

TRABAJO FIN DE GRADO

Título: *Estrategias de asignación de recursos para el DL en LTE.*

Autor: *Víctor Cornejo García.*

Titulación: *Grado en ingeniería de sistemas de comunicaciones.*

Profesor: *Ángel María Bravo Santos.*

Fecha: *18-Junio-2013*



ÍNDICE:

1. Objetivo	3
2. Introducción	4
2.1. Sistemas 1G	4
2.2. Sistemas 2G 2.5G	5
2.3. Sistemas 3G 3.5G	6
2.4. Sistemas 4G	9
3. Sistema LTE	10
3.1. Interfaz Radio LTE	11
3.2. El canal LTE, MIMO	15
3.3. El enlace descendente y ascendente	17
4. Asignación de recursos en LTE	19
4.1. Fundamentos de la asignación de recursos en LTE	19
4.2. Asignación de recursos en el UL	23
4.3. Asignación de recursos en el DL	25
5. Estrategias de asignación de recursos en el DL	28
5.1. Algoritmos de asignación de recursos	32
5.2. Estudio comparativo de asignación de recursos	38
5.2.1. Estudio con máxima capacidad	43
5.2.1.1. Asignación a varios usuarios	43
5.2.1.2. Asignación a un único usuario	44
5.2.1.3. Usuarios satisfechos	45
5.2.2. Estudio con capacidad media	46
5.2.2.1. Asignación a varios usuarios	47

5.2.2.2.	<u>Asignación a un único usuario.....</u>	<u>48</u>
5.2.2.3.	<u>Usuarios satisfechos.....</u>	<u>49</u>
5.3.	<u>Propuesta de mejoras.....</u>	<u>50</u>
5.3.1.	<u>Asignación a varios usuarios.....</u>	<u>51</u>
5.3.2.	<u>Asignación a un único usuario.....</u>	<u>52</u>
5.3.3.	<u>Usuarios satisfechos.....</u>	<u>52</u>
6.	<u>Conclusiones.....</u>	<u>54</u>
7.	<u>Acrónimos.....</u>	<u>55</u>
8.	<u>Referencias.....</u>	<u>58</u>
9.	<u>Anexo A: Código utilizado en las simulaciones.....</u>	<u>61</u>
9.1.	<u>Main.....</u>	<u>61</u>
9.2.	<u>Proportional Fair.....</u>	<u>65</u>
9.3.	<u>Maximum C/I.....</u>	<u>66</u>
9.4.	<u>Función tasa instantánea.....</u>	<u>67</u>

1. Objetivo

El principal objetivo de este trabajo fin de grado es el estudio de las distintas estrategias utilizadas en la asignación de recursos en el enlace descendente del sistema Long Term Evolution (LTE).

Se pretende explicar tanto las particularidades de la asignación de recursos del sistema LTE, como algunos de los algoritmos que se emplean. Para ello se ha realizado una búsqueda en la literatura disponible y un posterior análisis y estudio de la información encontrada.

En la última parte del trabajo, se realizan simulaciones de algunos de estos algoritmos estudiados, con el fin de mostrar sus principales características y al mismo tiempo poder apreciar las principales diferencias en los resultados de la asignación de recursos al utilizar unos u otros algoritmos.

Al mismo tiempo, se pretende hacer un breve recorrido por los diferentes sistemas de comunicaciones móviles, describir su evolución hasta llegar a LTE y profundizar en ciertos aspectos destacados de este último.

2. Introducción

Para llegar a entender la cuarta generación de sistemas móviles (4G), en la cual se centra este trabajo, es conveniente hacer un breve recorrido para explicar la evolución de los distintos sistemas desde la primera generación (1G), pasando por la segunda (2G) y tercera (3G), hasta acabar en la cuarta generación.

2.1 Sistemas 1G.

Estos primeros sistemas de telefonía móvil se caracterizaban, principalmente, por el hecho de ser analógicos y fueron los primeros sistemas en utilizar una estructura celular.

A pesar de ser dispositivos de gran tamaño y con poca independencia, ya que las baterías no tenían una duración muy prolongada y la tecnología del momento no era la más sofisticada, tuvieron una gran aceptación social y pronto comenzaron a extenderse.

El primer sistema fue desarrollado en Japón, por la Nippon Telegraph and Telephone (NTT) [1]. Posteriormente surge el Nordic Mobile Telephony (NMT) en los países nórdicos, (Suecia, Noruega, Finlandia y Dinamarca), empleando la banda de los 450 MHz, conforme se fue extendiendo se utilizó la banda de los 900 MHz para aumentar la capacidad. Después fueron surgiendo diversos sistemas como el Advanced Mobile Phone Service (AMPS) [1] en EE.UU., el Total Access Communications System (TACS) en Reino Unido o el C-Netz en Alemania.

Todos estos sistemas compartían una serie de aspectos técnicos:

- Utilizaban una modulación analógica en FM para la transmisión de voz.
- El espectro se dividía en canales y éstos se asignaban a las estaciones base, se asignaban canales distintos para evitar las interferencias.
- Se utilizaba duplexado en frecuencia, Frequency Division Duplex (FDD), para el enlace ascendente y descendente, utilizando frecuencias próximas a los 900MHz.

El principal inconveniente que presentaban estos sistemas era, que a pesar de compartir una serie de aspectos técnicos, no era posible la interoperabilidad entre estos sistemas, es decir, un móvil que operaba en un sistema concreto, solo funcionaba en ese sistema y no podía ser utilizado en otro distinto.

2.2 Sistemas 2G y 2.5G.

Seguidamente a los sistemas 1G, comenzaron a desarrollarse los denominados sistemas de segunda generación (2G), debido principalmente a los problemas que presentaban los sistemas 1G, como el de la interoperabilidad entre distintos sistemas, sobre todo debido a que los sistemas 1G comenzaron a ser insuficientes en cuanto a capacidad y prestaciones.

Se desarrollaron diferentes sistemas de segunda generación [1], como el IS-54 en EE.UU., el Japan Digital Cellular (JDC) en Japón y el sistema GSM en Europa, siendo este último el más importante y exitoso ya que se extendió a un gran número de países, por ello a continuación quiero hacer especial hincapié en este sistema, ya que está considerado como uno de los mayores logros de la ingeniería europea, en cuanto a telecomunicaciones se refiere.

El principal objetivo era lograr un sistema global que permitiera la interoperabilidad entre distintos países. Para ello, la Conferencia Europea de Administraciones de Correos y Telecomunicaciones, “Conference of European Postal and Telecommunications” (CEPT), creó el grupo de trabajo Groupe Speciale Mobile (GSM) [15] encargado de desarrollar dicho sistema. Posteriormente, el European Telecommunication Standards Institute (ETSI), continuó con el proyecto de GSM hasta lograr su comercialización el 1992.

Pasemos ahora a describir los importantes avances técnicos de los sistemas 2G [2] respecto a los sistemas 1G:

- La principal característica de los sistemas 2G, es el paso de analógico a digital, lo que conlleva grandes ventajas respecto a los sistemas 1G, como una mayor eficiencia en la utilización del espectro y mayor robustez frente a las interferencias del canal.

- Incorpora diversas técnicas, como la codificación de canal, codificación de fuente, la modulación digital o las técnicas de entrelazado entre otras. La introducción de estas técnicas permitió notables mejoras en la calidad de las comunicaciones.
- Otro avance importante es el que se produjo en cuanto a la electrónica empleada en los terminales, lo que permitió reducir el tamaño de los terminales de forma considerable, además de aumentar la vida de las baterías y reducir los costes en su fabricación, lo que contribuyó a que resultaran más cómodos y atractivos para los usuarios.
- Otra de las principales novedades del sistema GSM, es la incorporación de las técnicas: Frequency Division Multiple Access (FDMA) y Time Division Multiple Access (TDMA) [15]. El sistema GSM utiliza FDD para el enlace ascendente y descendente y además incorpora FDMA para dividir el ancho de banda de 25MHz en 124 portadoras de 200 kHz, del mismo modo utiliza TDMA para dividir cada portadora de 200 kHz en 8 time-slots, es decir, divide cada canal en 8 intervalos de tiempo. La tasa máxima de estos 8 canales es de 270,83 kbps en una trama de 4,61 ms.
- Este sistema trabaja inicialmente en la banda de los 900 MHz, aunque posteriormente surgieron algunas versiones en distintos países que trabajan en otras bandas como la de 1800 o 1900 MHz. (GSM-1800 y GSM-1900).

El gran éxito cosechado por el sistema GSM fue debido, fundamentalmente, a que permitió la interoperabilidad de terminales en distintos países, utilizando para ello terminales tri-banda.

Con el paso del tiempo, los sistemas 2G evolucionan a los llamados sistemas 2.5G [2], un paso intermedio en la evolución entre los sistemas 2G y 3G.

Los sistemas 2G aportaron notables mejoras respecto a los sistemas 1G, pero eran principalmente utilizados para el tráfico de voz. Los sistemas 2.5G tratan de aportar mejoras en la transmisión de datos para aumentar la capacidad.

Las evoluciones más destacadas del sistema GSM fueron, General Packet Radio Services (GPRS) y Enhanced Data Rates for Global Evolution (EDGE) [15]. Se emplea conmutación de paquetes, lo que permite una mayor eficiencia espectral y al mismo tiempo permite un cambio en la tarificación muy atractivo para los clientes, ya que permite tarificar por datos transmitidos y no por tiempo de conexión como se hacía anteriormente.

Una de las principales novedades de GPRS [2], es que introduce distintos esquemas de codificación: CS1: 9.05 kbps, CS2: 13.4 kbps, CS3: 15.6 kbps y CS4: 21.4 kbps. Teniendo en cuenta la calidad del enlace, el tráfico de la celda o el tipo de terminal y permite usar varios time slots en una misma conexión. Todo esto permite lograr una velocidad máxima teórica de 171 kbps, utilizando el esquema de codificación CS4 y los 8 time slots disponibles. En cuanto a los cambios realizados en la arquitectura para introducir GPRS, únicamente se añaden dos nodos nuevos, el Serving GPRS Support Node (SGSN) y el Gateway GPRS Support Node (GGSN), para el tráfico de paquetes y se introduce la Packet Control Unit (PCU) en las Base Station Controller (BSC), para permitir la asignación dinámica de canales a GSM o GPRS.

Posteriormente surge EDGE [3]; su principal novedad con respecto a GPRS y GSM, es la introducción de nuevas técnicas de codificación y modulación adaptativa (8PSK), lo que permite un notable aumento de las velocidades, hasta alcanzar una velocidad máxima teórica de 384 kbps empleando 8 slots y el esquema de codificación MCS9. La arquitectura de la red no experimenta grandes cambios entre GPRS [15] y EDGE, únicamente se introduce una unidad transceptora en la estación base (TRU) y se realizan las actualizaciones de software pertinentes.

Como se puede apreciar, estos nuevos sistemas van introduciendo diversos cambios en la red original, que poco a poco va evolucionando hacia lo que serán las redes de tercera generación (3G).

2.3 Sistemas 3G y 3.5G.

Los sistemas 3G, surgen de la evolución de los sistemas 2.5G. De nuevo, igual que ocurriera con la evolución a los sistemas 2G, surgieron nuevas demandas por parte de los usuarios, las cuales era necesario satisfacer. En particular, se produjo un notable aumento en la demanda del tráfico de datos, para la cual los sistemas 2G no estaban capacitados y, al mismo tiempo, fueron creciendo las expectativas sobre nuevas aplicaciones multimedia para los terminales.

La Unión Internacional de las Telecomunicaciones, “International Telecommunications Union” (ITU) [3], fue la encargada de desarrollar este sistema, al cual llamó International Mobile Telecommunications (IMT-2000); realmente no fue un solo sistema, sino un conjunto de sistemas 3G utilizados en diferentes países. Esto fue debido a las diferencias entre los intereses de los distintos países participantes en el proyecto, que al no ponerse de acuerdo optaron por sistemas diferentes, pero pertenecientes a un mismo conjunto, IMT-2000. Por ejemplo, en Europa se adoptó el sistema Universal Mobile Telecommunications System (UMTS) [16] o en EE.UU el sistema Code Division Multiple Access (CDMA-2000).

El lanzamiento de UMTS fue en el año 1999, (Release 99). Este sistema puede trabajar en modo FDD o TDD, utilizando para ello dos técnicas de acceso, bien Wideband CDMA (W-CDMA) [16]; o bien Time Division CDMA (TD-CDMA).

UMTS permite alcanzar una velocidad de hasta 2 Mbps en entornos de baja movilidad, lo que permite cumplir con la demanda de los usuarios en cuanto a aplicaciones multimedia, acceso a internet, transferencia de archivos de mayor o menor tamaño e incluso transmisión de datos en tiempo real como video-llamadas.

Los principales cambios que experimenta la arquitectura de la red GSM/GPRS para hacer posible el despliegue de la red UMTS, es la introducción de los NodosB que sustituyen a las anteriores Base Transceiver Station (BTS) de GSM y las Radio Network Controller (RNC), que tienen funciones similares a las BSC de GSM.

El avance más importante, llega con la aparición de High Speed Downlink Packet Access (HSDPA) [2] y posteriormente de High Speed Uplink Packet Access (HSUPA). Ambos se combinan y evolucionan a High Speed Packet Access (HSPA) [2]. Las

novedades introducidas por HSPA [16] permiten lograr grandes avances en la transmisión de datos, debido, por ejemplo, a la utilización de nuevos esquemas de modulación o la introducción de mecanismos de retransmisión híbrida, “Hybrid Automatic Repeat reQuest” (HARQ), todo ello para lograr unas velocidades de hasta 14 Mbps para el DL y hasta 6 Mbps para el UL. Posteriormente surge una nueva evolución, HSPA+, cuya principal novedad es que es un sistema Multiple Input Multiple Output (MIMO), lo que le permite obtener velocidades entorno a los 80 Mbps para DL y 20 Mbps para el UL. HSPA y HSPA+ podrían considerarse como sistemas intermedios entre 3G y 4G conocidos como 3.5G.

2.4 Sistemas 4G.

Finalmente, llegamos al último punto en la evolución de los sistemas de comunicaciones móviles, llegamos a los sistemas 4G [3], que es además el tema principal que nos ocupa en este trabajo en concreto, el sistema LTE [17]. De nuevo, al igual que en las evoluciones anteriores, se busca un tipo de sistema capaz de satisfacer las nuevas necesidades que van surgiendo, ya que cada vez son más exigentes en cuanto a las capacidades de los sistemas se refiere, se exige sobre todo, cada vez más capacidad para el tráfico de datos.

La rama de la ITU especializada en radiocomunicación (ITU-R), fue la institución encargada de establecer los nuevos requisitos de los sistemas de cuarta generación [4], siendo el principal de estos requisitos el hecho de que fueran redes completamente basadas en conmutación de paquetes, utilizando para ello el protocolo IP. Se evoluciona hacia las llamadas redes “todo IP”. Con estas nuevas redes se pretenden obtener unas velocidades de hasta 1Gbps en entornos de baja movilidad y hasta 100Mbps con alta movilidad.

Tres son las organizaciones que se encargan de desarrollar estos nuevos sistemas: el Third Generation Partnership Project (3GPP) [18], el 3GPP2 y por otro lado el Institute of Electrical and Electronics Engineers (IEEE) [19].

El 3GPP, realizó la primera especificación del sistema LTE, que después evolucionó a LTE-Advanced [17]. El 3GPP2, comenzó a desarrollar el sistema Ultra Mobile

Broadband (UMB), con la intención de convertirlo en el próximo sistema 4G, pero finalmente el proyecto no llegó a consolidarse y se centraron en el desarrollo de LTE. Por último, el IEEE creó la familia de estándares 802.16 (WiMAX).

Finalmente, la ITU reconoce LTE-Advanced como una de las dos tecnologías, junto con WirelessMAN-Advanced (802.16m), que pueden considerarse oficialmente como redes 4G.

En la figura 1, se muestra un esquema en el que se resume brevemente la evolución descrita anteriormente:

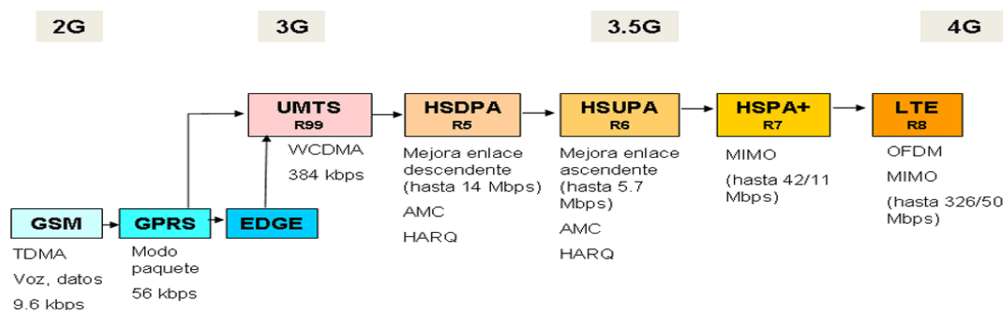


Figura 1: Evolución de los sistemas de comunicaciones móviles.

Fuente: C.M.T.

Con esto, concluye la introducción que describe la evolución de los distintos sistemas de comunicaciones móviles. A continuación, se profundizará en el sistema LTE, ya que es el sistema en el que se centra fundamentalmente este trabajo.

3. Sistema LTE

Como he comentado anteriormente, el sistema LTE fue desarrollado por el 3GPP para satisfacer las nuevas necesidades que fueron surgiendo y que no eran cubiertas por los sistemas anteriores. Tiene unas velocidades objetivo máximas de 1Gbps para el DL y en torno a los 50Mbps para el UL, las cuales resultan más que suficientes para cubrir las expectativas. Por último, recordar que la principal característica de este sistema es que es una red basada por completo en IP.

3.1 Interfaz radio LTE.

La interfaz radio [8], es la interfaz final del sistema LTE, que separa el medio de transmisión de las técnicas de procesamiento en banda base anteriores a la transmisión. El procesamiento de la señal de Radio Frecuencia (RF) tiene ciertos inconvenientes o dificultades, debido a que la interfaz utilizada es el aire, “interfaz aire”, y es un medio compartido por múltiples portadoras de RF. En concreto, para LTE, los requisitos exigidos a los transceptores no son mucho más complejos de los requeridos para UMTS, por ello, muchos de los requisitos de RF para LTE se derivan de los ya establecidos para UMTS.

Sin embargo, hay una serie de diferencias importantes entre la interfaz radio de LTE [17] y la de UMTS [16], las cuales pasaré a describir a continuación:

- La primera de estas diferencias es que en LTE se utiliza un ancho de banda de canal variable, hasta un máximo de 20 MHz. En general, suelen ser 1.4, 3, 5, 10, 15 y 20 MHz. Siendo esta última configuración la que permite alcanzar velocidades de hasta 100 Mbps para el DL. Esto implica que son necesarios unos requerimientos diferentes de RF para cada ancho de banda, a diferencia del modo FDD de UMTS, en el que solo se utilizaba un ancho de banda de canal de 5MHz. Esto representa un nuevo desafío en el diseño de la interfaz radio de LTE, ya que implica que los transceptores de LTE deben ser más flexibles que en los sistemas

anteriores. La segunda diferencia de LTE con respecto a UMTS, es que en LTE se asume que el User Equipment (UE) tiene dos antenas receptoras, lo que significa que permite múltiples trayectos de la señal.

- En tercer lugar, LTE es más flexible en cuanto a las velocidades de datos que soporta, para adaptarse a las diferentes condiciones de Relación Señal a Ruido más Interferencia, “Signal to Interference plus Noise Ratio” (SINR). Además, tiene la capacidad de variar el ancho de banda para un determinado usuario, lo que implica un gran número de modos de operación y una mayor flexibilidad en el manejo de la señal.
- Cuarta, la propia forma o estructura de la señal de LTE, que favorece la alteración de los aspectos que resultan más críticos en las transmisiones de radiofrecuencia. En LTE, se emplean dos técnicas de acceso, Orthogonal Frequency Division Multiple Access (OFDMA) [17] para el DL y Single Carrier Frequency Division Multiple Access (SC-FDMA) [17] para el UL, lo que garantiza una cierta robustez frente a la propagación multi-trayecto, esto implica que las distorsiones de amplitud y fase de los filtros del transmisor y receptor no son tan críticas como para UMTS, que emplea WCDMA [16]. Por otra parte, OFDMA requiere una mejor sincronización de frecuencia y es más sensible al ruido de fase.
- Por último, destacar que para WCDMA las especificaciones requeridas en RF son diferentes para los modos FDD y TDD, mientras que en LTE las similitudes entre ambos modos son tales, que permiten utilizar ambos modos bajo las mismas especificaciones.

A continuación se muestran las bandas de frecuencia empleadas por LTE y cómo comparten el espectro con las bandas ya utilizadas por otros sistemas, en especial con UMTS. LTE y UMTS están definidos para trabajar en un amplio rango de frecuencias. En la figura 2 se muestran las frecuencias utilizadas por LTE y UMTS en el modo FDD:

Band Number	Uplink (MHz)	Downlink (MHz)	Band Gap (MHz)	Duplex Separation (MHz)	UMTS Usage	LTE Usage
	$F_{UL, low}-F_{UL, high}$	$F_{DL, low}-F_{DL, high}$				
1	1920-1980	2110-2170	130	190	Y	Y
2	1850-1910	1930-1990	20	80	Y	Y
3	1710-1785	1805-1880	20	95	Y	Y
4	1710-1755	2110-2155	355	400	Y	Y
5	824-849	869-894	20	45	Y	Y
6*	830-840	875-885	35	45	Y	Y
7	2500-2570	2620-2690	50	120	Y	Y
8	880-915	925-960	10	45	Y	Y
9	1749.9-1784.9	1844.9-1879.9	60	95	Y	Y
10	1710-1770	2110-2170	340	400	Y	Y
11	1427.9-1447.9	1475.9-1495.9	28	48	Y	Y
12	698-716	728-746	12	30	Y	Y
13	777-787	746-756	21	31	Y	Y
14	788-798	758-768	20	30	Y	Y
17	704-716	734-746	18	30	N	Y
18**	815-830	860-875	30	45	N	Y
19**	830-845	875-890	30	45	Y	Y
20**	832-862	791-821	11	41	Y	Y
21**	1447.9-1462.9	1495.9-1510.9	33	48	Y	Y
23***	2000-2020	2180-2200	160	180	N	Y
24***	1626.5-1660.5	1525-1559	-135.5	-101.5	N	Y
25***	1850-1915	1930-1995	15	80	Y	Y
26****	814-849	859-894	10	45	Y	Y

* This band was defined in the context of Release 8; it is replaced by Band 19 for later releases (Release 9 and 10). Only legacy terminals would use band 6.

** These bands were specified in the timeframe of Release 9, although all bands are release-independent and can be implemented by UEs conforming to any release.

*** These bands were specified in the timeframe of Release 10, although all bands are release-independent and can be implemented by UEs conforming to any release.

**** This band was under consideration at the time of going to press.

Figura 2: Frecuencias modo FDD

Fuente: [4]

Misma relación de frecuencias pero para el modo TDD:

Band	$F_{\text{low}}-F_{\text{high}}$ (MHz)	UMTS	LTE
33	1900–1920	Y	Y
34	2010–2025	Y	Y
35	1850–1910	Y	Y
36	1930–1990	Y	Y
37	1910–1930	Y	Y
38	2570–2620	Y	Y
39	1880–1920	N	Y
40	2300–2400	Y	Y
41*	2496–2690	N	Y
42*	3400–3600	N	Y
43*	3600–3800	N	Y

* These bands were specified in the timeframe of Release 10, although all bands are release-independent and can be implemented by UEs conforming to any release.

Figura 3: Frecuencias modo TDD

Fuente: [4]

Otro aspecto a destacar en LTE, es que debido a su flexibilidad en la utilización de frecuencias, introduce dos nuevos conceptos, el primero es la configuración del ancho de banda de transmisión, que define el máximo número de Resource Blocks (RB) para el ancho de banda del canal. El segundo concepto es el ancho de banda de transmisión, que está estrechamente relacionado con la asignación de recursos, puede ser menor o igual que la configuración del ancho de banda de transmisión, como se muestra en la figura 4:

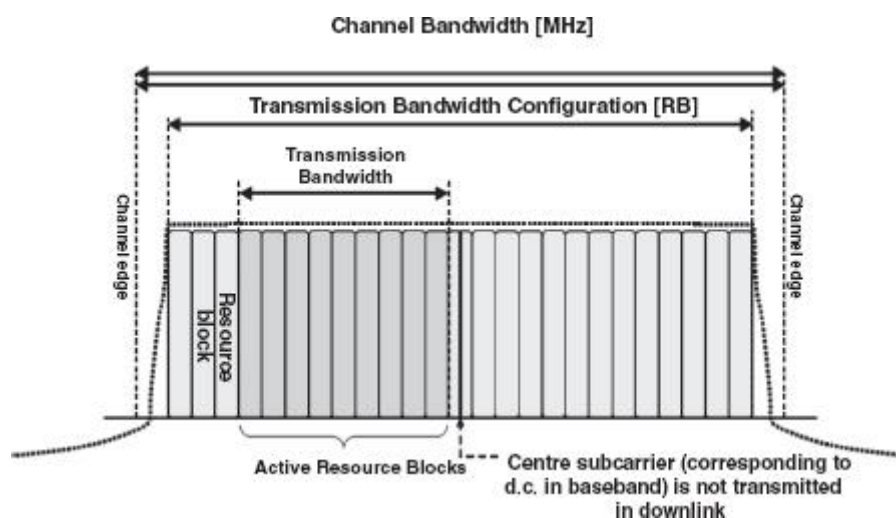


Figura 4: Configuración del ancho de banda de transmisión vs ancho de banda de transmisión.

Fuente: [4]

En cuanto a la arquitectura de la red LTE [4], esta está pensada teniendo en cuenta tres aspectos fundamentales: que tuviera un coste reducido, lograr una baja latencia y el principal objetivo, que fuera una red únicamente de conmutación de paquetes. Esta nueva arquitectura de la red LTE cuenta con una nueva red de acceso denominada Evolved-UMTS Terrestrial Radio Access Network (E-UTRAN) y una nueva red troncal llamada Evolved-Packet Core (EPC), ambas forman el Evolved Packet System (EPS). Como se observa, ya la propia denominación refleja la evolución hacia un “entorno de paquetes”.

La E-UTRAN [4], está formada por los denominados evolved-NodeB (eNB), con funciones similares a los NodosB de UMTS o BTS de GSM, que son las estaciones base de LTE. La principal diferencia de estos nuevos eNBs con respecto a las anteriores estaciones base, es que los eNBs realizan las mismas funciones que el conjunto formado por ejemplo, por BTS y BSC en GSM o RNC y NB en UMTS, es decir, se concentran las funciones de las estaciones base y los controladores en el eNB. Esto supone que los eNBs realizan un gran número de funciones: gestión de recursos radio, conectividad, seguridad, etc. Para llevar a cabo estas funciones, los eNBs se conectan entre sí a través de la interfaz X2. Del mismo modo, se conectan con los terminales de los usuarios a través de la interfaz Uu y con la red troncal mediante la interfaz S1.

La EPC [4], es la parte de la red encargada de conectar la red LTE con otras redes externas, proporciona conexión vía IP con otras redes o servicios. Está formada principalmente por el Mobility Management Entity (MME), el cual lleva a cabo las funciones de control entre el UE y la EPC. Y por otro lado el Serving Gateway (S-GW) y el Packet Data Network Gateway (P-GW); en estos se concentran las funciones del plano de usuario.

En la figura 5 se muestra un esquema simplificado de la arquitectura de la red LTE:

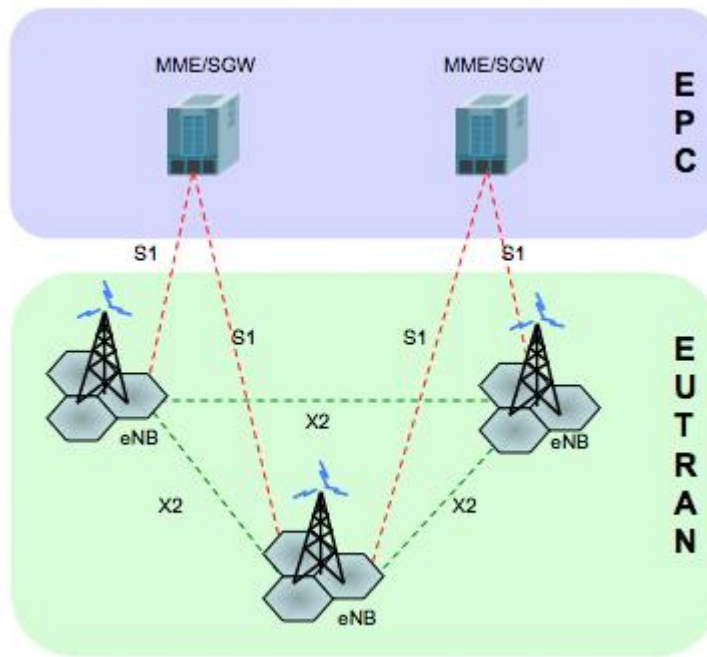


Figura 5: Arquitectura LTE

Fuente: ltemobilezone.wordpress.com

3.2 El canal LTE, MIMO.

Debido a los diferentes escenarios en los que se despliegan las redes LTE, se dan diferentes entornos de propagación, desde medios rurales a densas zonas urbanas o espacios tanto interiores (indoor) como exteriores (outdoor).

Cada uno de estos entornos tiene unas características concretas de propagación, al mismo tiempo, éstas se ven afectadas por las distintas frecuencias de portadora utilizadas. En el 2007, la ITU asignó nuevas bandas de frecuencia para las comunicaciones móviles (IMT) entre los 450 MHz y los 3.6 GHz. Dependiendo de la frecuencia de portadora utilizada, las características del canal LTE [4] pueden variar significativamente, incluso dentro de un mismo entorno. Al mismo tiempo, el hecho de utilizar una portadora en una frecuencia u otra también tiene un especial impacto en el desarrollo de técnicas MIMO, debido a que el tamaño de la antena depende de la longitud de onda de la portadora. La utilización de técnicas MIMO es una propiedad característica de LTE, mientras que las redes móviles convencionales están diseñadas para contrarrestar las pérdidas provocadas por el multi-trayecto, en LTE, la técnicas

MIMO permiten aprovechar el multi-trayecto en beneficio del sistema, por ejemplo para aumentar la capacidad. Por ello, es importante para el buen rendimiento del sistema LTE, crear modelos estandarizados de canales MIMO que puedan servir para evaluar el rendimiento de LTE [4] y realizar posibles mejoras en el futuro.

A continuación pasaremos a estudiar con más detenimiento el canal MIMO, ya que tiene una gran importancia en el sistema LTE.

Los canales MIMO, como su propio nombre indica, tienen múltiples antenas y por tanto múltiples trayectos de la señal. Se aprovechan las múltiples señales que llegan al receptor en beneficio del sistema, por ejemplo para aumentar la capacidad o reducir la tasa de errores. Esto es debido, fundamentalmente, a las reflexiones que se producen en la transmisión de la señal, que provocan el llamado multi-trayecto, el cual hace que lleguen varias señales con la misma información al receptor pero espaciadas en el tiempo y/o en el espacio.

Primero, destacar que el rendimiento de los sistemas MIMO depende en gran medida de las condiciones de propagación. En primer lugar, la SINR en los receptores de los sistemas MIMO determina la ganancia final de diversidad espacial. En segundo lugar, otro factor importante, es la correlación existente entre las señales de diferentes antenas, la separación de estas antenas condiciona en gran medida la correlación espacial de las señales. Las mejores ganancias en estos sistemas se obtienen en escenarios con una alta SINR y baja correlación espacial, algo que no suele ocurrir en la realidad. En el caso de LTE, en los eNBs se configuran las antenas con una polarización cruzada utilizando cuatro antenas, se tienen por tanto dos elementos con dos dobles polarizaciones con una inclinación de $+45^\circ$ y -45° y separados entre sí una distancia d , como se observa en la figura 6, esto permite al sistema beneficiarse de la diversidad espacial.

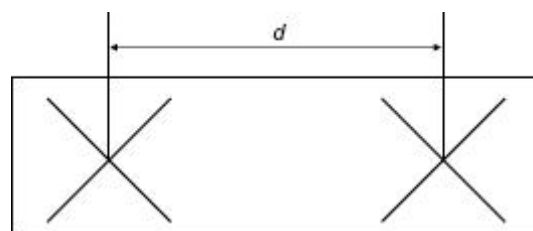


Figura 6: Antenas eNBs

Fuente: [4]

Normalmente se utilizan dos modelos principales de canales MIMO [4]:

- Modelos de canal basados en matriz de correlación.
- Modelos de canal basados en geometría.

Ambos modelos se utilizan en LTE. Por ejemplo, el Spatial Channel Model (SCM) del 3GPP, sigue un modelo estocástico basado en geometría, mientras que los modelos de la ITU, como EPA o EVA se basan en matriz de correlación.

Es importante especificar muy bien las condiciones de propagación para poder evaluar el rendimiento de LTE. Es necesario tener modelos que sean precisos con las distintas características espaciales de los canales MIMO, para lograr así los mejores esquemas de transmisión multi-antena.

3.3 Los enlaces descendente y ascendente.

En toda red de comunicaciones móviles que se precie, tiene una gran importancia la elección de la modulación y técnicas de acceso a utilizar si se quiere lograr un buen rendimiento del sistema. En el caso de LTE, esta elección cobra una mayor importancia, si cabe, al tratarse de una red completamente orientada a datos, en particular por el hecho de que los canales radio tienden a ser más dispersivos y variantes en el tiempo, por lo que se ha optado por utilizar una modulación multi-portadora. Los esquemas multi-portadora dividen el ancho de banda de los canales utilizados en un conjunto de sub-canales paralelos, como se observa en la figura 7 (a), esto tiene la ventaja de que el receptor puede compensar fácilmente las ganancias de cada sub-canal en el dominio de la frecuencia. En el caso concreto del DL en LTE, se emplea un esquema especial multi-portadora llamado OFDMA, que es una extensión del esquema Orthogonal Frequency Division Multiplexing (OFDM) para sistemas multiusuario. En la figura 7 se pueden observar las diferencias entre una modulación multi-portadora clásica (a) y OFDM (b):

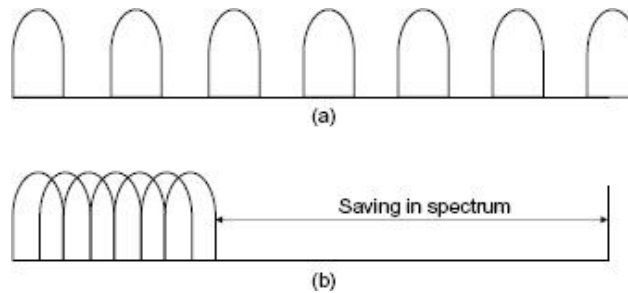


Figura 7: Espectro modulación multi-portadora vs OFDM

Fuente: [4]

Tanto OFDM como OFDMA son casos especiales de transmisión multi-portadora [4], donde los sub-canales se solapan pero son ortogonales entre sí. Esto evita tener que separar las portadoras mediante las denominadas bandas de guarda y, por lo tanto, hace que el sistema sea más eficiente espectralmente hablando. El espaciado entre los sub-canales permite que el receptor sea capaz de diferenciarlos sin problemas, lo que facilita la implementación del receptor. Este hecho hace que resulte especialmente atractivo para las transmisiones de datos de alta velocidad, como es el caso del DL de LTE. Además, en el DL se combina OFDMA con técnicas de codificación de canal y técnicas HARQ para combatir los desvanecimientos que se producen en los diferentes sub-canales.

Para terminar con lo referente al DL, únicamente señalar que OFDMA es una extensión de OFDM aplicada a sistemas multiusuario. Distribuye las portadoras entre varios usuarios de manera simultánea, lo que permite a varios usuarios recibir datos al mismo tiempo. Incorpora a OFDM los beneficios de la diversidad multiusuario y mejora su eficiencia espectral.

Pasemos ahora a comentar las principales características del UL [4]. Realmente no hay grandes diferencias con el DL en cuanto a los requerimientos se refiere:

- Necesita una transmisión ortogonal por parte de los diferentes usuarios para minimizar las interferencias dentro de la celda y aumentar la capacidad del enlace.

- Debe tener flexibilidad para adaptarse a diferentes tasas de transmisión de datos.
- Cuanto menor sea el Peak to Average Power Ratio (PAPR) de la onda transmitida menor será el tamaño y consumo del Power Amplifier (PA).
- Debe soportar técnicas multi-antena para beneficiarse de la diversidad espacial.

La técnica de acceso elegida para tratar de cumplir con estas y otras necesidades en el UL de LTE es SC-FDMA.

La mayor ventaja que presenta esta técnica frente a las utilizadas por ejemplo en UMTS, Direct Sequence Code Division Multiple Access (DS-CDMA), es que consigue ortogonalidad dentro de la celda, incluso en canales selectivos en frecuencia. SC-OFDMA reduce considerablemente las interferencias dentro de la celda, aumentando así la capacidad de la misma.

Esta técnica es especialmente apropiada para el UL de LTE [4], ya que es común con las técnicas empleadas en el DL, es decir, son técnicas similares para el DL y UL. Igual que ocurriera con OFDM/OFDMA, SC-FDMA divide el ancho de banda en varias sub-portadoras ortogonales entre sí. Sin embargo, la principal diferencia es que, en SC-FDMA, la señal modulada en una sub-portadora es una combinación lineal de todos los símbolos transmitidos en el mismo instante de tiempo, por lo que en cada periodo de símbolo, todas las sub-portadoras transmitidas en la señal llevan una componente de cada símbolo modulado. Esto permite tener un PAPR más bajo que otros esquemas multi-portadora como OFDM, lo cual es crucial en el UL, ya que como he comentado anteriormente, permite reducir tanto el tamaño como el consumo de potencia del PA del terminal.

4. Asignación de recursos en LTE

Como se ha comentado anteriormente, las redes 4G están totalmente orientadas a la transmisión de paquetes, lo cual supone un nuevo desafío a la hora de asignar los recursos radio disponibles. Son necesarias mejoras en las técnicas existentes o nuevas técnicas para lograr una asignación de recursos lo más eficiente posible y permitir dar la mejor calidad a cada uno de los servicios demandados por los usuarios.

4.1 Fundamentos de la asignación de recursos en LTE:

La asignación de recursos consiste en asignar RBs entre los usuarios que van a transmitir en cada Transmission Time Interval (TTI). En LTE, el eNB es el encargado de la planificación y asignación de recursos, tanto para el DL como para el UL. Su principal objetivo es satisfacer las necesidades del mayor número de usuarios posible, teniendo en cuenta las diferentes demandas de calidad de servicio para las distintas aplicaciones que estén utilizando.

En LTE, el conjunto de técnicas encargadas de gestionar los recursos radio se engloba en el Radio Resource Management (RRM) [5]. Éste se encarga de diversas funciones, entre ellas está la de “planificar” la asignación de recursos, es lo que se conoce como scheduling. Éste decide qué usuarios son los que van a enviar o recibir la información en cada momento, es lo que se conoce como asignación dinámica de recursos.

Una de las principales ventajas de este tipo de asignación de recursos es que contribuye a que el sistema LTE pueda tener un factor de reutilización de frecuencias de 1 [8], es decir, todas las frecuencias están disponibles para cada celda. Pero hay que tener en cuenta que los usuarios del borde de la celda sufren niveles de interferencias muy elevados y por ello deben adoptarse estrategias de compensación para la reducción de dichas interferencias. Algunos ejemplos de estas técnicas son la aleatorización de la interferencia intercelular o la coordinación de interferencias (ICIC), siendo esta última la más importante ya que es imprescindible para el buen funcionamiento del sistema.

Esta función de planificación de recursos radio o scheduling [4], es llevada a cabo por los eNBs y se basa principalmente en tener siempre controlados los buffer de

transmisión y recepción para saber en todo momento qué terminales son los que tienen que enviar o recibir información y asignarles los recursos necesarios. Otro aspecto importante que se tiene en cuenta para la asignación es la SINR, que permite obtener información del estado del canal para cada usuario. De igual modo, se tiene en cuenta la QoS [7] necesaria para cada usuario, ya que dependiendo del tipo de servicio que demande el usuario se le asignarán unos recursos u otros.

La asignación de recursos en LTE [6] se realiza en cada intervalo de transmisión o TTI, lo que corresponde con la duración de una sub-trama, que a su vez se compone de dos slots de 0,5 ms cada uno. Éstos están formados por un conjunto de 12 sub-portadoras separadas entre sí 15 kHz, lo que equivale a un conjunto de 180 kHz.

Por otra parte, destacar que el mínimo recurso que se puede asignar a un usuario es un RB, que son 12 sub-portadoras de una sub-trama, manteniendo las mismas sub-portadoras en los dos slots que forman la sub-trama o bien cambiando las sub-portadoras en cada slot mediante frequency hopping. Esto permite distinguir entre servicios que mantienen un mismo RB en dos slots consecutivos (1 sub-trama) o bien los servicios en los que se cambian los RB en cada slot. Por último, se denominan Resource Block Group (RBG) a un conjunto consecutivo de RB.

En la figura 8, se muestra una imagen aclaratoria que permite entender mejor lo explicado anteriormente:

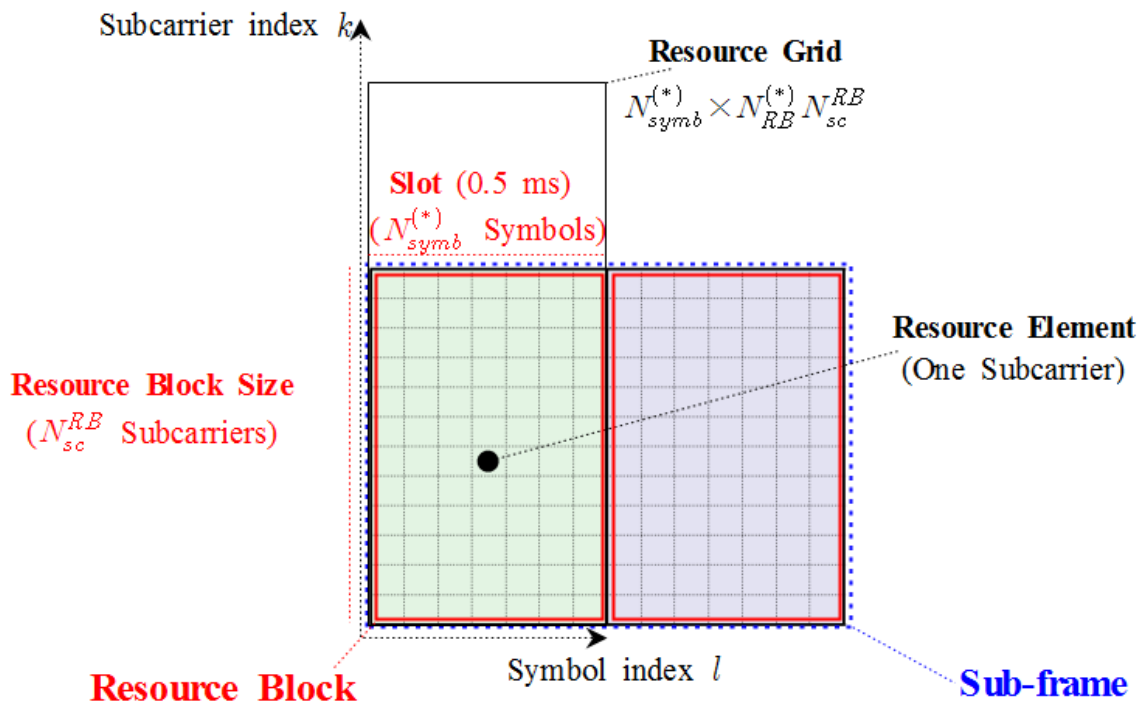


Figura 8: Esquema asignación de recursos en LTE

Fuente: [6]

En cuanto a las diferencias entre la asignación de recursos para el DL y UL [7], podría decirse que no son muy significativas, solo hay que tener en cuenta algunas pequeñas diferencias. Por ejemplo, en el enlace ascendente, al tratarse de SC-FDMA, las subportadoras deben asignarse de forma contigua, esto no tiene por qué ser así para el enlace descendente. Aunque la principal diferencia reside en que estas funciones de asignación de recursos (scheduling) se llevan a cabo en los eNBs, por lo que para el enlace ascendente es necesaria una mayor señalización de cada UE y por tanto hace este proceso más complejo para este enlace.

En lo referente a la señalización [9], los algoritmos de asignación de recursos necesitan información proveniente de diversas medidas para que la toma de decisiones se haga de la forma más eficiente posible. Éstas suelen medirse sobre el estado del canal, el volumen de tráfico o el estado de las colas.

Los indicadores del estado del canal (CQI) [4] se utilizan en el DL y son necesarios para elegir los mejores RB para cada usuario. El eNB configura los mensajes CQI para que se utilicen determinadas sub-bandas o RB. En cuanto al UL, las condiciones de canal se obtienen de forma diferente, en este caso se utilizan señales de referencia llamadas SRS. La SRS puede configurarse para transmitir con un ancho de banda determinado, cuanto mayor sea el ancho de banda sobre el que se tiene que informar menor será la potencia disponible por RB en el UE.

Es importante destacar que en LTE cada canal lógico lleva asociada una QoS [8], la cual se actualiza en función de las condiciones de canal, de tráfico o del estado de las colas. Estos datos los obtienen, o bien los eNBs, o bien el propio terminal de usuario UE, que posteriormente los envía al eNB a través del tráfico de señalización. Es importante destacar que siempre existe una estrecha relación entre el tráfico de señalización y el dedicado a datos, ya que un exceso de señalización provocaría una disminución considerable de la eficiencia en el tráfico de datos. Siempre existe un cierto compromiso entre ambos.

4.2 Asignación de recursos en el UL:

Como se ha comentado con anterioridad, el proceso de asignación de recursos es bastante similar para el DL y UL [10], pero existen ciertas diferencias que se deben tener en cuenta. Las principales a destacar son:

- El UE tiene ciertas limitaciones de potencia, lo que implica también ciertas limitaciones para el UL. Esto se traduce en que los usuarios del borde de la celda tienen ciertas limitaciones para usar un mayor ancho de banda para transmitir y compensar así las interferencias, ya que un ancho de banda mayor implica utilizar más potencia.
- En el UL se utiliza SC-FDMA [4], lo que obliga a asignar los RBs de forma contigua a los usuarios, esto impide explotar la diversidad multiusuario del canal.
- Otra diferencia importante a tener en cuenta, es que en el UL la asignación de recursos la lleva a cabo el eNB para todos los usuarios de la celda. Por lo que es necesaria una mayor señalización que para el DL.
- Los terminales necesitan un cierto tiempo de procesado. Este tiempo, necesario para el terminal, introduce un retardo de unos 3 a 4 ms, por lo que cuando se asigna la transmisión en un TTI n , la transmisión se realiza en el instante $n+3$ o $n+4$.
- En el UL cada usuario tiene varias colas con datos para transmitir, cada una se corresponde con canales lógicos de subida diferentes, cada uno con distintos retardos y distintas limitaciones de capacidad. En cuanto al DL, el eNB puede mantener varios buffers con tráfico de datos dedicado para cada usuario, cada una de estas colas tiene unas limitaciones de calidad de servicio distintas.
- Por último, en el UL se registra un elevado margen dinámico en el nivel de interferencias entre TTI consecutivos, lo que obliga a utilizar técnicas rápidas, que se basan en el nivel de señal deseado en lugar de basarse en el nivel de interferencias.

El proceso de asignación de recursos se divide en dos partes, asignación en el dominio de la frecuencia y asignación en el dominio del tiempo.

- **Asignación en el dominio de la frecuencia**, Frequency Domain Packet Scheduling (FDPS). Ésta se basa en el canal para realizar la asignación óptima para cada usuario, utiliza la calidad de canal de cada usuario en el instante de tiempo determinado para realizar la asignación.
- **Asignación en el dominio del tiempo**, Time Domain Packet Scheduling (TDPS). En este caso se trabaja con TTI, se detectan los usuarios que pueden transmitir en el próximo TTI, teniendo en cuenta para ello la información de los bloques de gestión como el DRX, HARQ [10], el gestor de colas, etc.

A continuación, pasaré a describir brevemente los distintos bloques que llevan a cabo las funciones de gestión de recursos (RRM) [8] en el UL:

- **Adaptive Modulation and Coding** (AMC): da información sobre el estado del canal para un usuario determinado y permite utilizar las técnicas más convenientes en cada situación.
- **Discontinuous Reception and Transmission** (DRX/DTX): Bloque encargado de la transmisión/recepción discontinua.
- **Buffer Status Report** (BSR): Informa sobre el estado de las colas de los usuarios.
- **Multiple Input Multiple Output** (MIMO): permite asignar recursos a usuarios distintos sin tener por ello que utilizar señalización extra.
- **Hybrid ARQ** (HARQ): consiste en un esquema de retransmisión híbrido síncrono.
- **Power Control** (P-C): Bloque de control de potencia, limita la potencia para intentar controlar la interferencia entre celdas.
- **Quality of Service** (QoS): Este bloque es el encargado de que se cumpla con los requerimientos de calidad que desea el usuario para un determinado servicio.

Para llevar a cabo la asignación de recursos en el UL se utilizan principalmente dos tipos de algoritmos [4]:

- **Fixed Transmission bandwidth (FTB)**: Se fija un número consecutivo de RBs a cada usuario. El número de usuarios es fijo.
- **Adaptive Transmission Bandwidth (ATB)**: Igual que en el caso anterior pero ahora el algoritmo considera las métricas de cada UE y RB. Se asigna un RB al usuario pero la asignación se expande a los RB contiguos, siempre que el usuario tenga la mejor métrica.

En la figura 9 se puede apreciar la estructuración típica de una celda de un sistema radio, la asignación de recursos se basa en esta estructura:

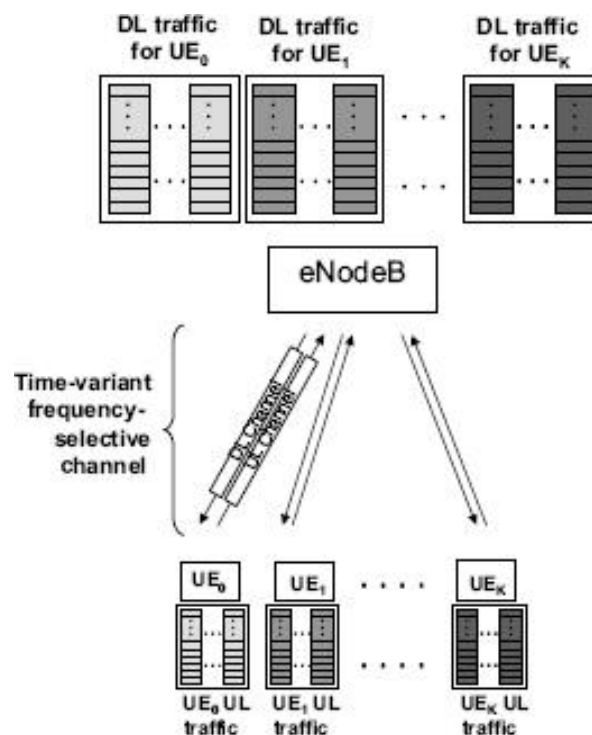


Figura 9: Estructuración típica del tráfico de una celda.

Fuente: [4]

4.3 Asignación de recursos en el DL:

En el sistema LTE se asignan recursos en cada TTI a los diferentes usuarios, se envía a cada usuario un conjunto de bits que representan los RBs asignados. Enviar a cada usuario un conjunto de bits, donde cada bit represente un RB, supondría demasiado tráfico de señalización para anchos de banda elevados y no sería viable con la eficiencia que se busca del sistema. Para evitar este problema se utilizan distintas técnicas [8]:

- **Mapa de bits de tipo 0:** Lo que se hace es que en lugar de asignar un único RB a un usuario (UE), se le asocia un conjunto de RBs, son los llamados RBG, el tamaño de estos puede variar entre 2, 3 o 4, en función de los RBs a asignar y del ancho de banda disponible. (Ver tabla 1.)
- **Mapa de bits de tipo 1:** Es similar a la estrategia anterior, con la diferencia de que permite asignar un único RB de manera individual en cada RBG. Esta técnica solo se utiliza en el DL.
- **Mapa de bits de tipo 2:** Se indica el RB inicial y el número de RB que se quieren asignar de forma contigua. Se utiliza tanto para UL como para DL.
- **Mapa de bits directo:** Se asigna un bit por RB, pero como hemos comentado anteriormente no es viable con anchos de banda elevados.
- **Ubicación distribuida:** Permite asignar RB de forma independiente aunque no sean contiguos. Sólo se utiliza en el DL.

La tabla 1 recoge los valores de RBG necesarios en función del ancho de banda utilizado:

Ancho de banda DL N_{RB}	Tamaño RBG (P)
0 <= 10	1
11-26	2
27-63	3
64-110	4

Tabla 1: RBG necesarios por ancho de banda utilizado.

Fuente: [8]

Comentar que las anteriores son técnicas de asignación a nivel físico y existen a su vez dos técnicas principales a nivel MAC [4] y [8]:

- Asignación basada en peticiones dinámicas de recursos: pensada para la transmisión de datos.
- Asignación estática: especialmente pensada para el tráfico de voz ya que al ser fija se reduce el tráfico de señalización.

Al igual que en el UL, en el DL también tenemos asignación de recursos en el dominio de la frecuencia y en el dominio del tiempo [8].

La asignación de recursos en el dominio de la frecuencia, FDPS, se basa en asignar RB a los usuarios teniendo en cuenta las condiciones de canal de cada uno. Asigna recursos a los usuarios con mejores condiciones de canal. Es importante tener en cuenta que cuando se utiliza asignación en frecuencia, la capacidad media de la celda es inversamente proporcional a la velocidad de los usuarios, es decir, la capacidad de la celda disminuye cuando aumenta la velocidad de los terminales, ya que no pueden seguirse las variaciones del canal.

Como se ha comentado en apartados anteriores, LTE permite, teóricamente, una reutilización de frecuencias de 1, en la práctica no suele darse este caso, debido principalmente a las interferencias que sufrirían los usuarios del borde de la celda. Aún así se pueden establecer distintos esquemas de carga fraccional [11] clasificando a los usuarios en distintos grupos y aplicando a cada grupo un esquema de reutilización distinto.

En cuanto a **la asignación en el dominio del tiempo**, TDPS, también explota la diversidad multiusuario como en la asignación en el dominio de la frecuencia aunque a un menor nivel, debido fundamentalmente a la coexistencia con otras técnicas de diversidad, como la diversidad debida al multi-trayecto, la diversidad de los terminales o diversidad en la transmisión en las estaciones base.

El proceso para realizar la asignación en el dominio del tiempo y la frecuencia se divide en tres pasos fundamentales [8]:

- En primer lugar, se hace la asignación en el dominio temporal, se eligen los usuarios que van a transmitir en el próximo TTI, teniendo en cuenta para ello diversos aspectos ya preestablecidos, como pueden ser las condiciones del canal para el usuario, el ancho de banda necesario, QoS necesaria, etc. Por otra parte, destacar que se da siempre máxima prioridad a usuarios en retransmisión [10], que transmiten los mismos paquetes que intentaron transmitir en el primer intento, por tanto utilizan el mismo número de RBs, lo que permite conocer los recursos necesarios con anterioridad, ya que son los mismos que se necesitaron en el primer intento de transmisión.
- El segundo paso es comprobar si hay recursos suficientes para todos los usuarios, teniendo en cuenta tanto a los usuarios que intentan hacer su primera transmisión, como a los que tienen que hacer retransmisiones. En caso de que no hubiera recursos suficientes para todos, se establece un subconjunto del total de usuarios que recibirá la asignación.

- En tercer y último lugar, se lleva a cabo la asignación en frecuencia, es decir, se asignan los RBs a los usuarios seleccionados previamente. Se asignan los mejores bloques a los usuarios que transmiten por primera vez y los restantes a los usuarios que retransmiten. De esta forma se logra máxima capacidad en la celda al asignar los mejores recursos a los usuarios que transmiten por primera vez y al mismo tiempo se reducen los retardos para los usuarios que tienen que retransmitir.

5. Estrategias de asignación de recursos en el DL

A continuación, nos centraremos en el estudio de las distintas estrategias empleadas para asignar los recursos en el enlace descendente del sistema LTE. Son una serie de algoritmos que permiten establecer diferentes prioridades o estrategias a la hora de asignar los recursos disponibles a los usuarios. Es importante comentar que los algoritmos utilizados están estrechamente relacionados con el esquema AMC utilizado y con el modulo HARQ de retransmisión [4], debido, primeramente, a que la asignación dinámica de recursos se utiliza junto con las condiciones del canal para adaptar la modulación y la codificación (AMC) y, en segundo lugar, el funcionamiento de las colas, que condiciona fuertemente tanto la tasa de transmisión como los retardos y depende estrechamente del protocolo HARQ y del tamaño de los bloques de transporte.

Por otra parte, los algoritmos de scheduling [5] pueden obtener la información que utilizan para tomar sus decisiones de planificación de dos maneras distintas. Obtienen información a través del estado del canal y/o mediante diferentes medidas del tráfico. Dicha información se obtiene o bien de los eNBs o mediante feedback de los canales de señalización, también pueden combinarse ambas formas. Hay que tener en cuenta la cantidad de información a transmitir por los canales de señalización para poder mantener una buena tasa de transmisión de información.

Como se ha comentado anteriormente, las funciones de planificación y asignación de recursos [7] se llevan a cabo en el eNB, por tanto es quien se encarga de controlar los requerimientos de cada usuario de cada celda que dependa de él y se asegura de que se asignan los recursos suficientes para proporcionar a cada usuario la QoS necesaria.

En cuanto a los algoritmos de scheduling propiamente dichos, existen dos grandes grupos en los que se los podría dividir, dos grupos extremos, por una parte los denominados oportunistas (opportunistic scheduling) [4] y por otra parte los denominados justos (fair scheduling) [4].

Los primeros, están pensados para proporcionar la máxima velocidad a cada usuario, aprovechándose de que usuarios diferentes tienen diferentes condiciones de canal en cada instante de tiempo y por lo tanto tendrán las mejores condiciones de canal en

instantes de tiempo diferentes y para frecuencias diferentes. Aprovechando esta propiedad se puede obtener un aumento considerable en la tasa total de todos los usuarios que transmiten y más aumenta cuanto mayor es el número de usuarios. El principal problema de este tipo de algoritmos es que son “injustos” con algunos usuarios y tienen cierta dificultad para cumplir con la calidad de servicio necesaria. Por ejemplo, hay usuarios que no pueden esperar a que sus condiciones de canal sean lo suficientemente buenas para transmitir, especialmente en canales poco variantes.

El segundo tipo de algoritmos es la familia de los considerados “justos”. Éstos son más eficientes en cuanto a latencia y se centran más en conseguir una tasa mínima para cada usuario antes que en la tasa total como en el caso anterior. Este tipo de características son muy apropiadas para las aplicaciones que necesitan tráfico en tiempo real, como videoconferencias o voz sobre IP, que necesitan tener siempre un mínimo de tasa para poder transmitir, independientemente del estado del canal. Por otra parte, comentar que éstas son aplicaciones especialmente propicias para LTE, considerando sus características de tasas de transmisión y retardos.

Realmente en la práctica, los algoritmos utilizados están entre estos dos grupos y combinan características de los dos para satisfacer las diversas necesidades de calidad de servicio de los usuarios.

En la figura 10, se describe gráficamente el procedimiento de scheduling en LTE:

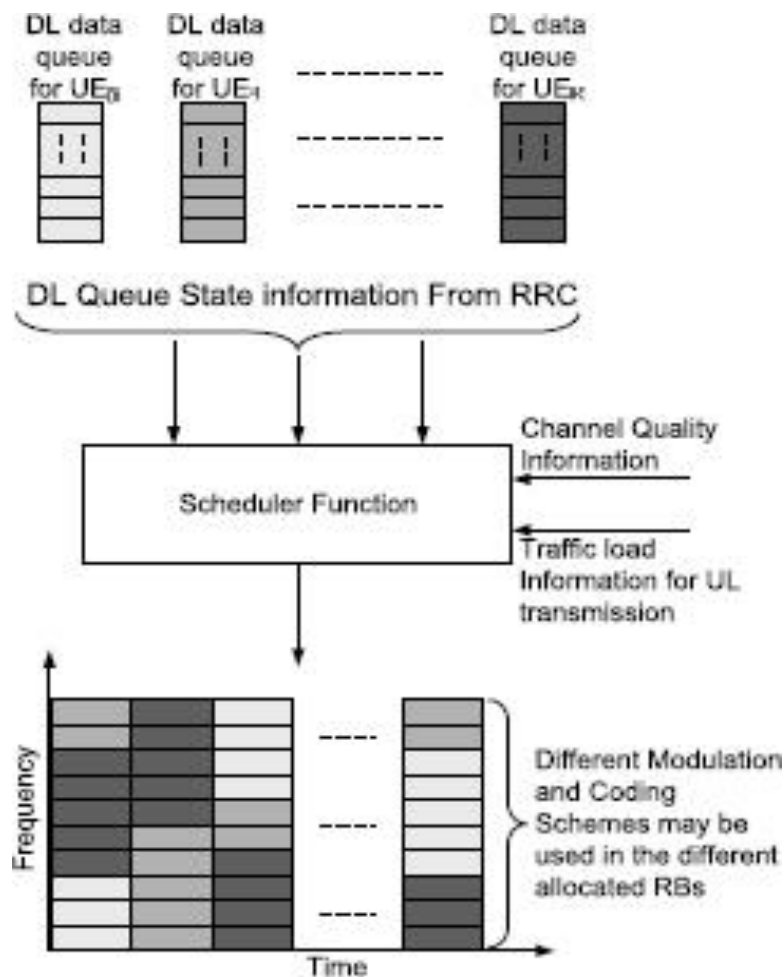


Figura 10: Scheduler LTE

Fuente: [4]

Como se ha comentado con anterioridad, el proceso de scheduling pretende maximizar la capacidad de los usuarios, aplicando para ello una asignación de recursos lo más eficiente posible. Por eso, antes de establecer un algoritmo se debe establecer la métrica de capacidad correspondiente, que después se trata de optimizar con las distintas soluciones disponibles de asignación de recursos, teniendo siempre en cuenta el hecho de cumplir con las restricciones existentes, ya sean restricciones físicas, como el ancho de banda o la potencia, o restricciones referidas a QoS. Algunos ejemplos de estas métricas de capacidad son: la capacidad ergódica o la capacidad con límite de retardo (delay-limited capacity) [4].

La capacidad ergódica o capacidad de Shannon [14], es la máxima tasa de transmisión que se puede enviar por un canal, con una baja probabilidad de error y promediada sobre un desvanecimiento. Esta métrica considera la tasa media a largo plazo que puede asignarse a un usuario, que por otra parte, no tenga restricciones de latencia. En este caso, la potencia de transmisión y la información mutua entre transmisor y receptor varían con los desvanecimientos del canal para tratar de lograr las mejores tasas medias de transmisión.

El segundo ejemplo, delay-limited capacity, se define como la tasa de transmisión que puede mantenerse durante los desvanecimientos cuando existen limitaciones de potencia a largo plazo. Se utiliza en los casos en los que hay una cierta “injusticia” en la asignación de recursos del sistema. Algunas aplicaciones no pueden soportar esta carencia de recursos debido a que tienen ciertas restricciones en los retardos. Esto es debido a la utilización de ciertos algoritmos de scheduling que por su naturaleza son “injustos” con algunos usuarios.

Destacar que, mientras que en la capacidad ergódica la información mutua entre transmisor y receptor va variando conforme varía el canal, en delay-limited capacity se mantiene siempre constante dicha información, incluso durante los desvanecimientos del canal.

Una vez se han establecido las métricas correspondientes de capacidad, se trata de optimizar al máximo esta capacidad utilizando las técnicas adecuadas de asignación de recursos. Para ello en LTE los eNBs utilizan diferentes algoritmos dependiendo del tipo de optimización que se requiera, es decir, de los criterios en los que se base esa optimización.

A continuación nos centraremos en el estudio de los diferentes algoritmos empleados en LTE para la asignación de recursos. Se explicará en que se basa su funcionamiento, en qué casos se utilizan, las principales diferencias entre unos y otros y cómo influye en el rendimiento de la red el hecho de utilizar uno u otro algoritmo.

Como se ha comentado con anterioridad, los algoritmos se dividen en dos grandes grupos, los orientados a maximizar la capacidad de los usuarios y los algoritmos utilizados a minimizar la latencia, es decir, los retardos de los usuarios. Los que se preocupan por la capacidad, utilizan las condiciones de canal para cada usuario en cada

instante. Mientras que los que tratan de minimizar retardos, tienen en cuenta otros factores como la prioridad que puedan tener determinados servicios o el estado de las colas de cada usuario.

5.1 Algoritmos de asignación de recursos.

Principales algoritmos de scheduling utilizados [8] [5] y [4]:

Maximum C/I:

Pertenece al grupo de algoritmos que utilizan la información del estado del canal en cada instante para determinar el usuario al que se le deben asignar los recursos. En este caso se asignan recursos a los usuarios con mejores condiciones de canal, es decir, aquellos que dispongan de una mejor relación Carrier/Interference (C/I) en el momento de la transmisión. Serán, por tanto, los usuarios que dispondrán de mayores velocidades de transmisión. El usuario al que se le asignan recursos viene definido por la siguiente expresión:

$$k^*(n) = \underset{k}{\operatorname{argmax}} \{r_k(n)\}$$

Donde $r_k(n)$, es la tasa de transmisión del usuario k en el instante de tiempo n . Los usuarios con mejores condiciones de canal son los que obtienen más y mejores recursos, en la expresión el usuario $k(n)$. Por el contrario, los usuarios que en ese instante tienen un canal peor reciben menos recursos; están, por tanto, en desventaja respecto a los usuarios con mejores condiciones, por eso se le considera un algoritmo “injusto”. Estas diferencias suelen producirse entre usuarios del borde de la celda y el resto de usuarios, ya que en el borde de la celda, entre otros inconvenientes, hay más interferencias y, por tanto, las condiciones de canal son peores por ejemplo, que las de un usuario de la zona más céntrica de la celda. En definitiva, lo que pretende este algoritmo es obtener una tasa media de transmisión lo más elevada posible. Es decir, que la suma de velocidades de todos los usuarios de la celda sea lo mayor posible. No se preocupa tanto de cumplir con los requerimientos de calidad de servicio de los usuarios. Para lograr éstas tasas

máximas se aprovecha de la diversidad multiusuario provocada por las variaciones de los canales de los distintos usuarios, los cuales varían de forma independiente.

Maximum-Largest weighted Delay First (M-LWDF):

M-LWDF es un algoritmo especialmente pensado para soportar varios usuarios con tráfico en tiempo real. Permite asignar recursos a varios usuarios con diferentes requerimientos de QoS. Para ello, utiliza tanto la información referente a las variaciones instantáneas del canal, como la información relativa a los retardos. Este algoritmo trata de equilibrar los retardos entre los distintos paquetes y trata de aprovechar la información del estado del canal de la forma más eficiente posible.

Viene definido por la siguiente expresión:

$$k^*(n) = \underset{k}{\operatorname{argmax}} \left\{ a_k W_k(n) \frac{r_{k(n)}}{\bar{r}_k} \right\}$$

Donde $a_k = -\log(\delta_k)/D_{ik}$, D_{ik} es el máximo retardo que un usuario puede soportar, siendo δ_k la probabilidad de que ese retardo máximo sea superado. $W_k(n)$ es el Head Of Line (HOL), representa el retardo del paquete. Los usuarios que reciben la asignación de recursos (k^*) son aquellos que maximizan esta expresión sobre el total de usuarios (k). $r_{k(n)}$ representa la tasa correspondiente para el estado del canal del usuario k en el instante de tiempo n . Por último, $\bar{r}_k$ es la tasa media que es capaz de soportar el canal. La principal característica de este algoritmo es que logra el máximo throughput para cada usuario, ya que es capaz de mantener las colas estables. Por último, destacar que siempre trata de asignar mayor prioridad a los usuarios que estén ejecutando servicios en tiempo real, en perjuicio de los usuarios con servicios considerados Best Effort (BE). Es decir, se establece una clasificación entre estos dos tipos de usuarios, ya que los requerimientos para unos u otros servicios son completamente diferentes.

Proportional Fair (PF):

Este algoritmo trata de proporcionar a cada usuario una tasa “justa” respecto a la tasa de transmisión total disponible. Es un tipo de algoritmo adecuado para tráfico en tiempo real, ya que garantiza siempre una tasa mínima de transmisión. Se basa en medidas de la calidad del canal para cada usuario en cada instante de tiempo y asigna recursos al usuario con la mayor relación entre su tasa de transmisión instantánea disponible y su tasa de transmisión real, a lo largo de un intervalo de tiempo definido. Como todos los algoritmos, trata de maximizar la capacidad de la red pero, como se ha comentado anteriormente, garantizando una tasa mínima para cada usuario. Para ello, asigna una prioridad más baja a los usuarios que ya tienen una tasa de transmisión media más alta, de esta forma permite que los usuarios con peores condiciones de canal puedan transmitir, aplicando así una cierta “justicia” característica de este algoritmo.

Responde a la siguiente expresión:

$$k^*(n) = \underset{k}{argmax} \left\{ \frac{r_k(n)}{T_k(n)} \right\}$$

Donde k^* es el usuario óptimo que recibe la asignación entre el conjunto de k usuarios. T_k corresponde a la tasa media del usuario k -ésimo en el time slot n y r_k es la tasa del usuario k en el instante de tiempo n .

Exponential (EXP):

Es un algoritmo muy parecido al M-LWDF, pero la expresión que proporciona el usuario al que se le asignan recursos varía ligeramente:

$$k^*(n) = \underset{k}{argmax} \left\{ a_k \frac{r_k(n)}{\bar{r}_k} \exp \left(\frac{a_k W_k(n) - \overline{aW}}{1 + \sqrt{aW}} \right) \right\}$$

Como se aprecia, los parámetros utilizados son los mismos que para M-LWDF, pero se introduce una función exponencial en la expresión, donde el numerador de dicha

exponencial representa el retardo ponderado. La principal diferencia con M-LWDF es que es más “justo” con los usuarios de servicios considerados BE, debido a que el peso de la exponencial solo es mayor que el término que depende del estado del canal, $\frac{r_k(n)}{\bar{r}_k}$, cuando el retardo del paquete, $W_k(n)$, es mucho mayor que el retardo medio ponderado, $\overline{aW}$.

En ocasiones, este algoritmo se combina con PF, se denomina EXP/PF y está pensado para permitir que un usuario pueda utilizar tanto servicios en tiempo real como servicios Non Real Time (NRT), aplicando, para ello, una mayor prioridad a los servicios en tiempo real que al resto. Cuando los retardos, $W_k(n)$, de los usuarios son muy parecidos entre sí, la exponencial es prácticamente 1 y en este caso el algoritmo exponencial se comporta como un algoritmo PF. Si, por el contrario, el retardo de alguno de los usuarios es mucho mayor que el resto, el término de la exponencial prácticamente anula el término referente al estado del canal y el usuario obtiene cierta prioridad sobre el resto.

Round Robin:

Es uno de los algoritmos más básicos. Divide el canal por igual entre los usuarios, es decir, asigna de manera alternativa el mismo número de recursos a cada usuario. Sigue una secuencia de asignación cíclica, sin tener en cuenta las condiciones del canal o las posibles prioridades entre usuarios. Por lo que, por una parte, podría considerarse un algoritmo justo, ya que asigna el mismo número de recursos a cada usuario. Pero al mismo tiempo, los usuarios con peores condiciones de canal dispondrán de menores tasas de transmisión, en igualdad de recursos asignados, que los usuarios que dispongan de mejores condiciones en ese momento. No es un algoritmo que suele utilizarse en la práctica. Su uso es más como referencia a la hora de evaluar otros algoritmos. La probabilidad de que un usuario sea asignado es:

$$p_k = \frac{1}{N}$$

Donde N, es el número de usuarios totales a tener en cuenta en la asignación. Como se ha comentado antes y se deduce de la expresión, es un algoritmo justo, aunque el

principal problema es que no tiene en cuenta las condiciones de canal, motivo por el cual no suele utilizarse en la práctica.

Max-Min:

Este algoritmo, asigna recursos en primer lugar a los usuarios que menos demanda de recursos tienen. Es decir, trata de lograr que el valor mínimo de la tasa de transmisión sea lo mayor posible. El procedimiento que sigue es sencillo, como he comentado. En primer lugar, asigna recursos al usuario que menos demande. Una vez asignados estos recursos, continúa asignando los recursos restantes a los demás usuarios siguiendo la misma regla; es decir, dando prioridad de asignación siempre a los usuarios que menos recursos demanden. En otras palabras, este algoritmo se basa en garantizar que en cada instante de tiempo el valor mínimo de la ganancia de canal de los usuarios sea lo mayor posible.

Sigue la siguiente regla de asignación:

$$R = \frac{\text{Recursos totales disponibles} - \sum \text{Recursos asignados}}{\sum \text{Usuarios insatisfechos}}$$

Donde R es la tasa de transmisión resultante para cada usuario. De esta forma siempre se garantiza una tasa de transmisión mínima para cada usuario.

Existe una variante denominada **Max-Min Weighted**, cuyo principio de asignación es el mismo, pero teniendo en cuenta ciertos pesos para cada usuario. Igual que en el caso anterior, se asignan recursos, en primer lugar, a los usuarios que menos demandan, pero normalizados respecto a sus respectivos pesos. Ningún usuario recibe más recursos de los que demanda y los usuarios insatisfechos comparten los recursos restantes en proporción a sus pesos.

En definitiva, como se observa, se dispone de diversos algoritmos de scheduling para llevar a cabo la asignación de recursos. Se utilizarán unos u otros dependiendo de los criterios utilizados para establecer las prioridades, es decir, según se quiera dar más o

menos prioridad a un tipo de servicio u otro de los requeridos por los usuarios. Por ejemplo, cuando se pretende priorizar el tráfico de datos, se suelen tener en cuenta las distintas clases de tráfico, especialmente de cara a lograr maximizar el throughput, tanto para aplicaciones sensibles al retardo, en las que es necesaria una cierta QoS, como para las aplicaciones que tienen una cierta tolerancia al retardo. Se pretende conseguir así, repartir la tasa disponible entre estos dos tipos de aplicaciones de manera justa y, sobre todo, en función de las necesidades de cada tipo.

Por otra parte, es importante señalar la importancia de la diversidad multiusuario a la hora de aplicar cualquiera de los algoritmos anteriores en LTE. Ésta se basa en seleccionar, entre todos los candidatos, al usuario con mejores condiciones para la transmisión. Por eso es especialmente importante cuanto mayor sea la densidad de usuarios, en cuyo caso la ganancia por diversidad multiusuario permite lograr una gran capacidad, incluso con restricciones de retardo muy exigentes.

Por último, es importante destacar que no se pueden considerar de forma aislada las celdas o los eNBs, ya que para lograr optimizar el sistema se requiere un cierto grado de coordinación entre celdas y eNB para evitar la interferencia entre celdas, ya que dicha interferencia puede convertirse en un factor limitante para el sistema. De esta forma, cuando se considera el sistema completo, se suelen obtener los mejores resultados simplemente con una asignación “on-off” de bloques de recursos. Esto consiste en que el eNB no transmite ciertos bloques de recursos usados por los nodos vecinos para los usuarios que se encuentran en el borde de la celda.

5.2 Estudio comparativo de asignación de recursos.

Para complementar la información expuesta anteriormente en este trabajo, se han realizado una serie de simulaciones de algunos de los algoritmos estudiados en el punto anterior. Para ello, se ha utilizado el programa **Matlab**, ya que ha sido el principal programa utilizado durante la carrera para la realización de diversas prácticas y simulaciones, motivo por el cual ya se ha adquirido una cierta facilidad en el manejo del mismo y me permite realizar estas simulaciones de una manera más sencilla. En

concreto, se ha simulado el proceso de asignación de recursos en el DL, utilizando los algoritmos *PF* y *Maximum C/I*.

El principal motivo de la elección de estos dos algoritmos, no es otro que el hecho de ser dos algoritmos con principios de funcionamiento muy diferentes. Se pretende así, mostrar claramente las diferencias en los resultados obtenidos al utilizar uno u otro algoritmo, ya que al ser algoritmos muy diferentes, las diferencias son mayores y se pueden observar mejor los efectos de la utilización de uno u otro esquema en el rendimiento del sistema.

Otro motivo por el cual me he decantado por hacer una comparativa entre estos dos algoritmos, es que son algoritmos que tienen utilidades muy diferentes, según el tipo de servicio que se desee proporcionar al usuario o el tipo de servicio al que se le quiere dar una cierta prioridad a la hora de realizar la asignación. Es decir, si por ejemplo se quiere dar cierta prioridad a los usuarios con mejores condiciones de canal se utilizará *Maximum C/I*, si, por el contrario, se pretende establecer una cierta igualdad entre los distintos usuarios y garantizar siempre una tasa mínima para cada usuario, se utilizará *PF*. Podría decirse que son algoritmos pertenecientes a extremos opuestos, en lo que a objetivos se refiere.

Posteriormente, se propondrán algunas posibles mejoras que permitan obtener resultados más eficientes; es decir, lo que podríamos traducir como un mejor rendimiento del sistema.

Para la realización de estas simulaciones, se generan aleatoriamente una serie de requerimientos para cada usuario. Del mismo modo, se generan de manera aleatoria lo que serían las muestras tomadas del canal, relaciones S/N. Se consideran 40 usuarios en la celda y se realizan 50 iteraciones del proceso, es decir, 50 asignaciones. El parámetro variable es el que define la capacidad, se define en términos de RBGs, los cuales tienen su equivalente en ancho de banda definido en número de RBs o en Hz. Esto permite observar el comportamiento de los algoritmos ante estas variaciones de capacidad y ver las consecuencias tanto, en las tasas asignadas a los usuarios, como en el número de usuarios satisfechos, es decir, aquellos a los que se les asignan tasas iguales o superiores a las que demandan.

5.2.1 Estudio con máxima capacidad

5.2.1.1 Asignación a varios usuarios

En esta simulación, se comparan los recursos asignados por cada algoritmo, en términos de velocidad asignada a cada usuario. Se muestran las tasas asignadas a cada uno de los 40 usuarios en una de las 50 asignaciones que se realizan. Se utiliza una capacidad de 4 RBGs, lo que equivaldría a un canal estándar de LTE de 20MHz.

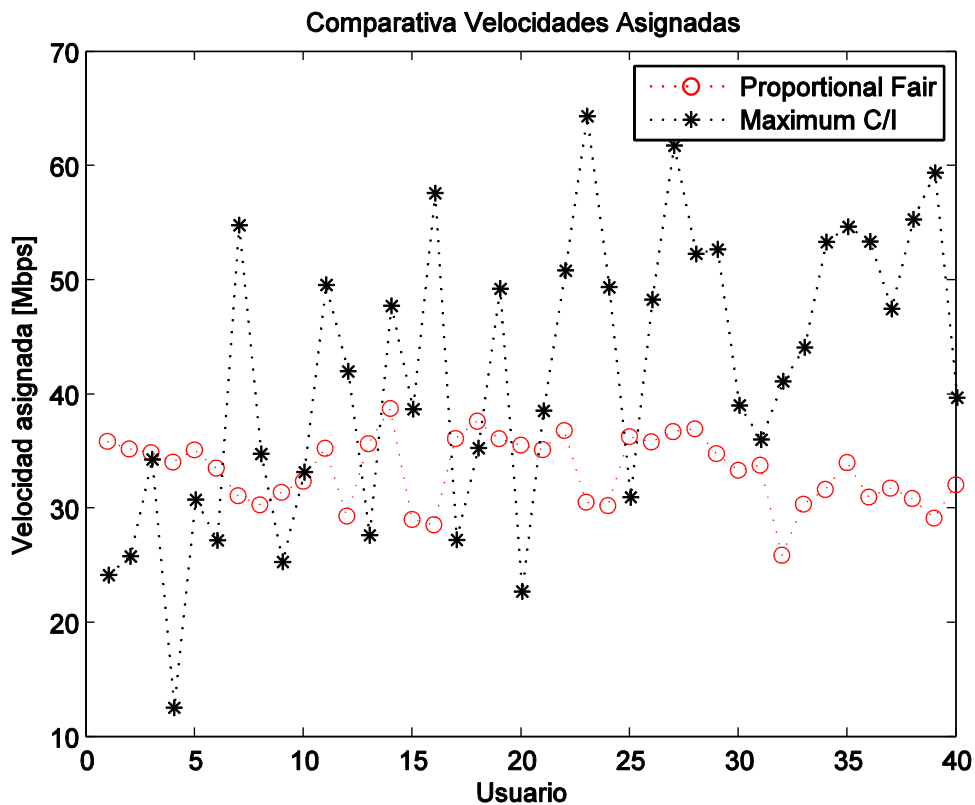


Figura 11: Tasas asignadas a los distintos usuarios con máxima capacidad.

Como se observa, se obtienen tasas muy elevadas de transmisión para cada usuario. Esto se debe a que se está utilizando el máximo ancho de banda, 20 MHz, 4 RBGs.

En la figura 11, se observan claramente las diferencias en cuanto a las tasas que proporciona cada algoritmo a cada usuario. Vemos como hay casos en los que Maximum C/I asigna tasas mucho mayores que PF para un mismo usuario, pero también casos en los que ocurre lo contrario. Para los casos en los que se obtienen mejores tasas con Maximum C/I, los usuarios que obtienen dichas tasas de transmisión,

son aquellos que tienen buenas condiciones de canal en el momento de la asignación. Por el contrario, aquellos para los que el algoritmo ha considerado una mala condición de canal, obtienen tasas bajas, que en ocasiones se aproximan a las proporcionadas por PF o incluso están por debajo de las mismas. Por el contrario, las tasas asignadas por PF siguen un patrón más o menos constante para todos los usuarios, sin “picos” que destaquen sobre el resto de velocidades. Trata de establecer una cierta igualdad entre los recursos asignados a los usuarios, intenta ser “justo”. Estas diferencias se deben a los distintos principios de funcionamiento de cada algoritmo, explicados en el apartado anterior. Se utilizará uno u otro dependiendo de los intereses que se tengan en la transmisión. Si por ejemplo, para este caso concreto, quisiéramos que el usuario que transmite, lo haga a la máxima velocidad posible, utilizaríamos Maximum C/I. Si por el contrario, lo que se quiere es garantizar una velocidad mínima para los usuarios y que todos dispongan de tasas similares para transmitir, se utilizará PF.

5.2.1.2 Asignación a un único usuario

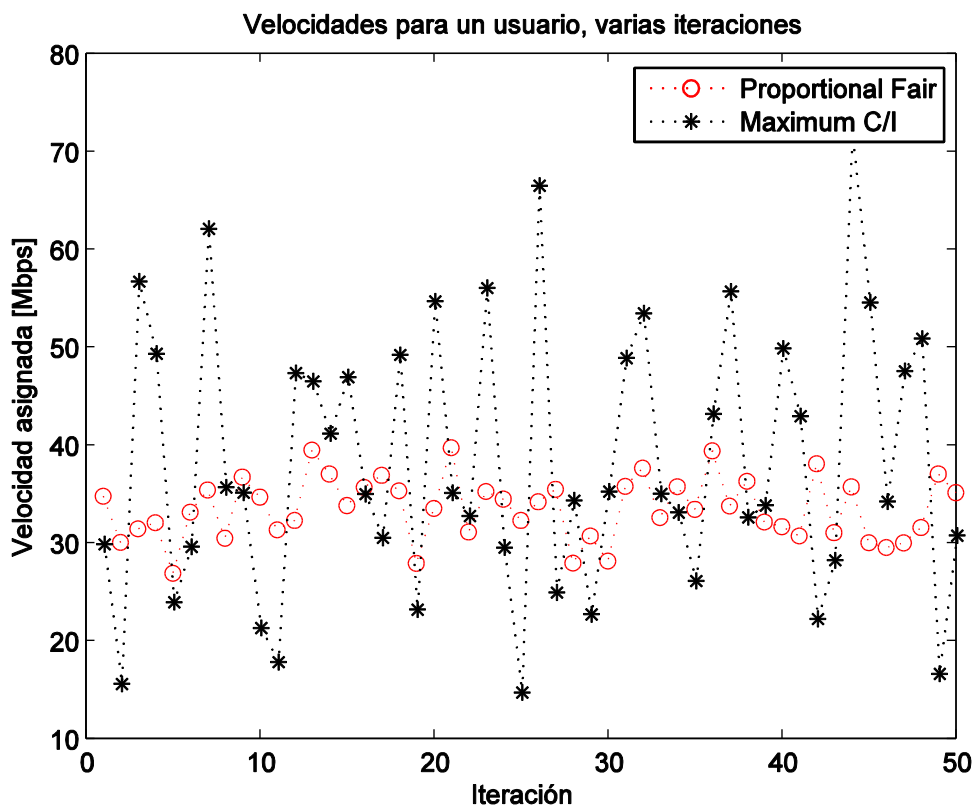


Figura 12. Tasas asignadas a un mismo usuario durante varias iteraciones con máxima capacidad.

En esta segunda simulación se mantienen las mismas condiciones, en cuanto a capacidad, número de usuarios y requerimientos de los mismos que en la anterior. En la figura 12 se muestran las tasas asignadas para un solo usuario a lo largo de las 50 asignaciones que se realizan. Al igual que en el caso anterior, se aprecian claramente las diferencias en las tasas asignadas por cada algoritmo. Vemos la variación, en los flujos asignados, que puede sufrir un mismo usuario a lo largo de los diferentes procesos de asignación que se realizan, dependiendo del algoritmo utilizado. De nuevo, se observa que PF sigue una tendencia más homogénea que Maximum C/I, es decir, hay menos variaciones entre las tasas asignadas en cada una de las iteraciones para el usuario. Por el contrario, Maximum C/I proporciona tasas muy elevadas en varias iteraciones que no se alcanzan con PF, pero al mismo tiempo no da esa cierta seguridad de tener una tasa más o menos constante garantizada como PF, ya que con Maximum C/I puedes tener una tasa máxima en una iteración y en la siguiente tener una mucho menor o incluso no tener recursos para transmitir. Por eso es tan importante la elección de unos tipos de algoritmos u otros, en función del servicio que se necesite proporcionar al usuario.

5.2.1.3 Usuarios satisfechos

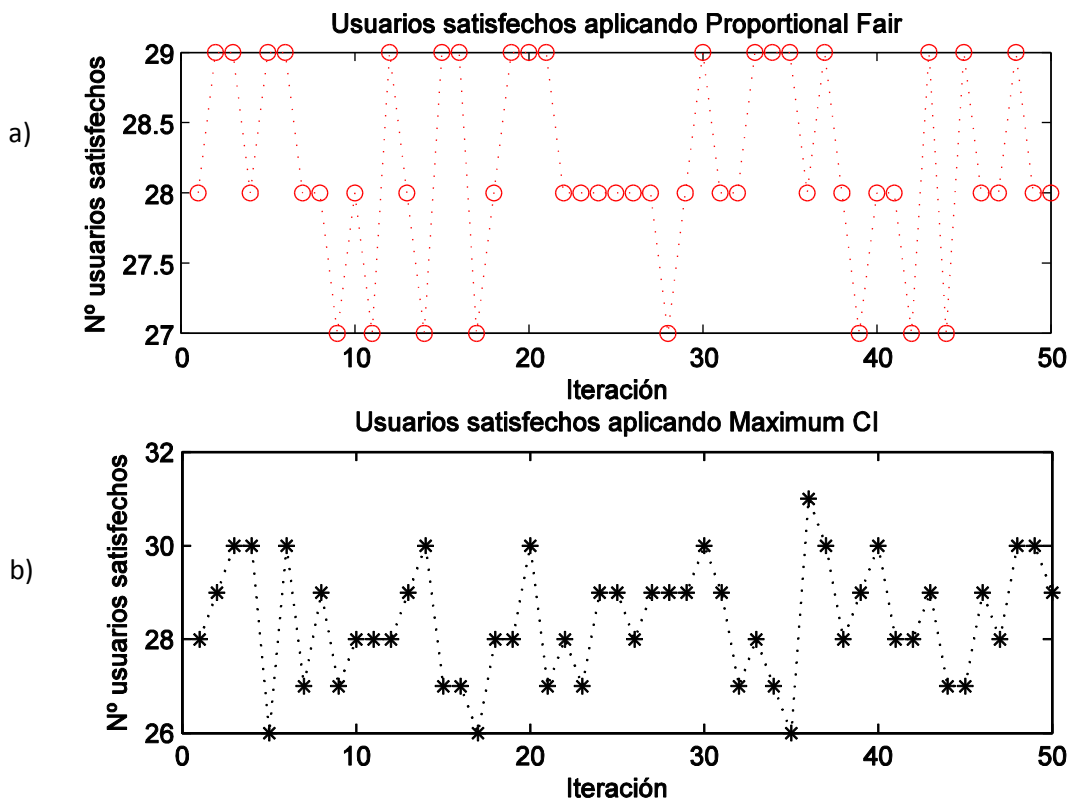


Figura 13: Usuarios satisfechos con cada algoritmo con máxima capacidad.

En esta tercera simulación, se muestra el número de usuarios satisfechos para cada algoritmo, en cada iteración.

Se aprecia cómo, igual que ocurriera con las tasas asignadas, hay iteraciones en las que se obtienen más usuarios satisfechos para PF que para Maximum C/I y viceversa. En este caso, en cada iteración los requerimientos de los usuarios son los mismos para los dos algoritmos, es decir, ambos tienen que cubrir las mismas demandas por parte de los usuarios para el mismo instante de tiempo. Por lo que las diferencias en el número de usuarios satisfechos se deben exclusivamente al funcionamiento de los algoritmos y la forma que tiene cada uno de gestionar los recursos disponibles. Así, por ejemplo, para una misma iteración, se tienen diferencias importantes en el número de usuarios al utilizar uno u otro algoritmo. Que por otra parte, no tiene por qué ser siempre el mismo; es decir, PF puede proporcionarte un mayor número de usuarios satisfechos en la iteración n , mientras que en la iteración $n+1$ ocurre al contrario y es Maximum C/I el que te proporciona una mayor eficiencia, en cuanto a usuarios satisfechos se refiere. Por ejemplo, es lo que ocurre en las iteraciones 35 y 36 de figura 13, en la iteración 35 PF proporciona 29 usuarios satisfechos mientras que Maximum C/I obtiene 26, en la siguiente iteración es PF el que obtiene 28 usuarios frente a los 31 de Maximum C/I.

5.2.2 Estudio con capacidad media

A continuación, se muestran los resultados de las mismas simulaciones que en el caso anterior pero reduciendo la capacidad a 3RBGs, en lugar de los 4 utilizados anteriormente.

5.2.2.1 Asignación a varios usuarios

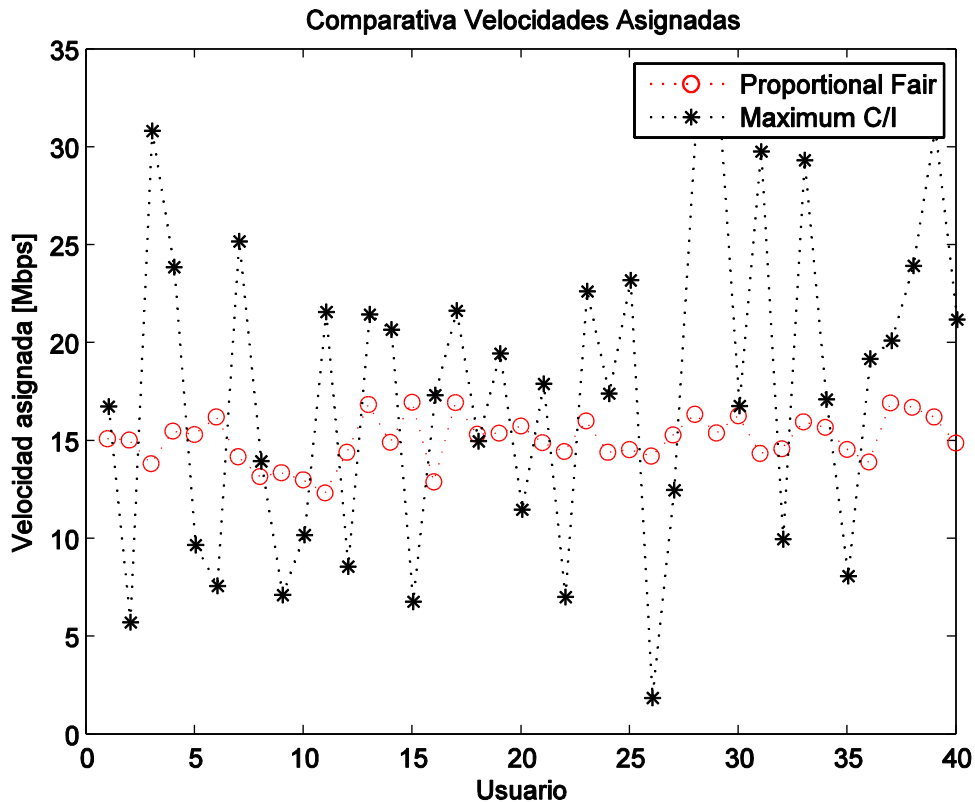


Figura 14: Tasas asignadas a los distintos usuarios con capacidad media.

En la figura 14 se muestran de nuevo, las tasas asignadas para cada usuario en una de las 50 iteraciones realizadas. Se observan claramente los efectos provocados por la reducción del ancho de banda, de 4 a 3 RBGs. Como consecuencia, las tasas se reducen notablemente respecto al caso anterior. Esto es debido, como se ha explicado en el punto 4.3, a que se asignan menos recursos que en el caso anterior por cada usuario, es decir, mientras que con 4 RBGs se asignan entre 64 y 110 RBs por usuario, con 3 RBGs se pueden asignar entre 27 y 63 RBs por usuario. Así ocurriría a medida que fuéramos disminuyendo el número de RBGs a 2 ó incluso 1, las tasas de los usuarios disminuirían considerablemente y por tanto, la eficiencia del sistema.

En cuanto al comportamiento de los algoritmos por separado, no se observan cambios significativos, es decir, se aprecian las mismas diferencias entre ambos algoritmos que en caso anterior; eso sí, las tasas asignadas son menores para ambos algoritmos, debido a la reducción del ancho de banda.

5.2.2.2 Asignación a un único usuario

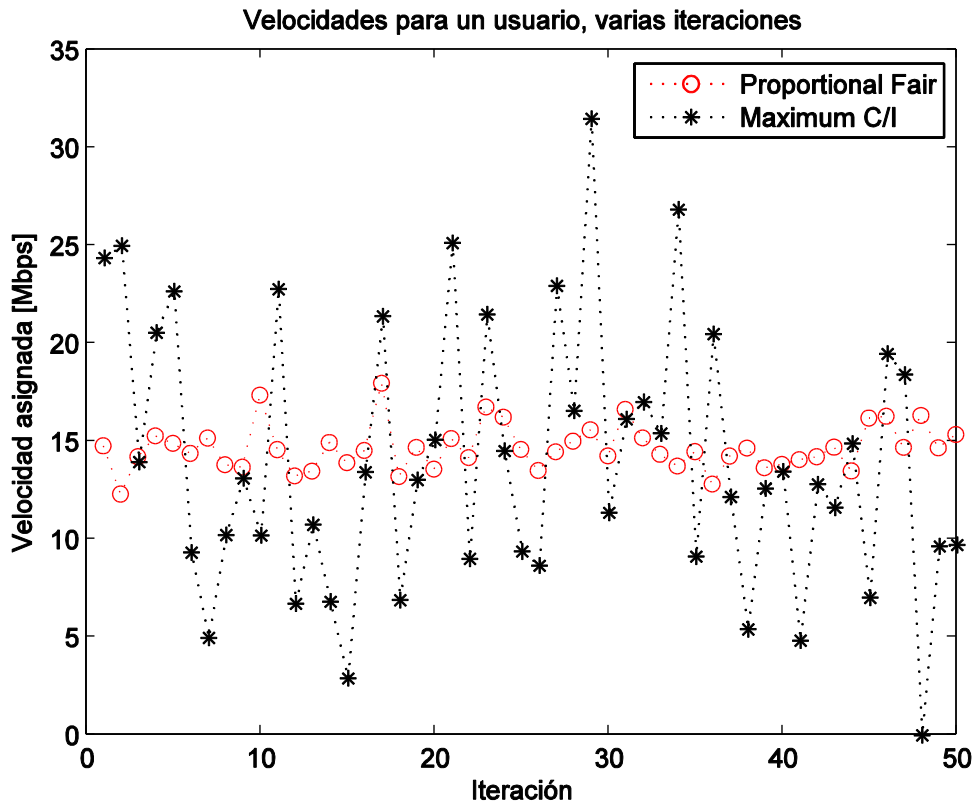


Figura 15: Tasas asignadas a un mismo usuario durante varias iteraciones con capacidad media.

En figura 15 se muestran de nuevo, las diferentes tasas asignadas a un mismo usuario a lo largo de las 50 iteraciones realizadas. Como se puede apreciar en la figura 15, no hay cambios significativos en el comportamiento de los algoritmos. El principal cambio respecto al caso anterior, es la reducción de las tasas debida al paso de 4 a 3 RBGs en la asignación, por la que se ven afectados ambos algoritmos. Un detalle importante a destacar en esta simulación es que se puede observar como este usuario, para una de las iteraciones, no recibe recursos al utilizar Maximum C/I. Es decir, no podría transmitir. Vemos que con PF esto no ocurre en ninguna de las iteraciones, esto podrá suponer un mayor o menor inconveniente en función de los servicios o aplicaciones que esté utilizando el usuario en ese determinado instante. Por ejemplo, en caso de una aplicación en tiempo real, en la que los datos tienen unas restricciones temporales muy estrictas y cualquier información recibida con un retardo superior al establecido por dichas restricciones no es válida para el usuario. Es el caso, por ejemplo, de una videoconferencia o una emisión en directo. En este tipo de aplicaciones, el hecho de que

el usuario no pueda transmitir supone un mayor problema que en el caso de que fuera un servicio NRT.

5.2.2.3 Usuarios satisfechos

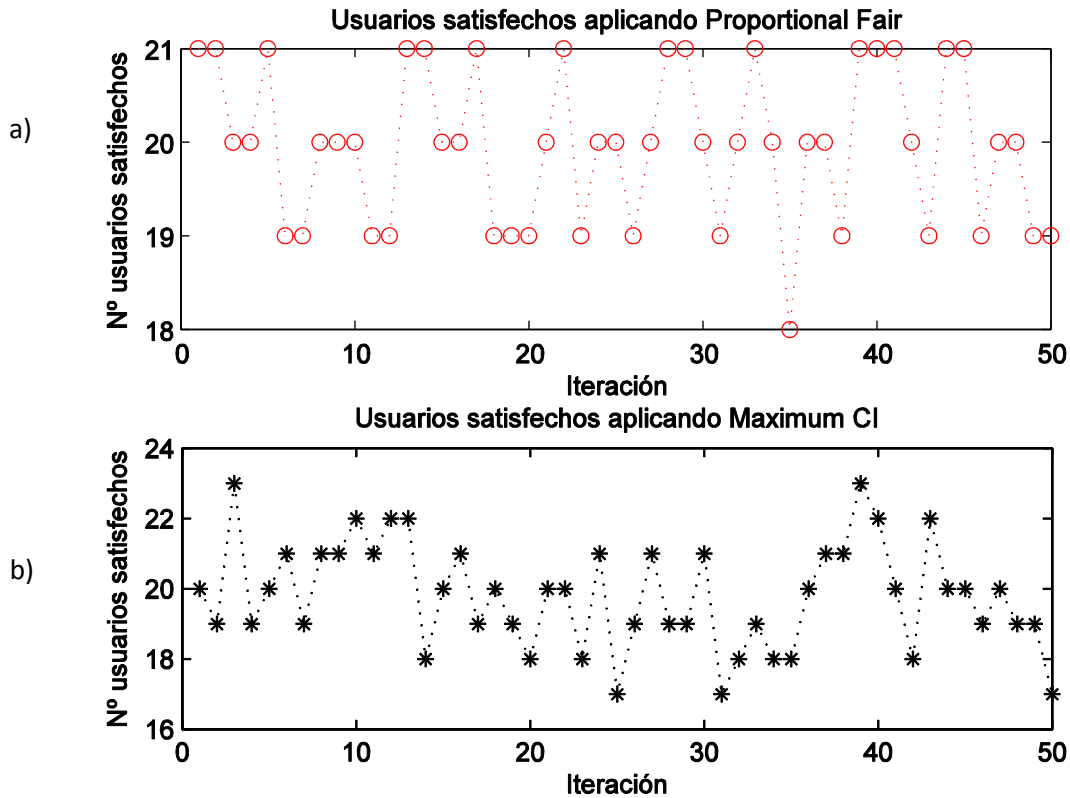


Figura 16: Usuarios satisfechos con cada algoritmo con capacidad media.

Por último, en la figura 16 se muestran los usuarios satisfechos para cada algoritmo. Como era de esperar, el número de usuarios satisfechos disminuye respecto al caso anterior para ambos algoritmos, esto se debe a que al reducir el ancho de banda, se dispone de menos recursos para asignar a los usuarios, lo que se traduce en peores tasas para cada usuario y, por tanto, en un peor rendimiento para el sistema. De esta forma se muestra, también, la importancia de contar con un buen ancho de banda. De lo contrario, podremos utilizar un algoritmo de asignación u otro, que siempre se verá afectado por esta limitación y, por muy eficiente que resulte el algoritmo para nuestro sistema, nunca se obtendrá el máximo beneficio de él.

5.3 *Propuesta de mejoras.*

En vista de los resultados obtenidos en estas simulaciones, me gustaría comentar algunas posibles mejoras que podrían aplicarse en el sistema para obtener una mayor eficiencia del mismo.

La principal estrategia que me gustaría comentar, como posible mejora a lo que se ha comentado hasta ahora, es el hecho de utilizar varios de estos algoritmos de manera simultánea; es decir, no utilizar únicamente un algoritmo de asignación, si no varios en función del tipo de tráfico o las condiciones del canal en cada instante de tiempo. Como hemos podido observar con los algoritmos de las simulaciones, no hay uno mejor o peor que el otro en términos generales. Hay iteraciones, en las que un algoritmo proporciona mejores tasas que el otro para esa misma asignación, o uno logrará un mayor número de usuarios que el otro para la misma iteración y en igualdad de condiciones, mientras que en la iteración siguiente puede ocurrir lo contrario y el algoritmo que en la iteración anterior proporcionaba peores resultados, en el instante actual se comporte mejor que el otro. Esto puede ocurrir fácilmente, por ejemplo, como hemos visto en las simulaciones, con los algoritmos que basen su asignación en condiciones del canal, las cuales pueden variar con cierta facilidad y con ellas los recursos asignados al usuario. Estamos de acuerdo en que hay algoritmos más indicados que otros en función de las limitaciones que se le impongan al sistema, tipo de tráfico, QoS, o determinadas condiciones de canal. Pero si combinamos varios de estos algoritmos “específicos” para una determinada limitación, lograremos una mejora notable, no solo en las tasas asignadas a cada usuario, sino en la tasa media del conjunto de todos los usuarios.

Volviendo al caso de las simulaciones, consiste en combinar ambos algoritmos, PF y Maximum C/I. El resultado que se obtiene es, para cada iteración, la mejor entre las dos tasas asignadas por los algoritmos. Es decir, si en la iteración n , de las dos tasas asignadas por cada algoritmo la mejor es la de PF, se asigna esa tasa al usuario y de igual modo si la mejor tasa es la de Maximum C/I. El resultado es, un conjunto formado por las mejores tasas para cada usuario, en este caso, la mejor entre las asignadas por estos dos algoritmos en cada iteración. De esta forma, siempre se obtiene la tasa óptima por usuario e iteración. Esta estrategia tiene el inconveniente de que probablemente, pueda resultar más compleja que el hecho de implementar sólo un algoritmo para la asignación, pero a cambio se obtienen mejores tasas de transmisión. Hay que estudiar,

en cada caso, que resulta más beneficioso para el sistema, es decir, si a pesar de aumentar la complejidad a la hora de asignar los recursos, las tasas asignadas son tales, que merece la pena incurrir en dicho aumento de complejidad. Todo dependerá de las condiciones particulares para cada escenario en el que se vaya a desplegar el sistema y si resulta o no más eficiente el hecho de aplicar este esquema de asignación.

A continuación, se muestran los resultados correspondientes a las simulaciones realizadas aplicando la estrategia anteriormente comentada.

5.3.1 Asignación a varios usuarios

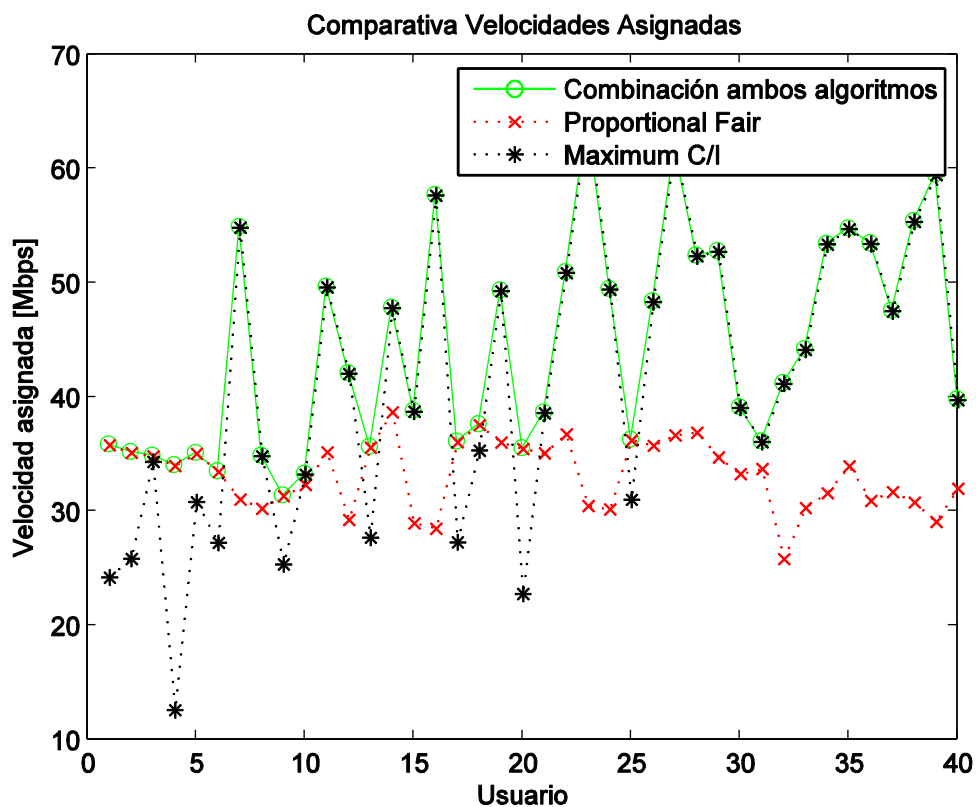


Figura 17: Tasas asignadas a los distintos usuarios aplicando las mejoras propuestas.

En la figura 17 se aprecian, en verde, las tasas asignadas a cada usuario aplicando la combinación de los dos algoritmos. Vemos como de esta forma el usuario siempre obtiene la mayor tasa entre las dos posibles.

5.3.2 Asignación a un único usuario

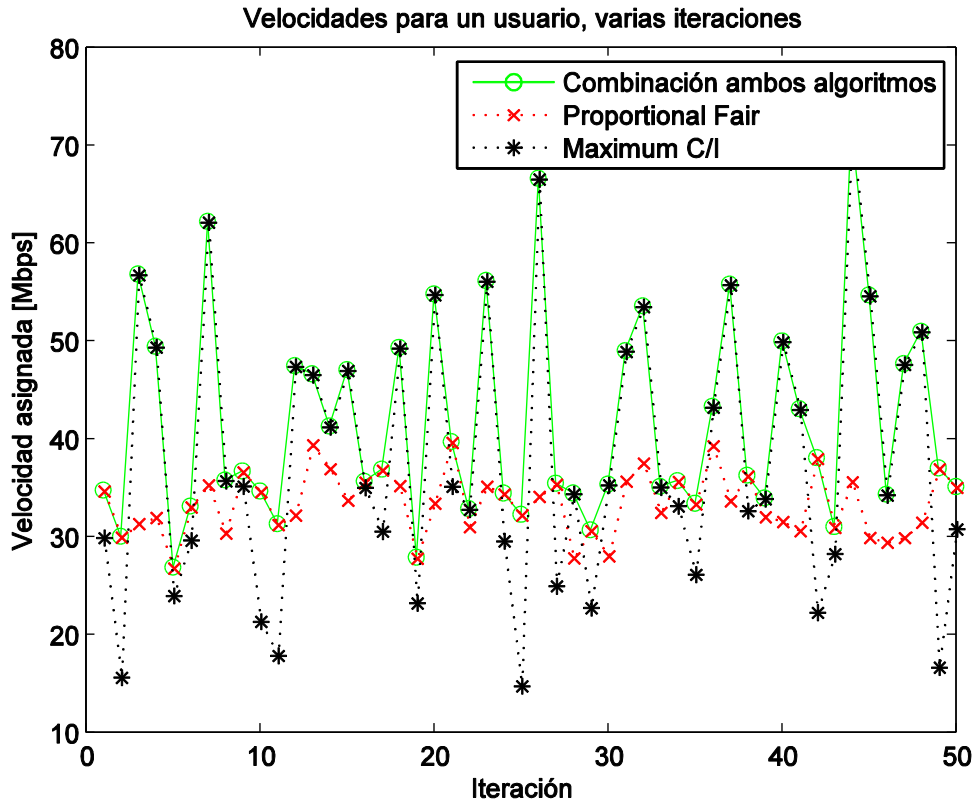


Figura 18: Tasas asignadas a un mismo usuario durante varias iteraciones aplicando las mejoras propuestas.

En este caso, siguiendo el mismo orden que en las simulaciones anteriores, tenemos las tasas asignadas a un único usuario a lo largo de las 50 iteraciones realizadas. Igual que en el caso anterior, observamos cómo a este usuario siempre se le asigna, en verde, la mejor tasa disponible para él en cada momento.

5.3.3 Usuarios satisfechos

En la figura 19, se muestran los usuarios satisfechos para cada algoritmo comparándolos con el resultado de combinar ambos algoritmos. Una vez más, se observa, en color verde, como hay iteraciones en las que se obtendría un mayor número de usuarios satisfechos al utilizar ambos algoritmos frente a utilizar uno u otro de forma individual. De esta forma, se obtiene el mayor número posible de usuarios satisfechos en cada iteración, es decir, el mayor rendimiento de nuestro sistema en estos términos.

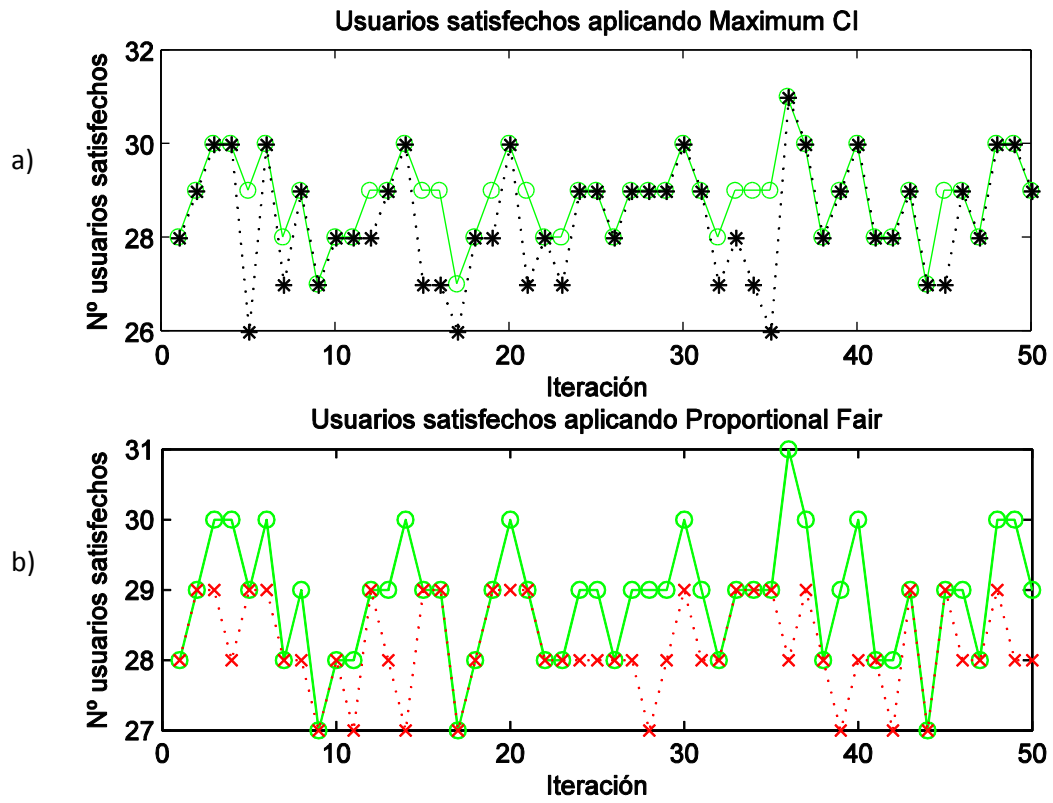


Figura 19: Usuarios satisfechos con cada algoritmo aplicando las mejoras propuestas

6. Conclusiones

Me gustaría cerrar este trabajo destacando las principales ideas que se recogen tras los estudios realizados, las ideas más importantes obtenidas tras su realización.

Se pretende demostrar con resultados reales, la influencia que puede tener en el sistema la utilización de uno u otro algoritmo. Este es el principal objetivo de las simulaciones, no tanto el ver el funcionamiento de un algoritmo u otro en concreto, que por otra parte también se puede apreciar claramente, si no poder apreciar las diferencias en las tasas asignadas entre ellos, sean cuales sean estos algoritmos. En mi caso, he elegido estos dos en concreto, PF y Maximun C/I, porque son dos algoritmos con principios de funcionamiento muy diferentes y permiten observar mejor las diferencias en los resultados. Por otra parte, en vista de los resultados, podríamos llegar a la conclusión de que no existen algoritmos mejores o más eficientes que otros, hablando en términos generales. Se debe juzgar la eficiencia del algoritmo en base al escenario en el que se este utilizando y del tipo de tráfico o servicio para el que se esté utilizando: algoritmos muy eficientes en determinados escenarios pueden resultar nefastos para otros.

Destacar, por otra parte, los resultados obtenidos en las simulaciones tras aplicar las mejoras propuestas. Se demuestra como combinando las propiedades de dos tipos de algoritmos distintos se pueden obtener mejores resultados, lo que se traduce en un sistema más eficiente. Estas mejoras permiten asignar mejores tasas de transmisión a los usuarios y en consecuencia se obtiene un mayor número de usuarios satisfechos.

En cuanto a mi experiencia personal, este trabajo me ha ayudado a aclarar conceptos referentes a los diferentes sistemas de comunicaciones móviles y a profundizar, en concreto, en el sistema LTE. Lo cual, considero de gran importancia debido a la relevancia que tiene y tendrá este sistema para las telecomunicaciones y, en concreto, en lo que a las comunicaciones móviles se refiere, ya que es hacia donde se orienta principalmente este mercado en la actualidad.

7. Acrónimos

1G.....	Primera Generación.
2G.....	Segunda Generación.
3G.....	Tercera Generación.
3GPP.....	Third Generation Partnership Project
4G.....	Cuarta Generación.
AMC.....	Adaptive Modulation and Coding
AMPS.....	Advanced Mobile Phone Service
ATB.....	Adaptive Transmission Bandwidth
BE.....	Best Effort
BSC.....	Base Station Controller
BSR.....	Buffer Status Report
BTS.....	Base Transceiver Station
CDMA.....	Code Division Multiple Access
CEPT.....	Conference of European Postal and Telecommunications
CQI.....	Channel Quality Indicator
DL.....	Down Link
DRX/DTX.....	Discontinuous Reception and Transmission
DS-CDMA.....	Direct Sequence Code Division Multiple Access
EDGE.....	Enhanced Data Rates for Global Evolution
eNB.....	evolved-NodeB
EPC.....	Evolved-Packet Core
EPS.....	Evolved Packet System
ETSI.....	European Telecommunication Standards Institute
E-UTRAN.....	Evolved-UMTS Terrestrial Radio Access Network

FDD.....	Frequency Division Duplex
FDMA.....	Frequency Division Multiple Access
FDPS.....	Frequency Domain Packet Scheduling
FTB.....	Fixed Transmission bandwidth
GGSN.....	Gateway GPRS Support Node
GPRS.....	General Packet Radio Services
GSM.....	Groupe Speciale Mobile
HARQ.....	Hybrid Automatic Repeat reQuest
HSDPA.....	High Speed Downlink Packet Access
HSPA.....	High Speed Packet Access
HSUPA.....	High Speed Uplink Packet Access
IEEE.....	Institute of Electrical and Electronics Engineers
IMT.....	International Mobile Telecommunications
ITU.....	International Telecommunications Union
JDC.....	Japan Digital Cellular
LTE.....	Long Term Evolution.
MIMO.....	Multiple Input Multiple Output
MME.....	Mobility Management Entity
NMT.....	Nordic Mobile Telephony
NRT.....	Non Real Time
NTT.....	Nippon Telegraph and Telephone
OFDM.....	Orthogonal Frequency Division Multiplexing
OFDMA.....	Orthogonal Frequency Division Multiple Access
PA.....	Power Amplifier
PAPR.....	Peak to Average Power Ratio
P-C.....	Power Control
PCU.....	Packet Control Unit

P-GW.....	Packet Data Network Gateway
QoS.....	Quality of Service
RB.....	Resource Blocks
RBG.....	Resource Block Group
RF.....	Radio Frequency
RNC.....	Radio Network Controller
RRM.....	Radio Resource Management
SC-FDMA.....	Single Carrier Frequency Division Multiple Access
SCM.....	Spatial Channel Model
SGSN.....	Serving GPRS Support Node
S-GW.....	Serving Gateway
SINR.....	Signal to Interference plus Noise Ratio
TACS.....	Total Access Communications System
TD-CDMA.....	Time Division CDMA
TDD.....	Time Division Duplex
TDMA.....	Time Division Multiple Access
TDPS.....	Domain Packet Scheduling
TTI.....	Transmission Time Interval
UE.....	User Equipment
UL.....	Up Link
UMB.....	Ultra Mobile Broadband
UMTS.....	Universal Mobile Telecommunications System
W-CDMA.....	Wideband CDMA

8. Referencias

- [1] Farley Tom. Mobile Telephone History. Telektronikk Volumes. 03/04/2005.
Page(s) 22-34.
Link: http://www.telektronikk.com/volumes/pdf/3_4.2005/Page_022-034.pdf
- [2] Hernando Rábanos, José María. Comunicaciones Móviles. Madrid: Centro de estudios Ramón Areces. (2004).
- [3] Gorricho, Mónica, & Gorricho, Juan Luis. Comunicaciones Móviles. Barcelona: Ediciones UPC. (2002).
- [4] Sesia, Stefania, Toufic, Issam, Baker, Mateo. LTE: UMTS Long Term Evolution: From Theory To Practice. John Wiley & Sons Ltd. (2011) 2ªed.
- [5] Iturralde, M. ; Ali Yahiya, T. ; Wei, Anne ; Beylot, A.-L. “Performance study of multimedia services using virtual token mechanism for resource allocation in LTE networks”. Vehicular technology conference (VTC Fall). (2011).Page(s): 1-5
- [6] Adaptive signal processing and information theory research group. Drexel University. “Resource allocation in LTE”. Electrical and computer engineering at Drexel. Gwanmo Ku. November (2011).
Link: <http://www.ece.drexel.edu/walsh/Gwanmo-Nov11-2.pdf>
- [7] Kan Zheng ; Fanglong Hu ; Wenbo Wang ; Wei Xiang ; Dohler, M. “Radio resource allocation in LTE-Advanced cellular networks with M2M communications”. Communications magazine. IEEE.
Volume: 50, Issue: 7. (2012). Page(s): 184 – 192
- [8] Narcís Cardona, Juan José Olmos, Mario García, José F. Monserrat. 3GPP LTE: Hacia La 4G Móvil. Marcombo universitaria. (2011).

- [9] A.Simonsson, “Frequency reuse and intercell interference co-ordination in E-UTRA”. Vehicular technology conference (VTC 2007 Spring). Ireland. April (2007) Page(s): 3091 - 3095.

- [10] H. Holma y A. Toskala, LTE for UMTS: OFDMA and SC-OFDMA Based Radio Access. John Wiley & Sons Ltd. (2009).

- [11] A. Pokhariyal, “Frequency domain packet scheduling under fractional load for the UTRAN LTE downlink”. Vehicular technology conference (VTC 2007 Spring). Ireland. April (2007). Page(s): 699 – 703.

- [12] Kwan, R. ; Leung, C. ; Jie Zhang. “Resource Allocation in an LTE Cellular Communication System”. Communications, ICC '09. IEEE International Conference. (2009) Page(s): 1 – 5.

- [13] S. Shakkottai and A.L. Stolyar, “Scheduling algorithms for a mixture of real-time and non-real-time data in HDR”. Proceedings of the 17th International Teletraffic Congress - ITC-17, Salvador da Bahia, Brazil, 24-28 September, (2001). Page(s). 793-804.

- [14] P. Visawanath, D. N. C. Tse, y R. Laroia, “Oportunistic beamforming using dumb antennas”. IEEE Transaction on Information Theory, vol. 48, no. 6, (2002). Page(s). 1277-1294.

- [15] 3GPP. Especificaciones 3GPP para GSM y EDGE.
Link: <http://www.3gpp.org/ftp/Specs/html-info/45-series.htm>.

- [16] 3GPP. Especificaciones 3GPP para W-CDMA y UTRAN.
Link: <http://www.3gpp.org/ftp/Specs/html-info/25-series.htm>.

- [17] 3GPP. Especificaciones 3GPP para LTE y LTE Advanced.
Link: <http://www.3gpp.org/ftp/Specs/html-info/36-series.htm>.

- [18] 3GPP Official website. The mobile broadband standard.
Link: <http://www.3gpp.org/>.
- [19] IEEE Official Website. The world's largest professional association for the advancement of technology.
Link: <http://www.ieee.org/index.html>.

9. Anexo A: Código utilizado en las simulaciones

9.1 Main.

```
clc;
clear all;
close all;

N=40; %número de usuarios
T=100; %tamaño de la ventana
C=4; %capacidad del sistema,tamaño de RBG(1,2,3,4) en función del
ancho de banda
time= 50; %iteraciones

req=((linspace(0.5,112,N)).*rand(1,N))'; %Requerimientos de los
usuarios (Mbps)

for n=1:time

    M=linspace(1,100,N*T).*rand(1,N*T); %Se obtienen muestras del
canal (S/N)
    M1=reshape(M,N,T); %Se crea una matriz con las muestras

    [R_final1]=proportional_fair(N,T,C,M1);
    [R_final2]=maximum_CI(N,T,C,M1);

    M_final1(:,n)= R_final1;
    M_final2(:,n)= R_final2;

    comp1= R_final1-req;
    comp2= R_final2-req;

    us_satisf1(n,1)=(sum(sign(comp1))+N)/2; %Usuarios satisfechos
Proportional Fair

    us_satisf2(n,1)=(sum(sign(comp2))+N)/2; %Usuarios satisfechos
Maximum CI

end

%Mejoras propuestas: Combinación de los dos algoritmos, tasas
asignadas.

for j=1:time;
```

```

for i=1:N;
if M_final1(i,j)> M_final2(i,j)

    M_max(i,j)=M_final1(i,j);
else
    M_max(i,j)=M_final2(i,j);
end
end
end

%Mejoras propuestas: Combinación de los dos algoritmos, tasas
asignadas a un usuario.

for i=1:N;
if R_final1(i,1)> R_final2(i,1)

    M_max2(i,1)=R_final1(i,1);
else
    M_max2(i,1)=R_final2(i,1);
end
end

%Mejoras propuestas: Combinación de los dos algoritmos, usuarios
satisfechos.

for i=1:time;
if us_satisf1(i,1)> us_satisf2(i,1)

    M_max3(i,1)=us_satisf1(i,1);
else
    M_max3(i,1)=us_satisf2(i,1);
end
end

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% Gráficas %%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%

% Usuarios satisfechos con cada algoritmo

figure;

subplot(2,1,1);
plot(us_satisf1(:,1),'ro:');
title('Usuarios satisfechos aplicando Proportional Fair');
xlabel('Iteración');
ylabel('Nº usuarios satisfechos');

subplot(2,1,2);
plot(us_satisf2(:,1),'k*:');
title('Usuarios satisfechos aplicando Maximum CI');
xlabel('Iteración');
ylabel('Nº usuarios satisfechos');

```

```

%Comparativa de tasas asignadas por los dos algoritmos
figure;

plot(R_final1,'ro:');
title('Comparativa Velocidades Asignadas');
xlabel('Usuario');
ylabel('Velocidad asignada [Mbps]');

hold on;

plot(R_final2,'k*:');
xlabel('Usuario');
ylabel('Velocidad asignada [Mbps]');

legend ('Proportional Fair','Maximum C/I');

% Comparativa usuario concreto, varias iteraciones
figure;

plot(M_final1(5,:), 'ro:'); %Tasa usuario 5
title('Velocidades para un usuario, varias iteraciones');
xlabel('Iteración');
ylabel('Velocidad asignada [Mbps]');

hold on;

plot(M_final2(5,:), 'k*:'); %Tasa usuario 5
xlabel('Iteración');
ylabel('Velocidad asignada [Mbps]');
legend ('Proportional Fair','Maximum C/I');

%Comparativa de tasas asignadas aplicando las mejoras propuestas
figure;

plot(M_max2, 'g-o');
title('Comparativa Velocidades Asignadas');
xlabel('Usuario');
ylabel('Velocidad asignada [Mbps]');

hold on;

plot(R_final1, 'rx:');
title('Comparativa Velocidades Asignadas');
xlabel('Usuario');
ylabel('Velocidad asignada [Mbps]');

hold on;

plot(R_final2, 'k*:');
xlabel('Usuario');

```

```

ylabel('Velocidad asignada [Mbps]');
legend ('Combinación ambos algoritmos','Proportional Fair','Maximum
C/I');
%Comparativa usuario concreto, varias iteraciones, aplicando las
mejoras propuestas
figure;

plot(M_max(5,:), 'g-o'); %Tasa usuario 5
title('Velocidades para un usuario, varias iteraciones');
xlabel('Iteración');
ylabel('Velocidad asignada [Mbps]');

hold on;

plot(M_final1(5,:), 'rx:'); %Tasa usuario 5
title('Velocidades para un usuario, varias iteraciones');
xlabel('Iteración');
ylabel('Velocidad asignada [Mbps]');

hold on;

plot(M_final2(5,:), 'k*:'); %Tasa usuario 5
xlabel('Iteración');
ylabel('Velocidad asignada [Mbps]');
legend ('Combinación ambos algoritmos','Proportional Fair','Maximum
C/I');

% Usuarios satisfechos aplicando las mejoras propuestas
figure;

subplot(2,1,1);
plot(M_max3(:,1), 'g-o');
title('Usuarios satisfechos cambiando ambos algoritmos');
xlabel('Iteración');
ylabel('Nº usuarios satisfechos');

hold on

plot(us_satisf2(:,1), 'k*:');
title('Usuarios satisfechos aplicando Maximum CI');
xlabel('Iteración');
ylabel('Nº usuarios satisfechos');

hold on;

subplot(2,1,2);
plot(M_max3(:,1), 'g-o');
title('Usuarios satisfechos cambiando ambos algoritmos');
xlabel('Iteración');
ylabel('Nº usuarios satisfechos');

hold on;

plot(us_satisf1(:,1), 'rx:');
title('Usuarios satisfechos aplicando Proportional Fair');
xlabel('Iteración');
ylabel('Nº usuarios satisfechos');

```


9.2 Proportional Fair.

```
%-----Función Proportional Fair-----%

function [R_final]=proportional_fair(N,T,C,M1)

planif=zeros(N,T); %Inicializo matriz de asignación
x=ones(C,1);
r_promedio=ones(N,1);

r=tasa_instantanea(N,T,C,M1); %Obtengo la tasa instantánea

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% Matriz de asignación %%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%

for k=1:T
    r_inst=r(N*(k-1)+1:N*k);%Obtengo tasas de N en N
    razon=r_inst./r_promedio;
    [Y,I]=sort(razon,'descend');%Ordeno de mayor a menor las velocidades(Y
    velocidades, I indices)
    J=I(1:C); %Obtengo los C índices útiles en función de la capacidad.
    for i=1:N
        s=(i==J);
        if (x(s))
            planif(i,k)=1;
        end
    end
    r_promedio=(r_promedio+r_inst.*planif(:,k))./k; % Actualizo la tasa
    promedio
end

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% Vector de asignación %%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%

as=zeros(N*T,1);
for m=1:T
    as(N*(m-1)+1:N*m)=planif(:,m);
end

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% Tasa promedio %%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%

I=eye(N); %Matriz identidad
for i=1:T-1
    I=[I;eye(N)];
end

R=r;
for j=1:N-1
    R=[R r];
end

H=R.*I; %Se asignan las tasas correspondientes a cada usuario
H=H';
H=H./T;
R_final=H*as;
end
```

9.3 Maximum C/I.

```
%-----Función Maximum C/I-----%

function [R_final]=maximum_CI(N,T,C,M1)
planif=zeros(N,T);
x=ones(C,1);

r=tasa_instantanea(N,T,C,M1);

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% Matriz de asignación %%%%%%%%%%
for k=1:T
[Y,I]=sort(M1(:,k),'descend');%Ordeno de mayor a menor las tasas(Y
tasas, I indices)
J=I(1:C); %Recupero los C indices útiles, teniendo en cuenta la
capacidad
for i=1:N
s=(i==J);
if (x(s))
planif(i,k)=1;
end
end
end

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% Vector de asignación %%%%%%%%%%
as=zeros(N*T,1);
for m=1:T
as(N*(m-1)+1:N*m)=planif(:,m);
end

%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%%% Tasa promedio %%%%%%%%%%

I=eye(N); %Matriz identidad
for i=1:T-1
I=[I;eye(N)];
end

R=r;
for j=1:N-1
R=[R r];
end

H=R.*I; %Se asignan las tasas correspondientes a cada usuario
H=H';
H=H./T;
R_final=H*as;
end
```

9.4 Función tasa instantánea.

```
%-----Función tasa instantanea-----  
  
%Se obtienen las tasas instantáneas que serán asignadas a los  
usuarios.  
  
function r=tasa_instantanea(N,T,C,M1)  
  
r=zeros(N*T,1);  
  
for k=1:T  
for i=1:N  
  
r(i+(k-1)*N)=C*(C*5)*log2(1+M1(i,k)); %Teorema capacidad  
  
end  
end  
  
end
```