posibilitar el estudio y desarrollo de las técnicas DSL.

La planta de bucles de abonado que se encuentra instalada actualmente se desplegó hace muchos años, y ha sido objeto de numerosos estudios a lo largo del tiempo. Pero en realidad, las características básicas de esos bucles han cambiado muy poco, por lo que todos los estudios realizados anteriormente siguen teniendo validez y han debido ser rescatados para el desarrollo de las tecnologías DSL. Todo nuestro estudio se basa en [14].

Básicamente podemos decir que un bucle de abonado se compone de un par de hilos aislados de cobre, que conectan el equipo terminal de la central telefónica con el teléfono del abonado, situado en su domicilio. Esta elección del medio de transmisión puede parecer demasiado simple, pero cumple (y ha cumplido durante muchos años) su función: es un medio de transmisión barato (hay que tener en cuenta que para un operador de telefonía el despliegue de su red de acceso supone una gran parte del coste total de la inversión) y por otra parte, es un medio que es capaz de transportar sin problemas señales vocales con una calidad más que aceptable. Como cada bucle sólo tiene que transportar un único canal vocal (bidireccional, y en banda base), no es mucho el ancho de banda necesario en el medio de transmisión.

Sin embargo, al reutilizar esta planta de cableado ya instalada, surgieron problemas debido a la necesidad de transportar señales de un

ancho de banda considerablemente superior (un canal vocal puede ocupar sólo hasta 4 KHz de ancho de banda, y por el contrario, ahora con ADSL hay que llegar hasta 1.1 MHz o hasta los 15 MHz de VDSL). Este hecho provoca además que los fenómenos de acoplamiento entre distintos pares se hagan más evidentes (la cantidad de interferencia recibida depende del nivel de potencia que circule por el cable interferente, pero también de la frecuencia de la señal interferente).

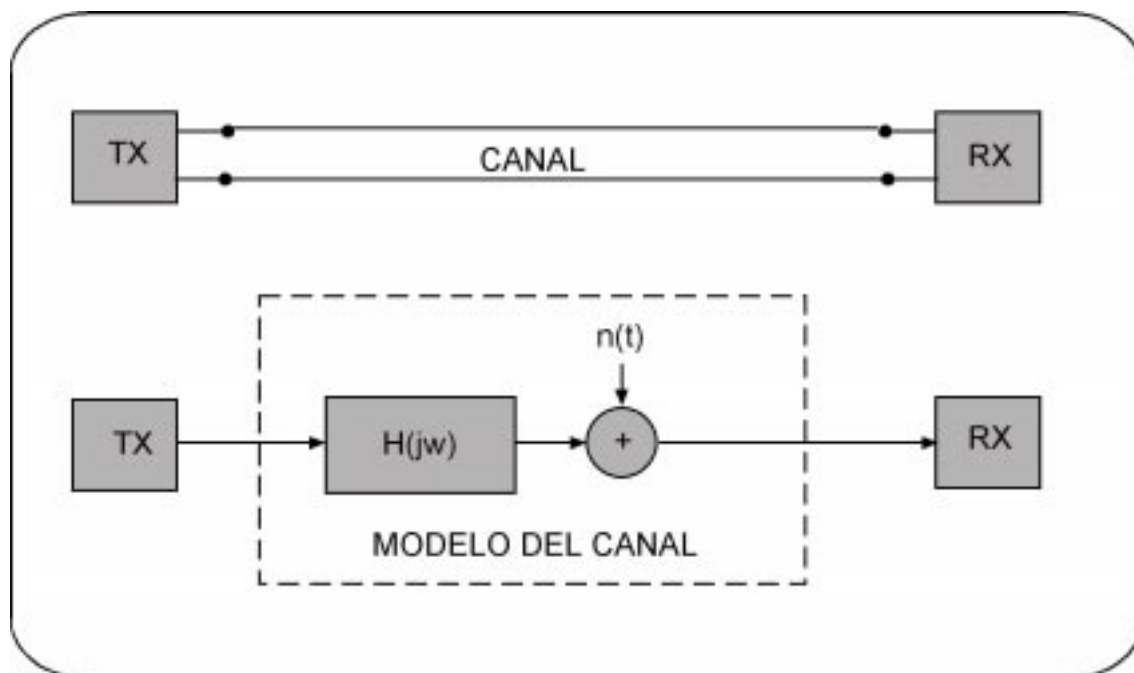


Figura 3.1 : Modelo matemático del canal

El objetivo de este capítulo es buscar un modelo matemático para el canal de comunicaciones, partiendo de las características físicas del par metálico que constituye el medio de transmisión. La idea se muestra en la figura 3.1. Para ello, debemos proporcionar una función de transferencia y un modelo para el ruido aditivo, que en nuestro caso, no va a ser un ruido blanco de fondo, sino que va a ser un ruido provocado por las interferencias procedentes de otros pares.

3.2 – PRIMER ACERCAMIENTO AL PAR DE COBRE

Como ya hemos adelantado, el medio de transmisión se compone de dos cables metálicos (de cobre) de calibre oscilando entre 26 AWG* y 19 AWG (aproximadamente entre 0.4 y 0.91 mm) y aislados. El aislante utilizado en la mayoría de los casos es el polietileno, aunque en instalaciones antiguas aún pueden encontrarse cables aislados mediante papel.

La forma usual de distribuir los pares hasta los hogares de los abonados es la siguiente: de la central telefónica parten cables multipar, cada uno de los cuales puede contener hasta 50 grupos (conjunto de pares empaquetados, aislados y apantallados del resto). Cada uno de estos grupos puede contener 10, 25 o 50 pares (según normalizaciones). A medida que la red de despliegue avanza por la ciudad, se van separando los cables cada vez más, hasta llegar al último tramo del bucle en el cual el par ya viaja solo desde el cable principal hasta el local del abonado. Este último tramo es lo que se llama "cable de bajada" (*drop-wire*), el cual no suele ser un par trenzado, sino una simple línea bifilar.

Dentro de estos grupos, los cables de cobre van aislados, pero sin apantallar. Pero cada par se trenza sobre sí mismo, formando lo que se llama un cable UTP (*Unshielded twisted pair*). Se suelen utilizar cables UTP de categoría 3 o 5; la principal diferencia se encuentra en el número de vueltas por unidad de longitud. La mayor parte del cableado existente es del tipo UTP-3, para el que el *pitch* o paso de trenzado es inapreciable en el primer vistazo. El tipo UTP-5 tiene un paso de trenzado más pequeño, y sí se puede apreciar.

Como la transmisión de señales por estos cables se lleva a cabo en modo balanceado, conviene evitar en todo momento el modo común inducido por las interferencias, por lo que un parámetro importante es la relación λ/p , donde p es el *pitch* (vueltas/metro) y λ es la longitud de onda de la señal. Por ello, las transmisiones llevadas a cabo sobre pares de categoría 5 son de mejor calidad, ya que tienen un relación λ/p más pequeña.

* AWG es una unidad anglosajona para expresar el calibre (grosor) de los cables de cobre.

A su vez, y para intentar promediar las interferencias, se agrupan los cables trenzados en grupos de cuatro (por poner un número), y se vuelven a trenzar otra vez. La idea es que dos pares no estén juntos dentro del grupo durante un trayecto demasiado largo, puesto que así, si este cable se convierte en interferente, repartirá su interferencia entre un mayor número de pares.

3.2.1 – La longitud del par

El parámetro utilizado por un operador para informar de este dato es el porcentaje de abonados que se encuentran cubiertos en función de la longitud del bucle. Podemos decir por ejemplo que, con una longitud de 5 Km tenemos cubierto al 80 % de los abonados. Cuando el ANSI formó un grupo para estandarizar el ADSL (el T1.413), se supuso que el 20% de abonados restante se encontraban en zonas rurales, por lo que su demanda de servicios de datos a alta velocidad será escasa.

Como ya adelantamos en el primer capítulo del proyecto, la principal diferencia entre ADSL y VDSL es su diferencia de velocidad. Y esta diferencia de velocidad se alcanza gracias a que VDSL va a ser implementado sobre bucles de abonado muy cortos (<1Km). Con un bucle corto, limitamos la atenuación (que aumenta con la longitud del bucle), con lo que podemos aprovechar un mayor ancho de banda.

3.2.2 – Balance

Todas las señales que viajan por el bucle de abonado son transportadas en modo diferencial, en el cual, la corriente que circula por un hilo está compensada por una corriente igual pero de sentido contrario que viaja por el otro hilo del par. Se hace un gran esfuerzo (tanto en el proceso de fabricación del cable como en el diseño de los equipos de línea terminales) para que la componente en modo común sea lo más pequeña posible. Los transmisores deberían alcanzar una razón modo diferencial/modo común de al menos 55 dB en toda la banda.

Sin embargo, a causa del no-balanceamiento de los dos hilos con respecto a "tierra" (representada principalmente por los otros pares del grupo), siempre existe una parte del modo diferencial que pasa a ser común.

3.2.3 – Calibre de los hilos y cambios de calibre

Un parámetro importante que controla la posibilidad para la central telefónica de utilizar señalización para mantenimiento y diagnóstico es la resistencia DC del bucle, medida en la central con un cortocircuito en el domicilio del abonado. Por ejemplo, en EE.UU. este valor está limitado a 1500Ω . Esta limitación implica que para mayores distancias, debe ser instalado un par con un calibre mayor, con el fin de mantener la resistencia DC dentro del límite (la resistencia es inversamente proporcional al calibre).

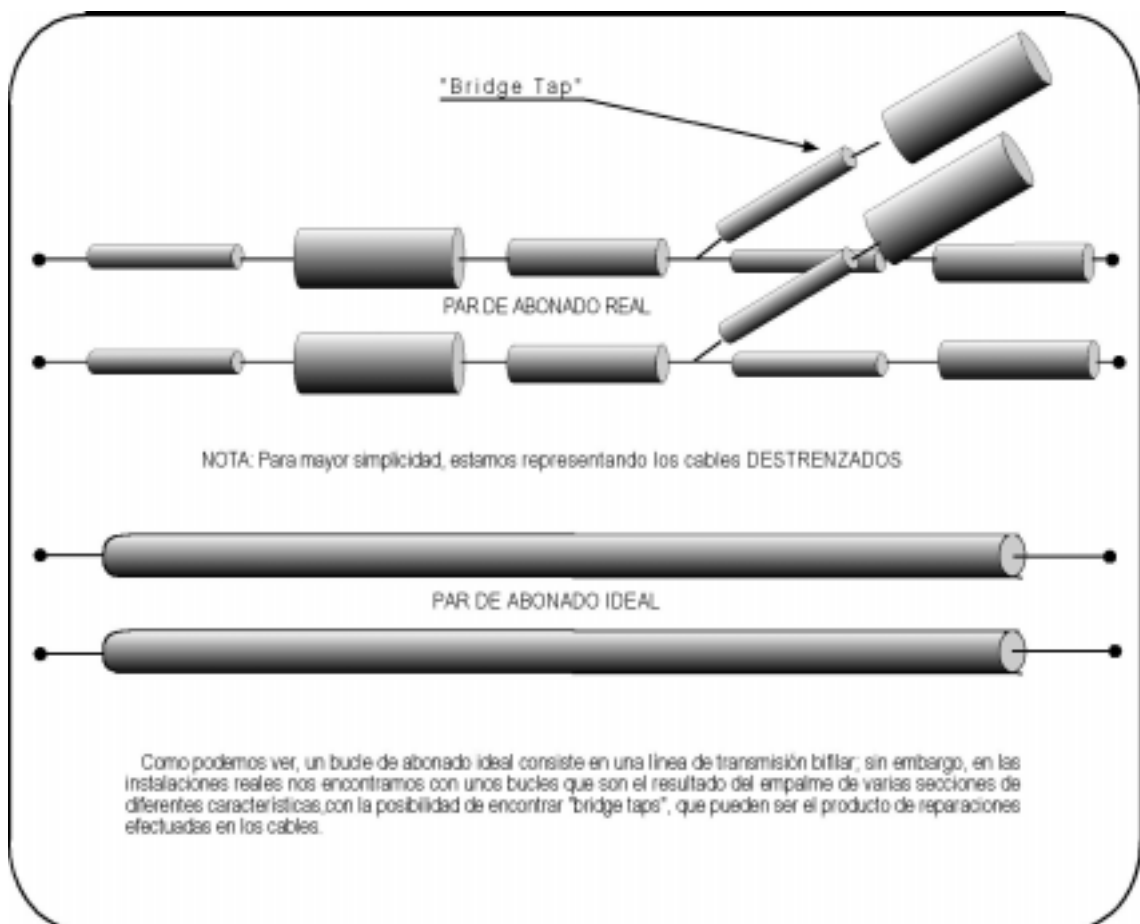


Figura 3.2 : Detalle de los cambios de calibre que tienen lugar en un bucle. También se incluye un bridge-tap

Pero esto no puede conseguirse en la práctica, ya que dentro de un grupo, todos los cables han de tener el mismo calibre. La solución que se adopta entonces es comenzar el despliegue desde la central con cables de calibre pequeño, e incrementar el calibre conforme nos vamos alejando de la central. Estos cambios de calibre conseguidos mediante empalmes, suponen de alguna manera una discontinuidad en las características de la línea de transmisión, y por tanto, van a afectar a la respuesta en frecuencia del bucle de abonado. En la figura 3.2 se muestra un esquema de los cambios de calibre que tienen lugar en una línea de transmisión.

3.2.4 – "Bridge taps"

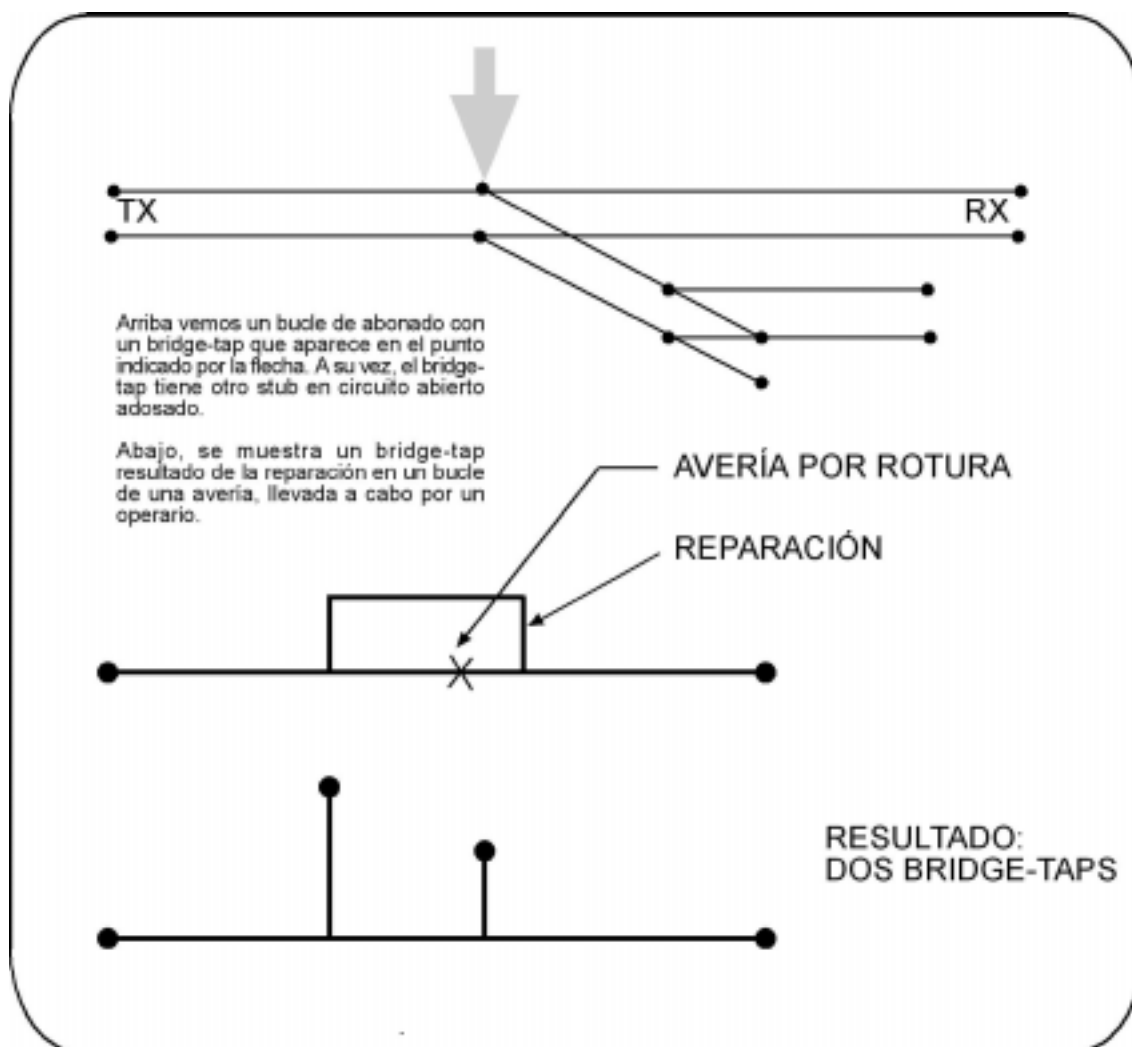


FIGURA 3.3 : Bridge-tap y una posible forma de aparición

Los *bridge taps* son tramos de cable UTP en circuito abierto que están conectados al par en consideración a lo largo de su longitud. Son consecuencia de algunas prácticas habituales en la red:

- Líneas compartidas: Durante la primera etapa de desarrollo de las redes telefónicas, era práctica habitual que varios abonados compartieran el uso del mismo par. Posteriormente, cuando se instalaban más cables y la privacidad se hacía más abordable, los cables de bajada hasta los otros abonados eran simplemente desconectados de los usuarios, dejando los *stubs* en circuito abierto conectados al par principal. formando un *bridge-tap*
- Reparaciones: Si se produce una rotura (discontinuidad en el paso de la corriente) en un cable, el operario puede simplemente empalmarlo sin desconectar las secciones rotas, formando así un *bridge-tap*.
- Múltiples tomas de teléfono en el local del abonado: Esta es la configuración más típica del cableado dentro de la casa, ya que así se permite tener más de un aparato (teléfono supletorio), o equipo de datos (ordenador conectado a internet). Aunque estos son *bridge-taps* de corta longitud, pueden ser muy significativos a frecuencias de VDSL.

En la figura 3.3 se muestran esquemas de líneas con “bridge-taps”. Y en la figura 3.2 se puede ver un esquema más realista.

3.2.5 – Bobinas de carga

A menudo se piensa que los bucles de abonado tienen un ancho de banda de sólo 4 KHz, pero como ya señalamos anteriormente, esta limitación viene impuesta por los equipos de multiplexación de la red telefónica, no es una limitación del bucle en sí. En los orígenes de la red telefónica, se comenzó a utilizar la multiplexación por división de frecuencia basada en múltiplos de

una banda base de 4 KHz. Por tanto, si el par de abonado va a ser utilizado como medio de acceso a la RTB, no nos va a ser necesario un ancho de banda mayor que 4 KHz.

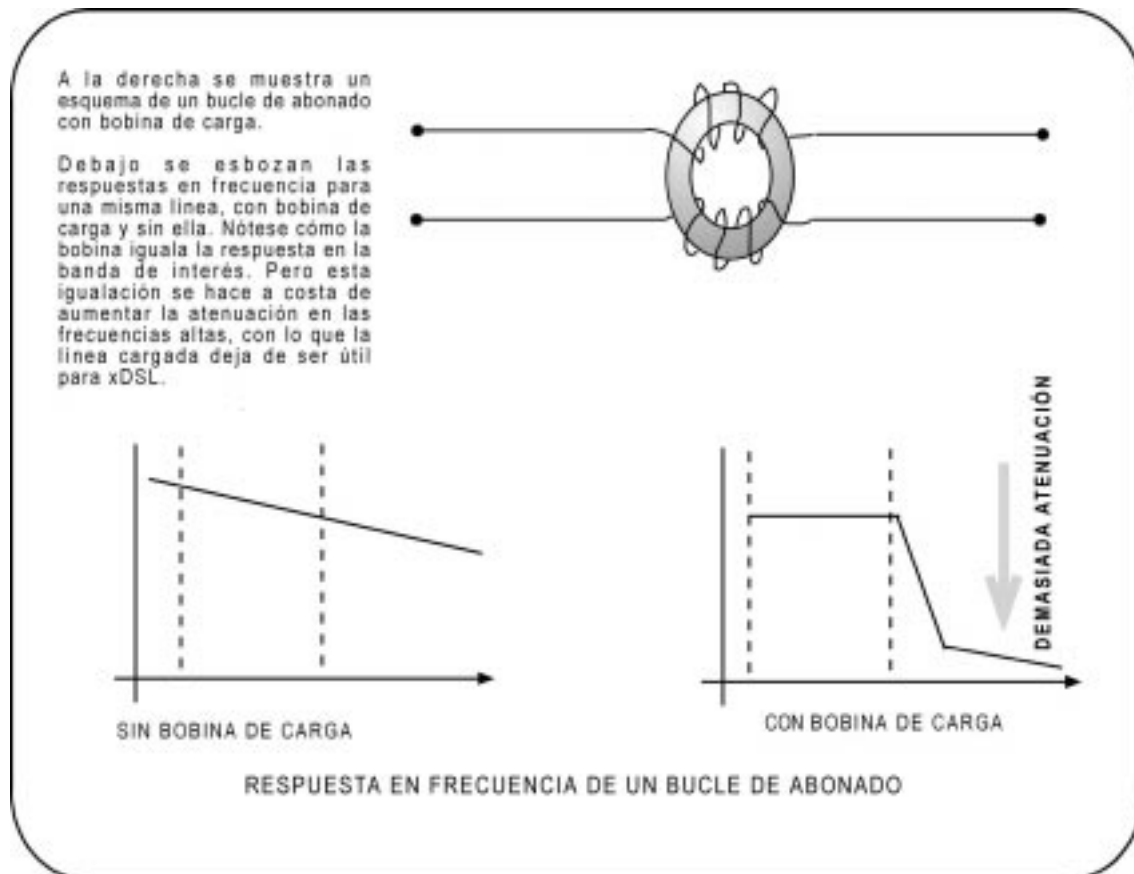


Figura 3.4 : Línea cargada. Detalle de la bobina de carga

Para frecuencias bajas, un UTP (unshielded twisted pair) se comporta como un circuito RC distribuido, y su respuesta en frecuencia decae a lo largo de los 4 KHz de la banda vocal. (como mucho, unos 12 dB en pares largos). Esta caída degrada la calidad de las transmisiones (fenómeno de distorsión lineal), así que Heaviside propuso la inserción de inductores (lo que llamamos "bobinas de carga") en intervalos regulares a lo largo de toda la longitud del bucle (véase la figura 3.4 para entender mejor este concepto). Una configuración usual en los Estados Unidos es la 26H88 (bobinas de 80 mH insertadas cada 6000 pies en un UTP del tipo 26-AWG). El caso es que las bobinas de carga convierten a la red RC en una red máximamente plana/Paso-bajo con una frecuencia de corte de unos 4 KHz. Pero para aplanar la respuesta en la banda vocal (banda de interés) las bobinas atenúan mucho la señal en la zona por encima de los 4 KHz; así que deben ser

eliminadas (o cortocircuitadas) si queremos que los servicios de banda ancha operen sobre el bucle.

Resulta paradójico ver cómo un elemento que se añadió a las líneas para mejorar sus prestaciones ahora tenga que ser eliminado para mejorarlas más aún.

3.2.6 – El cable de bajada ("drop-wire")

Cuando un cable emerge del cable de distribución de pares, se conecta al abonado mediante un cable de bajada. El término hace referencia a la bajada desde un poste o torreta, pero se utiliza incluso para el caso de que el cable de distribución sea subterráneo. Las bajadas de cable suelen ser de cobre, acero o aleación. Pueden ser trenzados o no trenzados, y su balance es mucho peor que cualquier otro tramo del bucle. Esto puede resultar en la mayor recogida de interferencias de radiofrecuencia. La impedancia característica de esta sección es típicamente más alta que en un UTP, y las consecuencias de la desadaptación de impedancias pueden ser muy significativas en el rango de uso de VDSL (longitudes de bucle muy cortas y frecuencias de señal elevadas). Este tramo del bucle fue ignorado en la definición de los bucles de test para ADSL, pero está siendo incluido en los tests para VDSL.

3.3 – MODELO EN ESCALERA DE UN CABLE UTP

Un cable UTP consta de inductancia y resistencia distribuidas, en serie, y de capacitancia y conductancia en paralelo. Estos son los llamados parámetros "primarios" del cable. Un esquema se muestra en la figura 3.5.

La capacitancia por unidad de longitud viene dada por

$$C_{per} = \frac{\pi k \epsilon_0}{\cosh^{-1}(D/d)} \text{ (F / m)}, \quad (3.1)$$

donde k es la constante dieléctrica del medio y ϵ_0 es la permitividad del vacío, igual a 8.85×10^{-12} F/m. En esta expresión se ha asumido que el medio de transmisión es homogéneo, pero en la práctica no lo es. Como vemos, hay dos capas de aislante, pero entre ellas, y fuera, hay una mezcla de aire y de aislante de otros hilos. La constante dieléctrica del polietileno es de 2.26, pero el valor efectivo de k varía ligeramente, y parece ser que es de unos 2.05.

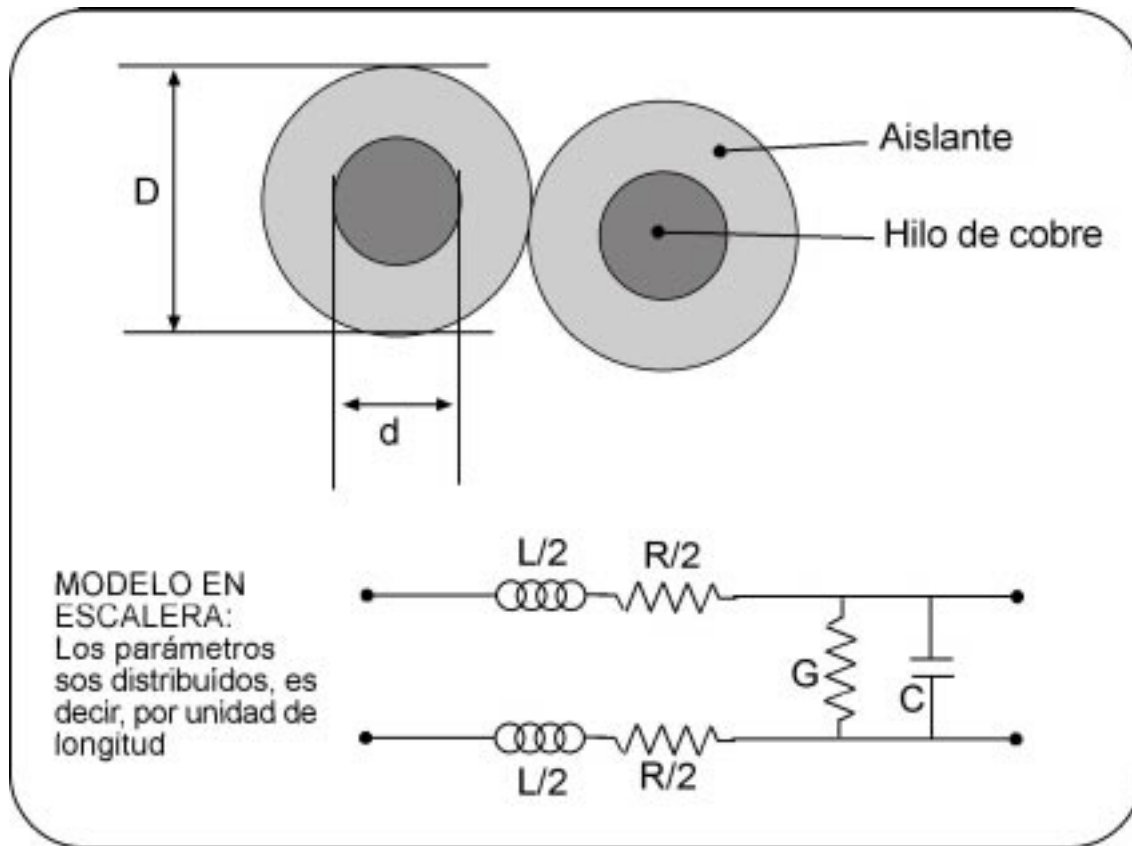


FIGURA 3.5 : Sección de un par y modelo en escalera del mismo

La inductancia por unidad de longitud a altas frecuencias (donde la corriente es transportada principalmente por la superficie de los hilos, y no por el interior), viene dada por

$$L_{per} = \frac{\mu_0}{\pi} \cosh^{-1}(D/d) = \frac{\mu_0}{\pi} \ln \left[\frac{D}{d} + \sqrt{\left(\frac{D}{d}\right)^2 - 1} \right] \text{ (H / m)}, \quad (3.2)$$

donde μ_0 es la permeabilidad del vacío, que es igual a $4\pi \times 10^{-7}$ H/m. De este modo, la impedancia característica queda

$$Z_0 = \sqrt{\frac{L_{per}}{C_{per}}} = \sqrt{\frac{\mu_0}{k\epsilon_0}} \frac{\cosh^{-1}(D/d)}{\pi} \quad (\Omega). \quad (3.3)$$

En la literatura se han usado otras fórmulas para C_{per} y L_{per} , con el *arccosh* reemplazado por $\ln(2D/d-1)$ o $\ln(2D/d)$. Ambas aproximaciones son válidas si $D/d \gg 1$, pero esta hipótesis no es cierta en un cable UTP.

El valor de L_{per} dado por (3.2) recibe normalmente el nombre de *externo* porque resulta del flujo de corriente por el exterior del cable. En las bajas frecuencias, cuando la corriente fluye por toda la sección del hilo, hay también una inductancia *interna*,

$$L_{int} = \frac{\mu_0}{4\pi}, \quad (3.4)$$

y para bajas frecuencias, la inductancia distribuida es la suma de la interna con la externa.

Hasta ahora, hemos presentado los parámetros "primarios" de la línea: R , C , L y G . Pero la línea también puede caracterizarse por un par de parámetros llamados "secundarios", que son la impedancia característica (Z_0) y la constante de propagación ($\gamma = \alpha + j\beta$). Ambos son magnitudes complejas, por lo que su completa definición exige el conocimiento de cuatro cantidades reales. Sin embargo, usar estos parámetros para calcular la respuesta exacta de una línea en la que cambia varias veces el calibre, y en la que están presentes *bridge-taps* (como ocurre en el caso que nos ocupa) resulta una tarea tediosa. Así que el uso habitual de estos parámetros es el de utilizarlos como paso intermedio antes de calcular otros parámetros "terciarios", como son los elementos de la *matriz de parámetros en cadena*, los cuales son más útiles en el análisis de líneas con desadaptaciones a lo largo de su longitud. No obstante, un conocimiento explícito de la constante de propagación fue muy útil en las primeras etapas del DSL, por varias razones:

- La tasa de cambio de la constante de fase, β (parte imaginaria de la constante de propagación) define el retardo de propagación de una sección en cada frecuencia:

$$\tau = \frac{d\beta}{d\omega} \text{ (s/m)} \quad (3.5)$$

- Una estimación de la constante de atenuación, α (parte real de la constante de propagación) va a ser muy útil en el cálculo de los efectos de un *bridge-tap* o del *crosstalk*. Por encima de los 300KHz, la constante de atenuación puede aproximarse por

$$dB = 8.686\alpha(f) \approx \alpha_1 \sqrt{f} . \quad (3.6)$$

Valores de α_1 (normalizados a 1MHz) para diversos tipos de cables europeos o americanos pueden encontrarse en tablas en [14].

3.4 – CROSSTALK

El *crosstalk* entre pares de un cable multipar es el principal inconveniente en cualquier implementación de un sistema DSL. La razón de que aparezca es la existencia de acoplamientos inductivos y capacitivos (o más precisamente, el no-balanceamiento de dichos acoplamientos). Para intentar cancelar sus efectos, es necesario hacer un estudio de las funciones de transferencia del crosstalk par-a-par.

Si el par bajo consideración es considerado como el interferente, las corrientes y voltajes inducidas en los otros pares viajan en ambas direcciones; aquellas que continúan en la misma dirección que la señal interferente reciben el nombre de FEXT (far-end crosstalk), y aquellas que regresan en dirección a la fuente del interferente son llamadas NEXT (near-end crosstalk). Esto puede comprenderse a la luz de la figura 3.6. Para aquellos sistemas en los que ocurran a la vez los dos fenómenos, el NEXT será en general mucho más

importante. El NEXT se incrementa con la frecuencia, y para frecuencias de VDSL (hasta los 15 MHz) puede hacerse intolerable.

Hata ahora, el análisis del caso más desfavorable en la función de transferencia del *crosstalk* para un cable multipar ha sido uno de los principales objetivos. El conocimiento del caso más desfavorable proporcionaba a los ingenieros la posibilidad de predecir (y de garantizar) las tasas máximas alcanzables en los sistemas DSL. Ahora, sin embargo, cobra mayor importancia el uso de promedios estadísticos (más reales), que son los que se utilizarán en este proyecto a la hora de hacer simulaciones.

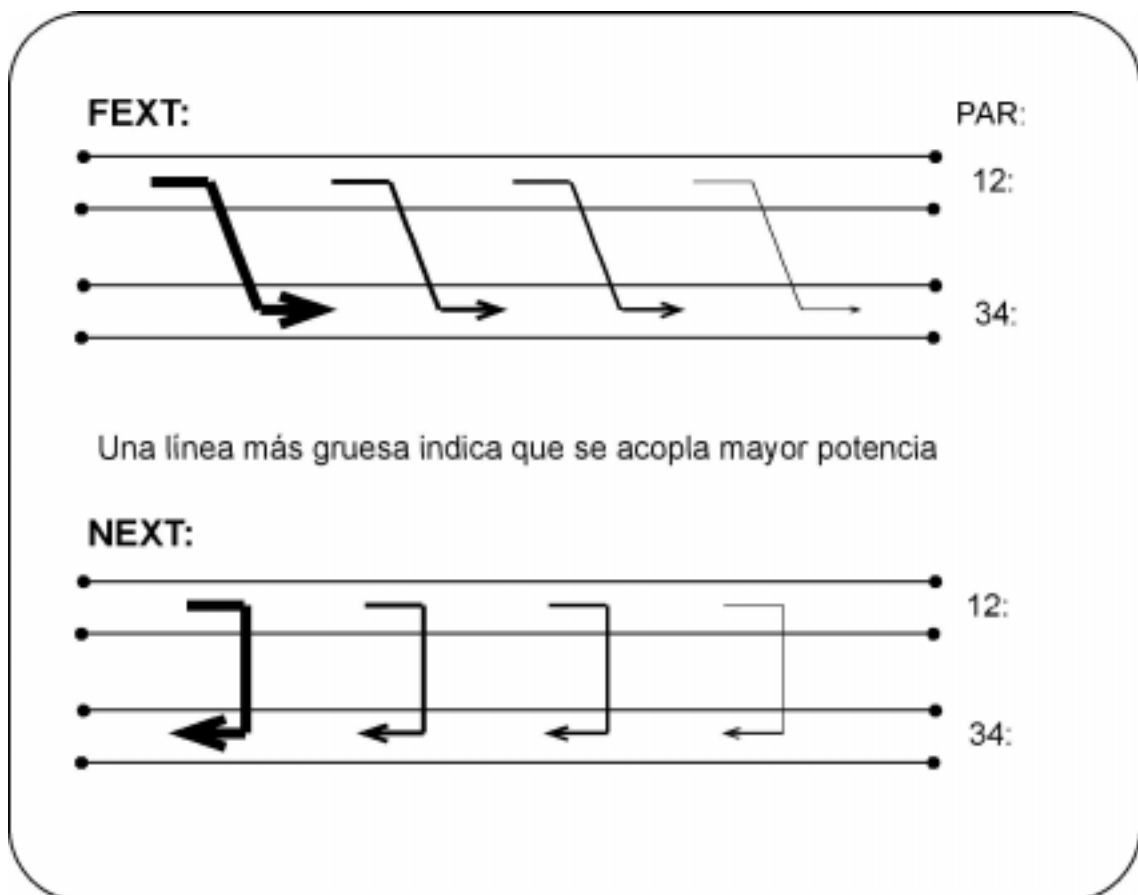


FIGURA 3.6 : Clasificación del crosstalk en NEXT y FEXT

3.5.1 – NEXT

Como se muestra en la Figura 3.4, el NEXT desde el par 12 hasta el 34 es la suma de un número infinito de señales que se propagan por el par 12, se acoplan al par 34 y se propagan hacia atrás en el 34, esto es,

$$H_{NEXT}(l, f) = \int_0^l H_{f12}(\lambda, f) H_{XT}(\lambda, f) H_{b34}(\lambda, f) d\lambda, \quad (3.7)$$

donde $H_{f12}(\lambda, f)$ es la función de transferencia hacia delante de una sección de longitud λ del par 12, $H_{XT}(\lambda, f)$ es la función de transferencia del acoplamiento a una distancia λ de la entrada y $H_{b34}(\lambda, f)$ es la función de transferencia hacia atrás del par 34.

La ecuación (3.7) puede ser simplificada si los pares 12 y 34 no tienen *bridge-taps* (o tienen la misma configuración de *bridge-taps*); entonces sus características de transmisión son las mismas y por tanto,

$$H_{f12}(\lambda, f) = H_{b34}(\lambda, f) = H(\lambda, f) \approx e^{-\gamma\lambda}. \quad (3.8)$$

Entonces,

$$H_{NEXT}(l, f) = \int_0^l H(2\lambda, f) H_{XT}(\lambda, f) d\lambda. \quad (3.9)$$

Si escribimos H_{XT} como el producto adimensional de una cierta susceptancia de acoplamiento por una cierta (inespecificada) impedancia de carga, la función de transferencia de potencia del NEXT puede ser escrita como

$$|H_{NEXT}(l, f)|^2 \approx \int_0^l |H(2\lambda, l)|^2 (\omega CR)^2 d\lambda. \quad (3.10)$$

Y si ahora consideramos la atenuación en la línea como proporcional a $f^{1/2}$,

$$|H_{NEXT}(l, f)|^2 \approx \int_0^l (2\pi CR)^2 f^2 e^{-4\alpha\sqrt{f}l} dl = \frac{(2\pi RC)^2}{4\alpha} f^{1.5} (1 - e^{-4\alpha\sqrt{f}l}). \quad (3.11)$$

El último término del paréntesis en la ecuación (3.11) tiende a 1 para valores grandes de la longitud l ("grande" significa aquí "longitud para la cual la atenuación es grande"). Esto ocurre para todas las longitudes en el rango de VDSL, por lo que podemos afirmar que el NEXT es independiente de la longitud del bucle y proporcional a la frecuencia, mediante el término $f^{1.5}$.

3.5.2 – FEXT

La forma más simple de FEXT, llamada FEXT de igual nivel (EL-FEXT), se muestra en la Figura 3.7. Los transmisores 12 y 34 están situados en el mismo lugar, y lo mismo ocurre con los receptores 12 y 34. Puede verse que todas las contribuciones al FEXT recibidos en el receptor 34 (víctima de las interferencias) han de recorrer toda la longitud del bucle: un tramo por el par 12 (culpable) y el resto por el par 34 (víctima). Esto significa que

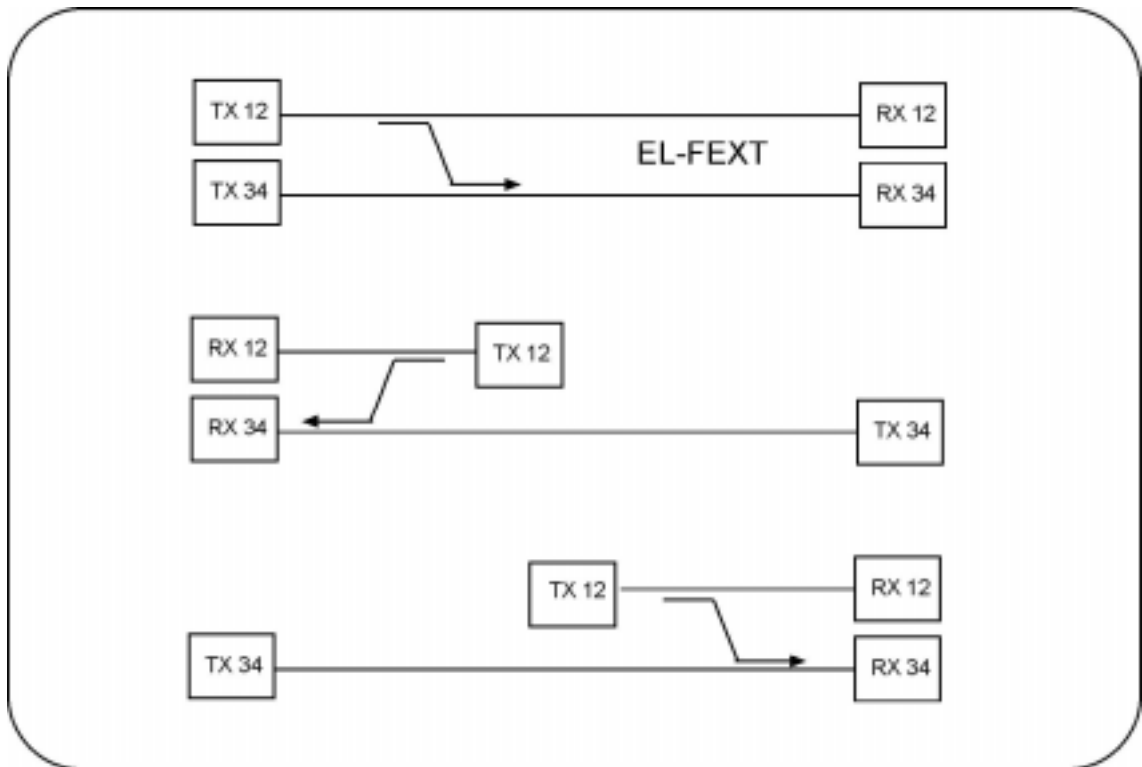


Figura 3.7 : EL-FEXT y otros tipos de FEXT

$$H_{FEXT}(l, f) = \int_0^l H_{12}(\lambda, f) H_{XT}(\lambda, f) H_{34}(l - \lambda, f) d\lambda. \quad (3.12)$$

Entonces, mediante las mismas etapas de simplificación llevadas a cabo en el análisis del NEXT,

$$H_{12}(\lambda, f) H_{34}(l - \lambda, f) = H(l, f) \quad (3.13)$$

$$H_{FEXT}(l, f) = H(l, f) f \int_0^l H_{XT}(\lambda) d\lambda. \quad (3.14)$$

Puede demostrarse [14] que la última integral equivale a una constante que varía según el par de pares que estemos analizando. De esta forma, la función de transferencia de potencia del FEXT puede ponerse en la forma

$$|H_{FEXT}(l, f)|^2 = |H(l, f)|^2 k_{FEXT} l f^2, \quad (3.15)$$

donde k_{FEXT} es un coeficiente constante que tendrá un valor distinto en función del par de pares que estemos considerando. También debemos notar que ahora la cantidad de potencia que se acopla sí depende de la longitud del bucle en cuestión.

3.5.3 – Análisis estadístico de los fenómenos NEXT y FEXT

Durante la década de los 70 y principios de los 80, se hicieron multitud de medidas de crosstalk en bucles de abonado. La mayoría fueron realizadas sobre cables de pares de 50 bucles. Aunque para VDSL suelen utilizarse cables de 25 pares, las medidas y sus conclusiones son aplicables también a este tipo de cables. Dichas medidas eran de dos tipos:

- Medidas par a par. Para cada cable de N pares, deben hacerse $N(N-1)/2$ medidas, a cada frecuencia de interés.

- Medidas $(N-1)$ pares a un par. Aquí se considera la influencia de $N-1$ interferentes sobre un par víctima. Se considera que todos los pares culpables interfieren de forma simultánea, independiente y aleatoria. El resultado es hacer N sumas de potencia por cada cable de pares.

De ambos tipos de medidas, se extrajeron modelos para las funciones de densidad de probabilidad de la potencia de las interferencias. Una de las conclusiones que podemos extraer de estos trabajos es una expresión para los promedios estadísticos del crosstalk. Esta expresión devuelve la función de transferencia de la potencia de crosstalk, en decibelios.

$$10 \log |XT|^2 = K_1 + K_2 \log(N) + K_3 \log(f) + K_4 \log(l), \quad (3.16)$$

donde l está expresada en kft, f expresada en MHz, y K_1 , K_2 , K_3 y K_4 se indican en la tabla adjunta (tabla 3.1). N es el número de fuentes interferentes.

	<i>NEXT(dB)</i>		<i>FEXT(dB)</i>	
	Promedio	Peor caso	Promedio	Peor caso
K_1	-75	-51	-75	-51
K_2	16	6	16	6
K_3	15	25	20	20
K_4	0	0	10	10

Tabla 3.1 : Valores de los parámetros en la expresión (3.16)

Es una casualidad que coincidan algunos valores para el NEXT y para el FEXT, ya que los coeficientes son el resultado de elegir para la expresión una frecuencia normalizada de 1MHz y una longitud normalizada de 1Kft. Otra elección en estas normalizaciones conduciría a valores distintos. Obsérvese cómo el coeficiente K_4 se anula en el caso del NEXT, ya que en este caso no existe dependencia con la longitud total del bucle.

3.5 – OTRAS FUENTES DE RUIDO EN EL BUCLE

En esta última sección del capítulo, se van a mencionar otras dos posibles contribuciones modelables como ruido en el canal. Dichas contribuciones son:

- Ruido térmico: Modelado como ruido blanco de fondo, no se va a considerar en las simulaciones, al ser su potencia despreciable frente a la de las interferencias de tipo NEXT y FEXT.
- Interferencias de radiofrecuencia (RFI): Se producen debido a que las emisiones radioeléctricas del exterior (radioaficionados, estaciones de AM, etc.) se acoplan con distinta potencia en cada uno de los hilos constituyentes del par. Este tipo de interferencias (sobre todo las procedentes de radiodifusión AM) se caracterizan por estar perfectamente localizadas en frecuencia, lo cual permite al sistema de entrenamiento VDSL detectarlas con facilidad y tener este hecho en cuenta en el algoritmo de carga de los subcanales (esto se estudiará en el capítulo 4). No se van a considerar en nuestro modelo, por simplicidad.
- Otras fuentes: En este apartado, tienen cabida el resto de fuentes de ruido que pudieran aparecer en el transcurso de una transmisión DSL. Esto incluye ráfagas, fenómenos transitorios,

etc. Suelen modelarse de forma conjunta como *ruido impulsivo*. Tampoco se considerarán.

3.6 – OBTENCIÓN DE UN MODELO DE CANAL

Como se vio al principio de este capítulo, el objetivo del mismo era el de determinar un modelo matemático del canal de comunicaciones en un enlace VDSL. Para ello, hemos realizado un estudio de las características tanto de la planta de distribución de cableado de un operador como de la línea de cobre en sí. Ahora se trata de aplicar todo lo que hemos visto para obtener un modelo válido en una situación real. Para la determinación del modelo, debemos proporcionar:

- Una función de transferencia de la línea.
- Un modelo del ruido aditivo que se añade a la salida del canal.

Para el cálculo de la función de transferencia tenemos disponibles algoritmos para el cálculo de la misma por ordenador, en función de las características (agrupamiento, calibre, *bridge-taps*, etc) de cada sección del bucle de abonado. Uno de estos algoritmos puede encontrarse en código MATLAB en [14]. Otro de ellos, también utilizado, es el programa LINEMOD, realizado en la Universidad de Stanford. También existen modelos en la bibliografía ([7], [9], [13]).

En cuanto al ruido aditivo, sólo se van a considerar en su modelo las contribuciones del NEXT y del FEXT (aún sabiendo que el primer efecto es dominante sobre el segundo). El resto de las contribuciones estudiadas no se va a tener en cuenta en nuestras simulaciones, dada su poca importancia frente a las ya incluidas. Las potencias de ruido se calcularán a partir de las expresiones obtenidas en el apartado 3.5.3 (resumidas en la expresión (3.16)).

Teniendo en cuenta las consideraciones anteriores, y si suponemos que el transmisor emite una potencia máxima $P_{\text{máx}}$, las densidades espectrales de potencia de las señales que entran en juego en el sistema son las siguientes:

$$PSD_{dB(\text{señal})}(f) = 10 \log_{10} \left(\frac{10^{0.1P_{\text{máx}}}}{BW} \right) + H_{dB(\text{canal})}(f) \quad (\text{dBm/Hz}) \quad (3.17)$$

$$PSD_{dB(\text{NEXT})}(f) = 10 \log_{10} \left(\frac{10^{0.1P_{\text{máx}}}}{BW} \right) + H_{dB(\text{NEXT})}(f) \quad (\text{dBm/Hz}) \quad (3.18)$$

$$PSD_{dB(\text{FEXT})}(f) = 10 \log_{10} \left(\frac{10^{0.1P_{\text{máx}}}}{BW} \right) + H_{dB(\text{FEXT})}(f) + H_{dB(\text{canal})}(f) \quad (\text{dBm/Hz}). \quad (3.19)$$

donde BW es el ancho de banda que se utiliza en el canal para la transmisión.

4.1 - INTRODUCCIÓN

En este capítulo vamos a introducir la modulación multiportadora (MCM). Este estudio es importante, puesto que los conceptos de MCM son la base de la implementación del sistema VDSL que vamos a simular en este proyecto. Veremos que básicamente se trata de una multiplexación por división de frecuencia. Sin embargo, se habla de modulación multiportadora porque en este caso no tenemos varias fuentes independientes que se multiplexan para transmitirse en una única señal, sino que partiendo de una sola fuente de datos, se extraen de ella varios flujos binarios y son estos flujos los que se multiplexan para formar la señal agregada. Explicaremos las diferentes técnicas utilizadas para extraer los flujos binarios del flujo principal, es decir, qué porcentaje de la carga total va a transportar cada subcanal y en función de qué criterios. Posteriormente, veremos técnicas para la separación de las señales transportadas por las diferentes portadoras en el receptor; y veremos que podemos hacer un uso más eficiente del ancho de banda si permitimos que se solapen los espectros de los subcanales (con ciertas condiciones).

El capítulo finaliza con la presentación del sistema DMT; se describirá un sistema básico y se analizarán cada uno de los bloques del transceptor. Comienzan a vislumbrarse algunos efectos no deseados que dificultan las transmisiones, lo cual puede ser solucionado mediante el uso de igualadores.

Con el término "modulación multiportadora" designamos al conjunto de técnicas de transmisión de datos que dividen el flujo principal de bits en varios flujos entrelazados que se transmiten en paralelo por el canal, utilizando cada uno de estos flujos para modular a una portadora distinta. Está claro entonces que podemos ver a la modulación multiportadora como una forma de FDM (Multiplexación por división de frecuencia).

Antes de profundizar más en el tema, vamos a indicar una de las ventajas de la MCM frente a los sistemas de portadora única. Un canal de comunicaciones real presenta fenómenos de ruido aditivo y de interferencia entre símbolos (ISI), y el grado en el que se manifiestan estas interferencias limita la máxima tasa de datos transmisible. En los sistemas de portadora única es habitual reducir la ISI mediante el empleo de técnicas de igualación en el receptor; lo que ocurre es que esta reducción de la ISI se hace a costa de aumentar la potencia de las señales interferentes [5]. En la figura 4.1 se ilustra este hecho.

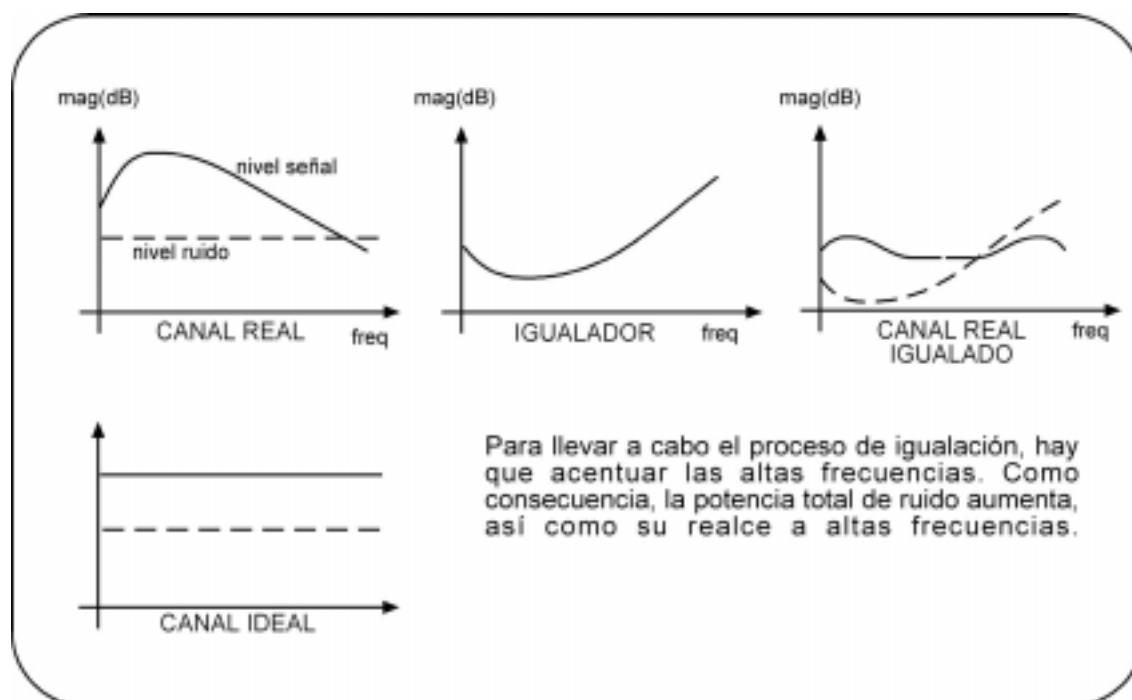


Figura 4.1 : Ilustración del aumento de la potencia de ruido causado por la igualación

Por el contrario, en un sistema de modulación multiportadora, podemos a priori suponer que si el número de sub-bandas en las que se divide el canal es elevado, la función de transferencia del canal se mantiene constante en el ámbito de cada subcanal, por lo que se elimina la posibilidad de ISI (dentro del propio subcanal). Por supuesto, esto está supeditado a que en el sistema MCM se puedan separar perfectamente los subcanales, lo cual nos llevará a efectuar los procesos de demodulación y detección por separado para cada subcanal. La idea se refleja en la figura 4.2.

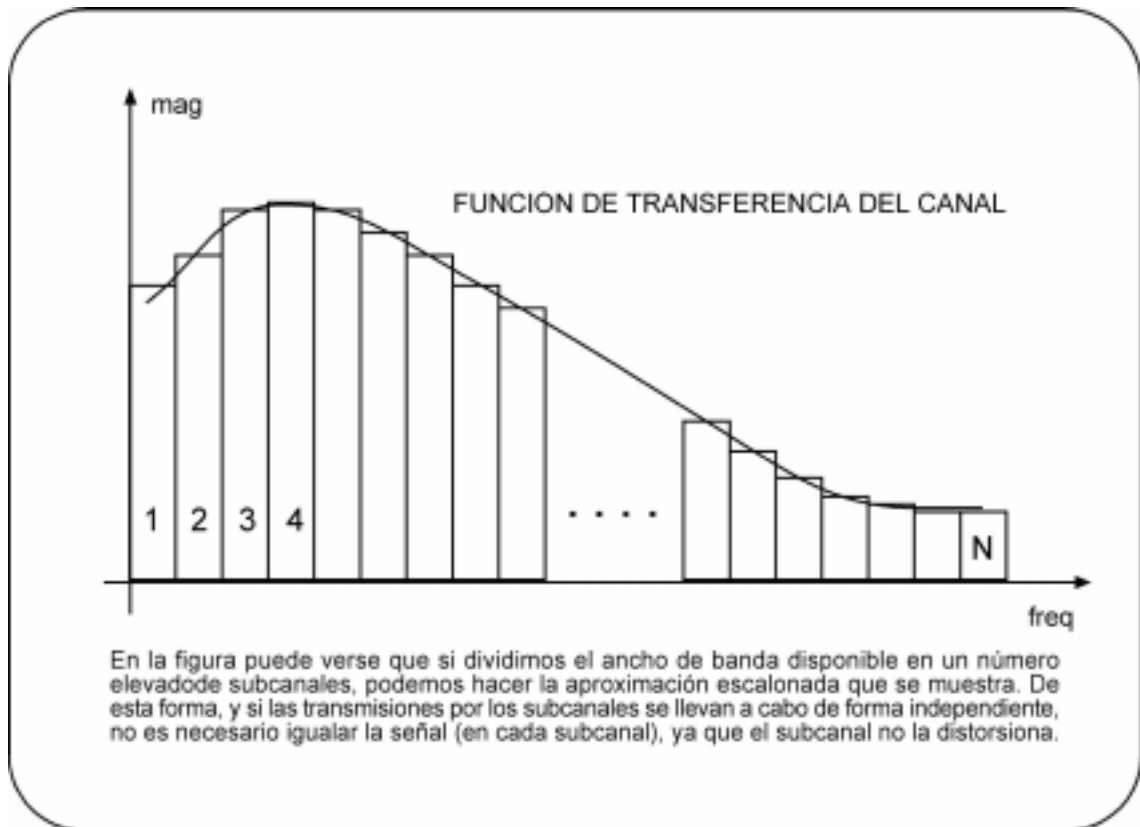


Figura 4.2 : Aproximación escalonada del canal, fruto de la división en subcanales

4.2 - DESCRIPCIÓN DE UN SISTEMA BÁSICO

En la figura 4.3, mostramos el esquema básico de una transmisión MCM. Casi todo el capítulo se basa en [5], [7-9] y [13].

La fuente de información genera datos a una velocidad de R_b b/s. Lo primero que hace el transmisor es agrupar esta información en bloques de M bits; cada uno de estos bloques es lo que constituirá un símbolo en la transmisión MCM. De esos M bits constituyentes del bloque, los m_1 primeros los asignará al subcanal 1, los m_2 siguientes al subcanal 2 y así sucesivamente hasta asignar los m_n últimos bits al subcanal n . Se observa así que las velocidades de transmisión de información (bps) de los distintos subcanales pueden ser desiguales. Ya estudiaremos más adelante cómo esta propiedad nos va a permitir transmitir a velocidad más lenta por los subcanales que presenten una mayor atenuación y/o estén más afectados por

ruido e interferencias. Existe incluso la posibilidad de no utilizar aquellos subcanales en los que haya una SNR muy degradada.

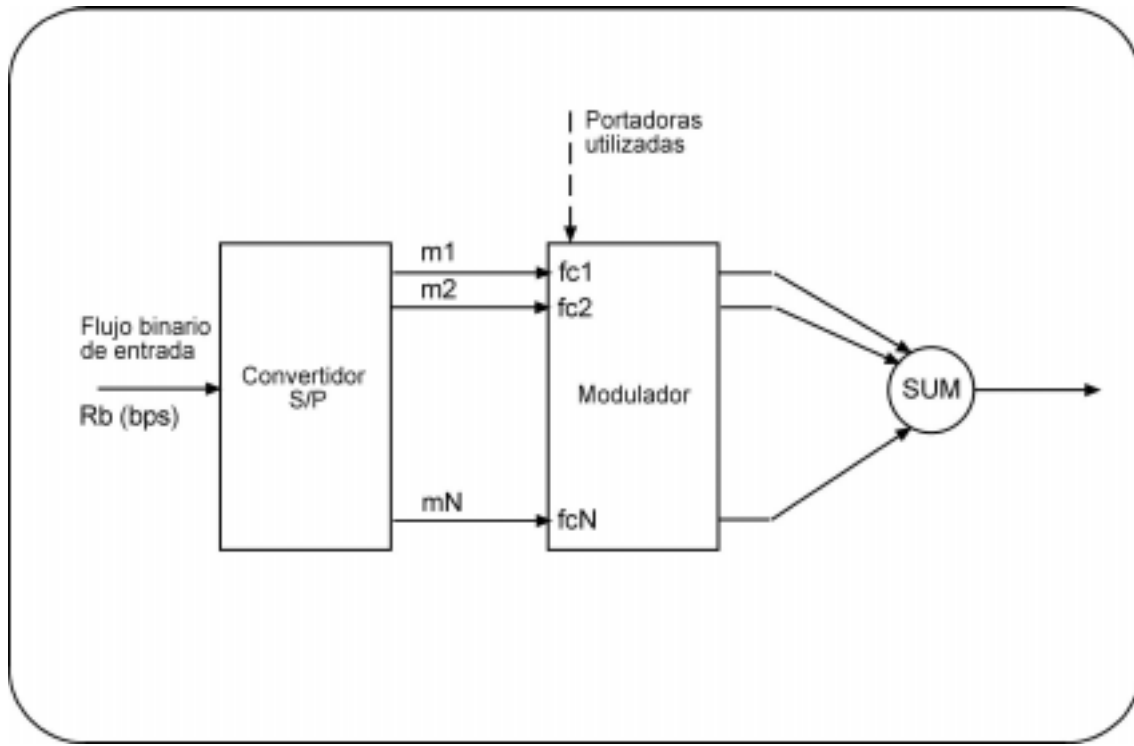


Figura 4.3 : Esquema básico de una transmisión MCM

Llegados a este punto, podemos afirmar que son ciertas las siguientes premisas:

- Un símbolo en el sistema MCM está formado por M bits, siendo M la suma de los m_i bits que transporta cada subcanal, y que van a venir determinados por un algoritmo de asignación.
- La duración de un símbolo es de T_s . Si R_b es la velocidad de transmisión de información se cumple la relación

$$T_s = \frac{M}{R_b}. \quad (4.1)$$

- Una vez que el algoritmo de asignación de bits ha terminado su trabajo (paso necesario antes de dar comienzo a la transmisión

propiamente dicha), conoceremos M , y, por tanto, conoceremos la tasa binaria R_b a la que trabajará el sistema.

Los diversos subcanales, trabajan en paralelo, por lo que la velocidad de transmisión (baudios) es igual en todos ellos. No ocurrirá lo mismo con la velocidad binaria, que dependerá también de los coeficientes m_i (número de bits por símbolo). Por tanto, cada subcanal debe transmitir en cada intervalo de duración T_s un total de m_i bits.

4.3 - INDEPENDENCIA ENTRE SUBCANALES

Uno de los aspectos más importantes cuando se diseña una aplicación MCM es el de la independencia entre los subcanales. Este concepto se ocupa de que las transmisiones efectuadas entre dos subcanales distintos no se interfieran entre sí, de que lo que se transmite por un subcanal no afecte a la señal que se propaga por los demás. Además, una vez que se ha logrado este objetivo, el sistema MCM debe proporcionar mecanismos para que el receptor pueda identificar y separar fácilmente las señales que se propagan por los diferentes subcanales.

Puede parecer a primera vista que la única forma de conseguir la separación perfecta de los subcanales en el receptor es la de separar los subcanales en el dominio de la frecuencia con una banda de guarda entre ellos. Sin embargo, existen técnicas de modulación ortogonal que permiten que los espectros de los subcanales se solapen, pero garantizando su perfecta separabilidad en el receptor. Este tipo de técnicas de modulación se estudiarán con más profundidad en este proyecto, cuando se aborde el estudio del sistema DMT (Discrete Multitone Transmission). Con este tipo de modulaciones se hace un uso espectral más eficiente del canal de comunicaciones, y podemos obtener rendimientos incluso mayores que los conseguidos en sistemas de portadora única. Describiéndolos con más detalle, tres han sido los métodos históricos usados hasta ahora para la separación de los subcanales [5]:

- Al principio, los modems MCM utilizaron la multiplexación FDM, y se usaron filtros para separar por completo las sub-bandas. Como consecuencia de la dificultad de implementación de filtros muy selectivos (abruptos) las señales tenían que utilizar un ancho de banda $(1+\alpha)f_s$, donde f_s es el ancho de banda mínimo exigido por el criterio de Nyquist. Como consecuencia, la eficiencia en el uso del espectro pasa a ser

$$\eta = \frac{1}{1+\alpha} \quad (4.2)$$

- Posteriormente se utilizaron sistemas que incrementaron esta eficiencia hasta casi el 100%, utilizando técnicas de modulación SQAM (Staggered QAM, QAM escalonada). Los espectros de cada subcanal utilizan un exceso de ancho de banda, pero el solape en las frecuencias de 3dB hace cumplir la condición de suma plana y, por tanto, se asegura la independencia. La ortogonalidad entre los subcanales se logra transmitiendo las componentes en fase y en cuadratura de las señales QAM con un offset de medio periodo de símbolo con respecto a los subcanales adyacentes. Se requiere una calidad de filtrado menor que con el FDM (separación completa), pero aún es considerable. Esto provoca que el sistema sólo pueda ser utilizado cuando el número de portadoras es del orden de 20 o inferior.
- El tercer sistema que se ha utilizado es la modulación QAM, en la que las portadoras son "generadas" por los datos. Los espectros individuales son ahora funciones sinc. No son funciones limitadas en banda, pero como ya veremos, la separación no se consigue por filtrado paso de banda, sino por procesamiento en banda base. La gran ventaja de esta técnica es que tanto el emisor como el receptor pueden implementarla de forma eficiente con el uso de FFT's.

4.4 - MEDIDA DE PRESTACIONES DE UN SISTEMA DE MODULACIÓN MULTICANAL

Para medir la calidad de un sistema de transmisión multicanal, es evidente que nos tenemos que fijar en parámetros distintos que en los sistemas de portadora única. El objetivo último que intentaremos conseguir es el de lograr un mayor rendimiento que en un sistema de portadora única. Para ello, definiremos parámetros que permitan comparar las dos formas de modulación.

En el análisis que vamos a realizar en este apartado, vamos a suponer que se utiliza un codificador QAM para la transmisión de datos en cada subcanal. Se elige este tipo de codificación porque la QAM es una solución muy versátil, dado que se adapta muy bien a las transmisiones de un número variable de bits (recordemos que en un sistema MCM cada uno de los subcanales transporta un número distinto de bits por símbolo). Vamos a comenzar con un pequeño repaso de la modulación QAM.

4.4.1 – Breve repaso de la modulación QAM

La modulación de amplitud en cuadratura es un método de modulación muy empleado para codificar un número variable de bits, en una señal que está modulada tanto en fase como en amplitud. Para la misma frecuencia, se emiten dos portadoras en cuadratura (ortogonales), por lo que se nos permite transmitir símbolos complejos, proyectando la parte real sobre la componente en fase y la parte imaginaria sobre la componente en cuadratura. Como sabemos, el conjunto de todos los símbolos complejos posibles se denomina constelación del codificador QAM.

En nuestras simulaciones hemos diseñado un codificador QAM sencillo, dado que se limitará a utilizar constelaciones rectangulares. Pero para determinadas aplicaciones comerciales, se utilizan constelaciones más complicadas que permiten obtener mayores prestaciones en el sistema (como por ejemplo, menores requerimientos de potencia o mayor inmunidad frente al

ruido). En las figuras 4.4 y 4.5 se muestran esquemas de algunas de las constelaciones utilizadas.

Una de los parámetros que hay que tener en cuenta en la implementación de un codificador QAM es la energía media disponible que vamos a tener para realizar las transmisiones. Esta energía promedio transmitida está íntimamente relacionada con la distancia entre los puntos de la constelación. Por ejemplo, si nos fijamos en un canal que opera a una frecuencia ω_0 , y llamando $X_i=(I_i, Q_i)$ al símbolo complejo a transmitir, la señal que genera el codificador será de la forma:

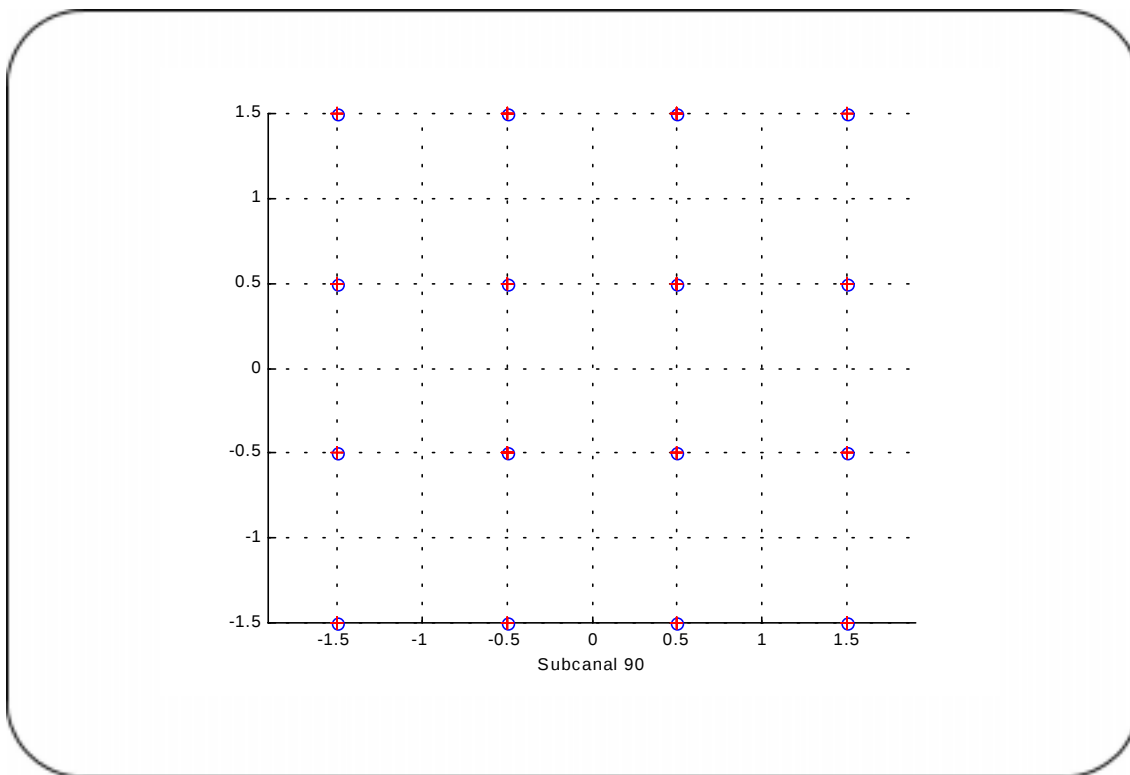


Figura 4.4 : Constelación para un número par de bits/símbolo

$$s(t) = kI_i \cos(\omega_0 t) + kQ_i \sin(\omega_0 t), \quad (4.3)$$

donde el factor k se utiliza para ajustar la distancia entre los puntos de la constelación, d . Este parámetro se relaciona con la energía promedio transmitida, ε , según la expresión [13]

$$d^2 = \frac{6\varepsilon}{M-1} \quad (4.4)$$

El término $M=2^b$ se corresponde con el número total de símbolos en la constelación, necesarios para transmitir b bits por símbolo. También asumiremos que todos los M símbolos de la constelación son igualmente probables.

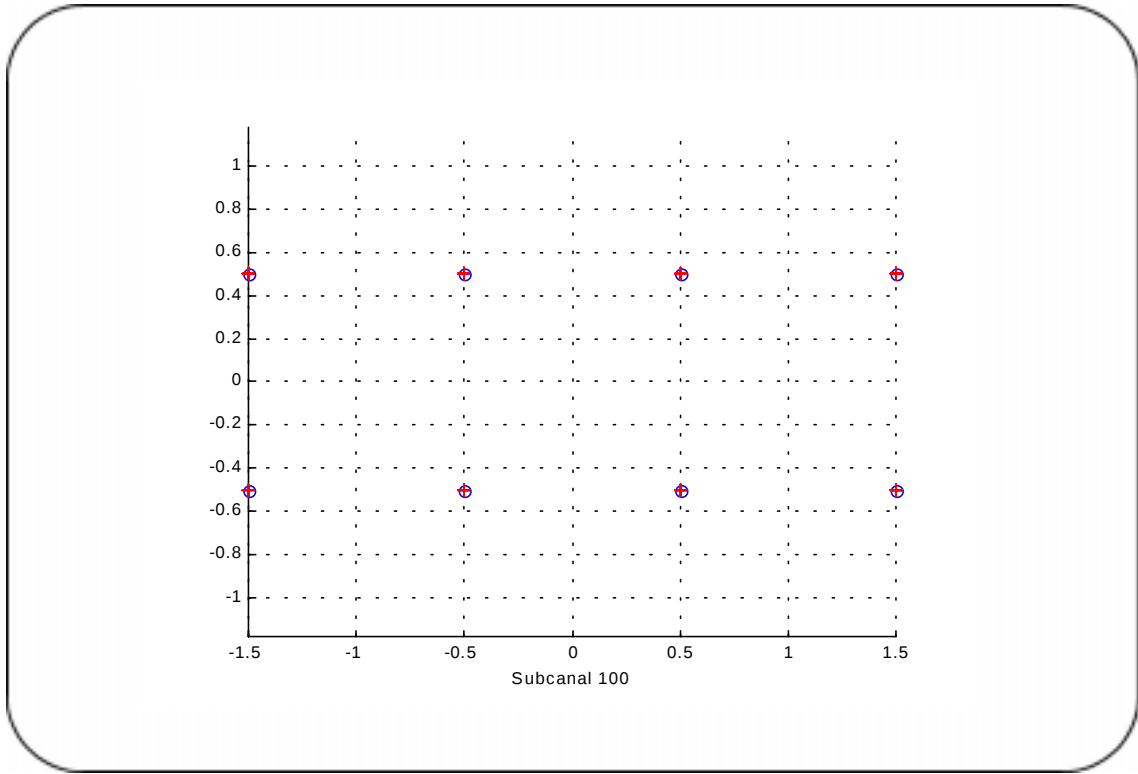


Figura 4.5 : Constelación para un número impar de bits/símbolo

Para un canal 2-D de ganancia compleja H_i (nótese que estamos trabajando con ganancia discretizada en el dominio de la frecuencia, de modo que H_i representa la ganancia compleja en el subcanal i), y con una densidad espectral de potencia de ruido por dimensión igual a σ_i^2 , tenemos la siguiente expresión (válida si la SNR es elevada),

$$SNR_i = \frac{\epsilon |H_i|^2}{2\sigma_i^2}. \quad (4.5)$$

Y la probabilidad de error de símbolo QAM viene dada por

$$P_{e,i} = K_b Q \left[\frac{|H_i|d}{2\sigma_i} \right]. \quad (4.6)$$

El factor K_b es un factor regulador de la tasa de error en función de b . Un estudio realizado por Irving Kalet muestra lo ajustado de la aproximación [13], y justifica el empleo de $k_b=4$ para las aplicaciones que trabajen con un número de bits que oscile entre 2 y 12 por símbolo QAM [13]. Como ya veremos, este perfil de hipótesis se adaptará perfectamente a nuestro caso.

En la situación en la que nos encontramos ahora, podemos calcular el número máximo de bits por símbolo que puede ser transmitido por un canal, cuando la p_e del sistema y la energía disponible del transmisor están fijadas. Igualando la distancia en la constelación en la expresión (4.4) con la de la distancia según (4.6) podemos llegar a

$$b_i = \log_2 \left(1 + \frac{3SNR_i}{\left[Q^{-1} \left(\frac{P_e}{K_b} \right) \right]^2} \right). \quad (4.7)$$

Como vemos, cuanto mayor SNR tengamos o mayor probabilidad de error permitamos, mayor número de bits podremos transmitir con la misma energía en el transmisor. La tasa de bits en cada subcanal vendrá dada por la relación

$$R_b = \frac{b_i}{T_s}. \quad (4.8)$$

donde T_s es el tiempo de duración de símbolo QAM.

4.4.2 - Análisis de un canal del sistema. El "gap" y el "margen"

Ahora vamos a analizar el caso de un canal de comunicaciones que se divide mediante MCM en N subcanales complejos, 2-D. Nuestra forma de trabajar será partir del análisis de un solo canal, de forma independiente, para después extrapolar para el sistema multicanal completo.

Así pues, vamos a tratar por el momento con un subcanal de ganancia compleja H_i y caracterizado por una densidad espectral de potencia de ruido de σ_i^2 . Como sabemos, la SNR a la entrada del receptor viene dada por la expresión

$$SNR_i = \frac{\epsilon |H_i|^2}{2\sigma_i^2}. \quad (4.9)$$

Según el teorema de Shannon, la capacidad por dimensión de un canal (o máxima información que el canal es capaz de transmitir, es decir, la información mutua entre la entrada y la salida, medida en bits), viene dada por ([8], [13])

$$\overline{C}_i = \frac{1}{2} \log_2(1 + SNR_i). \quad (4.10)$$

Cualquier sistema realizable tiene que funcionar transmitiendo una cantidad de información inferior a su capacidad. Esta diferencia es la que nos lleva a definir el "gap" de un sistema, el cual se define como

$$\Gamma = \frac{2^{2\overline{C}} - 1}{2^{2\overline{b}} - 1} = \frac{SNR}{2^{2\overline{b}} - 1}. \quad (4.11)$$

La existencia de un "gap" de X dB en un sistema equivale a pensar que el sistema está transmitiendo una cantidad de información igual a la capacidad que tendría un sistema con una SNR X dB menor

En general, podemos afirmar que el "gap" de un sistema depende sólo del método de codificación empleado, y no del número de bits por símbolo. Así, cuando empleamos codificaciones mejores, lo que hacemos es reducir el "gap" del sistema por un factor que denominamos "ganancia del código". Cuanto menor sea el "gap", mejor es el sistema, ya que nos acercamos a su capacidad de transmisión máxima teórica. Por tanto, si se emplea un código con ganancia γ_c dB, el "gap" pasa ser $(\Gamma - \gamma_c)$ dB. Esta interpretación del "gap" se ve claramente al rescribir la expresión anterior en la forma:

$$\bar{b} = \frac{1}{2} \log_2 \left(1 + \frac{SNR}{\Gamma} \right) \quad (4.12)$$

Para calcular el "gap" con nuestro sistema de codificación QAM, basta con combinar (4.7) y (4.12), teniendo en cuenta que al ser un canal bidimensional,

$$b = 2\bar{b} . \quad (4.13)$$

Por tanto,

$$\log_2 \left(1 + \frac{3 \text{ SNR}}{\left[Q^{-1} \left(\frac{P_e}{K_b} \right) \right]^2} \right) = 2 \cdot \frac{1}{2} \log_2 \left(1 + \frac{SNR}{\Gamma} \right), \quad (4.14)$$

de donde podemos deducir que la expresión para el "gap" es:

$$\Gamma = \frac{1}{3} \left[Q^{-1} \left(\frac{P_e}{K_b} \right) \right]^2 . \quad (4.15)$$

Podemos apreciar claramente cómo una vez que hemos seleccionado la probabilidad de error para el sistema, el gap se mantiene constante para todos los valores de b (la expresión no depende de este parámetro). Esto resulta particularmente útil para los sistemas de transmisión multicanal, que

tienen que enviar un número de bits diferente por cada subcanal pero que emplean el mismo esquema de codificación y la misma probabilidad de error deseada.

Ahora vamos a hablar del "margen", que es otro parámetro importante a tener en cuenta en el análisis de este tipo de sistemas. Si el "gap" era una medida de lo que nuestro esquema de codificación nos separaba del máximo teórico, el "margen" va a ser un parámetro que introduzca a propósito el transmisor para proteger el mensaje frente a ruido inesperado. Vamos a definir el margen como la cantidad en la que puede ser reducida la SNR y aun así, mantener la probabilidad de error dentro del límite fijado. En un sistema con "margen", podemos escribir que

$$\bar{b} = \frac{1}{2} \log_2 \left(1 + \frac{SNR}{\gamma_m} \right), \quad (4.16)$$

donde γ_m es el "margen" utilizado. Nótese que un margen negativo (en dB) indicaría que la SNR en el sistema debe ser aumentada para poder cumplir con los requisitos de p_e . Con el "gap" y el "margen" expresados en dB, tenemos en conjunto que

$$\bar{b} = \frac{1}{2} \log_2 \left(1 + \frac{SNR}{\Gamma \gamma_m} \right) \quad . \quad (4.17)$$

El margen también se puede entender como una medida de la calidad del sistema. En efecto, si forzamos a que nuestro sistema trabaje con una tasa binaria fija (fijada de antemano), el margen mediría la cantidad en que puede ser empeorada la SNR para seguir teniendo la probabilidad de error dentro del límite prefijado.

4.4.3 – El sistema multicanal

En este apartado vamos a extrapolar el análisis hecho hasta el momento para un único canal al caso N -dimensional. En el supuesto de que estemos garantizando la independencia entre los subcanales del sistema de modulación multiportadora, podemos ver al sistema como un conjunto de N subcanales independientes entre sí (sin ISI). Este hecho facilita el análisis.

Tratándose de un sistema de transmisión digital, la figura que utilizaremos para medir las prestaciones del mismo será la probabilidad de error de símbolo, que calcularemos como la media de las probabilidades de error en cada uno de los subcanales. Como ya vimos, la minimización de la probabilidad de error del sistema pasa por forzar en todos los subcanales la misma probabilidad de error, y este será el criterio que impondremos en nuestras simulaciones [13]. Sin embargo, se está investigando sobre las ventajas que aportaría un reparto desigual de las probabilidades de error en los subcanales; la idea es utilizar los subcanales con mejores probabilidades de error para el transporte de la información más crítica.

Una forma de conseguir probabilidades de error iguales es utilizar el mismo tipo de código, lo cual implica, según hemos visto, un “gap” constante. En este caso, podemos definir un parámetro característico de los sistemas multicanal conocido como relación señal a ruido geométrica, que puede compararse con la relación señal a ruido presente en un sistema de portadora única con igualación.

Sea un conjunto de N canales complejos bidimensionales. El número total de bits transmitidos en un símbolo es

$$b_{tot} = \sum_{i=1}^N b_i = \sum_{i=1}^N \log_2 \left(1 + \frac{SNR_i}{\Gamma_i} \right). \quad (4.18)$$

Se va a definir la SNR geométrica, SNR_{geo} , como el valor que cumple

$$b_{tot} = N \log_2 \left(1 + \frac{SNR_{geo}}{\Gamma} \right). \quad (4.19)$$

Nótese que hemos supuesto que $\Gamma_i = \Gamma$ para todos los subcanales. Ya hemos visto que en nuestro sistema, se cumple esta condición. Igualando (4.18) y (4.19) llegamos a

$$SNR_{geo} = \Gamma \left\{ \left[\prod_{i=1}^N \left(1 + \frac{SNR_i}{\Gamma} \right) \right]^{\frac{1}{N}} - 1 \right\}, \quad (4.20)$$

desde donde despreciando los términos $1 +$ y -1 , obtenemos

$$SNR_{geo} = \left[\prod_{i=1}^N (SNR_i) \right]^{\frac{1}{N}}, \quad (4.21)$$

expresión que justifica la denominación de relación señal a ruido geométrica. La SNR_{geo} obtenida puede ser comparada directamente con la SNR detectada para un sistema de portadora única, como vemos en la forma de la expresión (4.19).

Al fin hemos conseguido un parámetro que permite hacer comparaciones entre los sistemas multiportadora y los sistemas de portadora única. No obstante, nosotros vamos a encontrarle otra utilidad adicional: su uso en el estudio de uno de los igualadores que vamos a estudiar.

4.5 - LA CARGA DE LOS SUBCANALES

4.5.2 - Introducción

En este apartado vamos a estudiar cómo se lleva a cabo el proceso de asignación de bits a los diferentes subcanales. Por supuesto, existen muchas formas de llevar a cabo este proceso, pero sólo una será la óptima para

cumplir un determinado objetivo. En nuestro caso, el objetivo que se debe cumplir es el siguiente: maximizar la tasa binaria en el sistema manteniendo la probabilidad de error de bit inferior a un cierto valor p_e fijado de antemano. Además, para llegar a una solución coherente, también debemos limitar la energía máxima disponible en el transmisor (para cada símbolo) a $E_{m\acute{a}x}$.

Para algunos de los sistemas de comunicación, de los que utilizan modulación multiportadora en su implementación, es posible calcular la carga de los subcanales en función de las características del canal. Esta situación es la que se va a dar en los sistemas xDSL en los que se realiza una fase de entrenamiento antes de cada transmisión para determinar con precisión la capacidad que puede ser transportada en cada subcanal. Para que esto sea posible, debe poder habilitarse un camino de realimentación desde el receptor al transmisor, con el fin de poner en conocimiento de este último los parámetros medidos durante la fase de entrenamiento (que serán la atenuación que presenta el medio de transmisión en cada frecuencia y la potencia de ruido y su distribución en el espectro). Sin embargo, para otras aplicaciones (como las de radiodifusión) no es posible habilitar este canal de realimentación, por lo que no podemos utilizar un algoritmo de carga adaptativo. En estos casos, la estrategia consistiría en configurar la transmisión de una forma suficientemente robusta frente al ruido y las interferencias por propagación multitrayecto que pudieran generarse.

Las conclusiones a las que llegamos en el apartado anterior nos permiten deducir que en un sistema MCM son dos los parámetros con los que tenemos que jugar para llevar a cabo el proceso de asignación de bits: por un lado, elegir los coeficientes m_i (número de bits por símbolo que se transmiten por cada subcanal); y por otro, elegir los coeficientes γ_i correspondientes a las fracciones de la $E_{m\acute{a}x}$ que se utilizan para transmitir en cada subcanal. Puede demostrarse que la tasa de bits del sistema MCM se maximiza cuando las variables (m_i, γ_i) se escogen de tal forma que las probabilidades de error de bit en cada subcanal sean iguales.

Por otra parte, la probabilidad de error de bit en un subcanal depende tanto de la SNR presente en el subcanal como del número de

símbolos (y la energía media con la que son transmitidos) disponibles para la transmisión. Por tanto, en un canal con menor SNR, tenemos que aumentar el coeficiente γ_i (enviar los símbolos con una mayor energía promedio) o bien disminuir m_i (enviar menos bits por símbolo en el subcanal). De todas formas, nosotros no podemos conocer las SNR en los subcanales a priori, ya su valor dependerá de los coeficientes γ_i ; por tanto, tenemos que trabajar con un parámetro llamado SNR normalizada, que viene dado por

$$SNR_{i-norm} = \frac{|H_i|^2}{\sigma_i^2} . \quad (4.22)$$

En la ecuación (4.22) H_i representa la ganancia compleja del canal en el subcanal i , y σ_i es la densidad espectral de potencia de ruido en ese subcanal.

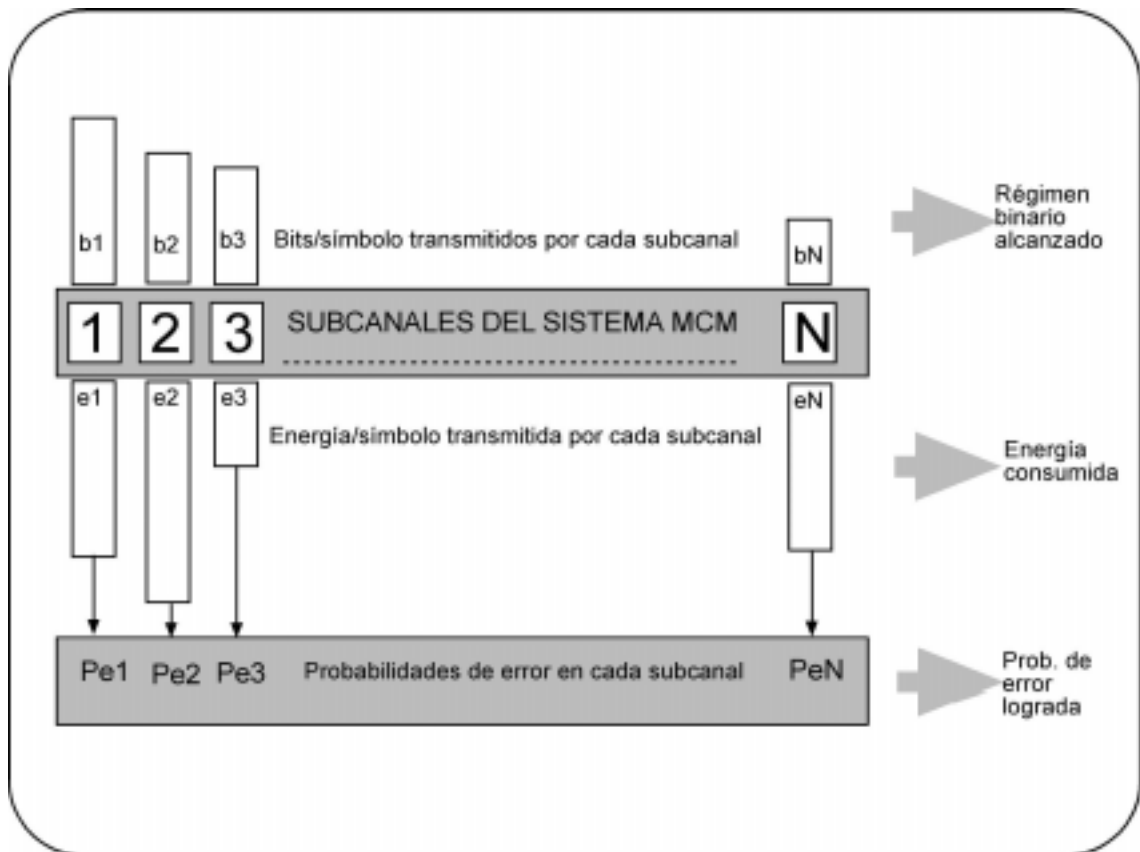


Figura 4.6 : Resultado de un algoritmo de asignación

En la figura 4.6 podemos ver el resultado de un algoritmo de asignación. Teniendo en cuenta la SNR normalizada existente en cada

subcanal, se generan dos conjuntos de valores: por un lado, los bits por símbolo que se van a transmitir por cada subcanal, y por otro, la energía media que se va a emplear en la transmisión por dicho subcanal. En cada subcanal, podemos imponer una probabilidad de error deseada. Cuando concluye el algoritmo, podemos conocer la tasa binaria alcanzable en el sistema, el coste en energía que dicho proceso supone y la probabilidad de error teórica del sistema multiportadora, que puede definirse como la media de las probabilidades de error alcanzadas en cada subcanal.

Los pasos concretos de un algoritmo de asignación se estudiarán más adelante, en el apartado 4.5.3, cuando se aborde el estudio del algoritmo que se empleará en las simulaciones. Dicho algoritmo es el llamado “Algoritmo óptimo discreto”, que es un algoritmo de carga adaptativa caracterizado por proporcionar soluciones “admisibles”, es decir, un algoritmo que siempre asigna valores enteros a los valores de bits/símbolo. (El algoritmo óptimo recibe el nombre de “water filling”, y en general no proporciona valores enteros, por lo que su implementación, aunque posible, es muy costosa).

4.5.3 – El Algoritmo Óptimo Discreto

Como ya hemos adelantado, la asignación de bits a subcanales en el sistema DMT se llevará a cabo con el llamado “algoritmo óptimo discreto”. Consideremos un cierto subcanal componente de un sistema de transmisión DMT, en el cual existe una cierta SNR normalizada (SNR cuando el transmisor genera una señal de energía unidad).

Una vez que hemos fijado el esquema de codificación (en nuestro caso QAM), y teniendo en cuenta el valor de la SNR normalizada del subcanal, existe una relación unívoca entre el número de bits por símbolo que se transmite por ese subcanal y la energía promedio con la que hay que transmitir esos símbolos, siempre que tengamos fijada la probabilidad de error de símbolo a un cierto valor. En concreto, si queremos aumentar el número de bits por símbolo, tenemos que transmitir con una mayor energía promedio (la constelación QAM aumentaría de tamaño, ya que para garantizar la misma

probabilidad de error habría que mantener las distancias relativas entre sus puntos).

Cuando en un sistema DMT sea necesario transmitir un bit más por cada símbolo DMT, el algoritmo óptimo discreto asignará dicho transporte al subcanal que requiera un menor incremento de energía para transportarlo. Veamos, no obstante, los pasos concretos del algoritmo.

4.5.3.1 – El algoritmo óptimo discreto

Definamos para un sistema multicanal de N canales el vector de asignación de bits como

$$\mathbf{b} = [b_1 \ b_2 \ b_3 \ \dots \ b_N] , \quad (4.23)$$

siendo cada b_i el número de bits por símbolo que transporta el subcanal i . Es obvio que la suma de todos los elementos de bits es igual al número de bits que se transmiten en cada símbolo DMT.

Para una determinada probabilidad de error, la energía promedio que hay que transmitir por un subcanal es una función monótona creciente con el número de bits por símbolo que se transmiten por ese subcanal. Es decir, cuantos más bits por símbolo queramos transmitir, mayor energía necesitaremos en el proceso. Podemos pues, expresar

$$e_i = e_i(b_i) , \quad (4.24)$$

donde e_i es la energía empleada en el subcanal i .

Definamos ahora otra función: el incremento de energía. Para un determinado subcanal, el incremento de energía para transportar b_i bits es la cantidad de energía adicional necesaria para transportar esos b_i bits, con respecto a la energía necesaria para transportar b_i-1 bits. Podemos pues, expresar

$$inc_i(b_i) = e_i(b_i) - e_i(b_i - 1). \quad (4.25)$$

Recordando la expresión (4.22) para la SNR normalizada, podemos construir ya una SNR (sin normalizar), multiplicándola por la energía transmitida. Así pues, para cada subcanal podemos escribir,

$$SNR_i = \frac{e_i |H_i|^2}{2 \sigma_i^2} = \frac{e_i g_i}{2}, \quad (4.26)$$

donde g_i es la SNR normalizada. El factor 2 en el denominador se debe a que estamos considerando ahora un canal de transmisión bidimensional. El número máximo de bits que pueden transmitirse por ese subcanal complejo, para una probabilidad de error y una SNR determinadas, viene dado por la expresión (4.7),

$$b_i = \log_2 \left(1 + \frac{3SNR_i}{\left[Q^{-1} \left(\frac{P_e}{K_b} \right) \right]^2} \right). \quad (4.27)$$

Recordamos también la expresión (4.15), que relacionaba el “gap” en el sistema con la probabilidad de error deseada,

$$\Gamma = \frac{1}{3} \left[Q^{-1} \left(\frac{P_e}{K_b} \right) \right]^2. \quad (4.28)$$

De las expresiones anteriores, (4.26), (4.27) y (4.28), obtenemos la energía necesaria para transportar b_i bits por el subcanal i ,

$$e(b_i) = 2 \frac{\Gamma}{g_i} (2^{b_i} - 1). \quad (4.29)$$

Nótese que si el valor de b_i es elevado, para transportar un bit más por el subcanal tenemos que duplicar la energía invertida en el proceso. En cuanto a los incrementos de energía, podemos afirmar que

$$inc_i(b_i) = \frac{\Gamma}{g_i} 2^{b_i} = 2inc_i(b_i - 1), \quad (4.30)$$

es decir, que la transmisión de un bit más por un subcanal exige un aumento de 3 dB en la energía empleada.

Una vez que se ha escogido el sistema de codificación y se ha fijado la probabilidad de error deseada para cada uno de los subcanales, se procede a calcular los incrementos de energía para todos los subcanales y para todos los números de bits posibles. Dichos incrementos se tabularán para su uso en el algoritmo de asignación. Los bits se van asignando de uno en uno, escogiendo en cada momento el subcanal que requiera un menor incremento de energía para transportarlo. La asignación finaliza cuando se acaba la energía disponible en el transmisor, parámetro que debe estar fijado de antemano. Veamos cómo se desarrolla una iteración del algoritmo:

1. Si aún se dispone de energía en el transmisor, se acude a las tablas de incrementos de energía, y se busca el subcanal que requiera un menor incremento de energía para transportar un bit más de los que ya transporta. Supongamos que es el subcanal n .
2. Si el incremento de energía requerido en el subcanal n es inferior a la energía disponible en el transmisor, E_{dis} , se asigna un bit más a ese subcanal

$$b_n \leftarrow b_n + 1 \quad (4.31)$$

y se actualiza el valor de la energía disponible en transmisión

$$E_{dis} \leftarrow E_{dis} - inc_n(b_n). \quad (4.32)$$

En la expresión (4.32), b_n ya ha sido actualizado. Se regresa de nuevo al paso 1.

3. Si, por el contrario, en el transmisor no tenemos energía suficiente para cubrir el incremento necesario, se da por finalizado el proceso de asignación.

Como vemos, este algoritmo garantiza la asignación de un número entero de bits a todos los subcanales. Aunque este algoritmo no es el óptimo, es cierto que la restricción de transmitir en cada sub-símbolo un número entero de bits no supone una degradación importante en las prestaciones del sistema.

4.5.3.2 – Restricciones impuestas al algoritmo

Antes de cada transmisión, el modem DMT tiene que llevar a cabo el proceso de asignación de bits, de acuerdo con el algoritmo anterior. Para ello, se supone conocida la existencia de la siguiente información:

- La energía máxima disponible en la transmisión de cada símbolo (o bien la potencia de transmisión)
- La probabilidad de error que se desea en el sistema, así como el margen de protección con el que se va a garantizar. Añadir un margen supone forzar en el sistema una SNR mayor que la necesaria, para compensar la aparición de fenómenos de ruido no considerados.
- Las SNR normalizadas de cada subcanal, que, como sabemos, son una medida de la calidad de los subcanales. Los subcanales que presenten menor atenuación y/o menor potencia de ruido podrán transportar más bits por símbolo con menor cantidad de energía, manteniendo constante la probabilidad de error deseada.

Además, vamos a imponer unas restricciones en cuanto al número de bits que se transmiten por cada subcanal.

- El mínimo número de bits por símbolo que transporta cada subcanal es 2. Por tanto, cuando se asignan bits a un subcanal por primera vez, o se le asignan 2 bits de golpe o no se le asigna ninguno. Esta condición es un requisito de la codificación QAM, en la que se necesitan al menos $4=2^2$ símbolos para formar una constelación.
- El máximo número de bits por símbolo que va a transportar un subcanal es 16. Si en el desarrollo del algoritmo, en algún momento el mínimo incremento de energía se encuentra sobre un subcanal que ya está transportando 16 bits, se ignorará dicho mínimo, y se buscará el siguiente, con el fin de que no haya subcanales transportando 17 o más bits. Este límite superior se toma considerando los diversos estudios que se han hecho sobre la implantación de sistemas DMT sobre líneas de cobre (ADSL, VDSL) [13].

Como consecuencia de los puntos anteriormente señalados, el receptor comienza el proceso de asignación, y genera dos vectores de la forma

$$\mathbf{b} = [b_1 \ b_2 \ \dots \ b_N] \quad (4.33)$$

$$\mathbf{e} = [e_1 \ e_2 \ \dots \ e_N], \quad (4.34)$$

que son los vectores de asignación de bits y de distribución de energías, respectivamente. Obviamente, habrá valores de b_i iguales a cero. Esto ocurrirá cuando no se transmita información por los subcanales, bien porque forcemos nosotros esa condición (caso de los primeros subcanales), bien porque el subcanal no alcanza la calidad suficiente.

Estos dos vectores son enviados al transmisor por el camino de realimentación, a una velocidad suficientemente lenta como para garantizar

su perfecta recuperación (es información de importancia capital para el correcto funcionamiento del sistema).

Nótese cómo la suma de los elementos que componen el vector **b** nos sirve para calcular la tasa binaria a la que funcionará el sistema. En efecto, el tiempo de símbolo es T_s , la tasa binaria a la que funcionará el sistema será

$$R_b = \frac{\sum b_i}{T_s}. \quad (4.35)$$

4.6 – EL SISTEMA DMT

Estamos considerando el problema de la transmisión de datos sobre un canal con una considerable distorsión en frecuencia. En un sistema de portadora única, para que esta transmisión se lleve a cabo de una forma fiable, es necesario introducir en el receptor un proceso de igualación. Dicho proceso no produce mejora si el ruido presente en el canal es muy potente, ya que lo realza en las altas frecuencias (un igualador de este tipo, para el canal telefónico, tiene una característica HP). Además, si la atenuación en las frecuencias altas es elevada, la implementación del igualador puede ser muy difícil.

Como alternativa, en este capítulo se mostrada la modulación multiportadora. Al dividir el ancho de banda en un gran número de subcanales, podemos asumir que la respuesta del canal es plana en el ámbito de cada uno de ellos, con lo que se elimina la necesidad del igualador (para cada subcanal). Es evidente que con la modulación multiportadora hemos solucionado este problema; pero, sin embargo, hemos creado otro nuevo: ahora para que todo esto funcione hay que asegurar la perfecta independencia entre los subcanales. En el sistema multitono básico que se presentó al principio del capítulo, si las portadoras utilizadas eran ortogonales, se aseguraba su perfecta separación en el receptor. Puede demostrarse que con

unas portadoras de este tipo, combinadas con el algoritmo “water filling” (algoritmo óptimo para cargar los subcanales) y el empleo de un número infinito de subcanales, se consigue llegar al punto óptimo de funcionamiento.

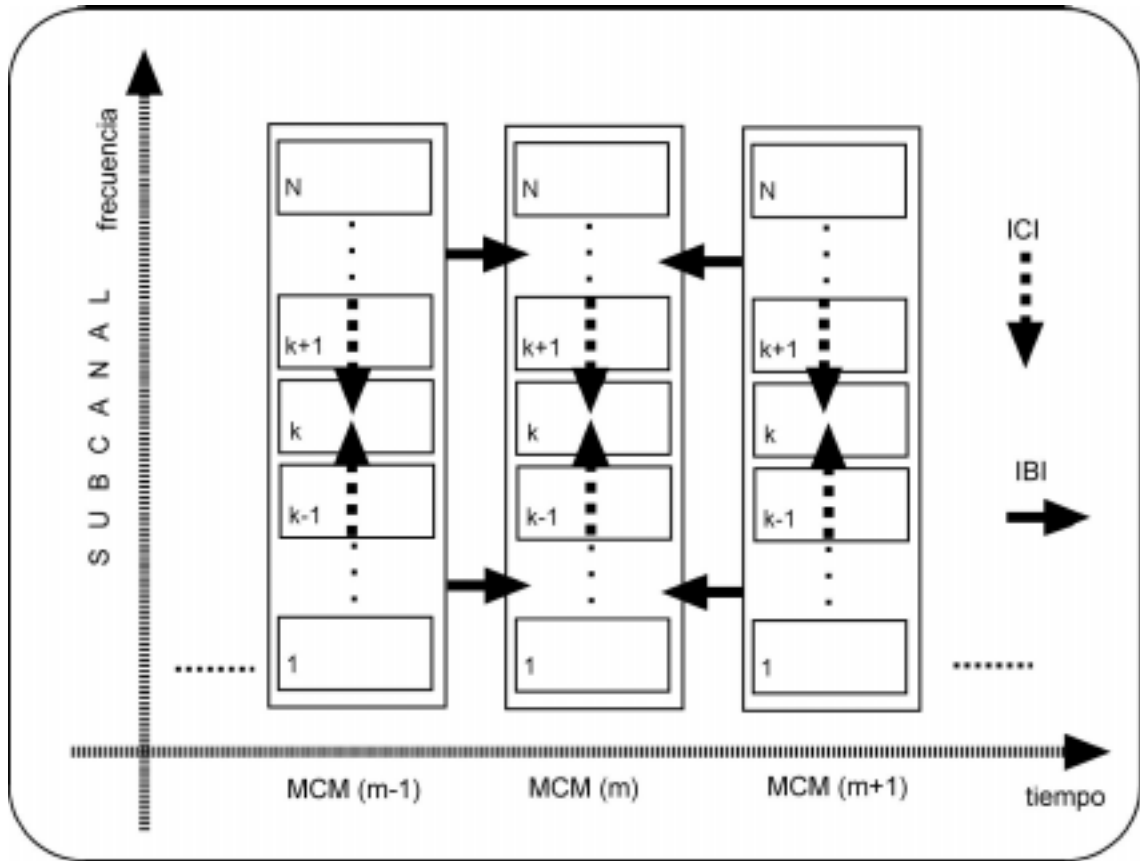


Figura 4.7 : Ilustración de los posibles tipos de ISI

Otro de los efectos que se producen en un sistema multitono es la llamada IBI (interferencia entre bloques). En un sistema MCM, se transmiten secuencialmente símbolos, cada uno de los cuales se compone de M bits, distribuidos sobre un conjunto de N subcanales. Llamaremos ICI al fenómeno de interferencia entre portadoras, es decir, la interferencia que un canal sufre procedente de otro, pero dentro del mismo símbolo MCM. Por otro lado, llamaremos IBI al fenómeno de interferencia entre símbolos MCM, es decir, la interferencia que un símbolo MCM sufre procedente de los símbolos inmediatamente anterior y posterior. En la figura 4.7 se muestran gráficamente los distintos tipos de ISI presentes en un sistema MCM [3].

Para evitar la ICI, tendremos que confinar el espectro de los subcanales en una banda estrecha. Esto implica extenderlos sobre un gran

intervalo en el dominio del tiempo, por lo que provocarán mucha IBI. Y si restringimos la IBI usando portadoras que se extiendan sólo durante el periodo de duración del símbolo MCM, provocaremos un aumento de su ancho de banda, con lo que obtendremos ICI. Además, estamos obligados a alejarnos del caso ideal, puesto que no podemos utilizar un número infinito de subcanales ni el algoritmo “water filling”. Todas estas consideraciones, son las que desembocan en el desarrollo de un método de división del canal que funcione de forma eficiente y que elimine la posibilidad de ISI* (en forma de IBI o ICI).

Diversos estudios han demostrado que el método óptimo utiliza “portadoras” que son dependientes del canal de transmisión, y están relacionadas con los autovectores de una cierta matriz de covarianza del canal. En su versión discreta, este método recibe el nombre de VC (vector coding). El estudio de los pormenores de este método de división del canal no entra dentro de los objetivos del proyecto y no se va a llevar a cabo. Lo que sí se va a estudiar (profundizando sólo lo necesario) es DMT, que es una variación de VC añadiendo ciertas restricciones que permiten simplificar su implementación. Los dos pilares sobre los que se apoya DMT son:

- El uso de la DFT como herramienta para modular los datos (con el consiguiente uso de algoritmos FFT para su cálculo). Esto consiste en diseñar las constelaciones QAM de cada subcanal en el dominio de la frecuencia, y posteriormente, hacer la IFFT de los símbolos complejos de la constelación para pasarlos al dominio del tiempo y transmitirlos.
- El uso de un CP (*cyclic prefix* , prefijo cíclico) en cada símbolo DMT transmitido. Sea $2N$ el número de muestras temporales necesario para transmitir un símbolo DMT (formado por M bits distribuidos sobre N subcanales, ya veremos por qué se cumple esta relación). El CP consiste en seleccionar las v últimas muestras temporales del símbolo (ya veremos como ajustar el valor de v) y anteponerlas a su transmisión. De esta manera, la

* En realidad, lo más que podemos hacer es minimizar su efecto.

transmisión de un símbolo exige ocupar $2N+v$ muestras temporales. Es obvio que usar el CP equivale a desperdiciar ancho de banda y energía (es una redundancia que se transmite), pero es el precio que tenemos que pagar si queremos controlar los fenómenos de IBI. Estudiaremos, no obstante, cómo las técnicas de igualación permiten disminuir la ineficiencia producto del empleo del CP.

3.6.1 – Breve repaso de la DFT

La DFT (*Discrete Fourier Transform*) es una herramienta que permite el análisis en frecuencia de señales discretas. Si $\mathbf{x}$ es una señal discreta en el dominio del tiempo,

$$\mathbf{x} = \begin{bmatrix} x_0 \\ x_1 \\ \dots \\ x_{N-1} \end{bmatrix}, \quad (4.36)$$

formada por N muestras; al aplicar la DFT sobre ella obtenemos otra señal $\mathbf{X}$ también discreta, también formada por N muestras:

$$\mathbf{X} = \begin{bmatrix} X_0 \\ X_1 \\ \dots \\ X_{N-1} \end{bmatrix}, \quad (4.37)$$

y donde cada componente de $\mathbf{X}$ se calcula como [13]

$$X_n = \frac{1}{\sqrt{N}} \sum_{k=0}^{N-1} x_k e^{-j \frac{2\pi}{N} kn}, \quad n = 0, 1, \dots, N-1. \quad (4.38)$$

La correspondiente transformada inversa se calcularía como

$$x_n = \frac{1}{\sqrt{N}} \sum_{k=0}^{N-1} X_k e^{j \frac{2\pi}{N} kn}, \quad n = 0, 1, \dots, N-1. \quad (4.39)$$

Esto puede verse como una transformación lineal, si expresamos las anteriores ecuaciones en la forma matricial (3.24).

$$\mathbf{X} = \mathbf{Q}\mathbf{x} \quad (4.40)$$

$$\mathbf{x} = \mathbf{Q}^H \mathbf{X}, \quad (4.41)$$

donde $\mathbf{Q}$ es una matriz ortogonal que viene dada por

$$\mathbf{Q} = \frac{1}{\sqrt{N}} \begin{bmatrix} 1 & 1 & 1 & \dots & 1 \\ 1 & e^{-j \frac{2\pi}{N}} & e^{-j \frac{2\pi}{N} 2} & \dots & e^{-j \frac{2\pi}{N} (N-1)} \\ 1 & \vdots & \vdots & \vdots & \vdots \\ 1 & e^{-j \frac{2\pi}{N} (N-2)} & e^{-j \frac{2\pi}{N} 2(N-2)} & \dots & e^{-j \frac{2\pi}{N} (N-1)(N-2)} \\ 1 & e^{-j \frac{2\pi}{N} (N-1)} & e^{-j \frac{2\pi}{N} 2(N-1)} & \dots & e^{-j \frac{2\pi}{N} (N-1)(N-1)} \end{bmatrix}. \quad (4.42)$$

Algunos autores [16] eliminan el factor $1/\sqrt{N}$ de las expresiones (4.25) y (4.26), y en cambio, incluyen en (4.26) un factor $1/N$. Ambas posibilidades son equivalentes. La notación utilizada aquí facilita la expresión de la DFT como transformación lineal.

Ambas señales ($\mathbf{x}$ y $\mathbf{X}$) pueden considerarse como un periodo dentro de una secuencia N -periódica, si se quieren generar muestras fuera del intervalo $[0, N-1]$.

Los valores tanto de $\mathbf{x}$ como de $\mathbf{X}$ pueden ser complejos; pero si nos ceñimos al caso de señales reales en el dominio del tiempo, las componentes de $\mathbf{X}$ cumplen una propiedad de simetría, de modo que [16]

$$X_{N-n} = X_n^*, \quad n = 0, 1, 2, \dots, N-1. \quad (4.43)$$

Añadamos antes de finalizar este apartado un comentario más sobre la DFT. Las operaciones DFT-IDFT (nótese que son equivalentes) pueden calcularse de forma muy eficiente si hacemos uso de los algoritmos FFT (Fast Fourier Transform) y el número de muestras de las señales es una potencia entera de dos.

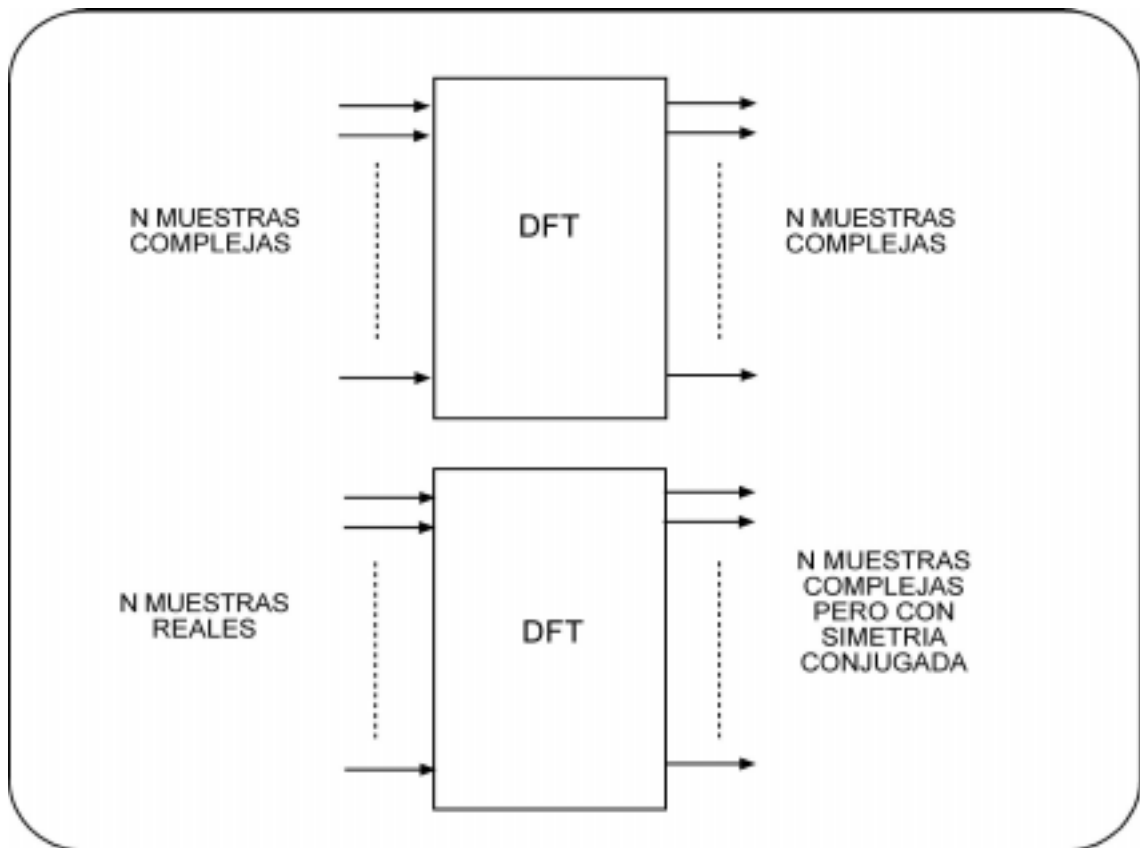


Figura 4.8 : Propiedad de la DFT que va a explotarse en DMT

3.6.2 – Descripción de un enlace DMT completo

En este apartado vamos a analizar ya un sistema xDSL completo, igual que los utilizados en la realidad. Vamos a ir haciendo una breve descripción de lo que hace cada bloque, con el fin de adquirir una rápida visión que nos permita entender los capítulos posteriores, donde complicaremos el sistema con la introducción de los igualadores.

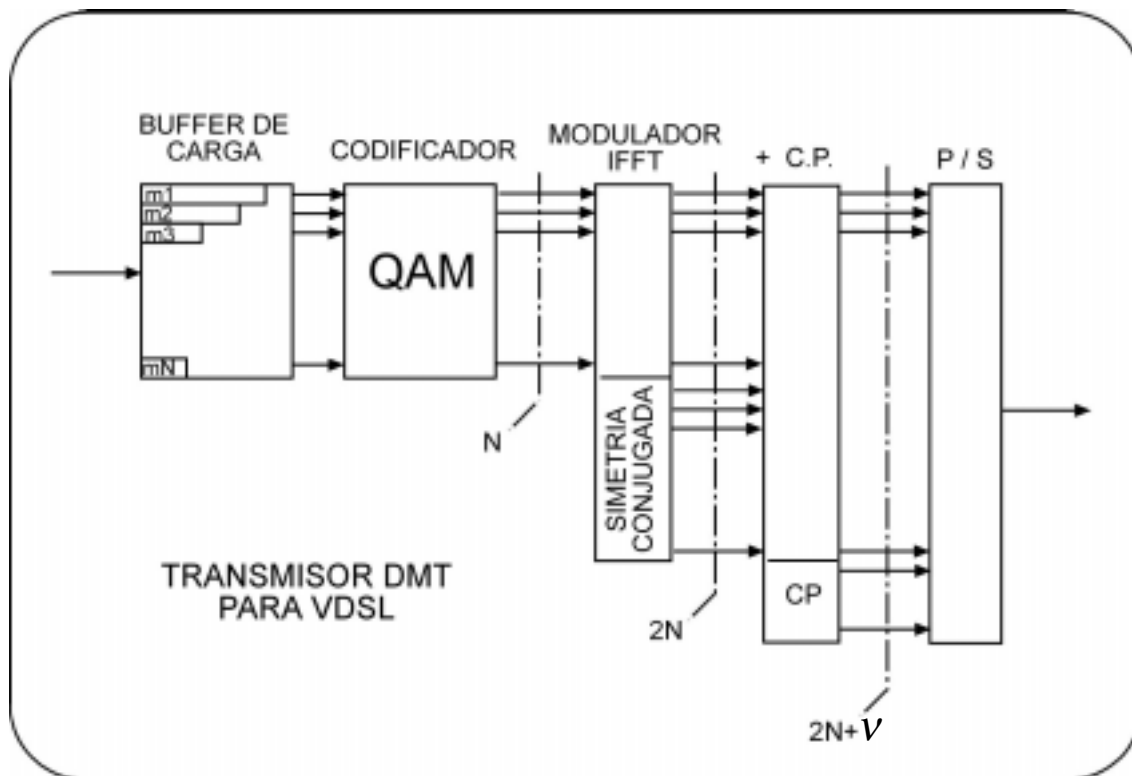
(A) TRANSMISOR

Figura 4.9 : Diagrama de bloques de un transmisor DMT

El sistema DMT que vamos a presentar se muestra en la figura 4.9. La entrada al sistema es la salida de una fuente de datos que genera un flujo de información binaria, una cadena de bits que llegan en serie a la entrada del modem xDSL. A la salida, lo que obtenemos es otro flujo binario, también en serie. Este flujo de salida es lo que el modem va a enviar al canal. Vamos a trabajar con un canal discreto, por lo que obviaremos la descripción de los filtros conformadores a la entrada y a la salida del canal, y el proceso de muestreo que tiene lugar en el receptor. Utilizaremos, por tanto, una función de transferencia discreta (opera directamente sobre las muestras) y consideraremos el ruido como un proceso discreto (muestras que se suman a las de señal). En el modem, entre los flujos binarios de entrada y de salida tienen lugar una serie de procesos que se llevan a cabo sobre la señal que son los que caracterizan al sistema DMT, y que se van a explicar a continuación.

3.6.2.1 – Buffer de carga

El proceso que lleva a cabo en este bloque es una conversión serie/paralelo de los datos. El buffer se va llenando hasta alcanzar los M bits constituyentes del símbolo a transmitir. Estos M bits se dividen en N “paquetes” de b_i bits, con $i=1,2,\dots,N$ (N es el número de subcanales utilizados). El número de bits que constituye cada grupo, b_i , se determina en la fase de entrenamiento, y permanece constante para todos los símbolos DMT de la transmisión. Cada paquete constituye la información que se transmite por cada subcanal. Una vez que el buffer se ha llenado, se pasan los datos al siguiente bloque y se libera el buffer, para ir dando cabida a los nuevos datos que van entrando.

3.6.2.2 – Codificador QAM

Es un bloque que trabaja en paralelo, y de forma independiente para cada una de sus N entradas/salidas. Es decir, el dato generado en la salida i sólo depende del dato presente en la entrada i .

Una vez que ha finalizado la fase de entrenamiento, se han determinado tanto el número de bits por símbolo que se transmite por cada subcanal, b_i , como la energía empleada para efectuar dicha transmisión, e_i . Con esta información, se diseñan N constelaciones QAM, una para cada subcanal. Con el valor de b_i se calcula el tamaño (número de puntos) que tiene que tener la constelación, y el valor e_i nos indica la distancia (separación) de los símbolos en la constelación. Cuanto más separados estén los puntos, más difícil será para el ruido mover el símbolo detectado fuera de la correspondiente región de decisión, por lo que el sistema será más robusto.

Así pues, para cada subcanal se codifica el símbolo binario de entrada, asignándole un punto de la constelación QAM. El valor que se pone en la salida es un número complejo, resultado de proyectar el punto de la constelación sobre los ejes I (Inphase, parte real) y Q (Quadrature, parte imaginaria).

De esta forma, podemos ver que la salida de este bloque es un conjunto de N números complejos totalmente independientes, que no guardan ninguna relación de simetría; en realidad, podemos decir que trabajamos con $2N$ cantidades reales.

3.6.2.3 – Modulador IFFT

Este bloque también recibe N datos en paralelo (cantidades complejas precedentes del codificador QAM); pero genera a su salida $2N$ datos reales. El proceso que se sigue es muy sencillo: Se toman los N datos complejos de entrada y se generan artificialmente otros N datos (aplicando simetría conjugada sobre los datos de entrada). Ahora consideramos los $2N$ datos resultantes como un único vector de tamaño $2N$ que cumple las condiciones de simetría conjugada. Se realiza una operación IFFT sobre este vector y el resultado es otro vector del mismo tamaño pero con componentes reales. Las condiciones de simetría conjugada que debe cumplir el vector generado artificialmente (para obtener datos reales al antitransformar) son las que resultan de aplicar la ecuación (4.43).

Este bloque genera un vector de $2N$ componentes, correspondientes ya a las muestras de la señal a transmitir en el dominio del tiempo.

3.6.2.4 – Adición del “cyclic prefix”

Este bloque lo que hace es introducir un prefijo cíclico (CP). El proceso consiste en formar un nuevo vector de tamaño $2N+v$, insertando antes de las $2N$ muestras originales las v últimas muestras del bloque [13]. En el apartado 3.2.7 estudiaremos los motivos por los que es necesario hacer este procesado, así como la forma de elegir el valor de v .

3.2.6.5 – Bloque 5: Convertidor paralelo/serie

En este último bloque se realiza sobre los datos una simple conversión paralelo/serie. El vector de $2N+v$ muestras se pone en formato serie, y se van transmitiendo secuencialmente todas sus componentes. En la figura 4.10 se muestra el formato de la señal generada por este bloque, así como su transmisión por el canal.

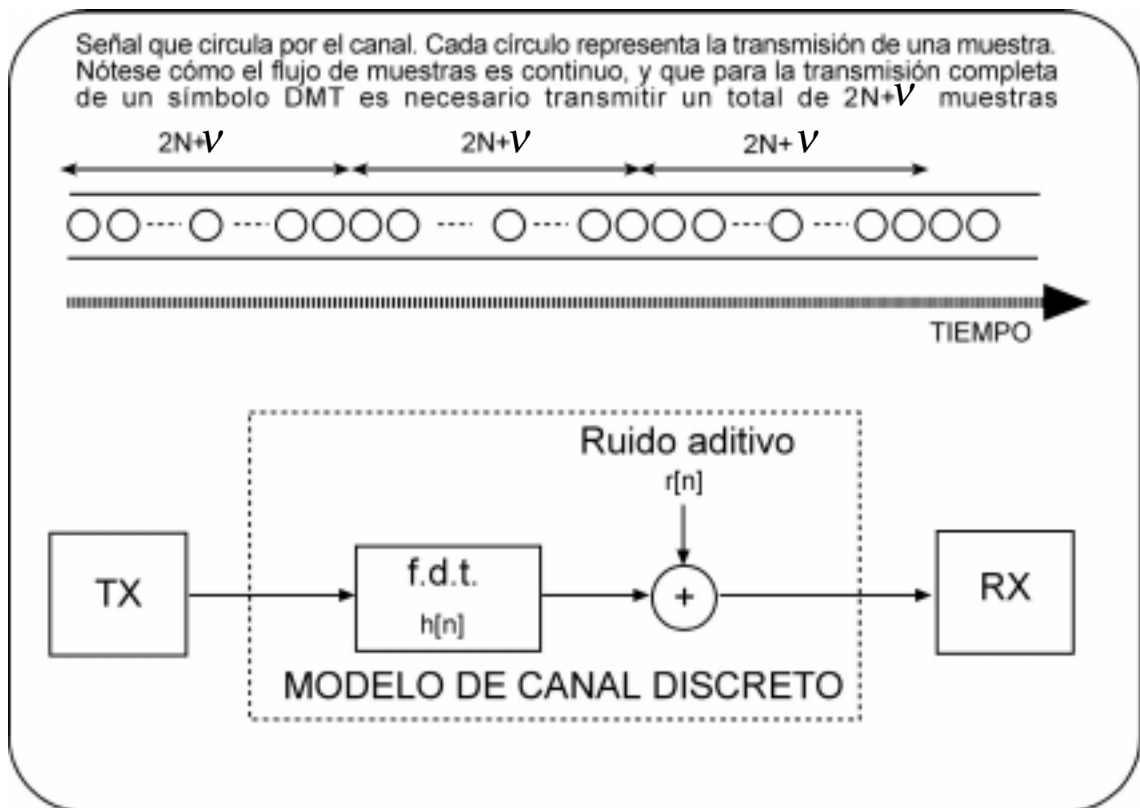


Figura 4.10 : Modelo discreto del canal, y formato de la señal que por él se transmite

Así pues, si observamos una transmisión xDSL a la salida de este bloque, lo que vemos es un flujo binario, en serie, constante, y consistente en la concatenación de bloques de tamaño $2N+v$, cada uno de los cuales transporta un símbolo DMT.

(B) CANAL

3.2.6.6 – Canal discreto

Este bloque se compone de dos sub-bloques:

En el primero, se tiene en cuenta sólo la función de transferencia. El flujo de datos serie procedente del modem transmisor es convolucionado con la respuesta al impulso $h[n]$ del canal de transmisión. Esta $h[n]$ tendrá asociada la función de transferencia del canal $H(z)$, que supondremos conocida (ya estudiamos en el capítulo 3 la forma de obtenerla).

El segundo sub-bloque lo que hace es incluir ruido aditivo en las muestras una vez que han sido convolucionadas con la $h[n]$. Dicho ruido, será generado para que se ajuste a los resultados obtenidos en el capítulo 3, dedicado al estudio del bucle de abonado.

La salida de este bloque es un flujo binario constante, consistente en bloques de tamaño $2N+v$.

(C) RECEPTOR

3.2.6.7 – Convertidor serie/paralelo

Este es el primer bloque del receptor. Recibe como entrada un flujo constante de bits, el que se obtiene a la salida del canal. Dicho flujo se caracteriza por ser de muy baja potencia, dada la gran atenuación que presenta el bucle de abonado.

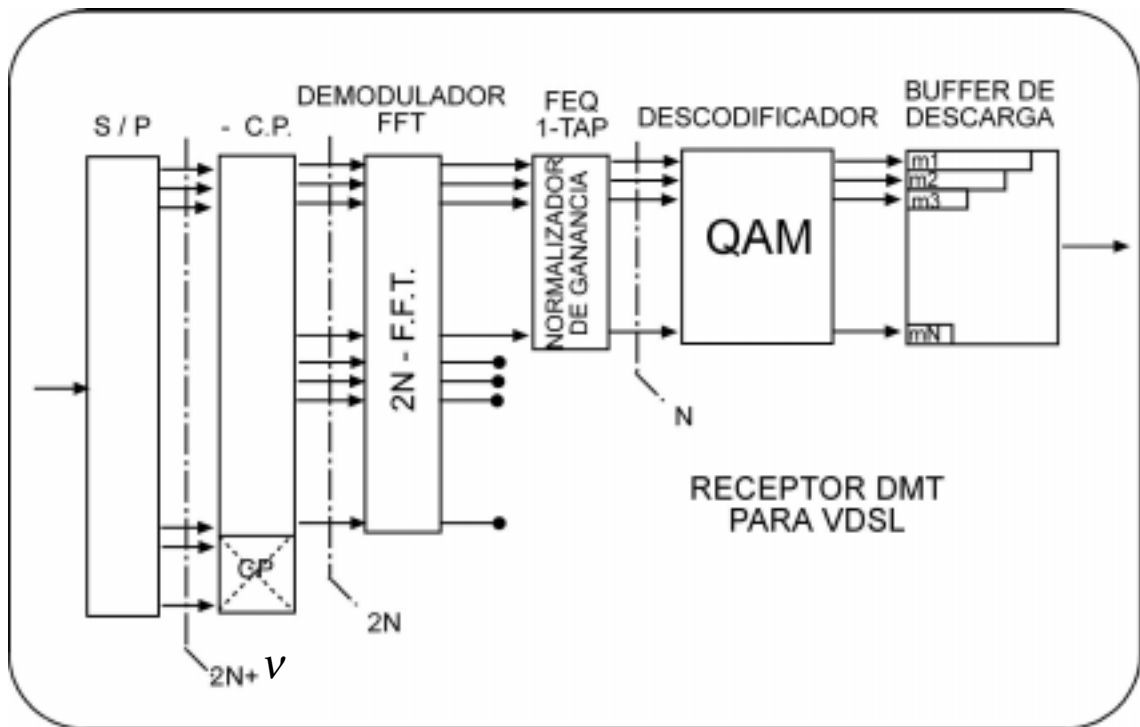


Figura 4.11 : Diagrama de bloques de un receptor DMT

Como sabemos, este flujo está estructurado en tramas de longitud $2N+v$ muestras. Este bloque posee un buffer con capacidad para $2N+v$ muestras temporales (reales), el cual se va llenando a medida que avanza la transmisión. Cuando se llena totalmente, se considera que se han capturado todas las muestras correspondientes a un símbolo DMT, las cuales se ponen a la salida del convertidor en formato paralelo. El buffer se libera para dar paso a las nuevas muestras, correspondientes ya al siguiente símbolo.

3.2.6.8 – Bloque 8: Eliminación del “cyclic prefix”

Este bloque recibe en su entrada datos en formato paralelo, compuestos por $2N+v$ componentes. Este bloque se encarga de eliminar las v primeras componentes del dato de entrada, con lo cual, el bloque pone en su salida un dato en formato paralelo, con $2N$ componentes (las últimas $2N$ componentes del vector de entrada).

3.2.6.9 – Bloque 9: Demodulador FFT

En este bloque se lleva a cabo el proceso de demodulación de las muestras temporales. Ahora los datos están agrupados en bloques de longitud $2N$ muestras, las cuales llegan al demodulador en formato paralelo.

El proceso que se lleva a cabo en este bloque es el de realizar una operación FFT de tamaño $2N$. Este proceso conduce a la obtención de otras $2N$ muestras en el dominio de la frecuencia. Dichas muestras, serán por lo general complejas. Sin embargo, el hecho de que las $2N$ muestras de entrada sean reales, fuerza a que el vector de salida cumpla las condiciones de simetría conjugada vistas anteriormente. Por tanto, nos basta con tomar las N primeras componentes del vector transformado (nótese que al ser muestras complejas, estamos tomando el equivalente a $2N$ muestras reales).

De esta forma, lo que obtenemos a la salida de este bloque es un vector, en formato paralelo, compuesto por N muestras complejas. Dichas muestras representan la información transmitida por cada uno de los N subcanales. La parte real de cada componente está asociada con la componente en fase del subsímbolo transmitido, y la parte imaginaria con la componente en cuadratura.

3.2.6.10 – Normalizador de ganancia

Estamos trabajando con un canal no ideal. Dicho canal da un tratamiento desigual a cada uno de los subsímbolos, en función del subcanal por el que se transmiten. Sin embargo, dividiendo el ancho de banda disponible en un número elevado de subcanales, podemos suponer que el canal se comporta de forma ideal en el ámbito de cada subcanal. Bajo esta hipótesis, podemos afirmar que, en el dominio de la frecuencia, un subsímbolo (complejo) que atraviesa el canal ve modificada a la salida del mismo su magnitud y su fase.

Si calculamos la respuesta en frecuencia (magnitud y fase) del canal, y tomamos en dicha respuesta N muestras equiespaciadas, obtenemos N valores complejos, correspondientes a la ganancia del canal en cada uno de los subcanales.

Dado que la transmisión por los subcanales se realiza de forma independiente, el receptor puede compensar por separado este efecto, multiplicando cada subsímbolo recibido por el inverso de la ganancia compleja del canal a la frecuencia correspondiente. Dicho proceso (multiplicación por un número complejo) provoca una expansión de la constelación (cambio en la magnitud) y un giro de la misma (cambio en la fase); estos cambios contrarrestan los introducidos por el canal.

Este proceso puede entenderse como un proceso muy simple de igualación en el dominio de la frecuencia, por lo que a este bloque también se le conoce como FEQ-1 tap (Frequency-domain Equalizer - 1 coeficiente) [5]

Teniendo en cuenta lo anterior, es claramente visible que la salida de este bloque son otras N muestras complejas, resultantes de aplicar una corrección en magnitud y en fase a las muestras de entrada. Como hemos visto, dicha corrección se realiza de forma independiente en cada subcanal.

3.2.6.11 – Bloque 11: Decodificador QAM

En este bloque tiene lugar la detección de los símbolos recibidos. El procesado se lleva a cabo en paralelo y de forma independiente para cada uno de los subcanales. Para detectar un subsímbolo, se compara cada dato recibido con la constelación QAM utilizada para el subcanal correspondiente. Se considera que el símbolo transmitido (estimado) es el correspondiente a la región de decisión en la que se encuentre el símbolo recibido. Al ser una constelación QAM, las fronteras de las regiones de decisión forman una rejilla rectangular, y se pueden detectar por separado las componentes en fase y en cuadratura.

Una vez que se han detectado los N subsímbolos, se procede a su decodificación. De esta forma, cada salida es un conjunto de m_i bits, verificándose que

$$m_1 + m_2 + \dots + m_N = M , \quad (4.44)$$

donde M es el número de bits que compone un símbolo DMT.

3.2.6.12 – Bloque 12: Buffer de descarga

Este bloque recibe de golpe los N grupos de m_i bits cada uno, y los va descargando para proporcionar a la salida un flujo serie de bits, y continuo. Dicho flujo debe ser el mismo que se introdujo en la cabecera del modem transmisor. Cuando se ha descargado totalmente, recibe de golpe los bits del siguiente símbolo DMT y se procede a una nueva conversión serie/paralelo.

3.2.7 – El “cyclic prefix” (CP, prefijo cíclico)

En este apartado se pretende comprender la utilidad (la necesidad) de un prefijo cíclico en el sistema DMT. Hasta ahora, se ha hablado del CP como lo que realmente es, una sobrecarga que hay que transmitir junto con los datos para que el sistema funcione de manera correcta. La inclusión de esta sobrecarga provoca una disminución del rendimiento del sistema, puesto que del flujo de información que se transmite por el bucle de abonado, el rendimiento que se obtiene es

$$\eta = \frac{2N}{2N + v} \leq 1. \quad (4.45)$$

En el análisis del enlace DMT completo que se ha realizado, se ha visto cómo los procesos de introducción y eliminación del CP se realizan de forma sencilla y natural, y que apenas introducen complejidad computacional en el proceso completo.

Ahora es el momento de presentar la razón por la que es necesario incluir este elemento en el sistema: asegurar que los símbolos DMT se transmitan de forma independiente unos de otros, es decir, que cada bloque de $2N$ muestras temporales de señal útil no influya sobre los bloques correspondientes a los símbolos DMT transmitidos anterior y posteriormente. La utilidad del CP, por tanto, es la de eliminar la IBI (*Inter-Block Interference*).

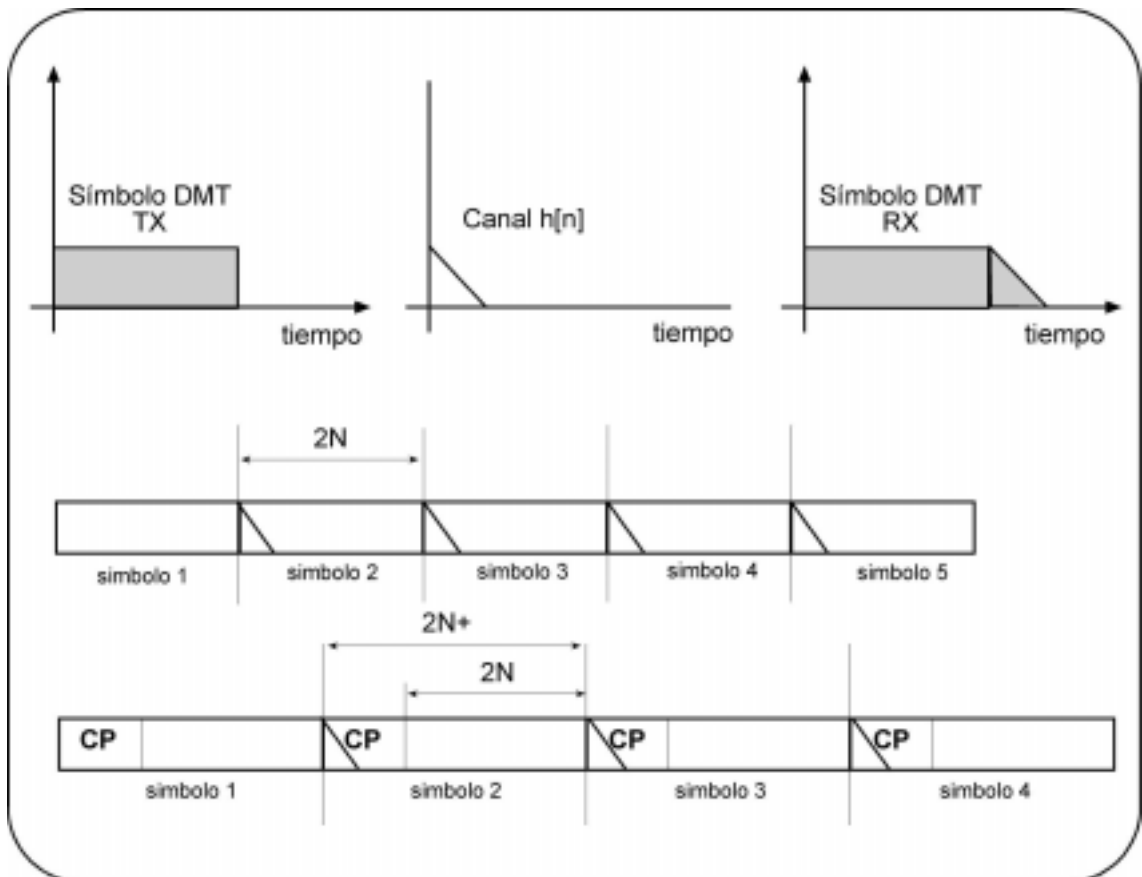


Figura 4.12 : Comparación de las transmisiones sin CP y con él

Considérese la transmisión de una señal compuesta únicamente por un solo símbolo DMT, sin la inclusión del CP. Teniendo en cuenta todos los pormenores del procesamiento de señal llevado a cabo en el transmisor, ya estudiado en el apartado anterior, la señal que se envía al canal se compone de $2N$ muestras reales. No se va a considerar el ruido aditivo, puesto que no es necesario para comprender el CP. Supongamos ahora que el canal de transmisión (canal digital, como venimos considerando hasta ahora) tiene una respuesta impulsiva de longitud L muestras. Por un simple cálculo de convolución, sabemos que la señal a la salida del canal tendrá una longitud de

$2N+L-1$ muestras. Como la transmisión se compone de un único símbolo DMT, el receptor no tiene problemas en realizar la asociación muestras_recibidas/símbolo_transmitido.

Considérese ahora otra segunda transmisión, realizada en las mismas condiciones que la anterior, pero transmitiendo una señal distinta, compuesta esta vez por 2 símbolos DMT. La señal que se envía al canal tendrá una longitud de $2N+2N=4N$ muestras. Por tanto, la señal de salida del canal tendrá una longitud de $4N+L-1$ muestras. ¿Cómo recupera ahora el receptor los 2 bloques de $2N$ muestras cada uno que han sido transmitidos? Resulta evidente que se ha producido un fenómeno de ISI entre ambos bloques de señal (esta es la forma de ISI que se conoce como IBI), y que la asociación muestras_recibidas/símbolos_transmitidos no puede hacerse con normalidad.

Surge entonces la necesidad de introducir en el sistema un periodo de guarda de longitud $L-1$ muestras (donde L es la longitud de la respuesta impulsiva del canal). Una alternativa posible es el llamado procedimiento de *zero padding*, que en nuestro caso equivaldría a anteponer a las $2N$ muestras del símbolo DMT un conjunto de muestras nulas de longitud $v=L-1$ muestras. De esta forma, por el canal se transmitirían bloques de $2N+v$ muestras. El receptor, identificaría perfectamente esos bloques, y desecharía de cada uno de ellos las v primeras muestras, que son las que se encuentran contaminadas por las muestras del símbolo anterior. Las $2N$ muestras restantes del bloque (desde la muestra $v+1$ hasta la muestra $2N+v$) no se encuentran contaminadas, puesto que su transmisión ha sido precedida por la de $L-1$ ceros.

Sin embargo, en el sistema DMT se opta por una solución distinta (teniendo siempre en cuenta que la longitud del prefijo que se añade tiene forzosamente igual a $v=L-1$ muestras). Dicha solución consiste en, como ya se ha comentado, anteponer al bloque de $2N$ muestras un bloque de longitud $L-1$ correspondiente a las últimas $L-1$ muestras del bloque de $2N$, formando así un bloque de tamaño $2N+L-1=2N+v$.

Vamos a dar unas ideas acerca de por qué se utiliza un CP en DMT y no se recurre al simple relleno con ceros (*zero padding*). Trabajando con señales continuas, un resultado conocido es que la convolución de dos señales en el dominio del tiempo se corresponde con el producto de sus transformadas de Fourier, dominio de la frecuencia. En el dominio discreto, el producto de dos DFTs se corresponde con la convolución circular de las secuencias en el dominio temporal. Si anteponemos un prefijo cíclico como el utilizado en DMT aseguramos por una parte la independencia entre los bloques y por otro, forzamos una especie de convolución circular (estamos haciendo a la señal parecer periódica).

CAPÍTULO 5:

IGUALACIÓN EN DMT APLICADA A VDSL. EL IGUALADOR EN EL DOMINIO DEL TIEMPO (TEQ) Y EL IGUALADOR EN EL DOMINIO DE LA FRECUENCIA (FEQ)

5.1 – INTRODUCCIÓN A LAS TÉCNICAS DE IGUALACIÓN.

En este capítulo se van a estudiar dos posibilidades de igualación para un sistema DMT. En un sistema DMT sin igualación se producen fenómenos de interferencia entre símbolos DMT. La forma de evitar que esta interferencia afecte a la calidad de la transmisión es utilizar un prefijo cíclico en cada símbolo. Ahora, con la igualación, vamos a minimizar la cantidad de interferencia que se produce, por lo que vamos a poder hacer funcionar al sistema con un CP menor. En concreto, se van a estudiar dos igualadores: uno que funciona en el dominio del tiempo (TEQ, Time-domain EQualizer) y otro en el dominio de la frecuencia (FEQ, Frequency-domain EQualizer)

En la mayoría de los diseños de MCM se utiliza el CP. Pero su uso acarrea los problemas ya vistos de disminución del rendimiento. En particular, en VDSL se explota un gran ancho de banda en el canal, y ello supone transmitir la señal con un periodo de muestreo muy pequeño. Esto provoca que la duración de la respuesta impulsiva del canal sea muy grande en número de muestras (aunque la duración en el tiempo de dicha respuesta impulsiva es la misma, el tener que muestrear cada menos tiempo hace que se alargue la respuesta discreta). Si la respuesta impulsiva se alarga demasiado, nos vemos obligados a introducir un CP muy largo, lo cual degrada bastante el rendimiento del enlace en el nivel físico, según vimos en el capítulo 4.

Una solución a este problema puede consistir en dividir el canal en un número muy grande de subcanales. Así, el efecto de v puede llegar a hacerse tan despreciable como queramos (ecuación (4.32)).

Sin embargo, el uso de un elevado número de subcanales presenta graves inconvenientes tanto desde el punto de vista computacional (gran tamaño de las FFTs) como del de los requerimientos de memoria (almacenamiento de tablas de asignación de bits, resultados intermedios, buffers, etc.). Para evitar la división del canal en número elevado de subcanales, y evitar también la aparición de los problemas antes señalados, se han propuesto unos esquemas de igualación en el sistema DMT, que son especialmente útiles en el caso de VDSL, y que se estudiarán a lo largo del capítulo.

Hasta ahora, el único mecanismo del que hemos hablado en el sistema DMT para reducir la IBI ha sido la inserción de un periodo de guarda entre los sucesivos símbolos DMT. Como ya vimos, la solución que se adoptaba era la de la inserción de un prefijo cíclico; con esta solución garantizábamos la independencia entre los subcanales.

Por su parte, el único mecanismo de igualación que se ha habilitado es el proceso de igualación en la frecuencia mediante el llamado “FEQ 1-tap”. En dicho proceso (llevado a cabo de forma independiente para cada subcanal) se divide el coeficiente complejo recibido por cada subcanal por la ganancia compleja del canal de transmisión (evaluada en la frecuencia del subcanal correspondiente), de forma que compensamos (en magnitud y en fase) el efecto del canal.

La combinación de las dos técnicas anteriores (inserción de CP y normalización de ganancia) da buenos resultados, siempre que la longitud de la respuesta impulsiva del canal (en número de muestras) sea inferior a la longitud del CP insertado. En los casos en los que el canal presenta una respuesta impulsiva muy larga (como en VDSL), la inserción de un prefijo cíclico adecuado supone un desperdicio tanto en ancho de banda como en energía empleada en la transmisión, y debemos buscar otras técnicas que palien este problema. Además, el problema se agrava mucho más en los casos en los que el número de subcanales no es elevado.

El proceso de igualación de la señal tiene lugar en el modem receptor. Vamos a estudiar dos tipos de igualadores, uno que funciona en el dominio del tiempo (TEQ, anterior a la demodulación vía FFT) y otro en el dominio de la frecuencia (FEQ, posterior a la FFT). La introducción de mecanismos de igualación implica un aumento de la complejidad computacional, pero debe provocar también un aumento de las prestaciones del sistema.

5.2 – IGUALACIÓN EN EL DOMINIO DEL TIEMPO

El objetivo de este igualador es provocar un acortamiento sustancial en la respuesta impulsiva del canal. Para ello, se inserta un filtro FIR (igualador) a la entrada del receptor, diseñado especialmente para este menester [1].

Si atendemos a la división por bloques que se hizo en el capítulo 4 para el enlace DMT completo, en esta parte nos vamos a ceñir exclusivamente al bloque correspondiente al canal de transmisión y a un nuevo bloque que vamos a introducir (el igualador TEQ) situado justo a la entrada del receptor.

Tenemos que recordar también el formato de la señal que se propaga por el canal. Como ya vimos, se trata de un flujo continuo de muestras, pero que están agrupadas en bloques de tamaño $2N+v$ (donde v es la longitud del CP y debe ser superior a la longitud de la respuesta impulsiva del canal).

Asumiremos que el canal es de tipo FIR, y que tiene una respuesta impulsiva de longitud $v+1$, esto es,

$$H(z) = h_0 + h_1 z^{-1} + \dots h_v z^{-v}, \quad (5.1)$$

donde los h_i son los coeficientes de la respuesta impulsiva hasta el índice $n=v$. Vamos a definir el vector CIR (*Channel Impulse Response*) como un vector de longitud $v+1$, en el que almacenaremos los coeficientes de la anterior respuesta impulsiva,

$$\mathbf{h} = [h_0 \ h_1 \ h_2 \ \dots \ h_v]. \quad (5.2)$$

Con el igualador que pretendemos implementar (TEQ), trataremos de lograr que la combinación canal+igualador tenga una respuesta impulsiva mucho más corta. A esta respuesta combinada del canal con el igualador vamos a denominarla TIR (*Target Impulse Response*, respuesta al impulso objetivo). La idea se muestra en la figura 5.1. Vamos a imponer que la TIR tenga una longitud de N_b+1 muestras (eligiendo N_b considerablemente menor que v). Si conseguimos un igualador que logre esta TIR, el enlace podrá funcionar con un CP de longitud de tan sólo N_b , en lugar de v .

Tenemos varias opciones a la hora de determinar cómo calcular los parámetros del TEQ a partir de la CIR y la TIR. La forma más natural sería decir que la TIR es igual a la CIR (pero haciendo los coeficientes cero a partir del índice N_b) y a partir de aquí, ajustar los coeficientes del igualador (obtención de la TIR por truncamiento de la CIR). Sin embargo, esta forma no es la más eficiente de resolver el problema. Como alternativa, existen dos técnicas óptimas para determinar tanto los coeficientes del igualador (que es un filtro FIR como ya hemos indicado) como los coeficientes de la TIR (que no tienen por qué coincidir con los de la CIR). Una técnica se basa en ajustar los coeficientes del TEQ según una minimización de cierto error cuadrático medio (MMSE-TEQ) y la otra se basa en una maximización de la relación señal a ruido geométrica en el canal (G-TEQ). En este proyecto, sólo nos centraremos en la primera de ellas.

5.3 – INICIALIZACIÓN DEL IGUALADOR TEQ: EL MMSE-TEQ

Considérese el esquema de la figura (5.1). En él está representados varios bloques: la función de transferencia del canal o CIR (**h**), la adición de ruido aditivo a la salida (**n**), el igualador (**w**) y la TIR (**b**).

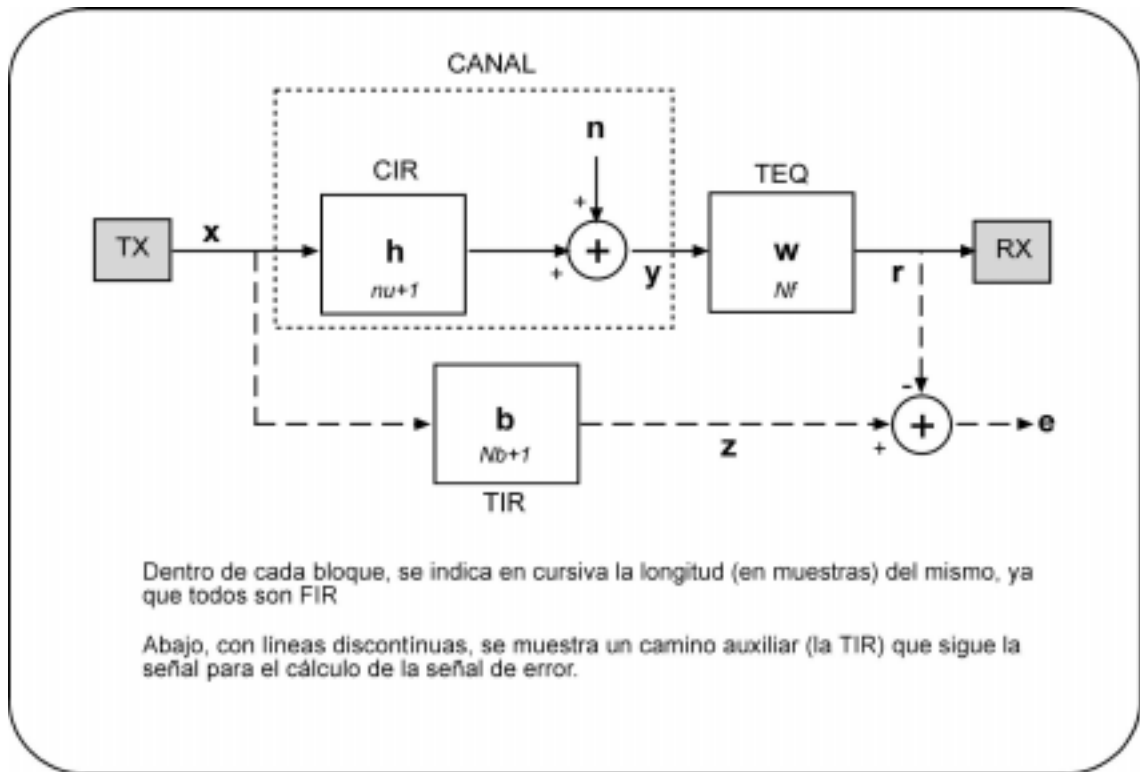


Figura 5.1 : Esquema del igualador TEQ

La señal que entra en el esquema (**x**) sería la señal tal y como se genera en el transmisor. Tras atravesar el bloque de la CIR se suma con el ruido, formando así la señal **y**. En un sistema sin igualación, ésta sería la señal que llega al receptor DMT. Sin embargo, ahora se da un paso más: la señal **y** se filtra con el igualador antes de ser procesada por el receptor. Como resultado de este proceso de igualación, se genera la señal **r**. Supondremos que el igualador es un filtro FIR de N_f taps,

$$\mathbf{w} = [w_0 \ w_1 \ w_2 \ \dots \ w_{N_f-1}]^t \quad (5.3)$$

donde $(.)^t$ denota la trasposición.

Por otro lado, formamos la señal $\mathbf{z}$, que es el resultado de convolucionar la señal transmitida $\mathbf{x}$ con el bloque de la TIR.

Al final del esquema de la figura 5.1 podemos ver un bloque de cálculo de una secuencia de error:

$$\mathbf{e} = \mathbf{z} - \mathbf{r}. \quad (5.4)$$

En el diseño del MMSE-TEQ [1] trataremos de minimizar la potencia de esta secuencia de error, es decir, minimizaremos el error cuadrático medio entre la señal igualada y la señal convolucionada con la TIR. Un error pequeño significa que la combinación de canal+igualador presenta una respuesta impulsiva muy similar a la TIR, y que, por tanto, tenemos luz verde para utilizar un CP de longitud menor.

En particular, definimos la TIR como un vector de longitud N_b+1

$$\mathbf{b} = [b_0 \ b_1 \ \dots \ b_{N_b}], \quad (5.5)$$

al cual vamos a imponer una restricción: que tenga energía unidad, es decir,

$$\sum_{i=0}^{N_f} b_i^2 = 1. \quad (5.6)$$

Esta restricción recibe el nombre de UEC (*Unit Energy Constraint*), y por eso, al igualador así diseñado también se le conoce como MMSE-UEC TEQ. La UEC es necesaria para evitar la solución trivial de $\mathbf{b}=\mathbf{0}$ y $\mathbf{w}=\mathbf{0}$, ya que esta solución es la que provoca el óptimo $\mathbf{e}=\mathbf{0}$.

La forma de actuar será la de buscar de forma conjunta la solución óptima para $\mathbf{b}$ y para $\mathbf{w}$. En primer lugar, se determinará la TIR óptima, y luego, con esta TIR, se calcula el igualador (resolviendo un problema de mínimos cuadrados). Si suponemos que las muestras de señal a la entrada no

están correlacionadas, puede demostrarse que la TIR óptima, $\mathbf{b}_{\text{opt}}$, es igual al autovector (normalizado) correspondiente al mínimo autovalor de la matriz $\mathbf{R}_\Delta$ (dependiente tanto del ruido como de la función de transferencia del canal) definida como [1]

$$\mathbf{R}_\Delta = \begin{bmatrix} \mathbf{0}_{(Nb+1) \times \Delta} & \mathbf{I}_{Nb+1} & \mathbf{0}_{(Nb+1) \times s} \end{bmatrix} \quad (5.7)$$

$$\left(\frac{1}{S_x} \mathbf{I}_{N_f+v} + \mathbf{H}^* \mathbf{R}_{nn}^{-1} \mathbf{H} \right)^{-1} \begin{bmatrix} \mathbf{0}_{\Delta \times (Nb+1)} \\ \mathbf{I}_{Nb+1} \\ \mathbf{0}_{s \times (Nb+1)} \end{bmatrix},$$

donde $\mathbf{0}_{m \times n}$ es la matriz nula de tamaño $m \times n$, $\mathbf{I}_m$ es la matriz identidad de tamaño m , $(.)^*$ denota la matriz traspuesta conjugada, Δ es el retraso del igualador,

$$0 \leq \Delta \leq N_f + v - N_b - 1, \quad (5.8)$$

s se define como

$$s = N_f + v - \Delta - N_b - 1, \quad (5.9)$$

S_x es la energía promedio de los símbolos de entrada, $\mathbf{R}_{nn}$ es la matriz de correlación del ruido, N_f -dimensional, y $\mathbf{H}$ es una matriz de tamaño $N_f \times (N_f + v)$ que viene dada por

$$\mathbf{H} = \begin{bmatrix} h_0 & h_1 & \cdots & h_v & 0 & \cdots & 0 \\ 0 & h_0 & h_1 & \cdots & h_v & 0 & \cdots \\ \vdots & \ddots & \ddots & \ddots & & \ddots & \vdots \\ 0 & \cdots & 0 & h_0 & h_1 & \cdots & h_v \end{bmatrix}. \quad (5.10)$$

Como decimos, hay que formar la matriz $\mathbf{R}_\Delta$ (que es de tamaño $(N_b+1) \times (N_b+1)$), calcular sus autovectores y elegir el autovector correspondiente al mínimo autovalor. Dicho autovector, normalizado, será la TIR óptima ($\mathbf{b}_{\text{opt}}$)

del sistema. Una vez que se ha determinado la TIR, los coeficientes del igualador se calculan como

$$\mathbf{w}_{opt}^* = \begin{bmatrix} \mathbf{0}_{1 \times \Delta} & \mathbf{b}_{opt}^* & \mathbf{0}_{1 \times s} \end{bmatrix} \mathbf{H}^* \left(\mathbf{H} \mathbf{H}^* + \frac{\mathbf{R}_{nn}}{S_x} \right)^{-1}. \quad (5.11)$$

Mediante esta solución se optimizan los coeficientes del igualador y de la TIR, con lo que se logra una mejora en el rendimiento del sistema (dado que se puede eliminar la IBI con un prefijo cíclico mucho más corto). Además, la carga computacional no aumenta demasiado: no se ha modificado la longitud de la FFT y, por lo general, no hará falta un igualador de alto orden [3].

5.4 – IGUALACIÓN EN EL DOMINIO DE LA FRECUENCIA

Una vez que se han estudiado las técnicas de igualación en el dominio del tiempo, en los siguientes apartados vamos a presentar otro igualador, esta vez funcionando en el dominio de la frecuencia. El objetivo sigue siendo el mismo: queremos utilizar un prefijo cíclico de menor longitud sin empeorar las prestaciones del sistema. Si nos limitamos a acortar el prefijo cíclico y no hacemos nada más, es inevitable la aparición de fenómenos de IBI e ISI. La misión de los igualadores es precisamente paliar esos fenómenos de interferencia de manera que el sistema mantenga sus prestaciones (en cuanto a probabilidad de error y capacidad alcanzada).

Ya hemos estudiado un igualador que funciona en el dominio del tiempo (TEQ). Ahora vamos a introducir un esquema que funciona sobre la señal ya demodulada, es decir, la igualación se hace después de la FFT (de ahí el nombre de “igualación en el dominio de la frecuencia”). Como en el sistema DMT la detección de los símbolos se hace en el dominio de la frecuencia, es de esperar que este tipo de igualadores funcione mejor. Además, esta igualación presenta otra ventaja adicional: ya que las transmisiones se llevan a cabo de

forma independiente para cada subcanal, podemos ajustar los coeficientes del igualador por separado para cada uno de ellos. Es decir, si en un sistema DMT se usan N subcanales, podemos diseñar de forma independiente N igualadores, y hacer que cada uno actúe sobre cada subcanal. Sin embargo, es evidente que el igualador TEQ, con el esquema que hemos propuesto, no puede realizar una optimización separada para cada subcanal, ya que en el dominio del tiempo no podemos discriminar la señal que viaja por cada uno de los subcanales. Una idea de lo que se pretende se muestra en la figura 5.2.

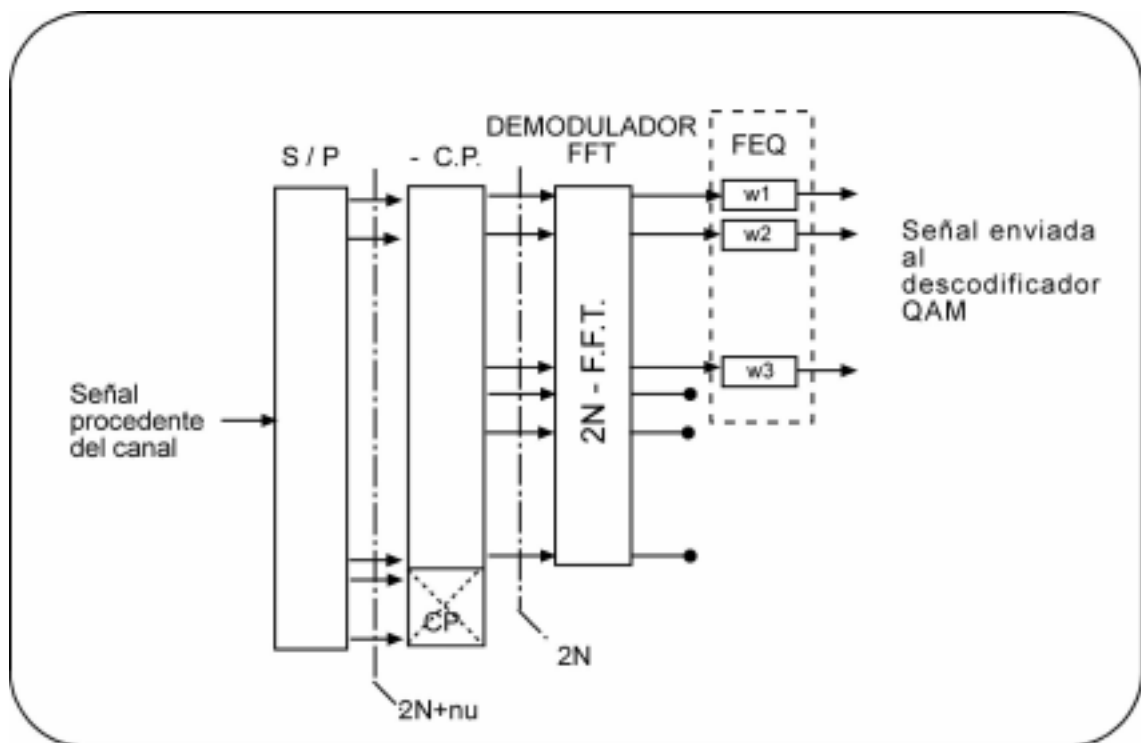


Figura 5.2 : Situación del igualador FEQ dentro del receptor.

5.5 – EL IGUALADOR FEQ

5.5.1 – Descripción de la notación utilizada

Vamos a introducir en este apartado la notación utilizada en el estudio de este igualador, la cual difiere un poco de la empleada hasta ahora, pero facilita el análisis que vamos a realizar.

Llamaremos N al número de muestras temporales de que consta un símbolo DMT (sin incluir el CP); $X_i^{(k)}$ será el subsímbolo complejo (incluyendo las muestras que por simetría se generan artificialmente) que el transmisor emite por el subcanal i en el símbolo DMT numero k ; $Y_i^{(k)}$ será la salida demodulada para el tono i (después de la FFT); y $Z_i^{(k)}$ será la salida final (después de la igualación). Nótese que tanto las X como las Y deben cumplir la relación de simetría conjugada

$$X_i^{(k)} = X_{N-(i-2)}^{*(k)} \quad i = 2, \dots, \frac{N}{2}. \quad (5.12)$$

Denotaremos con v la longitud del prefijo cíclico utilizado en la transmisión. Por tanto, $s=N+v$ denotará la longitud total de un bloque DMT en el tiempo. El vector

$$\mathbf{h} = [h_L \quad \dots \quad h_0 \quad \dots \quad h_{-K}] \quad (5.13)$$

denota la respuesta impulsiva del canal en orden inverso; n_l será el ruido aditivo de la l -ésima muestra a la salida del canal; e y_l será la l -ésima muestra de la señal recibida.

Para describir el modelo de datos consideraremos tres símbolos sucesivos $X_{1:N}^{(c)}$, con $c=k-1, k, k+1$, respectivamente, donde

$$X_{1:N}^{(c)} = [X_1^{(c)} \quad X_2^{(c)} \quad \dots \quad X_N^{(c)}]^T. \quad (5.14)$$

El símbolo de interés es el símbolo k . Los símbolos $k-1$ y $k+1$ se corresponden con el símbolo inmediatamente anterior e inmediatamente posterior, respectivamente, y van a ser los causantes de la IBI.

Con esta notación, la señal recibida (en el tiempo) puede expresarse de la siguiente forma [2]:

$$\begin{aligned}
 & \overbrace{\begin{bmatrix} y_{ks+v-T+2+\delta} \\ \vdots \\ y_{(k+1)s+\delta} \end{bmatrix}}^{\mathbf{y}} \\
 &= \begin{bmatrix} \mathbf{O}_{(1)} & \left| \begin{array}{ccc} \mathbf{h} & 0 & \dots \\ \vdots & \ddots & \ddots \\ 0 & \dots & \mathbf{h} \end{array} \right| & \mathbf{O}_{(2)} \end{bmatrix} \cdot \begin{bmatrix} \mathbf{P} & \mathbf{O} & \mathbf{O} \\ \mathbf{O} & \mathbf{P} & \mathbf{O} \\ \mathbf{O} & \mathbf{O} & \mathbf{P} \end{bmatrix} \\
 & \cdot \begin{bmatrix} \mathfrak{I}_N & \mathbf{O} & \mathbf{O} \\ \mathbf{O} & \mathfrak{I}_N & \mathbf{O} \\ \mathbf{O} & \mathbf{O} & \mathfrak{I}_N \end{bmatrix} \cdot \underbrace{\begin{bmatrix} X_{1:N}^{(k-1)} \\ X_{1:N}^{(k)} \\ X_{1:N}^{(k+1)} \end{bmatrix}}_{\mathbf{X}} + \underbrace{\begin{bmatrix} n_{ks+v-T+2+\delta} \\ \vdots \\ y_{(k+1)s+\delta} \end{bmatrix}}_{\mathbf{n}}
 \end{aligned} \tag{5.15}$$

$\mathbf{O}_{(1)}$ y $\mathbf{O}_{(2)}$ son matrices nulas de tamaño $(N+T-1) \times (N+v-T+1-L+v+\delta)$ y $(N+T-1) \times (N+v-K-\delta)$ respectivamente. El parámetro δ indica el alineamiento de símbolos en el receptor, es decir, qué grupo de s muestras se corresponde exactamente con un símbolo transmitido. La matriz $\mathbf{P}$ es de la forma

$$\mathbf{P} = \left[\begin{array}{c|c} \mathbf{O} & \mathbf{I}_v \\ \hline & \mathbf{I}_N \end{array} \right], \tag{5.16}$$

y representa la adición del CP. Las matrices $\mathbf{I}_N$ son matrices IFFT de tamaño $N \times N$. $\mathbf{I}_N$ es la matriz identidad de tamaño N . En el vector $\mathbf{h}$, el índice cero se coloca de forma que se maximice la energía contenida en las muestras $[h_0 \ h_1 \ \dots \ h_v]$. Esto hace que quede una cabeza $[h_K \ \dots \ h_1]$ y una cola $[h_{v+1} \ \dots \ h_L]$, que provocan interferencias con los símbolos anterior y posterior.

5.5.2 – Igualación por tonos en un módem DMT

Para comprender mejor el funcionamiento de este igualadore, se emplearán grafos de flujo de señal (SFGs), contruidos con bloques elementales: sumador, multiplicador-sumador, elemento de retraso y submuestreador. El esquema de estos bloques constructores se muestra en la figura 5.3.

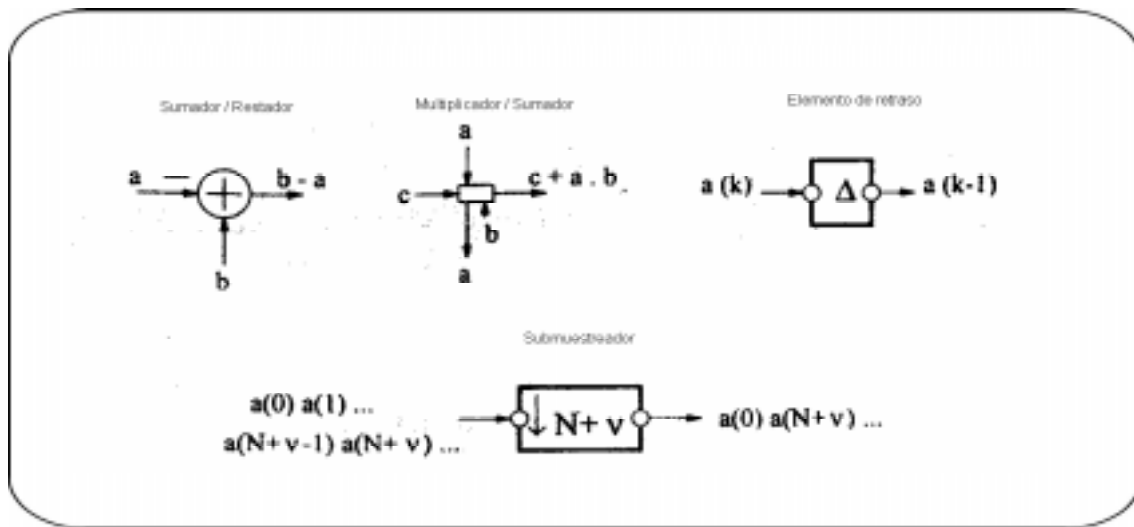


Figura 5.3 : Celdas individuales para el SFG

Para implementar el igualador en la frecuencia [2], partiremos del igualador en el tiempo, mostrado en la figura 5.4 y estudiado al principio de este capítulo. La relación matemática que se cumple en el receptor con TEQ implementado estudiado en el capítulo anterior es la siguiente:

$$\begin{bmatrix} Z_1^{(k)} \\ \vdots \\ Z_N^{(k)} \end{bmatrix} = \begin{bmatrix} D_1 & 0 & \dots \\ 0 & \ddots & 0 \\ \vdots & 0 & D_N \end{bmatrix} \cdot F_N \cdot (\mathbf{Y} \cdot \mathbf{w}), \quad (5.17)$$

donde la matriz $\mathbf{Y}$ es

$$\mathbf{Y} = \begin{bmatrix} y_{ks+v+1} & y_{ks+v} & \cdots & y_{ks+v-T+2} \\ y_{ks+v+2} & y_{ks+v+1} & & y_{ks+v-T+3} \\ & \ddots & \ddots & \\ y_{(k+1)s} & y_{(k+1)s-1} & \cdots & y_{(k+1)s-T+1} \end{bmatrix} \quad (5.18)$$

y el vector $\mathbf{w}$ sería el igualador TEQ, el cual viene dado por

$$\mathbf{w}_{T \times 1} = [w_0 \ w_1 \ \cdots \ w_{T-1}]^T. \quad (5.19)$$

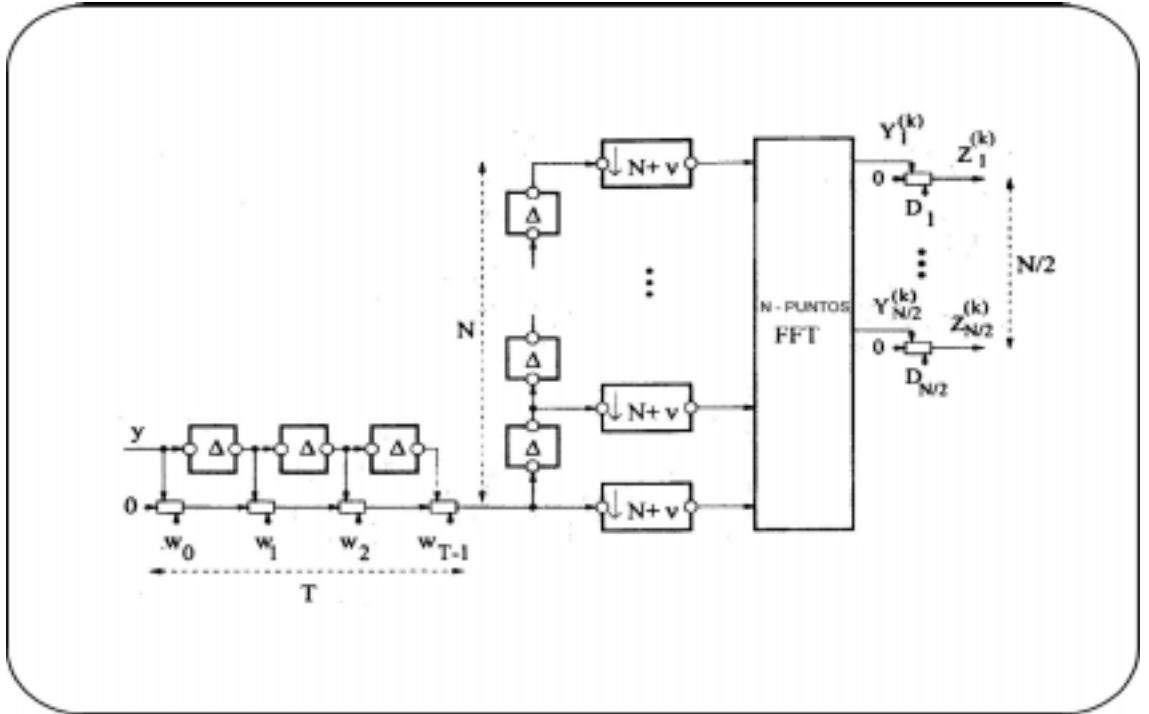


Figura 5.4 : Esquema del igualador TEQ

F_N sería la matriz FFT de tamaño N y D_i sería cada uno de los elementos del bloque llamado “normalizador de ganancia”, el cual ya vimos que podía verse como un igualador FEQ de 1 coeficiente. Para no complicar demasiado la notación, se ha omitido el retardo de sincronización δ en la expresión de $\mathbf{Y}$. Como vemos, $\mathbf{Y}$ es una matriz Toeplitz que contiene las mismas muestras que el vector $\mathbf{y}$ de la ecuación (5.15).

El FEQ que vamos a estudiar se basa en el TEQ ya estudiado, pero trasladando las operaciones en el dominio del tiempo al dominio de la frecuencia. De hecho, para cada tono i podemos reescribir la ecuación (5.7) como

$$Z_i^{(k)} = D_i \cdot \text{fila}_i(F_N) \cdot (\mathbf{Y} \cdot \mathbf{w}) = \underbrace{\text{fila}_i(F_N \cdot \mathbf{Y})}_{T \text{ FFTs}} \cdot \underbrace{\mathbf{w} \cdot D_i}_{\mathbf{w}_i} \quad (5.20)$$

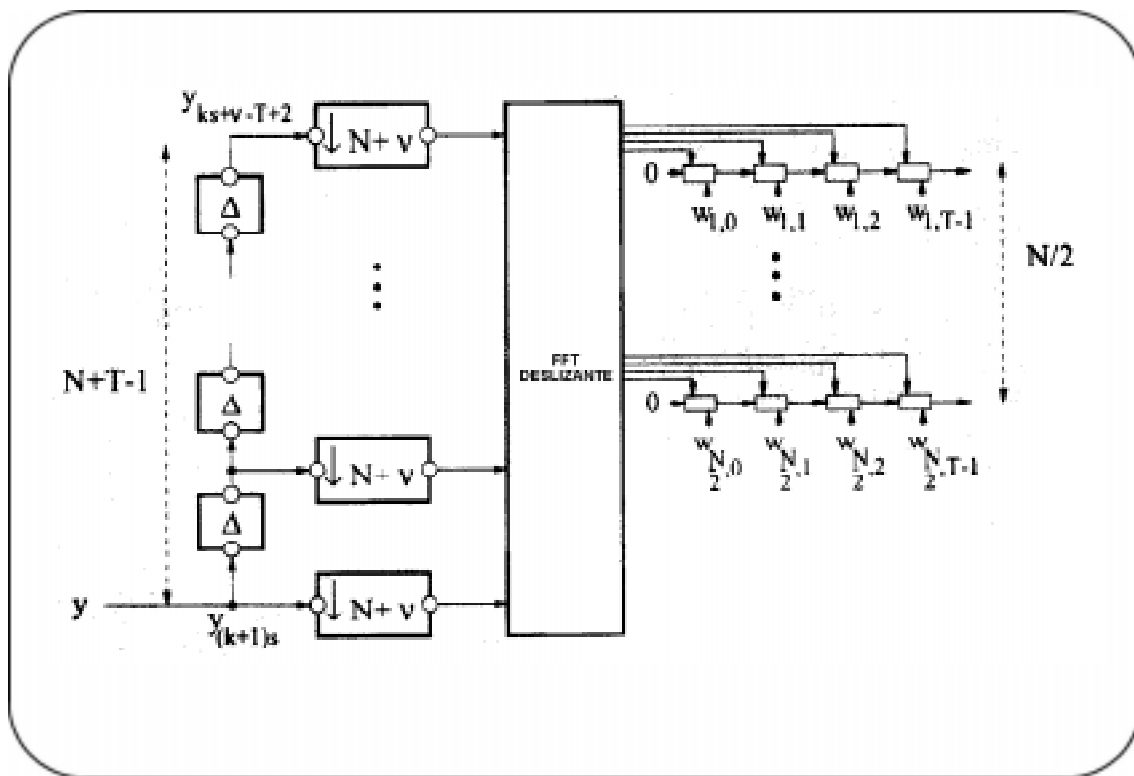


Figura 5.4 : Esquema del igualador FEQ

Colocando D_i a la derecha, obtenemos otro igualador $\mathbf{w}_{i,Tx1} = \mathbf{w} D_i$, el cual puede verse como un igualador de T taps (complejos) en el dominio de la frecuencia, para el tono i . Como vemos, ahora cada tono puede tener su propio igualador. El próximo paso es permitir una optimización por separado de los igualadores para cada tono. La estructura del receptor utilizando este tipo de igualación se muestra en la figura 5.4.

En la expresión (5.20) podemos observar que, en lugar de una operación FFT por símbolo DMT (como era necesario en el receptor con

igualación TEQ), ahora son necesarias T operaciones FFT (una para cada columna de la matriz $\mathbf{Y}$). Sin embargo (y afortunadamente) debido a la estructura Toeplitz de la matriz $\mathbf{Y}$, estas FFTs pueden ser calculadas de forma eficiente con una operación llamada “FFT deslizante”. Tan sólo es necesario calcular una FFT completa, pudiendo ser calculadas las $T-1$ restantes a partir de ésta. En efecto,

$$\begin{aligned} \mathbf{F}_N(:, m+1) &= \\ &= \mathbf{F}_N(:, m) \bullet \mathbf{p} + \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix}_{N \times 1} \cdot (y_{ks+v-(m-1)} - y_{ks+s-(m-1)}) \end{aligned} \quad (5.21)$$

donde el operador $\bullet$ indica una multiplicación componente a componente, y el vector $\mathbf{p}$ viene dado por

$$\mathbf{p}^T = [\alpha^0 \ \alpha^1 \ \dots \ \alpha^{N-1}] \quad \text{con} \quad \alpha = e^{-j\frac{2\pi}{N}}. \quad (5.22)$$

$\mathbf{Y}(:, m+1)$ es la $(m+1)$ -ésima columna de $\mathbf{Y}$, $y_{ks+v-(m-1)}$ es su primer elemento. $\mathbf{Y}(:, m)$ es la m -ésima columna de $\mathbf{Y}$, $y_{ks+s-(m-1)}$ es su último elemento. Para el tono i , todos los elementos relevantes de su FFT pueden calcularse a partir del tono i de la primera columna y de una serie de $T-1$ términos en diferencia.

Para reducir aún más la complejidad, estas combinaciones lineales (de coeficientes de la columna anterior y de términos en diferencia) pueden incluirse en los coeficientes del igualador, de tal forma que el FEQ global para cada tono i tenga como entradas la i -ésima componente (compleja) de la FFT (no deslizante) y $T-1$ términos en diferencia (reales) entre muestras de la señal recibida. Llamaremos $\mathbf{v}$ a este igualador modificado, el cual puede obtenerse a partir del $\mathbf{w}$ según

$$\begin{bmatrix} v_{i,0} \\ v_{i,1} \\ \vdots \\ v_{i,T-1} \end{bmatrix} = \begin{bmatrix} 1 & \alpha^{i-1} & \dots & \alpha^{(i-1)(T-1)} \\ 0 & 1 & \ddots & \vdots \\ \vdots & \ddots & \ddots & \alpha^{i-1} \\ 0 & \dots & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} w_{i,0} \\ w_{i,1} \\ \vdots \\ w_{i,T-1} \end{bmatrix}. \quad (5.23)$$

Ahora, el igualador $\mathbf{v}_i$ tiene dos misiones: por un lado, iguala el tono i (recordemos que hay que calcular un igualador para cada tono), y por otro, incluye la FFT deslizando para el tono i . Como podemos ver en la figura 5.5, la igualación se realiza con $\mathbf{v}$, por lo que no es necesario un cálculo explícito de $\mathbf{w}$.

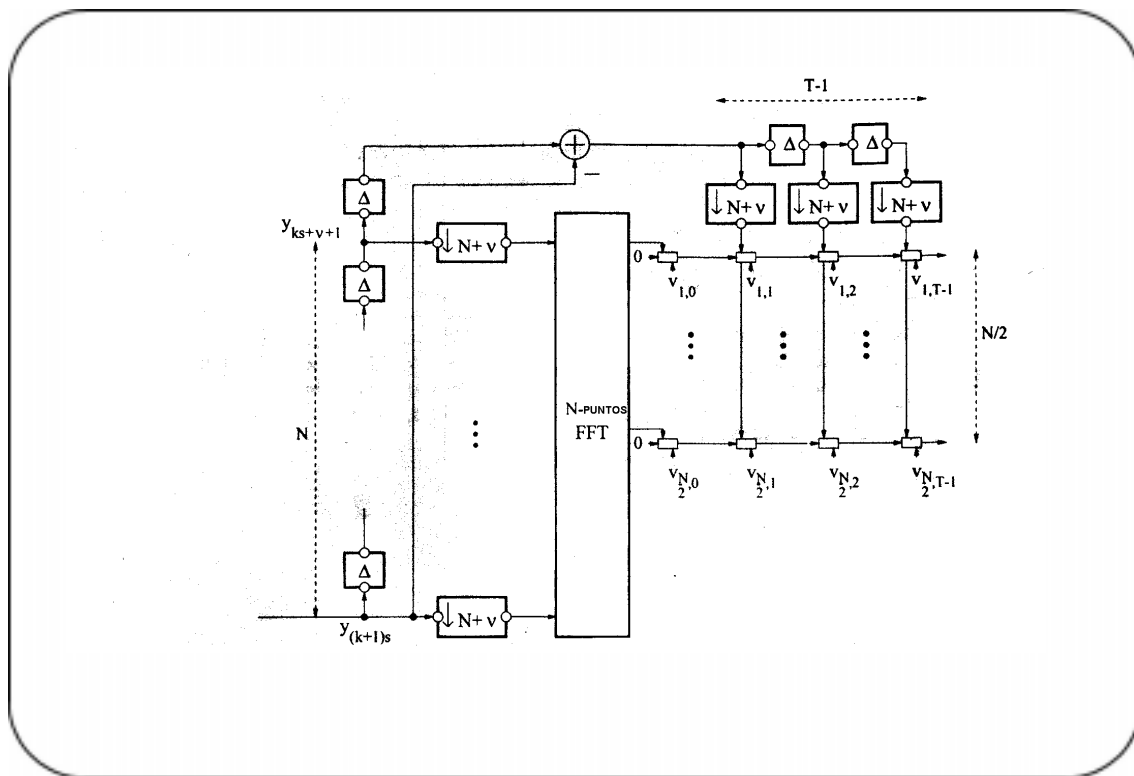


Figura 5.5 : Igualación en la frecuencia sin necesidad de FFT deslizando

5.6 – INICIALIZACIÓN DEL FEQ

Para cada uno de los tonos usados, encontraremos el MMSE-FEQ (para una elección particular del retraso δ) mediante la minimización de la siguiente función de coste [2]:

$$\begin{aligned}
 & \min_{\mathbf{w}_i} J(\mathbf{w}_i) \\
 &= \min_{\mathbf{w}_i} \varepsilon \left\{ \left| Z_i^{(k)} - X_i^{(k)} \right|^2 \right\} \\
 &= \min_{\mathbf{w}_i} \varepsilon \left\{ \left| \overline{\mathbf{w}}_i^T \cdot \begin{bmatrix} \mathbf{F}_N(i,:) & 0 & \dots \\ \vdots & \ddots & \vdots \\ 0 & \dots & \mathbf{F}_N(i,:) \end{bmatrix} \cdot \mathbf{y} - X_i^{(k)} \right|^2 \right\}
 \end{aligned} \tag{5.24}$$

con el vector

$$\overline{\mathbf{w}}_i^T = [w_{i,T-1} \ \dots \ w_{i,0}]. \tag{5.25}$$

(el operador “barra” significa poner los coeficientes del vector con los índices en orden inverso). $\mathbf{F}_N(i,:)$ es la fila i -ésima de la matriz FFT N -dimensional $\mathbf{F}_N$ (4.42). Este MMSE-FEQ optimiza la SNR por separado para cada tono.

De forma similar, podemos obtener el problema para la optimización del igualador $\mathbf{v}$ [2]:

$$\min_{\mathbf{v}_i} \varepsilon \left\{ \left| \overline{\mathbf{v}}_i^T \cdot \underbrace{\begin{bmatrix} \mathbf{I}_{T-1} & | & \mathbf{O} & | & -\mathbf{I}_{T-1} \\ \mathbf{O} & | & \mathbf{F}_N(i,:) \end{bmatrix}}_{\mathbf{F}_i} \cdot \mathbf{y} - X_i^{(k)} \right|^2 \right\} \tag{5.26}$$

siendo

$$\bar{\mathbf{v}}_i^T = [v_{i,T-1} \ \cdots \ v_{i,0}]. \quad (5.27)$$

La primera fila de bloques de la matriz $\mathbf{F}_i$ se ocupa de extraer los términos en diferencias, mientras que la última fila de bloques es la que implementa la FFT.

En el caso de que conozcamos más datos acerca del sistema, podemos llegar a una expresión cerrada el problema de optimización. En concreto, necesitamos conocer la función de transferencia del canal y las matrices de correlación de la señal y del ruido presentes en el sistema. Tomemos entonces las siguientes definiciones. El vector $\mathbf{e}_i^{(k)}$ será

$$\mathbf{e}_i^{(k)} = [00 \cdots 010 \cdots 0]_{1 \times 3N} \quad (5.28)$$

donde el 1 está colocado en el índice $N+i$. $\mathbf{R}_x = \mathcal{E}\{\mathbf{X}\mathbf{X}^H\}$ es la matriz de correlación de la señal, y $\mathbf{R}_n = \mathcal{E}\{\mathbf{n}\mathbf{n}^H\}$ es la matriz de correlación del ruido. Por su parte, $\mathbf{H}$ es la matriz que se definió en la ecuación (5.15), y como se vio, incluye la convolución con el canal, la modulación mediante IFFT y la adición de un CP.

Por tanto, mediante la combinación de las expresiones (5.15) y (5.26), y teniendo en cuenta todo lo anterior, podemos decir que el problema de inicialización del igualador se reduce a la minimización de [2]

$$\begin{aligned} J(\mathbf{v}_i) &= \mathcal{E} \left\{ \left| \bar{\mathbf{v}}_i^T \mathbf{F}_i (\mathbf{H}\mathbf{X} + \mathbf{n}) - \mathbf{e}_i^{(k)} \right|^2 \right\} \\ &= \mathcal{E} \left\{ \left\| \begin{bmatrix} \mathbf{X}^H & \mathbf{n}^H \end{bmatrix} \cdot \begin{bmatrix} \mathbf{H}^H \mathbf{F}_i^H \bar{\mathbf{v}}_i^* - \mathbf{e}_i^{(k)H} \\ \mathbf{F}_i^H \bar{\mathbf{v}}_i^* \end{bmatrix} \right\|^2 \right\} \\ &= \left\| \begin{bmatrix} \mathbf{R}_x^{1/2} & \mathbf{O} \\ \mathbf{O} & \mathbf{R}_n^{1/2} \end{bmatrix} \cdot \begin{bmatrix} \mathbf{H}^H \mathbf{F}_i^H \bar{\mathbf{v}}_i^* - \mathbf{e}_i^{(k)H} \\ \mathbf{F}_i^H \bar{\mathbf{v}}_i^* \end{bmatrix} \right\|_2^2 \\ &= \left\| \underbrace{\begin{bmatrix} \mathbf{R}_x^{1/2} \mathbf{H}^H \mathbf{F}_i^H \\ \mathbf{R}_n^{1/2} \mathbf{F}_i^H \end{bmatrix}}_{\mathbf{U}} \bar{\mathbf{v}}_i^* - \underbrace{\begin{bmatrix} \mathbf{R}_x^{1/2} \mathbf{e}_i^{(k)H} \\ \mathbf{O} \end{bmatrix}}_{\mathbf{d}} \right\|_2^2. \end{aligned} \quad (5.29)$$

Este problema de ajuste por mínimos cuadrados debe ser resuelto de forma separada para cada tono i . Por ejemplo, puede utilizarse una factorización QR para la matriz compuesta $[\mathbf{U} \ \mathbf{d}]$. Con los coeficientes calculados de este modo, la energía de las componentes de señal, interferencia y ruido puede calcularse según las expresiones

$$\begin{aligned} S_i &= \left| \left(\mathbf{R}_x^{1/2} \mathbf{H}^H \mathbf{F}_i^H \right)_{(N+i,:)} \bar{\mathbf{v}}_i^* \right|^2 \\ I_i &= \left\| \mathbf{R}_x^{1/2} \mathbf{H}^H \mathbf{F}_i^H \bar{\mathbf{v}}_i^* \right\|^2 - S_i \\ N_i &= \left\| \mathbf{R}_n^{1/2} \mathbf{F}_i^H \bar{\mathbf{v}}_i^* \right\|^2. \end{aligned} \quad (5.30)$$

La SNR para cada tono puede calcularse como $\text{SNR}_i = (S_i / (I_i + N_i))$.

5.7 – COMENTARIOS ACERCA DE LOS IGUALADORES

5.7.1 – Comentarios acerca del TEQ

El primero de los igualadores estudiados, operaba sobre la señal en el dominio del tiempo. Consistía en un filtro FIR (el igualador) que se insertaba a la salida del canal, convirtiéndose así en la primera etapa del receptor.

El objetivo del igualador es conseguir que la respuesta combinada del sistema canal-igualador sea significativamente más corta que la del canal sin igualación. Para ello debemos desarrollar un método para ajustar los coeficientes del igualador (ciñéndonos a los parámetros de diseño que imponamos) de forma que se cumpla el objetivo marcado.

Se ha desarrollado un método para inicializar los coeficientes del igualador. El método, aunque no es el óptimo, es el mas natural; se basa en la minimización de un cierto error cuadrático medio. De hecho, puede demostrarse que la pérdida de eficiencia con respecto al óptimo es escasa [1]. Se ha visto cómo con este sistema se puede lograr de manera sencilla el acortamiento de la respuesta impulsiva.

En cuanto a la complejidad computacional que la inserción del igualador supone para el sistema, debemos distinguir dos aspectos: en primer lugar, debemos tener en cuenta los cálculos necesarios para inicializar los coeficientes del igualador, que incluyen el cálculo de autovectores de una matriz para resolver un ajuste por mínimos cuadrados; en segundo lugar, también debemos considerar la complejidad que para el modem receptor supondría el filtrado en tiempo real de las muestras recibidas (proceso de igualación). La parte más compleja es la inicialización, la cual introduciría un retardo en el sistema, pero que luego se ha de compensar con la mejora del rendimiento.

Vamos a añadir también unos comentarios acerca del efecto que pueden tener los parámetros de diseño (N_f , N_b) sobre la solución del igualador. La elección de N_f (número de taps del igualador) es un compromiso entre la complejidad que se introduce en el sistema y el grado de ajuste de la respuesta global con la TIR [5]. Por otro lado, una elección de N_b (máximo índice de la TIR) demasiado pequeña provoca una mejora notable del rendimiento (el CP necesario se reduce), pero para una misma N_f , el MSE se incrementa notablemente, con lo que el sistema se degradará (la respuesta global se ajustará muy mal a la TIR).

5.7.2 – Comentarios acerca del FEQ

En este capítulo hemos estudiado otro esquema de igualación para DMT. Se ha partido del igualador en el dominio del tiempo (TEQ), y se ha llegado al igualador en el dominio de la frecuencia, mediante el trasvase de operaciones de uno a otro dominio.

Como resultado, de un solo igualador TEQ se obtienen muchos igualadores FEQ, uno para cada subcanal. Este hecho nos permite optimizar cada subcanal por separado, cosa que no podía hacer el TEQ. Aparentemente, esto supone un aumento de la complejidad computacional, pero hemos estudiado cómo se pueden realizar los cálculos de forma eficiente sin tener que recurrir de forma explícita al cálculo de la FFT deslizante. Esto quiere decir que durante la fase de transmisión (después del arranque del modem, en el que se calculan los parámetros a utilizar y se inicializan los igualadores) la complejidad computacional necesaria (en el sentido de número de operaciones por unidad de tiempo) es muy similar a la requerida con el TEQ.

Sin embargo, sí es cierto que los requerimientos de memoria aumentan con respecto al caso del TEQ, ya que ahora no tenemos uno, sino muchos igualadores, cada uno con T coeficientes. No obstante, esto no está considerado como un inconveniente prohibitivo a la hora de la implementación.

CAPÍTULO 6:

TÉCNICAS ESTADÍSTICAS PARA LA ESTIMACIÓN DE LA PROBABILIDAD DE ERROR EN UN SISTEMA DIGITAL

6.1 – INTRODUCCIÓN

Este capítulo nos va a servir para describir con detalle las técnicas que se utilizarán en las simulaciones para calcular la probabilidad de error. Se comenzará estudiando el método de Monte Carlo; después introduciremos otro método para intentar acortar los tiempos de simulación. Se estudiará la forma de construir estos estimadores y sus propiedades más importantes.

En un sistema de comunicación digital, el parámetro de más relevancia a la hora de evaluar su calidad es el comportamiento en cuanto a ocurrencia de errores. Dependiendo de la situación, nos interesarán unos aspectos u otros de ese comportamiento. Pero no cabe duda que la situación que más veces se nos presenta es aquella en la que se transmiten símbolos pertenecientes a un alfabeto de tamaño M (usualmente $M=2^k$) y el parámetro de interés es el promedio de la ocurrencia de errores para una secuencia de símbolos transmitidos indefinidamente larga. Por supuesto, en otros contextos, puede ser útil no sólo el promedio de los errores, sino también la forma en lo que estos ocurren: de forma aislada, en ráfagas, etc. En estos casos tienen cabida otras definiciones como el porcentaje de bloques libres de errores.

En nuestro análisis nos centraremos en el caso habitual, es decir, en determinar (estimar) la probabilidad de error de símbolo en un sistema digital. En concreto, vamos a analizar dos posibilidades de realizar dicha estimación [17]:

- Método de Monte Carlo.

- Extrapolación de cola (pseudoumbrales).

Para centrar las ideas, hagamos una breve descripción de lo que pretendemos conseguir. Supongamos que tenemos un sistema de comunicación digital. Como sabemos, los sistemas de este tipo se caracterizan por ceñirse exclusivamente a la utilización de símbolos pertenecientes a un conjunto finito. Si el transmisor transmite de forma continua una secuencia de tales símbolos, a la salida del receptor obtenemos una cadena de símbolos, que debe coincidir con la generada por el transmisor, salvo algunos símbolos que se habrán detectado de forma errónea. Pues bien, es precisamente la probabilidad de error lo que queremos estimar. Se supondrá que los símbolos son transmitidos de forma independiente unos de otros, y que, por tanto, los errores ocurren también de forma independiente.

6.2 – EL MÉTODO DE MONTE CARLO

6.2.1 – Definición del estimador

El término “Monte Carlo” es meramente una etiqueta para la implementación de una secuencia de ensayos de Bernouilli. Esto quiere decir que contamos el número de “éxitos” (en nuestro estudio, un “éxito” sería la detección de un símbolo erróneo) y lo dividimos entre el número de ensayos realizados.

Este método no requiere hacer hipótesis acerca de la distribución estadística del ruido que genera los errores. Matemáticamente, podemos expresar

$$p = \int_{v \in D_0} f_v(v) dv, \quad (6.1)$$

donde f_v es la función de densidad de probabilidad para el símbolo de estudio y D_0 es la región de v que se corresponde con la detección de un error. Si definimos la función indicador de error

$$H(v) = \begin{cases} 1, & v \in D_0 \\ 0, & v \notin D_0 \end{cases}, \quad (6.2)$$

podemos expresar la ecuación (6.1) en la forma

$$p = \int_{-\infty}^{\infty} H(v) f_v(v) dv. \quad (6.3)$$

Esto es equivalente a decir que

$$p = E[H(v)], \quad (6.4)$$

donde el operador E representa la esperanza. Como un estimador de la esperanza es la media muestral, podemos escribir que

$$\hat{p} = \frac{1}{N_0} \sum_{i \in I_0} H(V_i), \quad (6.5)$$

donde $H(V_i)$ actúa como un detector de error (suma) y N_0 como promediador. Si el número de símbolos observados es muy grande, el estimador (la media muestral) converge a la media poblacional.

6.2.2 – Propiedades del estimador

La estimación por el método de Monte Carlo es uno de los pocos casos en los que podemos obtener una distribución exacta del estimador. Esto se debe a que, para un N dado, el número de errores en una secuencia de ensayos de Bernouilli sigue una distribución binomial. Sin embargo, la

manipulación de las expresiones de dicha distribución resulta tediosa, por lo que se recurre a aproximaciones. Las más usadas son la aproximación por medio de la distribución de Poisson y de la normal.

Dada su importancia y su amplio conocimiento, la aproximación por medio de la distribución normal es la más utilizada. De su estudio se llega a que, para estimar una probabilidad p , es necesario realizar un número de ensayos igual a

$$\#ensayos = \frac{1}{p} \times k, \quad (6.6)$$

donde k es un parámetro que varía según la varianza que deseemos para este estimador. Normalmente se considera que una varianza aceptable es la correspondiente a un error de $\pm 5\%$ en la estimación, lo cual implica que $k=10^2$.

Es evidente que cuando las probabilidades de error que tenemos que estimar son muy pequeñas el uso de este método implica alargar mucho el tiempo de simulación. Sin embargo, es un método sencillo de utilizar (es el más natural).

6.3 – EL MÉTODO DE EXTRAPOLACIÓN DE COLA (“TAIL EXTRAPOLATION”)
--

6.3.1 – Definición del estimador

Este método sólo es aplicable si la función de densidad de probabilidad de la señal se ajusta a la llamada “exponencial generalizada”, que viene dada por

$$f_v(x) = \frac{v}{2\sqrt{2}\sigma\Gamma(1/v)} e^{-\left|\frac{x-\mu}{\sqrt{2}\sigma}\right|^2}, \quad (6.7)$$

donde $\Gamma(.)$ es la función gamma. Nótese que tomando $v=2$ obtenemos la función de densidad de probabilidad de una distribución gaussiana, lo cual se adapta perfectamente al tipo de ruido que estamos generando.

Si el umbral de detección de un error es t , la probabilidad de error viene dada por

$$p(t) = \int_{\mu+t}^{\infty} f_v(x) dx. \quad (6.8)$$

Para t suficientemente grande, se verifica que

$$\ln(-\ln p(t)) \approx v \ln\left(\frac{t}{\sqrt{2}\sigma}\right), \quad (6.9)$$

lo cual significa que el doble logaritmo de la probabilidad de error es asintóticamente lineal con el logaritmo del umbral.

La relación lineal anterior es la que va a ser aprovechada en este método. En el decisor de nuestro receptor tenemos que ajustar un umbral a partir del cual los símbolos que se detectan son considerados como error. Pero si definimos un t menor (un pseudo-umbral), lo que obtenemos es una probabilidad de error mayor (una pseudo-probabilidad de error).

Con dos pseudo-umbrales, podemos obtener dos pseudo-probabilidades de error, con lo que podemos estimar la pendiente de la recta de la ecuación (6.9). Una vez que se conoce esta pendiente, podemos hacer una extrapolación y obtener el valor de la probabilidad de error en el verdadero umbral. Lo que se hace normalmente es calcular tres o cuatro

pseudo-probabilidades y mediante una regresión lineal extrapolar hasta el verdadero umbral.

Analizando la expresión (6.9), observamos que para tomar pseudo-umbrales adecuados debemos elegirlos equiespaciados (en escala logarítmica). Es decir, una buena elección sería tomar como pseudo-umbrales puntos x_1 , x_2 y x_3 tales que se encuentren, por ejemplo, a 1, 2 y 3 dB por debajo del verdadero umbral (respectivamente).

En la figura 6.1 podemos observar una pdf gaussiana y una posible elección de los pseudo-umbrales. Como vemos, se desea hacer una estimación de la probabilidad de error con umbral de 0.5. Dicha estimación se hace extrapolando los resultados obtenidos para los pseudo-umbrales, según la expresión (6.9). Con las áreas de las colas que quedan a la derecha de los pseudo-umbrales tenemos que estimar el área de la cola que queda a la derecha del verdadero umbral.

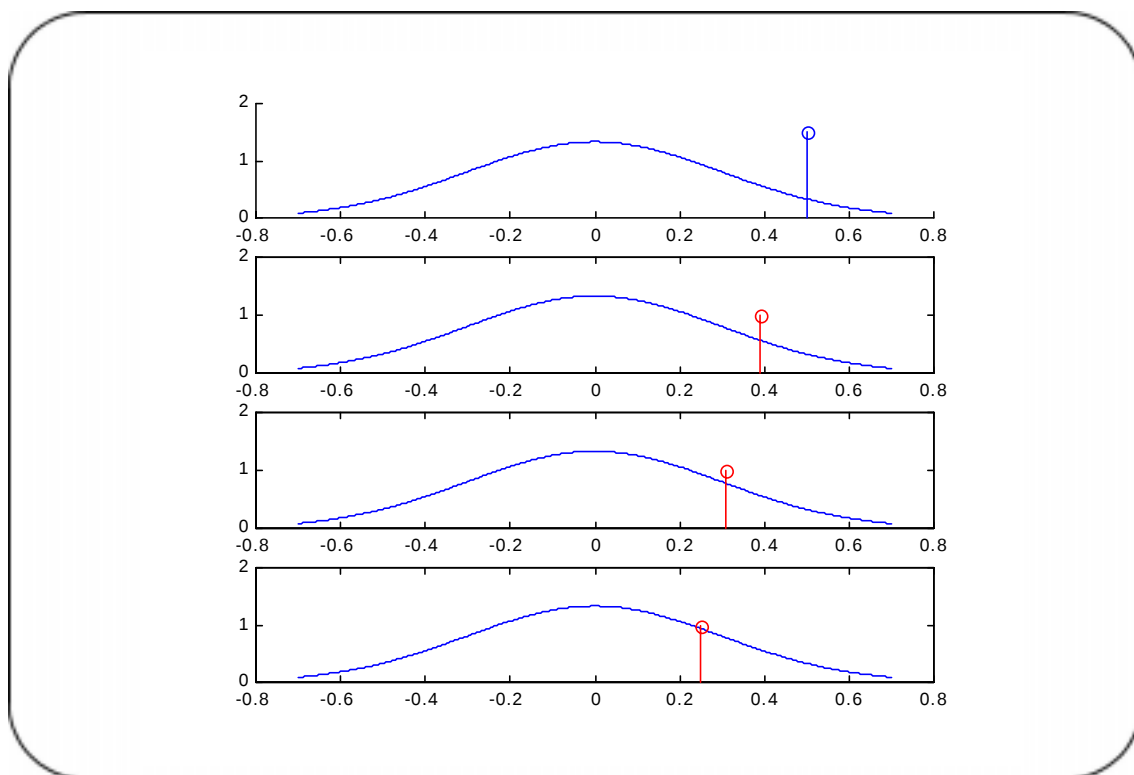


Figura 6.1 : Posibilidad de elección de los pseudo-umbrales

Las ventajas de este sistema son evidentes. Al utilizar estimaciones con pseudo-umbrales, las pseudo-probabilidades de error que se obtienen son muy grandes (relativamente grandes); por ello, una simulación de Monte Carlo en cada pseudo-umbral resulta totalmente factible. El inconveniente es que la validez de la aproximación lineal (ec. (6.9)) se restringe al caso en el que el verdadero umbral es muy alto, i.e., que la probabilidad de error que se desea estimar es muy pequeña. Afortunadamente, la región de “buen comportamiento” de este sistema coincide con la región de Monte Carlo en la que dejan de ser factibles las simulaciones, debido a la gran duración de las mismas.

6.3.2 – Propiedades del estimador

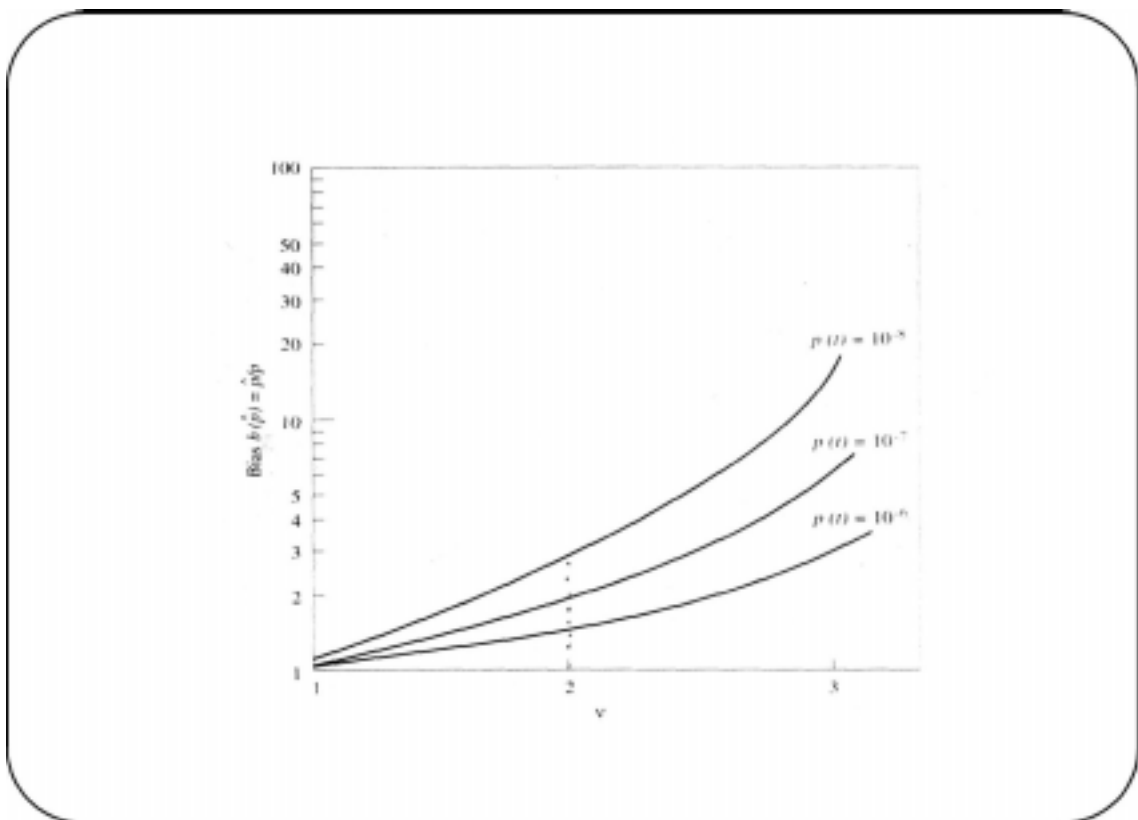


Figura 6.2 : Gráfica del sesgo (“bias”) del estimador

En la construcción del estimador se ha hecho la hipótesis de que para realizar la estimación se han utilizado tres pseudo-umbrales. En realidad, si estamos en la zona de validez de la aproximación (6.9) tan sólo nos

hacen falta dos pseudo-umbrales, ya que dos puntos proporcionan toda la información de la recta de extrapolación. Sin embargo, el tener un punto más nos permite hacer un ajuste por mínimos cuadrados, por lo que el estimador así logrado será mejor.

La forma exacta del estimador depende del número de umbrales que se utilizan, y del nivel en el que se sitúan. Si la distribución del estimador que resulta de aplicar la regresión lineal fuera conocida, podríamos calcular todas sus propiedades. Desafortunadamente, esto no es posible. Dado que las probabilidades de error en los distintos pseudo-umbrales tienen una cierta variabilidad estadística (se calculan con una simulación de Monte Carlo), el estimador de la probabilidad de error real (el extrapolado) presenta un sesgo (“bias”) inherente al ajuste por mínimos cuadrados. Definiremos el sesgo como

$$b(\hat{p}) = \frac{\hat{p}}{p}, \quad (6.10)$$

donde $\hat{p}$ es el estimador de la probabilidad de error y p es la probabilidad de error que se desea estimar (la teórica). Si elegimos los pseudo-umbrales a 1 dB, 2 dB y 3 dB por debajo del verdadero umbral, una gráfica del sesgo frente al valor del parámetro v (ecuación 6.7) se encuentra representada en la figura 6.2. La gráfica se encuentra parametrizada para diferentes valores de p . Sin embargo, como veremos, esto no será un problema grave en las simulaciones que realizaremos.

7.1 – INTRODUCCIÓN

En este capítulo vamos a estudiar la implantación de un sistema VDSL sobre un bucle de abonado. Para ello debemos implementar tanto el transmisor como el receptor DMT.

Se implementarán todos los bloques componentes del sistema DMT estudiados en el capítulo 4, particularizando los parámetros característicos del sistema para el caso de VDSL. Además, realizaremos también simulaciones con igualación.

El medio de transmisión será, como hemos dicho, un bucle de abonado. Utilizaremos modelos realistas de función de transferencia y ruido, aunque con algunas simplificaciones que iremos comentando sobre la marcha.

En cuanto al resto de los parámetros necesarios para configurar la transmisión (probabilidad de error, potencia de transmisión, anchura de banda ocupada, número de subcanales, etc.), se elegirán de la manera más acorde posible con los estándares publicados hasta la fecha ([9], [11], [12], [14]).

El objetivo final del capítulo es, en primer lugar, mostrar una idea clara y precisa del desarrollo de las etapas de procesamiento de señal, estudiadas en el capítulo 4, que tienen lugar en un modem DMT; en segundo lugar, se simularán sistemas con igualación, se presentarán los resultados obtenidos y se sacarán conclusiones.

7.2 – CONFIGURACIÓN DE LA TRANSMISIÓN

En este apartado vamos a concretar los parámetros básicos que se utilizarán en la transmisión. Para ello tenemos que proporcionar tanto las características del medio de transmisión (par trenzado de cobre) como las del equipo transmisor (modem VDSL).

7.2.1 – El transmisor

Vamos a utilizar unas especificaciones de transmisión acordes con las utilizadas en los estándares:

- *Potencia máxima transmitida: $P_{\text{máx}}=11.5 \text{ dBm}$.*
- *Ancho de banda de la transmisión: $BW=15 \text{ MHz}$.*
- *Número de subcanales: $N=256$ (los 7 de menor frecuencia no son usados).*
- *Probabilidad de error de símbolo deseada: Aunque realmente es $p_e=10^{-7}$, nosotros usaremos $p_e=10^{-3}$ y $p_e=10^{-1}$ para acortar los tiempos de simulación a niveles razonables.*
- *Margen de protección utilizado: $\gamma_m=6 \text{ dB}$, aunque en las simulaciones en las que se pretenda detectar la probabilidad de error no se incluirá este parámetro en el sistema.*
- *Cyclic prefix utilizado: $CP=60$.*

Con la potencia máxima elegida, el transmisor emite una densidad espectral de potencia de -60 dB/Hz . Los 15 MHz de banda utilizados implican enviar muestras a una tasa igual o superior a los 30 Mhz , para asegurar su correcta transmisión. Los subcanales más bajos no se utilizan con idea de no perturbar el resto de los servicios (POTS, RDSI) que puedan coexistir en el bucle. El margen de 6 dB ha sido elegido según las recomendaciones de otros estudios realizados. La elección del CP se justificará en el apartado 7.2.2, dedicado al medio de transmisión.

7.2.2 – El medio de transmisión

Vamos a utilizar un par trenzado de cobre, situado en el seno de un cable de pares de 25 pares. Dicho cable tiene una longitud de 1.5 kft*. Su modelo en el dominio z viene dado por [13]

$$H(z) = \frac{0.1 - 0.1z^{-2}}{1 - 1.5z^{-1} + 0.54z^{-2}}. \quad (7.1)$$

Su respuesta al impulso se encuentra representada en la figura 7.1. Como puede verse, no se ha considerado en el modelo el retardo de propagación (la respuesta impulsiva comienza a tomar valores en el índice cero), ya que lo único que provoca son problemas de sincronización que no tienen mayor relevancia en las simulaciones que vamos a realizar. Por otra parte, obsérvese cómo a partir del índice 60, aproximadamente, las muestras son todas despreciables. Por tanto, queda justificado (en este caso particular) el uso de tal valor para el CP.

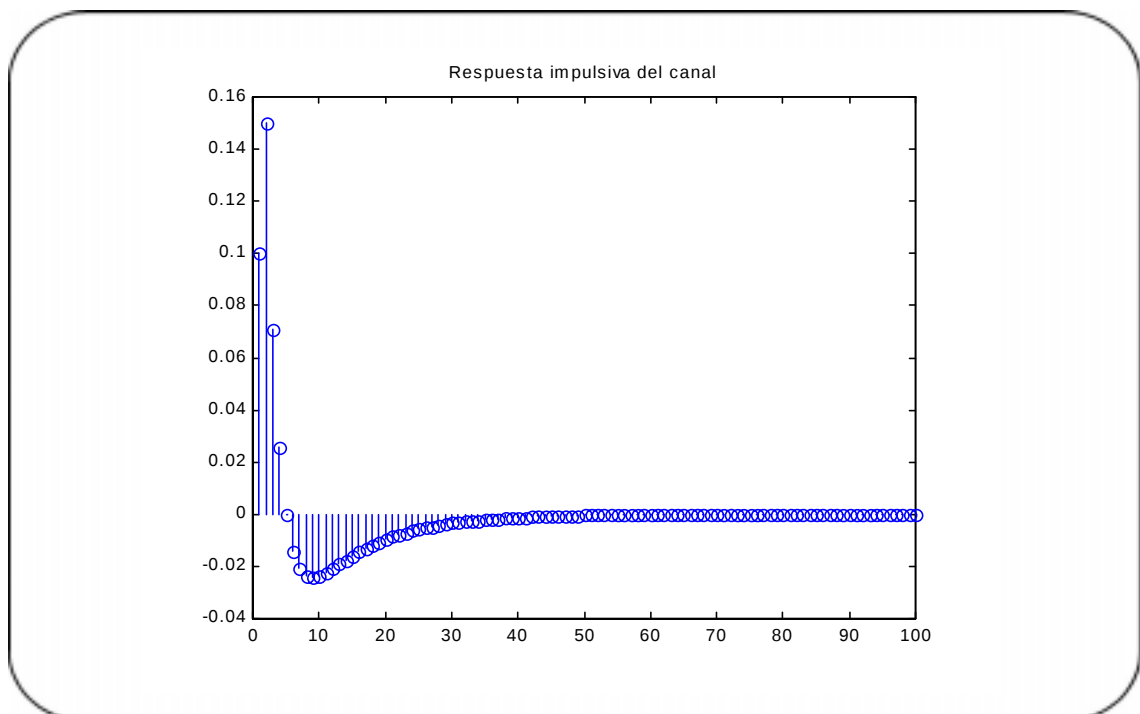


Figura 7.1: Modelo de canal utilizado en las simulaciones. Respuesta al impulso

* 1 kft equivale aproximadamente a unos 300 m

En cuanto al ruido, podemos calcularlo a partir de las expresiones obtenidas en el capítulo 3. Rescatando la ecuación (3.16), la función de transferencia de potencia para el crosstalk (en dB) venía dada por

$$10\log|XT|^2 = K_1 + K_2 \log(N) + K_3 \log(f) + K_4 \log(l), \quad (7.2)$$

donde $K_1, \dots, K_4$ vienen dados en la tabla 3.1 (capítulo 3). Tomaremos los valores correspondientes al caso promedio. Una vez que se ha calculado el parámetro XT (tanto para el NEXT como para el FEXT), las densidades espectrales de potencia se calculan con las expresiones (3.17), (3.18) y (3.19), y quedan como se ve en la figura 7.2.

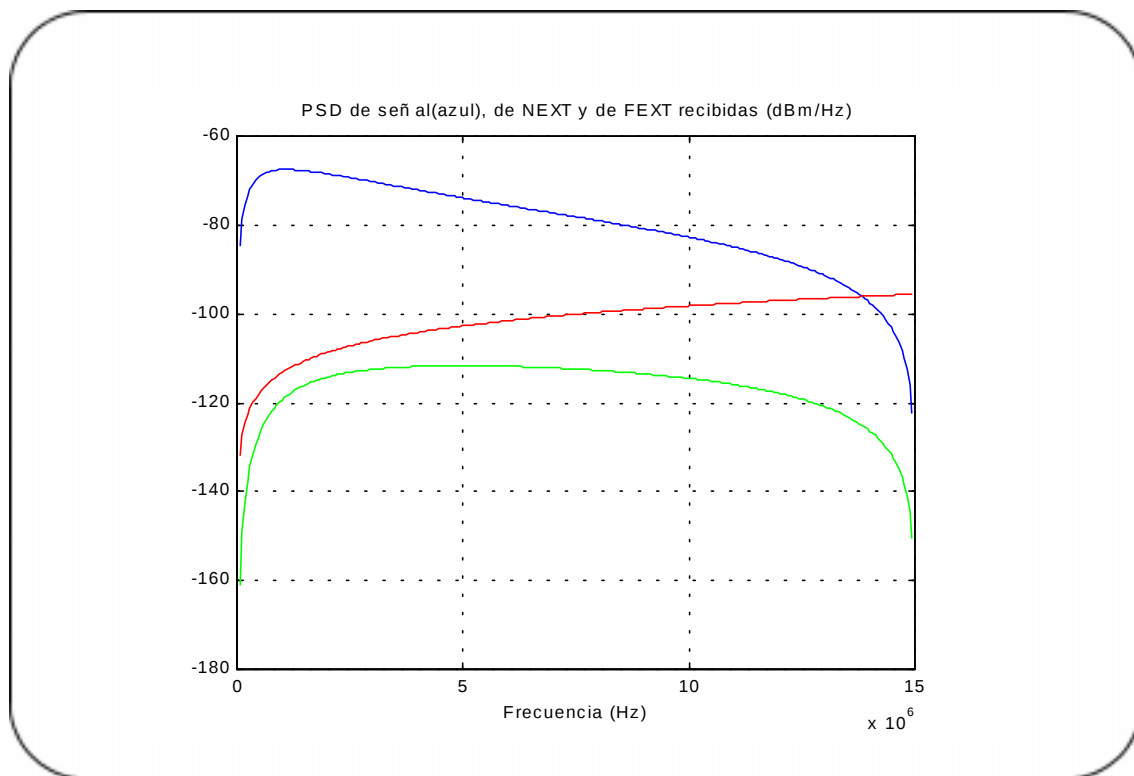


Figura 7.2: Situación de las señales en el canal: Señal útil, ruido NEXT y ruido FEXT

En la figura anterior tenemos una representación de las tres densidades espectrales de potencia anteriores (se han tomado $N=256$ muestras equiespaciadas en la frecuencia, desde 0 hasta 15 MHz). Como puede observarse, la potencia del ruido tipo FEXT (verde) es despreciable frente a la del NEXT para altas frecuencias. Si sumamos (en unidades naturales) las PSD

de NEXT y de FEXT obtenemos la PSD de ruido XT total. Cuando hagamos simulaciones, a la salida del canal tenemos que añadir un ruido que se ajuste a esta PSD total.

Por otro lado, la resta (en dB) de la potencia de la señal recibida (ecuación (3.17)) menos la potencia total de ruido ((3.18) + (3.19)) es lo que hemos llamado SNR normalizada. Esta SNR es la que se va a utilizar en la fase de entrenamiento para calcular la asignación de bits a los diferentes subcanales.

7.3 – FASE DE ENTRENAMIENTO. ASIGNACIÓN DE BITS A SUBCANALES

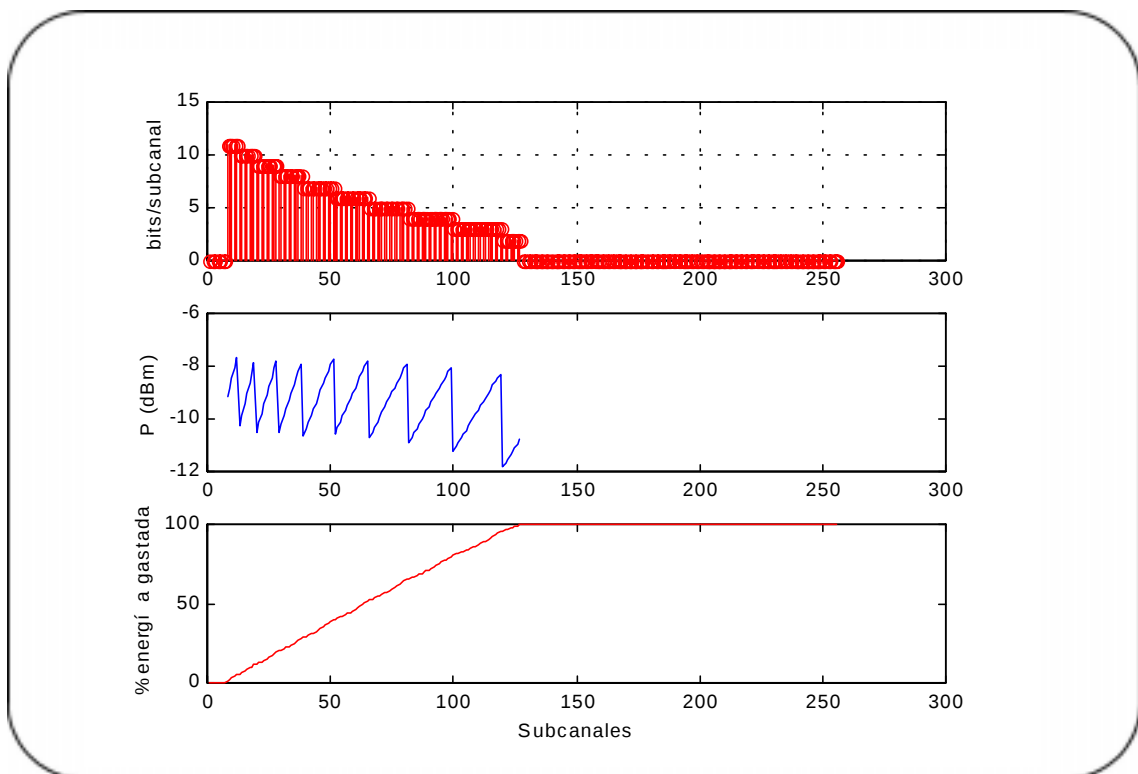


Figura 7.3 : Resultado del algoritmo de asignación

En la figura 7.3 se ha representado el resultado del algoritmo óptimo discreto aplicado al caso particular que estamos analizando. En la gráfica superior se representa el número de bits por símbolo que se asigna a cada

subcanal; en la gráfica del medio se representa, en dBm, la potencia que el transmisor pone en cada subcanal; y en la inferior, el porcentaje de energía gastada en el proceso de asignación. Como vemos, cuando se gasta el cien por cien de la energía, el algoritmo de asignación acaba y ya no podemos seguir asignando bits a subcanales.

Si nos fijamos en la gráfica del medio, observaremos que tiene forma de diente de sierra, y que las discontinuidades son saltos de 3 dB. El motivo ya se ha visto: en un sistema QAM como el que estamos simulando, mandar un bit menos por un subcanal equivale a ahorrar 3 dB de potencia (la nueva constelación tiene ahora la mitad de los puntos que la anterior). Sin embargo, como la calidad de los canales empeora a medida que aumenta su índice, cada vez tenemos que gastar más potencia para garantizar la probabilidad de error (zonas lineales de la gráfica en diente de sierra). Cuando la potencia exigida aumenta demasiado, el algoritmo opta por asignar un bit menos al siguiente subcanal, lo cual provoca el ahorro de 3 dB ya comentado (zonas abruptas de la gráfica).

Ya se dijo en el capítulo 3 que no se iban a considerar otras fuentes de ruido diferentes de las de NEXT y FEXT ya incluidas en el modelo; pero ni que decir tiene que si se hubieran incluido fuentes de ruido como interferencias de banda estrecha (provocadas por emisiones de radiodifusión) o zonas de gran atenuación en la función de transferencia (splitters, filtros de separación de canales ascendente y descendente, etc.), el algoritmo habría detectado la mala calidad de los correspondientes subcanales y habría hecho una asignación acorde con estas características. Si hubiera sido necesario, el algoritmo habría inhibido la transmisión por tal clase de subcanales.

7.4 – GENERACIÓN DE SÍMBOLOS DMT Y TRANSMISIÓN

Llegados a este punto, tras haber finalizado el proceso de asignación de bits, y una vez que el transmisor conoce toda la información, se procede al inicio de la transmisión propiamente dicha.

En el caso que venimos analizando, tras el proceso de asignación se obtiene la siguiente información (aparte de los vectores de asignación de bits y de distribución de energías):

```
>>>>EJECUTANDO ALGORITMO DE CARGA<<<<
Tasa binaria (Mb/s): 3.613636e+001
El rendimiento del enlace es: 8.951049e-001
Un simbolo DMT se compone de 689 bits
SUBCANALES USADOS: 120
>>>>FIN DEL ALGORITMO DE CARGA<<<<
```

El rendimiento del enlace se calcula teniendo en cuenta las pérdidas por la inserción del CP, como se vio en el capítulo 4. Ahora nos interesa el dato del número de bits por símbolo DMT, 689. Se prepara entonces el buffer de carga del transmisor para almacenar estos 689 bits, los cuales se agrupan en paquetes según el vector de asignación de bits.

El codificador QAM debe generar tantas constelaciones como subcanales utilizados. El vector de asignación de bits es utilizado para determinar el tamaño de la constelación (por ejemplo, un subcanal con 3 bits por símbolo implica la utilización de una constelación 8-QAM). Por su parte, el vector de distribución de energías se utiliza para determinar la distancia entre los puntos de la constelación, la cual es función de la energía promedio con la que son transmitidos los símbolos.

Aquí es necesario hacer una importante apreciación, ya que las constelaciones QAM que se diseñan en este codificador son un poco diferentes de las habituales. Las constelaciones suelen diseñarse en el dominio del tiempo, es decir, las proyecciones de cada punto de la constelación sobre los ejes nos indican la amplitud con la que se modulan las portadoras temporales.

En este caso, esto no es cierto, ya que la constelación se diseña en el dominio de la frecuencia (antes de hacer la IFFT). Si recordamos, en la expresión (4.4)

$$d^2 = \frac{6\varepsilon}{M-1}, \quad (7.3)$$

el parámetro ε se corresponde con la energía promedio, si ha sido medida en el dominio del tiempo (como la suma de los cuadrados de las muestras temporales). Como nosotros tenemos que diseñar la constelación antes de la IFFT, tenemos que hacer unos cambios en la expresión (7.3), la cual queda como

$$d^2 = \frac{6\varepsilon}{M-1} \times \frac{2 \cdot N}{2}. \quad (7.4)$$

El factor que se añade al numerador responde al hecho de que hay una FFT de por medio, de tamaño $2 \times N$. Por su parte, el 2 añadido al denominador se debe a la duplicación de muestras enviadas (la generación de N muestras artificiales con simetría conjugada [16]).

En las figuras 7.4 y 7.5 se muestran constelaciones generadas y transmitidas para dos subcanales (se representan con una distancia entre símbolos normalizada a $d=1$). Se puede notar que los símbolos transmitidos (en rojo, asteriscos) se ajustan perfectamente a los puntos de la constelación (azules, círculos). En el receptor, los símbolos llegarán con ruido, por lo que formarán una nube alrededor del punto de la constelación. Esto se mostrará en posteriores figuras.

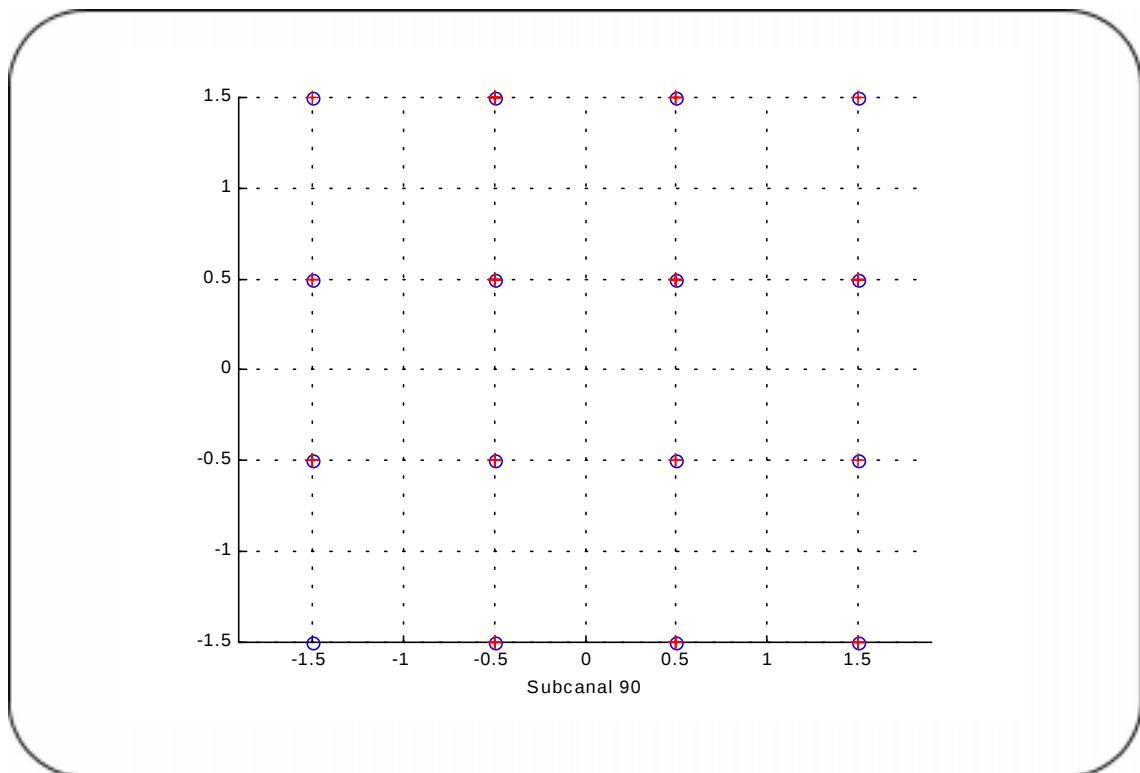


Figura 7.4 : Constelación para un número par de bits/símbolo

Otro aspecto importante a comentar es el del mapeado de los símbolos en la constelación. Es decir, cómo asignamos a cada símbolo posible de la constelación su palabra binaria correspondiente. En realidad, no nos hemos preocupado mucho de cuidar este aspecto, ya que no es relevante en el estudio de un sistema DMT. Nos hemos limitado a hacer una asignación sencilla. Para ello, utilizamos la primera mitad de los bits de la palabra binaria para codificar la componente en fase (para constelaciones con número impar de bits/símbolo, redondearemos el número de bits por arriba), y la segunda mitad para codificar la componente en cuadratura. Como sabemos, en una constelación QAM, las componentes en fase y en cuadratura se detectan de forma independiente.

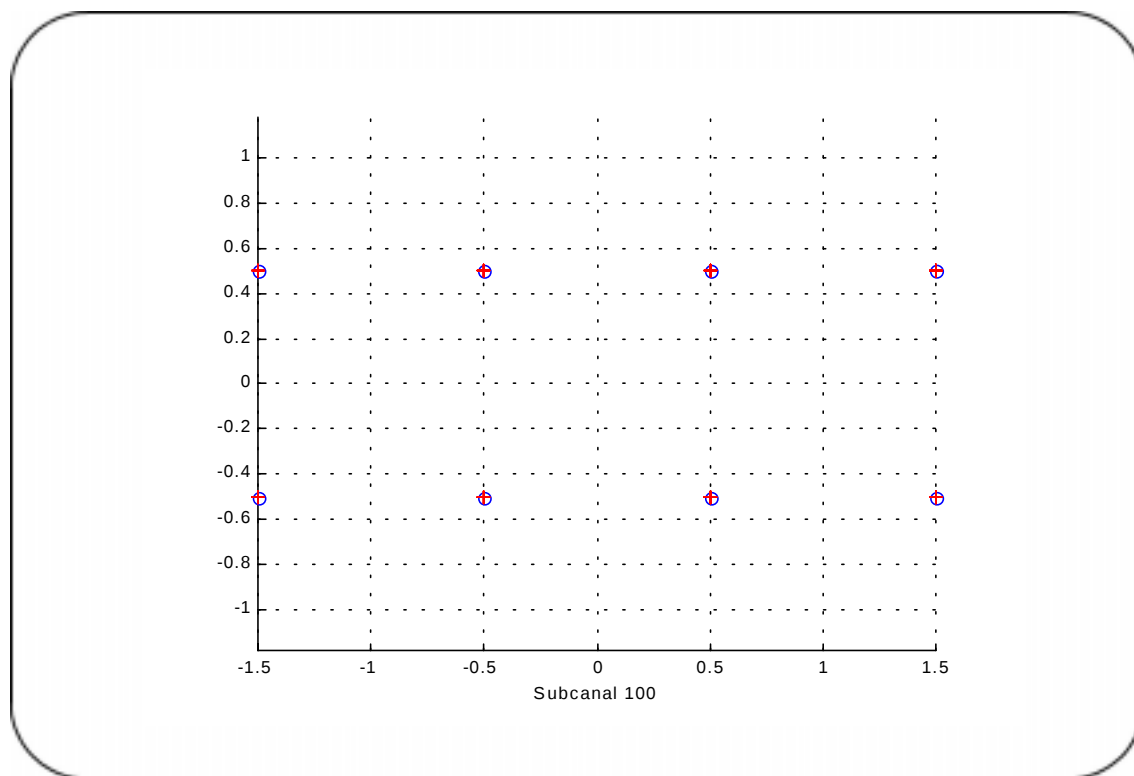


Figura 7.5 : Constelación para un número impar de bits/símbolo

Continuando ya con el procesado en transmisión, comentamos por último que una vez que se han generado en paralelo los símbolos complejos correspondientes a cada símbolo DMT (salida del codificador QAM), se procede a aplicar la simetría conjugada y a antitransformar, para obtener las muestras que se envían hacia el receptor a través del canal.

Eso sí, como ya sabemos, antes de la transmisión por el canal se procede a la inserción del prefijo cíclico. Esta operación se realiza en todos los símbolos DMT que se transmiten, y garantiza (siempre que se escoja la longitud adecuada) la separación correcta de los subcanales en recepción.

7.5 – EFECTO DEL CANAL SOBRE LA SEÑAL

En este apartado vamos a discutir lo que ocurre con la señal al atravesar el canal. Para ello, nos vamos a abstraer de la estructura de la señal (transmisión continua y consecutiva de tramas de longitud $2N+v$ muestras) y nos vamos a centrar en la forma en la que vamos a simular el canal de transmisión.

Para dicha simulación, se supone conocido:

- El modelo en el dominio z del bucle de abonado, o, alternativamente, su respuesta impulsiva.
- La densidad espectral de potencia de ruido presente en el canal, para generar a partir de ella muestras temporales que se sumen a la señal.

En este caso, disponemos de la función de transferencia en z (recuérdese que fue utilizada al principio del capítulo para poder ejecutar el algoritmo de carga adaptativa, ec. (7.1)).

En cuanto al ruido, lo que haremos es sumar las potencias de NEXT y de FEXT, de forma que obtengamos para cada subcanal un valor de la PSD de ruido. A partir de esta PSD generaremos un proceso aleatorio cuyas muestras iremos sumando con las de la señal, a la salida del canal.

Pero vayamos por partes. En primer lugar, lo que hay que hacer es aplicar sobre la señal el efecto de la función de transferencia. La manera de incluir este efecto es convolucionar la respuesta impulsiva del canal (de longitud generalmente menor de 200 muestras) con la señal de entrada (cuya longitud en una transmisión típica es de miles o millones de muestras). Debido a la gran diferencia en la longitud de ambas señales, en la simulación

se ha utilizado para el filtrado la función de MATLAB *fftfilt*, que implementa internamente una técnica del tipo *overlap & save*.

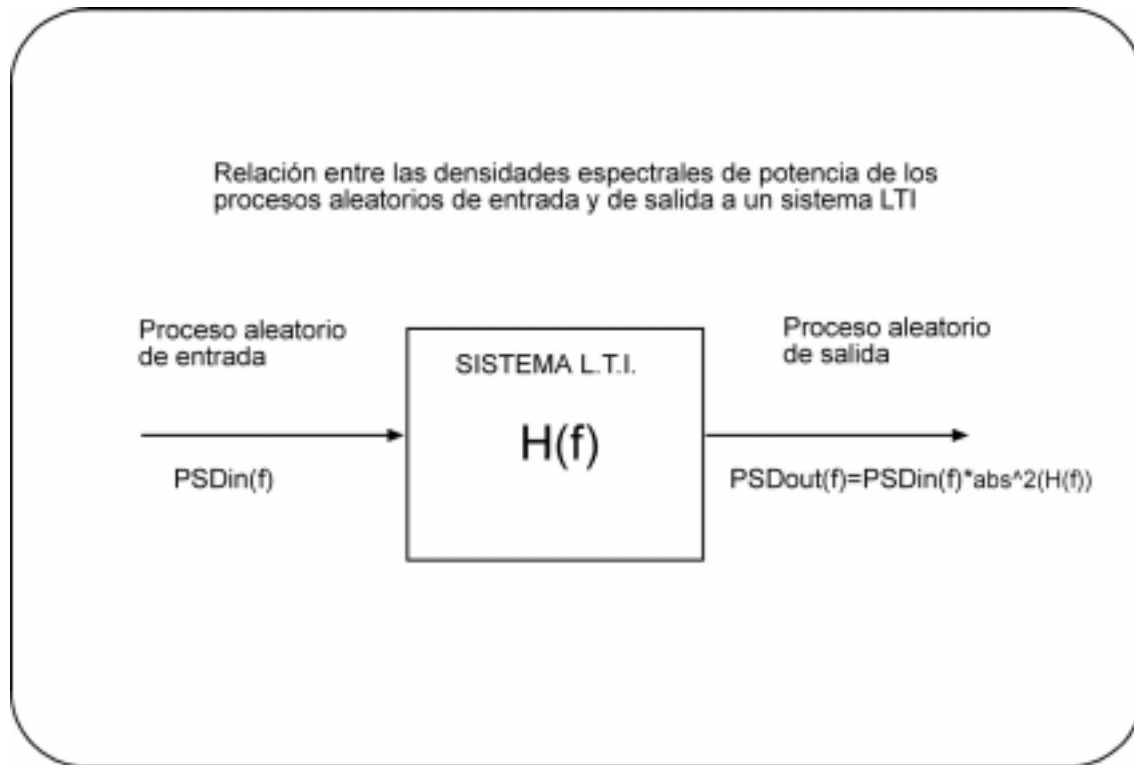


Figura 7.6 : Propiedad utilizada para generar el ruido

El ruido se ha generado pasando ruido blanco gaussiano por un filtro para conseguir ruido NEXT+FEXT. La idea es la siguiente: si a la entrada de un filtro con función de transferencia $H(j\omega)$ ponemos un proceso aleatorio con una cierta densidad espectral de potencia, a la salida obtenemos otro proceso aleatorio cuya densidad espectral de potencia es la misma que la del proceso de entrada, pero multiplicada por el cuadrado de la función de transferencia del canal. Un esquema de esta propiedad puede verse en la figura 7.6.

Para generar el ruido, tenemos que seguir entonces los siguientes pasos:

- Consideramos un vector de N muestras reales, una para cada subcanal, donde tenemos almacenada la PSD del ruido objetivo.

- Extraemos la raíz cuadrada a cada una de las muestras del vector anterior.
- Formamos un nuevo vector de longitud $2xN$, extendiendo simétricamente el vector del paso 2.
- Calculamos la IFFT del vector anterior. Obtenemos $2xN$ muestras reales y simétricas (el vector antes de la IFFT era real y par) correspondientes a la respuesta impulsiva de un filtro, al que vamos a denominar filtro auxiliar.
- Generamos un proceso aleatorio gaussiano de potencia unidad, y lo pasamos a través del filtro auxiliar.
- Se desechan las primeras muestras de salida del filtro auxiliar, para evitar los efectos del transitorio.

El proceso de salida tiene una PSD igual a la PSD objetivo, en virtud de la propiedad comentada anteriormente. En la figura 7.7 se muestran unas gráficas de la PSD objetivo (en dBm/Hz, correspondientes al NEXT+FEXT) comparada con otra correspondiente a la PSD del ruido generado. Como podemos observar, el ajuste es casi perfecto.

Ya que sabemos cómo generar las muestras de ruido, lo que hacemos es crearlas e ir las sumando a la señal. De esta forma, obtenemos finalmente las muestras de entrada al receptor.

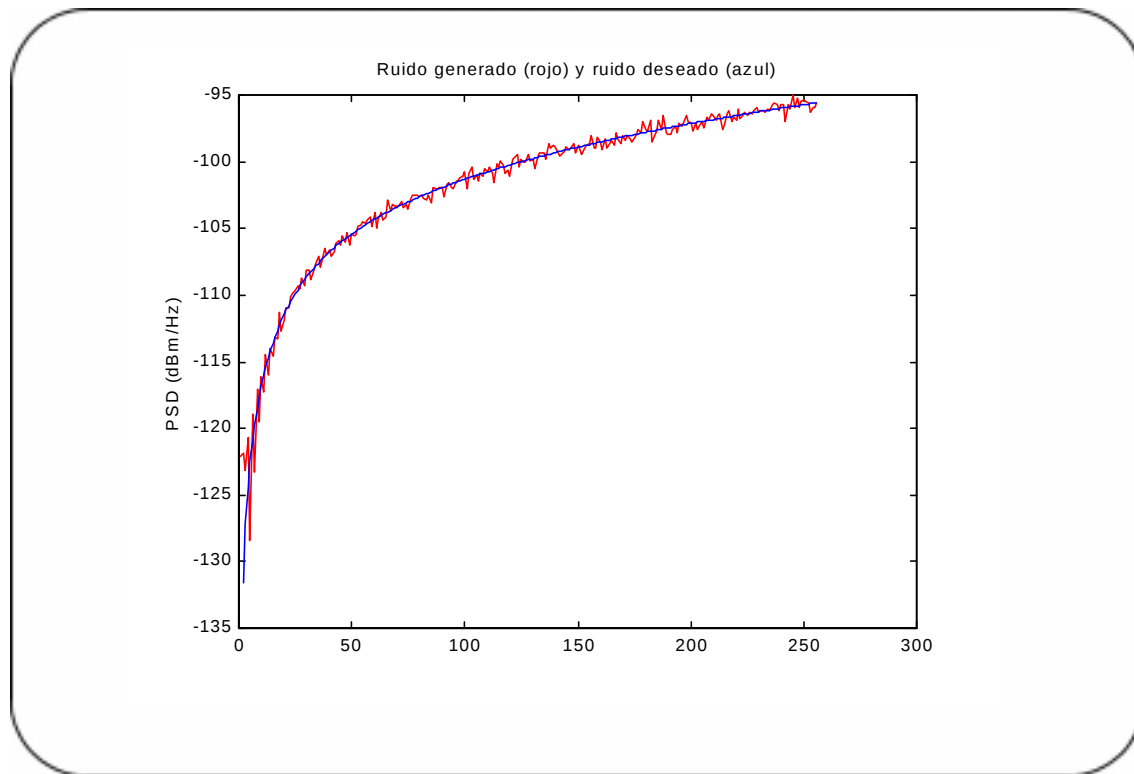


Figura 7.7 : PSD objetivo junto con la PSD generada

7.6 – RECEPCIÓN DE SÍMBOLOS DMT Y DETECCIÓN.

En esta sección mostraremos cómo se ha llevado a cabo el procesamiento de la señal en el modem del receptor. Por la forma en que se ha realizado la transmisión de los símbolos DMT tenemos garantizada la independencia entre ellos, por lo que podemos hacer detecciones símbolo a símbolo.

En primer lugar, la señal recibida se va agrupando en símbolos DMT, cuya longitud es de $2xN+v$ muestras. Posteriormente, se va eliminando el CP de cada bloque, según la metodología explicada en el capítulo 4. De todas formas, la verdadera complejidad del procesado en recepción comienza con el bloque de demodulación.

Como sabemos, este bloque lo que hace es tomar las $2 \times N$ muestras reales de cada símbolo DMT y aplicar sobre ellas una transformación FFT. Como las muestras de entrada son reales, las de salida guardan una relación de simetría. Por este motivo, se desechan las últimas N muestras, y nos quedamos con las N primeras, complejas. Con estos N números, tenemos que hacer N estimaciones, para averiguar el símbolo que ha sido transmitido por cada subcanal. Para ello, antes de efectuar la detección, se utiliza un proceso de normalización de ganancia.

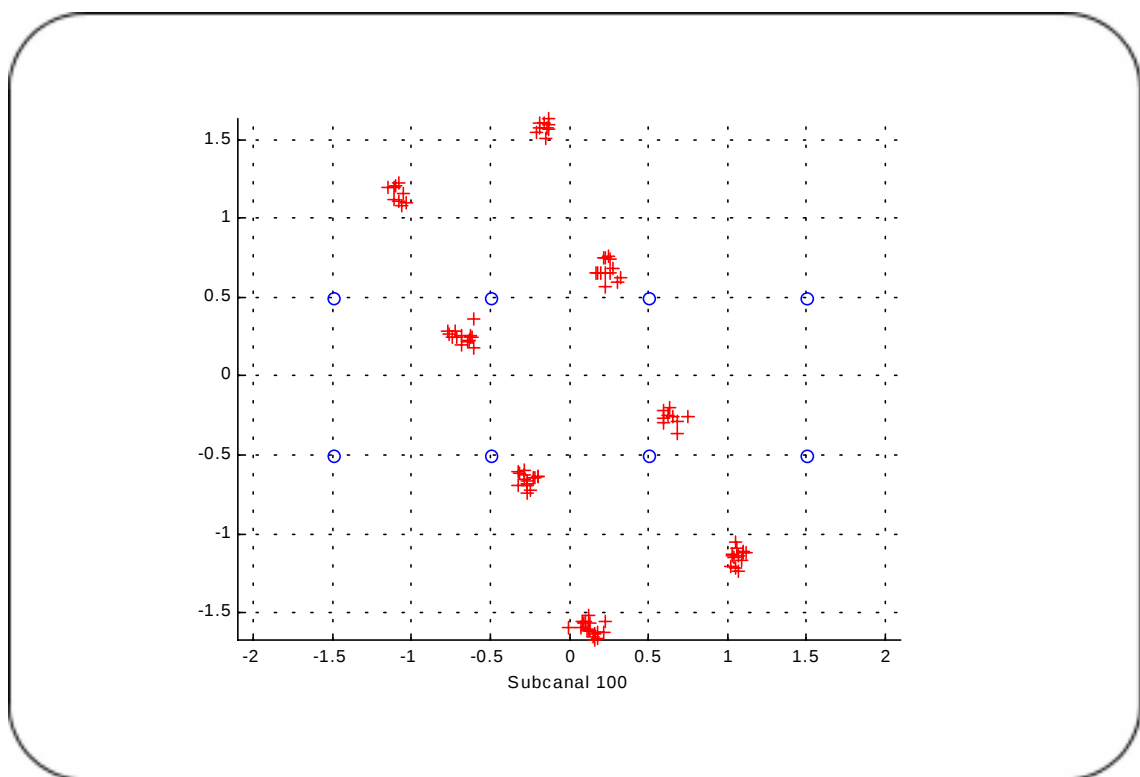


Figura 7.8 : Constelación recibida si no se corrige la fase en la normalización

El proceso de normalización de ganancia se lleva a cabo para compensar el cambio en amplitud y en fase que sobre cada subsímbolo provoca el canal. Como ya se estudió en el capítulo correspondiente, este proceso se hace inmediatamente después del bloque demodulador FFT, y se lleva a cabo de forma independiente para cada subcanal. La normalización se hace multiplicando la salida de cada subcanal por un número complejo. Por ello, su único efecto sobre la señal es cambiar la amplitud y la fase de la misma. Se supone que el canal es ideal en el ámbito de cada subcanal, y que,

por tanto, no es necesario ningún proceso de igualación dentro del mismo, siendo suficiente con la multiplicación por una ganancia compleja. Como muestra de la necesidad de este procesado, se muestra en la figura 7.8 una constelación correspondiente a un subcanal donde se normaliza por un coeficiente real (en lugar de complejo). Este hecho tiene como consecuencia que no se compense la distorsión de fase introducida por el canal, lo cual se manifiesta en un giro de los puntos de la constelación recibida en torno al origen. De paso, podemos observar cómo el ruido ha hecho que los puntos recibidos se dispersen formando una nube.

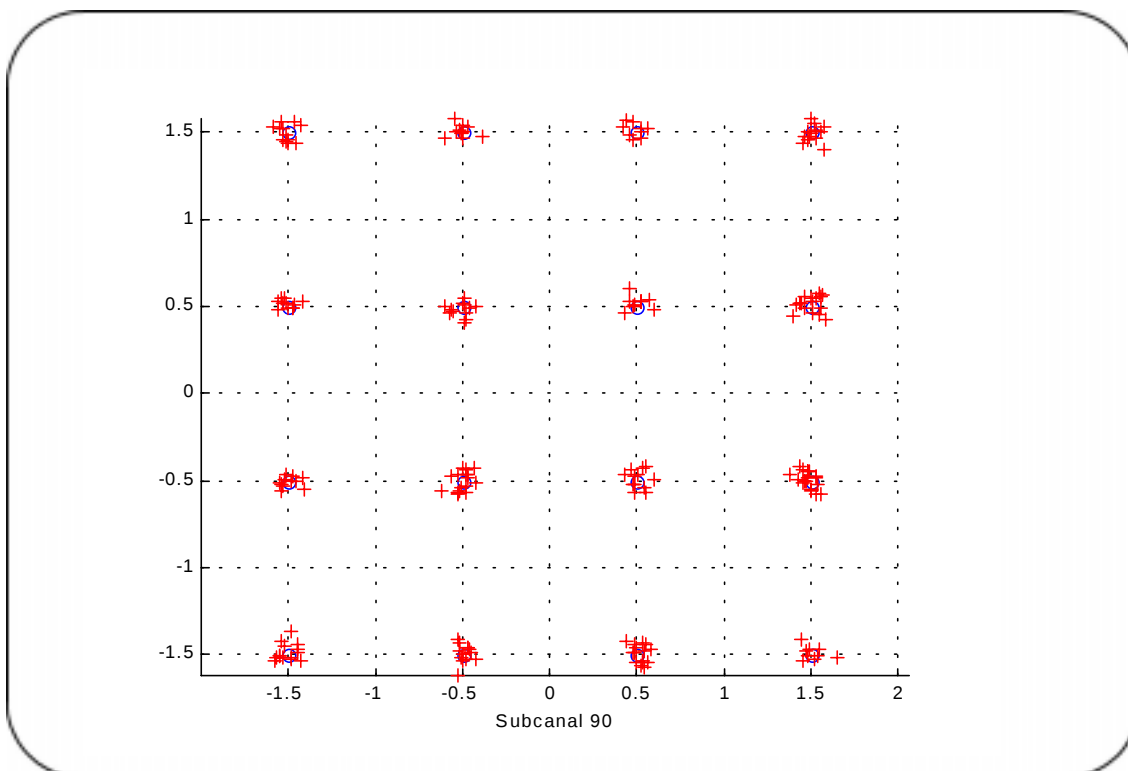


Figura 7.9 : Constelación recibida en un subcanal. Efecto del ruido

Con posterioridad a la normalización se procede a la detección de los subsímbolos, la cual se realiza también de forma independiente para cada subcanal. El decodificador tiene que decidir a partir de cada subsímbolo recibido cuál ha sido el subsímbolo transmitido. Para ello, se sigue un criterio MMSE, es decir, a cada subsímbolo recibido se le asigna el punto de la constelación más próximo en el sentido de los mínimos cuadrados. Una

constelación QAM, además, nos permite realizar la detección de las componentes en fase y en cuadratura de forma independiente.

Una vez que se han detectado (de forma simultánea) todos los subcanales, se puede afirmar que se ha recibido completamente un símbolo DMT. Este símbolo DMT es susceptible de ser comparado con el transmitido, con el fin de evaluar la cantidad de errores que el sistema ha introducido. Si se realiza el proceso generación+transmisión+recepción+detección para muchos símbolos DMT, la estimación de la probabilidad de error de símbolo se puede hacer con mayor fiabilidad. En la figura 7.9 mostramos una constelación típica recibida. También podemos ver cómo el ruido dispersa los puntos alrededor del valor teórico.

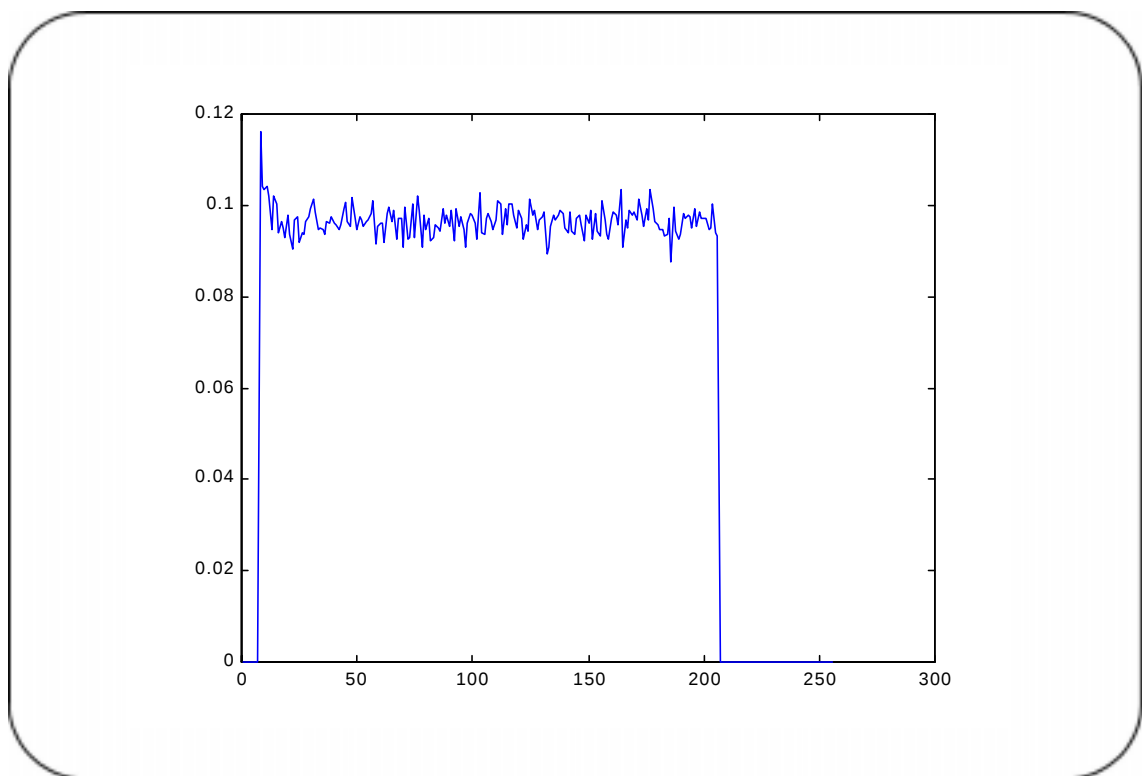


Figura 7.10 : Distribución de la probabilidad de error a lo largo de los subcanales.

Como ilustración, se muestra en la figura 7.10 la gráfica de distribución de probabilidad de error a lo largo de los diferentes subcanales utilizados, medidas en la transmisión de una gran cantidad de símbolos DMT (cantidad lo bastante grande como para asegurar que la varianza de la

probabilidad de error obtenida es suficientemente pequeña). La figura 7.10 se ha generado con la eliminación del margen de 6 dB y aumentando la probabilidad de error a un valor de 10^{-1} , con el fin de provocar un mayor número de errores que permitan hacer una mejor estimación de la p_e .

Podemos observar que la distribución de errores es uniforme a lo largo de todos los subcanales que componen la transmisión. Como se han cambiado algunos parámetros (margen y probabilidad de error deseada), el algoritmo de asignación ha llegado a una solución distinta. En este caso, ahora se han utilizado más subcanales en la transmisión. En cuanto a la uniformidad de la distribución de los errores, recordemos que fue uno de los requisitos que impusimos al algoritmo de asignación.

Vamos ahora a observar qué ocurre si acortamos el CP y no utilizamos igualación para compensarlo. Para ello mostramos las gráficas 7.11, 7.12 y 7.13 donde se representa la distribución de errores a lo largo de los diferentes subcanales para con un CP insuficientemente largo.

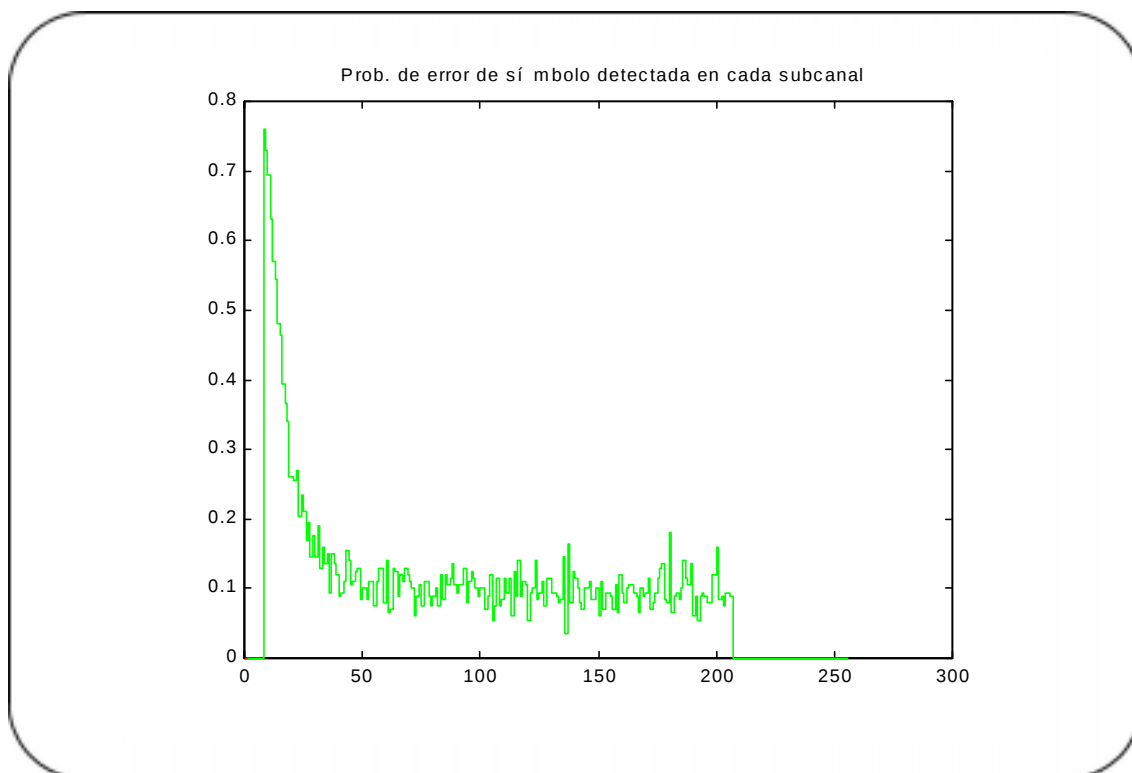


Figura 7.11: Sin igualación; CP=20, distribución de errores

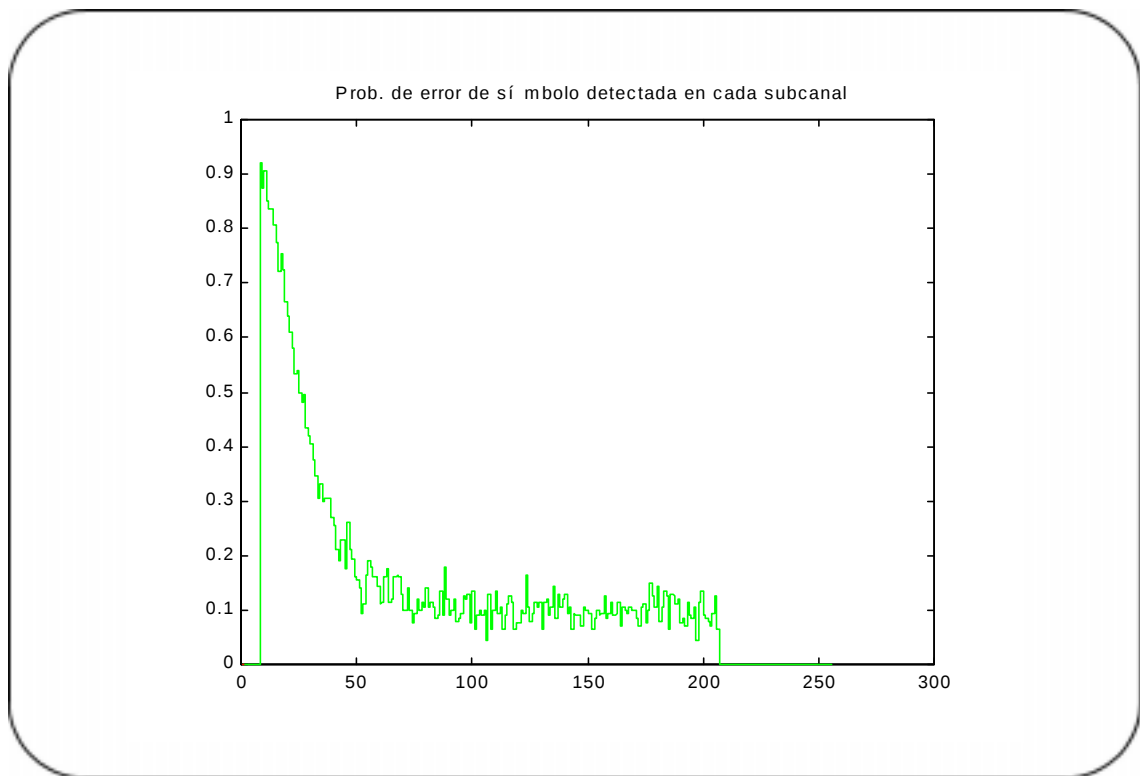


Figura 7.12: Sin igualación; CP=10, distribución de errores

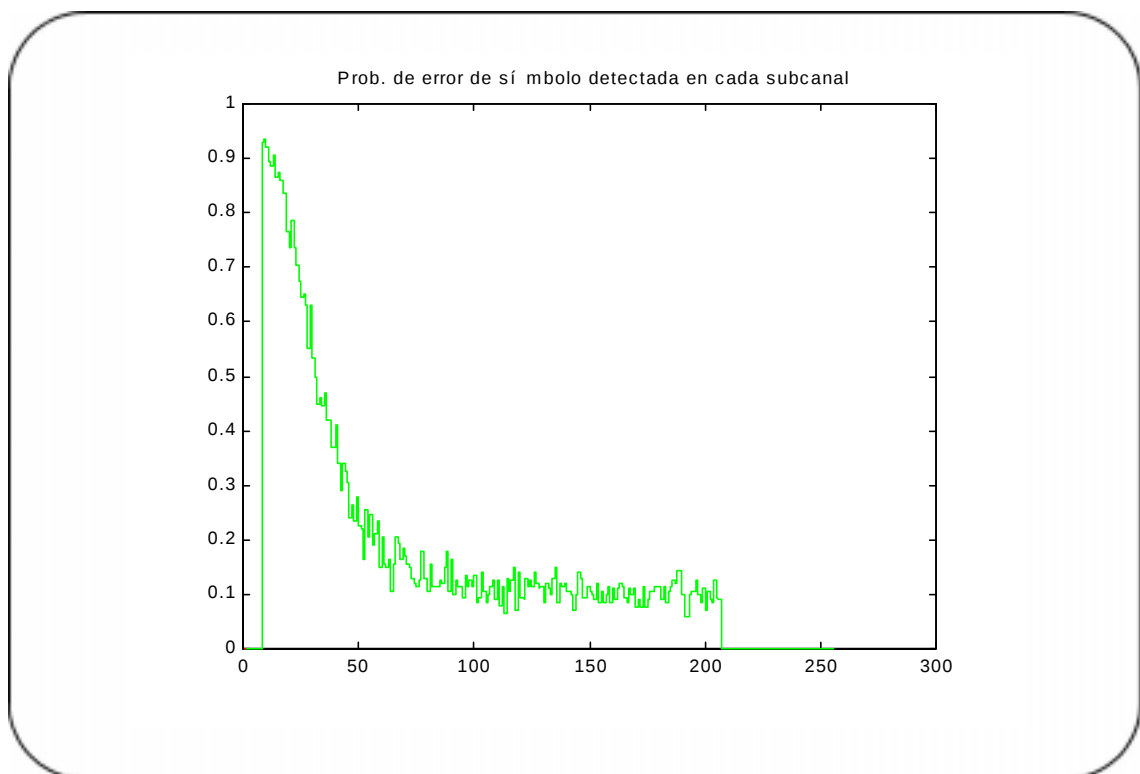


Figura 7.13: Sin igualación; CP=5, distribución de errores

Como podemos observar, la utilización de un CP inadecuado provoca una acumulación de errores en los primeros subcanales, de manera que la distribución de errores por subcanales deja de ser uniforme. Este efecto es más acusado cuanto menor sea el CP que utilicemos. En efecto, dado que la misión del CP es eliminar el efecto de la IBI (interferencia entre bloques DMT), cuanto menor sea el CP, mayor será la cantidad de interferencia.

Sin embargo, uno podría pensar que la IBI es un fenómeno que debería afectar por igual a todos los subcanales, dado que es algo que ocurre en el dominio del tiempo. Vamos a tratar de explicar por qué los errores debidos a la IBI sólo se manifiestan en los subcanales bajos. Como hemos adelantado, la IBI es un fenómeno que ocurre en el dominio del tiempo; luego en principio, su efecto es una especie de “ruido” que afecta por igual a todos los subcanales, y que provoca una degradación en la SNR de cada subcanal. Lo que ocurre es que los canales que están transportando mayor número de bits por símbolo (que son los primeros, como ya hemos visto) son más sensibles a esta degradación (funcionan con una SNR mayor), y, por tanto, acusan más la IBI que el resto de los subcanales. Dicho de otra forma, ya que se añade una cantidad de ruido aproximadamente igual en todos los subcanales, se verán más afectados aquellos con ruido inicialmente menor (ya que en términos relativos la adición de ruido es mayor).

Acabamos de ver cómo el uso de un CP de longitud insuficiente provoca un aumento drástico en el número de errores, debido a la interferencia. Por tanto, si se quiere acortar el CP debe usarse algún mecanismo de igualación que permita eliminar la IBI.

7.7 – UTILIZACIÓN DE UN IGUALADOR TEQ

En este apartado vamos a estudiar el comportamiento del igualador en el dominio del tiempo (TEQ) expuesto en el desarrollo teórico que se hizo en los apartados 5.2 y 5.3.

Lo primero que hay que hacer es inicializar el igualador. Para ello necesitamos conocer la función de transferencia del canal, la distribución estadística del ruido y la energía promedio con la que se lleva a cabo la transmisión de cada símbolo. Los datos anteriores son todos conocidos, y emplearemos los mismos que se están utilizando a lo largo del capítulo. También necesitamos imponer los parámetros de diseño, que son N_b (longitud de la TIR deseada), N_f (número de coeficientes del igualador) y Δ (retraso relativo entre la señal igualada y la señal convolucionada con la TIR). En las sucesivas simulaciones que se realicen con el TEQ, se trabajará con una probabilidad de error de símbolo de 10^{-1} . De esta forma obtendremos un número de errores lo bastante grande para estudiar su comportamiento con fiabilidad. También mantendremos fijo el parámetro de diseño $\Delta=0$. Recordemos brevemente que para inicializar un TEQ realmente son necesarios dos pasos: primero, la obtención de la TIR a partir de la CIR, y segundo, el cálculo de los coeficientes del TEQ propiamente dicho.

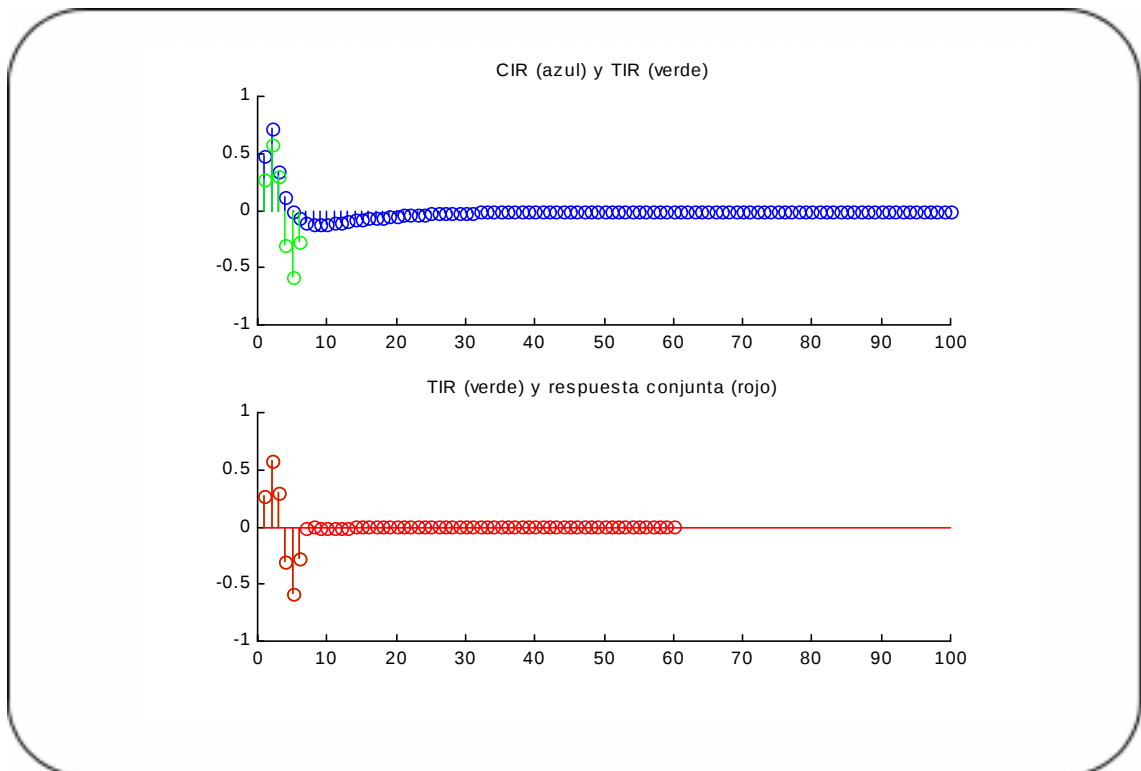


Figura 7.14 : Igualación TEQ; $N_f=8$, $N_b=5$. Respuesta impulsiva

A continuación se mostrarán gráficas de algunos de los TEQ's construidos, para ilustrar mejor el proceso de inicialización y su

funcionamiento. Como primer ejemplo vamos a usar un igualador con $N_f=8$ taps, tratando de obtener una TIR de longitud 6 ($N_b=5$). Mostramos en la figura 7.14 una gráfica en la que se representan varias cosas. En la parte superior están dibujadas la CIR y la TIR (normalizadas para que tengan energía unidad, con el fin de establecer comparaciones con más facilidad), para ilustrar el acortamiento de la respuesta impulsiva que lleva a cabo el igualador. En la gráfica inferior, se representan la TIR junto con la respuesta conjunta del canal y el igualador. Puede observarse que el igualador hace que la respuesta conjunta se adapte perfectamente a la TIR, y, de hecho, ambas gráficas se solapan.

En la figura 7.15 se representan otras tres gráficas asociadas a este igualador. Se trata de las respuestas en frecuencia (en magnitud) del canal original (gráfica superior), del igualador TEQ (gráfica del medio) y de la respuesta conjunta (gráfica inferior), que, como comprobaremos en breve, van a ser muy importantes a la hora de enjuiciar el comportamiento del igualador.

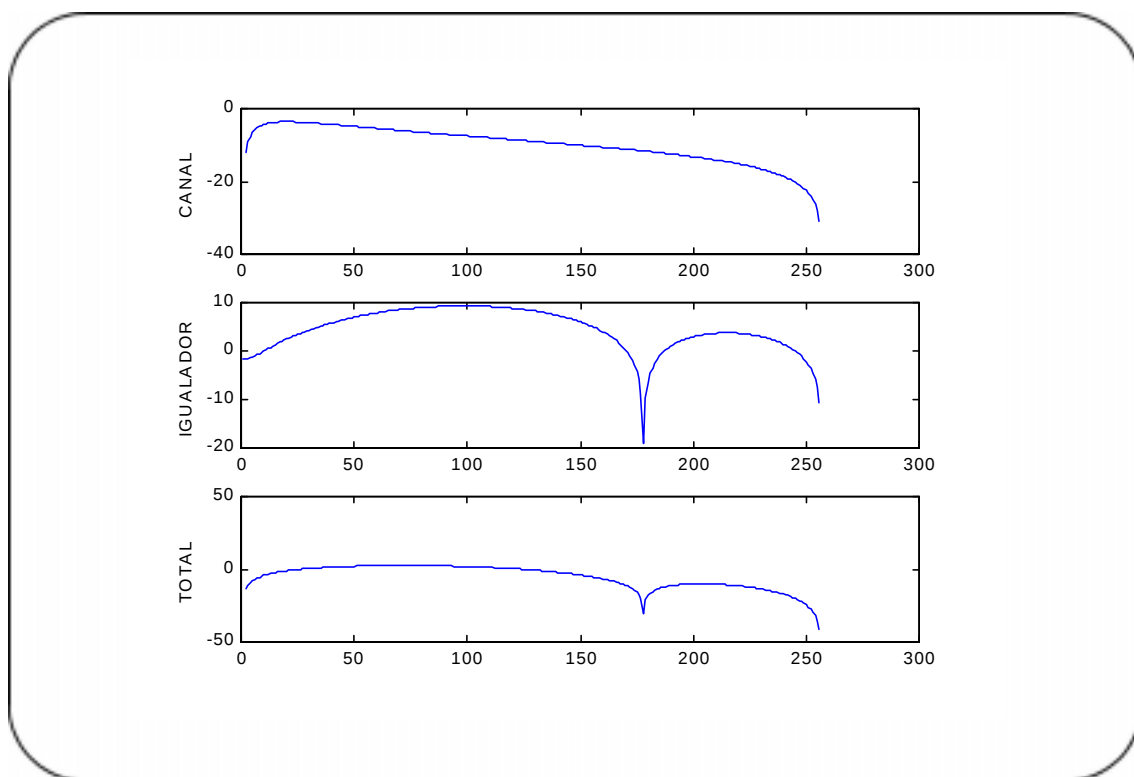


Figura 7.15: Igualación TEQ; $N_f=8$, $N_b=5$. Respuesta en frecuencia

Como vemos, los resultados obtenidos se ajustan a los esperados. Con la inserción del igualador la longitud de la respuesta impulsiva conjunta del canal y el igualador es considerablemente menor que la del canal sin igualación. Este hecho es el que nos va a permitir realizar transmisiones con un CP más corto (ahora el CP necesario es de N_b).

Sin embargo, los resultados obtenidos en cuanto a probabilidad de error no son tan satisfactorios. De hecho, con la inserción de un igualador TEQ hemos comprobado que la probabilidad de error detectada en el sistema es mayor que la esperada. En el caso que venimos analizando ($N_f=8$, $N_b=5$) la probabilidad de error es 1.30×10^{-3} , que es 1.3 veces mayor que la nominal (valor obtenido con el método de Monte Carlo).

La primera explicación que a uno se le ocurre es pensar que siguen apareciendo fenómenos de interferencia; pero se ha observado que la probabilidad de error en el sistema igualado no mejora ni incluso seleccionando un CP de 60 (valor requerido por el sistema sin igualación), luego esta hipótesis queda descartada.

Realizando un análisis más minucioso de la situación, mostramos en la gráfica 7.16 la distribución de errores a lo largo de los subcanales para el sistema igualado con el TEQ que venimos estudiando.

Si comparamos la figura 7.16 con la 7.15 observamos que el aumento de la probabilidad de error se debe a la acumulación de errores en los subcanales afectados por los ceros del igualador. Parece ser que las zonas del espectro en las que la respuesta en frecuencia varía muy rápidamente, como ocurre en estos ceros tan abruptos, el sistema tiene problemas para detectar los símbolos de forma correcta. Por este motivo, la probabilidad de error se dispara en los subcanales cercanos a los ceros, que son las zonas en las que la pendiente es más abrupta.

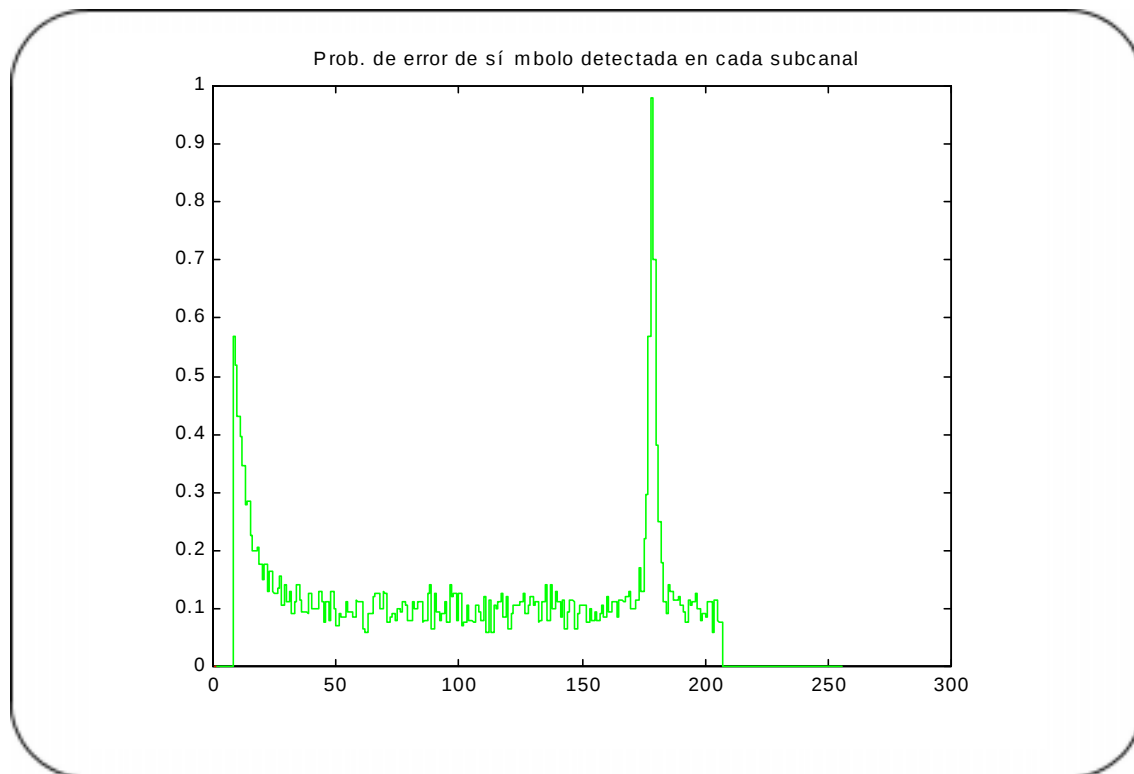
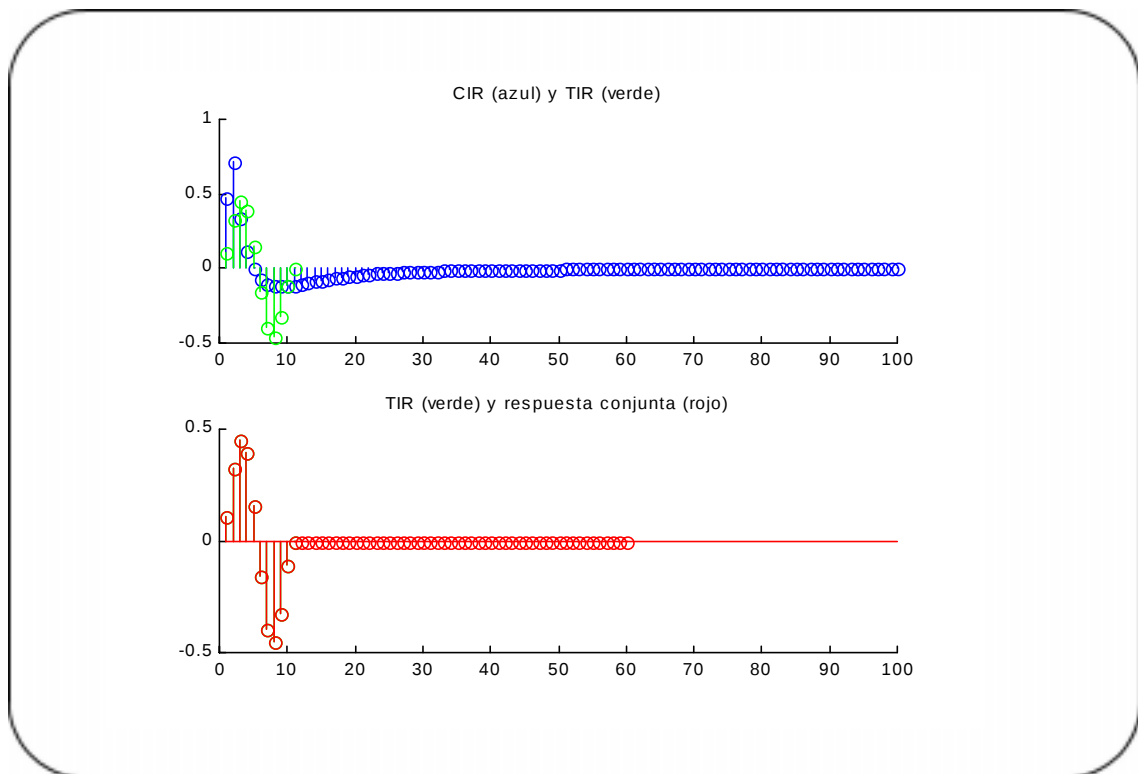
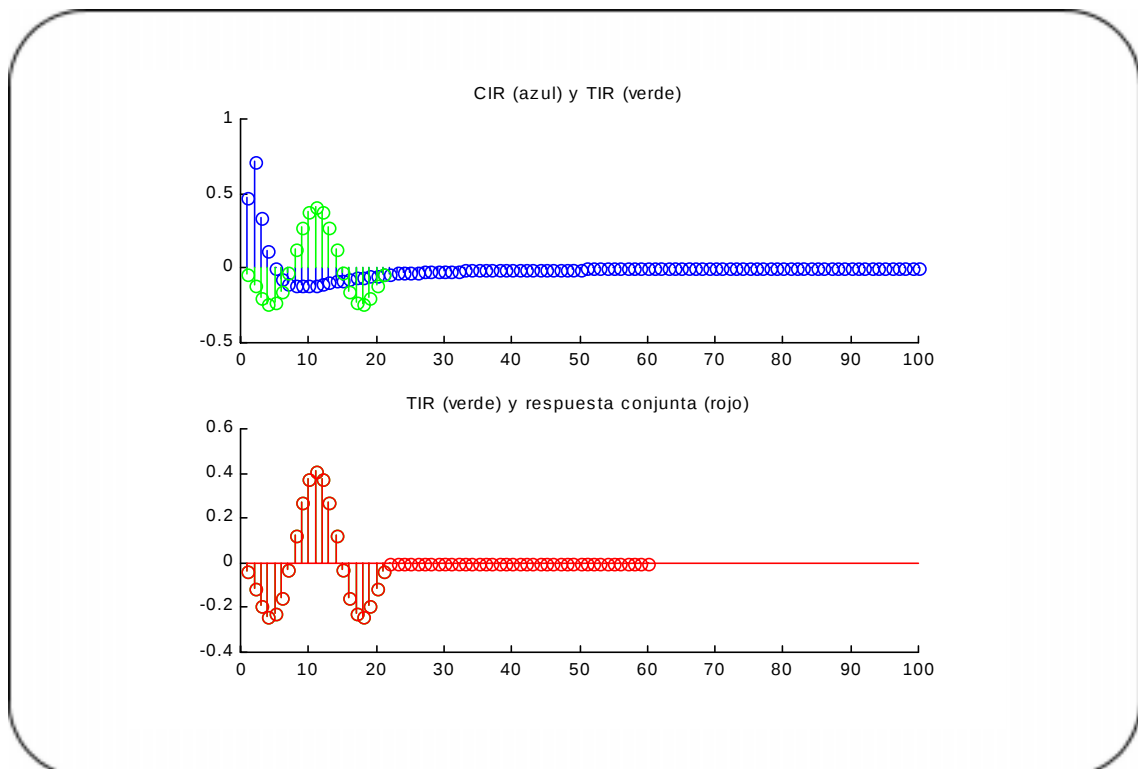


Figura 7.16: Igualación TEQ; $N_f=8$, $N_b=5$. Distribución de errores

Para comprobar que de hecho es así, vamos a construir dos igualadores más, forzando a que tengan distintas características en cuanto a posición de los ceros. En concreto, vamos a construir un igualador con $N_f=10$ y $N_b=10$, y otro con $N_f=30$ y $N_b=20$, para los cuales vamos a construir gráficas similares a las anteriores, representadas en las figuras 7.17-7.22.

¿Tiene justificación este aumento del número de errores? Una explicación puede estar relacionada con las hipótesis que se hicieron en el capítulo 4 cuando se introdujeron las técnicas MCM. En concreto, en un sistema DMT se supone que la función de transferencia es plana en el ámbito de cada subcanal. En nuestro caso, parece que esta hipótesis deja de cumplirse en los subcanales cercanos a los ceros. Como el receptor normaliza cada subcanal multiplicando por un solo coeficiente (FEQ 1-tap), no está compensando de forma adecuada la distorsión introducida por el canal. Se han hecho diversas simulaciones orientadas a la obtención de resultados en este sentido, y todo parece apuntar a que es ésta la causa. Esto justificaría


 Figura 7.17 : Igualación TEQ; $N_f=10$, $N_b=10$. Respuesta impulsiva

 Figura 7.18 : Igualación TEQ; $N_f=30$, $N_b=20$. Respuesta impulsiva

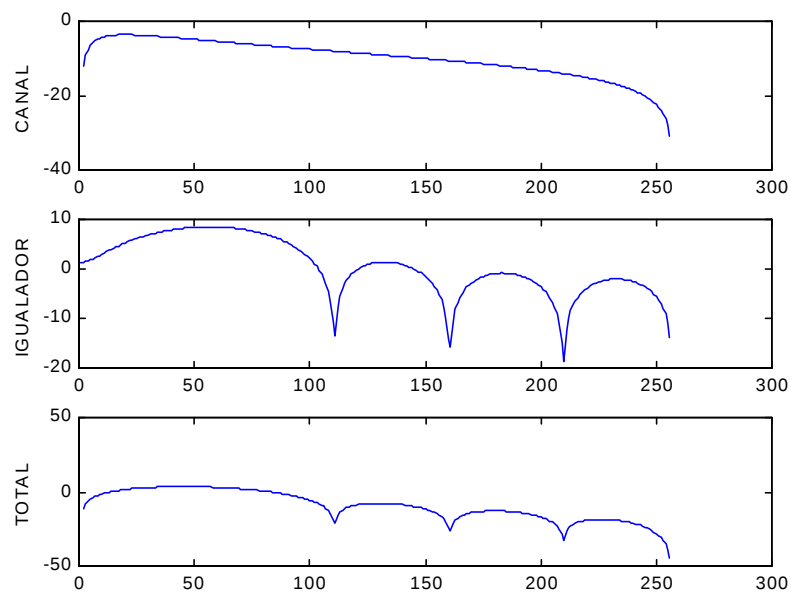


Figura 7.19 : Igualación TEQ; $N_f=10$, $N_b=10$. Respuesta en frecuencia

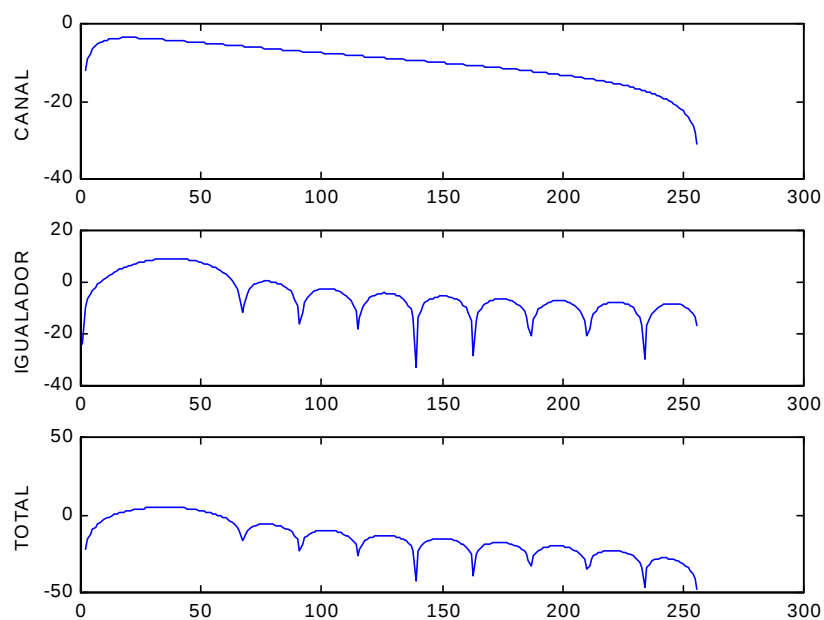


Figura 7.20 : Igualación TEQ; $N_f=30$, $N_b=20$. Respuesta en frecuencia

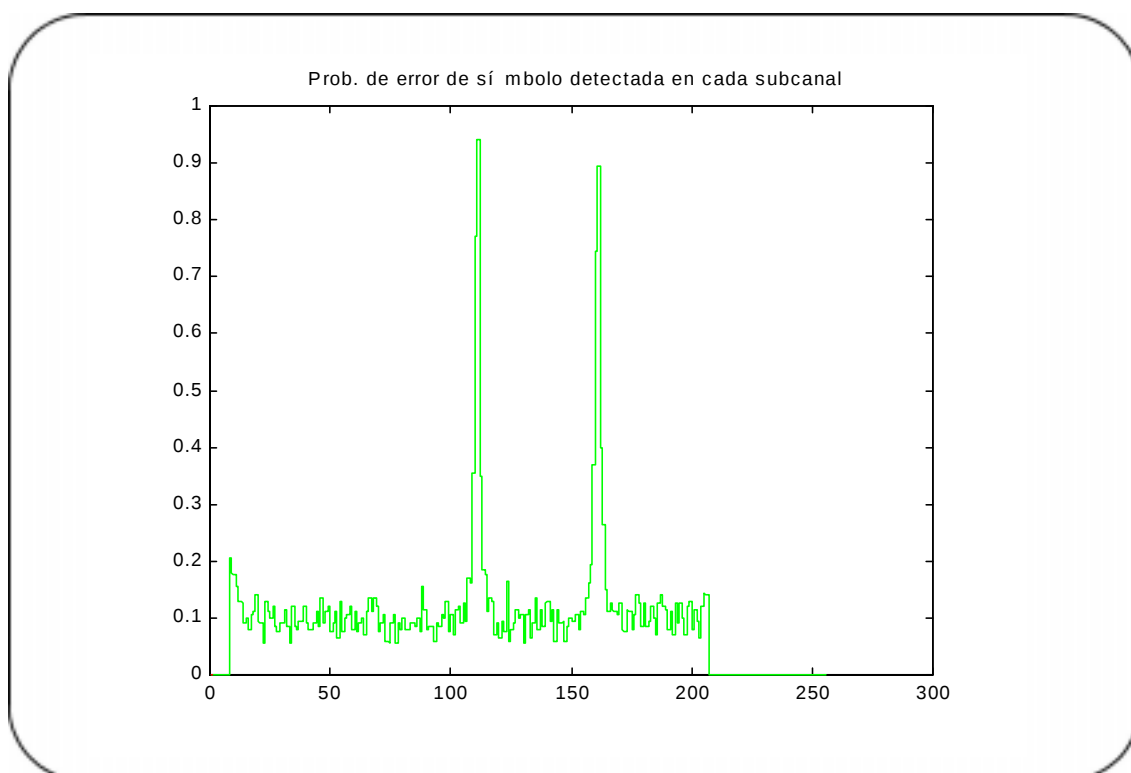


Figura 7.21 : Igualación TEQ; $N_f=10$, $N_b=10$. Distribución de errores

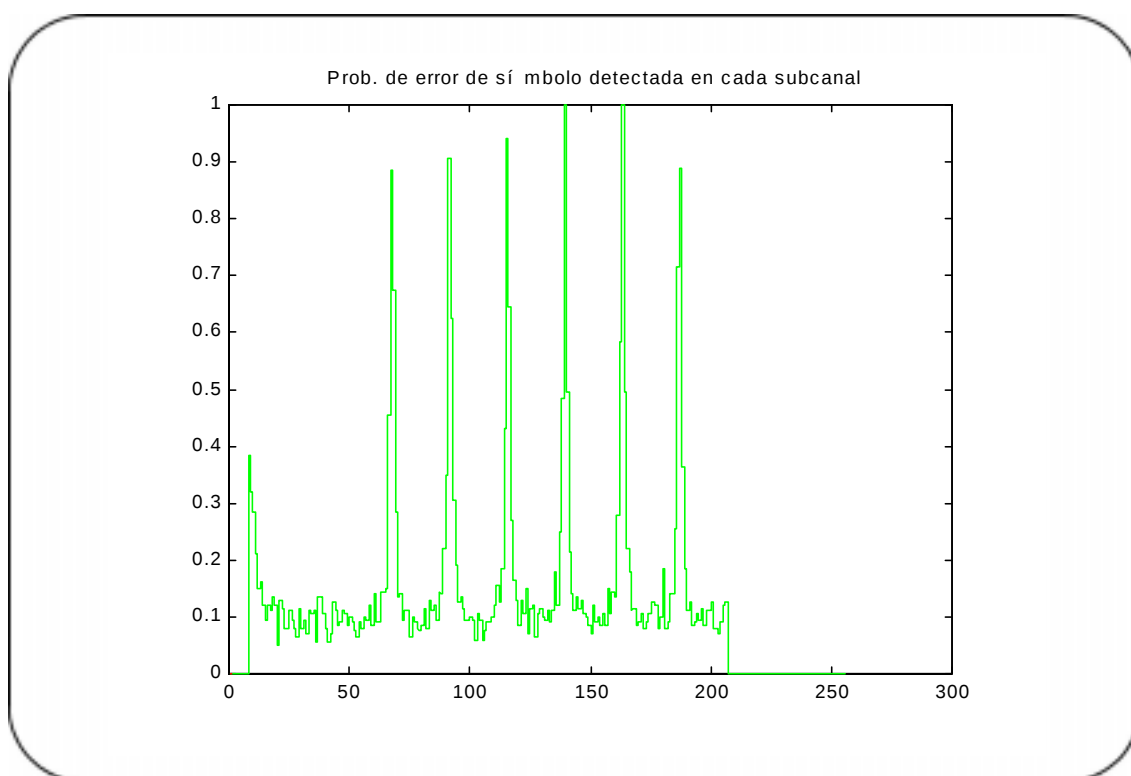


Figura 7.22 : Igualación TEQ; $N_f=30$, $N_b=20$. Distribución de errores

también que ni siquiera la elección de un CP más largo que N_b solucione el problema. Con esta justificación, queda claro que la bondad del MMSE-TEQ no va a depender del error de ajuste entre las respuestas impulsivas del canal igualado y la TIR, sino que va a estar condicionada a la situación de los ceros del igualador.

Se han hecho también simulaciones utilizando un mayor número de subcanales; pero los subcanales cercanos a los ceros siguen presentando este efecto. Además, un aumento del número de subcanales implica, como ya vimos en el capítulo 4, un aumento tanto de la complejidad computacional (las FFT's que se tienen que realizar son de mayor tamaño) como de los requerimientos de memoria; aparte de que con un número elevado de subcanales las pérdidas por inserción del CP son escasas y cada vez resulta menos interesante la introducción de un igualador. Por tanto, no parece una solución interesante.

Otra posible solución puede ser la realización de una búsqueda exhaustiva de igualadores (variando los parámetros de diseño) hasta encontrar uno que no presente ceros tan abruptos o bien que los presente fuera de los subcanales utilizados para la transmisión. Sin embargo, esta solución resulta muy poco práctica, ya que el proceso de inicialización puede hacerse excesivamente largo (se han realizado muchas simulaciones y se ha comprobado que resulta difícil dar con un igualador con estas características).

Otra alternativa sería inhibir la transmisión por los subcanales cercanos a los ceros. Pero si esos subcanales se encuentran en la zona “buena” del canal (zona en la que podemos enviar muchos bits por símbolo), evitarlos en el algoritmo de asignación supondría invertir esa energía en subcanales “peores”, con lo que no podríamos enviar tantos bits por símbolo y disminuiríamos la velocidad de transmisión. Cabe preguntarse entonces qué es peor, si reducir el rendimiento a causa del CP o reducir la tasa binaria como consecuencia de evitar ciertos subcanales.

En definitiva, hemos visto que el uso de un igualador MMSE-TEQ acarrea una serie de inconvenientes que no hacen aconsejable su utilización. Frente a su uso, existen dos alternativas: el TEQ geométrico (que utiliza otro criterio en la inicialización distinto al MMSE) y el FEQ (igualador en el dominio de la frecuencia).

La utilización de un FEQ se analiza en el apartado 7.8. Sin embargo, la gran complejidad que entraña el TEQ geométrico ha hecho que no lo incluyamos en nuestro estudio. Una descripción detallada de dicho igualador puede encontrarse en [1]. Señalemos aquí brevemente que el criterio que se sigue en su inicialización es el de maximizar la SNR geométrica (4.21), ya que es éste el parámetro que verdaderamente maximiza la tasa de datos en el sistema (4.19). Al final, hay que resolver un complejo problema de optimización no lineal con restricciones, y se deben usar técnicas iterativas del cálculo numérico.

7.8 – UTILIZACIÓN DE UN IGUALADOR FEQ

En este apartado se pretende estudiar el efecto del igualador FEQ sobre el sistema. Se presentan los resultados de una gran cantidad de simulaciones realizadas con distintas longitudes del CP y del igualador. La probabilidad de error se mantiene fija a 10^{-3} (para no alargar en exceso los tiempos de simulación, y a cambio, poder desarrollar un mayor número de simulaciones) y el retraso se mantiene en $\Delta=0$. Lo primero que debemos decir es que la inicialización es más compleja que en el caso del TEQ, y además se lleva a cabo de una forma menos intuitiva (no podremos obtener gráficas de una TIR ni una respuesta en frecuencia del igualador).

Tanto es así que, debido a la excesiva complejidad computacional necesaria para la inicialización del igualador, hemos tenido que optar por trabajar con un sistema un poco distinto al del estándar VDSL del capítulo 5. La modificación que hemos introducido consiste en reducir el número de

subcanales a 32 (en lugar de 256). En la fase de entrenamiento el subcanal de frecuencia más baja no será utilizado.

Los resultados de las simulaciones se van a presentar en unas tablas. La forma de actuar ha sido la de fijar para cada simulación un CP a priori, e ir viendo cómo cambia la probabilidad de error en función del número de taps con que formemos el igualador FEQ. La columna que aparece en las tablas como TAPS=1 corresponde a las simulaciones realizadas sin igualador; en este caso, el único tap se corresponde con el normalizador de ganancia.

En cuanto a los datos que nos hacen falta para la inicialización (véase el desarrollo teórico efectuado en el capítulo 5), todos son conocidos excepto la matriz de correlación de las muestras de señal y la matriz de correlación de muestras de ruido, las cuales han sido calculadas a partir de promedios realizados en una fase previa de la transmisión.

Para cada caso, se han realizado simulaciones con MATLAB, cada una de ellas basándose en el método de Monte Carlo (50000 símbolos DMT en cada ensayo). Recordemos también que existe una limitación en el número de taps del igualador en función del CP que se utilice en la transmisión, que proviene de imponer que las matrices involucradas en la inicialización (ver apartado 5.5) tengan dimensiones no nulas. De hecho, la longitud del igualador sólo puede superar en una unidad a la longitud del CP.

	CP=5	
TAPS	1	5
p_e	2.4e-1	5.4e-2

Tabla 7.1 : Probabilidades de error obtenidas con CP=5

	CP=10		
TAPS	1	5	10
p_e	1.8e-1	2.5e-2	7.8e-3

Tabla 7.2 : Probabilidades de error obtenidas con CP=10

	CP=15				
TAPS	1	5	10	14	15
p_e	1.2e-1	1.4e-2	3.4e-3	1.1e-3	7.3e-4

Tabla 7.3 : Probabilidades de error obtenidas con CP=15

	CP=20						
TAPS	1	5	10	11	12	15	20
p_e	8.4e-2	1.4e-2	1.2e-3	1.0e-3	6.2e-4	6.0e-4	3.9e-4

Tabla 7.4 : Probabilidades de error obtenidas con CP=20

	CP=25						
TAPS	1	5	10	11	15	20	25
p_e	5.4e-2	1.3e-2	1.1e-3	8.1e-4	4.5e-4	3.6e-4	3.0e-4

Tabla 7.5 : Probabilidades de error obtenidas con CP=25

	CP=30						
TAPS	1	5	10	15	20	25	30
p_e	3.3e-2	1.2e-2	1.0e-3	4.5e-4	3.4e-4	2.6e-4	2.3e-4

Tabla 7.6 : Probabilidades de error obtenidas con CP=30

	CP=35								
TAPS	1	5	9	10	15	20	25	30	35
p_e	1.8e-2	8.7e-3	1.1e-3	1.0e-3	4.3e-4	3.2e-4	2.7e-4	2.2e-4	2.0e-4

Tabla 7.7 : Probabilidades de error obtenidas con CP=35

	CP=40										
TAPS	1	5	8	9	10	15	20	25	30	35	40
p_e	7.1e-3	4.1e-3	1.9e-3	1.4e-3	8.3e-4	4.5e-4	3.8e-4	2.6e-4	2.2e-4	1.7e-4	1.6e-4

Tabla 7.8 : Probabilidades de error obtenidas con CP=40

	CP=45											
TAPS	1	5	7	8	10	15	20	25	30	35	40	45
p_e	2.5e-3	1.3e-3	1.3e-3	1.1e-3	8.0e-4	5.3e-5	3.4e-4	2.7e-4	2.2e-4	1.7e-4	1.6e-4	1.1e-4

Tabla 7.9 : Probabilidades de error obtenidas con CP=45

	CP=50											
TAPS	1	3	4	5	10	15	20	25	30	35	40	50
p_e	1.5e-3	1.1e-3	7.4e-4	6.3e-4	6.0e-4	5.1e-4	3.7e-4	2.8e-4	2.1e-4	1.7e-4	1.4e-4	1.1e-4

Tabla 7.9 : Probabilidades de error obtenidas con CP=50

Recordemos que en las simulaciones se han realizado imponiendo una probabilidad de error máxima de 10^{-3} . Analizando los resultados de las

simulaciones podemos ver cómo un aumento de la longitud del CP (manteniendo constante el número de taps) conduce a una disminución de la probabilidad de error. De la misma forma, manteniendo constante el CP, un aumento del número de taps también provoca una disminución de dicha probabilidad.

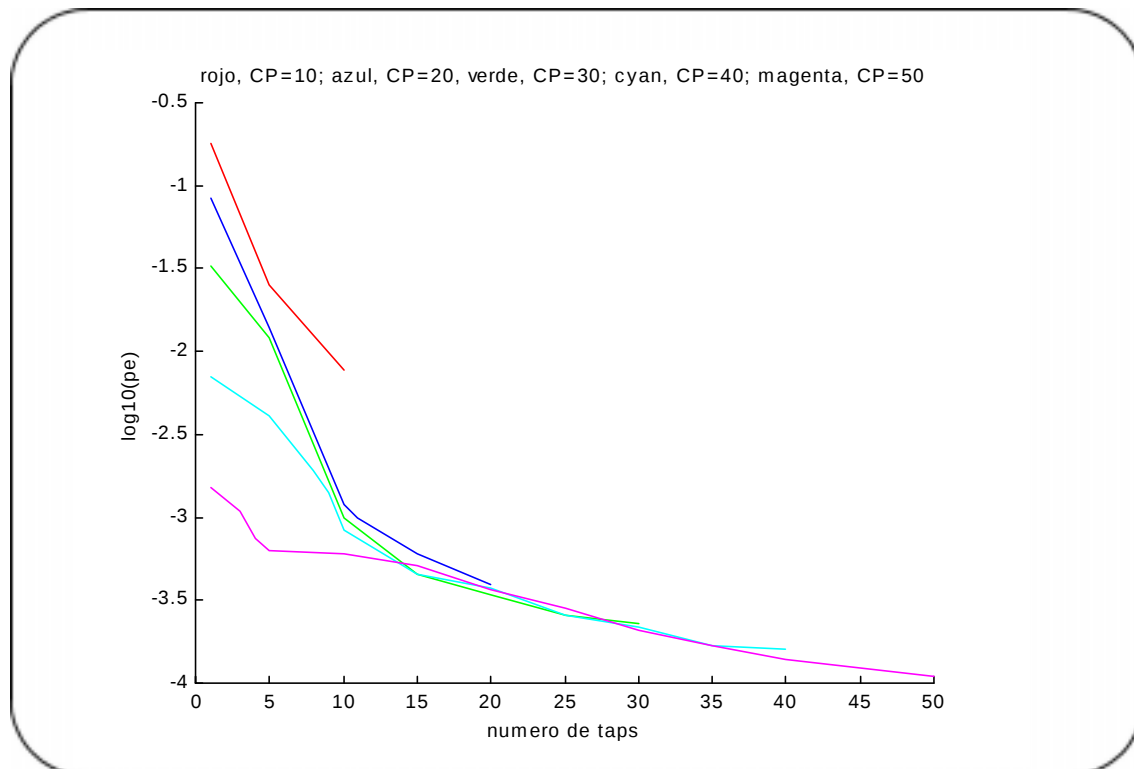


Figura 7.23 : Evolución de p_e en función del número de taps

Para tener una idea más clara de lo que ocurre, se han representado gráficamente los valores obtenidos en algunas de las tablas en la figura 7.23. Como vemos, si utilizamos un más CP largo, nos hacen falta menos taps para alcanzar la misma probabilidad de error.

En la figura 7.24 mostramos una gráfica cuyos ejes son el CP utilizado y el número de taps del igualador, y en la que se representan los puntos obtenidos para una probabilidad de error de 10^{-3} . Debido a que este igualador no puede ser más largo que el CP, nunca vamos a alcanzar este valor de probabilidad de error con un CP corto (5 o 10), porque en esta zona no podemos utilizar igualadores largos. En las simulaciones que hemos hecho,

se alcanza la probabilidad de error deseada a partir de un CP de 15 muestras, en estas condiciones son necesarios 14 taps para cancelar la interferencia. Evidentemente, según aumentamos la longitud del CP va disminuyendo el número de taps necesario para lograr la igualación, aunque no parece existir ninguna relación concreta (lineal o de otro tipo).

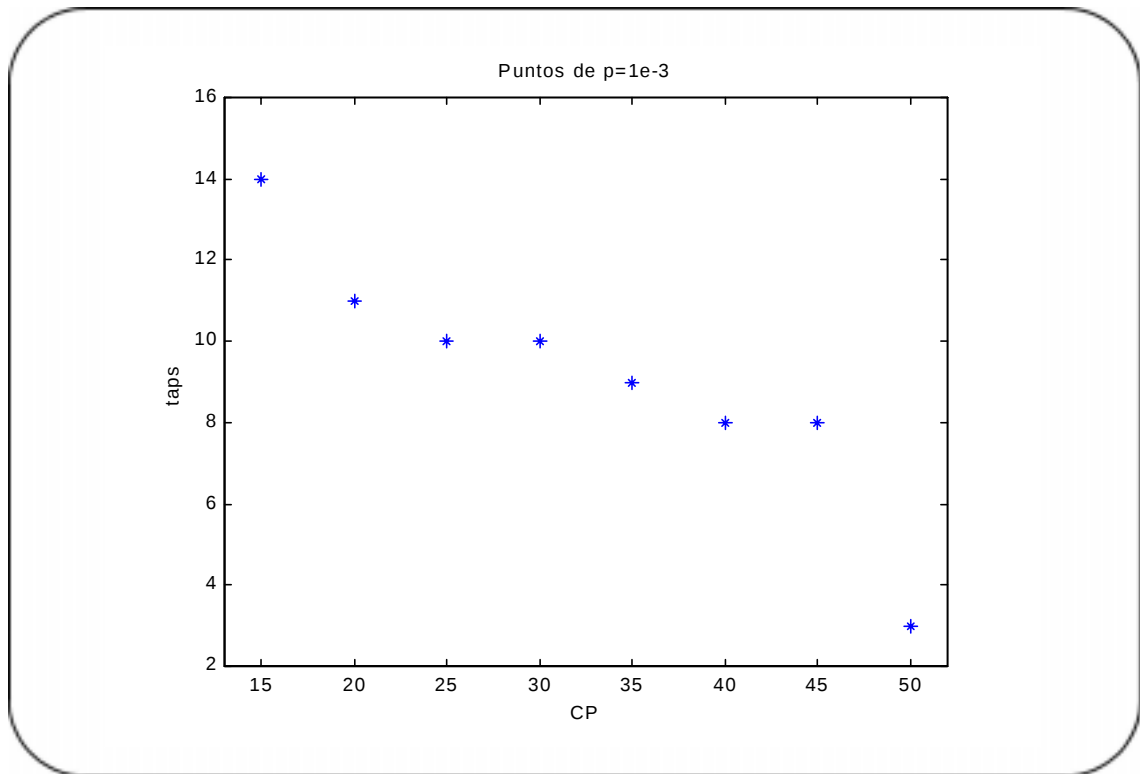


Figura 7.24 : Relación entre el número de taps y el CP para $p_e=10^{-3}$

Por último, y para terminar con el análisis de los resultados obtenidos con el FEQ, vamos a proporcionar una explicación al hecho constatado de que, para una longitud suficientemente larga, el FEQ es capaz de hacer que el sistema funcione por debajo de la probabilidad de error nominal. Este hecho no deja de ser curioso, ya que el FEQ se diseña para eliminar la componente de interferencia y no la de ruido.

Podemos encontrar una posible explicación si nos fijamos con atención en la forma en la que se lleva a cabo la inicialización del FEQ. Para determinar cada uno de los igualadores (recordemos que se diseña uno para cada subcanal) se resuelve un problema de mínimos cuadrados (véase el desarrollo teórico efectuado en el capítulo 5). Al hacerlo, se intenta maximizar

la SNR en cada subcanal. La componente de interferencia puede calcularse a partir de la cabeza y la cola de la respuesta impulsiva (los trozos de la respuesta que no son abarcados por el CP), y la de ruido puede calcularse a partir de la matriz de correlación. Ya que introducimos en el problema tanto interferencia (usamos un CP de longitud insuficiente) como ruido, el algoritmo va a maximizar la SNR en cada subcanal teniendo en cuenta ambos efectos. En principio, el efecto de la interferencia es dominante sobre el del ruido, puesto que no se está compensando con el CP adecuado. Por tanto, va a ser el primer efecto que elimine el FEQ. Sin embargo, si utilizamos más taps de los necesarios para eliminar la componente de interferencia, el igualador va a comenzar a realizar filtrado del ruido, siguiendo su criterio de maximizar la SNR. Esta eliminación del ruido es perfectamente posible, ya que se conoce su matriz de correlación (sería un filtrado de ruido igual que el realizado en un filtro de Wiener).

Por todo lo anterior, podemos ver que el FEQ es un igualador que funciona bastante bien, ya que consigue que el sistema DMT funcione incluso por debajo de la probabilidad de error nominal. Es cierto que su inicialización es más costosa, pero una vez que se calculan los coeficientes para todos los subcanales, requiere aproximadamente la misma complejidad computacional que el TEQ para realizar el filtrado.

7.9 – SIMULACIÓN DE UN SISTEMA CON PROBABILIDAD DE ERROR REAL

En este apartado vamos a simular un sistema con la probabilidad de error de símbolo definida en los estándares de VDSL. En concreto, vamos a simular un sistema DMT funcionando con un CP de 30 y con un igualador de 10 taps.

La probabilidad de error de símbolo definida en los estándares es de 10^{-7} . Si utilizásemos el método de Monte Carlo para determinar la probabilidad

de error, deberíamos hacer una simulación muy grande (si queremos obtener una varianza aceptable). Por eso, nosotros utilizaremos ahora el método de extrapolación de cola, estudiado en el capítulo 6.

Como primer paso para implementar una simulación usando este método, debemos definir los valores de los pseudoumbrales. Nuestro detector trabaja con constelaciones QAM normalizadas, de forma que la distancia entre dos puntos es igual a la unidad. Se produce entonces un error cuando el desplazamiento de la componente en fase o en cuadratura (o ambos) es mayor que 0.5, ya que entonces el valor recibido invadiría una región de decisión vecina. Por lo tanto, el verdadero umbral será 0.5, lo cual se corresponde con -3dB.

Según el criterio del método, vamos a escoger 3 pseudoumbrales situados 1dB, 2dB y 3dB por debajo del verdadero umbral (es decir, en -4, -5 y -6 dB). Tenemos que realizar una simulación de Monte Carlo en cada uno de estos puntos, y los valores de probabilidad de error que obtengamos nos deben servir para extrapolar el valor de la probabilidad de error en el verdadero umbral. En la tabla 7.11 se muestran los resultados del proceso.

Con los valores obtenidos en los pseudoumbrales, se hace una estimación (mediante regresión lineal) del valor de la probabilidad de error. En la figura 7.27 se muestra una gráfica con la regresión realizada. En el verdadero umbral, obtenemos el punto 2.735. Si deshacemos el doble logaritmo, llegamos a una probabilidad de error igual a 2.3671×10^{-7} .

	Punto a estimar	1er punto de la regresión	2º punto de la regresión	3er punto de la regresión
Umbral (dB)	-3	-4	-5	-6
p_e	(desconocida)	1.12e-5	1.58e-4	1.59e-3
$\ln(-\ln p_e)$	(desconocida)	2.433	2.169	1.863

Tabla 7.10 : Resultado de las medidas en los pseudoumbrales

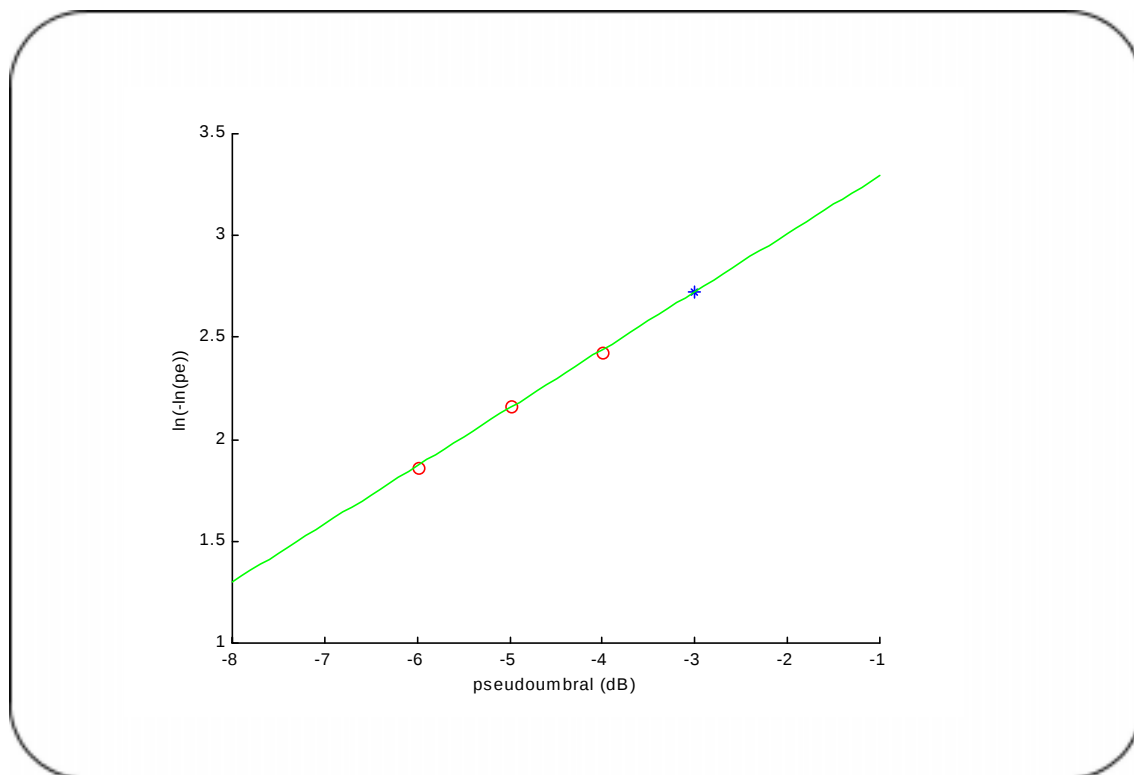


Figura 7.25 : Regresión realizada para estimar la probabilidad de error

Con los valores obtenidos en los pseudoumbrales, se hace una estimación (mediante regresión lineal) del valor de la probabilidad de error. En la figura 7.25 se muestra una gráfica con la regresión realizada. En el verdadero umbral, obtenemos el punto 2.735. Si deshacemos el doble logaritmo, llegamos a una probabilidad de error igual a 2.3671×10^{-7} .

La probabilidad de error nominal del sistema era de 10^{-7} , por lo que parece que el sistema no está funcionando de forma correcta. Pero en realidad lo que está ocurriendo es que la estimación que se ha realizado posee un sesgo

que es inherente al método de extrapolación de cola, como se mostró en la figura 6.2.

Por tanto, si corregimos el valor obtenido por el sesgo (aproximadamente 2), podemos ver que el sistema se comporta adecuadamente con esta probabilidad de error.

8.1 – CONCLUSIONES

En este proyecto se han analizado con detalle dos esquemas de igualación para sistemas DMT aplicados a VDSL, uno que opera en el dominio del tiempo (TEQ) y otro en el de la frecuencia (FEQ).

La forma de llevar a cabo el proyecto ha sido la siguiente:

En primer lugar, se analizó el problema de conseguir que un usuario de una red de telefonía fija sea capaz de utilizar la red para transmitir o recibir datos a gran velocidad (servicios de banda ancha). Vimos que en la red de un operador el cuello de botella (en cuanto a capacidad de transmisión) es el bucle de abonado. Así pues, para aumentar la velocidad de acceso a la red de un usuario, debemos centrar los esfuerzos en el desarrollo de técnicas que permitan un aprovechamiento óptimo de la capacidad de transmisión del canal. Se estudiaron brevemente las diversas soluciones que históricamente han ido resolviendo este problema y nos centramos especialmente en el servicio VDSL, que a fecha de hoy, es el que más capacidad ofrece.

El problema de la optimización de la capacidad de transmisión del canal no es abordable si no partimos de un estudio del medio de transmisión. Está claro que para lograr la optimización debemos desarrollar un método que se adapte a las características particulares del canal. En nuestro caso, tras el estudio del bucle de abonado que se ha realizado, queda claro que el uso de una técnica de modulación multiportadora junto con un algoritmo de carga adaptativa es la elección más acertada.

Dentro del estudio de las técnicas de modulación multiportadora, se ha estudiado el sistema DMT, que es el que se utiliza en los estándares de VDSL. Se ha visto cómo la utilización de este sistema en un canal con una respuesta impulsiva larga nos obliga a trabajar con una tasa de datos inferior a la teórica, es decir, con una disminución del rendimiento del sistema debido a la inserción del prefijo cíclico; esta es la forma más simple de evitar que los fenómenos de interferencia degraden la señal.

Las técnicas de igualación aparecen como una forma eficaz de aumentar el rendimiento, a la par que mantienen controlados los fenómenos de interferencia. El igualador DMT, entonces, es una herramienta muy útil a la hora de optimizar la capacidad de transmisión del bucle de abonado. En el desarrollo teórico de los dos procesos de igualación, hemos visto cómo se resuelve el problema.

Para la obtención de conclusiones, se han implementado en MATLAB un transmisor y un receptor DMT completos, en los que se han realizado todas las etapas de procesamiento de señal que el sistema requiere. Aunque VDSL es una tecnología que no ha madurado del todo, hemos procurado ajustarnos a los estándares publicados hasta la fecha. Sobre el sistema DMT implementado se han realizado multitud de simulaciones que ponen de manifiesto la necesidad de igualación si queremos maximizar la tasa de datos.

Para estudiar el comportamiento de los igualadores, se ha programado el código necesario en MATLAB para inicializarlos; y se han incluido en el sistema DMT. Los resultados obtenidos en los ensayos realizados varían de un igualador a otro.

En las pruebas realizadas con el TEQ hemos visto cómo la igualación se hace de una forma sencilla, y cómo el igualador es capaz de acortar la respuesta impulsiva del canal sin dificultad. Sin embargo, al insertarlos en el sistema los resultados no eran los esperados. La probabilidad de error detectada en los sistemas con igualación TEQ era siempre ligeramente superior a la nominal. Analizando con más detalle la situación, se vio que la

clave estaba en la respuesta en frecuencia del igualador. Mientras que en el dominio del tiempo se comportaba perfectamente, en el dominio de la frecuencia presentaba zonas de pendiente muy abrupta alrededor de sus ceros. En el análisis teórico que se hizo de la modulación multiportadora se vio cómo estas zonas abruptas podían dar problemas. Se comprobó que, en efecto, el aumento de la probabilidad de error detectada se debía a la acumulación de errores en los subcanales situados en dichas zonas.

Por suerte, los ensayos realizados con el FEQ fueron totalmente satisfactorios, ya que en algunos casos se detectaron probabilidades de error incluso por debajo del valor nominal.

Ahora es el momento de extraer conclusiones acerca de la configuración óptima de un sistema DMT para VDSL. En primer lugar, tenemos que decidir si utilizar o no algún mecanismo de igualación. Es evidente que la inclusión de un igualador en el sistema conlleva un aumento de la complejidad computacional. Sin embargo, si el objetivo último que perseguimos es lograr una optimización de la capacidad de transmisión del canal, siempre va a ser conveniente un mecanismo de igualación.

Pero la decisión de qué tipo de igualador debemos utilizar es más delicada. El igualador TEQ es más sencillo de inicializar, pero su uso acarrea los problemas anteriormente descritos. Pero el FEQ, aunque funciona de forma satisfactoria, eleva enormemente el coste computacional, debido al complejo proceso de inicialización. Teniendo en cuenta entonces el resultado de las simulaciones, llegamos a la conclusión de que el igualador que mejor se adapta a este caso es el FEQ, ya que es el único que nos asegura que la probabilidad de error no supere la nominal.

8.2 – LÍNEAS FUTURAS DE INVESTIGACIÓN

Una posible línea de investigación puede ser la de investigar con más profundidad los problemas del MMSE-TEQ, intentando descubrir cómo evitar el efecto de los ceros.

Otra posibilidad sería implementar el TEQ geométrico (comentado en el apartado 7.7).

También se podría realizar un estudio del proceso de inicialización del FEQ e intentar simplificarlo, ya que es el principal inconveniente de este igualador.

Otra posibilidad sería investigar acerca del efecto que pudiera tener la elección de un retraso del igualador distinto de cero.

GLOSARIO

ADSL Asymmetric Digital Subscriber Line
AM Amplitude Modulation
AMI Alternate Mark Inversion
AN Access Node
ANSI American National Standards Institute
ATM Asynchronous Transfer Mode
CIR Channel Impulse Response
CP Cyclic Prefix
DFT Discrete Fourier Transform
DMT Discrete Multitone Transmission
DSL Digital Subscriber Line
EL-FEXT Equal Level Far End Crosstalk
FDM Frequency Division Multiplex
FEQ Frequency Domain Equalizer
FEXT Far End Crosstalk
FFT Fast Fourier Transform
FIR Finite duration Impulse Response
FTTCab Fiber To The Cabinet
FTTEx Fiber To The Exchange

FTTH Fiber To The Home
G-TEQ Geometric TEQ
HDSL High data rate DSL
IBI Inter-Block Interference
ICI Inter-Carrier Interference
IDFT Inverse Discrete Fourier Transform
IFFT Inverse Fast Fourier Transform
IP Internet Protocol
ISI Inter-Symbol Interference
ITU-T International Telecommunication Union - Telecommunications
LAN Local Area Network
MCM Multi-Carrier Modulation
MMSE-TEQ Minimum Mean Square Error TEQ
NT Network Termination
NEXT Near End Crosstalk
OLT Optical Line Termination
ONU Optical Network Unit
PDF Probability Density Function
POTS Plain Old Telephone Service
PSD Power Spectrum Density

QAM Quadrature Amplitude Modulation
RDSI Red Digital de Servicios Integrados
RFI Radio Frequency Interference
SDSL Single Line DSL
SFG Signal Flow Graph
SNR Signal to Noise Ratio
SQAM Staggered QAM
STM Synchronous Transfer Mode
TEQ Time Domain Equalizer
TIR Target Impulse Response
UEC Unit Energy Constraint
UTP Unshielded Twisted Pair
VCR Video Cassette Recorder
VDSL Very High data rate DSL
VTU VDSL Transmission Unit
VTU-O VTU ONU
VTU-R VTU Remote
WAN Wide Area Network
XT Crosstalk

REFERENCIAS

- [1] Naofal Al-Nadir and J. M. Cioffi, "Optimum Finite-Length Equalization for Multicarrier Transceivers, *IEEE Transactions of Communications*, Vol 44, No. 1, January 1996, pp 56-64.
- [2] Katleen Van Acker, Geert Leus, Marc Moonen, Olivier van de Wiel and Thierry Pollet, "Per Tone Equalization for DMT-Based Systems", *IEEE Transactions of Communications*, Vol 49, No. 1, January 2001, pp. 109-119.
- [3] Thierry Pollet, Miguel Peeters, Marc Moonen and Luc Vandendorpe, "Equalization for DMT-Based Broadband Modems", *IEEE Communications Magazine*, May 2000, pp. 106-113.
- [4] Vladimir Oksman and Jean-Jacques Werner, "Single-Carrier Modulation Tecnology for Very High-Speed Digital Subscriber Line", *IEEE Communications Magazine*, May 2000, pp. 82-89.
- [5] John A. C. Bingham, "Multicarrier Modulation for Data Transmission: An Idea Whose Time Has Come", *IEEE Communications Magazine*, May 1990, pp. 5-14.
- [6] Thierry Pollet and Miguel Peeters, "Synchronization with DMT Modulation", *IEEE Communications Magazine*, April 1999, pp. 80-86.
- [7] Daniel Bengsson and Daniel Landström, *Coding in a Discrete Multitone Modulation System*, Master's Thesis, Luleå University of Technology, 1996.
- [8] J. M. Cioffi, "EE 479 Course Notes". Stanford University, 1998
- [9] VDSL Alliance, "DMT VDSL Draft Standard Proposal", www.VDSLAlliance.com.
- [10] ADSL Forum, "General Introduction to Copper Access Technologies", 1998.
- [11] ETSI TS 101 270-2 V1.1.1 (2001-02) Technical Specification, *Transmission and Multiplexing (TM); Access transmission systems on metallic access cables; Very high speed Digital Subscriber Line (VDSL); Part 2: Transceiver specification*
- [12] Study Group 15, *Physical Medium Specific Specification for G.VDSL*, ITU – T Standarization Sector, Temporary document, May 1999 (documento provisional)
- [13] Clara Valle Carreras, *DMT: El Código de Línea Estándar para ADSL*, PFC Escuela Superior de Ingenieros (Telecomunicación), Junio 1999.
- [14] John A. C. Bingham, *ADSL, VDSL, and Multicarrier Modulation*. Palo Alto, California. John Wiley & Sons, Inc., 2000.
- [15] John G. Proakis, *Digital Communications*, Third Edition, Northeastern University, McGraw-Hill, Inc., 1995.

[16] Alan V. Oppenheim, Ronald W. Schaffer,
Discrete-Time Signal Processing, Prentice-Hall
International, Inc., 1989.

[17] M. C. Jeruchim, P. Balaban, K. Shanmugan,
Simulation of Communication Systems, Plenum
Press, New York, 1992.

SCRIPT 1: SIMULACIÓN DE UNA TRANSMISIÓN DMT

```

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%
%SCRIPT 1: dsl.m
%
% FUNCIONES QUE UTILIZA:
%     aprend.m
%     asigna.m
%     ser_par.m
%     codif.m
%     modul.m
%     canal.m
%     rx.m
%     represen.m
%
% ES UTILIZADO POR:
%     ---
%
% VARIABLES "GLOBALES":
%     Pmax
%     fs
%     pe
%     NSUB
%     CP
%     margen
%     ganancia
%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%

% Función que simula una transmisión de datos utilizando la tecnología
% xDSL y el sistema DMT
%

warning off;          % Evita los WARNING MSGS.

% INICIO VARIABLES GLOBALES:
Pmax=11.5;            % Potencia máxima de transmisión, dBm
fs=30e6;              % Frecuencia de muestreo, en Hz
pe=1e-7;              % Probabilidad de error de símbolo deseada
CP=60;               % CYCLIC_PREFIX utilizado en la simulación
NSUB=256;            % Número de subcanales
margen=6;             % Margen deseado, en dB
ganancia=0;          % Ganancia del código, en dB
% FIN VARIABLES GLOBALES

% VARIABLES AUXILIARES
FP=(2*NSUB)/(2*NSUB+CP);          % Factor de pérdidas por CP
tasa_simbolo=fs/(2*NSUB);         % Tasa de símbolos (baudios)
% nota: se han incluido los efectos del CP en cuanto a tiempo de tx
PSD=10^(0.1*Pmax)/(fs*1e3)*2;     % PSD W/Hz tx.
BW=fs/2;                          % BW utilizado en la tx
espaciado=BW/NSUB;               % Espaciado entre subcanales
Edis=10^(Pmax/10)/(1e3*tasa_simbolo)*FP; % Ener/símbolo disp en el tx (J)
frecuencias=[0:espaciado:BW-espaciado]; % Frec. de "portadora"
% FIN VARIABLES AUXILIARES
UM=0;
while UM~=0 & UM~=1
    UM=input('¿Usar mascara? (0=NO, 1=SI):\n');
end

% COMIENZO DE LA SIMULACIÓN
% FASE DE APRENDIZAJE
% Obtención de los parámetros del canal.
[Hnum, Hden, H, h, HxtNEXTdb, HxtFEXTdb, w]=aprendizaje(fs, NSUB);

% Varias gráficas de la situación del canal
hasta=NSUB;
Hdb=20*log10(abs(H(1:NSUB)))';

```

```

figure(1);figure.figure(1)); clf;
plot(frecuencias(1:hasta),Hdb(1:hasta)+10*log10(PSD*1e3));
title('PSD de señal(azul), de NEXT y de FEXT recibidas (dBm/Hz)');
xlabel('Frecuencia (Hz)');
grid;
hold on;
plot(frecuencias(1:hasta),HxtNEXTdb(1:hasta)+10*log10(PSD*1e3),'r');
plot(frecuencias(1:hasta),HxtFEXTdb(1:hasta)+Hdb+10*log10(PSD*1e3),'g');
hold off;
senal=Hdb(1:NSUB)+10*log10(PSD*1e3);

save senal senal;
pause;

% Paso de la ganancia a naturales
% Paso del crosstalk a naturales
% Paso del crosstalk a energía
Hnat=abs(H(1:NSUB))';
xtNEXTdb=HxtNEXTdb+10*log10(PSD*1e3);
xtFEXTdb=HxtFEXTdb+Hdb+10*log10(PSD*1e3);
xtNEXT=10.^(0.1.*xtNEXTdb).*1e-3;           % WATIOS/Hz
xtFEXT=10.^(0.1.*xtFEXTdb).*1e-3;           % WATIOS/Hz
Potxt=xtNEXT+xtFEXT;                         % WATIOS/Hz   Pot total de ruido
Enxt=Potxt/tasa_simbolo;

% FASE DE ASIGNACIÓN
% Ahora llamamos a la función de asignación, que utilizando el algoritmo
% óptimo discreto nos devuelve el vector de asignación de bits junto con
% el vector de energías asociadas a cada subcanal.
MSG=sprintf('>>>>EJECUTANDO ALGORITMO DE CARGA<<<<');
disp(MSG);
[b pn]=asignacion(Hnat, Potxt, Edis, pe, NSUB, tasa_simbolo, ganancia, margen, UM);
en=pn;
rendi=2*NSUB/(CP+2*NSUB);
tasa=sum(b)*tasa_simbolo*rendi;
tasaMBPS=tasa/1e6;

%
% Tasa binaria neta que conseguimos en la TX (en Mbps)
%
MSG1=sprintf('Tasa binaria (Mb/s): %d',tasaMBPS);
disp(MSG1);
MSG=sprintf('El rendimiento del enlace es: %d',rendi);
disp(MSG);
MENSAJE=sprintf('Un simbolo DMT se compone de %d bits',sum(b));
MENSAJE2=sprintf('SUBCANALES USADOS: %d',nnz(b));
disp(MENSAJE);
disp(MENSAJE2);
MSG=sprintf('>>>>FIN DEL ALGORITMO DE CARGA<<<<');
disp(MSG);
pause;

% GENERACIÓN DE SÍMBOLOS DMT
% Sea la siguiente cadena de bits a transmitir
genera=input('¿Cuántos símbolos DMT quiere generar?: \n');

MSG=sprintf('>>>>GENERACIÓN ALEATORIA DE SEÑALES Y RUIDO.TX Y RX<<<<');
disp(MSG);
datos=(rand(sum(b),genera));
s=round(datos); %Este es el tren de bits que tenemos que transmitir

M=ser_par(s,b, NSUB);

% CODIFICADOR QAM
% Ahora tenemos que pasar esta matriz de símbolos al cod. QAM
[simbolos distancia]=codificador(M,b,en,NSUB);

% MODULADOR IFFT
% Ahora enviamos los sub-símbolos DMT al modulador IFFT
simbolo_DMT=modulador(simbolos, NSUB);

```

```

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%
%%INICIO TX
%%
%% PROCESO DE TRANSMISIÓN POR UN CANAL DIGITAL

[salida2]=canal(simbolo_DMT,h,Potxt,0,0,0,CP, NSUB);

% ELIMINAMOS EL CYCLIC_PREFIX
container=0.*simbolo_DMT;
for k=1:size(simbolo_DMT,2)
    delta=(k-1)*(size(simbolo_DMT,1)+CP);
    sal=salida2(delta+1+CP:delta+NSUB*2+CP);
    container(:,k)=sal.';
end

salida1=container;
%%
%%FIN TX
%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%

% GUARDAMOS PARÁMETROS EN DISCO: param2.mat
save param2 simbolo_DMT simbolos distancia genera en b H CP NSUB salida1 salida2 tasa
Edis;

% PROCESADO EN EL RECEPTOR DE LA CADENA DE MUESTRAS RECIBIDAS
[simbolo_rx_norm simbolo_rx rx_cfase rx_ccuad]=rx(salida1,H,distancia,genera,b, NSUB);
% Ahora, con la información de estas dos componentes, tenemos que
% detectar el símbolo que hemos recibido.

% Como sabemos el número de bits que está transportando el símbolo,
% sabemos que la mitad mayor (si la hubiera) corresponde a la
% componente en fase.

% REPRESENTACIÓN GRÁFICA DE CONSTELACIONES Y CÁLCULOS FINALES
NO=0;
y=representa(simbolo_rx,simbolos,distancia,en,b,NSUB,NO);

```

SCRIPT 2: INICIALIZACIÓN DEL MMSE-TEQ

```

% Script para calcular el TEQ (TIME-DOMAIN EQUALIZER)

% Se cargan de ficheros diferentes parámetros de
% transmisiones previas.
load param;
load param2;
CYCLIC_PREFIX=CP;
sec=sali;
clear output CP;
H1=H;
Potxt=xt;
% Contiene la variable Potxt (PSD del ruido)

% Cálculo de la CIR (Channel Impulse Response)
cir=h;
% CIR normalizada
cir_n=cir./norm(cir,2);
bm=zeros(1,length(b));
bf = find(b>0);
bm(bf)=ones(1,length(bf));

% Cálculo de la autocorrelación del ruido
auxi=Potxt(2:length(Potxt));
psd=[Potxt 0 flipr(auxi)];
auto=corre;

```



```

title('CIR (azul) y TIR (verde)');
subplot(2,1,2); hold on;
stem([zeros(1,delay) bopt.'],'g');
stem(total(1:60)./norm(total,2),'r');

plot(zeros(1,100),'r');
title('TIR (verde) y respuesta conjunta (rojo)');
pause; hold off;

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%
% Simulación con el TEQ calculado
%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%
%Actualizamos el CYCLIC PREFIX
CYCLIC_PREFIX2=Nb;

% Enviamos la señal al canal, con el nuevo CP
[output ousen ourui]=canal(simbolo_DMT,h,xt,0,0,0,CYCLIC_PREFIX2,NSUB);
corre=xcorr(output(1001:length(output)),MUL*NSUB-1,'biased');

% Diversas representaciones gráficas

% PSD del ruido_generado y del ruido_deseado
senal_db=10*log10(abs(fft(corre)).*1e3);
senal_db=senal_db(1:MUL:length(senal_db));
%ruido_db=(abs(fft(xcorr(ruido_final,'biased'))));
%figure(1);
figure(2);clf;
plot(senal_db(1:NSUB),'r');
hold on;
%plot(10*log10(abs(xt)));
%plot(10*log10((xt.*1e3)));
title('señal generada (rojo)');
ylabel('PSD (dBm/Hz)');
hold off;
pause;
load cosa;
load senal, load ruido;
plot(senal_db(1:NSUB)-cosa),hold on,plot(senal_db(1:NSUB),'r'),
plot(cosa,'g'),plot((senal-ruido_teo).*bm,'c'),hold off,
title('SNR (azul), señal (rojo) y ruido (verde)')
pause;

% Datos que llegan al receptor
%UTILIZACIÓN DEL IGUALADOR
salida_ig=fftfilt(wopt(delay+1:Nf),output);
salida_s=fftfilt(wopt(delay+1:Nf),ousen);
salida_r=fftfilt(wopt(delay+1:Nf),ourui);

save ourui ourui;

%salida_ig1=fftfilt(wopt.',output);
%salida_ig=[salida_ig1(1+delay:length(salida_ig1)) zeros(1,delay)];

%load param;
corre_s=xcorr(salida_s(1001:length(output)),MUL*NSUB-1,'biased');
corre_r=xcorr(salida_r(1001:length(output)),MUL*NSUB-1,'biased');

% Diversas representaciones gráficas

% PSD del ruido_generado y del ruido_deseado
senals_db=10*log10(abs(fft(corre_s)).*1e3);
ruidos_db=10*log10(abs(fft(corre_r)).*1e3);
senals_db=senals_db(1:MUL:length(senals_db));
ruidos_db=ruidos_db(1:MUL:length(ruidos_db));
%ruido_db=(abs(fft(xcorr(ruido_final,'biased'))));
%figure(1);
figure(6);clf;
plot(senals_db(1:NSUB),'r');
hold on;
%plot(10*log10(abs(xt)));
%plot(10*log10((xt.*1e3)));
title('señal generada (rojo)');
ylabel('PSD (dBm/Hz)');

```

```
hold off;
load cosa;
load senal, load ruido;
plot(senals_db(1:NSUB)-ruidos_db(1:NSUB),hold on,
plot(senals_db(1:NSUB),'r'),plot(ruidos_db(1:NSUB),'g'),
plot((senal-ruido_teo).*bm,'c'),hold off,
title('SNR (azul), señal (rojo) y ruido (verde)')
pause;

% Eliminamos el CP
container=zeros(NSUB*2,genera);
for k=1:genera
    desp=(k-1)*(NSUB*2+CYCLIC_PREFIX2);
    sal=salida_ig(desp+CYCLIC_PREFIX2+1:desp+CYCLIC_PREFIX2+NSUB*2);
    container(:,k)=sal';
end

sal=container;

% Cálculo de la respuesta en frecuencia total (canal+TEQ)
tot=(H1).*(H2);
figure(5); clf;
subplot(3,1,1);
plot(10.*log10(abs(H1(1:NSUB))));
ylabel('CANAL');
subplot(3,1,2);
plot(10.*log10(abs(H2(1:NSUB))));
ylabel('IGUALADOR');
subplot(3,1,3);
plot(10.*log10(abs(tot(1:NSUB))));
ylabel('TOTAL');
pause;

% Codigo para calcular la derivada de la respuesta en frecuencia
re=10*log10(abs(tot(1:NSUB)));
re2=zeros(size(re)).';
re2(1)=0;%abs(re(1)-re(2));
re2(NSUB)=abs(re(NSUB)-re(NSUB-1));
for t=2:NSUB-1
    re2(t)=mean([abs(re(t)-re(t-1)) abs(re(t)-re(t+1))]);
    if re2(t)==Inf
        re2(t)=10;
    end
end

mask=zeros(1,NSUB);
for t=1:NSUB
    if re2(t)>0.1
        mask(t)=1;
    end
end
save mask mask;

% Enviamos la señal igualada y sin CP al detector
[simbolo_rx_norm2 simbolo_rx rx_cfase rx_ccuad]=rx(sal,tot,distancia,genera,b,NSUB);

% Función para la representación grafica de constelaciones
% y calculo de probabilidades de error.
y=representa(simbolo_rx,simbolos,distancia,en,b,NSUB,0);

% Salvamos a disco las variables del TEQ
save teq Nf Nb tot wopt bopt mask delay;

% Información de utilidad sobre el TEQ construido
ceros=angle(roots(wopt));
subcanales=(ceros/pi*NSUB)+1;
%MSG2=sprintf('Subcanales afectados por ceros: %d',ceros);
rendi2=2*NSUB/(2*NSUB+CYCLIC_PREFIX2);
tasa_n=tasa*rendi2;
```

```
MSG=sprintf('Tasa del sistema: %d', tasa_n);
disp(MSG);
disp('Subcanales afectados por ceros y magnitud del cero correspondiente');
[mod(subcanales.*(sign(subcanales)+1)/2, NSUB) abs(roots(wopt))]
```

SCRIPT 2: INICIALIZACIÓN DEL FEQ

```
% SCRIPT QUE IMPLEMENTA EL IGUALADOR EN EL DOMINIO DE LA FRECUENCIA
MSG=sprintf('Iniciando la fase preliminar');
disp(MSG);
token1=clock;
```

```
% Cargamos parámetros de transmisiones anteriores
load param;
load param2;
CYCLIC_PREFIX=CP;
```

```
h=h(1:59);
```

```
% PARÁMETROS DE DISEÑO
```

```
N=NSUB*2; % Tamaño de las FFTs
nu=CP; % CP de la CIR
T=7; % Número de taps del igualador
delta=0; % Retardo de sincronización
L=length(h)-1; % Longitud no causal de la CIR
K=0; % Longitud causal de la CIR
s=N+nu;
```

```
% Cálculo de las matrices previas
fila=[fliplr(h) zeros(1,N+T-2)];
columna=[h(length(h)) zeros(1,N+T-2)];
```

```
H11=toeplitz(columna, fila);
O1=zeros(N+T-1, N+nu-T+1-L+nu+delta);
O2=zeros(N+T-1, N+nu-K-delta);
H1=[O1 H11 O2];
```

```
O=zeros(nu, N-nu);
I1=eye(nu);
I2=eye(N);
P1=[O I1; I2];
O3=zeros(size(P1));
PPP=[P1 O3 O3; O3 P1 O3; O3 O3 P1];
```

```
I3=mafft(N);
O4=zeros(size(I3));
III=[I3 O4 O4; O4 I3 O4; O4 O4 I3];
```

```
H=H1*PPP*III;
```

```
% Cálculo de la autocorrelación de la señal
% almacenada en la variable SIMBOLOS [N/2 x 1 x #simbolos]
```

```
estim=200; %Número de símbolos para realizar la estimación
```

```
simbolo_grande=zeros(NSUB*2, size(simbolos, 2));
simbolo_grande(1,:)=real(simbolos(1,:));
simbolo_grande(2:NSUB,:)=simbolos(2:NSUB,:);
simbolo_grande(NSUB+1,:)=imag(simbolos(1,:));
simbolo_grande(NSUB+2:NSUB*2,:)=conj(simbolos(NSUB:-1:2,:));
```

```
coger=N+T-1;
coger_total=estim*coger;
```

```
noise=reshape(ruido_final(1:coger_total),coger,estim);
signal=zeros(3*N,estim);

for n=1:estim
    signal(:,n)=reshape(simbolo_grande(:,(3*n-2):3*n),3*N,1);
end

Rx=signal*signal'./estim;
Rn=noise*noise'./estim;

tiempo=etime(clock,token1)/60;
MSG=sprintf('Finalizada fase preliminar. Tiempo (min): %g', tiempo);
disp(MSG);
% Cálculo del ecualizador MMSE-FEQ
V=calculafeq(Rx,Rn,H,N,T); % Se resuelve el problema de min. cuadrados

% Salvamos en un archivo el igualador y su retraso
save feq V delta;
tiempo2=etime(clock, token1)/60;
MSG=sprintf('\nTiempo total invertido (min): %g',tiempo2);
disp(MSG);
```