

Aspectos prácticos de compatibilidad electromagnética e integridad de señal

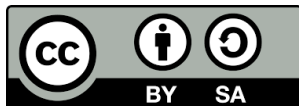
José F. Toledo Alarcón, jtoledo@eln.upv.es

Departamento de Ingeniería Electrónica

Universitat Politècnica de València

Septiembre de 2022, Versión 0.143

Podrás encontrar la versión más reciente en <http://personales.upv.es/jtoledo/>



Los texto e imágenes publicados en esta obra están sujetos -excepto que se indique lo contrario- a una licencia CC-BY-SA v3.0 o v4.0 Internacional de Creative Commons. Te ruego me hagas saber si estoy usando por error alguna imagen con derechos y la sustituiré.

Se puede reproducir la obra, distribuirla o comunicarla públicamente, así como crear obras derivadas, siempre que se cite el autor y la fuente (José F. Toledo Alarcón, jtoledo@eln.upv.es, Universitat Politècnica de València) y se distribuya bajo la misma licencia.

La licencia completa se puede consultar en <https://creativecommons.org/licenses/by-sa/4.0/legalcode.es>

Antes de comenzar

Esta obra no es un libro de texto sobre integridad de señal y compatibilidad electromagnética. Hay excelentes libros de texto publicados con los que no pretendo ni puedo competir.

Lo que tienes en tus manos es un curso, una guía rápida, una exposición de los qué, cómo y porqué que todo diseñador electrónico necesita en su mochila para poder entender y aplicar criterios correctos y buenas prácticas en su trabajo.

Porque no es tan peligroso lo que sabemos que ignoramos, como aquello que no sabemos que ignoramos. Asimilar los contenidos de esta obra permitirá al antes incauto darse cuenta de que se adentra en un campo de minas, y le proporcionará herramientas para detectarlas y no pisarlas. Incluso para desactivar las más sencillas. No se trata de llegar a ser artificiero diplomado: a menudo basta con aplicar el sentido común sobre una sólida base de principios y buenas prácticas. Dejo aquí la metáfora.

Esta obra nace de la necesidad de proporcionar algo más que diapositivas y un discurso hablado a mis alumnos de máster (Máster Universitario en Ingeniería de Sistemas Electrónicos en la Universitat Politècnica de València), en un momento concreto de 2020 en el que la docencia ha de ser *online* (corren los tiempos del COVID-19) y paradójicamente son más necesarios que nunca buenos materiales *offline* para el estudio.

Pero está escrita pensando también en ese profesional, muy posiblemente trabajando en una empresa dedicada a la ingeniería electrónica, que necesita una actualización de su formación.

Citando al gran Félix Rodríguez de la Fuente: “*Ustedes me obligan a estudiar, a leer libros y a hacer síntesis dentro de mi mente, me obligan a concentrar en media hora lo que quizá he tardado en leer tres meses*”. Y en palabras de Neruda: “*Para que tú me oigas/ mis palabras/ se adelgazan a veces/ como las huellas de las gaviotas en las playas.*”

Este ejercicio de síntesis, de destilado, puede llamar la atención a un lector experto en la materia que juzgará la exposición de algunos pasajes del texto como una simplificación excesiva. Mi objetivo es proporcionarte una base simplificada pero sólida, que puedas aplicar directamente a diseños reales sin tener que dedicar tres meses a leer obras especializadas y sin emplear una innecesaria complejidad matemática en el proceso. No obstante, recomiendo una vez completado el estudio de este librito, aprovechar unas vacaciones o un periodo de mayor disponibilidad para beber de un par de grandes obras (si me escribes, te podré recomendar una o dos adecuadas a tus necesidades). Porque nunca hay que parar de aprender.

~~Esta es una obra libre.~~ Te invito a difundirla y hacerme llegar tus opiniones, sugerencias, dudas o erratas a la dirección jtoledo@eln.upv.es

Vas a dedicar al estudio de esta obra horas de tu vida que no volverán. Es una responsabilidad y una confianza que agradezco. Intentaré estar a la altura. [Ahora, comencemos.](#)

Contenido

Día 1. Cosas que tal vez no te contaron sobre electromagnetismo	14
Olvídate de los electrones si quieres hacer carreras.....	17
...porque son demasiado lentos	17
Pero juntos sí son rápidos.....	18
Aunque la verdadera razón de que el efecto sea tan rápido es otra.....	18
Ondas electromagnéticas	19
Guiando ondas electromagnéticas.....	21
Todo es luz a tu alrededor.....	25
Entonces, si todo es "luz", ¿por qué hay diferentes especialidades técnicas en función del rango del espectro en el que se trabaja?.....	27
La velocidad de la luz ¿no? es constante	28
Caminar junto a la señal en un circuito eléctrico	29
Corriente de retorno	30
Caminar junto a la señal es arriesgado: ¡no ves lo que hay delante!.....	31
Campo cercano y campo lejano	32
¿Qué importancia tiene para nosotros saber si estamos en campo cercano o lejano?.....	34
Día 2. Pistas en circuitos impresos	35
¿Qué vamos a aprender hoy?	36
¿Qué vamos a aprender en los próximos días?	37
Una interconexión requiere dos conductores... ..	38
(dos conductores)... que forma un bucle con una determinada inductancia.....	39
Lo mínimo que deberías saber sobre líneas de transmisión	41
¿Cómo sabemos si una conexión es una línea de transmisión?	41
Parámetros básicos de una línea de transmisión	47
Reflexiones	48
Pistas en capas externas de un circuito impreso: líneas microstrip	52
Pistas en capas internas de un circuito impreso: líneas stripline.....	53
Corrientes de retorno en una stripline	53
Día 3. Solución al problema de las reflexiones.....	54
Efecto de las reflexiones en la integridad de señales digitales.....	54
Atenuando de las reflexiones	57
Terminación serie en la fuente	57
Heavy point-to-point.....	60
Terminación paralela en la carga	63
Terminación Thevenin	64
Terminación AC	65
Atenuación mediante resistencias serie en las cargas	65

Ejercicios de reflexiones y terminaciones.....	66
Primer caso práctico (es un caso real de diseño).....	67
Segundo caso práctico.....	72
Día 4. Líneas diferenciales.....	78
¿Por qué usamos líneas diferenciales?	79
La historia de los buses de altas prestaciones en física de altas energías.....	79
Resumiendo lo anterior	81
¿Qué es una conexión diferencial?	81
Un ejemplo de estándar diferencial: LVDS	82
CML (current mode logic).....	82
Cuando el reloj forma parte de los datos.....	82
¿Qué es una línea diferencia en un PCB?	84
Impedancia diferencial.....	89
¿Qué ocurre cuando las líneas del par se influyen?	89
Compensando el efecto de proximidad.....	91
¿Qué separación entre pistas es adecuada?	91
Los caminos de retorno en líneas diferenciales	92
Efecto de cambios de capa de la señal sobre el camino de retorno.....	93
Unas últimas consideraciones prácticas.....	95
Interferencia entre símbolos.....	97
Diagramas de ojos.....	98
Máscaras (eye patterns).....	98
Día 5. Radiación en un PCB	100
Resumen antes de la lección	101
Radiación en modo diferencial	102
Un ejemplo: radiación de una pista de reloj.....	103
Otra vuelta de tuerca sobre el ejemplo... Vamos a añadir terminaciones.....	106
Reducción de radiación diferencial.....	107
Visión general de la radiación por modo común	108
Radiación por conversión de modo diferencial a modo común.....	109
¿Cuál de los dos modos de radiación domina?	109
Reducir la generación de señales en modo común.....	110
Radiación por antenas parásitas	112
Antenas microstrip tipo parche	113
Radiación de una tarjeta pinchada en una placa base	115
Propagación de ruido en planos de masa y de alimentación.....	116
Te presento de nuevo a la guía biplaca	116
Vías como generadores de perturbaciones.....	117
Reducir la generación de ground bounce: condensadores de desacoplo.....	118
Reducir la propagación del ground bounce: aislamiento entre secciones del diseño.....	118
Radiación por bordes del PCB.....	120
Regla 20-H.....	120

PCB stitching	120
Resumen después de la lección	121
Recursos para aprender más	123
Día 6. Diseño de estructuras de PCB multicapa	124
Estructura de un PCB multicapa	125
Parámetros de fabricación de un PCB	127
Impedancia controlada	130
Consideraciones para el diseño del stack-up	131
¿Quién diseña el stack-up?	131
A medida de una señal pasa de una capa a otra, no debe notar un cambio de impedancia	131
Simetría	131
Si puedes permitirte, usa sólo planos de masa para los caminos de retorno de las señales	131
Conviene (dentro de lo posible) alternar capas de señal y planos de masa para reducir crosstalk entre capas de señal	133
Conviene que un plano de alimentación sea adyacente a otro de masa para mejorar el desacoplo	133
Ejemplo de diseño	134
Paso 1: Dibujar la estructura	134
Paso 2: Asignar funciones a las capas	135
Paso 3: Diseñar las capas externas	139
Paso 4: Diseñar las capas internas	140
Paso 5: Comprobaciones finales	140
Paso 6: Generar la especificación de la estructura	142
Para terminar, unas lecturas recomendadas	143
Día 7. Diseño de redes de desacoplo	144
Funciones de un condensador de desacoplo	145
Primera función: eliminar fluctuaciones de baja frecuencia en la alimentación	145
Segunda función: eliminar fluctuaciones de media frecuencia en la alimentación	147
Tercera función: proporcionar un camino de baja inductancia para las corrientes de retorno	149
Cuarta función: reducir la propagación de ruido	150
Comportamiento en frecuencia de un condensador real	150
Extensión del ancho de banda útil mediante combinación de condensadores	151
Tipos de condensadores de desacoplo	152
Inductancia de montaje	153
Inductancia de los planos de masa y de alimentación (plane loop inductance o spreading inductance)	155
Ubicación de los condensadores de desacoplo	156
La herramienta Altera PDN Design Tool	156
EJEMPLO 1: red de desacoplo para 64 ADCs	158
EJEMPLO 2: red de desacoplo para la interfaz LVDS	162
Día 8. Metodología de diseño para integridad de señal	167
Análisis de integridad de señal pre-layout	169
Análisis de integridad de señal post-layout	171

Metodología de diseño.....	172
Modelos IBIS.....	173
Evolución de la especificación IBIS	173
Ejemplo de fichero IBIS (.ibs): Buffer de reloj CY2305	174
Día 9. La Directiva Europea de EMC	180
Vocabulario básico en EMC	182
Solución a los problemas EMC.....	184
Diseñando para EMC.....	185
Marco normativo.....	186
Aclaremos un poco lo anterior	186
Directiva europea de EMC.....	187
Marcado CE	188
Declaración de conformidad	188
Directiva de Equipos de Radio (RED).....	191
Estándares armonizados.....	193
Aclaración sobre alta, baja y media frecuencia en EMC	195
Ensayos de inmunidad	197
Criterios de aptitud en los ensayos de inmunidad.....	200
Ensayos de emisión.....	201
Emisiones radiadas	201
Emisiones conducidas.....	203
Emisiones en baja frecuencia.....	204
Día 10. Diseño de protecciones frente a transitorios ESD	205
Una visión general de los transitorios de alta tensión.....	206
¿Con qué elementos podemos proteger nuestro diseño?.....	207
Normativa IEC para ESD	209
Inciso: simulación del pulso ESD	210
Procedimiento de ensayo ESD.....	211
Primera línea de protección: la caja	213
Cajas de plástico	213
Cajas metálicas	215
Conclusión: la caja no protege completamente.....	216
Diodos TVS: características, consideraciones y ejemplos de uso.....	217
Ejemplo de diodo TVS: On Semiconductor ESD5Z2.5T1G	218
Aplicaciones: protección para SD/μSD	219
Aplicaciones: protección para SD/μSD (un segundo ejemplo).....	220
Aplicaciones: Ethernet.....	221
Aplicaciones: USB 2.0.....	222
El TVS es sólo la mitad de la red de protección: añadiendo una impedancia en serie	224
¿Cómo sabemos cuánta corriente pueden soportar los diodos de protección del IC?.....	225
¿Qué tipo y valor de impedancia serie podemos añadir?	225
Simulaciones SPICE	226

¿Qué simulador?	226
Modelo de una ferrita.....	227
Modelo de un diodo TVS – ejemplo: SMAJ5.0A.....	228
Ejercicio.....	229
Solución	229
Protecciones basadas en aislamiento	231
Día 11. Diseño de protecciones para EFT y onda de choque.....	232
Normativa para EFT	233
Norma básica	233
Normas de familia de productos	235
Simulando pulsos EFT en SPICE	235
Ejercicio de descarga EFT	237
Normativa para onda de choque (Surge).....	238
Norma básica	238
Normas de familia de productos	238
Simulación de ondas de choque.....	239
Dispositivos de protección para ondas de choque	240
Diodos TVS de alta potencia	241
Ejemplo.....	243
Varistores.....	245
Tubos de descarga de gas (GDTs)	248
Día 12. Estudio de un caso de aplicación de normativa EMC: productos IoT.....	250
Extracto del informe (con comentarios añadidos) sobre los productos a certificar por Cute-4 S.L.....	252
Introducción.....	252
Directiva EMC europea para equipos de radio	253
Normativa FCC para equipos de radio.....	254
Recomendaciones de cara a la certificación de los cuatro productos basados en la plataforma #Producto1#.....	256
Recomendaciones de cara a la certificación de los dos productos basados en la plataforma #Producto2#	257
Recomendaciones de cara a la certificación de los dos productos basados en la plataforma #Producto3b#	258
Día 13. Blindajes: atenuando interferencias radiadas	259
¿Qué es un blindaje o pantalla?	260
¿Qué necesitas aprender de blindajes y pantallas?.....	261
Lo que ya sabías: apantallamiento electrostático	262
Reflexión en una superficie conductora.....	264
Absorción en un medio conductor	266
El efecto pelicular como fundamento de la absorción en un blindaje.....	266
Cálculo de blindajes.....	268
Efecto de las ranuras	270
Materiales empleados en los blindajes para RF	270
Ejemplo de cálculo de blindajes	272

<i>Vuelta a 1993</i>	273
Día 14. Interferencias en líneas de PCB, cables planos, coaxiales y pares trenzados	274
Diafonía capacitiva entre dos cables eléctricamente cortos	276
Diafonía inductiva sobre bucles eléctricamente cortos	280
Diafonía entre pistas de un PCB (líneas no eléctricamente cortas).....	282
Vocabulario básico: NEXT, FEXT, forward y backward.....	283
Un caso simplificado: estudio de crosstalk en líneas sin reflexiones.....	284
Estudio de crosstalk en un caso más real	285
Cables planos	287
Cables de pares trenzados	290
Cables coaxiales.....	292
Cable coaxial en un escenario ideal	292
Cable coaxial eléctricamente largo	292
Cable coaxial y bucles de masa en baja frecuencia.....	293
Conexión del blindaje en cables apantallados.....	294
Bibliografía	296

PARTE 1

Fundamentos de Integridad de Señal

Día 1. Cosas que tal vez no te contaron sobre electromagnetismo



Figura 1.1. Fuente: Wikipedia (https://en.wikipedia.org/wiki/File:The_Great_Wave_off_Kanagawa.jpg). Imagen de dominio público

La gran ola de Kanagawa. Una famosa estampa japonesa del siglo XIX representa una ola gigante de cerca de 12 metros (fenómeno casual que resulta de la superposición de varias olas) cerca de la bahía de Tokio. El Monte Fuji se puede apreciar al fondo y la perspectiva lo representa más pequeño que la gran ola con forma de garra, añadiendo dramatismo a la escena. Unas barcas de remo zozobranes, presumiblemente de pescadores, completan una escena que destaca por la belleza de sus azules.

Una ola es la propagación de una onda en la superficie del agua y está causada generalmente por el viento. La energía la transporta la onda, no el medio (agua). Las moléculas de agua, idealmente, vuelven a su posición inicial tras el paso de la ola. La realidad es que hay un pequeño desplazamiento, pero la idea es que la ola es la energía que se propaga.

Una **onda electromagnética** es “una combinación de campos eléctricos y magnéticos oscilantes, que se propagan a través del espacio transportando energía de un lugar a otro” (Wikipedia). Mutando esta definición, me parece más claro decir que una onda electromagnética es “**la propagación de cambios en campos eléctricos y magnéticos transportando energía de un lugar a otro**”.

Y es que para trabajar en EMC (acrónimo de *Electromagnetic Compatibility*, o Compatibilidad Electromagnética) e integridad de señal **debes pensar en términos de propagación de campos electromagnéticos** donde otros sólo ven corrientes eléctricas y diferencias de potencial eléctrico (tensión, si

lo prefieres). Puede parecer sutil, pero es la clave para entender qué está ocurriendo a partir de unos pocos kHz (kilohercios, miles de ciclos por segundo) de frecuencia.

Otra herramienta fundamental es dejar de pensar sólo en el dominio del tiempo y **usar también la perspectiva única que te da el dominio de la frecuencia**. Un osciloscopio te muestra cómo es una señal en el dominio del tiempo. Puedes observar la evolución temporal de su valor instantáneo y extraer información muy útil sobre la señal: su amplitud, si es pulsada o continua, si es repetitiva (periódica) o si parece darse aleatoriamente, si presenta cambios suaves o abruptos (y en este caso sabes que contendrá altas frecuencias).

El dominio de la frecuencia amplía esta información. Como ya te han hablado de la **Transformada de Fourier**, sabes que cualquier señal periódica puede descomponerse como una suma de funciones tipo seno de diferentes frecuencias y amplitudes (la fundamental y por lo general una larga serie de armónicos, múltiplos de la frecuencia fundamental).

En la práctica, con un número pequeño o razonable de armónicos basta para acercarse mucho a la señal real. **¿Qué información nos da el espectro?:** nos indica qué frecuencias contiene la señal y la amplitud de cada componente, lo que resulta extremadamente útil para entender las consecuencias que tendrá la propagación de la señal por la placa de circuito impreso (PCB) y su acoplamiento a otros circuitos (es decir, que la señal actúe como **interferencia** para otras partes de nuestro equipo o para otros equipos) y las medidas que podemos adoptar para reducir o evitar su efecto.

En la Figura 1.2 se muestra un gif animado. Si tienes una versión impresa o pdf de este documento no lo verás animado, pero puedes visualizarlo en la dirección que hay en el pie de la figura. Muestra una señal más o menos cuadrada (por ejemplo, una señal digital) y cómo se descompone en una suma de armónicos de amplitud decreciente. El resultado es una representación de su **espectro**.

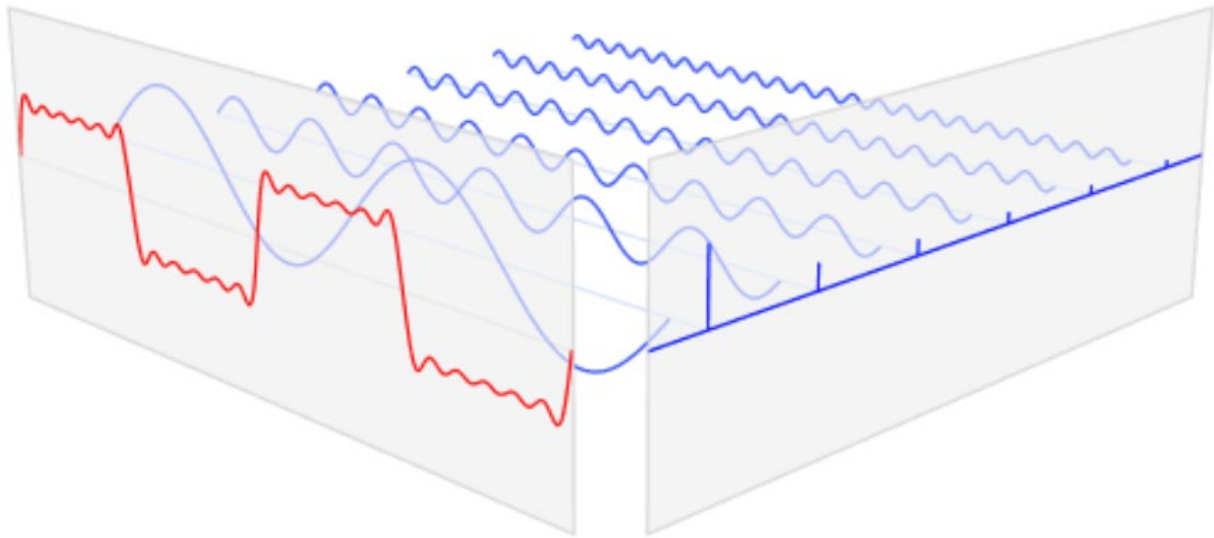


Figura 1.2. Fuente: [Wikipedia](https://commons.wikimedia.org/wiki/File:Fourier_transform_time_and_frequency_domains_(small).gif). Autor: Lucas V. Barbosa, cedida como dominio público. Gif animado.
[https://commons.wikimedia.org/wiki/File:Fourier_transform_time_and_frequency_domains_\(small\).gif](https://commons.wikimedia.org/wiki/File:Fourier_transform_time_and_frequency_domains_(small).gif)

Algunos de los que leéis estas páginas tenéis un *background* adecuado sobre electromagnetismo que os permite entrar en las disciplinas de la Compatibilidad Electromagnética (EMC) y de la Integridad de Señal con buen pie. Otros, en cambio, venís de titulaciones donde trabajáis (casi) exclusivamente con señales de muy baja frecuencia y no manejáis con soltura conceptos y herramientas tales como el espectro de una señal, el espectro electromagnético, la relación entre frecuencia, longitud de onda y velocidad de propagación, impedancia de una línea de transmisión, camino de las corrientes de retorno, propagación de ondas electromagnéticas, campo cercano y campo lejano.

En cualquier caso, un breve repaso de conceptos básicos no le vendrá mal tampoco a los más duchos en electromagnetismo, pues el enfoque que presentamos aquí puede ser distinto (y por tanto complementario y enriquecedor) al que estudiaron en su momento.

Comenzamos, en este Día 1, hablando sobre **corrientes eléctricas** y su inseparable contraparte, los **campos eléctricos**. Comentaremos también algo sobre **propagación de campos eléctricos**, en lo que afecta a la propagación de señales en un cable o en una pista de circuito impreso.

Expondremos el origen de la **radiación electromagnética** y presentaremos algunas formas de confinar, guiar la radiación, como por ejemplo en una **pista en un circuito impreso**.

Terminaremos explicando a qué se refieren los “telecos” cuando hablan de **campo cercano** y **campo lejano**. Con esto daremos por terminada la presentación de conceptos básicos, de ese *background* que te vendrá bien repasar para entrar con buen pie en la materia.

Y lo mejor de todo, vamos a intentar hacerlo sin ecuaciones. No digas que no intento ponerlo fácil...

Olvídate de los electrones si quieres hacer carreras...

Debes cambiar la forma en la que consideras un circuito eléctrico en conmutación, por ejemplo, con señales pulsadas como son las digitales o durante un transitorio. Debes **desaprender** la idea de que la señal (pulsos digitales o una perturbación) se propaga mediante una corriente de obedientes electrones que viajan en fila desde el polo negativo de la fuente de alimentación hasta la carga (por ejemplo, una resistencia) y de aquí al polo positivo de la fuente. Esto sólo te llevará a sacar conclusiones equivocadas. Vamos a explicarlo mediante un ejemplo, antes de que me taches de hereje.

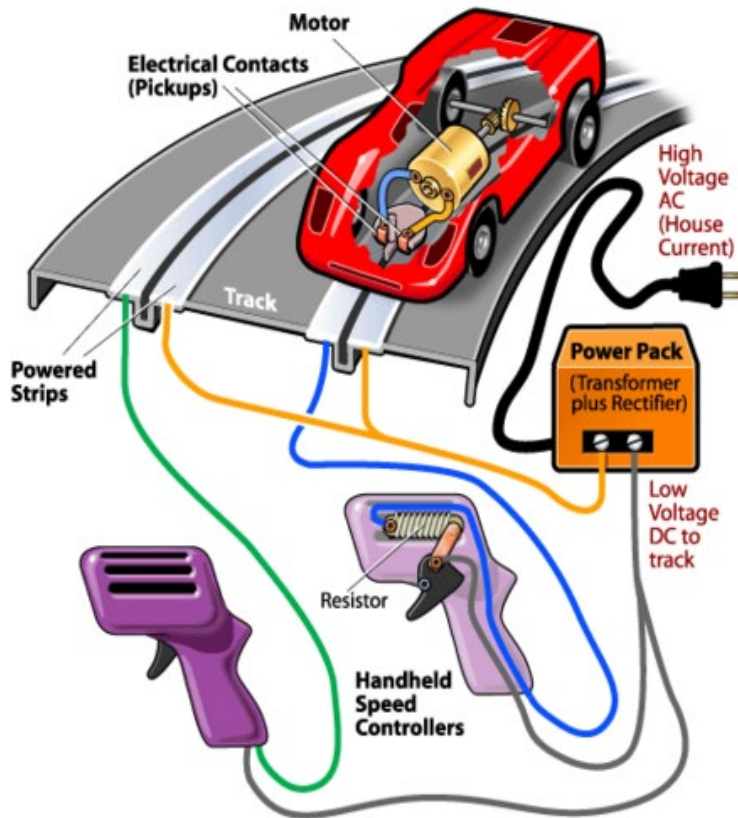


Figura 1.3. Fuente: Wikipedia

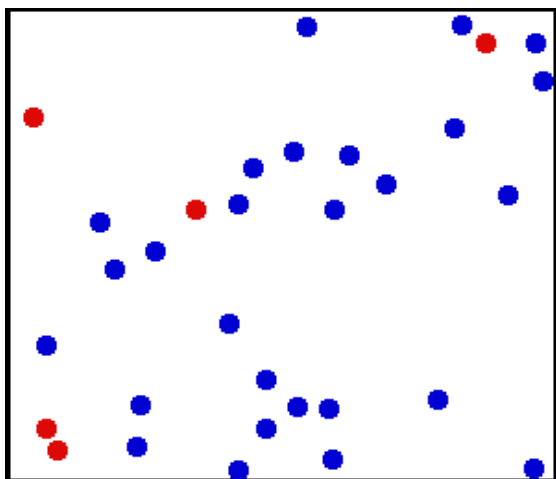
Considera un juego de coches de carreras como el de la figura. Fuente de alimentación de continua, mando (una resistencia ajustable, al fin y al cabo), cables y motor del coche forman un circuito de, digamos, dos metros de longitud (ida y vuelta).

Antes de apretar el mando, la resistencia es tan grande que vamos a considerar la corriente eléctrica nula a efectos prácticos. Por tanto, no hay corriente eléctrica para excitar el motor eléctrico del coche. Imaginemos ahora que mano y mando son tan rápidos que pasamos de una posición (mando suelto, resistencia eléctrica máxima) a la opuesta (mando totalmente apretado, resistencia eléctrica mínima) en un tiempo despreciable.

Si esperas que la Naturaleza te siga el juego y que una fila de obedientes electrones salga de la fuente de alimentación y llegue por el cable hasta el motor, un recorrido de un metro, debes esperar (en función de la sección del cable y de la corriente que pase por él) horas.

...porque son demasiado lentos

Aunque la verdad es que los electrones no son nada lentos. Se mueven a unos 1000 km/s en un cable de metal a temperatura ambiente. No está nada mal. En ausencia de campo eléctrico que los empuje, cada electrón se mueve en una dirección propia hasta que choca con un átomo del conductor, cambia su dirección y vuelve a chocar con otro átomo para volver a cambiar de dirección. El resultado no es muy diferente de lo que ocurre en un gas encerrado en un recipiente, donde cada molécula se mueve independiente de las demás en direcciones distintas y chocando sin parar entre sí y con las paredes (mira la animación en https://en.wikipedia.org/wiki/File:Translational_motion.gif).



En ausencia de campo eléctrico, la corriente eléctrica media en el cable es **cero**, porque en cualquier instante una sección del cable es atravesada por prácticamente la misma cantidad de electrones que van de derecha a izquierda como de izquierda a derecha. El valor instantáneo no nulo de la diferencia es simple agitación térmica que produce **ruido térmico**. Digamos, para simplificar, que la velocidad media en la dirección del cable es nula.

Ahora, cuando apretamos el mando, la resistencia eléctrica es lo bastante baja como para permitir a la fuente de alimentación empujar a los electrones hacia el motor por un hilo del cable y atraerlos por el otro. ¿Qué efecto tiene este "empujón"?

Figura 1.4. Fuente: [Wikipedia](https://commons.wikimedia.org/wiki/File:Translational_motion.gif). Licencia CC BY-SA 3.0. Gif animado.
https://commons.wikimedia.org/wiki/File:Translational_motion.gif

Pues ahora los electrones no se mueven en direcciones completamente aleatorias, sino que tienen una pequeña tendencia a ir en la dirección que marca el campo eléctrico que genera la fuente, y la velocidad media ha pasado de ser nula a ser de aproximadamente 0,1 mm/s.

Si quieres hacer el cálculo de la velocidad media de los electrones en la dirección del campo en función de la corriente eléctrica y de la sección del cable, te recomiendo leer la sección 25.1 de "Física para la ciencia y la tecnología", volumen 2 de Tipler y Mosca. Es un gran libro para repasar conceptos de física.

La corriente en un conductor de sección A , densidad volumétrica de carga n (aproximadamente $8,5 \cdot 10^{28}$ electrones/m³ para el cobre), con cargas de valor Q ($1,6 \cdot 10^{-19}$ C para el electrón) que se mueven a una velocidad v es: $I = n \cdot A \cdot Q \cdot v$. Lo que te sorprenderá es que para un cable de 1 mm de diámetro y una corriente de 1 A, la velocidad de los electrones es de aproximadamente 0,1 mm/s.

Los electrones siguen moviéndose a 1000 km/s, chocando y cambiando de dirección constantemente. Pero sus erráticas trayectorias los acerca al motor del coche a un promedio de unos 0,1 mm/s. Vamos, que para recorrer un metro necesitan 10^4 segundos. Casi tres horas.

Pero juntos sí son rápidos...

No obstante, todo el que ha jugado con Scalextric o encendido la luz de una habitación pulsando un interruptor sabe que el efecto no tarda tanto: es casi instantáneo.

El truco está en que un cable (un conductor de electricidad lleno de electrones libres) se parece mucho a una manguera (un conductor de fluidos lleno de agua), donde al abrir el grifo el agua que entra por un extremo empuja casi instantáneamente a toda la columna de agua. Pues simplificando, lo mismo.

Aunque la verdadera razón de que el efecto sea tan rápido es otra

El símil de la manguera no pasa de ahí. La verdad, porque de eso tratamos aquí, es que un transitorio (una variación instantánea de tensión y de corriente eléctrica como el que provocamos al apretar a fondo el mando) **se propaga por los cables como una onda electromagnética**, a la velocidad de la luz en el medio (ya hablaremos más tarde de a qué velocidad exactamente). Esta onda "empuja" los electrones a su paso. Esto requiere una reflexión por tu parte. Y algo más de explicación por la mía. Nos llevará varios apartados.

Ondas electromagnéticas

Hemos aprendido que, en un cable, los electrones son como las moléculas de agua en el mar: son agitadas por la ola (onda), pero no son quienes transportan la energía. El origen de las olas en el mar suele ser el viento (primero produce un rizado, que crece en altura ofreciendo más superficie de empuje al viento, en una realimentación positiva). Pero **¿cuál es el origen de las ondas en el caso del electromagnetismo?**

Si te digo que son **las cargas eléctricas aceleradas**, puedo provocarte confusión, porque parece que las cargas son a la vez causa y consecuencia. Pero es que es la verdad. Paso a paso. Vamos a hacer esto sin ecuaciones.

Una carga eléctrica estática crea un campo eléctrico estático y no un campo magnético. ¿Cierto?

Si te encuentras fuerte, lee el recuadro siguiente, si no, supón que la afirmación anterior es cierta y sáltatelo. De hecho, la primera vez que leas esta sección, sáltate el recuadro. Podrás volver a él más adelante. Recuerda, paso a paso.

Imagina un electrón estático creando un campo eléctrico estático en el espacio. Llámalo electrón A. Ahora imagina otro electrón B alejado del primero pero que modifica la distribución de campo eléctrico. Por efecto del campo eléctrico, ambos electrones se repelen con una fuerza en la dirección de la línea que les une. Si en lugar de ser dos electrones se trata de dos positrones o de un electrón y un positrón, las distribuciones de campo eléctrico son las siguientes. Nada que no supieras ya de electrostática, ¿verdad?

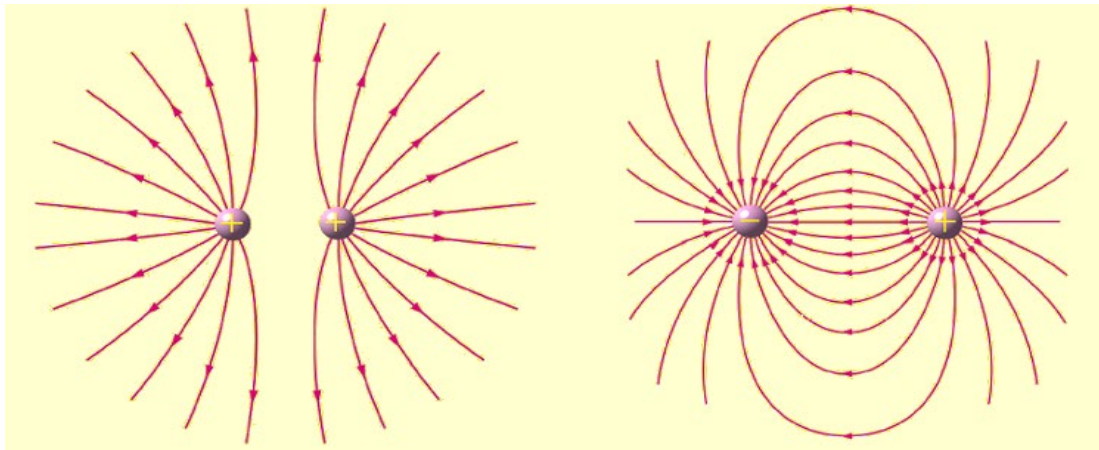


Figura 1.5. Fuente: [Wikipedia](https://es.wikipedia.org/wiki/Imagen:Electric_field_lines.png). Autor: <http://commons.wikimedia.org/wiki/User:Chanchocan>. Licencia Creative Commons Genérica de Atribución/Compartir-Igual 3.0

Ahora añadamos movimiento. El electrón A sigue estático, pero B se mueve a velocidad constante v . Como ya sabes, B en movimiento es una corriente eléctrica. Y eso crea un campo magnético. Como resultado, aparece ahora sobre A una fuerza en el plano ortogonal a la línea que une A y B. Hasta aquí todo bien. Desde el marco de referencia de A (observador quieto respecto a A), B crea un campo magnético (carga en movimiento) y A crea un campo electrostático.

Ahora viene lo bueno. Desde el marco de referencia de B (observador quieto respecto a B, percibe a A en movimiento), A crea un campo magnético. B crea sólo un campo eléctrico.

Las leyes de la Naturaleza no dependen del marco de referencia que tomemos. De modo que hay una única conclusión: campo eléctrico y campo magnético son en esencia manifestaciones de una misma realidad. El campo que observamos (y que sentimos) cambia con el marco de referencia que elijamos.

De hecho, el campo magnético (no generado por imanes, hablamos del producido por corrientes eléctricas) aparece como un efecto de la relatividad especial de Einstein cuando hay campos eléctricos [1].

Hemos dicho que una carga eléctrica estática crea un campo eléctrico estático (Ley de Gauss para campos eléctricos, recuerda). También sabes que una carga en movimiento (corriente eléctrica) crea también un campo magnético con líneas de campo que se cierran sobre sí mismas (Ley de Gauss para campos magnéticos). Si la velocidad de la carga (considerada como un vector, incluyendo por tanto la dirección) es constante, el campo es también constante y todo queda aquí.

¿Qué pasa si el movimiento de la carga es acelerado? Aquí se produce la magia. **Se produce radiación.** Una onda electromagnética. Podríamos justificarlo a partir de las ecuaciones de Maxwell, pero hay otra explicación intuitiva, que prefiero utilizar aquí.

En un artículo de 2001 con el prometedor título “*Electromagnetics without equations*” publicado en IEEE Potentials (Volume 20, Issue 2, Apr/May 2001), Edmund K. Miller parte de las siguientes ideas fundamentales:

1. Los campos electromagnéticos se propagan a la velocidad de la luz
2. Las líneas de campo eléctrico son continuas (no tienen interrupciones)

El autor plantea un estado inicial con una carga en reposo, creando un campo electrostático en el espacio (Figura 1.6 izquierda). Si le damos un brusco empujón a la carga, se desplazará y, por supuesto, seguirá creando un campo eléctrico, centrado en su nueva posición, que se propaga alejándose de la carga a la velocidad de la luz (Figura 1.6 centro). Aunque la carga deje de moverse, la perturbación en el campo eléctrico se propaga por el espacio (Figura 1.6 derecha). Como las líneas de campo han de ser continuas, aparecen unos tramos “radiales” (*kinks*) que unen las líneas de campo “viejas” con las “nuevas”.

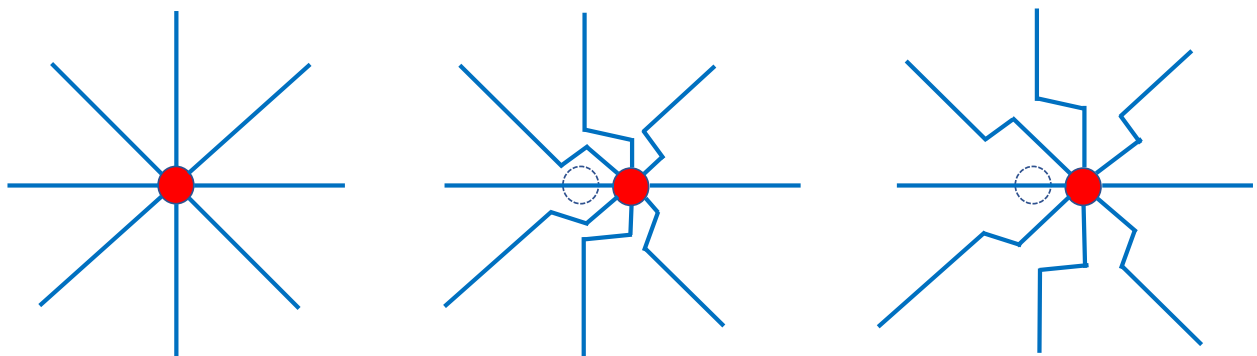


Figura 1.6. Generación y propagación de una onda electromagnética explicada sin ecuaciones. Fuente propia.

A medida que pasa el tiempo y el campo eléctrico se siga propagando a la velocidad de la luz, los *kinks* se propagan. **¡Bingo! ¡Tenemos una onda electromagnética en propagación!**

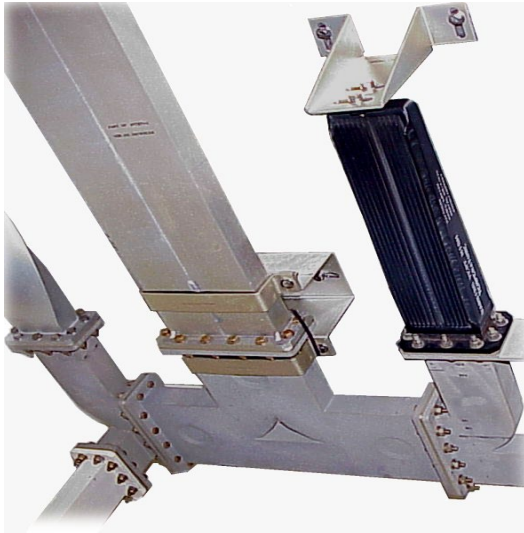
Hemos hablado sólo del campo eléctrico, pero algo similar sucede para el campo magnético. De hecho, ambos campos están trabados, no puede existir uno sin el otro, y por eso hablamos de campo o radiación electromagnéticos.

Quédate con la idea básica: la radiación electromagnética está causada por cargas aceleradas. Ya sea un electrón que salta de un nivel de energía a otro inferior (así funcionan los LEDs), una corriente eléctrica variable a través de un cable (lo que constituye un tipo de una antena y hará que tus cables no apantallados te den dolores de cabeza cuando intentes pasar los ensayos de EMC a tu nuevo producto), o una pista en capa externa de tu placa de circuito impreso.

Vamos a ver ahora formas de confinar la radiación en estructuras apropiadas, para llevarla donde nos interesa.

Guiando ondas electromagnéticas

Seguramente asocias el término "**onda electromagnética**" con las antenas. Porque asumes que se propaga por el aire y por el vacío. **Pero hay otras formas de propagar una onda electromagnética por un medio, especialmente si queremos contenerla, guiarla, por donde nos interese.**



Por ejemplo, un ingeniero de microondas usa una pequeña antena como fuente de señal y una tubería hueca para confinar y propagar una onda en su interior (lo que, en su argot, los “telecos” llaman **guía de onda**, que siempre queda mejor que “tubería”).

Una antena en un extremo de la tubería genera un campo eléctrico (antena de varilla) o magnético (antena de lazo) variable, dando lugar a una onda que se propaga rebotando en las paredes internas.

Eso sí, hay que usar una tubería pulida que produzca pocas pérdidas, ya que el movimiento de los electrones (corriente eléctrica) en el metal disipa la energía de la onda en forma de calor (choques de los electrones).

La Figura 1.7 muestra un componente de un radar de control de tráfico aéreo.

Figura 1.7. Fuente: Wikipedia. By Averse 19:31, 25. Jan. 2007 (CET) - Own work, CC BY-SA 3.0, <https://commons.wikimedia.org/w/index.php?curid=12708990>

Otro ejemplo de guía de onda es una **fibra óptica**.

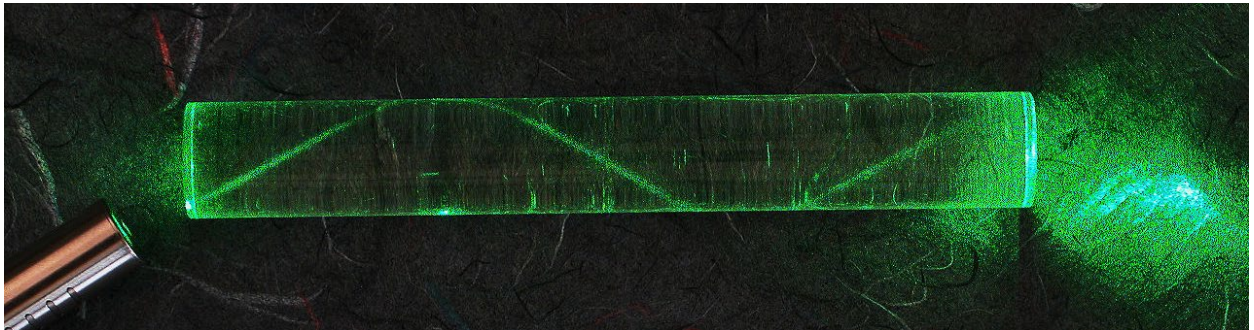


Figura 1.8. Fuente: Wikipedia. Autor: Timweather. Licencia, CC-BY-SA 3.0 Unported

En óptica se usa el término índice de refracción en lugar de permitividad: diferentes especialidades, diferente lenguaje. Si la guía de onda es un medio con permitividad superior a la del medio que la rodea (por ejemplo, fibra de cuarzo -una forma de SiO_2 - rodeada de aire) y el ángulo de incidencia es el adecuado, se produce el efecto de reflexión total interna y la luz queda confinada en el interior de la fibra, rebotando por las paredes y avanzando por la fibra a la velocidad de la luz en el medio.

Las **líneas de transmisión** son otra forma de guía de onda. En la forma de cables coaxiales o pistas de circuito impreso, definen una estructura en la que un campo electromagnético se propaga más o menos confinado entre dos planos conductores a las frecuencias que empleamos habitualmente en electrónica (desde continua hasta decenas de GHz).



Figura 1.9. Fuente: [Wikipedia](https://commons.wikimedia.org/w/index.php?curid=36701532). By Sbyrnes321 - Own work, CC0, <https://commons.wikimedia.org/w/index.php?curid=36701532>

En la imagen anterior, que es sólo un cuadro de una animación que te recomiendo visualizar, las flechas rojas representan la amplitud de campo eléctrico de una onda que se desplaza de izquierda a derecha, contenida (guiada) entre dos planos metálicos. Los puntos negros son electrones de los planos metálicos, empujados a la derecha y a la izquierda por efecto del campo. Algunos de mis alumnos tienden a creer que el campo se propaga **por** las superficies metálicas, cuando realmente lo hace **entre** ellas.

En un cable **coaxial**, el conductor interno hace las veces de uno de los dos planos y la malla metálica que lo rodea, del segundo.

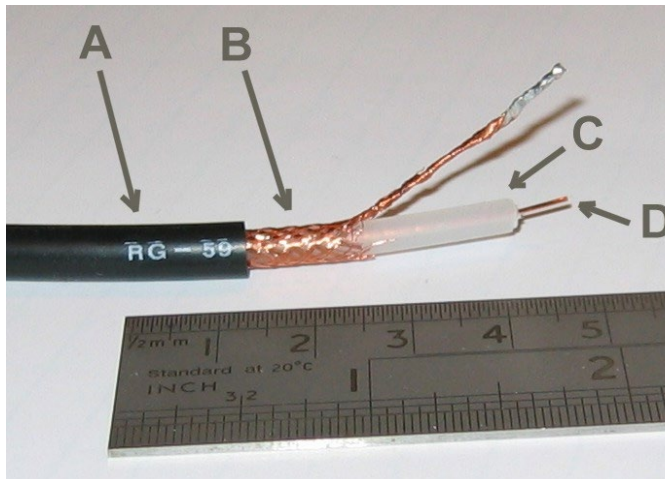


Figura 1.10. Fuente: [Wikipedia](#). CC BY-SA 3.0, <https://commons.wikimedia.org/w/index.php?curid=50762>

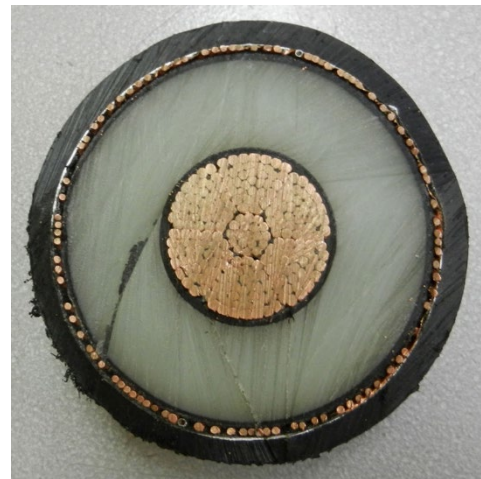


Figura 1.11. Fuente: [Wikipedia](#). Licencia CC BY-SA 3.0

En una **placa de circuito impreso**, uno de los planos metálicos es la pista, el otro es un plano de masa o de alimentación. Se denomina **microstrip** a esta estructura (Figura 1.12).

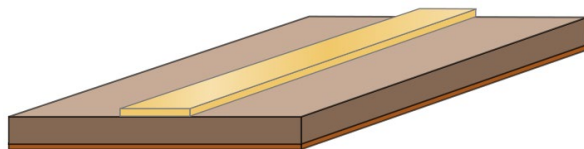


Figura 1.12. Fuente: [Wikipedia](#). By Courtesy of Spinningspark at Wikipedia, CC BY-SA 3.0, <https://en.wikipedia.org/w/index.php?curid=52984582>

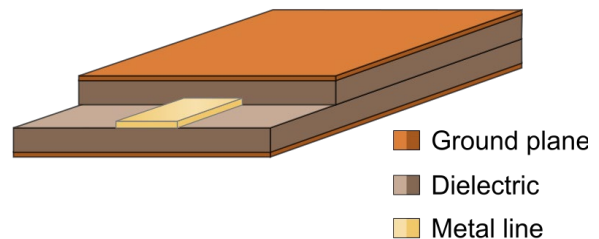


Figura 1.13. Fuente: [Wikipedia](#). Licencia CC BY-SA 3.0

Otra posibilidad es añadir un plano adicional, lo que ocurre en capas internas de un circuito impreso multicapa (Figura 1.13). A esta estructura le llamamos **stripline**.

ADVERTENCIA: lo que sigue es complicado, no te preocupes si no lo entiendes, es una muestra del argot y el nivel de comprensión que alcanzaremos (o eso espero) al final del camino

Hay un campo de aplicación a medio camino entre las guías de onda rectangulares y las placas de circuito impreso (PCB) convencionales, que permiten trabajar hasta 100 GHz. La siguiente figura es un filtro pasa banda (es decir, deja pasar a través de él las frecuencias comprendidas entre un límite inferior y otro superior) seguido de un filtro pasa bajo (deja pasar sólo aquellas frecuencias por debajo de un límite).

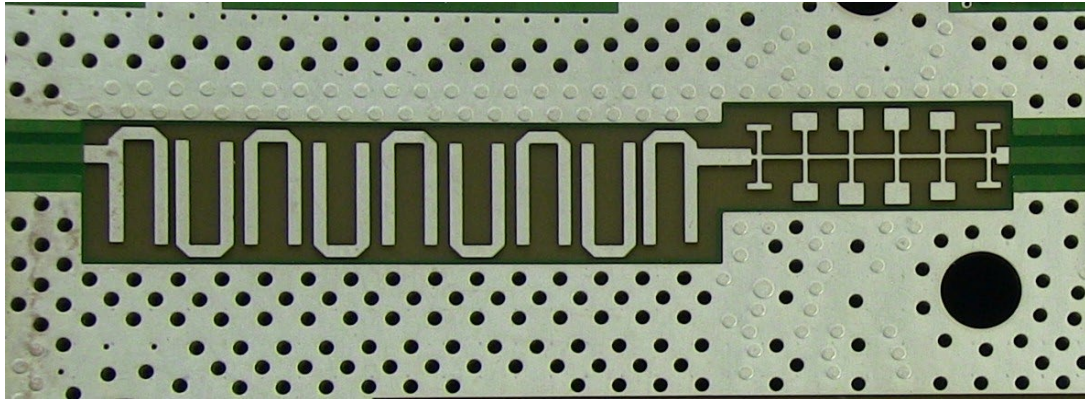


Figura 1.14. Fuente: [Wikipedia](#). Licencia CC-BY-SA-3.0. Autor: [Binarysequence](#)

Fíjate en los puntos de entrada y salida del filtro a la derecha y a la izquierda en la Figura 1.14, cubiertos de una máscara de soldaduras verde.

Los orificios son taladros pasantes metalizados (denominados **vías** en argot de los diseñadores) que cortocircuitan a distancias inferiores a $\lambda/10$ (la letra lambda se usa para representar la longitud de onda de la señal, ya veremos qué es esto) un plano de masa situado en la cara superior (cubierto de máscara de soldadura blanca) con otro en una capa interior o en la inferior. La longitud de onda a 5 GHz en un PCB (suponiendo que ésta sea la frecuencia central del filtro) es aproximadamente 30 mm. Las vías deberían estar separadas no más de 3 mm para formar una barrera eficaz e impedir su propagación.

¿Por qué se hace esto?

1. Para confinar la señal dentro de la “jaula” rodeada por las vías, evitando que escape
2. Para evitar que ruido y señales de otras partes de la PCB afecten (interfieran) a la señal a filtrar
3. Para evitar que los dos planos de masa que conforman la jaula creen una estructura resonante
4. Para evitar que un área de metal aislada (a veces el programa o el diseñador se dejan un área de cobre no conectada al resto) actúe como antena. Hablaremos sobre esto más adelante

Como ves, a poco que nos pongamos a comentar una imagen salen muchos temas importantes a la luz. Bueno, tenemos un curso entero para irlos presentando. Paciencia. Un pequeño paso detrás de otro permite, con el paso del tiempo, completar un viaje. Como dicen los italianos, *piano, piano, si va sano e lontano*. Y si eres más moderno, es el principio Kaizen de Lean: mejora continua a partir de pequeños pasos incrementales. Aprende algo cada día, aunque sea poco, se constante y con el paso del tiempo habrás llegado lejos. Poco a poco.

Volviendo a la propagación de campos, la verdad es que cada margen de frecuencias encuentra medios y geometrías de propagación más favorables que otras. Pero antes de continuar hablando de propagación de radiación electromagnética asegurémonos de comprender algunos conceptos sencillos en la siguiente sección.

Todo es luz a tu alrededor

Debo insistir: en un cable eléctrico o en una pista de circuito impreso, la energía no viaja en los electrones. Si excitas un conductor con una señal de 50 Hz, cuyo periodo es de 20 ms, los electrones se mueven durante 10 ms en una dirección y durante la otra mitad del ciclo en la contraria. Asumiendo la velocidad de vértigo de 0,1 mm/s que hemos comentado con anterioridad (y no es un valor exacto, es sólo un orden de magnitud, depende de la corriente eléctrica y de la sección del conductor), en 10 ms un electrón recorrería 1 μm (micra o micrómetro, 10^{-6} metros).

Es decir, los electrones se limitan a oscilar hacia adelante y hacia atrás una distancia irrisoria. [En este link \(https://en.wikipedia.org/wiki/Transmission_line#/media/File:Transmission_line_animation3.gif\)](https://en.wikipedia.org/wiki/Transmission_line#/media/File:Transmission_line_animation3.gif) encontrarás de nuevo la animación que ilustra cómo se propaga una onda electromagnética entre dos conductores (los puntos negros representan electrones y las flechas el campo eléctrico). Mírala de nuevo para fijar la idea.

No, la energía no la transportan los electrones. Siento despertarte del sueño: la energía la transporta un campo electromagnético que se propaga entre dos conductores. El campo electromagnético es la energía de la ola en el mar, y los electrones son las moléculas de agua, que se mueven arriba y abajo por efecto de la energía de la ola.

Pero es cierto que sin agua no hay ola. Del mismo modo, si resolvemos las ecuaciones de Maxwell para una pista de circuito impreso o un par de cables, vemos que:

1. Un campo eléctrico da lugar a una corriente eléctrica (empuja a los electrones)
2. La corriente eléctrica da lugar a un campo magnético
3. Si el campo magnético es variable (porque la corriente eléctrica también lo es), da lugar a su vez a un campo eléctrico, que... (vuelve al punto 1)

Vamos, que es condición necesaria para la propagación de la onda en el caso que consideramos (cable o circuito impreso, es decir, guiado por una estructura metálica) que haya corriente eléctrica variable y que haya campos variables. Sin variaciones, no hay propagación de campos electromagnéticos. Sólo hay campos estáticos (y ya te han hablado bastante en la carrera de electrostática y magnetismo).

Aquí queremos estudiar qué ocurre cuando hay cambio. Qué es lo que pasa al usar señales senoidales en comunicaciones por radiofrecuencia, pulsos en señales digitales y conmutación en fuentes de alimentación conmutadas. Por poner sólo algunos ejemplos.

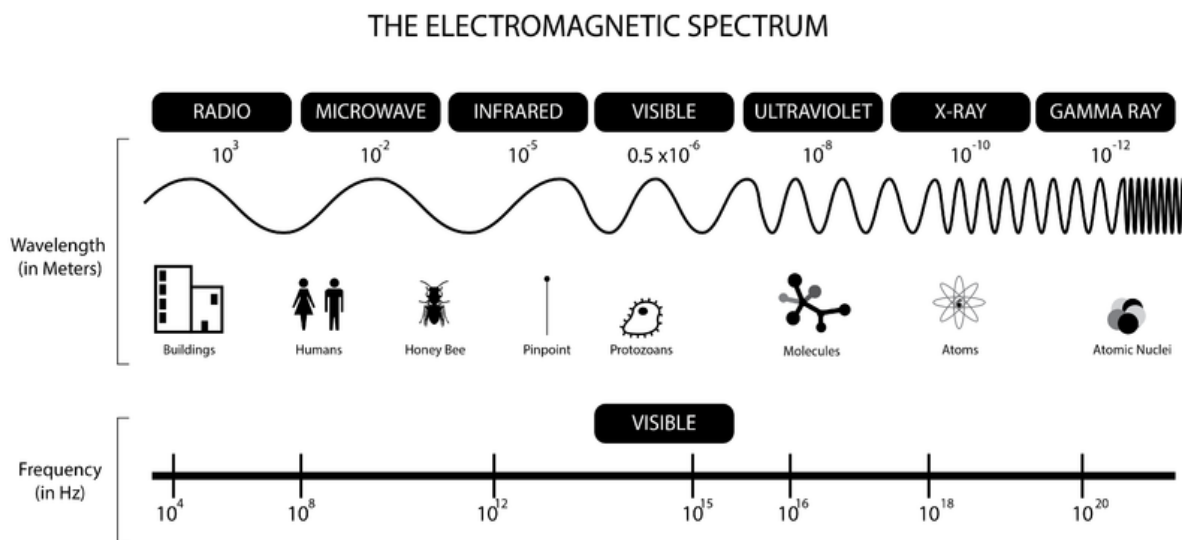


Figura 1.15. Fuente: Wikipedia (https://commons.wikimedia.org/wiki/File:BW_EM_spectrum.png) bajo licencia CC-BY-SA 3.0. Autor: Jonathan S Urie

Vamos a aclarar algunos **conceptos sobre ondas electromagnéticas**. Estudia la Figura 1.15. Desde campos estáticos hasta radiación gamma, todo son campos electromagnéticos. Podemos llamarlo a todo luz, para generalizar la idea. El eje horizontal de la figura no es el tiempo, como estás acostumbrado, sino frecuencia (en ciclos por segundo, o hercios, Hz).

Un **ingeniero eléctrico** está acostumbrado a trabajar con señales que se extienden desde continua hasta unos pocos kHz (kilohercios, ciclos por segundo). Excepto cuando analiza armónicos en la red eléctrica, prefiere contemplar la señal en el dominio del tiempo, como en un osciloscopio.

Un ingeniero eléctrico sabe que en la red eléctrica hay un máximo cada 20 ms (es idealmente una función seno de 50 Hz de **frecuencia**, luego su **periodo** es 1/50 Hz, 20 ms). Y si acepta que todo es luz, a 300.000 km/s (luego veremos que este valor será algo diferente) en 20 ms la señal eléctrica recorre 6.000 km. Es decir, en el cable encontramos que los máximos están separados 6.000 km. En un cable hasta la luna (384.000 km) caben 64 ciclos de la señal eléctrica, con sus 64 máximos y 64 mínimos. Recuerda que un electrón se mueve en promedio tal vez 1 μm en cada semiciclo de una señal de 50 Hz. ¿Enterramos ya definitivamente la idea de electrones que viajan por el cable? Simplemente vibran alrededor de una posición de equilibrio.

La longitud que ocupa un ciclo de la onda de la señal eléctrica de 50 Hz es de 6.000 km. Para abreviar, diremos que la **longitud de onda** a 50 Hz es de 6.000 km. La expresión que solemos usar para calcular la longitud de onda en metros es $\lambda = c/f$, siendo c la velocidad de la luz en el medio y f la frecuencia en Hz. Y usamos la letra griega λ "lambda" para referirnos a la longitud de onda.

En el rango de frecuencias usado en electricidad y en electrónica de baja frecuencia, como los cables son normalmente mucho más pequeños que la longitud de onda (6.000 km para una onda de 50 Hz), decimos que los cables son **eléctricamente cortos**. Eso quiere decir que todos los puntos de cable se encuentran a efectos prácticos al mismo potencial eléctrico. Por tanto, en un cable corto (de nuevo, a efectos prácticos) no cae tensión, es una conexión ideal y podemos trazar pistas en un circuito impreso y cablear equipos sin preocuparnos por nada.

Lo que acabo de describir es lo que ocurría hasta los felices años 80 del siglo XX, cuando los ingenieros electrónicos que no trabajaban con RF (radiofrecuencia) no tenían necesidad de vérselas con efectos indeseados (e indeseables) sobre los que versa buena parte de este curso

Un **ingeniero electrónico** actual amplía su margen de frecuencias de trabajo del **espectro electromagnético**, pues suele trabajar hasta al menos 1 GHz (a menudo hasta 10 GHz). El periodo de una señal de 1 GHz es $1/10^9 \text{ Hz} = 1 \text{ ns}$. En 1 ns, la luz recorre $300.000 \text{ km/s} \cdot 1 \text{ ns} = 30 \text{ cm}$. Por tanto, la longitud de onda es de 30 cm. En un cable de 30 cm de longitud cabe exactamente un ciclo de la señal de 1 GHz.

¡Espera! Entonces, un cable de 10 cm ya no es mucho más pequeño que la longitud de onda a 1 GHz, como ocurriría si la señal a considerar fuera de 50 Hz. No podemos considerar que todos los puntos del cable están a la misma tensión eléctrica. Se dice entonces que **el cable ya no es eléctricamente corto**.

Cuando un cable o una pista de circuito impreso no es eléctricamente corto (frontera que se sitúa en aproximadamente la décima parte de la longitud de onda, lo que verás por todas partes como $\lambda/10$), se manifiestan ciertos fenómenos que nos complican la vida, ya hablaremos largo y tendido sobre ello en el Día 2. Y esto hace que a medida que los ingenieros electrónicos hemos tenido que trabajar con *señales que contienen frecuencias cada vez más altas* (he elegido cuidadosamente las palabras en cursiva) hemos pasado de ser felices pensando que un cable no era parte de un circuito a tener dolores de cabeza con **líneas de transmisión** (que podemos definir como cables o pistas de circuito impreso que no son eléctricamente cortos).

Los **ingenieros de RF (radiofrecuencia)** y microondas trabajan típicamente entre 300 MHz y 300 GHz (longitudes de onda entre 1 m y 1 mm). El límite superior lo marca el comienzo de los infrarrojos lejanos. Y a partir de aquí comienza la óptica (infrarrojo cercano, luz visible y ultravioleta).

Entonces, si todo es "luz", ¿por qué hay diferentes especialidades técnicas en función del rango del espectro en el que se trabaja?

Ingenieros eléctricos, electrónicos, de microondas y los especializados en óptica, todos trabajamos con lo mismo. Ondas electromagnéticas. Luz. Pero no nos entendemos al hablar. ¿No te lo crees? He mantenido no pocas conversaciones frustrantes con físicos, ingenieros especialistas en antenas y en óptica.

En cada “especialidad” hemos desarrollado unos modelos diferentes, un argot específico, unas simplificaciones de la Naturaleza que nos permiten entender la fenomenología que encontramos en la porción del espectro que usamos en nuestro trabajo. Y estos modelos son distintos y con frecuencia contradictorios.

Que un conductor sea eléctricamente corto (ingeniería eléctrica y electrónica de baja frecuencia) permite ignorar varios efectos muy importantes que los ingenieros electrónicos de alta frecuencia y los ingenieros de microondas no pueden obviar (por ejemplo, que una pista de circuito de impreso de 1 cm se comporte como un condensador, una inductancia o un circuito abierto según la frecuencia de trabajo).

Que las propiedades de distintos materiales cambien bruscamente a altas frecuencias es algo que permite la óptica (no es casualidad que nuestros ojos sean sensibles al espectro visible, donde hay cambios bruscos en las propiedades ópticas de la materia).

Así, un ingeniero eléctrico no se olvida de "cerrar el circuito" en sus diseños (si no, no circula corriente eléctrica, ¿verdad?). Un ingeniero electrónico debe además mimar la corriente de retorno que se induce en un plano de masa bajo una pista, evitando discontinuidades en el camino si quiere no distorsionar la señal debido a reflexiones y a caminos adicionales de propagación, amén de provocar radiación indeseada que le haga incumplir la normativa sobre compatibilidad electromagnética.

En cambio, un ingeniero de radiofrecuencia que diseña una antena no suele pensar para nada en que existe la corriente de retorno inducida en el plano de masa, pues precisamente busca sin pararse a pensarlo una discontinuidad como forma de provocar la radiación.

Un óptico ya no emplea un par de conductores como medio de propagación, por lo que ya no es que no considere corrientes de retorno, es que no considera corrientes en absoluto.

Lo dicho, diferentes problemas a resolver dan lugar a diferentes formas de explicar (simplificar) la Naturaleza y a diferentes herramientas, lenguajes, y especialidades. Una buena base de la física subyacente a la electricidad y al electromagnetismo es la clave para ser un ingeniero transversal.

A través del buscador de la biblioteca de tu Universidad tendrás disponible como libro electrónico una excelente obra publicada en 2002, "[*Electromagnetics explained: a handbook for wireless/RF, EMC and high speed electronics*](#)", de Ron Schmitt. Te recomiendo sin reservas su lectura, si tienes disponibilidad y ganas de entender un poco mejor la naturaleza de los fenómenos electromagnéticos

La velocidad de la luz ¿no? es constante

¿300.000 km/s? Sólo a veces. Para ser más concreto, sólo allí donde el medio de propagación de la onda electromagnética tiene una permitividad (constante dieléctrica) igual a la del vacío. Que viene a ser igual a la del aire (bueno, la del aire es aproximadamente 1,0005 veces mayor que la del vacío, pero no vamos a pelearnos por esto, ¿verdad?). También llamada **permitividad**, a veces **constante dieléctrica**, representada por la letra griega ϵ , *epsilon*, determina la velocidad a la que se propaga la luz por el medio.

En una placa de circuito impreso la permitividad es aproximadamente 4 veces mayor que en el vacío (y decimos que su permitividad relativa es de 4). En una fibra óptica es también cercana a 4, y no por casualidad, sino porque ambos medios se basan en dióxido de silicio. En el agua, la permitividad relativa es casi de 80.

Y ocurre que la velocidad de propagación en el medio disminuye de forma inversamente proporcional a la raíz cuadrada de este valor ($c = c_0/\sqrt{\epsilon_r}$, siendo c_0 la velocidad de propagación en el vacío y ϵ_r la permitividad relativa del medio). Lo que nos deja (circuitos impresos y fibras ópticas) en aproximadamente la mitad de la velocidad en el vacío, es decir, 150.000 km/s. De 30 cm/ns en el aire pasamos a 15 cm/ns en un módulo electrónico, un cable eléctrico o en una fibra óptica. De modo que en un circuito impreso una señal se aleja de tu nariz un poco menos de un palmo en un nanosegundo (un océano de tiempo, ya verás).

Un óptico ve esto de otra forma. Tirando de Wikipedia:

*El **índice de refracción** de un material indica cuán lenta es la velocidad de la luz en ese medio comparada con el vacío. La disminución de la velocidad de la luz en los materiales puede causar el fenómeno denominado refracción, como se puede observar en un prisma atravesado por un rayo de luz blanca formando un espectro de colores y produciendo su dispersión. Al pasar a través de los materiales, la luz se propaga a una velocidad menor que c , expresada por el cociente denominado «**índice de refracción**» del material.*

Lo que habíamos dicho antes: aunque todo sea luz, diferentes problemas a resolver dan lugar a diferentes aproximaciones, jergas y finalmente, especialidades. Unos hablamos de permitividad, otros de índice de refracción.

Pero... ¿qué realidad física subyace a estas jergas de especialista? La respuesta larga (porque te va a llevar más tiempo) es que leas un delicioso librito de bolsillo de 225 páginas escrito por el archifamoso Richard Feynmann [1]. Si quieres entender un poco mejor el mundo físico en el que vives, léelo. La respuesta corta es **la difusión de la luz**. Cundo un fotón se adentra en un medio que no es el vacío, puede ser absorbido por un electrón de medio. El electrón se desplazará una pequeña distancia y emitirá otro fotón. Esto lleva un tiempo, pequeño pero relevante, y tiene dos efectos:

1. Retrasar el avance de la luz, haciendo que parezca que los fotones son más lentos (lo que no es verdad)
2. Cambiar la dirección de la luz, dispersándola de forma selectiva con la frecuencia y dando lugar a cielos azules, atardeceres rojos y luz que se “doble” al atravesar un vaso de agua

Si quieres saber más, lee el librito de Feynmann. No te llevará más tiempo que el que dedicas a ver una o dos (si haces una lectura cuidadosa) temporadas de una serie: no tienes excusa.

Caminar junto a la señal en un circuito eléctrico

Volvamos a nuestro ámbito, jerga y forma de explicar la parte de la Naturaleza que nos incumbe: la electrónica.

A la onda electromagnética que se propaga la llamamos **señal**. El modo de propagación que utilizaremos para esta sección son dos conductores metálicos paralelos: una pista de circuito impreso sobre un plano, ya sea masa o alimentación (el valor de tensión en continua es irrelevante para la propagación de la señal). Entre ambos conductores, un dieléctrico (igual te sientes más cómodo si lo llamas FR4 o fibra de vidrio con resina epoxi) con permitividad relativa cercana a 4. Como ya sabes, has creado una estructura que presenta capacidad eléctrica (dos conductores entre los que hay una diferencia de tensión eléctrica) e inductancia (la corriente que circula crea líneas de campo magnético a su alrededor: la inductancia es proporcional al número de líneas creadas por unidad de longitud por amperio que circula por la estructura).

Imagina a la izquierda de la estructura tipo sándwich (Figura 1.9), una fuente de tensión que produce un escalón de 5 V a 0 V. A la derecha de la estructura imagina una carga (una resistencia). **Reloj a cero.** Comienza el escalón de tensión.

- A. La fuente de tensión, para reducir la diferencia de tensión entre ambos conductores (recuerda que queremos pasar de 5 V a 0 V), debe entregar una corriente eléctrica que descargue el condensador formado por el tramo inmediatamente cercano de pista-dieléctrico-plano. $I = C \cdot \delta V / \delta t$. Ya sabes. La ecuación del condensador.
- B. Además, el dieléctrico, que se polariza (se “orienta” en la dirección de campo eléctrico, si las moléculas presentan cierta asimetría en la disposición de las cargas), entra en este juego, ralentizando el efecto de la descarga. Poco a poco, el tramo de línea se va descargando por efecto de la corriente que entrega la fuente.

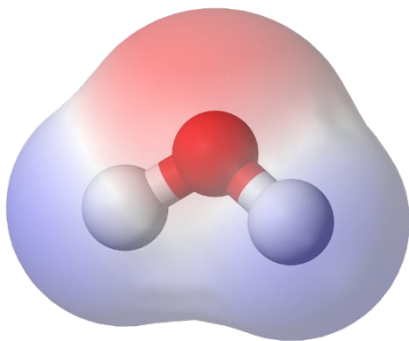


Figura 1.16. [fuente: Wikipedia](#). Imagen de dominio público.

En la Figura 1.16, una molécula de agua presenta una asimetría entre sus cargas positivas (átomos de hidrógeno) y negativas (oxígeno). Bajo el efecto de un campo eléctrico, la molécula se reorienta, aumentando la carga neta acumulada en el condensador y por tanto su capacidad (recuerda, $C = Q/V$). A altas frecuencias, la molécula vibra, disipando energía de la onda en forma de calor y por tanto atenuando la señal (es lo que hace un horno microondas). Por eso en alta frecuencia se usan a veces en circuitos impresos sustratos de bajas pérdidas, es decir, con un comportamiento polar reducido. Es algo que comentaremos otro día cuando hablemos de integridad de señal y pérdidas en alta frecuencia.

C. La onda electromagnética sigue propagándose hacia la derecha, recorriendo aproximadamente 15 cm cada nanosegundo. El frente de la onda es capaz de provocar una reducción de tensión eléctrica que depende de la corriente que sea capaz de entregar la fuente y de la capacidad eléctrica por unidad de longitud entre conductores. En otras palabras, la “ola” (**onda incidente**, **frente de onda**, le puedes dar diversos nombres) que se propaga a la velocidad de la luz en el medio no suele tener “fuerza” para conmutar de 5 V a 0 V. Necesita más tiempo. Ya veremos esto con más calma más adelante.

Si queremos observar el efecto no mediante un desplazamiento continuo del frente de ondas, sino a pequeños saltos (una forma de modelarlo y entenderlo mejor), diremos que, **caminando junto al frente de onda de la señal**, cada paso representa la descarga (parcial) de un nuevo tramo de línea, representado por un pequeño condensador (típicamente 1 pF por centímetro de línea). ¿Qué longitud debe tener un tramo para que nuestro modelo sea realista? Pues debe ser eléctricamente corto. Cuando más pequeño, más exacto. Pero basta con que sea un 5% de la longitud de onda para obtener buenos resultados.

Ya sabes que no pasa corriente realmente a través de un condensador (de ahí que se denomine corriente de desplazamiento), simplemente el efecto neto es el de exceso/déficit de electrones en cada conductor (pista y plano). Sólo es carga acumulada, no saltan electrones entre terminales.

Corriente de retorno

Pero ese movimiento de electrones es en definitiva una corriente eléctrica. Corriente de “ida” en la pista y de “retorno” en el plano (recuerda: estamos asumiendo en esta sección la propagación de la señal en una estructura formada por una pista sobre un plano de referencia). Esto es muy importante: a cada paso que damos, el frente de onda crea una corriente de ida en la pista, pero también otra de sentido contrario en el plano de referencia, a la que llamamos **corriente de retorno**.

Lo escribo de nuevo, para que no lo olvides: **corriente de retorno**. No tendrás problemas en reconocer que debes evitar que la pista esté cortada (interrumpiría la corriente de ida) o que cambie bruscamente de anchura (afectaría a la corriente de ida). Lo que igual te cuesta más es admitir que debes cuidar de igual manera la corriente de retorno, pues afecta de igual manera a la señal, deformándola y afectando a lo que llamamos integridad de señal. De hecho, parte de este curso versa sobre cómo mimar la corriente de retorno para mejorar la calidad de la señal y reducir la radiación.

¿Y qué pasa si tenemos un PCB a una o dos caras, si plano de referencia? En ese caso, la estructura por la que se propaga la señal no está bien definida y posiblemente observemos una serie de fenómenos que se manifiestan como distorsión de la señal.

La formación recibida en nuestros estudios previos hace que nos guste dibujarlo así, como un circuito eléctrico:

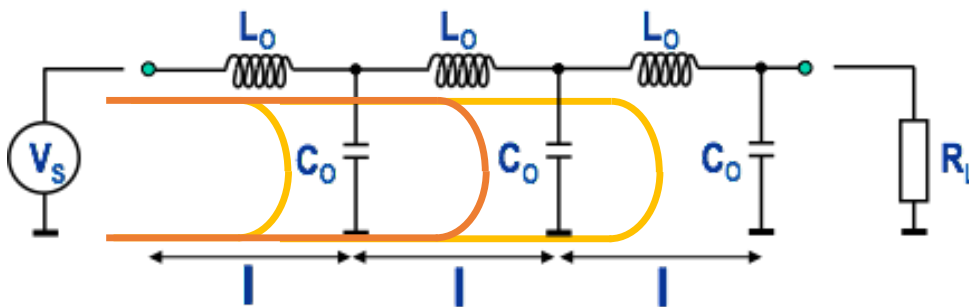


Figura 1.17. Caminando junto al frente de ondas de la señal, paso a paso en cada etapa LC. Fuente propia.

De forma que vamos caminando junto a señal, en cada paso nos encontramos con un nuevo y pequeño condensador (que representa la energía almacenada en el campo eléctrico) que comienza a descargar ($I = C \cdot \delta V / \delta t$). Para ser más realistas, añadimos una inductancia entre cada dos condensadores que representa la energía almacenada en el campo magnético creado por la corriente eléctrica. El **flanco** del escalón será más abrupto cuanto más corriente puede entregar la fuente, y ocupará varios milímetros o centímetros de la línea, propagándose siempre a la velocidad de la luz en el medio.

Caminar junto a la señal es arriesgado: ¡no ves lo que hay delante!

Nada puede viajar más rápido que la luz. Para ser exactos, debemos decir que nada puede viajar más rápido que la luz en el vacío, pero sí más rápido que la luz en un medio como el agua (cuando ocurre esto último, se produce una luz azulada, llamada luz de Cherenkov, que es responsable de que las fotografías del núcleo de reactores nucleares tengan un tono azulado y fue responsable de un espectáculo de luces azules en el cielo sobre el reactor 4 de la central de Chernóbil durante el fatídico accidente, pero dejemos esta digresión y volamos al tema: la propagación del frente de ondas).

Tampoco la información puede viajar más rápido que la luz en el vacío, de modo que el frente de la señal no “sabe” qué habrá 10 cm o 1 m por delante. Tampoco tú, que estás caminando a su lado. Ni el frente ni tú sabéis si un centímetro más adelante la línea acaba en 50 ohmios o en circuito abierto, si se prolonga dos metros más, o si llegará hasta la luna. Ni si en la luna la línea termina en 1 MΩ o en 1 kΩ.

¡Herejía! Desde primero de carrera sabemos que la fuente debe entregar una corriente tal que el cociente V/I sea igual a la carga, Z . La impedancia, ¿no? Pues no. Eso es sólo en estado estacionario, cuando ya ha pasado todo lo interesante, cuando ya no hay cambio. Es lo que ocurre el continuo y en baja frecuencia, es decir, en electricidad y en electrónica de baja frecuencia. No es que te engañaran, es que en primero de carrera considerabais unas condiciones sencillas.

Lo único que conoces a medida que caminas por la línea es un parámetro, que llamaremos **impedancia de línea** (V/I en el frente de la onda, que, te lo creas o no, coincide con el **cociente E/H** entre campo eléctrico y magnético), que no depende de la carga (impedancia) en el extremo de la línea. Porque la carga puede estar en la luna y de momento, no hay forma en este universo de que eso te afecte. Porque no puedes viajar en el tiempo: sólo veremos la carga cuando lleguemos a ella.

Mientras tanto, mientras caminas junto a la señal camino a la luna, son la geometría de la pista de circuito impreso y la permitividad del medio quienes determinan el cociente V/I (la impedancia de línea). De ahí que hablemos de coaxiales de 75 ohm o de pistas de circuito impreso de 50 ohm.

Explicaremos esto con más detalle en el segundo día. Abordemos ahora la última sección de hoy. Sí, está siendo un día espeso, pero nada te impide volver a revisar este inicio de curso mañana. Recuerda: no hay prisa. Avancemos un pequeño paso cada día.

Campo cercano y campo lejano

Vale. Las cargas aceleradas (basta una corriente eléctrica sinusoidal para ello, no busques mecanismos complejos ni extraños) producen radiación electromagnética (revisa lo que dijimos en torno a la Figura 1.6). Esta radiación puede propagarse por el espacio libre o por estructuras guía como cables coaxiales o pistas de circuito impreso. En cualquier caso, algo suele escapar de la estructura y se propaga por el espacio circundante.

Es importante aclarar un par de conceptos relacionados con la propagación de campos por el espacio. Lo primero que hay que entender es que:

Las fuentes de radiación son bien campo eléctrico, bien de campo magnético

Una espira, bucle, lazo produce radiación de campo magnético. Cerca de la espira, medimos mucho campo magnético (H) pero muy poco campo eléctrico (E). El cociente E/H es muy pequeño. En la página anterior hemos dicho que el cociente E/H era la **impedancia de línea**. Bien, en el espacio libre se usa el término **impedancia de onda**. Es lo mismo. Por tanto, cerca de una fuente de radiación de campo magnético, la impedancia de onda es baja.

Una antena de varilla (como la que está en tu coche) produce radiación de campo eléctrico. Cerca de la antena medimos mucho campo eléctrico (E) pero poco magnético (H). El cociente E/H es grande. Por tanto, cerca de una fuente de radiación de campo eléctrico, la impedancia de onda es alta.

Lejos de las fuentes de radiación, E/H es la impedancia de la onda para el medio. En el vacío o en el aire es de 377Ω

Cerca de la fuente de radiación, midiendo el cociente E/H, podemos saber si la fuente es de tipo bucle o de tipo varilla. A medida que nos alejamos de la fuente, el cociente E/H tiende a ser una característica del medio de propagación. Que en el espacio libre es de 377Ω .

Denominamos campo cercano a la región del espacio cerca de la fuente de radiación. Esta región se extiende hasta aproximadamente $\lambda/2\pi$ de la fuente

Esta es una región del espacio con características especiales. **Lo que coloquemos o hagamos en esta región afecta a la fuente de radiación.** Se almacena energía no radiante (se dice que es **campo reactivo**, para distinguirlo del **campo activo**, el que se propaga lejos de la fuente).

Déjame poner un ejemplo. Busca un tubo fluorescente de los que suele haber en el techo de la cocina en muchas casas. Busca una línea de alta tensión y colócate debajo. Si mantienes el tubo vertical, verás cómo se ilumina.

¿Qué ha pasado? Les estás robando energía a la línea de alta tensión, estás alterando (poco, pero estás haciéndolo) los parámetros V, I de la línea. **¿Aunque no la toques?** Pues sí, un abogado no tiene por qué saberlo, pero tú deberías: técnicamente, podrían multarte porque estás robando energía de la línea.

A 50 Hz, $\lambda/2\pi \approx 955 \text{ km}$. Así que tranquilo, todos estamos alterando el campo cercano de la línea de alta tensión. Pero si hablamos de frecuencias más habituales en radiocomunicaciones, a 2.4 GHz (donde tienes WiFi, Zigbee y Bluetooth), $\lambda/2\pi \approx 20 \text{ mm}$. Es decir, a menos de 2 cm de la antena, lo que pongamos o hagamos afecta a la transmisión (desintoniza la antena y aumenta o disminuye la ganancia de la antena). De hecho, evitamos elementos metálicos (taladros de fijación o incluso planos internos en el PCB) en las cercanías de una antena WiFi para evitar desintonizarla o cambiar sus características de radiación.

Denominamos campo lejano a la región del espacio lejos de la fuente de radiación. Esta región se extiende aproximadamente desde $2D^2/\lambda$ de la fuente, o en todo caso 2λ , siendo D la dimensión mayor de la antena

En esta región del espacio libre de campo activo, la antena no almacena energía y lo que coloquemos o hagamos no afecta a la fuente de radiación. La antena WiFi no cambiará su frecuencia de emisión ni verá su ganancia alterada por lo que tu hagas unos centímetros más allá de la antena (dos centímetros y medio para una antena de 4 cm de longitud).

Este hecho se usa legalmente para obtener energía de las ondas radioeléctricas presentes en el ambiente, almacenarla y alimentar nodos sensores de muy bajo consumo. Esta técnica, conocida como *energy harvesting*, es bien conocida y tranquilo, totalmente legal. ¡Estás en campo lejano!

Entre las zonas de campo cercano y lejano hay una zona de transición en la que la impedancia de onda va adaptándose a 377Ω

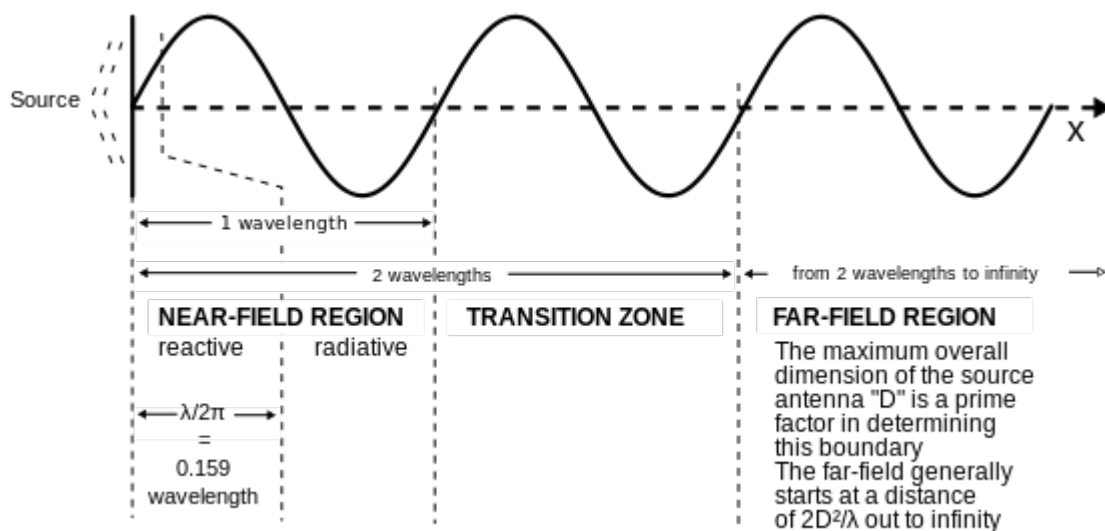


Figura 1.18. Fuente: [Wikipedia](#). Imagen de dominio público.

¿Qué importancia tiene para nosotros saber si estamos en campo cercano o lejano?

Simplificado, saber si la impedancia de onda es alta o baja. Porque si queremos **apantallar** un cable o un equipo frente a **interferencias**, escoger el material coloquemos como pantalla y su espesor requiere conocer si la impedancia de onda es alta o baja.

No te dejaré así, te daré una explicación breve. **¿Por qué la luz visible se refleja en un espejo?** Porque la delgada capa de metal pulido que hay tras el vidrio tiene una baja impedancia característica, bastante menor de 1Ω (E alto, H bajo). El aire es un medio con impedancia característica de 377Ω . Una diferencia de impedancias tan elevada hace que se refleje la luz.

De igual modo, una pantalla (caja metálica que rodea un equipo electrónico, también vale una caja de plástico con pintura conductora o incluso, para hacer pruebas, papel de aluminio) tiene una impedancia característica muy pequeña. Al incidir una onda electromagnética de impedancia alta (en campo lejano, aunque también cerca de una fuente de campo eléctrico), se produce una reflexión elevada: ¡estamos apantallando nuestro equipo! Es decir, evitamos que penetre la energía de la onda dentro de la caja y afecte a nuestro PCB (circuito impreso).

En campo cercano de fuentes de campo magnético, que son de impedancia baja, este principio no funciona y habrá que recurrir a otro principio, la absorción. Ya lo estudiaremos el día 12, no tengamos prisa.

Hemos alcanzado el final del primer día, completado el repaso de algunas cosas que tal vez no te contaron o que no recordabas sobre electromagnetismo. Lo normal es que si es la primera vez que te encuentras con estos conceptos necesites volver a leer el texto un par de veces, consultar las obras referenciadas o preguntar a quienes han creado un conjunto de simplificaciones y analogías útiles para creer que entienden la porción de la Naturaleza con la que trabajan (es decir, los expertos).

Tómate un descanso y si aún no has decidido que lo tuyo es otra cosa, como la contabilidad o el marketing, acompáñame en el segundo día del curso.

Día 2. Pistas en circuitos impresos

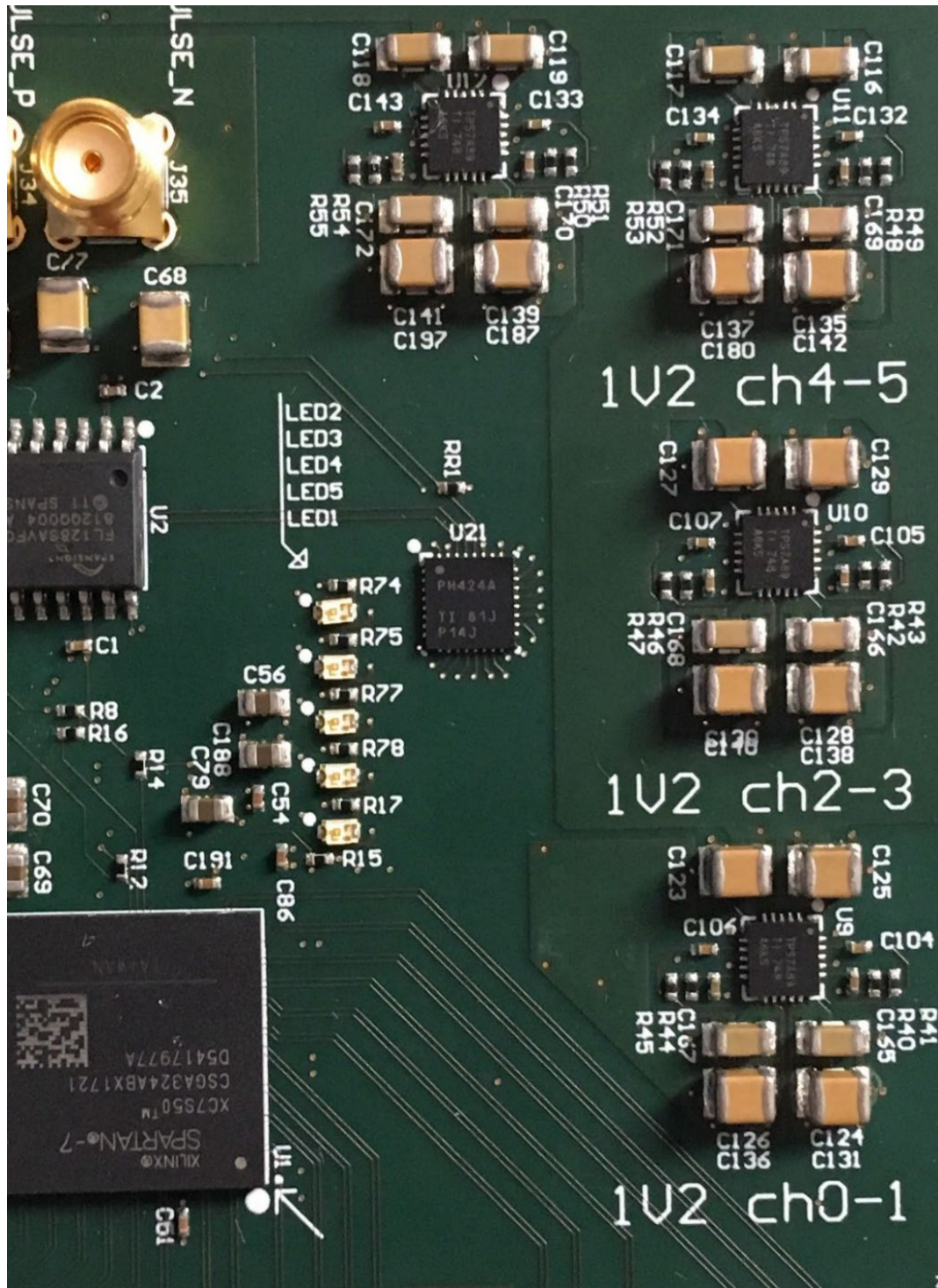


Figura 2.1. Fuente: elaboración propia (detalle de un diseño realizado por el autor en 2018)

Un circuito electrónico se compone básicamente de circuitos integrados, componentes pasivos e **interconexiones** que los comunican entre sí y con el exterior. Es lo que puedes ver en la fotografía. Un producto electrónico contiene otros subsistemas como pueden ser alimentación, refrigeración, envoltorio e interfaz con el usuario.

Las interconexiones constituyen el elemento que tiene un mayor impacto tanto en integridad de señal como en EMC. Esto es debido a que:

1. **Las interconexiones suelen ser la principal fuente de radiación** electromagnética a media frecuencia (cables de entrada/salida) y alta frecuencia (pistas en **placas de circuitos impresos -PCBs**). Es decir, se comportan como antenas. La pista en una PCB crea, junto con el camino de retorno (recuerda, hemos hablado sobre esto en el capítulo anterior) una espira, un bucle, que radia al entorno. Como diseñadores, está en nuestra mano minimizar la radiación siguiendo unas pocas consideraciones prácticas. La más importante de todas: cuidar los caminos de retorno. No te preocupes si todavía no entiendes bien este concepto, le dedicaremos mucho tiempo hoy.
2. Por reciprocidad, lo que es eficiente emitiendo campos electromagnéticos también lo es recibiendo. Por tanto, **las interconexiones suelen ser la principal vía de entrada de radiación** electromagnética no deseada de alta frecuencia en nuestro producto.
3. **Una interconexión mal diseñada puede provocar distorsión en la señal** (mala integridad de señal) y mayor radiación. Y recuerda, no sólo has de cuidar la pista, sino también el plano por el que circula la corriente de retorno.

Las interconexiones están presentes en un producto electrónico en forma de cables (cables planos, cables coaxiales, pares trenzados) y estructuras en el circuito impreso (pistas y cavidades). Hoy trataremos fundamentalmente con pistas en el circuito impreso, dejando para días posteriores el estudio de los otros tipos de interconexiones. Creo necesario hacer una consideración en este momento.

¿Qué vamos a aprender hoy?

En los primeros cursos de universidad, es posible que montaras en placa de prototipos un pequeño sistema digital como el de la figura. Seguimos guiando los campos electromagnéticos entre dos conductores, aunque no resulte evidente.

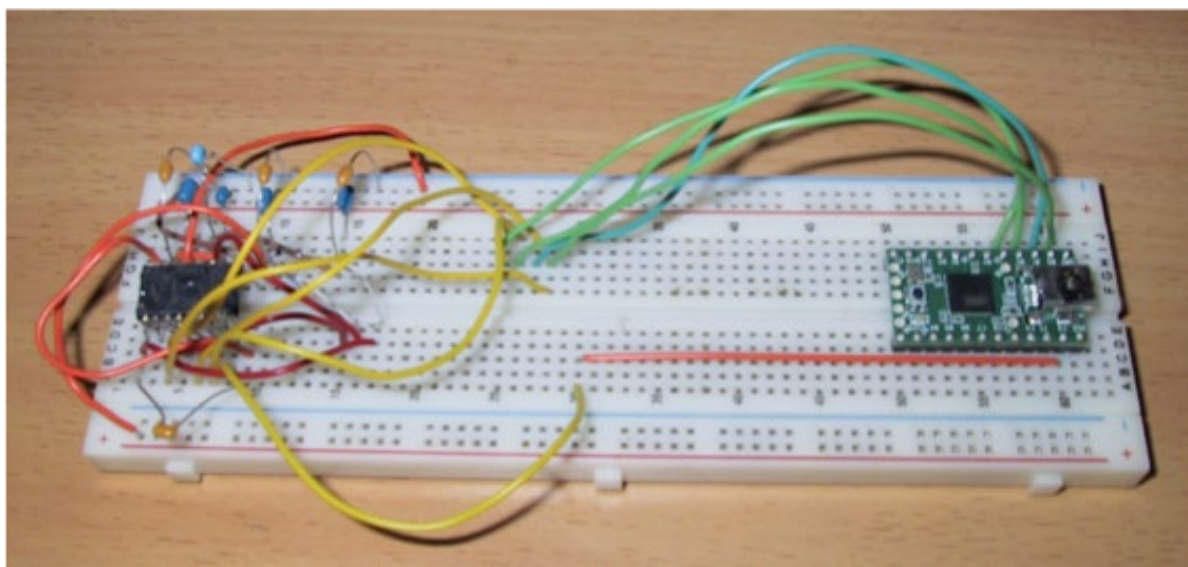


Figura 2.2. Típico montaje que funcionaba bien cuando estabas al inicio de tus estudios pero que raramente funcionará en un diseño real. ¿Por qué?

Si has intentado hacer algo parecido recientemente, usando componentes como los que encontrarías en un producto comercial, es fácil que el circuito falle. Y te habrás preguntado: *“Este tipo de cosas funcionaba en primero de carrera... ¿Por qué no funcionan ahora que realmente lo necesito?”*. Ya intuirás la respuesta, pero permite que dejemos aparcada un rato la discusión.

Vamos a comenzar el capítulo presentando los **conceptos mínimos de líneas de transmisión** que debes conocer. Recuerda de ayer que habíamos definido una línea de transmisión, en oposición a una interconexión ideal, como aquella que no es eléctricamente corta y por tanto no podíamos considerar que todos los puntos de línea están al mismo potencial eléctrico (si lo prefieres, no todos los puntos de la línea están en la misma fase de la señal). Esto puede tener un profundo impacto en la forma en la que se comporta tu circuito.

Hoy presentaremos uno de los dos efectos principales que se darán en las pistas de tu PCB: la **reflexión**.

A continuación, presentaremos los dos tipos de líneas de transmisión que usamos para diseñar PCBs: líneas en capas externas (**microstrip**) y en capas internas (**stripline**). Aprenderemos a diseñarlas, teniendo en cuenta las limitaciones que suelen imponer los fabricantes de PCBs y los objetivos que buscamos.

¿Qué vamos a aprender en los próximos días?

Dejaremos para el **tercer día** las soluciones asociadas a la problemática de las reflexiones en pistas en un PCB.

En el **cuarto día** abordaremos el estudio de las líneas diferenciales.

En el **quinto día** hablaremos de cómo radia una pista o un cable y expondremos técnicas para reducir este efecto.

Una interconexión requiere dos conductores...

Porque la idea es guiar, confinar la propagación de un campo electromagnético entre dos conductores. Ya vimos ayer que un campo electromagnético puede propagarse en el vacío o en un medio sin necesidad de conductores. Vimos el ejemplo de una fibra óptica, en la que confinamos la onda a base de hacer que rebote en las paredes interiores sin llegar a salir de la fibra. Pues bien, en electrónica, usamos el truco de confinar, guiar, conducir la onda entre dos conductores.

En el caso de un circuito impreso, dos planos adyacentes (sin importar su nivel de continua, es irrelevante para la alta frecuencia, aplicamos el principio de superposición que estudiaste en análisis de circuitos) permiten la propagación de un campo electromagnético en una estructura llamada **cavidad** (Figura 2.3). No es algo que hagamos intencionadamente: excitamos la cavidad cuando inyectamos ruido en los planos, por ejemplo, cuando una vía de señal atraviesa la cavidad y actúa en su interior como una pequeña antena. La propagación de esta perturbación se acopla a otras vías, llegando hasta las capas externas del PCB o bien simplemente se propaga hasta los bordes del PCB. En ambos casos se radia al entorno, por lo que deberemos aprender cómo reducir estos efectos.

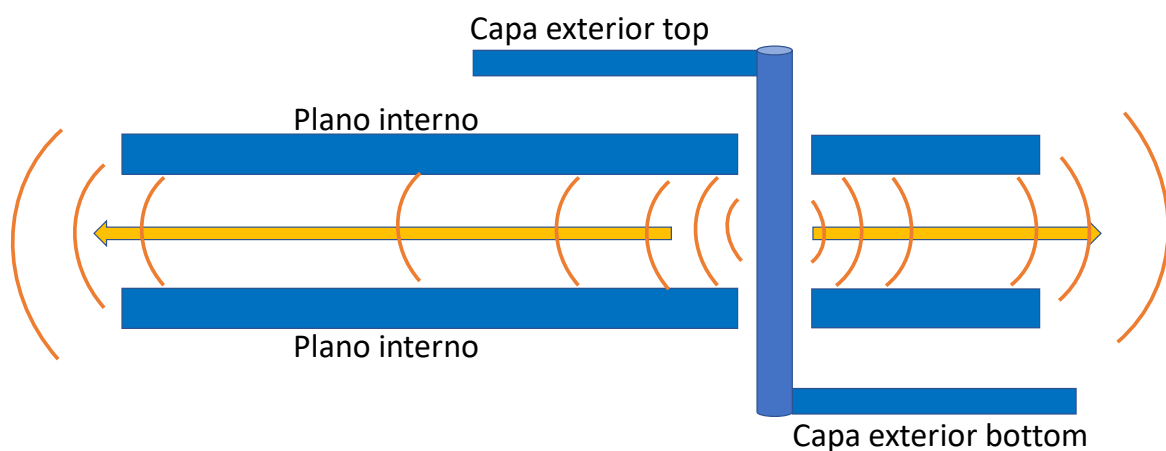


Figura 2.3. Dos planos en el interior de un PCB forman una estructura en la que pueden propagarse señales de (en la práctica) cualquier frecuencia. ¿Cómo se inyecta una señal en esta estructura? Ruido en los planos y ruido acoplado por una vía que atraviesa la estructura suelen ser dos de las principales causas. Cuando esta perturbación se propaga y llega al borde del PCB, parte se refleja al interior y parte es radiada al exterior. Fuente propia.

En una cavidad, ambos conductores tienen en principio en mismo o similar tamaño y no ocurre nada distinto en uno que en el otro, la distribución de tensiones y corrientes es la misma o similar, de modo que no hay razón para dar un nombre especial a cada plano conductor.

Justo lo contrario ocurre en una pista de circuito impreso, un tipo de interconexión donde llamamos **pista** al conductor de cobre delgado (17-35 micras de espesor típicamente) y estrecho (100-300 micras de anchura típicamente) y llamamos **plano de referencia** al también delgado (17-35 micras típicamente) pero ancho (generalmente centímetros) plano (en caso de pista en capa externa) o planos (en el caso de pista en capa interna) adyacentes.

En la Figura 2.4 (izquierda) encontramos el tipo de línea de transmisión típico de capas externas: la **línea microstrip**. La impedancia de línea queda determinada por la anchura de pista (w), altura del dieléctrico (h), constante dieléctrica del material y en menor medida por el espesor de la pista (t). Hablaremos más sobre líneas microstrip en la página 52.

En la Figura 2.4 (derecha) encontramos el tipo de línea de transmisión típico de capas internas: la **línea stripline asimétrica**. Se llama así porque $h_1 \neq h_2$. La impedancia de línea queda determinada por la anchura de pista (w), alturas de los dieléctricos (h_1 , h_2), constante dieléctrica del material y espesor de la pista (t). Hablaremos más sobre líneas stripline en la página 53.

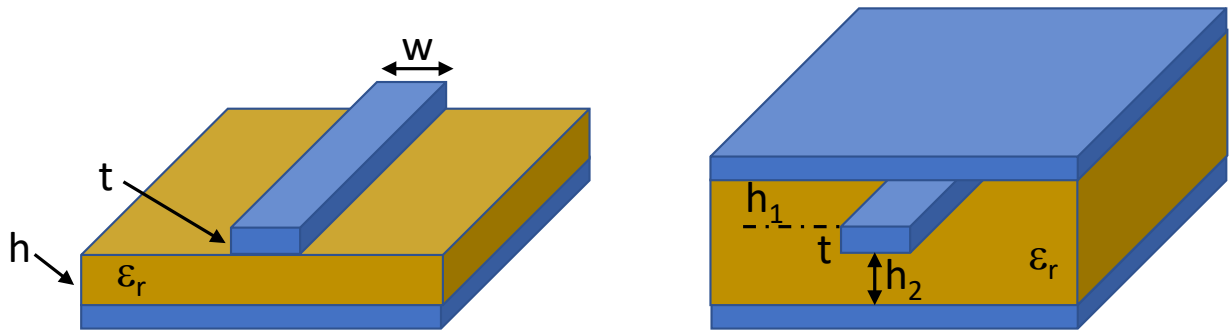


Figura 2.4. Pista en capa externa referida a un plano (línea microstrip) y pista en capa interna referida a dos planos (línea stripline). La geometría y la constante dieléctrica del material determinan la impedancia de línea. Fuente propia.

Ayer hablábamos de cómo se propaga una señal por una línea (un par de conductores). Recuerda que podíamos considerar una pequeña sección de la línea, en una primera aproximación, como un pequeño condensador. El frente de onda que se propaga carga uno detrás de otro estos pequeños condensadores. Y esto requiere corriente eléctrica simultáneamente tanto en un conductor (pista) como en el otro (plano de referencia). Recuerda que no hay circulación de corriente a través del dieléctrico, sino que se trata de corriente de desplazamiento, como en un condensador.

Pero aquí hay una asimetría: la corriente en la pista sólo puede circular por la geometría del conductor, está confinada, pero ¿por dónde circula la corriente en el plano de referencia? ¿Por toda la anchura del plano? La Naturaleza dice que no.

(dos conductores)... que forma un bucle con una determinada inductancia

La Naturaleza dice no a que la **corriente de retorno** circule por todo el plano de referencia porque:

1. un bucle tiene una inductancia proporcional a su área
2. la impedancia de un bucle es directamente proporcional a su inductancia (y a la frecuencia)
3. la corriente prefiere circular por el camino de menor impedancia (en este caso, inductancia)

Como conclusión de los tres puntos anteriores, la densidad de corriente en el plano de referencia (ya la hemos bautizado hace unas líneas: **corriente de retorno**) es máxima justo bajo la pista (forma el bucle de área mínima y por tanto de menor impedancia) y decrece según una ley cuadrática inversa a medida que nos alejamos de ésta.

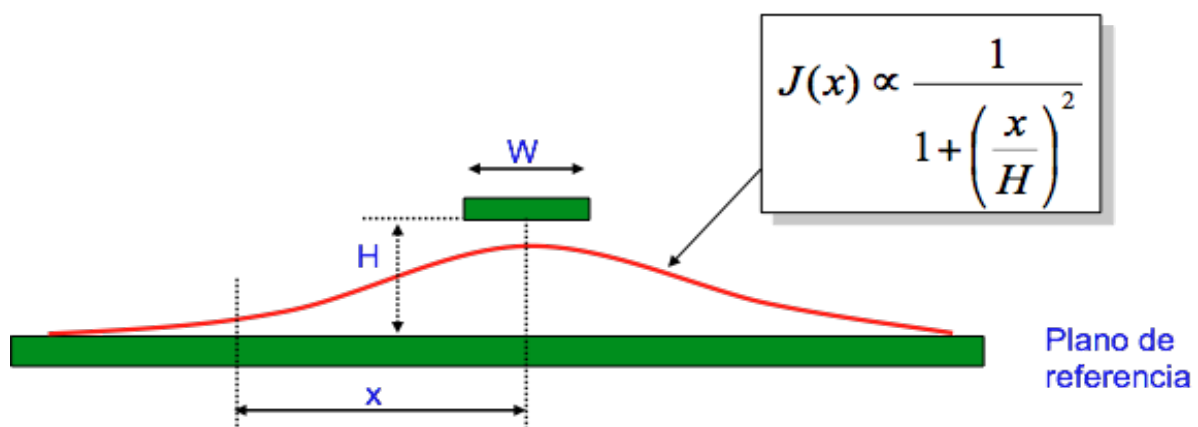


Figura 2.5. Distribución de la corriente de retorno en el plano de referencia en una línea microstrip. Fuente propia.



Figura 2.6. Camino de la corriente de retorno en alta y baja frecuencia. Fuente propia.

En continua (y muy baja frecuencia), el camino de menor impedancia coincide con el de menor resistencia, y como consecuencia la corriente retorna por un área amplia.

En frecuencias más altas, ya a partir de unos pocos kHz, el camino de menor impedancia coincide con el de menor inductancia, que es el de menor área de bucle: la corriente retorna siguiendo fielmente el camino de la pista, con una distribución de corriente que sigue lo indicado en la Figura 2.5.

Lo mínimo que deberías saber sobre líneas de transmisión

Una línea de transmisión no es una conexión ideal. Ya vimos en el capítulo anterior que, si la longitud de la línea no es pequeña comparado con la longitud de onda de la señal, no podemos considerar que todos sus puntos están al mismo potencial. Y eso trae consigo un par de desgracias (reflexiones y radiación), que contribuyeron a que, en la última década del siglo pasado, los diseñadores digitales tuvieran que aprender dos nuevas disciplinas de diseño: la **integridad de señal** y la **compatibilidad electromagnética**.

De modo que la diferencia es importante. Vas a trazar mentalmente una línea en el suelo y vas a colocar a un lado las pistas de tu diseño que consideres eléctricamente cortas, que se comportan casi como conexiones ideales. Al otro lado, las líneas de transmisión.

Las primeras no requieren ninguna atención por tu parte. Las segundas van a requerir que apliques lo que aprendas hoy y en los próximos dos días.

En 1999 estaba trabajando en mi tesis doctoral en el CERN, junto a Ginebra, en la frontera entre Francia y Suiza. Diseñaba partes de un sistema de adquisición de datos para un experimento de física de altas energías. Se trataba fundamentalmente de sistemas digitales formados por FPGAs, memorias y enlaces de E/S rápidos. Ese año ocurrió algo que supuso un punto de inflexión: un sistema digital que acababa de diseñar otro equipo de investigadores no funcionaba. Y no era debido a un error de diseño lógico: era un problema de Integridad de Señal (para ser más concreto, reflexiones en las líneas, lo que deforma las señales digitales produciendo falsos estados lógicos y falsos flancos de reloj, ya te imaginas su demoledor efecto en un sistema digital).

Por aquel entonces, la Integridad de Señal no era una disciplina conocida entre los diseñadores electrónicos que no trabajaran con radiofrecuencia o microondas. Los diseñadores digitales solíamos pensar que una pista era un camino conductor de cobre y punto. No se pensaba en que podía comportarse como algo distinto a una conexión más o menos ideal.

Ese hecho individual despertó en mí un interés enorme por aprender cómo incorporar la Integridad de Señal en mi metodología de diseño, por entender las causas, consecuencias y soluciones de los problemas que comenzaban a darse cada vez con mayor asiduidad en los nuevos diseños. Estando en el CERN, tuve acceso a herramientas avanzadas, pero la bibliografía disponible en el tema era muy escasa y tuve que aprender con cierta dificultad. Tú lo tienes más fácil. ¡Aprovechate!

¿Cómo sabemos si una conexión es una línea de transmisión?

Como consecuencia de lo anterior, lo primero que necesitas es un criterio claro para clasificar las pistas.

En una señal analógica suele ser fácil definir la longitud de onda de la señal. Por ejemplo, una señal WiFi estará centrada en unos de los 14 canales disponibles (el primero en 2,412 GHz, el último en 2,484 GHz) y tiene un ancho de banda de 20 MHz. Dado que el ancho de banda es pequeño comparado con la frecuencia del canal, puedes hacer la aproximación de suponer que toda la señal tiene la misma frecuencia. Por ejemplo, 2,412 GHz si estamos usando el canal 1. Calcular la longitud de onda (λ) es fácil teniendo en cuenta dos simples expresiones:

$$\lambda = \frac{c}{f}, \text{ siendo } c \text{ la velocidad de la luz en el medio y } f \text{ la frecuencia de la onda}$$

$$c = c_0 / \sqrt{\epsilon_r}, \text{ siendo } c_0 \text{ la velocidad de la luz en el vacío y } \epsilon_r \text{ la permitividad relativa al vacío}$$

En un PCB la permitividad relativa es cercana a 4. Para 2,412 GHz, la longitud de onda es $\lambda \approx 6,2 \text{ cm}$. En una pista menor que $\lambda/10 \approx 6 \text{ mm}$, los efectos de línea de transmisión serán poco acusados. Esto no quiere decir que no los haya, y por eso evitarás cambios de impedancia en el camino de la señal (es decir, entre la salida del RF *transceiver* y la antena o el conector de antena).

En una señal digital es más difícil determinar la longitud de onda asociada. Se trata de una secuencia más o menos aleatoria de pulsos más o menos trapezoidales con flancos de subida y de bajada abruptos,

normalmente por debajo del nanosegundo. Si representamos su espectro en frecuencia (Figura 2.7) observamos lo siguiente:

1. El espectro de la señal se extiende desde continua hasta frecuencias mucho mayores que la de la tasa binaria ($1/T$)
2. La amplitud de las componentes frecuenciales decrece como $1/f$ (en la Figura 2.7, como 20 decibelios por década, es decir, al aumentar 10 veces la frecuencia resulta $20 \cdot \log(1/10) = -20$ dB) hasta aproximadamente una frecuencia denominada frecuencia de codo $1/(\pi \cdot t_t)$. Curiosamente, en inglés se le denomina “*knee frequency*”, pero como “frecuencia de rodilla” suena mal, lo adaptamos como “frecuencia de codo”. Por cierto, t_t es el menor entre el tiempo de subida y de bajada de los flancos, definido como el tiempo de transición entre el 10% y el 90% de la amplitud del pulso.
3. A partir de la frecuencia de codo, la caída de amplitud de los armónicos es mayor, según $1/f^2$ (es decir, 40 dB/década, ya que $20 \cdot \log(1/100) = -40$ dB). Podemos decir por tanto que la mayor parte de la energía de la señal está comprendida entre continua y la frecuencia de codo.

La obsesión en ingeniería con expresar cambios en decibelios viene de antiguo y te supongo ya acostumbrado/a.

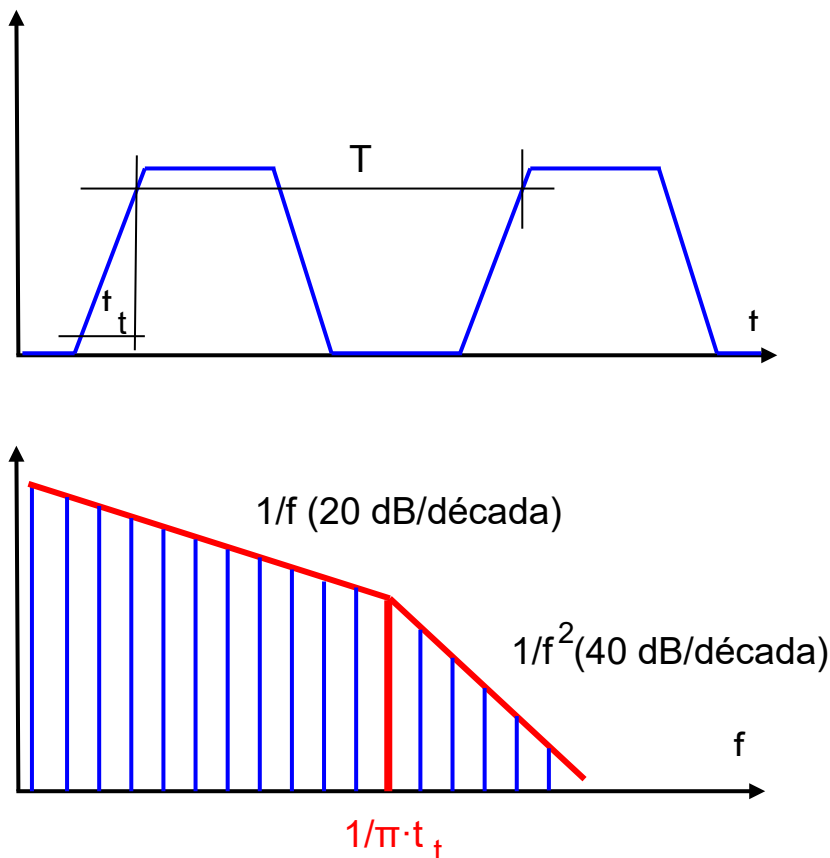


Figura 2.7. Espectro en frecuencia de una señal digital. Aprende bien esto: el ancho de banda no guarda relación con la frecuencia del tren de pulsos, sino con lo abruptos que sean los flancos.

La frecuencia de codo

De modo que el ancho de banda práctico de una señal digital se extiende hasta la frecuencia de codo. Aunque $1/\pi \approx 0,32$, para curarme en salud prefiero usar la siguiente expresión más conservadora donde, recordemos, t_t es el menor entre el tiempo de subida y de bajada de los pulsos:

$$f_{codo} \approx 0,5/t_t$$

¿Cómo es posible que el ancho de banda práctico de una señal digital no dependa directamente de la tasa binaria sino sólo de la duración del flanco? Parece contraintuitivo. Si te paras a pensarlo, el flanco presenta un cambio más rápido que la duración del pulso. Luego es lógico que sea determinante en el ancho de banda de la señal.

De aquí es inmediato derivar una expresión para la longitud de onda asociada a la frecuencia más alta que vamos a considerar para la señal digital:

$$\lambda_{codo} \approx c/f_{codo}$$

La longitud crítica

Limitando la longitud máxima de la línea a la décima parte de este valor, resulta la siguiente expresión para la **longitud crítica** de una pista en **centímetros**, para una señal con un tiempo de flanco expresado en **nanosegundos**, que de ser excedida consideraremos que se trata de una línea de transmisión:

$$L_c = \lambda_{codo}/10 \approx 3 \cdot t_t$$

Por ejemplo, si el flanco es de medio nanosegundo, una pista superior a 1,5 cm no puede ser considerada como eléctricamente corta, debemos ponerla en la lista de sospechosas y evaluar los efectos de línea de transmisión que pueden ocurrir en ella.

Un conductor de 1,5 mm (¡vaya! este es el espesor aproximado de buena parte de los circuitos impresos, de modo que una vía pasante tiene esta longitud) es potencialmente problemático para flancos de 50 ps o menos. Encontrarás flancos tan cortos en conexiones de E/S multigigabit (10 GbE, PCIe o SATA).

¿Cómo podemos asignar un significado físico a este resultado?

Si t_{pd} es el retardo de propagación en s/m, es decir, la inversa de la velocidad de propagación resulta que la longitud crítica se alcanza cuando el tiempo que tarda la señal en propagarse hasta la carga es mayor que 1/5 de la duración del flanco.

La Tabla 2.1 recoge las longitudes críticas para distintas familias lógicas, suponiendo $t_{pd}=60$ ps/cm, lo que es un valor cercano al que encontrarás en una pista de PCB. Cuando hice prácticas de electrónica digital en segundo de carrera (en 1991), usábamos integrados de la familia TTL-LS. Según la tabla, cualquier línea más corta de 33 cm se comportaba “casi” como una conexión ideal; las reflexiones y otros efectos indeseados eran casi despreciables.

Longitud crítica:

$$L_c \cdot t_{pd} = \frac{t_t}{5} \Rightarrow L_c = \frac{t_t}{5 \cdot t_{pd}}$$

Relación con el ancho de banda de la señal:

$$\lambda_{codo} = \frac{1}{t_{pd} \cdot f_{codo}} = \frac{t_t}{0,5 \cdot t_{pd}}$$

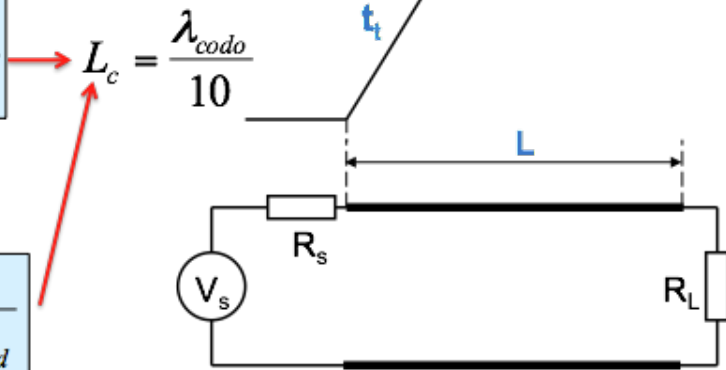


Figura 2.8. Cómo determinar si una línea de transmisión debe considerarse crítica o no. Fuente propia.

No es de extrañar que los montajes con cablecillos en placa de prototipos funcionaran bien y que el único reto fuera hacer bien la simplificación por tabla de Karnaugh y no equivocarse al hacer las conexiones.

Tabla 2.1. Tiempos de flanco y longitud límite para varias familias lógicas CMOS

Familia	t_r (ns)	t_f (ns)	$L_{límite}$ (cm)
LS	14	10	33,3
ALS	2.7	1.7	5,7
FAST	4.0	1.4	4,7
LVC	1.8	1.8	6
ALVC	1.2	1.1	3,7
LVT	0.8	0.6	1,8
ALVT	0.8	0.7	2,3
HC	2.9	2.9	9,7
AHC	2.1	1.6	5,3

Pero las longitudes críticas para familias CMOS más actuales (sombreadas en gris en la tabla) están por debajo de los 2-3 cm. Seguro que en estas condiciones las máquinas de estados y contadores que montaba en segundo de carrera fallarían de tanto en tanto, muy posiblemente debido a distorsión en las líneas de reloj. Es decir, en clase hacíamos trampa sin saberlo al usar familias lógicas con flancos lentos, perpetuando los felices años 80 del siglo pasado, donde las conexiones se consideraban ideales. ¡Qué tiempos aquellos!

Por cierto, verás que los resultados de la tabla son compatibles con la aproximación que dábamos en la página anterior para la longitud crítica (L_c) en centímetros, con los tiempos de transición (t_t) en nanosegundos:

$$L_c \approx 3 \cdot t_t$$

La Figura 2.9 muestra un ejemplo de línea corta y de línea larga. A la izquierda, una línea de 0,2 pulgadas (recuerda: una pulgada son 2,54 centímetros), a la derecha una línea diez veces más larga. En un extremo de la línea, un *driver* LVCMOS de 3,3V, rápido y de 16 mA de corriente de salida. A la derecha, una carga de unos pocos picofaradios (un modelo sencillo para representar un *buffer* de entrada digital).

Con un tiempo de transición rondando el medio nanosegundo, la longitud crítica es de aproximadamente 1,5 cm (0,6 pulgadas). Por tanto, en la figura de la izquierda, donde la línea es de 0,2 pulgadas de longitud, las formas de onda (rojo para el *driver*, verde para la carga) están muy poco deformadas. En cambio, en la figura de la derecha, donde la línea es de 2 pulgadas (más del triple de la longitud crítica), las formas de onda están gravemente deformadas por efecto de las reflexiones (**explicaremos este fenómeno a partir de la página 47**).

Observa de nuevo las formas de onda roja (driver) y verde (carga) de la figura de la derecha. ¿Queremos que las dos estén limpias? ¿Sólo una? ¿Cuál? Piensa un minuto antes de leer la respuesta en la página siguiente.

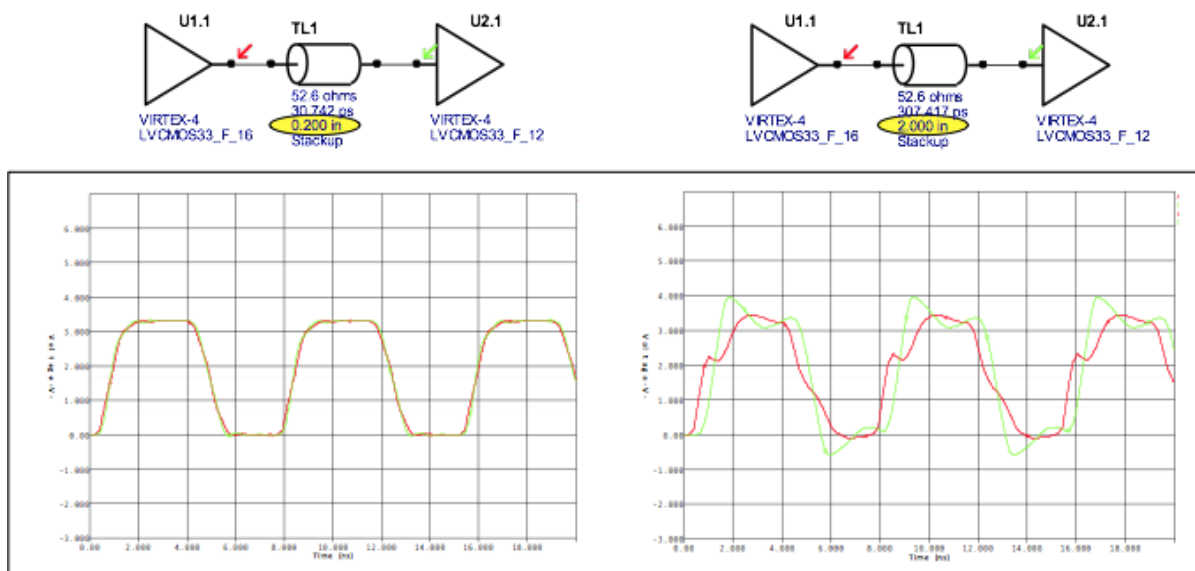


Figura 2.9. Parámetros: $t_r = 487$ ps; $f_{\text{codo}} = 1$ GHz; $t_{pd} = 60,4$ ps/cm; $L_c = 1,61$ cm (0,64 pulgadas). Fuente propia.

Para terminar este apartado, te voy a hacer una pregunta sobre el diseño de la Figura 2.10, en el que todas las líneas van conectadas a la FPGA (el encapsulado metálico en la parte superior de derecha de la placa). ¿Crees que hay líneas que podamos considerar eléctricamente cortas?

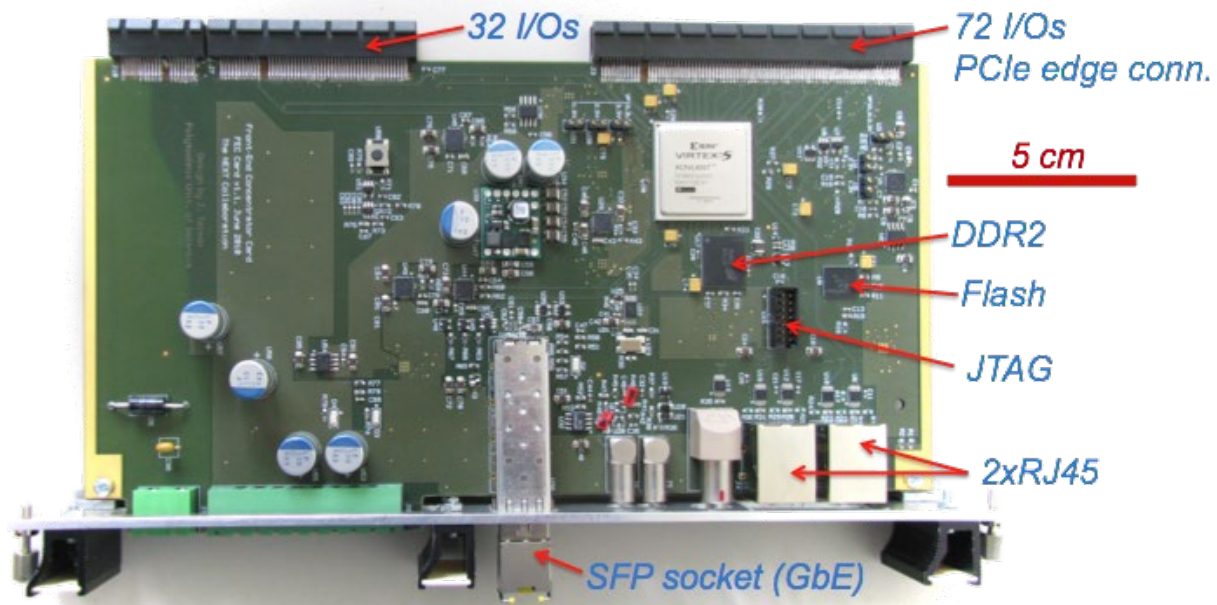


Figura 2.10. ¿Crees que hay líneas no críticas en este diseño? Fuente propia

La respuesta es un clamoroso no. **Espera, ¿qué hay del bus JTAG que no irá a más de 1 MHz? ¿No se trata de un bus lento?** Deberías saber ya que el contenido en frecuencia de una señal digital no depende la frecuencia del tren de pulsos sino de lo abrupto que sean los flancos. Recuerda, $f_{codo} \approx 0,5/t_t$. De modo que incluso un bus lento presentará efectos de línea de transmisión si su *driver* es rápido. Así que me temo que no hay líneas cortas, quitando aquellas inferiores a 2-3 cm y que además sean de drivers no muy rápidos.

El bus DDR2 funciona a 800 Mb/s y sus líneas tienen flancos de tal vez un par de cientos de picosegundos. El ancho de banda equivalente (frecuencia de codo) se situará en torno a 2,5 GHz y por tanto la longitud crítica estará en torno a los 6 mm. El par de centímetros que puede haber entre la FPGA y la memoria hacen de las líneas de este bus líneas de transmisión donde cualquier error en su trazado puede hacer que la memoria no funcione. Ya iremos detallando a lo largo del curso a qué me refiero con “**error en el trazado**”. Afortunadamente, para DDR2, DDR3, DDR4 y otras variantes de memorias rápidas existen reglas de diseño muy detalladas a modo de lista de comprobación (*checklist*) que si seguimos nos ayudan a no cometer errores.

En cambio, buses lentos como JTAG que no operan típicamente por encima de 1 Mb/s (si no sabes lo que es JTAG echa un rápido vistazo a <https://es.wikipedia.org/wiki/JTAG>) suelen sufrir el descuido de los diseñadores. Lo habitual es que sufran de exceso de reflexiones en la línea de reloj, produciendo dobles flancos o simplemente flancos monótonos (con cambios de pendiente).

Volvamos a la pregunta sobre la línea de dos pulgadas. Lo cierto es que desde el punto de vista de la integridad de señal sólo nos interesa que la señal esté limpia en la carga (forma de onda verde), que es donde queremos que los unos se interpreten correctamente como unos, los ceros como ceros y que los flancos de reloj definan con precisión el instante temporal de muestreo.

Si observamos la forma de onda verde, claramente la señal está por encima de 2 V para el uno lógico, por debajo de 0,8 V para el cero lógico, los flancos son monótonos (sin cambios de pendiente) y no hay dobles flancos. Por tanto, la integridad de señal es buena.

Tan sólo cabe comentar que las reflexiones producen sobreimpulsos positivos y negativos que podrían estar estresando las protecciones del buffer de entrada de la carga, lo que podría dar lugar a un fallo por destrucción al cabo de días, semanas, meses o años de uso. Además, esta distorsión aumenta el contenido en altas frecuencias y por tanto la radiación de las pistas.

Parámetros básicos de una línea de transmisión

Ya sabemos que una línea de transmisión es una pista o cable de suficiente longitud como para no poder ser considerado una conexión ideal. La línea puede modelarse como un circuito formado por múltiples etapas LC como las de la figura. Cada etapa representa una longitud de la línea inferior a la crítica, tal vez la décima o veinteava parte de la longitud de onda asociada al ancho de banda de la señal. Los valores L_o y C_o de cada etapa se obtienen a partir del conocimiento de que la línea presenta:

- una autoinductancia por unidad de longitud L (H/m), y representa la energía almacenada por el campo magnético en torno a la pista
- una capacidad por unidad de longitud C (F/m), y representa la energía almacenada por el campo eléctrico entre la pista y el plano o planos de referencia

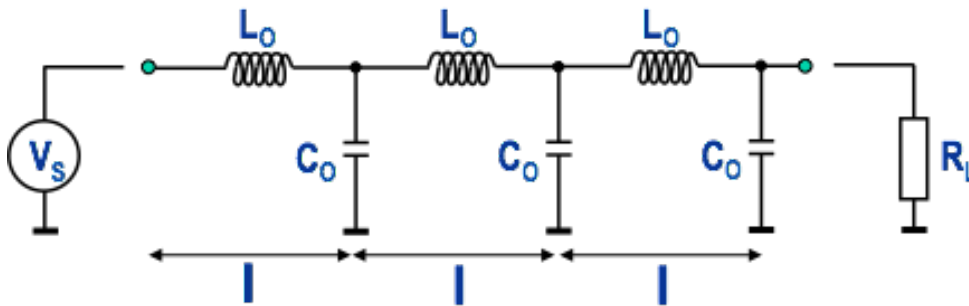


Figura 2.11. Modelo distribuido de una línea de transmisión sin pérdidas. Fuente propia

Estos valores determinan los dos parámetros más importantes de una línea de transmisión:

$$Z_o = \sqrt{L_o / C_o}$$

Impedancia característica de la línea (Ω)

$$t_{p_o} = \sqrt{L_o \cdot C_o}$$

Retardo de propagación (s/m)

En la práctica Z_o y t_{p_o} se obtienen a partir de la geometría de las pistas, utilizando programas de ordenador para realizar los cálculos. Valores típicos son 40-60 ohmios para Z_o y 60-70 ps/cm para t_{p_o} .

Podemos ignorar el efecto de inductancia y capacidad de la línea sobre señales estáticas o de lenta variación, pero no sobre el tipo de señal que nos interesa: los flancos de las señales digitales o señales analógicas de alta frecuencia.

Reflexiones

Análogamente a como la luz sufre un fenómeno de reflexión al cambiar de medio, una señal al propagarse por una pista sufre una reflexión si encuentra un cambio de impedancia en su camino. Este efecto debe considerarse sólo si la longitud de la línea es mayor que la crítica y sólo durante el transitorio.

Vamos a estudiar qué ocurre en el ejemplo de la Figura 2.12, con una fuente de impedancia de salida Z_S que genera un escalón de tensión V_S y una carga Z_L en el extremo de una línea de impedancia Z_O . **Consideramos Z_O , Z_L y Z_S distintas.**



Figura 2.12. Representación esquemática de una línea transmisión (tramo grueso) con una impedancia de carga Z_L y un driver con impedancia de fuente Z_S . Fuente propia

Primer paso: la fuente comienza a conducir la línea

Entre Z_S y Z_O se forma un primer divisor de tensión, por lo que sólo una parte de la tensión de la fuente se propaga hacia la carga. Usando valores usuales, con Z_S de unos 15 ohmios, Z_O de 50 ohmios y V_S de 3,3 V de amplitud, se propaga hacia la carga un flanco de unos 2,5 V.

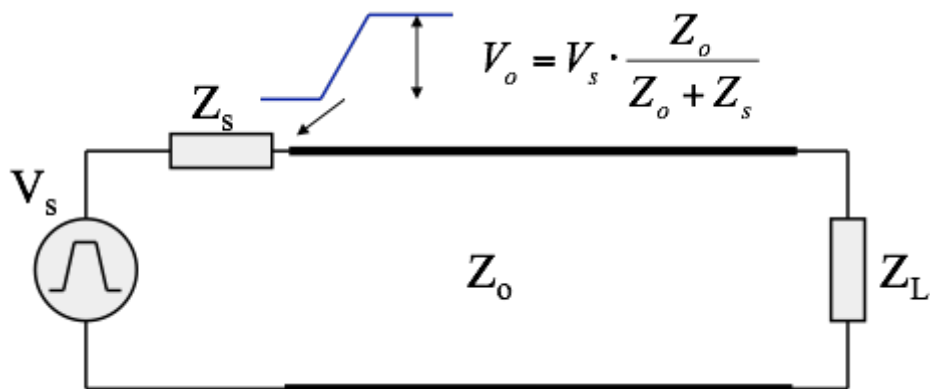


Figura 2.13. El flanco incidente, con amplitud atenuada, se propaga hacia la carga. Fuente propia

Recuerda: este flanco se propaga hacia la carga Z_L a la velocidad de la luz en el medio (aproximadamente a $1,5 \cdot 10^8$ m/s) como una onda electromagnética. A su paso, empuja a los electrones en la pista y en el conductor de retorno.

Segundo paso: el frente de ondas llega a la carga

Si $Z_L \neq Z_O$, parte de la onda se refleja de vuelta hacia la fuente, en una proporción determinada por el coeficiente de reflexión ρ_L (Figura 2.14). Como generalmente $Z_L > Z_O$, ρ_L es positivo. La carga (Z_L) representa la impedancia del buffer de entrada, típicamente de unos pocos picofaradios.

Como ejemplo, he buscado los parámetros de un microcontrolador Cortex M4 salido al mercado en 2017. La capacidad de entrada de los buffers es de 11 pF. Los flancos de salida son típicamente de 5 ns. Esto

implica un ancho de banda de 100 MHz. A esta frecuencia los 11 pF presentan una impedancia de unos 145 ohmios. Si la línea es de 50 ohmios, el coeficiente de reflexión es $\rho_L = \frac{145-50}{145+50} \approx 0,46$.

Es decir, si incide sobre la carga un flanco de 2,5 V de amplitud, en ella mediremos una amplitud igual a la onda incidente más la reflejada: $2,5 V \cdot (1 + 0,46) = 3,65 V$. Es decir, se producirá un pequeño sobreimpulso por encima de los 3,3 V nominales, lo que es algo habitual.

Es exactamente lo mismo que ocurre en un muelle de un puerto: un barco de pesca produce a su paso ondas de pequeña amplitud, pero cuando llegan al muelle su amplitud se duplica, ya que el coeficiente de reflexión es cercano a la unidad. **¿No te habías dado cuenta de este efecto?**

Y ocurre algo más: la onda reflejada viaja hacia la fuente. En este caso, con una amplitud $0,46 \cdot 2,5 V = 1,15 V$.

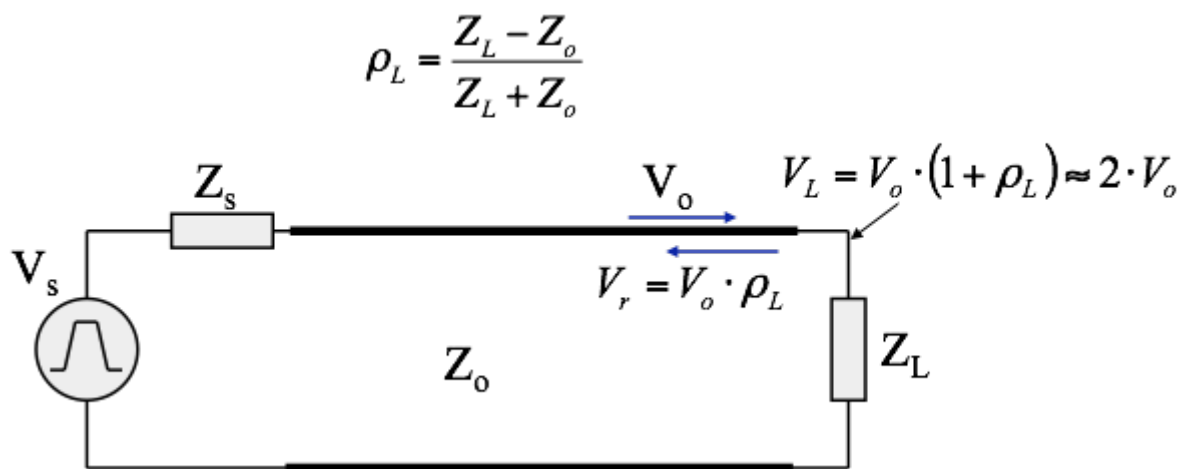


Figura 2.14. La onda incidente llega a la carga y se refleja. Esto provoca un sobreimpulso en la carga y la propagación de la onda reflejada de vuelta a la fuente (*driver*). Fuente propia

Tercer paso: la primera reflexión llega a la fuente

La onda reflejada de 1,15 V de amplitud llega a la fuente, donde de nuevo experimentará el efecto de un cambio de impedancia. Como generalmente $Z_s < Z_o$ (valores típicos son 15 y 50 ohmios respectivamente), el coeficiente de reflexión $\rho_s = \frac{15-50}{15+50} = -0,54$, negativo, indicando que hay inversión de la polaridad.

Como resultado, parte de la señal se refleja de vuelta a la carga según ρ_s y el resto es absorbido por la fuente. Un pulso de amplitud $-0,54 \cdot 1,15 V = -0,62 V$ es reflejado hacia la carga.

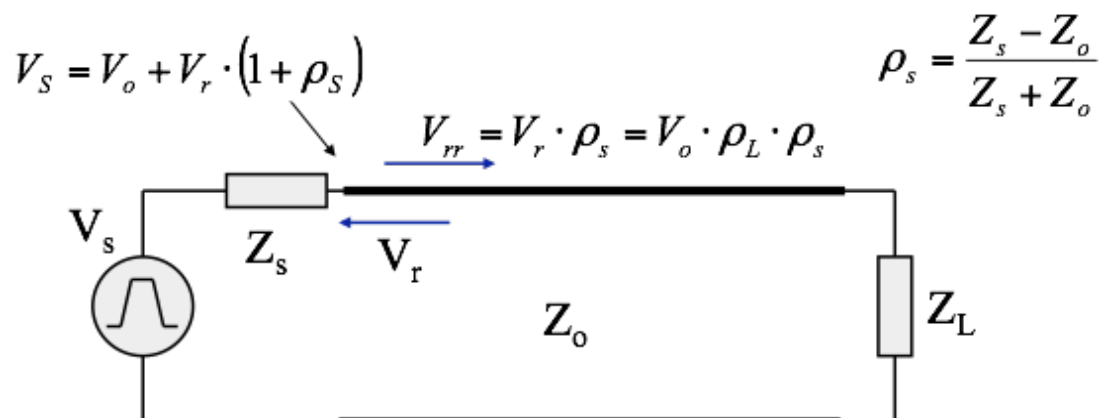


Figura 2.15. La primera reflexión llega a la fuente y es (normalmente) reflejada, atenuada y con polaridad invertida de vuelta a la carga. Fuente propia

Cuarto paso: la segunda reflexión llega a la carga

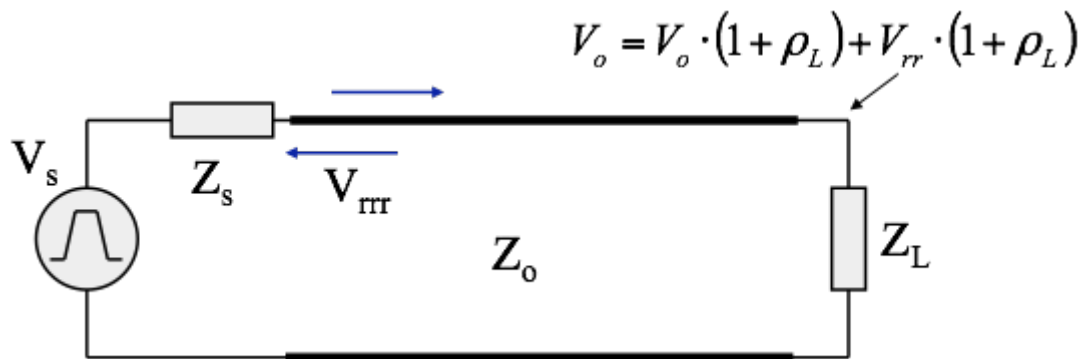


Figura 2.16. La reflexión que se generó en la fuente llega a la carga. Fíjate que en cada “rebote” la amplitud se reduce, de modo que pasados unos pocos ciclos de ida y vuelta se alcanza un estado que podemos considerar estacionario (es decir, termina el transitorio). Fuente propia

¿Recuerdas que onda incidente y reflejada se sumaron en la carga para dar lugar a un pulso de unos 3,65 V en nuestro ejemplo? Pues un tiempo después, lo que le lleva a la onda reflejada viajar hasta la fuente, rebotar allí con signo negativo y llegar de nuevo a la carga, la amplitud en la carga baja hasta $3,65 \text{ V} - 0,62 \text{ V}$, aproximadamente 3V.

Ya puedes observar una pauta:

- Como fuente y carga tienen impedancia distinta de la de línea, ambos extremos dan lugar a reflexiones. De signo positivo en la carga (porque su impedancia es por lo general mayor que la de línea), de signo negativo en la fuente.
- Los coeficientes de reflexión son inferiores a la unidad en valor absoluto, por lo que la onda que va rebotando de un lado a otro reduce su amplitud en cada paso. Es decir, la oscilación se va amortiguando.

El ping-pong continúa hasta que las múltiples reflexiones quedan tan atenuadas que las despreciamos y alcanzamos el estado estacionario, el que calcularías en un análisis en continua (ojo, ahora Z_L es mucho mayor, ¡es un condensador en continua!, por lo que $V_o = V_s = 3,3 \text{ V}$)

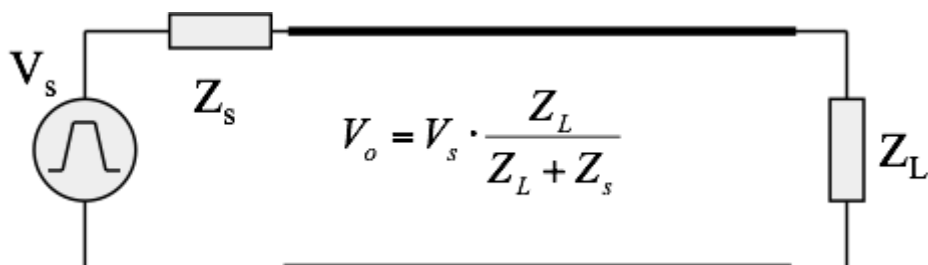


Figura 2.17. Pasado el transitorio, alcanzamos el estado estacionario y la amplitud de tensión en la carga que es la que te enseñaron a calcular en el primer curso de tus estudios. Fuente propia

El método de análisis descrito no refleja exactamente la realidad, ya que los flancos reales son continuos y en nuestro ejemplo hemos supuesto que eran escalones ideales. Pero es una buena aproximación a la solución

del problema. La nota de aplicación [AN-807](#) de 2004 de Texas Instruments describe el método de análisis en detalle.

Esta figura, extraída de una versión anterior de esta nota de aplicación (la prefiero porque está en color) muestra la evolución temporal de la forma de onda en la fuente (azul) y en la carga (rojo), aunque para unos valores concretos de Z_L y Z_S diferentes a los que hemos utilizado en este texto. En el eje de abscisas se indica el tiempo en número de retardos de propagación de la línea (τ).

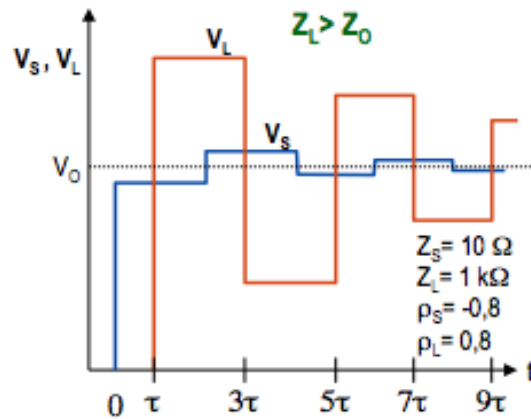


Figura 2.18. Extraída de una versión antigua de AN-807 de Texas Instruments

Pistas en capas externas de un circuito impreso: líneas microstrip

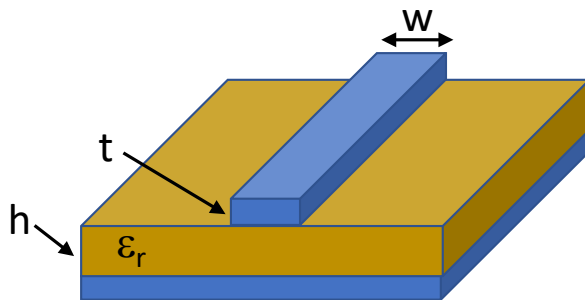


Figura 2.19. Línea microstrip. Fuente propia

En la página 38 comentábamos que la impedancia característica quedaba determinada por la anchura de pista (w), altura del dieléctrico sobre el plano (h), constante dieléctrica del material y en menor medida por el espesor del cobre de la pista (t).

El cálculo no lo hacemos a mano, ya que las expresiones aproximadas son complejas (más abajo reproducimos la que da Wikipedia en la entrada “microstrip”). Incluso esta expresión es incompleta, ya que no está considerando la máscara de soldaduras.

$$Z_{\text{microstrip}} = \frac{Z_0}{2\pi\sqrt{2(1+\epsilon_r)}} \ln \left(1 + \frac{4h}{w_{\text{eff}}} \left(\frac{14 + \frac{8}{\epsilon_r}}{11} \frac{4h}{w_{\text{eff}}} + \sqrt{\left(\frac{14 + \frac{8}{\epsilon_r}}{11} \frac{4h}{w_{\text{eff}}} \right)^2 + \pi^2 \frac{1 + \frac{1}{\epsilon_r}}{2}} \right) \right)$$

El cálculo lo hacemos con calculadoras online (te prevengo, muchas usan expresiones erróneas o inadecuadas) o con programas más fiables (como es el caso de la herramienta PCB Toolkit, que puedes descargarle sin coste de www.saturnpcb.com).

Idealmente, usaremos un *field solver 2D* que en lugar de usar expresiones aproximadas resuelve las ecuaciones de Maxwell y nos da un valor muy exacto de la impedancia.

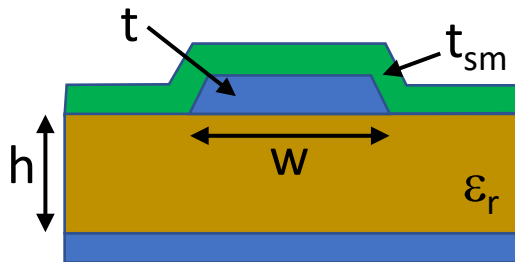


Figura 2.20. Microstrip con máscara de soldaduras. Fuente propia

Una línea microstrip más realista incluye la máscara de soldaduras (de espesor t_{sm}), que con un espesor y una constante dieléctrica relativa aproximados de $20\ \mu\text{m}$ y $3,5$ respectivamente, tiene una influencia de aproximadamente 1 o 2 ohmios en la impedancia de línea.

Otro pequeño error que se comete habitualmente, con una influencia cercana a 1 ohm, es no considerar que en capa externa al cobre base hay que sumarle el cobre añadido en el proceso de metalización de vías, del orden de $25\ \mu\text{m}$.

En realidad, las dos palancas que tenemos para diseñar líneas microstrip son w y h , ya que el resto de los parámetros quedan fijados una vez elegimos material base y fabricante del PCB. Valores habituales para obtener 50 ohm son h en torno a $100\text{-}150\ \mu\text{m}$ y w en torno a $200\text{-}250\ \mu\text{m}$. Puedes descargar el programa de Saturn y jugar con la calculadora de impedancias.

Pistas en capas internas de un circuito impreso: líneas stripline

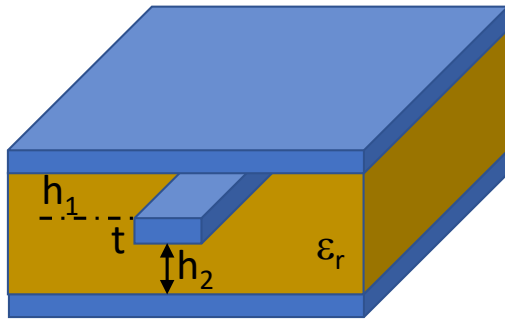


Figura 2.21. Línea stripline. Fuente propia

En la página 38 comentábamos que la impedancia característica quedaba determinada por la anchura de pista (w), distancias entre el dieléctrico y los planos de referencia (h_1 , h_2), constante dieléctrica del material y en menor medida por el espesor del cobre de la pista (t).

En realidad, las tres palancas que tenemos para diseñar líneas microstrip son w , h_1 y h_2 , ya que el resto de los parámetros quedan fijados una vez elegimos material base y fabricante del PCB o tienen una influencia menor (t).

El cálculo no lo hacemos a mano, ya que las expresiones aproximadas son incluso más complejas que las equivalentes para una línea microstrip.

El cálculo lo hacemos con calculadoras online (te prevengo, muchas usan expresiones erróneas o inadecuadas) o con programas más fiables (como es el caso de la herramienta PCB Toolkit, que puedes descargar sin coste de www.saturnpcb.com).

Idealmente, usaremos un *field solver 2D* que en lugar de usar expresiones aproximadas resuelve las ecuaciones de Maxwell y nos da un valor muy exacto de la impedancia.

Corrientes de retorno en una stripline

La proporción entre corrientes de retorno por cada plano, si h_1 y h_2 son las distancias a los planos de referencia, con $h_1 < h_2$, se calcula como [1]:

$$I_1 = \left(1 - \frac{h_1}{h_1 + h_2}\right) \cdot 100\%$$

$$I_2 = \left(\frac{h_1}{h_1 + h_2}\right) \cdot 100\%$$

Si $h_1 = h_2$ (stripline simétrica) las corrientes se reparten al 50%, lógicamente. Una relación 2:1 da lugar a un reparto 33%-67%. Una relación 3:1 a un reparto 25%-75%. Conclusión: no podrás despreciar la corriente de retorno por ninguno de los planos de referencia.

Día 3. Solución al problema de las reflexiones



Una ola rompiendo en un faro es una situación análoga a una carga de impedancia muy superior a la de línea reflejando el frente de ondas. [Imagen](#) de dominio público.

Efecto de las reflexiones en la integridad de señales digitales

Ayer estudiábamos el fenómeno de las reflexiones, que aparecen cuando las impedancias de fuente y de carga difieren de la impedancia característica de la línea. Vamos a comenzar la lección de hoy describiendo los efectos que las reflexiones provocan en las señales digitales.

Las reflexiones producen **sobreimpulsos** y **subimpulsos** que afectan a la forma de onda de las señales. Comencemos considerando, en la Figura 3.1, qué ocurre cuando la señal es una línea de datos o direcciones muestreada por el reloj CLKA por flanco de subida.

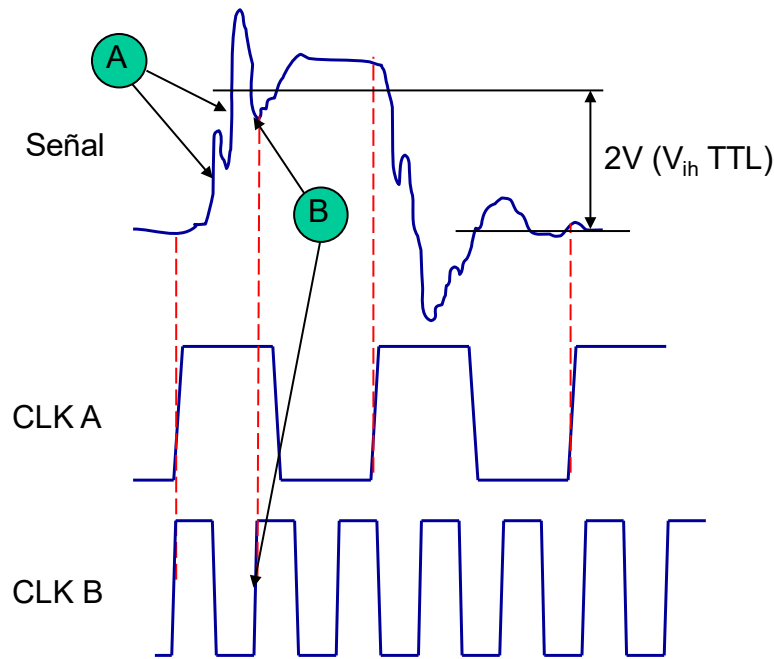


Figura 3.1. Efectos de las reflexiones en una señal digital. Fuente propia

Recuerda lo que te enseñaron sobre señales digitales síncronas: quien recibe la señal sólo la puede ver en una ventana de tiempo muy estrecha igual al tiempo de *setup* (t_{setup}) más el tiempo de *hold* (t_{hold}) del biestable que registra la entrada. Llamemos t_{flanco} al instante en el que se produce un flanco activo de reloj. El receptor sólo “ve” la señal en la ventana de tiempo comprendida entre $t_{\text{flanco}} - t_{\text{setup}}$ y $t_{\text{flanco}} + t_{\text{hold}}$. Y eso suele estar en el entorno de 1 ns.

Así que en la estrecha ventana sensible que rodea al primer flanco de subida de CLK A, el receptor “ve” un cero lógico perfecto. En el siguiente flanco activo, el receptor ve un uno lógico. Y en el tercer flanco activo, un cero lógico. Lo que ocurra fuera de las estrechas ventanas sensibles, es algo que no importa al receptor.

Bien, ¿qué ocurre si el reloj no es CLK A sino CLKB, de mayor frecuencia? En el primer flanco activo, la señal es un cero lógico. Pero en el segundo flanco activo la señal no alcanza el umbral del uno lógico: no es ni un 1 ni un 0. Lo que registre el receptor es algo impredecible: puede entrar en **metaestabilidad**. Y eso quiere decir que, tras un tiempo aleatorio, interpretará un valor 1 o 0 aleatoriamente. Imagina el efecto que esto puede tener en sistema digital. En función de los mecanismos de redundancia, comprobación y recuperación frente a errores que implemente el hardware y el software, el efecto puede variar entre nada, un comportamiento erróneo y un reinicio (si hay *watchdog*).

Vamos a resumir las primeras conclusiones:

Las reflexiones en una señal digital de datos o direcciones pueden producir falsos unos y ceros. El problema se agrava al aumentar la frecuencia de reloj, al no permitir que el transitorio producido por la reflexión llegue a estabilizarse.

Por tanto, es la combinación de flancos abruptos, líneas eléctricamente largas y frecuencias altas de reloj que provoca los errores.

Imaginemos ahora que la señal no es datos o direcciones, sino que se trata de una señal de reloj. Como los flancos no son limpios, es decir no son monótonamente crecientes o decrecientes, pueden producirse falsos flancos de reloj o adelantos/retrasos en la posición de los flancos. De nuevo, piensa en lo que esto puede producir en un sistema digital.

Las reflexiones en una señal de reloj pueden producir falsos flancos y cambios en la posición de los flancos. Este fenómeno no depende de la frecuencia del reloj, sino de lo abruptos de los flancos y de que la línea sea eléctricamente larga.

Hay otro efecto negativo adicional asociado a las reflexiones: los sobreimpulsos pueden producir tensiones lo suficientemente elevadas como para dañar los circuitos integrados. Expliquemos esto un poco mejor apoyándonos de la Figura 3.2.

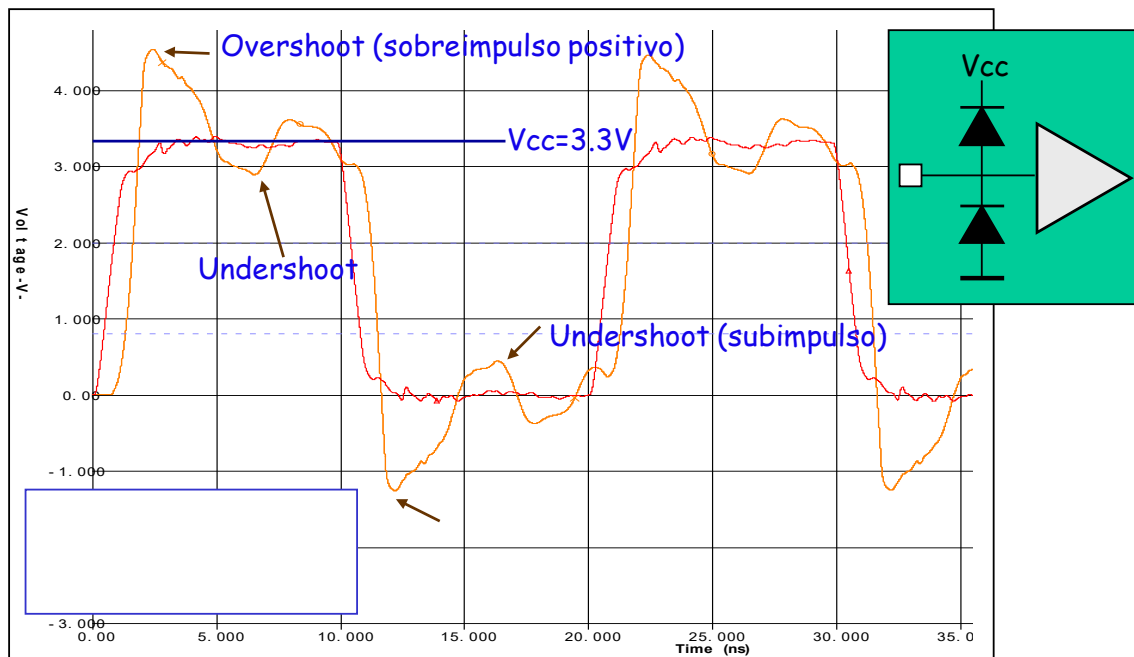


Figura 3.2. Nomenclatura (overshoot, undershoot) y su efecto en las protecciones de los circuitos integrados. Fuente propia

Un **sobreimpulso (overshoot)** consiste en una suma de onda incidente más onda reflejada de la misma polaridad (generalmente producida en una carga) que aumenta la tensión en valor absoluto. Todas las líneas de un circuito integrado incorporan internamente dos **diodos de protección**, uno a la línea de alimentación (V_{cc}), otro a masa, de modo que se ponen en conducción cuando la tensión en la línea supera $V_{cc}+0,7$ V o baja de $-0,7$ V.

Estos diodos tienen una capacidad limitada para conducir corriente eléctrica, determinada en gran medida por el área de la unión PN. Como el área de silicio en un circuito integrado es cara, por lo general estos diodos no son capaces de soportar corrientes de pico superiores a (esto es sólo un orden de magnitud, no lo tomes como un valor absoluto) 100 mA. Sólo ciertos tipos de integrados (como es el caso de un *transceiver* USB o 485) decidan mayor área a estos diodos de protecciones y pueden por tanto soportar mayores picos.

Bien, pues los sobreimpulsos pueden poner en conducción a los diodos de protección y en ciertos casos producir corrientes capaces de dañarlos. Una vez destruido un diodo de protección por sobrecorriente, el siguiente pulso destruirá transistores en el circuito integrado. Esta es otra razón para reducir la amplitud de los sobreimpulsos.

Un **subimpulso (undershoot)**, una reflexión de polaridad opuesta a la señal (generalmente en la fuente) es responsable de que no se alcancen unos y ceros lógicos (Figura 3.2).

Otro día hablaremos de cómo la radiación de las pistas de PCB y cables en un equipo electrónico aumenta cuando hay reflexiones, ya que los sobreimpulsos y subimpulsos aumentan el contenido de energía en altas frecuencias, que se radian con mayor eficiencia. **Las medidas de mitigación para las reflexiones tienen efectos beneficiosos tanto para la integridad de señal como para la EMC.**

A continuación, vamos a estudiar diferentes técnicas de reducción de la amplitud de las reflexiones, todas ellas basadas en un sencillo principio: acercar todo lo imposible la impedancia de fuentes y/o cargas en la línea al valor de la impedancia de línea.

Atenuando de las reflexiones

Como acabamos de comentar, la estrategia será siempre la misma: reducir la amplitud de las reflexiones por el “simple” hecho de igualar tanto como sea posible la impedancia de una o más cargas, o de la fuente, con la impedancia de línea. En función de la topología de la línea adoptaremos una de las siguientes técnicas:

- Terminación serie en la fuente
- Terminación paralela en la carga
- Terminación Thevenin
- Terminación AC (o circuito RC)
- Atenuación mediante resistencias serie en las cargas

Todas estas técnicas requieren añadir uno o más componentes pasivos (resistencias y en algún caso condensadores) a la línea. De acuerdo, las resistencias y condensadores SMD de pequeño valor pueden ser diminutos. Si sólo tuviéramos que **añadir terminaciones** (así es como llamamos al hecho de aplicar una de estas técnicas) a un puñado de líneas en un PCB no sería nada serio, pero con muchos cientos de nodos, un diseñador de PCBs ve incrementado significativamente su trabajo: hay que elegir con cuidado qué líneas requieren realmente terminaciones. Y eso requerirá sentido común y simulaciones.

Bien, ¿Cuándo debes usar cada tipo de terminación? Vamos a estudiar cada técnica por separado, analizando sus ventajas y desventajas y describiendo casos prácticos de uso.

Terminación serie en la fuente

Consiste en una resistencia serie (R_s en la Figura 3.3) cuyo valor, más la impedancia de salida de la fuente (R_d), es más o menos igual a la impedancia característica de la línea (Z_0). Valores típicos son 50 ohmios para Z_0 y 15-18 ohmios para R_d , aunque no es un valor exacto, sino que varía a lo largo del tiempo del flanco de la señal.

Por tanto, se suele escoger un valor entre 22 y 37 ohmios para R_s . El objetivo no es que no haya reflexiones, eso será imposible. El objetivo es que las reflexiones sean pequeñas. De modo que cualquier valor que elijas (22, 33 o 37 ohmios) funcionará.

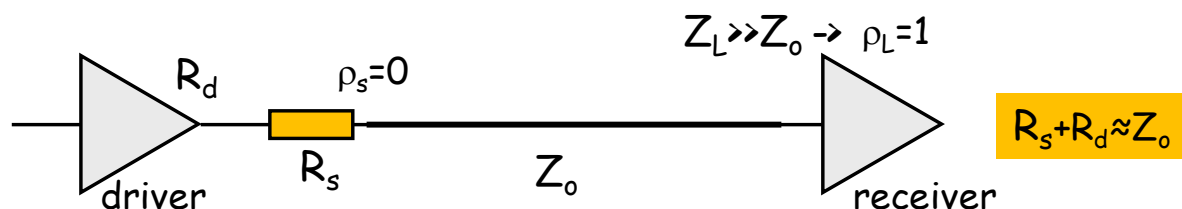


Figura 3.3. Terminación serie en la fuente. Fuente propia

Más importante que el valor concreto de la resistencia es colocarla lo más cerca posible del *driver* (es decir, del pin o bola de salida de circuito integrado que pone la señal en la línea). El conjunto debe ser eléctricamente corto, y eso implica limitarse a unos milímetros de distancia, no más.

¿Cómo evoluciona la señal en la línea?

Hay un divisor de tensión entre el *driver* más R_s (40-50 ohmios) y Z_0 (típicamente 50 ohmios). Es decir, un flanco de aproximadamente la mitad de la amplitud de señal se propaga hacia la carga. Al llegar a la carga, un coeficiente de reflexión alto casi duplicará la amplitud, devolviéndola así al valor que el *driver* pretendía colocar en la línea. Acabamos de resolver un problema: **el sobreimpulso no pasará de la tensión nominal** (ya sea 3.3V, 2.5V, 1.8V, por ejemplo).

Bien, el flanco reflejado en la carga viajará de vuelta a la fuente, pero al llegar ahí el coeficiente de reflexión será muy pequeño (recuerda, $R_s + R_d$ es cercano a 50 ohmios). **Idealmente ya no hay una nueva reflexión**. En la práctica, hay una pequeña reflexión que podemos ignorar: otro problema resuelto.

Y todavía hay más: R_s tumba el flanco, lo hace menos abrupto. Esto implica (recuerda la expresión para la frecuencia de codo que estudiamos en la página 43) que hay menos frecuencias altas en la señal, lo que **reduce la radiación** y por tanto mejora la EMC. Tres pájaros con una piedra.

Un pequeño ejemplo

El esquema de la Figura 3.4 representa un oscilador (U3) que proporciona una señal de reloj a una memoria Flash (U1).

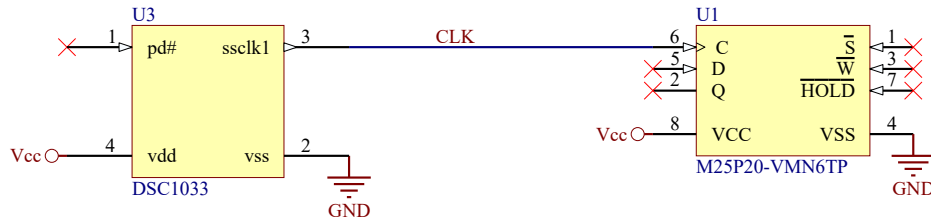


Figura 3.4. Esquema para el ejemplo de línea punto a punto: oscilador y memoria Flash. Fuente propia

En pocos minutos se puede crear el esquema en Altium o en otra suite de diseño de PCBs, definir un *stack-up* (estructura de capas del PCB) para rutar pistas a 50 ohmios, definir las reglas de diseño necesarias y rutar una línea de 10,7 cm entre oscilador y memoria (Figura 3.5).



Figura 3.5. Pista de 50 ohmios y 10,7 cm entre fuente y carga.

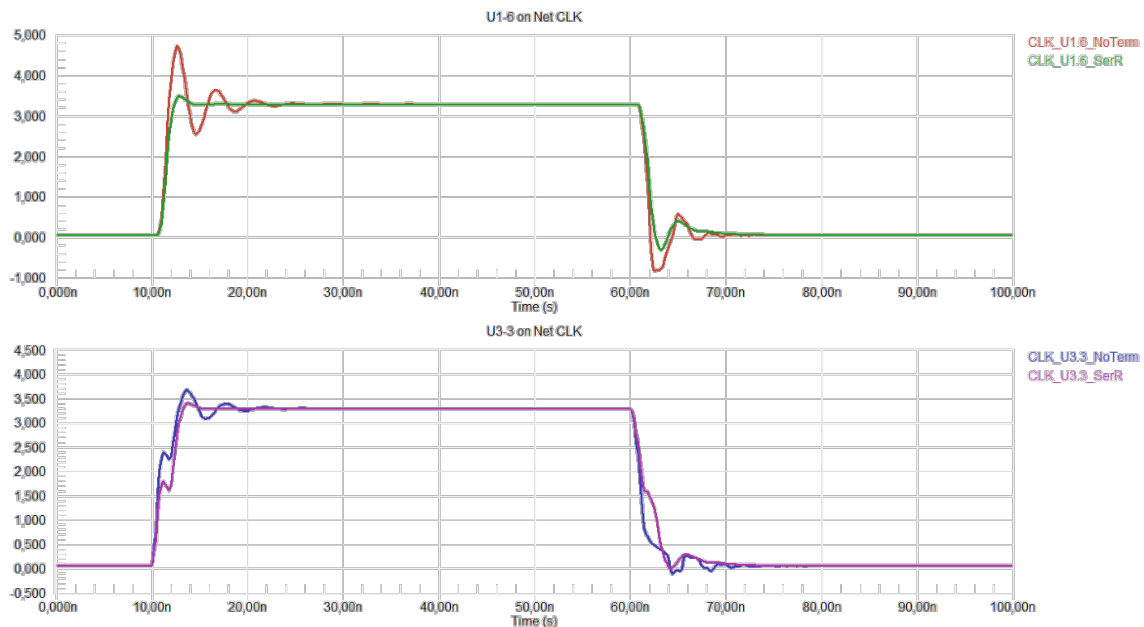


Figura 3.6. Resultado de la simulación en la carga (arriba) y en la fuente (abajo). La forma de onda con más sucia corresponde la simulación sin terminaciones. La forma de onda más limpia incorpora una resistencia serie en la fuente de 22 ohmios. Fuente propia

La simulación, realizada en Altium, (Figura 3.6) muestra la mejora en la integridad de señal en la carga (U1, la memoria Flash). ¿Es importante para las prestaciones del sistema la forma de onda en la fuente, U3? Desde el punto de vista de la integridad de señal, es irrelevante: sólo nos importa la forma de onda en la carga, que es donde debe llegar un reloj limpio.

Nos queda por hablar sobre el valor de la resistencia serie. Hemos utilizado 22 ohmios para obtener la Figura 3.6, pero ¿es este el valor óptimo? En la Figura 3.7 podemos observar que un valor óptimo estaría (buscando valores normalizados) en 27 o 33 ohmios.

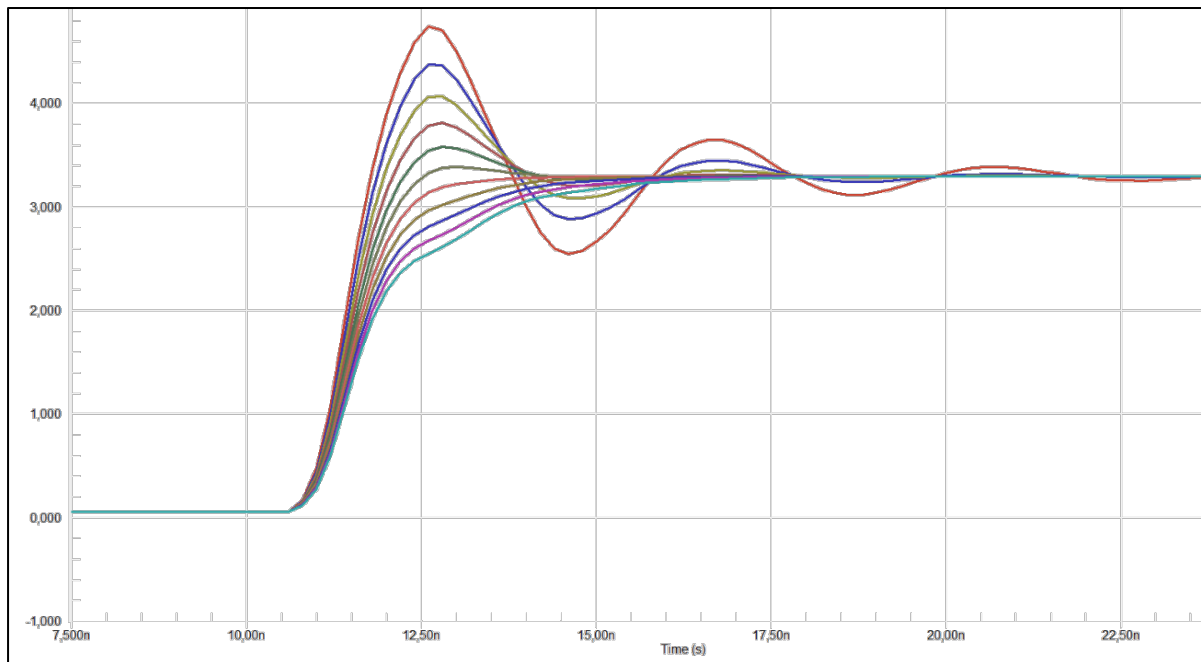


Figura 3.7. Barrido (*sweep*) para valores de R_s entre 0 y 50 ohmios, en pasos de 5 ohmios. Fuente propia

Si comparamos forma de onda y espectro de las señales en la carga sin terminación y con 33 ohmios de terminación serie en la Figura 3.8, apreciaremos lo mucho que pueden mejorar las cosas añadiendo 0,01 € de coste (una resistencia SMD tamaño 0402 o 0201).

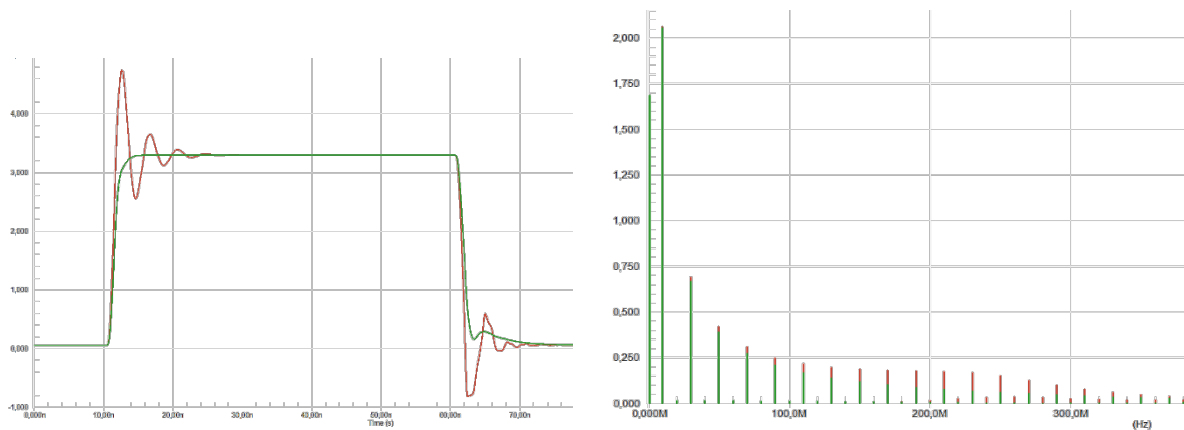


Figura 3.8. Reducción de componentes de alta frecuencia al añadir la resistencia de terminación (verde). Fuente propia

En líneas punto a punto (una fuente y una carga), añadir una pequeña resistencia serie de 22-37 ohmios junto a la fuente tiene el efecto de limpiar la señal y reducir la radiación de la pista. Añadiremos una pequeña resistencia serie inexcusablemente a la salida de cada oscilador y de cada señal activa por flanco. Puede ser necesario, si las simulaciones lo aconsejan, usar esta técnica también en otro tipo de señales para limitar los sobreimpulsos o cuando la señal vaya a salir del PCB y debamos reducir su radiación.

Heavy point-to-point

La terminación serie en la fuente funciona correctamente para señales punto a punto, pero ¿funcionará también en otras topologías? ¿Qué sucede si hay más de una carga en la línea? Considera la topología de la Figura 3.9.

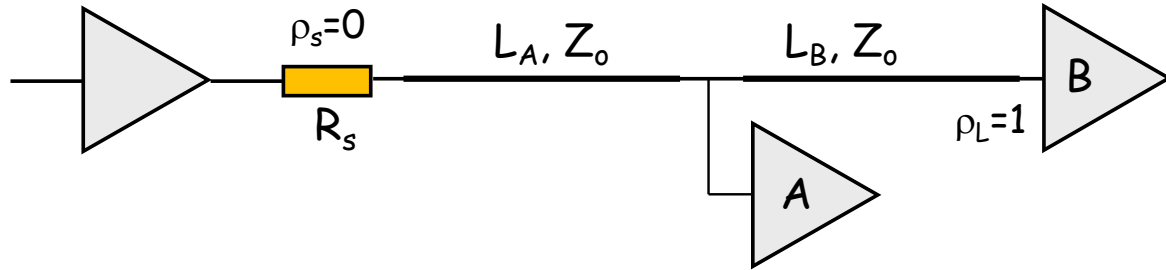


Figura 3.9. Topología *heavy point-to-point* (punto a punto pesada). Fuente propia

La onda incidente produce sólo la mitad de la amplitud en A, tal y como hemos razonado en el apartado anterior. Es necesario esperar a la onda reflejada que llega del extremo de la línea para alcanzar la amplitud total. Se observará un escalón en la forma de onda de A que será más acusado cuanto mayor sea la separación entre A y B. Esto produce flancos no monótonos y retardos elevados (si A y B están muy lejos), lo que no es apropiado para señales de reloj o cuando queremos minimizar retardos.

Por lo tanto, este tipo de terminación es adecuado para líneas punto a punto o, si hay más de una carga, cuando la distancia entre A y B es inferior a la crítica (especialmente en líneas de reloj). En este caso, la topología se conoce como *heavy point-to-point*.

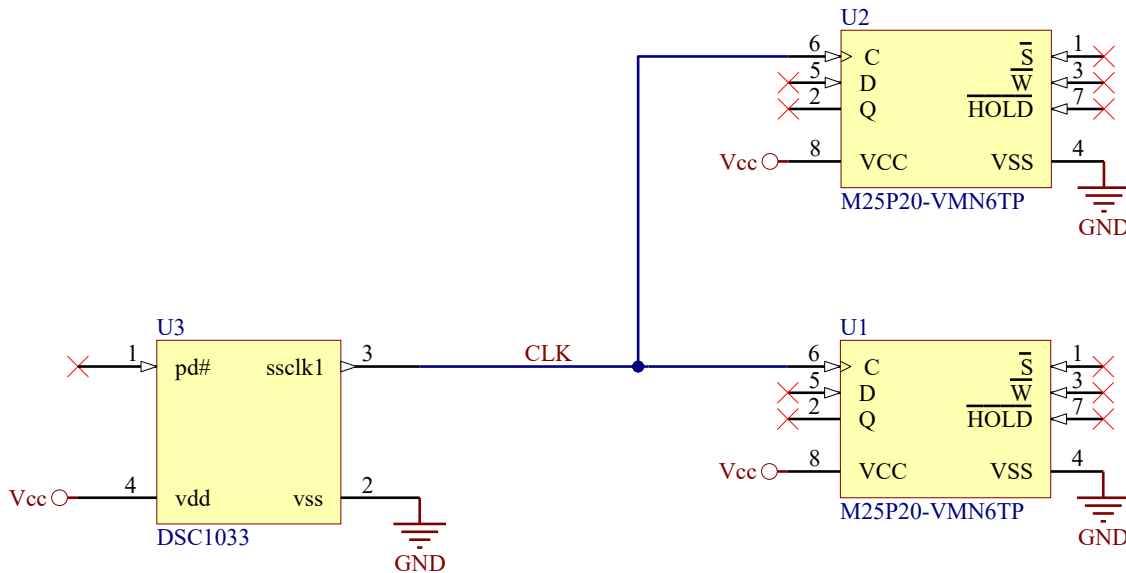


Figura 3.10. Esquema (una fuente y dos cargas) para estudiar un ejemplo de topología *heavy point-to-point*. Fuente propia

En Figura 3.10 y Figura 3.11 hemos definido (esquema y rutado) una topología de una fuente y dos cargas. La simulación de las formas de onda en las cargas se muestra en la Figura 3.13 a la izquierda y centro. Es fácil deducir qué forma de onda corresponde a la carga a mitad de línea: la que presenta un escalón en el flanco, tal y como hemos predicho en la página anterior.



Figura 3.11. Rutado con las cargas alejadas 5,5 cm entre sí. La longitud total de la línea es de 10,7 cm. Fuente propia

Cuando juntamos las cargas (Figura 3.12) desaparece el escalón (Figura 3.13, derecha).



Figura 3.12. Rutado con las cargas muy juntas. Fuente propia

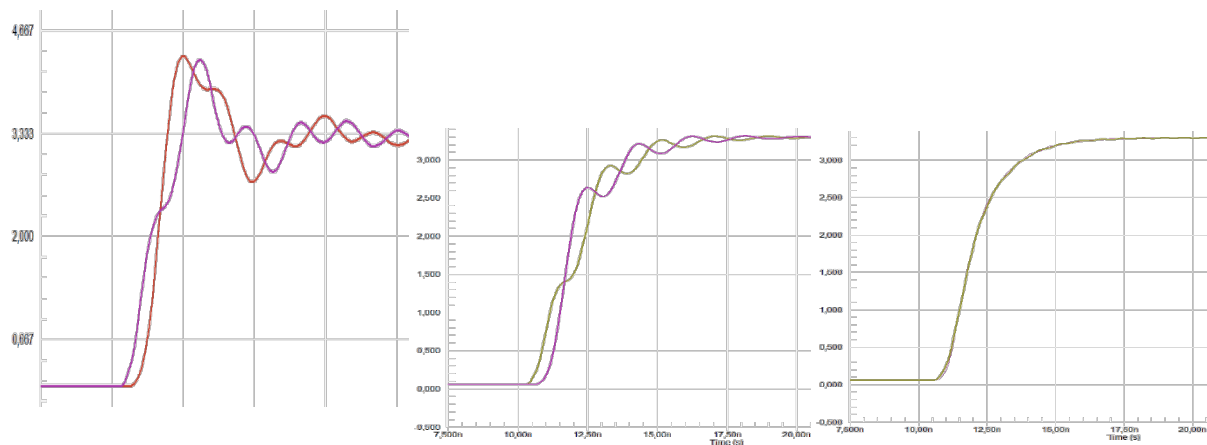


Figura 3.13. Formas de onda sin terminaciones (izquierda), con 33 ohmios (centro) y de nuevo con 33 ohmios, pero acercando las cargas tal y como aparece en la Figura 3.12. Fuente propia

En líneas heavy point-to-point la terminación serie en la fuente es también efectiva, siempre que la distancia entre las cargas sea inferior a la longitud crítica.

Ejemplo práctico

En un diseño antiguo, hacia el año 2000, me encontré ante la topología que aparece en la Figura 3.14. Es un caso más cercano a línea multipunto que a heavy point-to-point, si atendemos a las dimensiones (Figura 3.15). Pero la solución de una resistencia serie en la fuente funcionaba correctamente en simulación. El problema es que había que terminar 84 líneas, lo que ocupa bastante espacio de placa.

Opté por usar resistencias integradas, 4 por encapsulado, para ahorrar espacio. Si echas un vistazo de nuevo a la Figura 3.15 podrás ver los 16 packs de resistencias rodeando dos lados de la FPGA (el integrado marcado como “ORCA”).

Deberías haber caído en el hecho de que ciertamente relojes y direcciones son líneas unidireccionales y que por tanto la fuente es siempre la FPGA y es junto a ésta donde debemos colocar las resistencias. Pero ¿y las líneas de datos? ¿No son bidireccionales? Claro que lo son. La razón de que sólo colocara resistencias junto a la FPGA es que ésta es el *driver* más agresivo, con el flanco más abrupto, y por tanto la que crea más problemas. Las simulaciones, afortunadamente, daban formas de onda razonables cuando conducía las líneas una de las memorias.

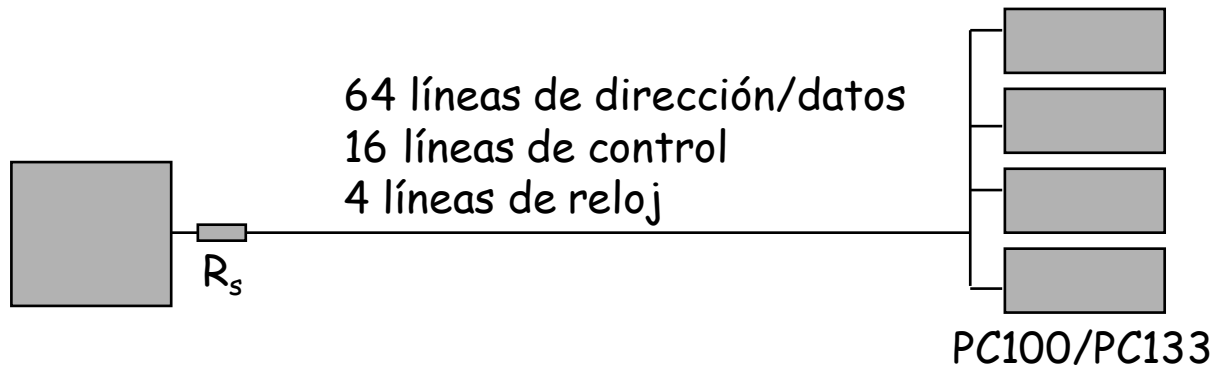


Figura 3.14. Topología multipunto que se resolvió tratándola como heavy point-to-point. Fuente propia

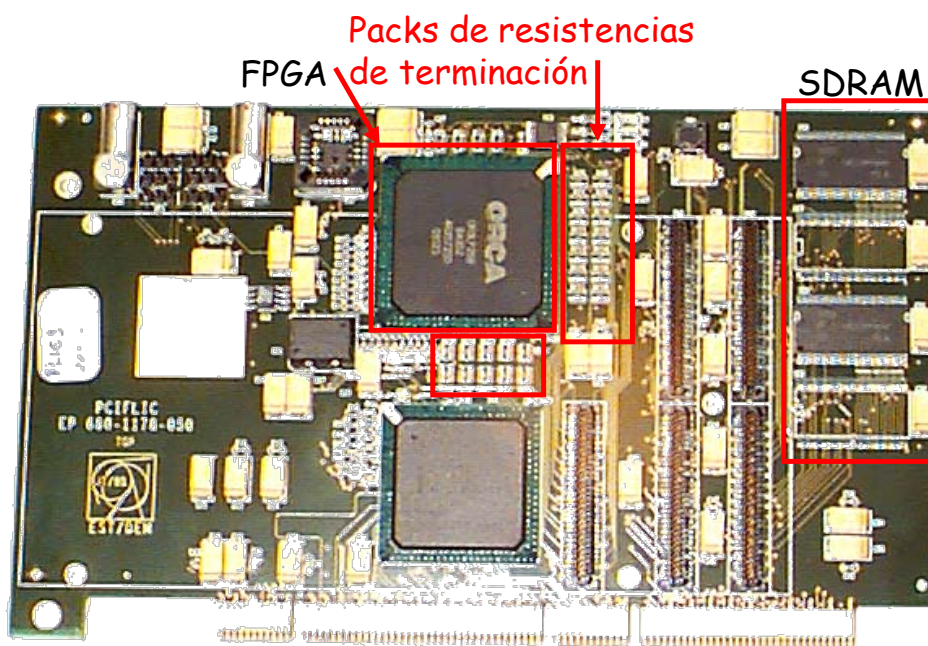


Figura 3.15. Diseño del ejemplo de terminación serie en líneas multipunto. Fuente propia

Terminación paralela en la carga

Este método es adecuado cuando la línea tiene cargas distribuidas, también llamada **línea multipunto**. Se basa en igualar la impedancia de la última carga (C, en la Figura 3.16) a la característica de la línea añadiendo una resistencia del mismo valor al plano de referencia, generalmente masa, aunque vale también llevar la resistencia a cualquier otro plano de alimentación si hay una impedancia baja entre éste y el plano de retorno (por ejemplo, con condensadores de desacoplo).

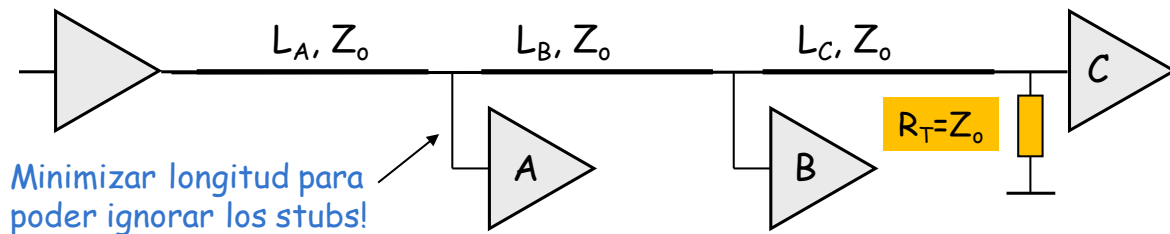


Figura 3.16. Terminación paralelo en la carga. Fuente propia

Al adaptar el extremo de la línea, no se producen reflexiones en C. Esto requiere además que la onda incidente tenga suficiente amplitud para conmutar la línea.

Te preguntarás si hay o no reflexiones en las cargas A y B. Fíjate en que he dibujado la conexión de la línea a A y B con trazo fino, queriendo indicar que son conexiones eléctricamente cortas. En cuanto a la presencia de A y B, su capacidad (típicamente por debajo de 10 pF) modifica localmente la impedancia de línea, pero su efecto es pequeño. Como conclusión, no se producen reflexiones significativas.

Por cierto, estas interconexiones de las cargas A y B a la línea, estas ramas que no forma parte del tronco (línea principal) reciben el nombre de *stubs*. Su presencia provoca reflexiones y siempre que sea posible debemos hacerlas tan cortas como sea posible.

Esta técnica tiene la ventaja de su sencillez, pero en la práctica no se puede utilizar (excepto excepciones) con tecnología CMOS por la elevada corriente requerida: para conducir a 2,5V una línea de 50 ohmios, el *driver* debe entregar una corriente de 50 mA, muy por encima de las posibilidades de la mayoría de los integrados.

Su empleo queda restringido a libros sobre integridad de señal y a tecnologías capaces de entregar corrientes elevadas. Un ejemplo son las memorias DDR2 y DDR3, que usan señalización SSTL (*stub series terminated logic*) y en ciertas condiciones de topología multipunto requieren terminación paralela.

Terminación Thevenin

Quid pro quo. Una cosa a cambio de otra. Vamos a suavizar el requisito de corriente elevada en el *driver* a cambio de complicar la terminación, que se hará con dos resistencias: una a la tensión de alimentación del driver, otra a masa, de modo que el equivalente de Thevenin de la carga (de aquí el nombre) sea las dos resistencias en paralelo, valor que debe ser cercano a la impedancia característica de la línea.

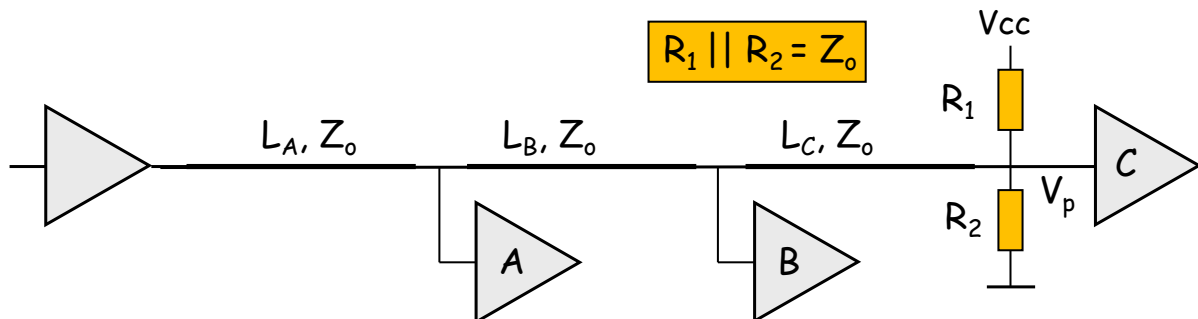


Figura 3.17. Terminación Thevenin. Fuente propia

Supongamos $R_1=R_2=100$ ohmios en la Figura 3.17. Si $V_{cc}=2,5$ V, para mantener un uno lógico en la línea (asumimos niveles lógicos TTL, siendo el umbral 2V), el *driver* debe entregar una corriente de salida de 15 mA. Para garantizar un 0 en la línea (el umbral es 0,8 V) el *driver* debe absorber 9 mA. Ambos valores son razonables para tecnologías CMOS, por lo que la terminación Thevenin sí es usada en la práctica.

Un ejemplo: la línea de reloj en un bus I2C

Podemos encontrar un ejemplo habitual de terminación Thevenin en la línea de reloj de un bus I2C. Se trata de un bus lento de configuración multipunto de dos líneas -SDA, datos y SCL, reloj-, generalmente no supera 1 MHz, pero con flancos lo suficientemente abruptos como para que el reloj quede degradado. En audio digital, se usa una variante denominada I2S.

No son pocos los diseños complejos (memorias DDR4, microprocesadores de última generación, interfaces SATA, ...) que simplemente no funcionan porque un humilde reloj de bus I2C no recibió la atención requerida.

En función de la topología, puede bastar una resistencia serie en la fuente o puede ser más adecuada una terminación Thevenin. Hay que simularlo para saberlo.

Terminación AC

Otra solución de terminación paralelo en la última carga que busca minimizar la corriente en continua demandada a la fuente. El truco, en este caso, está en reemplazar la resistencia de terminación por un circuito RC serie. En continua (un 1 o 0 estático) no requiere corriente circulando por la rama.

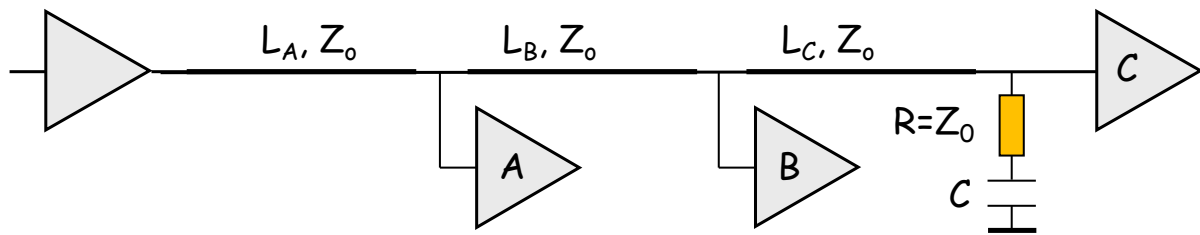


Figura 3.18. Terminación AC (o terminación RC). Fuente propia

La reactancia del condensador debe ser pequeña para las frecuencias altas de la señal, lo que fija un límite inferior al valor de C. Recuerda que la señal tendrá energía significativa hasta frecuencias tan altas como B_w :

$$B_w \approx \frac{0,5}{t_i}$$

Por otro lado, también hay un límite superior al valor de C, ya que queremos garantizar que la constante de tiempo del circuito RC es pequeño comparado con la duración del pulso digital:

$$R \cdot C \ll T_p$$

En la práctica, los valores ideales de R y C se obtienen iterativamente o mediante herramientas software.

Atenuación mediante resistencias serie en las cargas

En la práctica, en líneas punto a punto y *heavy point-to-point*, la terminación serie en la fuente funciona bastante bien. En topologías multipunto, con la enorme variedad de situaciones que pueden darse, es poco frecuente que una de las soluciones canónicas.

Hay una solución que siempre mejora las cosas y que en ocasiones no requiere el concurso de otras soluciones. Se trata de añadir pequeñas resistencias serie junto a las cargas. Te preguntarás de qué sirve colocar, por ejemplo 33 ohmios, en serie con una carga cuya impedancia es supuestamente alta. Una carga típica (una entrada en un circuito integrado) puede representarse como un condensador de 7 a 10 pF. A una frecuencia de 400 MHz, esta capacidad presenta una impedancia de 40 ohmios.

Ahora ves que las altas frecuencias “ven” un divisor de impedancia entre unas pocas decenas de ohmios (carga) y la resistencia serie junto a la carga (tal vez una trentena de ohmios). El sobreimpulso, fruto de la suma de onda incidente más reflejada en la carga, se ve atenuada. En cambio, frecuencias medias, bajas y continua no sufren esta atenuación. Este es el truco detrás de esta solución que, *a priori*, parecía no tener razón de ser.

Ejercicios de reflexiones y terminaciones

En cualquier diseño, una vez satisfechos con una primera versión del *floorplan* (plano de la ubicación de los componentes principales en el PCB) y una primera estimación de la estructura de capas del PCB, debemos realizar **simulaciones pre-layout** (antes de realizar el layout, es decir, antes de rutar) de los buses y líneas críticos. Entendemos por buses y líneas críticos:

- Todas las señales de reloj del sistema, sea cual sea la frecuencia de operación
- Una línea de datos, direcciones y control (ya que cada grupo puede tener una topología distinta) de todos los buses que superen la longitud crítica

Los resultados de cada simulación podrán ser satisfactorios o no satisfactorios. Serán no satisfactorios si:

- El *overshoot* es excesivo y puede dañar los *buffers* de los circuitos integrados (los límites se encuentran en las hojas de datos de los componentes, en la sección *de absolute maximum ratings*)
- En las líneas de reloj, se producen falsos flancos de reloj
- A la frecuencia de operación del bus la señal no se estabiliza a tiempo y en el siguiente flanco activo hay un nivel lógico equivocado o inválido
- Aunque no se produzca ninguna de las situaciones anteriores, nos encontramos cerca de alguna de ellas. Como normalmente simulamos en caso “typical”, si estamos cerca del límite repetiremos la simulación en “weak” y “strong” y es muy posible que el resultado sea no satisfactorio

Un resultado no satisfactorio deberá provocar alguna acción que corrija el problema, como:

- Modificar la longitud de los segmentos del bus y/o de los *stubs*
- Modificar la topología del bus
- Añadir terminaciones en las líneas
- Usar un buffer distinto en aquellos circuitos integrados que tengan buffers programables (generalmente FPGAs)

El proceso iterativo de simulación-modificación continúa hasta quedar satisfechos con el resultado.

HyperLynx es una herramienta muy útil, y podemos sentirnos tentados de dejar que piense por nosotros. Pero debemos usarla para ayudarnos a desarrollar nuestra intuición y experiencia como ingenieros; nunca para sustituirlas.

Vamos a estudiar dos casos prácticos, centrándonos únicamente en la problemática de las reflexiones y de cómo afecta a la integridad de la señal (niveles lógicos, falsos flancos y sobretensiones).

El primer caso práctico nos servirá para recordar que no hay que subestimar las líneas de baja frecuencia de operación.

El segundo caso práctico (útil para evitar que los alumnos más rápidos se aburran) nos permite estudiar un caso muy frecuente: una línea multidrop, en el que una señal es entregada a un conjunto de cargas siguiendo una topología de bus lineal.

Primer caso práctico (es un caso real de diseño)

Trabajas como ingeniero especialista en sistemas digitales y diseño de PCBs en una consultora. Un cliente acude hoy a tu empresa con un problema: han producido una preserie de 400 unidades de un nuevo producto basado en FPGA. Las unidades no funcionan.

El ingeniero asegura que no es posible configurar la FPGA ni su memoria a través del bus JTAG. Han revisado las soldaduras, las alimentaciones son correctas, parece no haber errores en el diseño del PCB... No entienden a qué es debido el problema.

Tras echar un vistazo a los esquemas, crees haber identificado el problema. Pides al cliente los ficheros del diseño y (aprenderás a hacer esto en otra sesión) extraes la topología del nodo TCK (ver fichero JTAG_TCK, que corresponde a un diseño real). Simulas la línea, identificas el problema y propones una solución al cliente.

Resolución del primer caso

La herramienta HyperLynx de Mentor Graphics permite, entre muchas otras funcionalidades, importar los ficheros de diseño de un PCB y extraer la topología de un nodo. En la Figura 3.19, a la izquierda, U15 representa el conector JTAG y el *driver* externo conectado a él, generalmente mediante un cable plano. Cierto, no estamos modelando el cable plano, podríamos hacerlo editando el esquema, pero para la resolución del ejercicio no será necesario.

HyperLynx representa tramos de líneas de transmisión como cilindros, drivers y cargas como triángulos, vías como carretes de hilo. Cada tramo de línea de transmisión tiene asociada una capa de rutado, una longitud y a partir del stack-up del diseño (que hay que importar también) una impedancia y un retardo de propagación. Ahora comprenderás un poco mejor la figura.

Tras una corta pista en capa top (menos de media pulgada) la señal pasa a capa 4, y tras aproximadamente media pulgada, una vía conecta la señal con la memoria Flash (U8) en capa top y con la FPGA en la cara opuesta del PCB (capa 12, el stack-up se muestra en la Figura 3.20). Antes de llegar a la FPGA, tres vías llevan la señal a un tercer dispositivo JTAG en la placa a través de las capas 10, 9 y 1 (top).

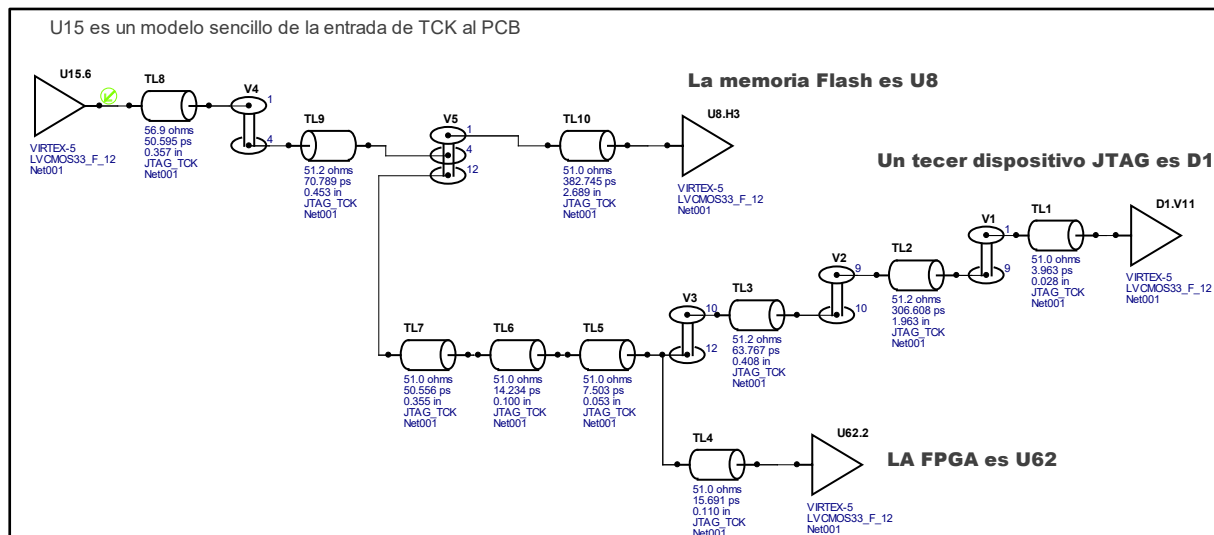


Figura 3.19. Esquema de la topología del nodo TCK, extraído automáticamente por la herramienta HyperLynx a partir de un PCB rutado. Fuente propia

Tras asignar modelos de simulación a los *buffers* (he usado modelos de 3,3V, de 12 mA y rápidos de la librería de FPGAs Virtex-5 a los cuatro elementos activos), podemos comenzar a analizar el funcionamiento de la línea. En el caso de la memoria Flash y del tercer dispositivo, que son entradas y nunca conducen la línea, su modelo no es crítico (entenderás cómo se modelan estos dispositivos en el Día 8). Para un análisis rápido está bien, pero para hacerlo perfecto deberíamos haber buscado los modelos de simulación que ofrece

cada fabricante para sus circuitos integrados. Para U15, lo cierto es que no sabemos de antemano qué dispositivo externo vamos a conectar: nuestro modelo de *driver* rápido de 12 mA puede ser una aproximación razonable, y de hecho sirve para poner de manifiesto los problemas del rutado de esta línea.

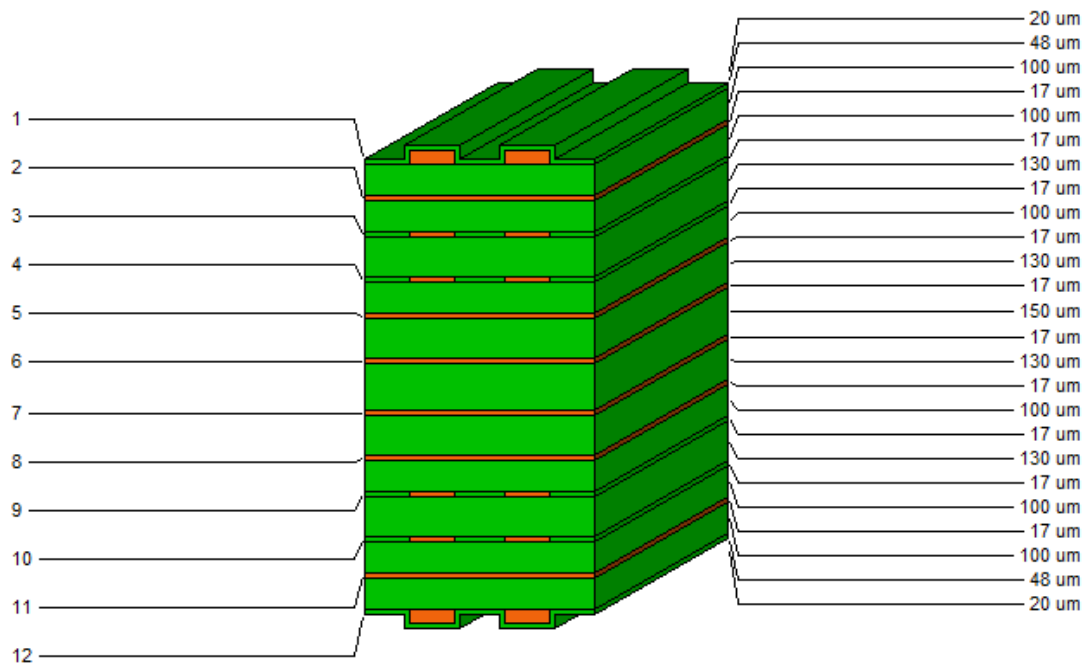


Figura 3.20. Stack-up del PCB del ejercicio. Fuente propia

La Figura 3.21 muestra claramente que el flanco de reloj que recibe la FPGA no es monótono, sino que presenta cambios de signo en la pendiente en torno al umbral de detección de reloj: [un desastre](#). El rizado que observamos tras el flanco no afecta a la calidad del reloj, pero contiene altas frecuencia que contribuirán a la radiación del PCB. Esta simulación es para el caso “typical”. Los casos “slow” y “fast” contemplan los casos extremos, dentro de la distribución estadística resultante de la dispersión de parámetros en fabricación. Es decir, la mayoría de los circuitos integrados se comportarán como el caso “typical”, pero unos pocos presentarán flancos más abruptos (fast) o lentos (slow) y has de verificar los tres casos en las simulaciones.

En el caso “fast” el diagnóstico es el que hemos comentado anteriormente (no muestro la figura para no recargar). En el caso “slow”, la memoria Flash se suma a la FPGA como receptores de flancos no monótonos cerca del umbral de transición.

¿Por qué se producen estas distorsiones en los flancos? Ya sabes que es debido a las reflexiones. La impedancia de línea en todas las capas de rutado difiere no más de $0,4 \Omega$ (cierto que en producción este valor será mayor), de modo que no hay que culpar a los saltos entre capas. Observando de nuevo la Figura 3.19, podemos identificar cuatro puntos en los que se producen reflexiones:

1. Vía V5: la señal viene por una línea de 50Ω y se divide en dos líneas también 50Ω , lo que resulta en 25Ω . Se produce una reflexión de señal de vuelta hacia la fuente y con inversión de signo.
2. U8: representa una impedancia elevada (con respecto a la de línea) tras una pista larga (sólo el tramo en capa top tiene ya 2,7 pulgadas). El coeficiente de reflexión es elevado y positivo (sin inversión de signo)
3. D1, por el mismo motivo que el caso anterior.
4. U15, el driver: su impedancia, menor a la de línea, hará que las reflexiones que llegan al driver sufran una nueva reflexión con inversión de signo.

Te preguntarás por qué me he olvidado de U62, la FPGA. ¿No se producirán reflexiones como en los casos 2 y 3 anteriores? [La respuesta es que la reflexión será muy pequeña](#). Fíjate en que el pequeño *stub* que conecta la rama principal con la FPGA tiene 0,11 pulgadas (como 2,8 milímetros), lo que será eléctricamente corto y no se comportará como línea de transmisión. La FPGA se comporta como una pequeña capacidad eléctrica, que afecta localmente a la impedancia, pero a un nivel mucho menor que los casos anteriores.

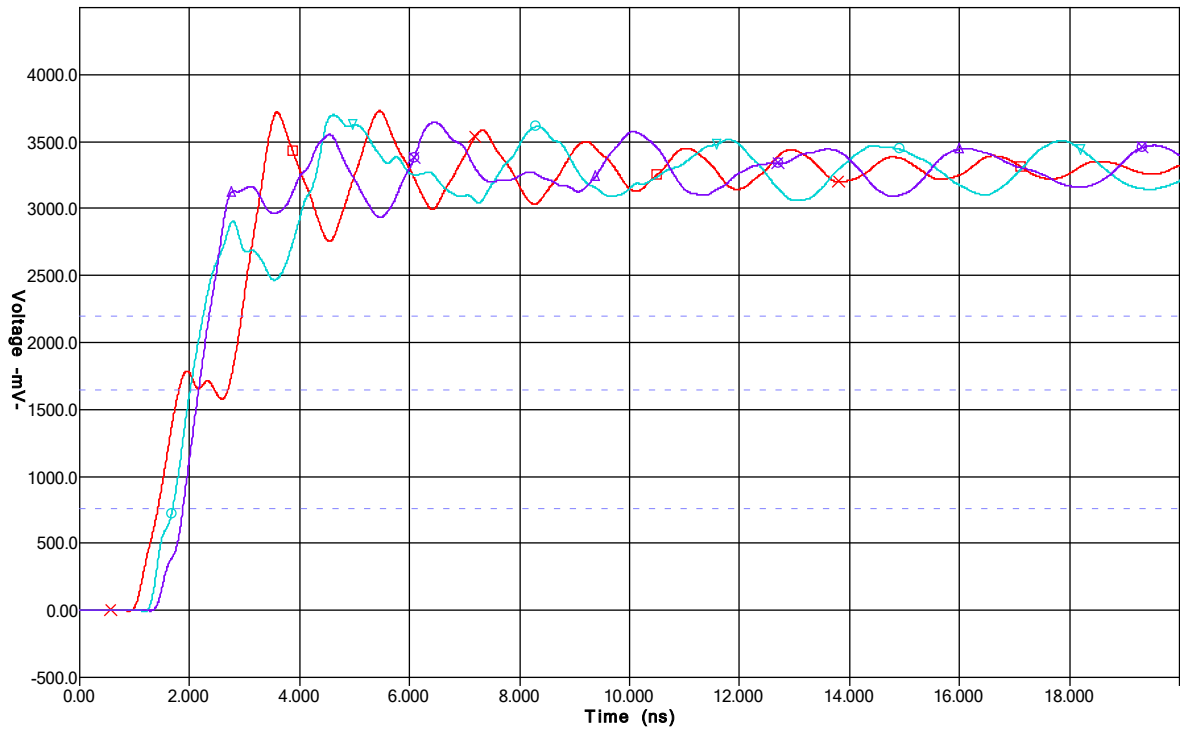


Figura 3.21. Simulación de un flanco de subida de reloj en el caso “typical” (podemos simular los casos slow, typical y fast). El umbral de detección del flanco (típicamente a $V_{cc}/2$, aprox. 1,65 V) es cruzado dos veces por la forma de onda roja (U62, la FPGA). No es de extrañar que la configuración de la FPGA no funcione, ¡el reloj que recibe no tiene flancos limpios! Fuente propia

La distorsión y el rizado que ves en la Figura 3.21 es fruto de la combinación de las múltiples reflexiones en los cuatro puntos que he indicado, y lo que debes hacer ahora es razonar qué debes tocar para mejorar la integridad de señal en la FPGA y en la memoria Flash.

A menudo no hay una solución única, sino varias y hemos de elegir la más sencilla. Las herramientas como HyperLynx permiten jugar a “¿qué pasaría si...?” sin riesgos y rápidamente. Pero podemos perder una tarde o sólo 10 minutos si pensamos un poco antes de tocar nada.

Y ahí va un pequeño truco que te será útil en la mayoría de las ocasiones. Pero primero mira el resultado:

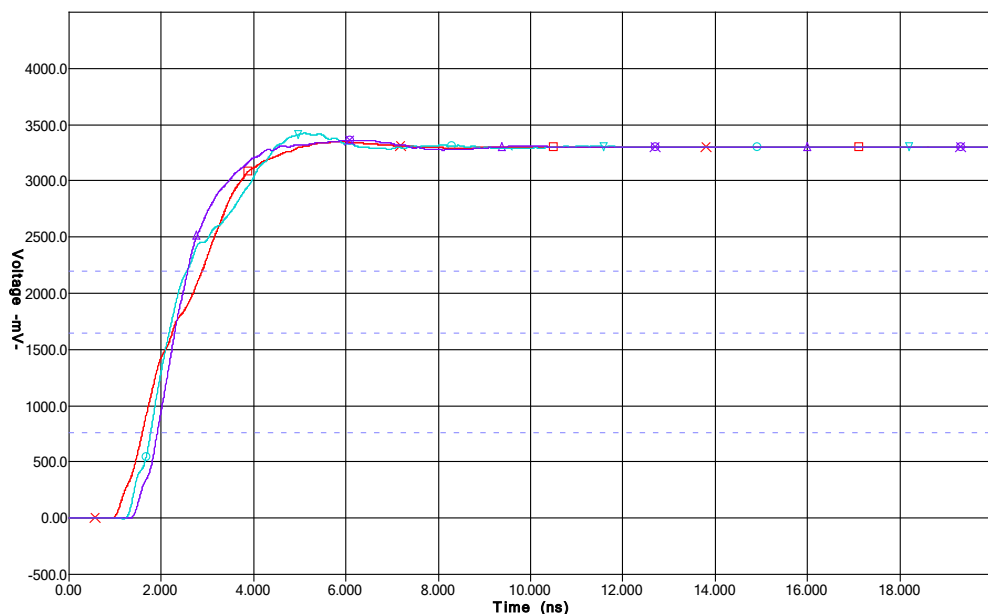


Figura 3.22. Flancos limpios y sin rizado en una simulación del caso “typical” aplicando un sencillo truco. Fuente propia

¿Cómo lo hemos logrado? Mira la figura Figura 3.23. Hemos añadido resistencias de $33\ \Omega$ en serie con cada una de las tres cargas. ¡Espera! Esta no es ninguna de las terminaciones que hemos estudiado y que encontrarás en la mayoría de los libros especializados. ¿Qué hemos hecho?

U15 es un modelo sencillo de la entrada de TCK al PCB

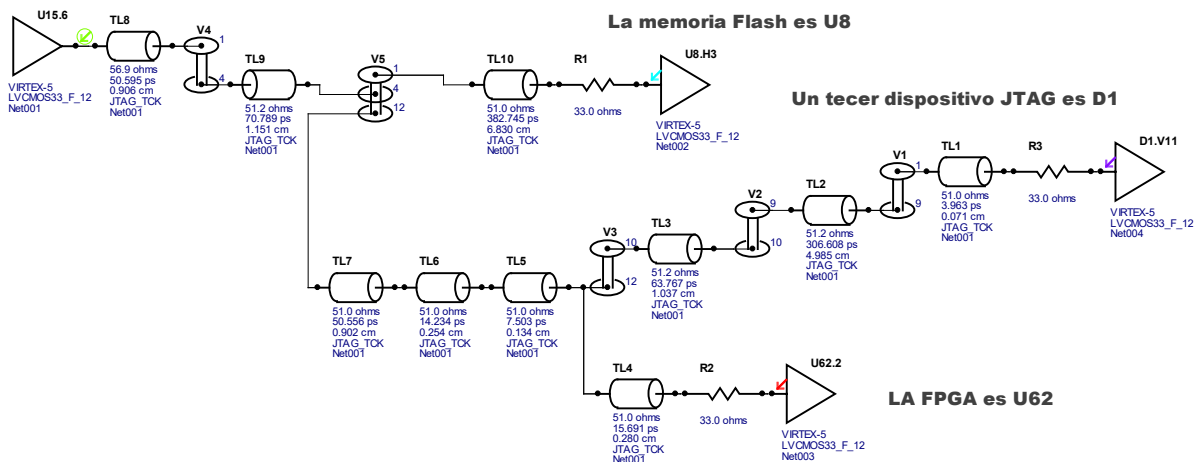


Figura 3.23. Con tres pequeñas resistencias limpiamos el reloj, mejorando integridad de señal y EMC del diseño. Fuente propia

Si una carga puede representarse por una pequeña capacidad del orden de $5\ \text{pF}$, a la impedancia del flanco (el flanco es aproximadamente $2,5\ \text{ns}$ en las figuras, lo que resulta en unos $200\ \text{MHz}$) estamos hablando de unos $160\ \Omega$. Añadir $50\ \Omega$ en serie crea un divisor de tensión ($33\ \Omega - 160\ \Omega$, es decir, una atenuación de un 17% de la amplitud de la onda incidente) sólo a estas altas frecuencias del flanco, ya que el resto del tiempo el buffer presenta una impedancia de entrada mucho mayor. Esto permite atenuar las reflexiones y limpiar las señales.

¿No podíamos haber resuelto el problema usando lo que hemos estudiado? Claro que sí. Mira lo que hemos logrado en la Figura 3.24.

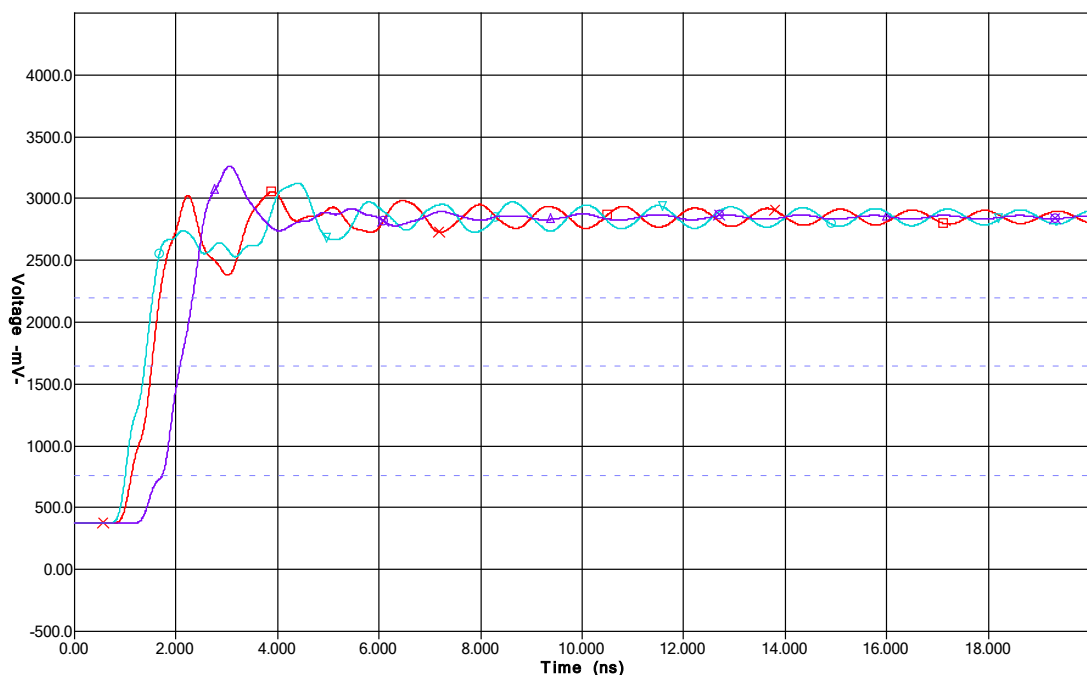


Figura 3.24. Ejercicio resuelto de una forma más ortodoxa. Fuente propia

Aunque no es tan bueno como la solución anterior, sirve. **¿Cómo lo hemos logrado?** Reduciendo las reflexiones en U8 (hemos reducido la distancia entre la vía y el circuito integrado hasta aproximadamente 0,7 cm) y en D1 (añadiendo una terminación Thevenin). Como ves, sólo hay que pensar un poco y tocar las palancas adecuadas.

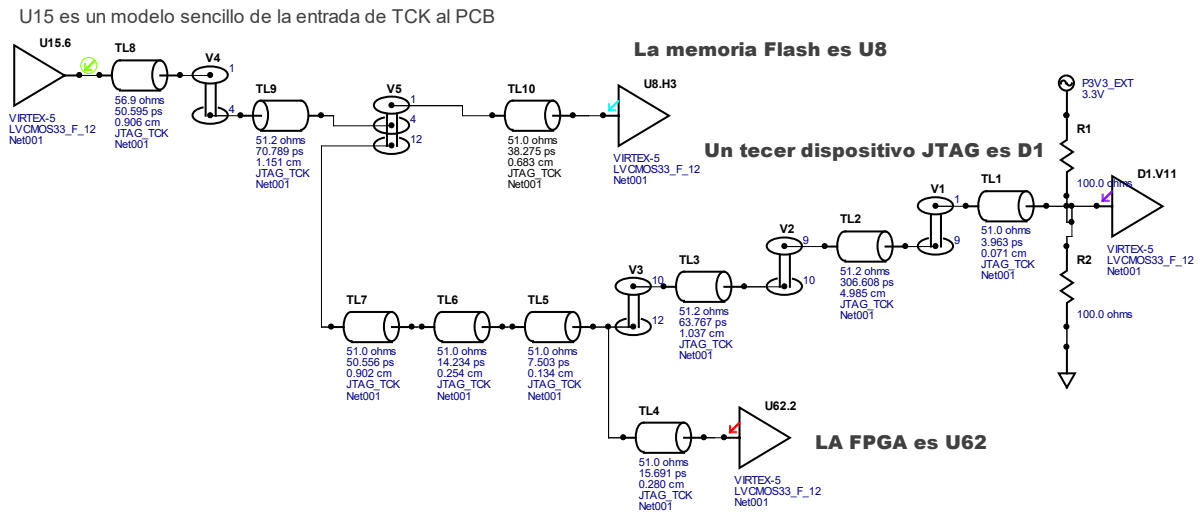


Figura 3.25. Resolución del problema atacando a las reflexiones en U8 y en D1. Fuente propia

Ya tienes dos soluciones que presentar al cliente. En función de cómo esté rutado el PCB, puede elegir la que sea más sencilla de realizar.

Segundo caso práctico

Una vez entregado el informe al cliente, queda tan satisfecho que te invita a comer y te pide ayuda para un nuevo diseño. En esta ocasión quiere que hagas un estudio previo de su red de distribución de reloj, que les da problemas incluso a frecuencias bajas en el prototipo. El ingeniero del cliente te dibuja un diagrama en una servilleta:

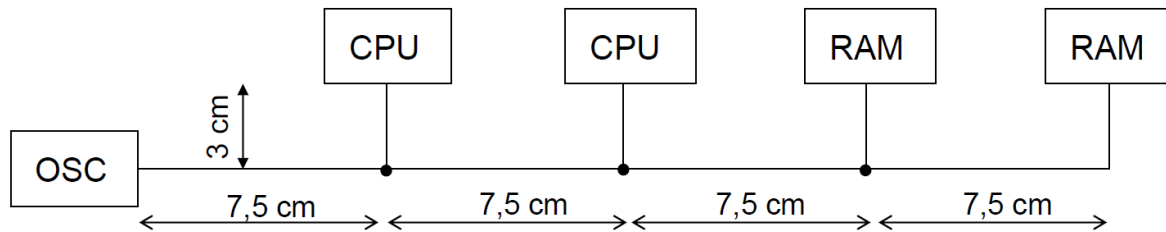


Figura 3.26. Segundo caso práctico de terminaciones. Fuente propia

La señal de reloj generada por un oscilador de 50 MHz de 3,3 V es entregada a dos CPUs y dos memorias RAM. Las longitudes aproximadas de los segmentos de la línea de reloj están indicadas en la Figura 3.26. Una vez más, usan 50 Ω como impedancia de línea y asumes (sabes que te equivocarás poco) que la señal se propaga por el PCB a aproximadamente la mitad de la velocidad de la luz en el vacío, es decir: 15 cm/ns. Preguntas por lo abruptos que son los flancos de las señales y te dicen que algo inferior a 1 ns. Ante su sorpresa, sin pedir los ficheros del diseño, prometes analizar la topología y entregar un informe al día siguiente.

Resolución del segundo caso

En este caso no partimos de un diseño rutado, sino que creamos un esquema de la topología del nodo manualmente (Figura 3.27). Buscando drivers con un tiempo de flanco entre 1 y 2 ns hemos escogido una vez más un modelo de la librería Virtex-5. No por tener un apego especial a esta tecnología, sino porque las FPGAs disponen de una amplia variedad de I/Os donde elegir. El modelo LVTTTL_F_24 sirve a nuestro propósito, pues producirá en nuestro nodo de reloj flancos de 700-800 ns de tiempo de subida. Los tramos de 7,5 cm tendrán un retardo de propagación de 500 ps y los de 3 cm de 200 ps.

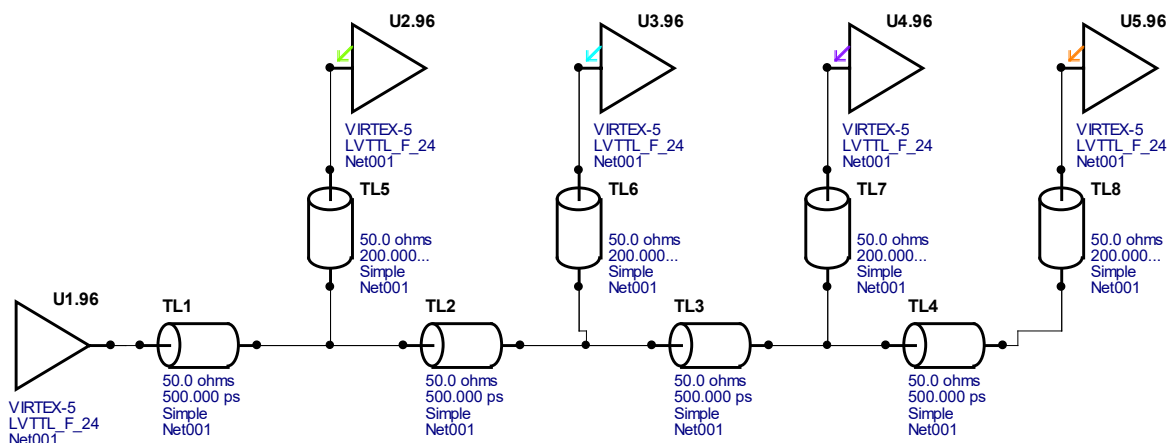


Figura 3.27. Topología del nodo del segundo caso práctico. Fuente propia

Antes de simular, hay que poner en marcha el motor del pensamiento crítico. Con flancos de 800 ps tendríamos un ancho de banda de aproximadamente $0,5/800\text{ps} = 625\text{ MHz}$. A esta frecuencia (y considerando una velocidad de propagación igual a la mitad de la de la luz en el vacío, ya que la constante dieléctrica del

dieléctrico en un PCB es aproximadamente 4) la longitud de onda equivalente a 625 MHz es de 24 cm. Por encima de $\lambda/10$ (2,4 cm) tenemos líneas de transmisión y se manifiesta el fenómeno de reflexión. Es decir, en los *stubs* entre la línea principal y cada carga tendremos algo de reflexión. Cabe considerar reducir estas distancias como estrategia para limpiar las señales, si fuera necesario.

Por otra parte, sabemos que habrá reflexiones tanto en la última carga como en la fuente de reloj: reducir este efecto en alguno de los dos extremos puede ser otra estrategia que considerar.

Si simulamos las formas de onda en el nodo, caso “typical”, ignorando U1 (es quien produce la señal, lo que nos interesa es la señal a la entrada de las cuatro cargas), obtenemos las siguientes gráficas, que comentaremos individualmente:

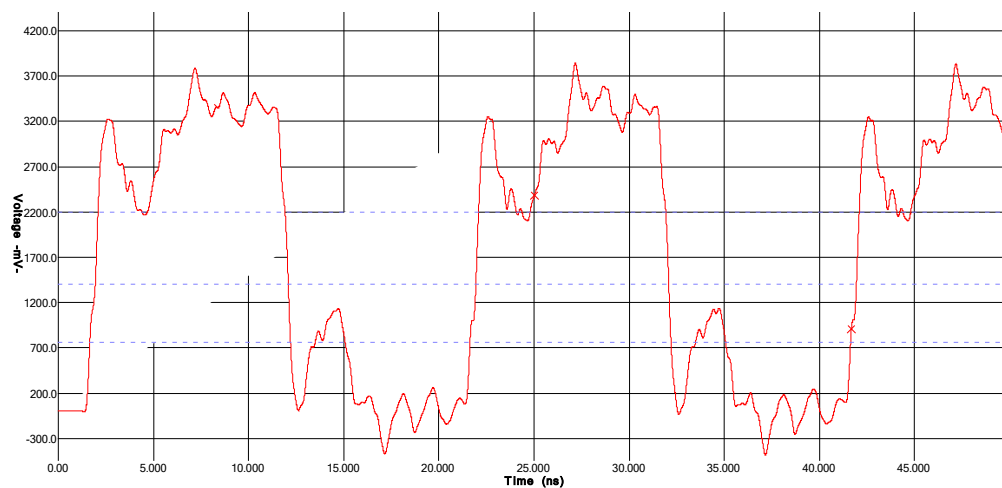


Figura 3.28. Reloj recibido en la primera carga (U2). Fuente propia

En la primera carga, U2, observamos un *undershoot* positivo y negativo elevado. Comprobando los casos fast y slow, resulta que el undershoot es incluso más acusado en el caso slow. Corremos el riesgo de detectar falsos flancos de reloj. El *overshoot* positivo y negativo está dentro de tolerancias (no supera $V_{cc}+0,7V$ o es menor de $-0,7V$). En definitiva, queremos un reloj más limpio.

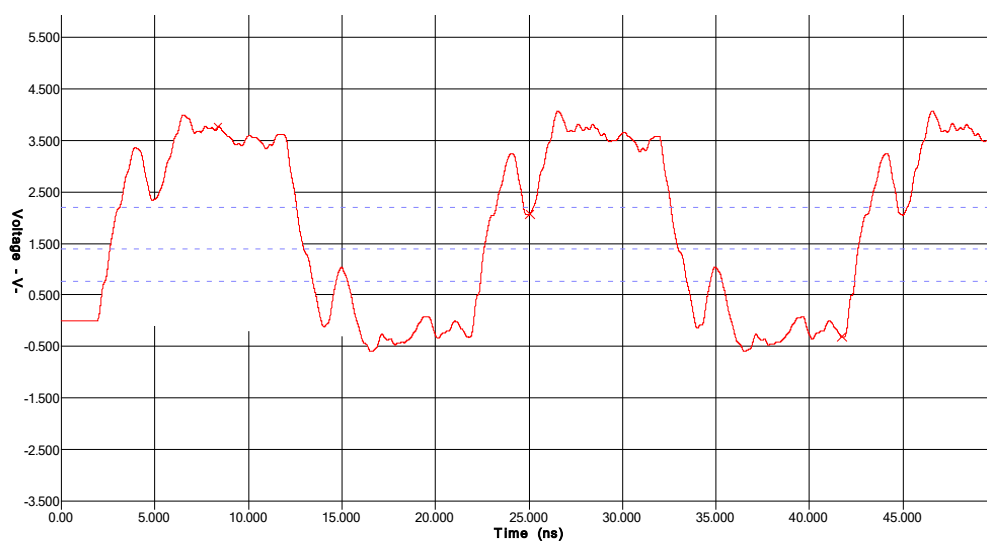


Figura 3.29. Reloj recibido en la segunda carga (U3). Fuente propia

Podemos extraer las mismas conclusiones de la forma de onda en U3, sólo que en este caso el *overshoot* es mayor y el *undershoot* no es tan acusado. Querríamos limpiar los flancos, igual que en caso anterior, ya que no son monótonos.

Las formas de onda en U4 y U5 son mucho más limpias. ¿Qué nos indica esto? Simplemente que el *undershoot* que observamos en U2 y U3 es fruto de una acumulación de reflexiones en los *stubs* de 3 cm.

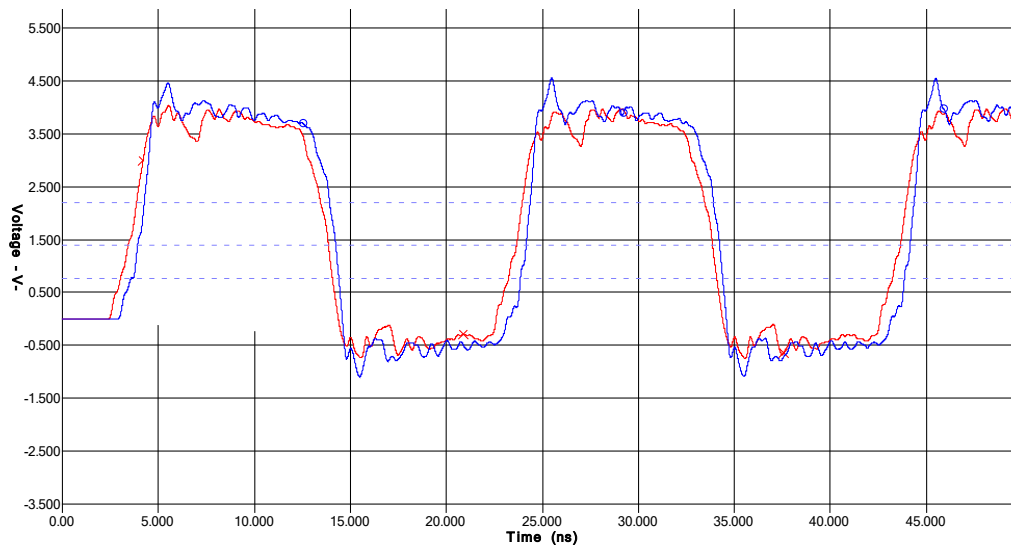


Figura 3.30. Reloj recibido en las dos cargas más lejanas (U4 y U5). Fuente propia

La onda incidente se propaga a lo largo de toda la línea. Pero fíjate en que la distorsión en la primera carga (U2, en verde) ocurre ya antes de que la señal llegue a la última carga (U5, azul). Lo que está ocurriendo es que cada pequeño stub de 3 cm y cada carga está reflejando parcialmente la señal. Parte de esta reflexión viaja hacia la derecha y parte hacia la izquierda, hacia la fuente. La señal se refleja con inversión de signo en la fuente, dando lugar a undershoot.

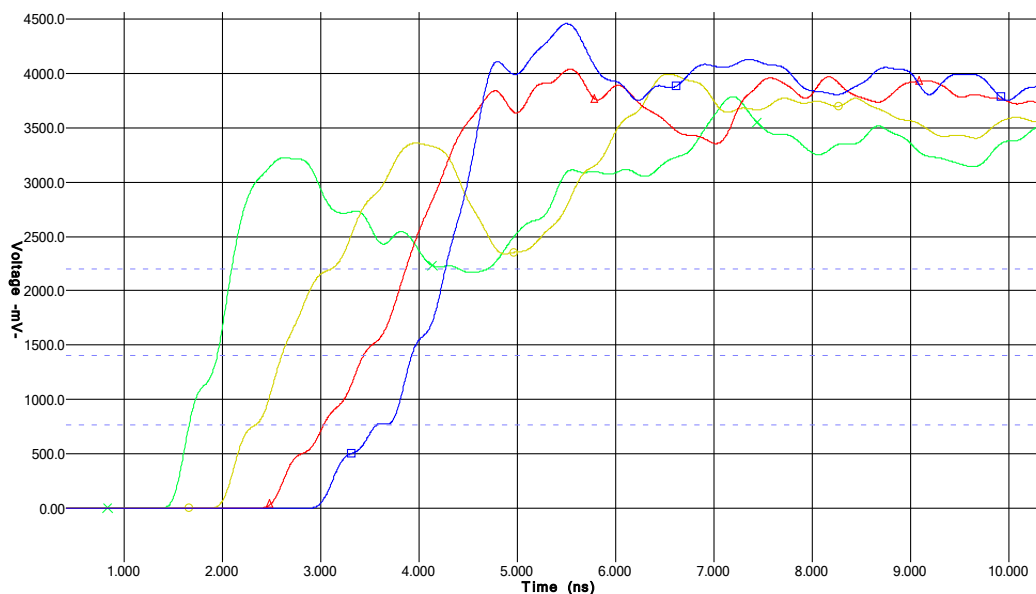


Figura 3.31. Vista ampliada del flanco de subida en las cuatro cargas. Fuente propia

El mismo efecto ocurre en todas las formas de onda, pero a medida que nos alejamos de la fuente, llega un momento en el que esta distorsión ya no afecta al flanco, simplemente porque llega más tarde (caso de U4 y U5). En U5 se presenta el problema de un *overshoot* por encima de $V_{cc}+0.7V$, lo que puede dañar los diodos de protección internos en los buffers. Esto nos da tres posibles estrategias para limpiar los flancos en U2 y U3:

- Reducir las reflexiones en la fuente
- Reducir las reflexiones en las cargas intermedias reduciendo el tamaño del *stub*.
- Reducir las reflexiones en las cargas intermedias añadiendo una pequeña resistencia serie en la carga

Para evitar el problema de *overshoot* excesivo en U5, podemos ensayar una terminación en la carga.

Atacando a la primera estrategia, vamos a ensayar añadir una resistencia serie en la fuente. Ya sabemos que no es garantía de éxito en una línea *multidrop* (con varias cargas), y si estudiemos su efecto vemos que empeora las cosas (Figura 3.32), ya que la reducción de amplitud del flanco incidente hace que el *undershoot* llegue más cerca del umbral de cruce del reloj (típicamente $V_{cc}/2$).

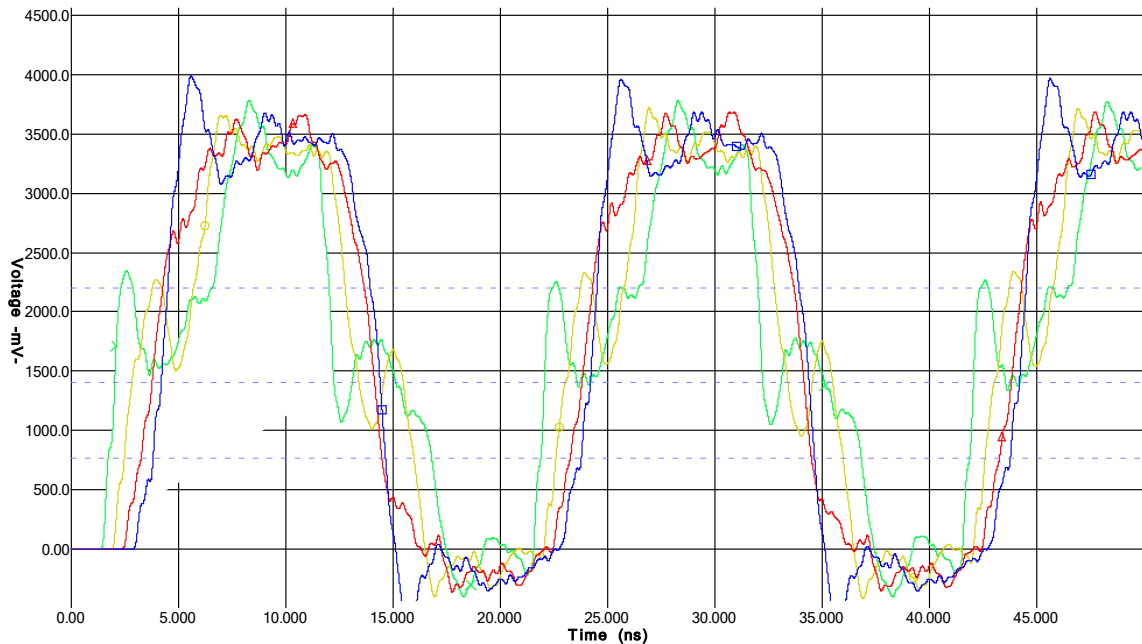


Figura 3.32. Primera estrategia fallida: añadir una resistencia serie en la fuente. Fuente propia

Así que eliminamos la resistencia serie en la fuente y atacamos la segunda estrategia: reducir los *stubs* por debajo de $\lambda/10$, que en este caso es de 2,5 cm. Vamos a dejarlo en 0,5 cm.

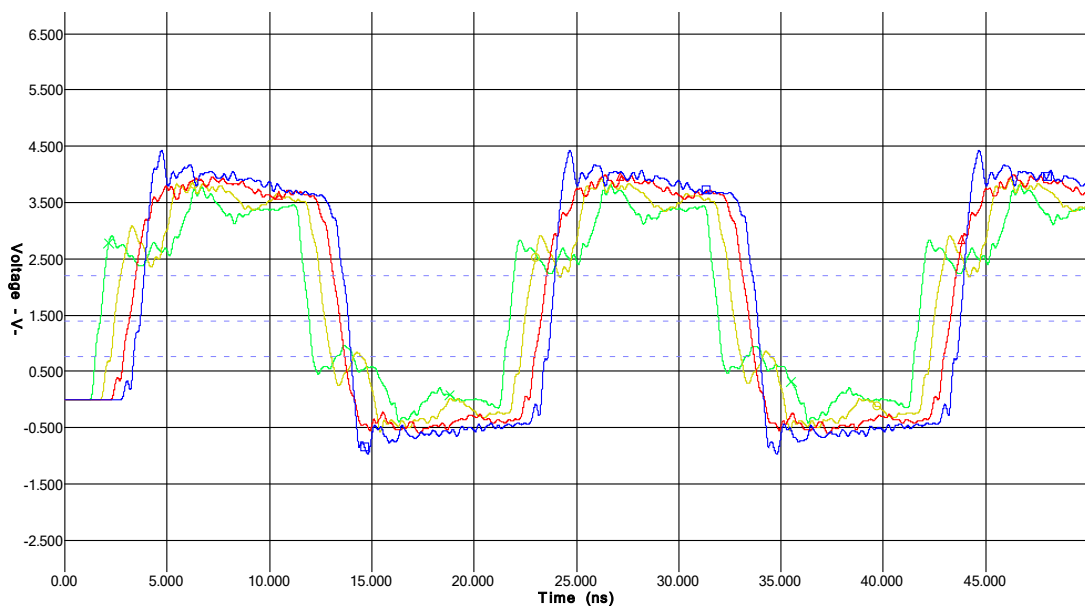


Figura 3.33. Segunda estrategia fallida: reducimos los *stubs* a 0,5 cm. Fuente propia

El resultado todavía presenta un *overshoot* excesivo y los flancos siguen siendo muy sucios.

Nos queda probar reducir la reflexión en cada carga mediante una resistencia serie de $100\ \Omega$. Los resultados (Figura 3.34) son aceptables, pues desaparece el *overshoot* excesivo en U5 y los flancos son mucho más limpios en U2 y U3.

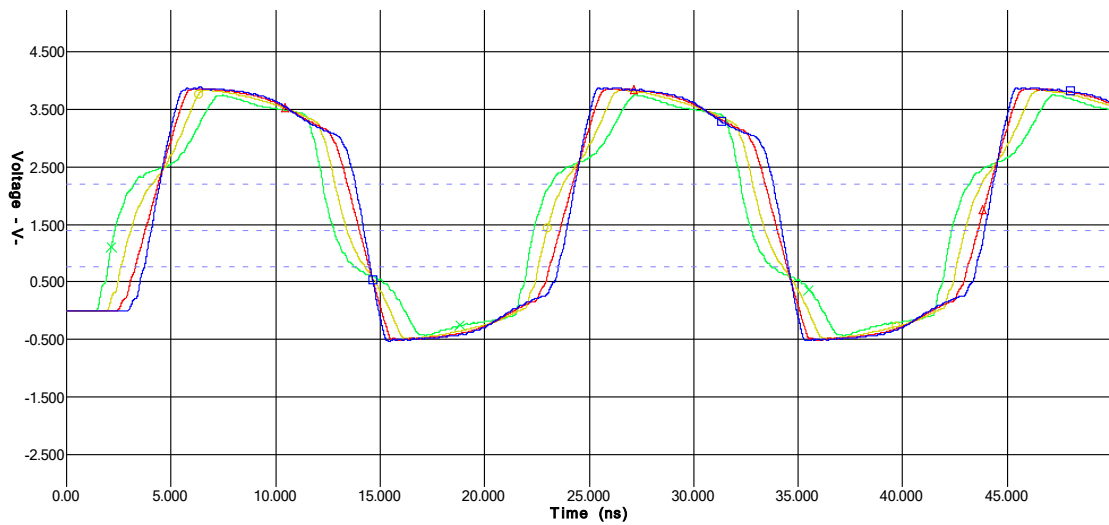


Figura 3.34. La tercera estrategia (añadir resistencias serie en las cargas) evita el overshoot excesivo. Fuente propia

Si añadimos a la estrategia anterior una terminación Thevenin en la última carga, conseguimos una pequeña mejora incremental (Figura 3.35).

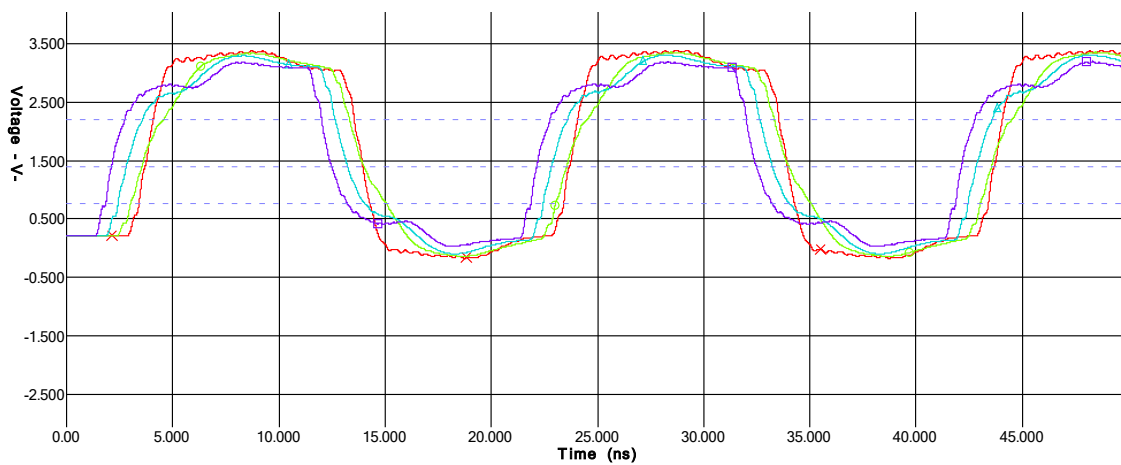


Figura 3.35. Añadimos una terminación Thevenin en la última carga. Fuente propia

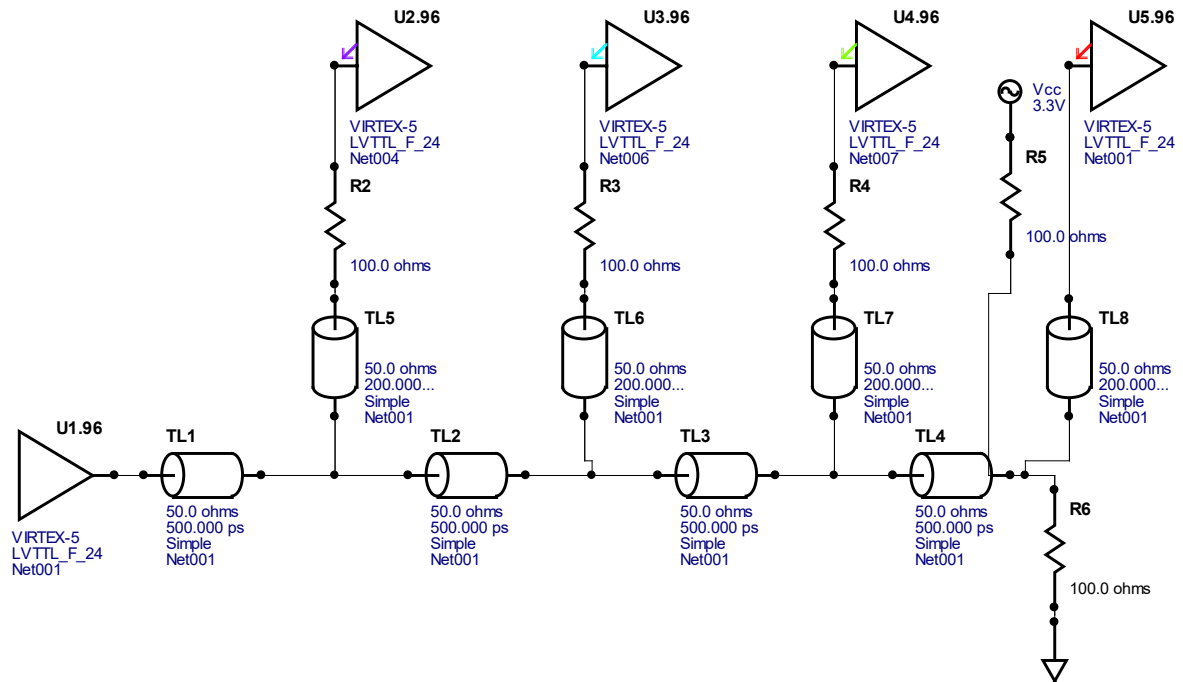
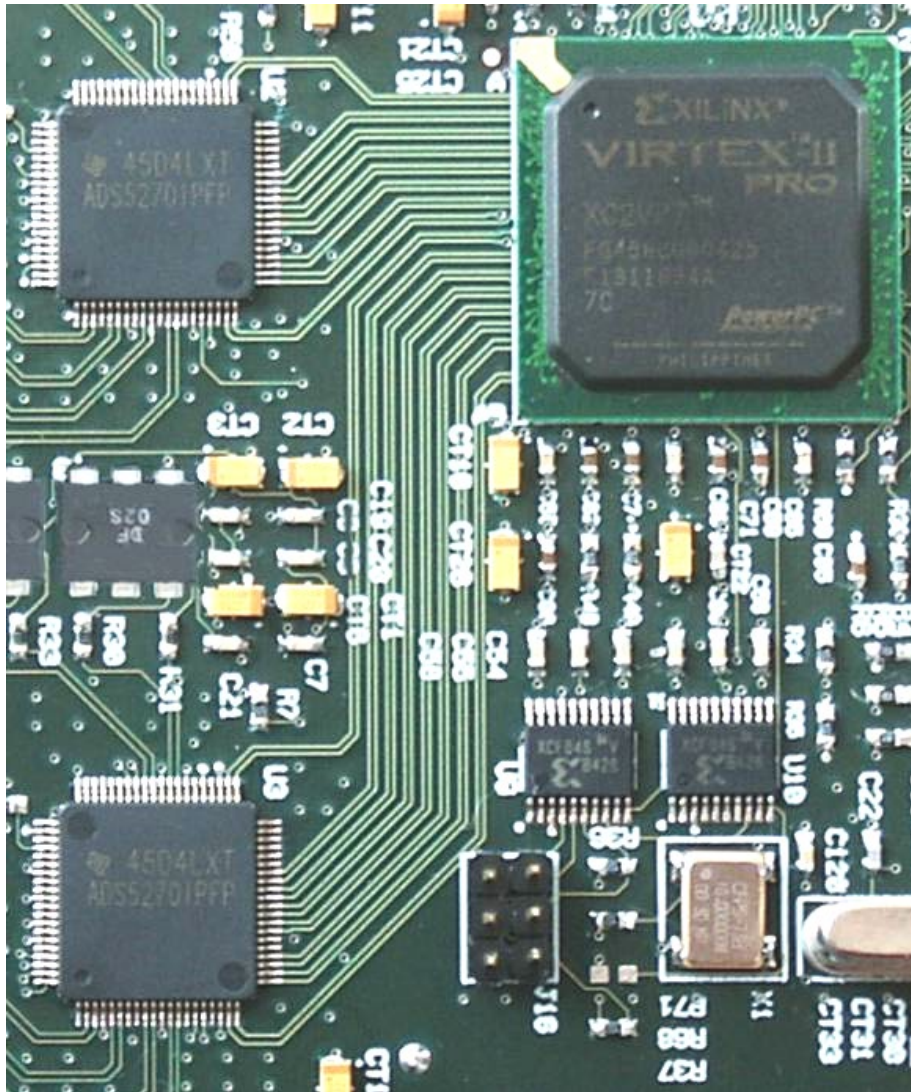


Figura 3.36. Solución final, que implica añadir cinco resistencias de 100 Ω . Fuente propia

Con un coste de un céntimo por resistencia, conseguimos limpiar la señal de reloj por cinco céntimos. No está nada mal.

Día 4. Líneas diferenciales



Ocho pares diferenciales entre cada uno de los dos ADCs (izquierda) y una FPGA (derecha) transportan y sincronizan 7 Gb/s ($7 \cdot 10^9$ bits por segundo) en un módulo de adquisición de datos para tomografía PET. Es parte de un diseño del autor, en los primeros años del siglo XXI (hacia 2005).

¿Por qué usamos líneas diferenciales?

En los años 90 casi no se usaban. Microprocesadores, memorias e interfaces de E/S se interconectaban mediante buses paralelo. En este esquema, un bus de datos de 8, 16, 32 o 64 líneas, junto a un bus de direcciones y varias líneas de control iban sincronizadas por un reloj. La frecuencia máxima del reloj estaba limitada por la topología del bus, que determinaba reflexiones, *skew* (diferencia de tiempos de llegada de una señal a las diferentes cargas) y *crosstalk* (acoplamientos entre líneas). Lo habitual era trabajar con buses a 40-60 MHz, a partir de ahí surgían problemas.

Un bus de 32 bit a 40 MHz permite transferir 1280 Mb/s, equivalente a (dividiendo por 1024) 1,25 Gb/s o (dividiendo por 8) 156,25 MByte/s (MB/s). Pero estos 1,25 Gb/s requerían un área elevada de rutado en el PCB, donde a las 32 líneas de datos había que sumar tal vez otras 20 entre líneas de control, direcciones y reloj. Dividiendo bits/s entre número de líneas resulta algo menos de 25 Mb/s por línea. Hoy en día, un solo par diferencial puede superar 10 Gb/s.

En sistemas de altas prestaciones no eran los microprocesadores sino los buses, con su topología multipunto, quienes imponían el límite a lo que se podía hacer. Vamos a hacer un recorrido histórico [2], que comienza en los años 60, por este camino que ha resultado ser una vía muerta y que ha sido reemplazado por conexiones punto a punto mediante líneas diferenciales de alta velocidad.

La historia de los buses de altas prestaciones en física de altas energías

Un campo donde se ha exprimido hasta el último bit por segundo de prestaciones a los buses digitales es la física de altas energías (HEP, *high-energy physics*). La elevada cantidad de datos generados por segundo, el gran número de canales de datos y la necesidad de llevar todos los datos a uno o unos pocos nodos centrales ha supuesto un reto no pequeño para los diseñadores de sistemas de adquisición de datos de altas prestaciones. Comencé mi carrera investigadora precisamente en este campo, a finales de los 90, en el CERN, usando tecnologías y soluciones que hunden sus raíces en los años 60.

Años 60 y 70

En un experimento de física de altas energías típico de finales de los 60 y años 70, todos los canales de datos se leían desde un único miniordenador, en una arquitectura que no permitía paralelismo ni tenía capacidad para ir más allá de unos pocos kByte/s. Los módulos y las interconexiones no estaban estandarizados, provocando confusión e ineficiencia. Esta situación llevó a un esfuerzo de estandarización, que comenzó por el *front-end* (electrónica cerca de los sensores). En 1964 aparece el [estándar NIM](#), que definía conectores, dimensiones de los módulos, niveles eléctricos de las señales y características de la fuente de alimentación. Pero no definía el *backplane* (panel posterior del chasis, que podría usarse para interconectar los módulos y facilitar la adquisición de datos). De modo que las interconexiones se realizaban por el panel frontal, entre módulos, mediante cables. NIM se usa todavía en algunos laboratorios, así de profunda es la huella que dejó.

Por si tienes curiosidad, NIM (tal y como se usa más habitualmente, con lógica rápida negativa) especifica impedancias de línea de 50 ohmios. Las cargas están terminadas a 50 ohmios para evitar reflexiones. El 1 y el 0 lógico se definen con tensiones de 0V y -0,8V y los flancos son tan breves como 1 ns. Los cables son apantallados y de 50 ohmios. Como ves, se mimaba la señal. Era vital no degradarla para poder utilizar la marca temporal que implica un flanco en NIM: la detección de un electrón, de un fotón o de un muon.

La estandarización llegó al *back-end* (adquisición de datos, interconexión de datos entre módulos y con el ordenador) en 1969 con el [estándar CAMAC](#) (*Computer Automated Measurement And Control*). En este estándar, un módulo controlador podía comunicarse con hasta 25 módulos en el mismo chasis interconectados por un *backplane*, que contenía un bus paralelo. A partir de los 80 se usó la modularidad de CAMAC para agrupar varios chasis en ramas y poder así aumentar en número de canales (sensores). La comunicación se basaba en comandos, que iban dirigidos (direccionados) a un módulo en concreto indicando rama, chasis, módulo en el chasis y número de función en el módulo.

En los años 70 era habitual usar NIM en el front-end (lectura de señales analógicas y digitalización) y CAMAC como bus de read-out para llevar datos a un ordenador. Las limitaciones de CAMAC (10⁶ palabras por segundo, siendo las palabras de 16 o 24 bit, lo que supone un máximo de 3 MByte/s) chocaron con la insaciable demanda de ancho de banda por parte de los físicos que diseñaban los experimentos. Además, la aparición de los primeros microprocesadores abrió las puertas al paralelismo (procesar los datos de forma distribuida en varios nodos a la vez), algo que antes era imposible, algo para lo que CAMAC estaba mal dotado.

Años 80

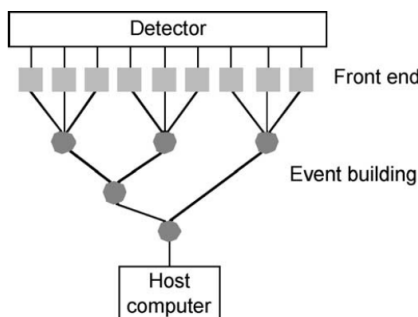


Figura 4.1. Arquitectura en árbol típica de los sistemas de adquisición de datos de los experimentos HEP de los años 80. Fuente propia

La necesidad de mayor ancho de banda (ya hemos comentado que en este aspecto los físicos son insaciables) y de realizar un procesamiento distribuido (lo que requiere no un máster, sino múltiples elementos en el bus que puedan iniciar transferencias), trasladó el protagonismo a estándares como **Fastbus** (1983) y **VME** (este último no se estandarizó hasta 1987). Por dar una referencia, VME64 (versión de 64 bit de ancho del bus de datos) permite alcanzar 40 MByte/s efectivos, doblando las prestaciones de versiones anteriores. Cabe recordar que, en los años 80, los microprocesadores disponibles (como la familia 68000 de Motorola en la que se basó el desarrollo de VME) eran de 16 bit, si bien posteriormente surgieron versiones de 32 bit.

Un *backplane* VME64 proporcionaba a cada módulo del chasis tres conectores de 64 líneas, si bien muchas eran de masa y alimentación. Y esto para permitir 40 MByte/s. Que es menos de la mitad de lo que consigues hoy en día con cable RJ45 (cuatro pares diferenciales) y Gigabit Ethernet.

Años 90

La década comienza con una coexistencia de CAMAC (había muchos expertos y hardware disponible), Fastbus y VME, con una marcada tendencia de dominio de este último. Pero **la arquitectura de bus había chocado con el muro**. Ya hemos comentado que un simple cable de red duplica (o triplica) lo

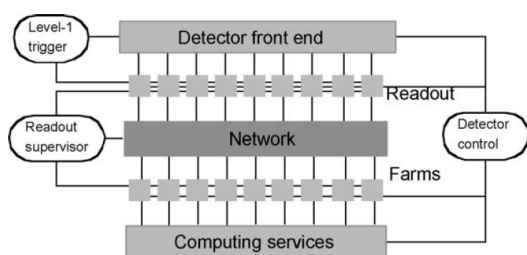


Figura 4.2. Arquitectura de red típica de los años 90. Fuente propia

que podemos obtener con el aparatoso conexionado VME64 (63x3 conexiones). Era necesario alcanzar varios Gb/s, y es en los 90 cuando comienza a abandonarse la arquitectura en bus típica de las décadas anteriores a favor de conexiones punto a punto de alta velocidad basadas en líneas diferenciales y conmutadores de red para reencaminar los paquetes de datos entre una fuente y un destino. Un grupo de PCs de coste moderado (*computing farm*) recibe los datos de la red.

Comienzos del siglo XXI

Obviamente las conexiones entre módulos consisten en líneas punto a punto diferenciales (Gigabit Ethernet y actualmente 10-Gigabit Ethernet) ópticas o de cobre. Entre los módulos de *front-end* y de *back-end*, uno o más niveles de conmutación (*switches*) se encargan de encaminar los datos.

Pero las líneas diferenciales no quedan limitadas a las conexiones a partir del *front-end*: el propio *front-end* (procesado analógico de señales, digitalización y formateo de datos) ha pasado de no contener líneas diferenciales a que éstas sean típicamente entre un 10% y un 50% en la mayoría de los diseños. Conexiones USB, SATA e interfaces diferenciales entre circuitos integrados son habituales en los PCBs actuales.

Cabe destacar el papel que el bus PCI jugó, sobre todo en las placas base de PC, en los años 90. Desde la especificación original (32 bit, 33 MHz, lo que supone un máximo teórico de 1 Gb/s) hasta sus últimos estertores (64 bit, 133 MHz, 8,31 Gb/s teóricos) dominó las placas base de ordenador durante algo más de una década. La integridad de señal era tan delicada que en su versión más potente sólo admitía el microprocesador y dos cargas, todo a pocos centímetros de distancia y por supuesto sin salir de la caja del ordenador.

En 2004 aparece PCI express: ¡muerte al bus paralelo! ¡larga vida al bus serie de alta velocidad con pares diferenciales! En la actualidad (2020) las placas base de ordenador están basadas en PCI express 3.0 y permiten varios slots x16, lo que equivale a 126 Gb/s por slot.

De este modo, los buses serie de alta velocidad, ¡no sólo son los reyes en las comunicaciones entre cajas, sino también dentro de la caja!

En la fotografía que abre la lección de hoy, dos convertidores A/D envían cada uno ocho pares diferenciales a una FPGA: 6 canales de datos (las palabras son de 12 bit y se envían en serie), un reloj (*bit clock*) y un reloj lento (marca el inicio de cada palabra de 12 bit). La tasa de datos generada por la digitalización de los 12 canales (muestreando cada canal a 50 MHz) supone 7.200 Mb/s, o 7 Gb/s. Y estamos utilizando únicamente 16 pares diferenciales, 32 pistas. No está nada mal. Sobre todo, si tenemos en cuenta que el equivalente con buses paralelo (12 líneas por canal, 12 canales, y un par de señales de reloj, todo ello a 50 MHz) supondría rutar 146 pistas en el PCB. Nueve veces más.

Resumiendo lo anterior

Una línea no diferencial en un bus en un *backplane* (bus multipunto que abarca entre 20 cm y 1 m de longitud con hasta 20 cargas aproximadamente) no suele ser viable por encima de 40-60 MHz de frecuencia de reloj. En interconexiones entre integrados en un mismo PCB, el límite está en torno a 120 MHz. Usando conexiones diferenciales punto a punto podemos alcanzar prestaciones mucho mayores, tanto en distancias de pocos centímetros como a varios metros. Y, además, reduciendo el consumo de potencia.

¿Qué es una conexión diferencial?

Comencemos repasando lo que ya sabemos, porque algo sobre transmisión diferencial te habrán contado en instrumentación electrónica cuando te hablaban de telemidida. Una transmisión diferencial requiere de tres elementos:

- **Un transmisor (*driver*) simétrico.** Es decir, que la señal de interés viaje en ambos hilos, pero con polaridad opuesta. Para mantener la señal siempre por encima de 0V (no solemos alimentar los integrados con tensiones negativas, ¿cierto?) se añade un nivel de continua (*offset*) sobre el que va superpuesta la señal de línea.
- **Un medio de transmisión simétrico** (mismo retardo y misma impedancia característica en cada línea)
- **Un receptor simétrico**, que realice la resta de las dos líneas y rechace el ruido acoplado en modo común.

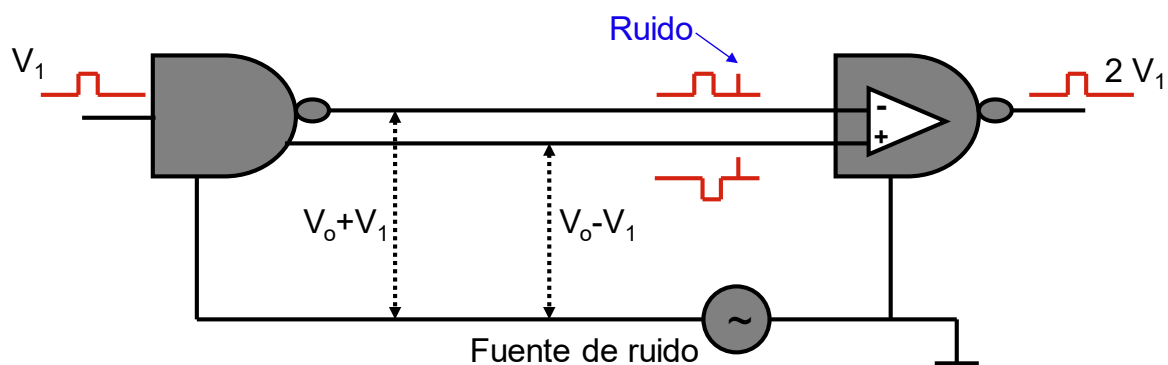


Figura 4.3. Concepto de transmisión simétrica o diferencial. Fuente propia

En la Figura 4.3, el *driver* tiene dos salidas separadas, con un offset fijo de tensión V_0 . El voltaje V_1 (señal de interés) se suma a una línea y se resta en la otra. La tensión recuperada en el receptor es: $V_{\text{diff}} = (V_0 + V_1) - (V_0 - V_1) = 2 \cdot V_1$. La tensión en modo común es V_0 .

Un ejemplo de estándar diferencial: LVDS

LVDS (*Low-Voltage Differential Signaling*) es un estándar (ANSI/TIA/EIA-644 de 2005, revisado en 2001) que define una transmisión balanceada de hasta 655 Mbps (actualmente se alcanzan de 1 a 3 Gb/s y se usa en enlaces de alta velocidad como SATA) en un par de pistas PCB o en un par de cables. La transmisión se realiza a corriente constante (aproximadamente 3,5 mA, si bien es configurable en muchos casos) y el sentido de la corriente diferencia al 0 del 1 lógico.

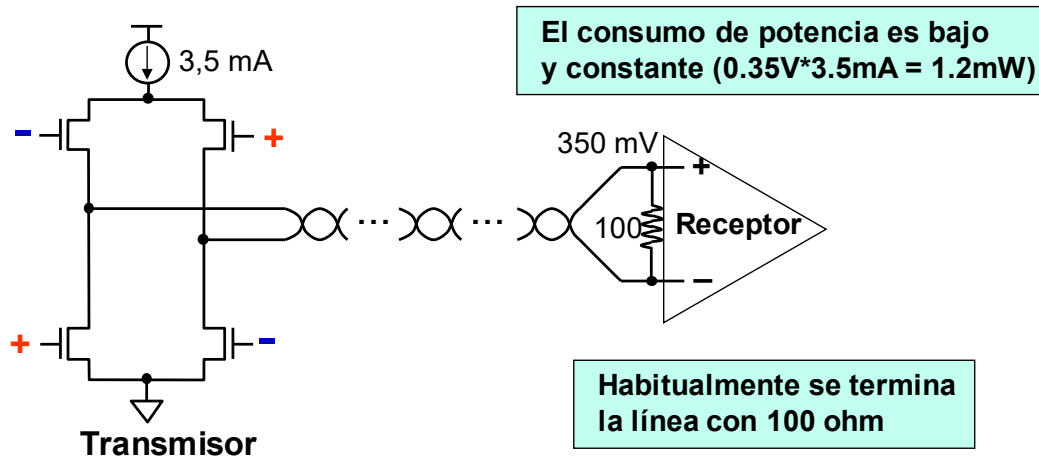


Figura 4.4. El estándar LVDS. Fuente propia

El *driver* LVDS contiene cuatro transistores. En todo momento dos de ellos están en corte y dos en saturación. Invertiendo el control de los transistores, se invierte la dirección de la corriente y por tanto la diferencia de tensión en la resistencia de terminación de 100 ohmios en el receptor. Esta resistencia tiene una doble función: genera una diferencia de tensión en el receptor e iguala la impedancia de la carga con la característica de la línea.

CML (current mode logic)

Se usa como capa física en interfaces como HDMI y DVI, diseñada para trabajar entre 312,5 Mb/s y 3,125 Gb/s. Se diferencia de LVDS principalmente en la estructura del *driver*, formada por dos transistores bipolares polarizados en la región activa, no en saturación como los transistores CMOS de LVDS, lo que permite una conmutación más rápida pese a tener una excursión de señal de aproximadamente el doble.

Desde el punto de vista del diseñador de sistemas, se trata exactamente igual que un par diferencial LVDS: la terminación es de una resistencia entre ambas líneas de valor igual a la impedancia diferencial de línea.

Es fácil conectar drivers y receivers CML y LVDS entre sí. Como usan diferentes tensiones de *offset* (o de modo común en continua, como prefieras), hay que usar condensadores de acoplamiento AC y si es necesario redes de resistencias en el receptor para sumar la tensión de offset necesaria. La nota de aplicación [3] trata sobre esta adaptación.

Cuando el reloj forma parte de los datos

Una técnica muy interesante y frecuente es codificar el reloj con los datos, de modo que el receptor pueda, a partir de la corriente de unos y ceros recibida, recuperar tanto uno como los otros. Como nada es gratis, hay que pagar un precio. En el caso más habitual, la [codificación 8b/10b](#), cada grupo de 8 bit que se quiera transmitir se sustituye por un código de 10 bit (llamado símbolo) que garantice suficientes transiciones 0-1 y 1-0 en la línea. Es decir, estamos aumentando un 25% los bits por segundo en la línea (no es un precio pequeño) a cambio de evitarnos un segundo par diferencial para el reloj.

¿Por qué queremos asegurar suficientes transiciones en la corriente de unos y ceros? Para que un PLL (*phase-locked loop*, bucle de enganche de fase) en el receptor sea capaz de recuperar el reloj original.

Obtenemos una ventaja adicional con la codificación 8b/10: el número de unos y ceros en la línea es el mismo, de modo que la señal tiene un valor medio constante, que no sufre variaciones. Esto permitirá introducir condensadores en serie sin producir problemas, como veremos más adelante. Por cierto, estos condensadores son obligatorios, por ejemplo, en las conexiones SATA, PCI express, HDMI, Gigabit Ethernet, USB 3.0, y DVI. Lo dicho, volveremos a esto antes de acabar la lección de hoy.

¿Cómo funciona la codificación 8b/10b?

Cada byte (256 combinaciones) se sustituye por un símbolo de 10 bit (1024 combinaciones). De modo que para cada byte hay cuatro posibles símbolos a elegir (si bien doce de los 1024 se reservan como símbolos de control).

La regla es evitar que en ningún momento haya más de 5 unos o ceros seguidos. Y que el número de unos y ceros en una cadena de 20 bit no difiera en más de dos.

¿Qué es una línea diferencial en un PCB?

¿Qué hace falta para tener un par diferencial en un PCB? Recordemos (página 81) que había tres condiciones a cumplir: *driver* simétrico, un par de líneas simétricas y un receptor simétrico. Un par de líneas simétricas son aquellas que tienen mismo retardo e impedancia característica de línea. Es decir, que la señal se propaga exactamente igual por cada línea del par.

Según lo anterior, dos líneas de 50 ohmios, rutadas en la misma capa o en capas diferentes (siempre que el retardo, es decir, el cociente entre su longitud y velocidad de propagación sea el mismo), juntas o separadas (como en la Figura 4.5), forman un par diferencial.

Quedan dos detalles para mantener la simetría:

- Si hay terminaciones de línea, deben ser iguales. En la Figura 4.5, cada línea lleva una terminación de 50 ohm en la carga a masa. Se mantiene la simetría.
- Recuerda que la pista es sólo el 50% del camino de la señal. El otro 50% es el camino de las corrientes de retorno. En la Figura 4.5 vemos un plano continuo bajo las pistas, sin discontinuidades que introduzcan asimetrías. Volveremos sobre este aspecto más adelante.

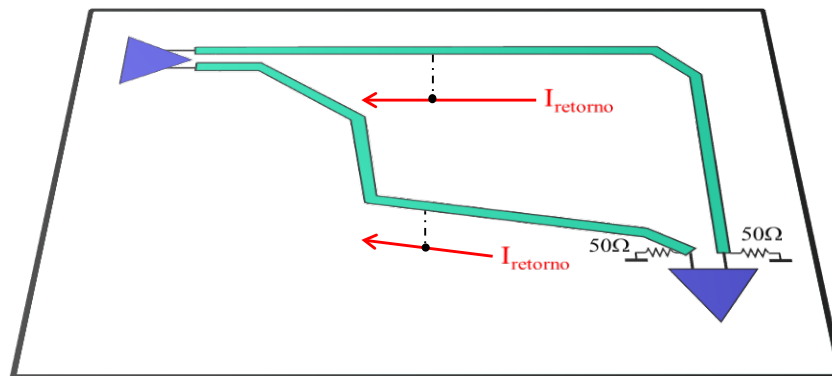


Figura 4.5. ¿Es esto un par diferencial? Fuente propia

Podemos concluir que lo que vemos en la figura anterior es ciertamente un par diferencial.

¿Entonces, por qué en un PCB siempre vemos las dos pistas juntas?

Porque así controlamos mejor la diferencia de retardos entre las dos líneas del par, reducimos el área que el par ocupa en el PCB y aseguramos que el ruido acoplado es el mismo en ambas líneas.

¿Y por qué terminamos cada línea con 50 ohm a masa en la Figura 4.5 si habitualmente lo que veo en un diseño son 100 ohmios entre las dos pistas del par?

Como líneas de 50 ohmios, lo correcto es terminar cada línea a 50 ohmios (Figura 4.6, izquierda). Que es exactamente lo mismo que la Figura 4.6, centro. Pero cabe hacer una simplificación, dado que las dos líneas del par conducen corrientes en sentido opuesto (e idealmente están alineadas). Es decir, la señal viaja en modo normal, no en modo común. La corriente que llega a la carga por una línea es la misma (idealmente) que la que vuelve por la segunda y no es necesaria la conexión a masa, porque no hay modo común. De este modo, basta (si la señal es diferencial) con la simplificación que aparece en la Figura 4.6, derecha. Que es más sencilla de rutar en el PCB.

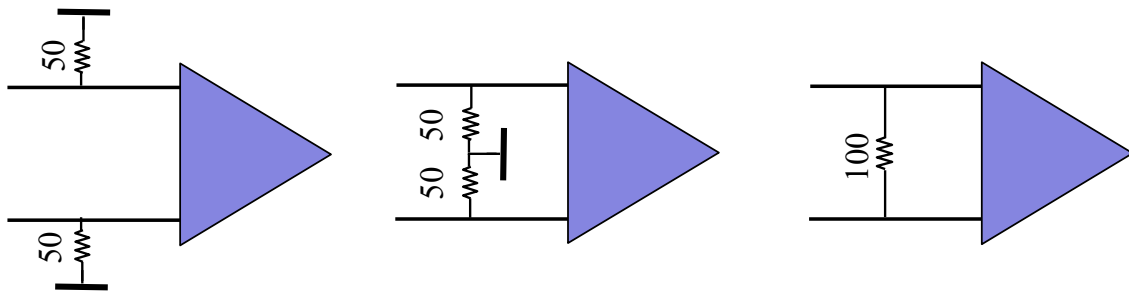


Figura 4.6. ¿Por qué se usa una única resistencia para terminar un par diferencial? Fuente propia

Pero esta solución no termina el modo común, que sí quiere cerrarse a través de masa. Y en las señales reales acaba habiendo modo común por **conversión de modo** (un término que aclararemos dentro de un poco, paciencia), queramos o no. Si hay suficiente conversión de modo, como no estará correctamente terminado el modo común, habrá reflexiones y la señal se distorsionará. Como este efecto suele ser pequeño, se opta por dejar una única resistencia de 100 ohmios. De modo que lo que vemos en un PCB se parece más a la Figura 4.7 (líneas del par juntas y terminadas con una única resistencia) que a la Figura 4.5.

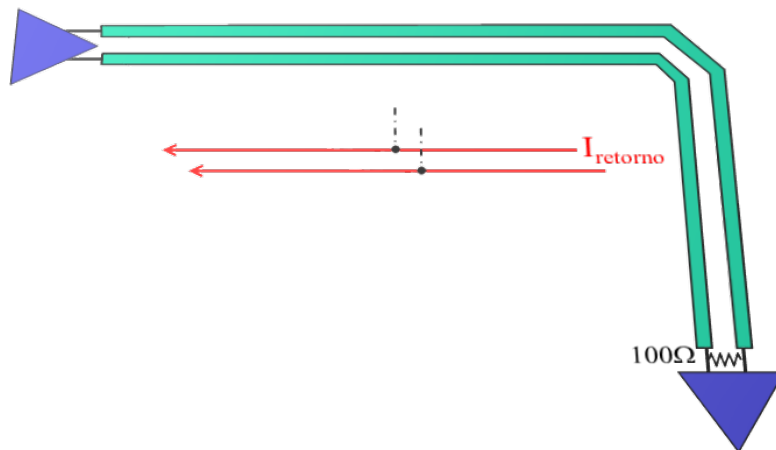


Figura 4.7. Algo más parecido que la Figura 4.5 a lo que vemos en un PCB . Fuente propia

¿Cómo de crítico es igualar el retardo de ambas líneas del par?

Ni las dos salidas del *driver* diferencial tienen los flancos perfectamente alineados, ni las dos pistas en el PCB tienen exactamente el mismo retardo. Tampoco el amplificador en el receptor está exento de asimetría. Caba apuntar dos causas más de diferencia en la posición de los flancos: la diferencia de ruido externo acoplado a cada pista (interferencia electromagnética) y la diferencia de ruido interno acoplado a cada pista (ruido que genera tu propia electrónica). Todas estas son causas de *skew* (diferencia de retardos) entre las líneas del par. Dejando el ruido aparte, lo único que puedes controlar es el rutado. Y con frecuencia compensarás los desajustes en *driver* y receptor con una ligera diferencia de retardo en el rutado.

Bien, ¿y cuánto *skew* es aceptable? Depende. Pero vamos a fijar un límite máximo que usaremos como referencia. Mira la Figura 4.8. Como punto de partida, el máximo *skew* permitido equivale a la duración del flanco en el receptor.

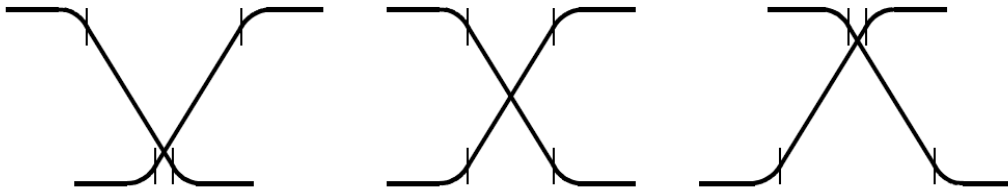


Figura 4.8. El máximo *skew* (diferencia de retardos) entre líneas del par diferencial corresponde a la longitud del flanco, si bien las recomendaciones de los fabricantes son a menudo mucho más restrictivas.

Por ejemplo, en una transmisión a 2 Gb/s, con un flanco de 100 ps (20% de bit a 2 Gb/s), suponiendo 15 cm/ns de velocidad de propagación, disponemos de hasta 1,5 cm de margen para igualar longitudes. Por tanto, el ajuste debe ser bastante mejor que $\pm 7,5$ mm. Recuerda, este es un valor máximo para la suma de todas las causas de *skew*.

Pero hay razones para ir más allá del criterio de la duración del flanco:

- Si las señales han de propagarse por un cable o conector sin masa. En este caso, la aparición de modo común debido a asimetrías en las líneas del par (causa de conversión de modo diferencial a modo común) puede bastar para que tu producto no supere los ensayos de EMC conducidas.
- Si trabajamos por encima de 1 Gb/s y usamos la típica terminación de 100 ohmios, cuando las dos señales no se cruzan en el punto medio del flanco, durante un pequeño tiempo debería circular corriente por masa. Al no poder hacerlo, uno de los flancos quedará retardado. Este efecto ha de tenerse en cuenta en enlaces más rápidos de 1 Gb/s.

Una posible solución a lo anterior consiste en reemplazar la terminación de 100 ohmios por dos resistencias de 50 ohmios unidas a masa o por un pequeño condensador.

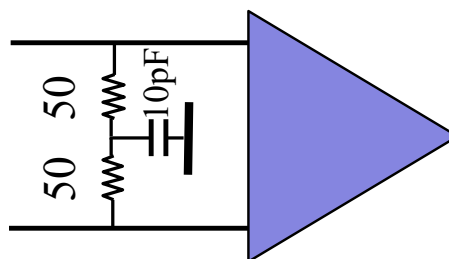


Figura 4.9. Una forma de minimizar el efecto de la presencia de señales de modo común de alta frecuencia en el par diferencial. Fuente propia

Generalmente, los fabricantes piden en hojas de datos y notas de aplicación márgenes mucho menores (de hasta 5 mils, que son 127 micras).

Skew de rutado para HDMI, PCI Express y SATA

En [4] se dan recomendaciones para el máximo skew entre líneas del mismo par para diferentes estándares. Para HDMI, el máximo es el 15% del tiempo de bit. Para PCI Express y SATA se recomienda no superar 5 mils (127 micras).

¿Cómo se ajusta el skew en el rutado?

En la Figura 4.7, el par hace un giro, provocando que una pista sea más corta. Si en el rutado hay más giros y no se compensan, puede ser necesario alargar la pista más corta. El modo correcto de hacerlo mediante serpentina se indica en la Figura 4.10 [4].

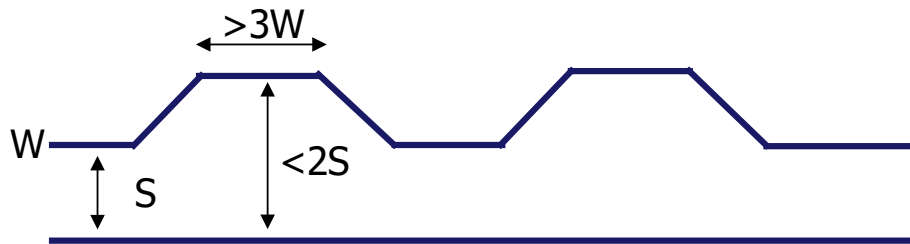


Figura 4.10. Ecuación de longitudes de líneas de un par mediante serpentina. Fuente propia

Hay tres reglas sencillas que debemos seguir:

- La serpentina debe ser rutada a 45° o con tramos curvos. Evitaremos tramos a 90°.
- La anchura de cada tramo debe ser de al menos tres veces la anchura de la pista. De este modo nos aseguramos de que la señal no salta de tramo en tramo (por efecto capacitivo) por estar demasiado juntos.
- La separación máxima entre pistas del par no debe superar el doble de la nominal. De este modo limitamos los cambios en la impedancia diferencial. Por cierto, esta regla está mal en [4], un error tipográfico.

¿Dónde se hace el ajuste de skew?

No vale hacerlo en cualquier punto del par. No sólo se trata de igualar retardos. Recuerda que queremos mantener la simetría en todo momento, de modo que idealmente, hacemos la corrección justo antes o justo después de la asimetría. En la Figura 4.11, la asimetría se produce en el extremo izquierdo de la línea. La corrección cabe hacerla en ese lado, tan cerca como sea posible, en lugar de hacerlo en el extremo opuesto.

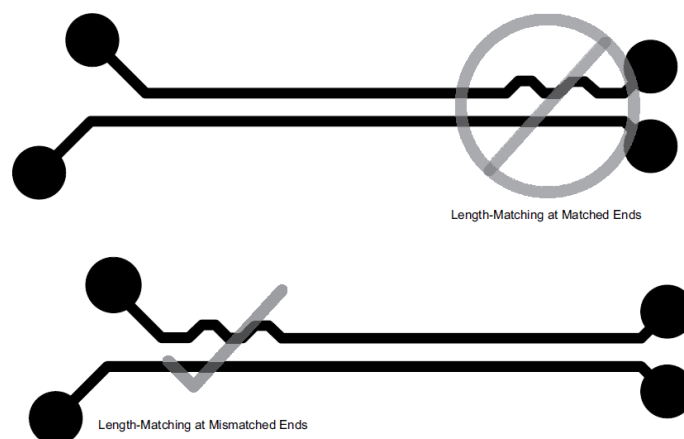


Figura 4.11. Figura 5 de [4].

¿Por qué y dónde se ponen los condensadores de bloque en DC?

Si los niveles de tensión de continua (*offset*) son distintos en transmisor y receptor, ya sea porque se trata de distintas tecnologías (como es el caso de LVPECL y LVDS) o porque estamos interconectando dos PCBs con referencias de masa ligeramente distintas, nos encontramos con un problema.

Pero con un problema cuya solución está prevista: se colocan unos condensadores en serie en ambas líneas, independizando el nivel de *offset* de ambos extremos de la línea.

No hay preferencia en cuanto a su ubicación (junto al receptor o al transmisor). En [5], figura 6, las curvas de pérdidas con condensadores y sin condensadores no reflejan ninguna diferencia significativa, indicando que estos condensadores no degradan apreciablemente la señal.

Pero (siempre hay un pero), eliminar el nivel de continua de la señal requiere que ésta tenga un nivel DC constante, lo que sólo se logra con algunas codificaciones como por ejemplo la 8b/10b (recuerda lo que hablamos en la página 82).

Impedancia diferencial

Llevamos varias páginas de la lección de hoy y hemos mencionado ya varias veces la impedancia diferencial. Vamos a definirla y ver cómo ajustar la geometría en el rutado para ceñirnos a un valor determinado, que por lo general será 90 ohmios (como USB 2.0) o 100 ohmios (Gigabit Ethernet).

Partimos de la impedancia característica de una línea no diferencial, $Z_0 = V/I$. Como tenemos dos líneas, que llamaremos 1 y 2, resulta que $Z_0 = V_1/I_1 = V_2/I_2$, y como en un par diferencial resulta que $I_1 = -I_2$, entonces $Z_0 = V_1/I_1 = -V_2/I_1$.

Llamamos **impedancia diferencial** al cociente $(V_1 - V_2)/I_1$, y como $V_1 = -V_2$, resulta $Z_d = 2 \cdot V_1/I_1 = 2 \cdot Z_0$.

En una línea diferencial en la que las dos líneas del par no se influyen (no están acopladas), la impedancia diferencial es dos veces el valor de la impedancia característica de cada línea del par: $Z_d = 2 \cdot Z_0$.

De esta forma, dos líneas de 50 ohmios crean un par diferencial de 100 ohmios de impedancia diferencial.

Fíjate en que hemos dicho “...en la que las dos líneas del par no se influyen...”. Para cumplir esta condición han de estar bastante separadas. Porque a medida que las aproximamos, la distribución de los campos eléctricos y magnéticos cambia, alternando el valor de la impedancia característica de cada pista y por tanto el valor de la impedancia diferencial.

Hay que hacer otra puntualización: estamos asumiendo que el par contiene sólo señal en modo diferencial, y sobre ella se define la impedancia diferencial. Si consideramos sólo el modo común, hay que hablar también de una **impedancia del par en modo común**, definida como $Z_{mc} = V_{mc}/(2 \cdot I_{mc}) = 1/2 \cdot V_1/I_1 = Z_0/2$.

En una línea diferencial en la que las dos líneas del par no se influyen (no están acopladas), la impedancia que ve el modo común es la mitad del valor de la impedancia característica de cada línea del par: $Z_{mc} = Z_0/2$.

¿Qué ocurre cuando las líneas del par se influyen?

Si las líneas se acercan, aparece un acoplamiento entre ambas. En la Figura 4.12 observamos que la impedancia en modo común aumenta y la impedancia en modo normal (diferencial) disminuye. Vamos a estudiar qué le ocurre al modo diferencial.

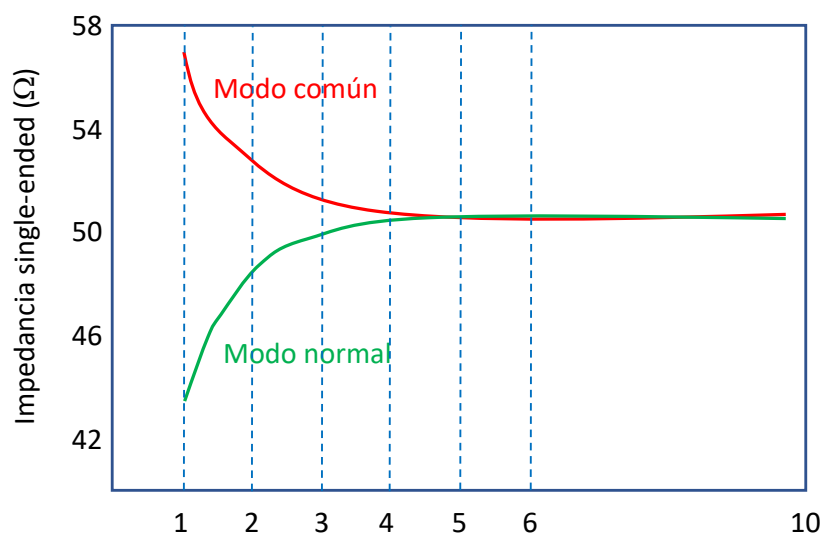


Figura 4.12. Efecto de acercar las pistas sobre las impedancias en modo común y modo diferencial. Eje horizontal: separación entre pistas (normalizada a la anchura de pista). Fuente propia

En la Figura 4.13, parte superior, partimos de dos pistas *stripline*, cada una de 162 micras de anchura y 50 ohmios de impedancia (Z_{od} , impedancia característica) separadas (borde a borde) 1 mm. Excitamos el par de pistas con una señal en modo diferencial. Podemos observar la distribución de campo eléctrico y de campo magnético. Vemos que sólo hay una línea de campo eléctrico entre las dos pistas: el acoplamiento es débil. La matriz de impedancia contiene 50 ohmios en la diagonal (Z_{11} y Z_{22}) y las impedancias cruzadas (Z_{12} , Z_{21} , que miden el efecto de proximidad entre pistas) son cero. De este modo, Z_{od} no baja de los 50 ohmios nominales y la impedancia diferencial, $2 \cdot Z_{od}$ es 100 ohmios.

Al reducir la separación entre pistas a 300 micras (Figura 4.13, centro), vemos que hay más líneas de campo eléctrico entre las pistas, evidenciando un mayor acoplamiento. La corriente eléctrica que va de cada pista a los planos de referencia encuentra así un camino adicional: hacia la otra pista del par y de ahí a los planos. Así se entiende que añadiendo una impedancia en paralelo baja la impedancia global. El efecto queda cuantificado mediante los elementos Z_{12} , Z_{22} de la matriz de impedancias, que valen 2,5 ohmios. De este modo, la impedancia de cada pista a los planos de referencia es de 47,5 ohmios ($50 - 2,5$ ohmios) y la impedancia diferencial (que es la que ve la señal en modo diferencial) es del doble de este valor, 95 ohmios.

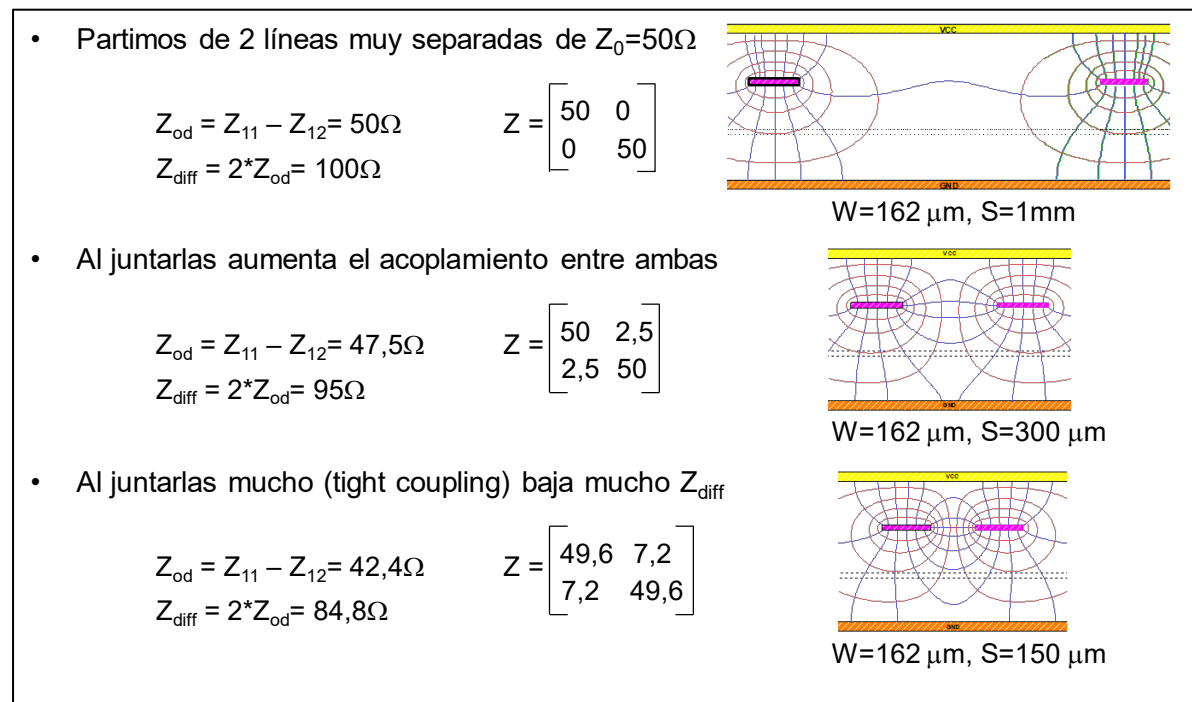


Figura 4.13. Efecto de proximidad de las pistas en modo diferencial (matriz de impedancias y distribuciones de campo). Fuente propia

En la Figura 4.13, parte inferior, aproximamos todavía más las pistas, hasta 150 micras entre bordes. El acoplamiento eléctrico es muy elevado, reduciendo la impedancia diferencial por debajo de 85 ohmios.

Lo que hemos descrito abre dos preguntas: ¿qué separación queremos para las dos líneas del par? ¿Cómo compensamos el efecto de proximidad y recuperamos 100 ohmios de impedancia diferencial?

Compensando el efecto de proximidad

Partiendo de una geometría concreta y suponiendo que excitamos el par en modo diferencial, para cada valor de separación entre pistas hay una reducción determinada de la impedancia característica Z_{od} de cada pista. Para compensar este efecto basta con reducir adecuadamente la anchura de pista, lo que aumenta Z_{od} , de modo que $2 \cdot Z_{od} = 100$ ohmios.

De modo que no hay que perder la cabeza hablando de impedancia diferencial. Se trata, simplemente, de mantener la impedancia de cada pista en 50 ohmios, reduciendo su anchura para compensar el efecto de proximidad. En la Figura 4.14, centro, reducimos la anchura de pista de 162 a 143 micras para recuperar los 100 ohmios. Y en la parte inferior de la figura nos vemos obligados a reducirla hasta 104 micras.

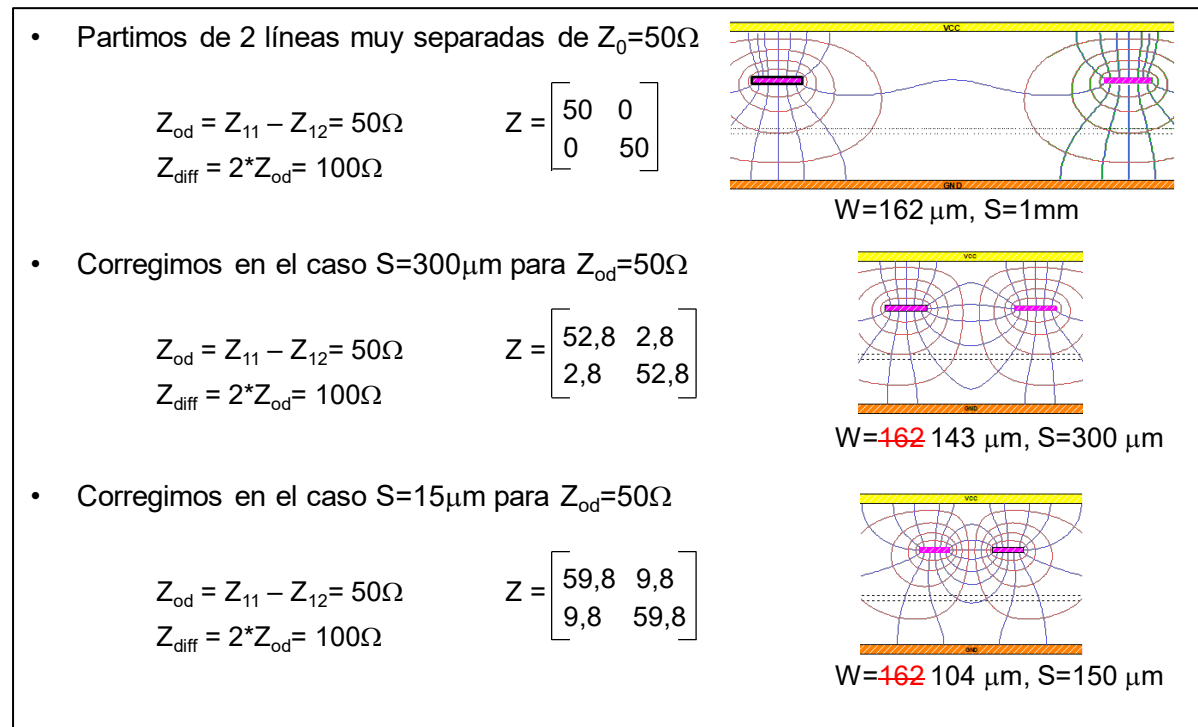


Figura 4.14. Compensamos el efecto de proximidad sobre la impedancia diferencial ajustando la anchura de las pistas. Fuente propia

¿Qué separación entre pistas es adecuada?

¿Es bueno o no rutar las pistas del par muy juntas? Si no están demasiado juntas (*loosely coupled*) podemos separarlas un poco más cuando sea necesario (por ejemplo, para entrar a una BGA, mira la Figura 4.15) sin alterar significativamente Z_{od} . Además, así podemos mantener una anchura de pista mayor y por tanto habrá menores pérdidas.

Si las pistas están muy juntas (*tightly coupled*) hay más desventajas que ventajas: mayores pérdidas resistivas al tener que reducir más la anchura de las pistas y mayor desadaptación al separar las pistas para evitar obstáculos o equalizar longitudes. La única ventaja es que el par ocupa algo menos de área en el PCB.

Si puedes permitirte, te recomiendo usar pistas débilmente acopladas. Para saber dónde se sitúa la diferencia entre débil y fuertemente acopladas, te dejo una regla sencilla: con separaciones borde a borde por encima de $2 \cdot W$, estamos en zona de acoplamiento débil.

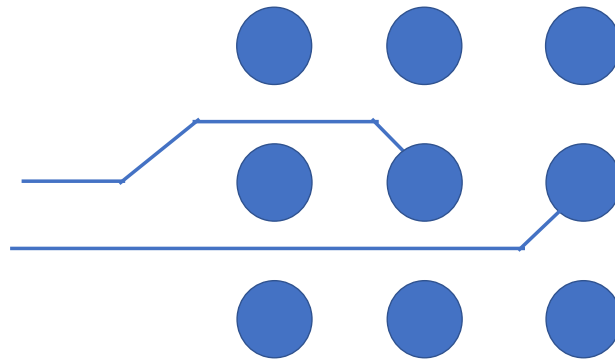


Figura 4.15. Entrada de un par diferencial a una BGA. Fuente propia

Los caminos de retorno en líneas diferenciales

Hemos comentado ya en varias ocasiones que la pista es la mitad de la señal; la corriente de retorno es la otra mitad. De modo que dejemos por un rato de pensar en las pistas y pensemos en lo que ocurre en el o los planos de referencia de la señal.

Debes comenzar por repasar lo que estudiamos el segundo día en la página 39 y siguiente. Y ahora vamos a extenderlo a un par diferencial. Mira la Figura 4.16. Cada pista de un par diferencial tiene corriente de retorno en el plano (o los planos) de referencia, con una distribución de densidad de corriente que sigue una ley cuadrática inversa con x/h , siendo h la altura de la pista sobre cada plano de referencia y x la distancia en el eje de abscisas desde el centro de la pista.

Sólo si las pistas están lo suficientemente juntas podrá haber solape entre las dos distribuciones y una cancelación parcial (ya que la corriente en una línea tiene polaridad opuesta respecto a la otra, recuerda: modo diferencial). Pero por lo general no esperes una cancelación de más del 10-15%.

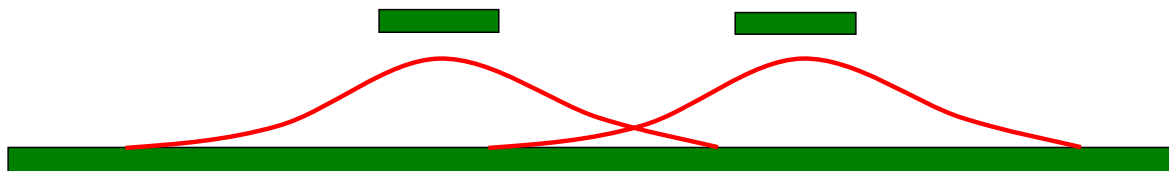


Figura 4.16. Densidad de corriente de retorno en un par diferencial *microstrip*. Fuente propia

En la *stripline* de la Figura 4.17, las densidades de corriente siguen una distribución que sigue la misma ley. La densidad de corriente es mayor en el plano de referencia superior, ya que el cociente $(x/h_1)^2$ será mayor que $(x/h_2)^2$. Si $h_2 > 3 \cdot h_1$, podemos suponer que casi toda la corriente de retorno fluye por el plano de referencia superior.

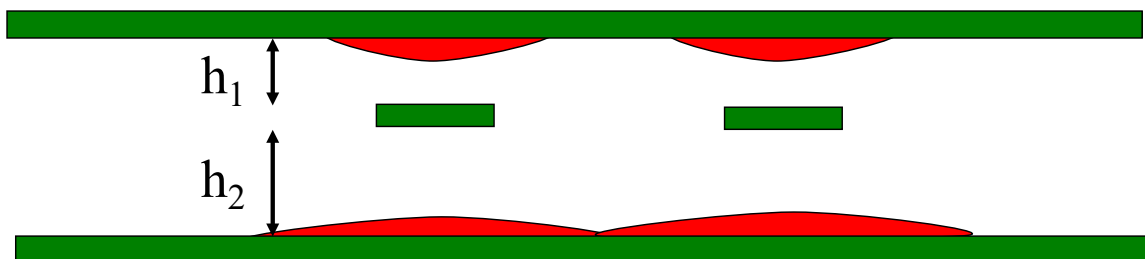


Figura 4.17. Densidad de corriente de retorno en un par diferencial *stripline* asimétrico. Fuente propia

Bien, ¿de qué nos sirve saber todo esto?

Para no degradar el camino de retorno de la señal, asegúrate de que, en los planos de referencia, a ambos lados de cada pista del par y a una distancia de al menos $4h$ (siendo h la distancia pista-plano), no hay ninguna discontinuidad.

En la Figura 4.18 izquierda, el par diferencial pasa sobre cuatro antipads en un plano de referencia, forzando a la corriente de retorno a dar un rodeo. Este aumento de distancia recorrida provocará una discontinuidad en la impedancia de línea, un mayor retardo de propagación y distorsionará la señal. Pero al menos el efecto es igual en ambas líneas de par: se mantiene la simetría.

El caso de la Figura 4.18, derecha es mucho peor: el efecto sobre ambas líneas del par no es simétrico. El resultado es que parte del modo diferencial se convertirá a modo común, creando otros problemas.

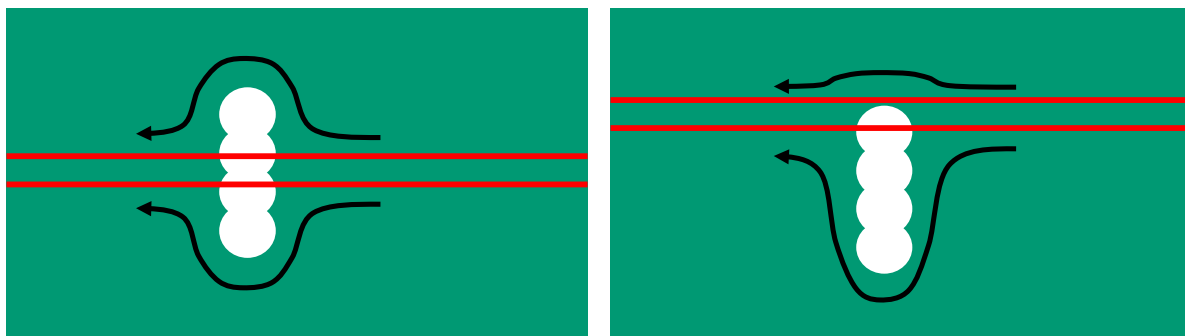


Figura 4.18. Efecto de una discontinuidad en la corriente de retorno. Fuente propia

Conversión de modo diferencial a modo común

Llamemos V_1 y V_2 a la señal en cada línea del par. La señal en modo diferencial queda definida como: $V_{md} = (V_1 - V_2)/2$. La señal en modo común se define como $V_{mc} = (V_1 + V_2)/2$.

Podemos reescribir lo anterior como: $V_1 = V_{mc} + V_{md}$, $V_2 = V_{mc} - V_{md}$.

Si $V_1 = -V_2$, está claro que $V_{mc} = 0$. Sólo aparece modo común si, por alguna asimetría, V_1 deja de ser igual a $-V_2$. El modo común es peligroso porque:

- Acoplado a un cable, puede hacer el producto no supere los ensayos de emisiones conducidas
- Distorsiona la señal y puede provocar que la transmisión de datos falle
- Aumenta la radiación del PCB
- Una terminación típica de 90 o 100 ohmios entre las dos líneas del par no termina el modo común, sólo el diferencial, provocando reflexiones en la línea

Así que, en la medida de lo posible, seguiremos una regla: si algo afecta al camino de retorno de una línea del par, modifica el rutado para intentar que el efecto sea el mismo en ambas líneas.

Efecto de cambios de capa de la señal sobre el camino de retorno

Cuando una señal cambia de capa, usas una vía para asegurar la continuidad entre pistas que están en diferentes capas. Muy bien, te has ocupado de la mitad de la señal. ¿Cómo aseguras la continuidad en el camino de retorno?

En la Figura 4.19, la señal pasa de capa *top* a *bottom*. En el tramo en capa *top*, el plano de referencia es L2 (GND). En el tramo en *bottom*, el plano de referencia es L5 (también GND). Damos continuidad a la corriente de retorno entre L2 y L5 con sendas vías de masa que unes los dos planos. Y, continuando con el principio de simetría, las vías deben colocarse simétricamente al par.

Cuidado con las vías en mitad del desierto

En la Figura 4.20 reproduzco un detalle diseño en el que me descuidé. En la esquina inferior izquierda, vemos un salto de capa entre *top* y *bottom* de un par diferencial. Pero olvidé añadir las vías que dan continuidad a las corrientes de retorno. El par diferencial en cuestión era una señal de reloj. La distorsión que la discontinuidad provocaba en la señal impedía al sistema trabajar a 600 Mb/s y tuve que bajar las prestaciones a 480 Mb/s. No puedo quejarme, todavía tuve suerte: aunque mermado, el sistema era funcional.

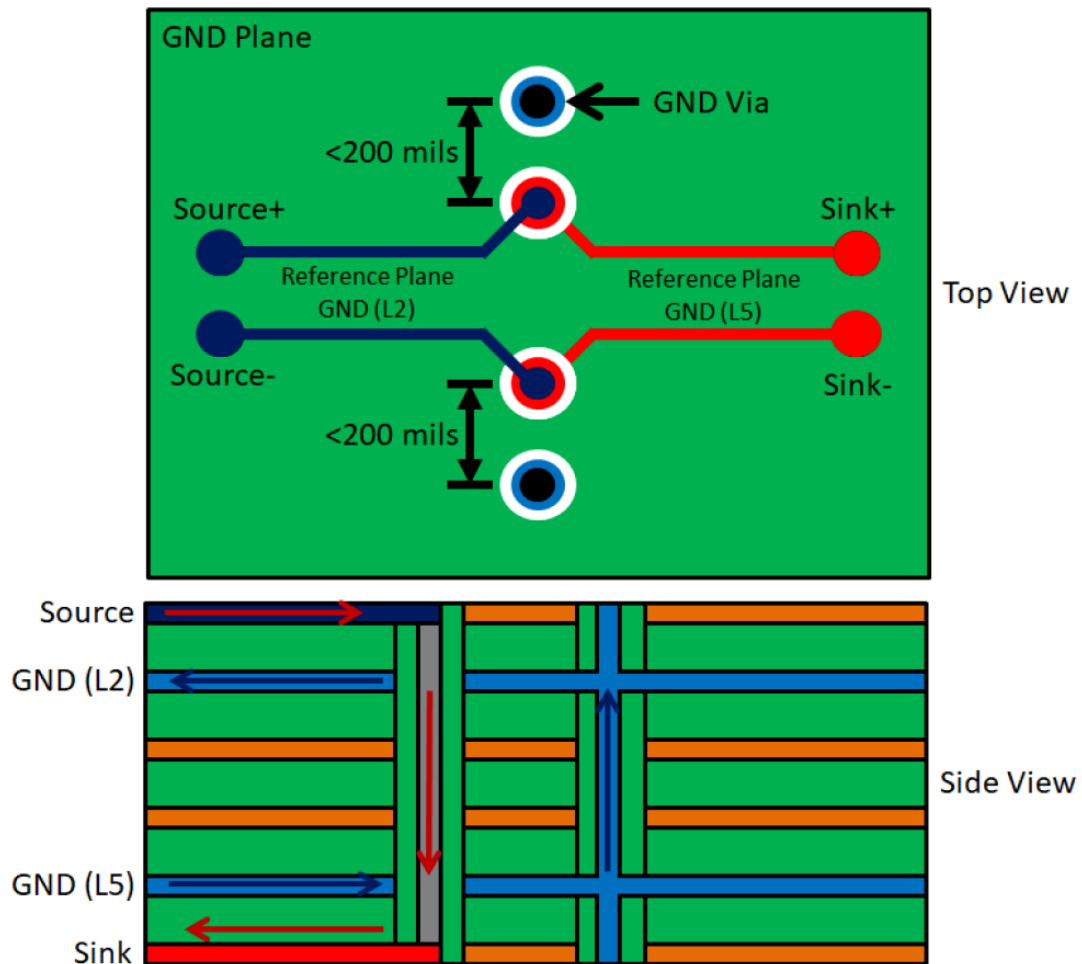


Figura 4.19. Forma correcta de asegurar la continuidad de la corriente de retorno [4]

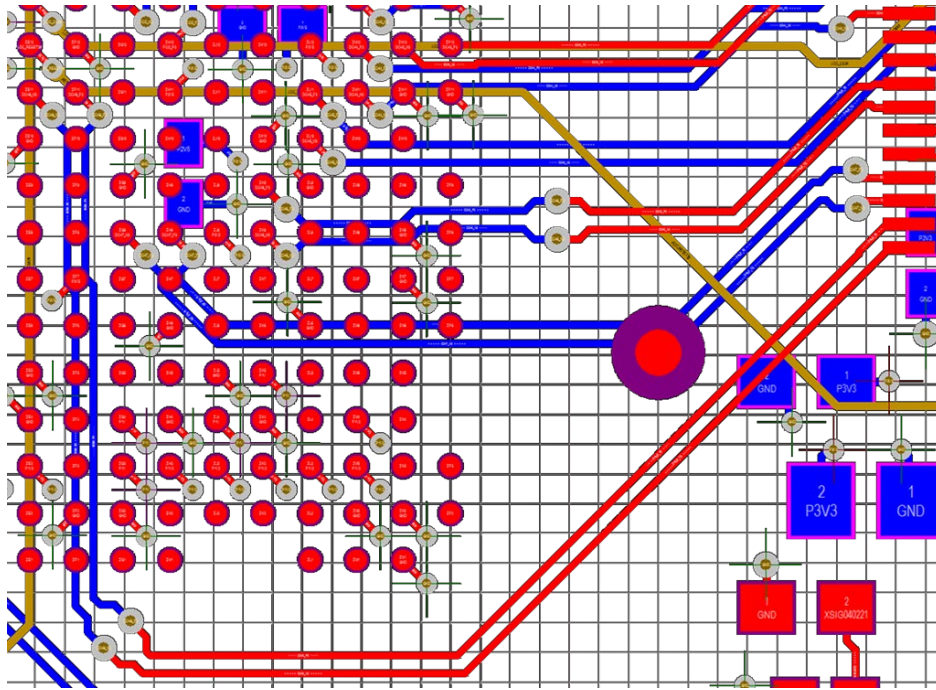


Figura 4.20. Un diseño en el que olvidé cuidar los caminos de las corrientes de retorno en un par diferencial importante. Fuente propia

Unas últimas consideraciones prácticas

Antes de dejar de hablar de diseño de (y con) pares diferenciales, vamos a completar el conjunto de reglas y principios básicos que necesitas llevar en tu mochila.

¿Debo dejar una separación mínima entre pares diferenciales?

Sí. Debes evitar que una línea de un par diferencial se acople con la línea más cercana de otro par diferencial. Este acoplamiento de energía, también llamando *crosstalk*, y del que nos ocuparemos dentro de un par de días, puede entenderse desde varios puntos de vista.

Tal vez el más sencillo e intuitivo podamos derivarlo de la Figura 4.16: el solape de corrientes de retorno de dos pistas cercanas equivale al paso de energía de una pista a otra. Recuerda que el perfil de la densidad de corriente seguía una ley cuadrática inversa con $(x/h)^2$, siendo x la distancia que me alejo del centro de la pista y h la altura de la pista sobre el plano de referencia. Pequeñas distancias entre pista y plano hacen que la corriente de retorno por el plano de referencia esté más concentrada, se extienda menos hacia los lados. Y como consecuencia, el *crosstalk* entre pistas cercanas será menor.

Una segunda forma de entender el *crosstalk* nos remite a la Figura 4.13 y Figura 4.14: el acoplamiento eléctrico y magnético entre pistas cercanas provocará el paso de energía entre ellas. Este efecto disminuye al aumentar la separación entre pistas.

Vale, sin duda todo esto es interesante, pero ¿hay alguna regla práctica que puedas seguir? Ahí va una: procura no rutar nada a una distancia $5 \cdot W$ de cada pista del par [4]. Otros autores reducen la recomendación de separación mínima a $3 \cdot W$ [6]. Y si se trata de relojes o señales periódicas, debes aumentar la separación, tal vez a $8 \cdot W$ si puedes permitirte. Ya sabes que usamos W para referirnos a la anchura (*width*) de una pista.

¿Podemos rutar un par diferencia con giros a 90°, o debemos usar giros a 45° o curvos?

Puedes rutar tranquilamente a 90°, excepto si vas a trabajar en microondas o con señales con contenido significativo por encima de 10 o 20 GHz [7]. Así que, en un diseño analógico o digital habitual, no debes preocuparte por usar giros a 90°.

Hace muchos años, cuando la tecnología de fabricación de PCBs estaba menos avanzada, los ángulos cerrados daban lugar a problemas de fabricación. Y como solución, se extendió el hoy ubicuo rutado a 45°.

Lo normal es que te encuentres con diseños actuales en los que las pistas son todas a 45°, excepto las de interfaces de radio (WiFi, Bluetooth, etc.), SATA, 10-Gbit Ethernet y otras interconexiones con anchos de banda bien por encima del GHz que se rutan con giros curvos. Si bien podrían rutarse también a 45°, hacerlo con curvas es una forma de significarlas de una forma muy visible y tener un cuidado especial con ellas en el diseño.

Personalmente, prefiero rutar a 45° antes que a 90°.

¿Degradan la señal las vías en el camino del par diferencial?

Poner dos o cuatro vías en el camino de cada línea del par (siempre que respeten la simetría, recuerda) no degradan apreciablemente la señal [8]. El problema está en que los cambios de capa llevan aparejados cambios de planos de referencia y si no proporcionamos continuidad a los caminos de las corrientes de retorno, degradaremos la señal. Siempre que mimemos los caminos de retorno, unas pocas vías no deberían ser un problema importante.

Interferencia entre símbolos

En una señal no diferencial (también llamada *single-ended*), con un tiempo de bit de 10 ns (lo que equivale a una tasa binaria de 100 Mb/s), el flanco puede durar 1 ns y el transitorio tras el flanco tal vez 2-3 ns. Es decir, cuando llega el próximo flanco el transitorio ha terminado y tenemos un 1 o un 0 estables.

En cambio, una línea a 500 Mb/s (que obviamente será un par diferencial) tiene un tiempo de bit de 2 ns, un flanco de alrededor de un 20% del tiempo de bit (400 ps) y un transitorio que no ha terminado cuando llega la próxima transición.

En el primer caso, los bits son independientes: la historia de la línea anterior al bit (símbolo) actual es irrelevante. En el segundo caso los bits no son independientes: el bit actual se ve influenciado por en qué nivel de tensión terminó el anterior (Figura 4.21). Que a su vez se vio influenciado por su predecesor. Es decir, los símbolos (bits) interfieren a los que vienen después. Podemos hablar de que hay [interferencia entre símbolos](#). Varios ceros seguidos dan lugar a un cero actual de amplitud elevada. Un cero que sigue a un uno será de menor amplitud.

Para entender este fenómeno debes pensar que la línea es un condensador que se carga durante el tiempo que dura un valor lógico y se descarga durante el valor contrario. Una transición completa (de 0V a 2,5V, por ejemplo) lleva un tiempo. Si este tiempo es mayor que el tiempo de bit, hemos cruzado el límite y nos encontramos con un problema. Porque la señal *single-ended* de 100 Mb/s se puede estudiar con un osciloscopio, pues todos los 1 lógicos tienen el mismo aspecto, igual que también lo tienen todos los 0 lógicos. Pero la señal de 500 Mb/s no se puede estudiar con un osciloscopio: la imagen no será estable.

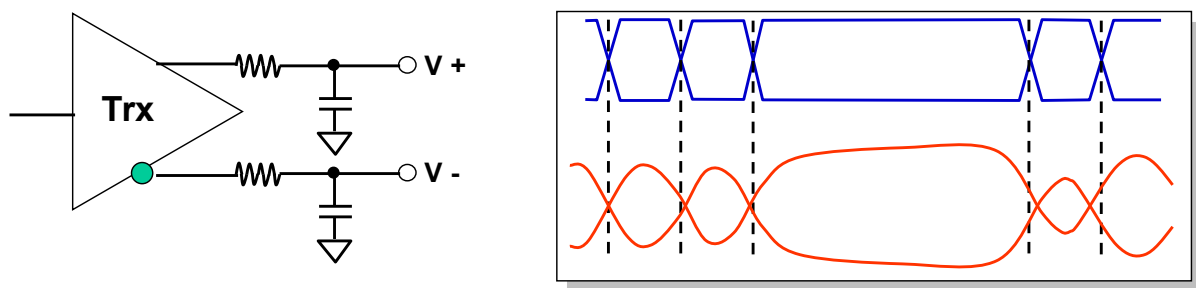


Figura 4.21. Interferencia entre símbolos. La amplitud que alcanza cada símbolo depende de los símbolos anteriores

La interferencia entre símbolos produce jitter en la señal

En la discusión anterior hemos puesto el énfasis en la amplitud de señal. Pero también se ve afectado el flanco (mejor dicho, el instante en el que se alcanza el umbral de detección de un cero, un uno o un flanco activo). Sabes que la variación aleatoria en la posición de los flancos se llama *jitter* y que tiene múltiples causas, entre las que cabe citar la precisión del oscilador de reloj, el ruido acoplado en la señal y la interferencia entre símbolos. A esta última componente, que es la que nos ocupa en esta última sección de hoy, solemos llamarle *data-dependent jitter*.

Para estudiar si la señal, ante cualquier patrón de datos, tendrá suficiente amplitud para permitir la detección de unos y ceros lógicos, y para estimar (o medir) el *jitter* producido por la interferencia entre símbolos, se ha ideado una nueva forma de representar las señales digitales que viajan en los pares diferenciales: el [diagrama de ojos](#).

Diagramas de ojos

En realidad, para construir un diagrama de ojos necesitamos un osciloscopio. Capturamos miles de bits (símbolos) con sus transiciones 0-1 y 1-0, cada símbolo con su forma de onda que está influenciada por los símbolos anteriores. El diagrama de ojos se forma superponiendo secuencias de señal de la misma duración, alineadas con los flancos de reloj. El resultado es un diagrama denominado ojo.

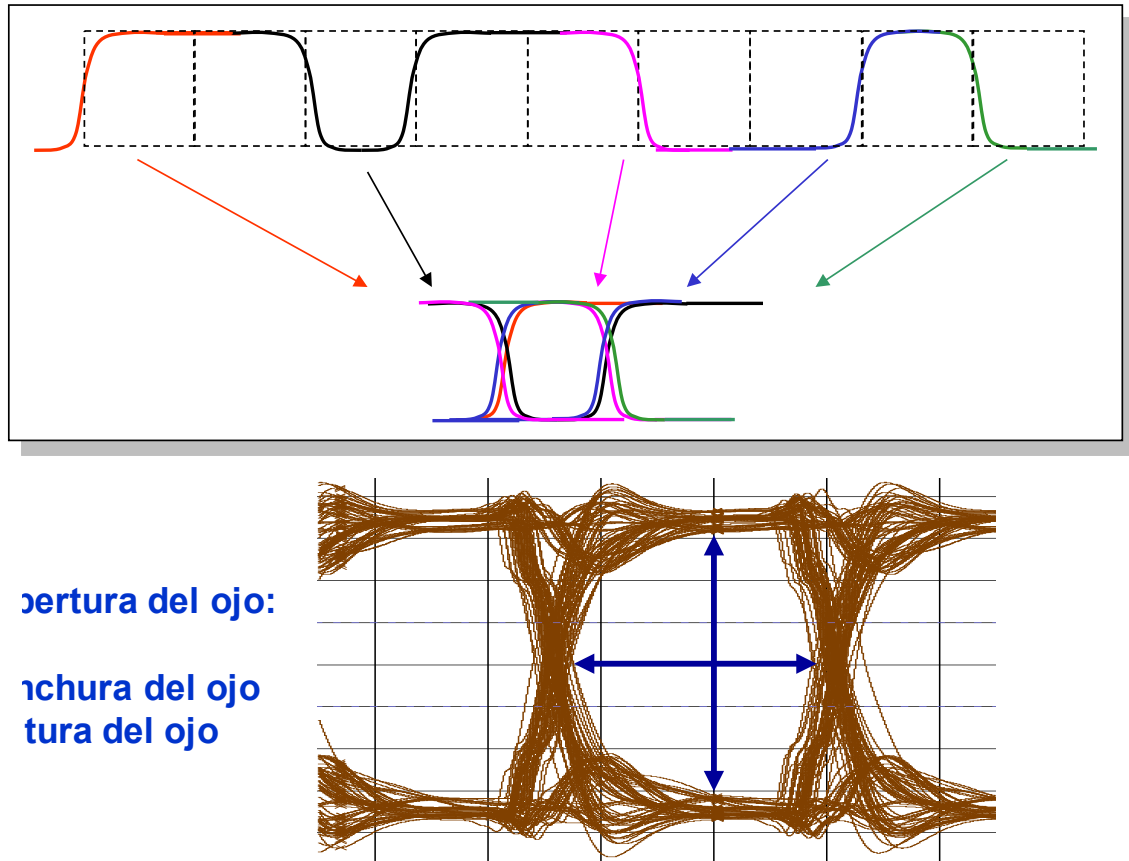


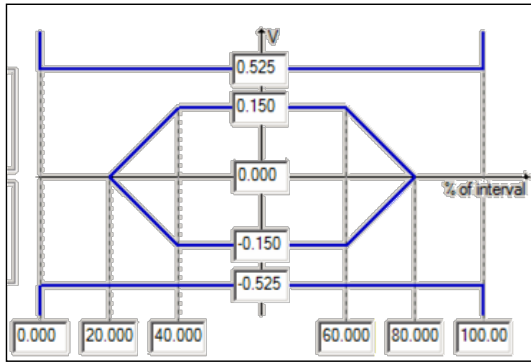
Figura 4.22. Formación de un diagrama de ojos. Fuente: una antigua nota de aplicación de Tektronix que no he podido localizar (pendiente de revisión)

De este ojo queremos medir su apertura, mediante varios parámetros que vamos a resumir en dos: su anchura y su altura (Figura 4.22, abajo). Un ojo cerrado nos habla de que a veces (o con frecuencia) se alcanzarán amplitudes pequeñas y/o el *jitter* será muy grande. Y eso equivale a errores en la transmisión de datos.

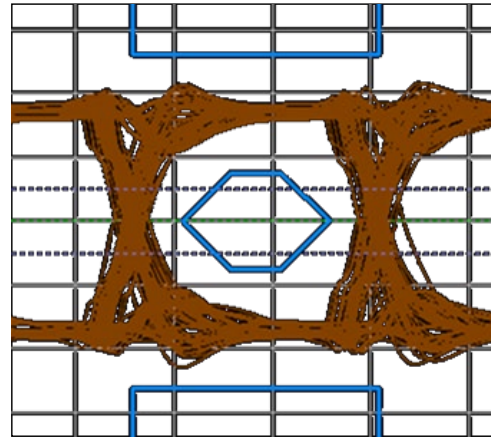
Máscaras (eye patterns)

Con el fin de comparar el ojo con un estándar, se definen, para cada tecnología (USB2.0, PCI Express, HDMI, etc.), tanto en transmisión como en recepción, áreas prohibidas que el ojo no debe cruzar para garantizar una adecuada calidad de la señal. Estas áreas son diferentes para cada estándar de transmisión. En la Figura 4.23 se muestra la máscara para USB 2.0 en recepción (izquierda) y la máscara superpuesta a la simulación de un diagrama de ojos (derecha).

Como el ojo no toca la máscara, decimos que la integridad de señal es lo suficientemente buena para cumplir con el estándar. En realidad, hay que demostrarlo con medidas, no con simulaciones, pero un estudio previo nos muestra si vamos por buen camino.



Máscara para USB 2.0 high-speed en recepción



El ojo no cruza la máscara: correcto!!
El enlace USB funciona...

Figura 4.23. Ejemplo de máscara. Fuente propia

Día 5. Radiación en un PCB

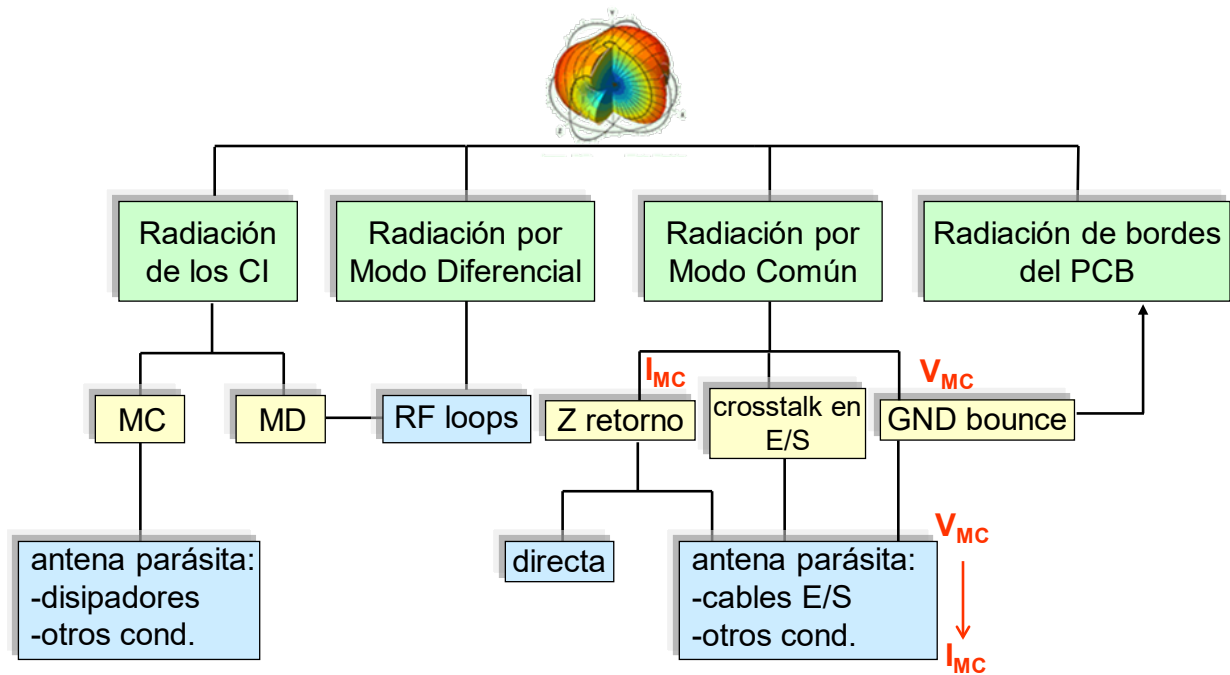


Figura 5.1. Diagrama resumen de las causas de radiación en un PCB. Fuente propia

Resumen antes de la lección

La Figura 5.1 no te dice nada todavía. De hecho, te resultará muy confusa ahora mismo, pero toda la lección de hoy gira a su alrededor. Si hago bien mi trabajo, acabarás comprendiéndola y te servirá de resumen.

En una frase, [la figura recoge las causas involuntarias por las que un PCB radia y que ponen en peligro que tu producto supere los ensayos de EMC de emisiones radiadas](#).

Estas causas se dividen en varios tipos:

- (1) [Radiación de las pistas del PCB y de las interconexiones dentro de los circuitos integrados](#). Esto es inevitable, responde a las leyes de la física. Ya sabes que la pista es la mitad de la historia: la otra mitad es la corriente de retorno. Ambos forman un bucle, una espira, que radia irremediablemente. Sólo podemos conocer y aplicar las buenas prácticas que reducen esta radiación.
- (2) Radiación por estructuras que, una vez excitadas, se convierten en [antenas tipo parche parásitas](#) que radian. Ejemplos son áreas de cobre sin conexión y disipadores sobre circuitos integrados. Podemos mitigar esta causa con conexiones adecuadas a masa que evitan la antena. También construyendo disipadores con materiales no conductores, y es una razón por la que a veces se usan disipadores cerámicos.
- (3) Radiación por estructuras que se comportan como [antenas monopolo parásitas](#). Ejemplos son tarjetas pinchadas sobre *backplanes* y placas madre.
- (4) Propagación de cualquier señal acoplada a una estructura biplaca (o cavidad, dos planos de masa o de alimentación contiguos), que acaban produciendo [radiación por los bordes del PCB](#). A este tipo de estructuras se acoplan tanto señales como ruido en masas y alimentaciones. Podemos emplear técnicas conceptualmente sencillas.
- (5) [Radiación por cables](#) a los que se acopla ruido en modo común, generalmente debido a ruido en los planos de masa.

Nuestro objetivo, durante el tiempo que te llevará leer las siguientes páginas, será entender causas y contramedidas. No reducirás la radiación de tu producto electrónico resolviendo ecuaciones, sino aplicando buenas prácticas. Por este motivo emplearemos muchas figuras y pocas ecuaciones.

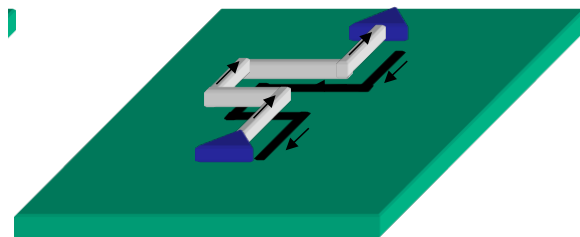
A la primera de las causas que hemos comentado le damos el nombre de **radiación por modo diferencial**. El resto quedan englobadas bajo el paraguas de la **radiación por modo común**. Tranquilo, iremos estudiando todo esto paso a paso.

Radiación en modo diferencial

Hemos comentado que es producida por la radiación de las pistas del PCB y de las interconexiones dentro de los circuitos integrados, y es inevitable. Pista y plano de referencia (por donde retornan las corrientes) discurren a corta distancia uno del otro y forman un bucle que radia de forma proporcional al área, corriente eléctrica y al cuadrado de la frecuencia.

Como el sentido de la corriente es opuesto en la pista y en el plano de referencia, buena parte de la radiación se cancela: **por este motivo, la radiación en modo diferencial es poco eficiente**. El término “diferencial” da lugar a confusión, sobre todo después de haber estudiado ayer líneas diferenciales. Debes entenderla como que es “diferente” en pista y plano de referencia (sentido opuesto).

En los libros de EMC -por ejemplo, en [9], sección 8.1.2- encontrarás la siguiente expresión, o una similar, válida para bucles pequeños, para la intensidad del campo eléctrico en la dirección de máxima radiación:



Alta frecuencia

La corriente de retorno circula bajo la pista de señal (mínima inductancia)

$$E = 263 \cdot 10^{-11} \cdot f^2 \cdot A \cdot I / r$$

Donde el campo eléctrico está expresado en V/m, el área (A) en cm², la frecuencia en MHz, la corriente (I) en mA, y la distancia (r) en metros.

El área será igual a la longitud de la pista multiplicada por su altura sobre el plano de referencia. En [10] se describe una calculadora para la expresión anterior, válida para pistas no mayores de un sexto de la longitud de onda a la frecuencia de interés.

Si queremos hacerlo en Matlab, el código siguiente genera una gráfica de la radiación en dBmV/m a 10 m de distancia, para una corriente de 2 mA, una pista de 6,35 cm de longitud a 140 micras del plano de referencia (Figura 5.2):

```
f=linspace(30,500,50);
i=2; %% Corriente en mA
long=6.35; %% Longitud de la pista en cm
t=0.014; %% Altura sobre el plano de referencia en cm
dist=10; %% Distancia a la que situamos la antena de medida virtual
rad=(263.2e-11*i*long*t/dist)*f.^2; %% Campo en V/m, en lineal
raddb=20*log10(rad*1e6); %% Convertimos V/m en uV/m y calculamos dBuV/m
plot(f,raddb,"LineWidth",3);
xlabel('Frecuencia en MHz');
ylabel('E (dBuV/m) a 10m');
grid on
```

En el código anterior faltaría imponer la condición de que, si **long** es mayor que la longitud de onda, se usa la longitud de onda como longitud, produciendo así una saturación.

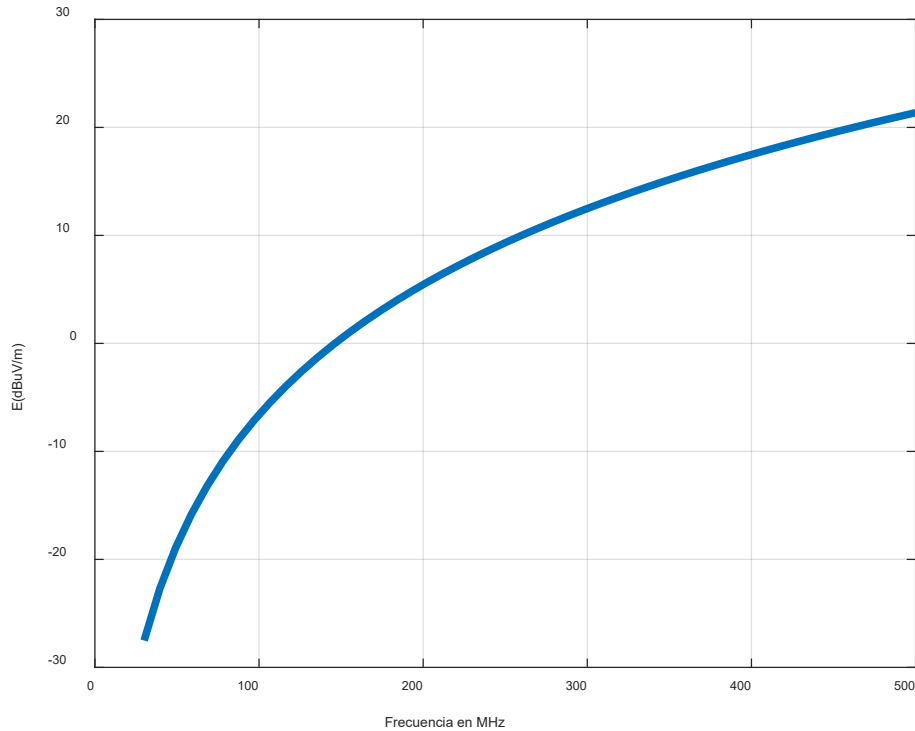
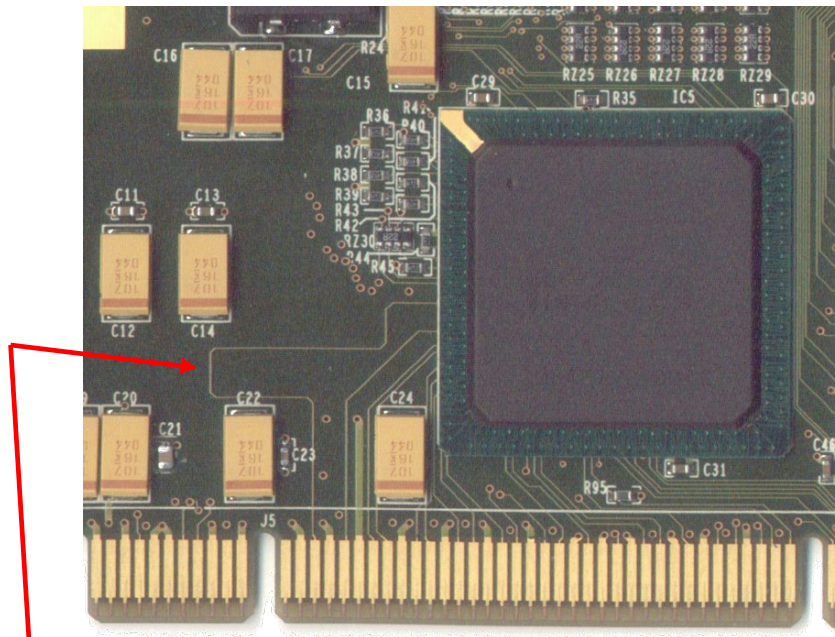


Figura 5.2. Gráfica obtenida con Matlab de la radiación de una pista de 6,35 cm, 140 micras sobre el plano de referencia, conduciendo una corriente de 2 mA. Fuente propia

Un ejemplo: radiación de una pista de reloj

En la Figura 5.3 podemos observar un detalle del rutado de una tarjeta PCI. Este estándar requería que la pista de reloj en una tarjeta, pinchada sobre el backplane, tuviera una longitud de dos pulgadas y media, con un margen de una décima de pulgada. Multiplicando esta longitud por el espesor del dieléctrico hasta el plano de referencia, obtenemos el área del bucle.

La Figura 5.4 es una simulación de esta pista de reloj a 33 MHz con la herramienta HyperLynx de Mentor Graphics. Podemos ver la forma de onda de tensión en fuente y carga, la corriente eléctrica en la línea, el espectro en frecuencia de dicha corriente y una estimación de la radiación de cada componente del espectro a 10 m (abajo a la derecha). La Figura 5.5 replica el estudio para un reloj a 66 MHz.



La longitud de la pista de reloj debe ser, conforme a la especificación PCI, $2.5'' \pm 0.1''$

Figura 5.3. Una pista de reloj sobre un plano de masa en una tarjeta PCI, forma un bucle que radia quieras o no. Fuente propia

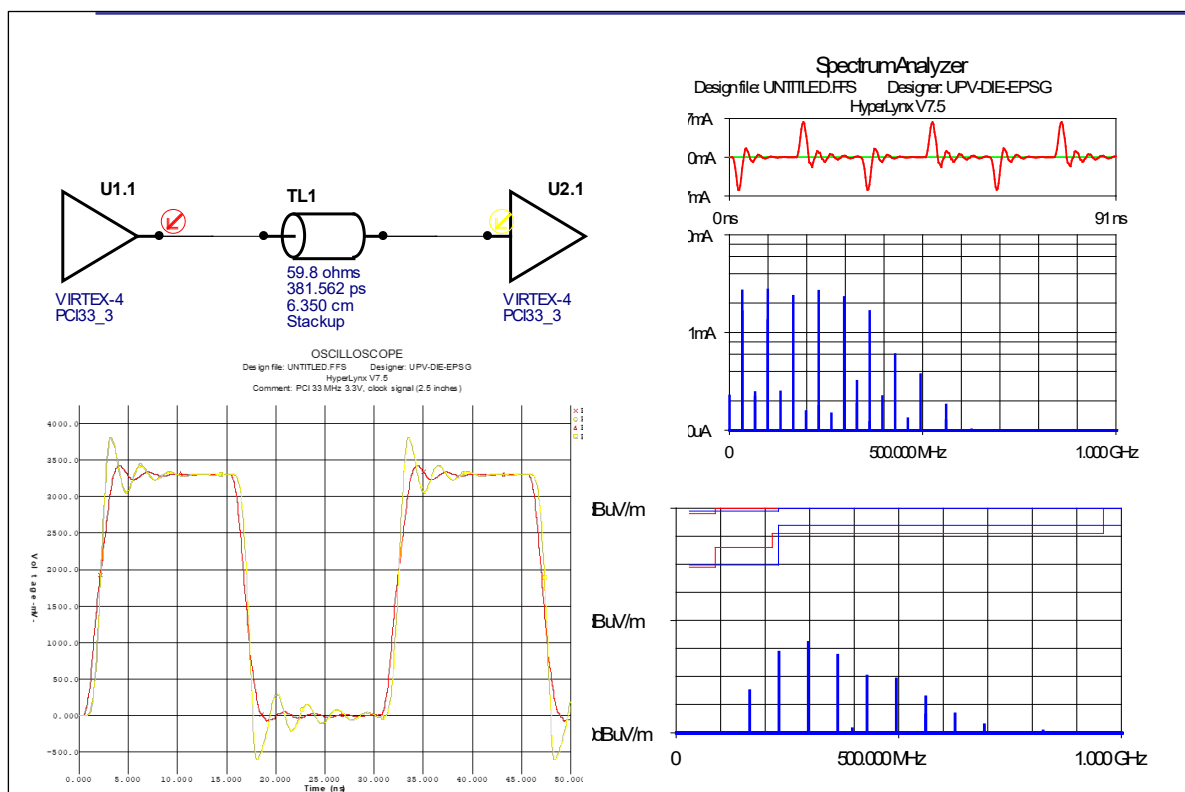


Figura 5.4. Simulación de una línea de reloj PCI a 33 MHz. Fuente propia

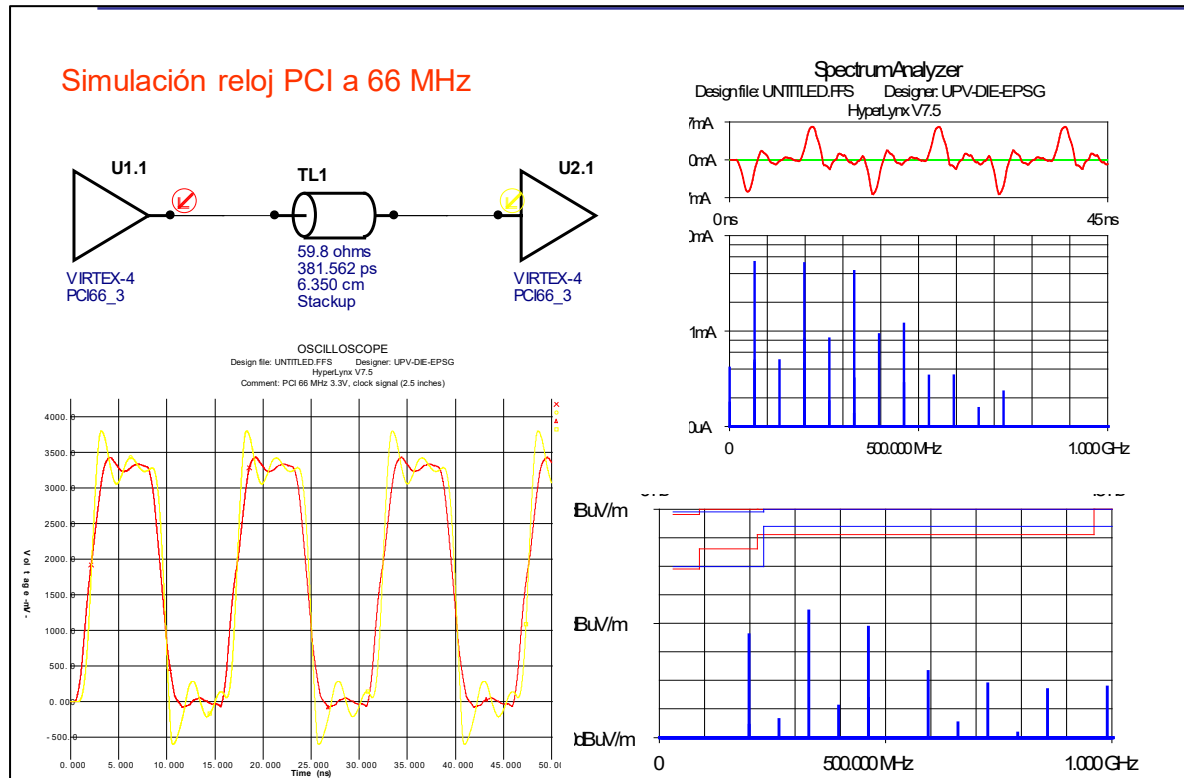


Figura 5.5. Simulación de una línea de reloj PCI a 66 MHz. Fuente propia

En el caso de la línea de reloj PCI a 66 MHz, siendo la pista de 6.35 cm (2,5") sobre plano de masa ($h=0,14$ mm), si consideramos la componente de corriente a 460 MHz de 2 mA, resulta según la expresión que hemos estudiado, resulta que produce 19,9 dB μ V/m a 10 m de distancia. Esto concuerda con el resultado de la simulación con HyperLynx (19 dBmV/m). ¿Quiere esto decir que HyperLynx usa también esta expresión para estimar el campo radiado? En cualquier caso, coincidiendo un cálculo a mano con una herramienta profesional, podemos concluir que no vamos muy desencaminados.

Bien, ¿y es mucho o poco 19 o 19,9 dB μ V/m? Hay que comparar este valor con los límites que permite la normativa EMC. En Europa, muchos equipos de electrónica que caen dentro de la categoría de tecnologías de la información (como puede ser una tarjeta PCI que pinchamos en un PC), que se rige por lo dispuesto en EN55032. La Figura 5.6 muestra los límites de radiación que impone la norma. Si miramos los valores de cuasi-pico para 10 m de distancia a 460 MHz, clase B (uso residencial), observamos que el límite es de 37 dB μ V/m. Para la radiación producida por la pista de reloj, tenemos un margen de 17 dB.

Esto quiere decir que con 7 pistas iguales radiando alcanzaríamos el límite de radiación. Pero un PCB tiene cientos de pistas. ¿Qué está pasando? Vamos a hacer una lista de motivos por los que no tenemos que preocuparnos excesivamente:

- La expresión calcula la radiación en la dirección máxima suponiendo una pista recta. En la Figura 5.3 podemos ver que la pista de reloj tiene tramos rutados en diferentes direcciones. Así que la radiación de la pista será inferior al cálculo, que es un peor caso. Además, el valor de cuasi-pico es algo inferior al de pico.
- Otras pistas en el PCB estarán rutadas en otras direcciones, llevarán señales distintas y en general no podemos suponer una suma de todas ellas en fase para decir que con 7 pistas alcanzamos el límite.
- Un producto electrónico hay que someterlo a ensayo en un *test setup*, que en el caso de una tarjeta PCI supone meterla dentro de un PC que tiene carcasa metálica y que, en mayor o menor medida, atenuará la radiación (ya hablaremos de esto otro día).

Lo anterior nos da un respiro, pero nos enseña una importante lección:

Las señales críticas (en cuanto a radiación) no las llevaremos, en la medida de lo posible, por capas externas. Veremos un poco más adelante cómo se reduce la radiación al llevarlas por capas internas y, en cualquier caso, al añadir terminaciones para reducir las reflexiones.

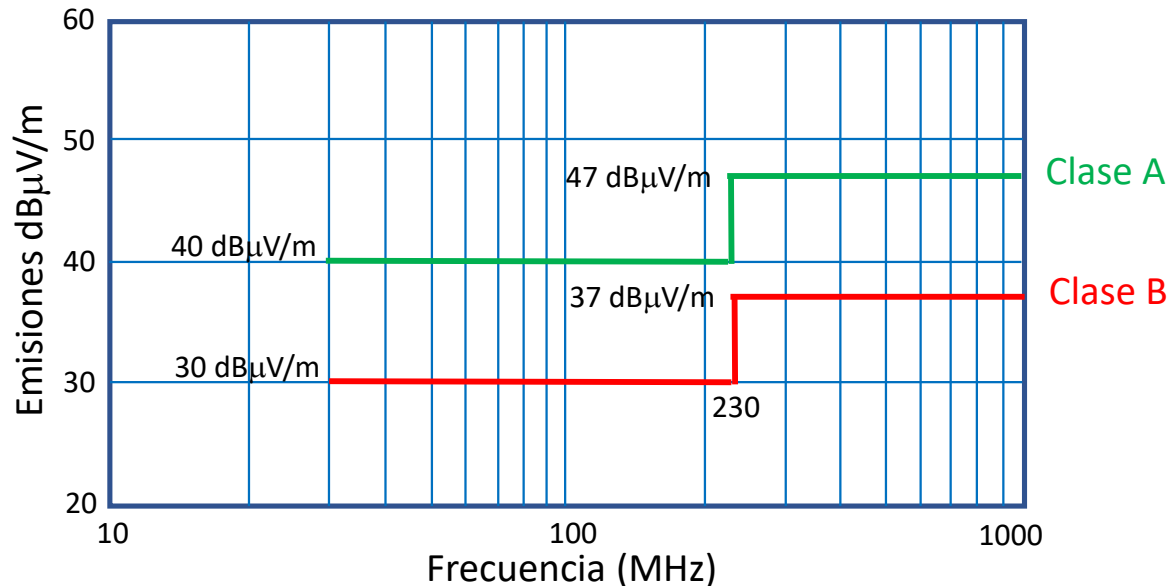


Figura 5.6. Límites de emisiones radiadas según EN55032 (compatible con CISPR32) a 10 m, en cuasi-pico (medido con RBW de 120 kHz). Fuente propia

Otra vuelta de tuerca sobre el ejemplo... Vamos a añadir terminaciones

¿Recordamos lo que hablábamos el tercer día sobre que añadiendo terminaciones reducíamos la presencia de altas frecuencias en la señal y por tanto a radiación? Lo mencionamos al hablar de la terminación serie en la fuente. Bien, vamos a cuantificar esto.

La Figura 5.7 compara la radiación de la pista de reloj PCI a 66 MHz del ejemplo con otra versión que añade una terminación serie en la fuente. Observamos que la radiación de la frecuencia de 460 MHz se ha reducido en aproximadamente 5 dB. Pero la componente de 600 MHz lo ha hecho en 12 dB. No está nada mal por un céntimo (el coste aproximado de una resistencia).

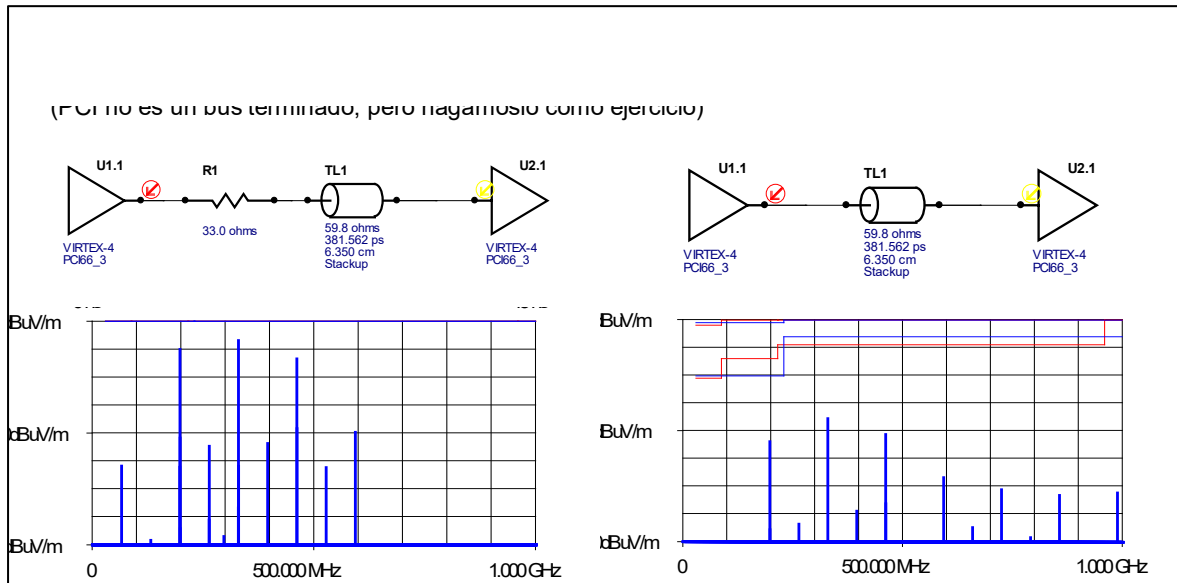


Figura 5.7. Reducción de la radiación al añadir terminaciones en la línea: lo que es bueno para la integridad de señal, es por lo general también bueno para EMC. Fuente propia

Reducción de radiación diferencial

Además de añadir terminaciones, podemos buscar una cancelación parcial de la radiación rutando las pistas críticas por capas internas. Al haber dos planos de referencia, se induce corriente de retorno por ambos planos, creando dos lazos con sentidos de corriente opuestos y que tienden a restarse. En la práctica, no conseguiremos cancelación porque eso implicaría una elevada simetría en ambos bucles, y la realidad (ya hablaremos otro día sobre estructuras de PCB multicapa) es que dicha simetría es harto complicada.

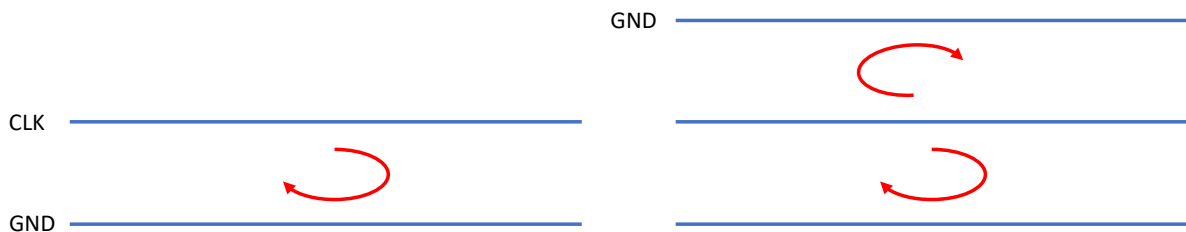


Figura 5.8. Cancelación (parcial) de radiación en capas internas del PCB. Fuente propia

Rutar en capas internas, además de conseguir una cancelación parcial por lo que acabamos de exponer, tiene otra ventaja: las capas de cobre sólido (planos) que hay entre la pista y el exterior van a atenuar (por absorción y por reflexión) el campo, de modo que podemos considerar a efectos prácticos que, exceptuando la radiación por los bordes del PCB (último punto a estudiar hoy), las señales en capas internas apenas radian.

Lo que hemos comentado hasta el momento sobre radiación por modo diferencial nos permite listar unas cuantas buenas prácticas para su reducción:

- Rutar señales críticas en capas internas (apantallamiento)
- Rutar pistas lo más cortas posible (reduce el área del bucle)
- Minimizar la distancia entre la pista externa y su plano de retorno (reduce el área del bucle)
- Usar buffers lo más lentos que permita la aplicación (menor contenido en altas frecuencias: recuerda que la radiación por modo diferencial es proporcional a f^2)

- Terminaciones en las líneas para reducir el contenido de armónicos (menor contenido en altas frecuencias)
- Usar líneas diferenciales (cancelación de flujos)
- Simular el PCB (la radiación por modo diferencial es predecible)

Visión general de la radiación por modo común

La sección anterior (radiación por modo diferencial) era sencilla de entender; ya que se basa en un único fenómeno. La radiación por modo común abarca una mayor variedad de causas y fenómenos y es por tanto más compleja y difícil de asimilar.

Vamos a estructurar el resto de la lección de hoy distinguiendo entre las siguientes, y por este orden:

- Radiación en pistas por conversión de modo diferencial a modo común
- Radiación por antenas parásitas
- Radiación debida al ruido en planos de masa y de alimentación
- Radiación por los bordes del PCB (una consecuencia del punto anterior)

En todos los casos, el apellido de familia “**por modo común**”, nos habla de que una misma señal (interferencia externa, ruido en planos de masa y alimentación o señal generada por nuestro producto que se introduce en otras líneas) está presente (es común) en varias estructuras radiantes, sin que se puedan producir efectos de cancelación como los que están presentes en la radiación por modo diferencial.

En la Figura 5.9 hemos sombreado las cajas relacionadas con la radiación por modo diferencial, destacando así lo que nos queda por abordar.

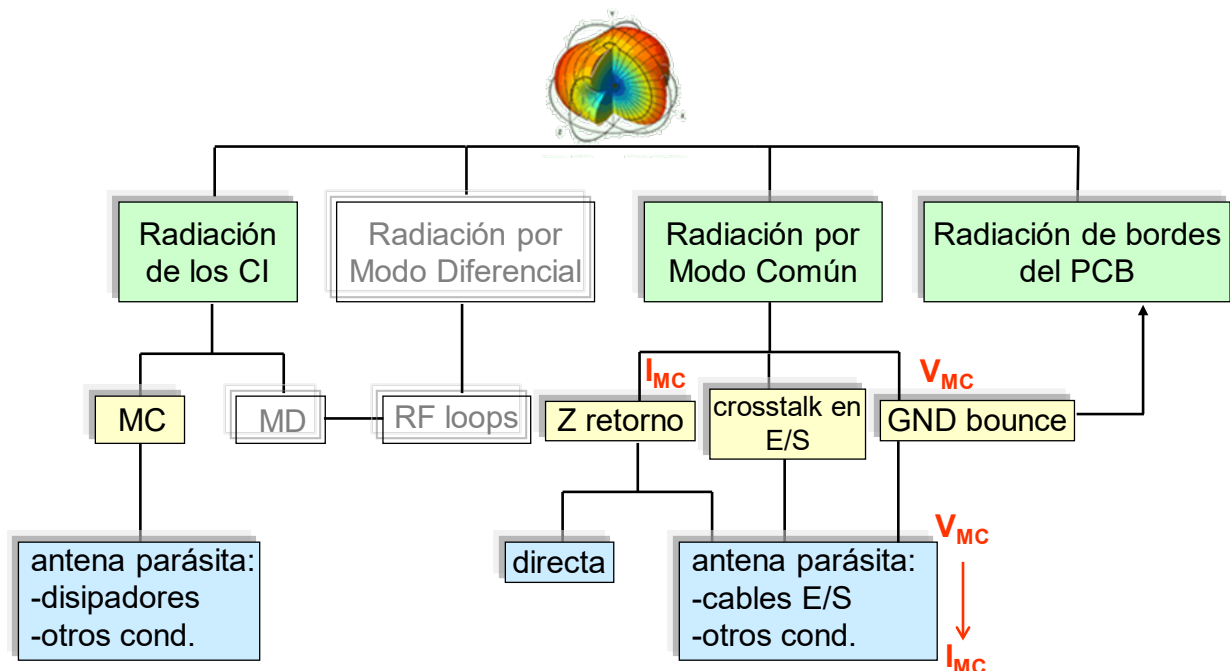


Figura 5.9. Exceptuando la radiación por modo diferencial, el resto de las causas se clasifican en la categoría de radiación por modo común

Radiación por conversión de modo diferencial a modo común

Vamos a reescribir un fragmento de la lección de ayer, donde hablábamos de conversión de modo en un par diferencial, para extenderlo también a señales *single-ended*:

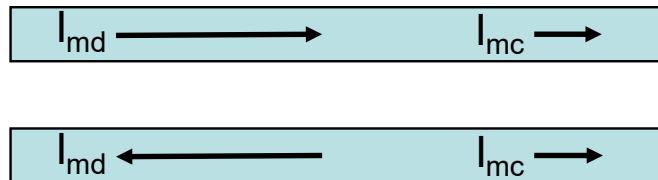
Conversión de modo diferencial a modo común en señales single-ended

Si afectamos de algún modo a la pista o al camino de retorno de corrientes, la corriente en ambos conductores dejará de estar en fase.

Llamemos I_1 y I_2 a la corriente en cada camino (pista y camino de retorno). La señal en modo diferencial queda definida como: $I_{md}=(I_1-I_2)/2$. La señal en modo común se define como $I_{mc}=(I_1+I_2)/2$.

Podemos reescribir lo anterior como: $I_1=I_{mc}+I_{md}$, $I_2=I_{mc}-I_{md}$.

Si $I_1 = -I_2$, como ocurre cuando la línea de transmisión no sufre discontinuidades, está claro que $I_{mc}=0$. Sólo aparece modo común si, por alguna asimetría, I_1 deja de ser igual a $-I_2$. El modo común es peligroso, desde el punto de vista de la EMC, porque se reduce la cancelación de la radiación que producen ambas corrientes.



Los libros sobre EMC -por ejemplo [9]- te darán expresiones similares a la siguiente para la radiación de la componente en modo común. Si quieres una derivación detallada de la expresión, puedes encontrarla también en la referencia [11]:

$$\text{Modo común: } E = 126 \cdot 10^{-9} \cdot f \cdot L \cdot I / r$$

$$\text{Modo diferencial: } E = 263 \cdot 10^{-11} \cdot f^2 \cdot A \cdot I / r$$

Donde E está expresado en V/m, A (área del bucle formado por pista y retorno) en cm^2 , L (longitud de la pista) en cm, f en MHz, I (corriente) en mA, y r (distancia a la antena virtual de medida) en metros.

¿Cuál de los dos modos de radiación domina?

Vamos a echar mano de Matlab una vez más para estudiar el siguiente supuesto. Vamos a completar el *script* de la página 102 para añadir el cálculo de la radiación por modo común, manteniendo los mismos datos del ejemplo, y suponiendo que sólo un 5% de la corriente se convierte a modo común.

```
f=linspace(30,500,50);
i=2; %% Corriente en mA en modo diferencial
long=6.35; %% Longitud de la pista en cm
t=0.014; %% Altura sobre el plano de referencia en cm
dist=10; %% Distancia a la que situamos la antena de medida virtual
coef=0.05; %% Porcentaje del modo común sobre el modo diferencial
```

```

rad_md=(263.2e-11*i*long*t/dist)*f.^2; %% Campo en V/m, en lineal, modo diferencial
rad_md_db=20*log10(rad_md*1e6); %% Convertimos V/m en uV/m y calculamos dBuV/m, modo diferencial
plot(f,rad_md_db,"LineWidth",3); %% Dibujamos la radiación en modo diferencial
hold;
rad_mc=126e-9*i*coef*long*f/dist; %% Campo en V/m, en lineal, modo común
rad_mc_db=20*log10(rad_mc*1e6); %% Convertimos V/m en uV/m y calculamos dBuV/m, modo común
plot(f,rad_mc_db,"LineWidth",3); %% Dibujamos la radiación en modo común
xlabel('Frecuencia en MHz');
ylabel('E (dBuV/m) a 10m');
grid on

```

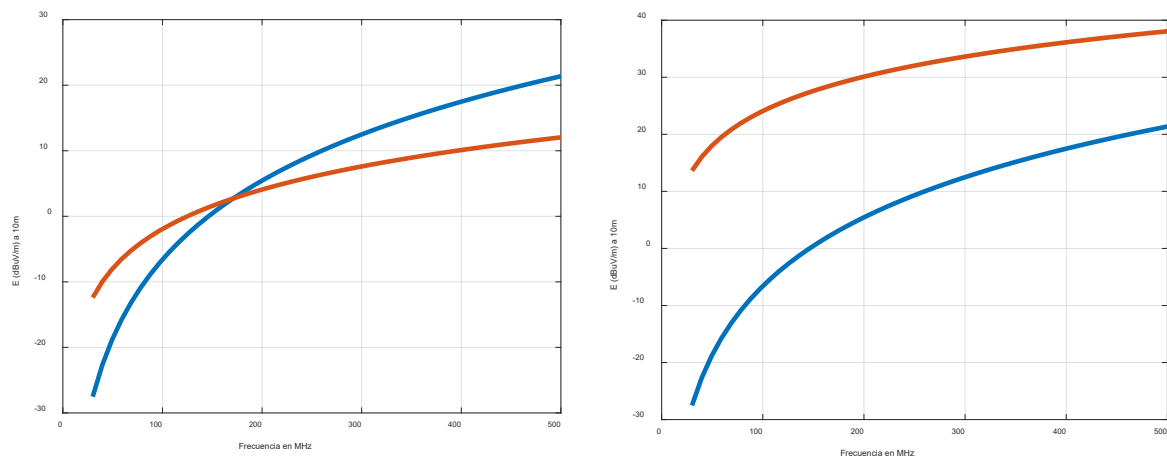


Figura 5.10. Radiación por modo común (curva naranja) y por modo diferencial (azul) en el ejemplo de la pista de reloj PCI. En la gráfica de la izquierda, la corriente en modo común es un 5% de la corriente en modo diferencial. En la figura de la derecha, son iguales. Vemos claramente que la corriente en modo común es más peligrosa. Fuente propia

Si sólo un 5% de la corriente se convierte a modo común, la radiación por modo diferencial sólo domina a partir de 174 MHz en el ejemplo anterior. La Figura 5.10, derecha, nos hace ver que, a igualdad de corriente, la radiación por común es bastante mayor.

Reducir la generación de señales en modo común

La Figura 5.11 ilustra un ejemplo de dos discontinuidades en el camino de retorno de la señal. A la izquierda, un grupo de *antipads* (círculos libres de cobre en un plano de masa o de alimentación, necesarios para evitar cortocircuito entre la vía y el plano) forman una ranura en el plano de referencia que obliga a la corriente de retorno a cambiar su camino, provocando así conversión de modo diferencial en modo común.

A la derecha, la pista cruza por encima de un *gap* entre planos. Si no hay más remedio que cruzar planos aislados, podemos utilizar un condensador (*stitching*) entre ambos planos que aseguren continuidad en alta frecuencia (realmente en un margen de frecuencias centrado en al de resonancia del condensador) y aislamiento en DC. La radiación de la ranura no queda eliminada, pero es una forma de mejorar la situación.

En ambos casos la radiación en modo común afectará no sólo a la integridad de señal, sino a los resultados de los ensayos de EMC de nuestro producto. **La lección es clara: debemos mimar los caminos de retorno de señal, evitando todo aquello que pueda provocar conversión de modo.**

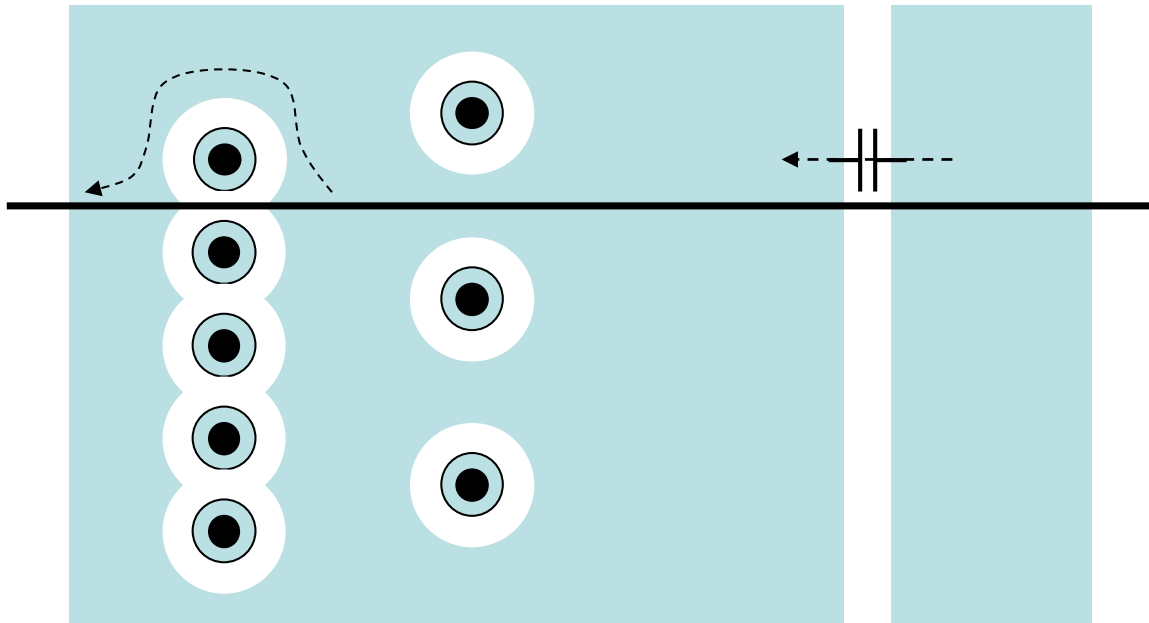


Figura 5.11. Discontinuidades en el camino de las corrientes de retorno que produce radiación por modo común. Fuente propia

Como ya hemos estudiado, para pistas críticas (por su alto contenido en altas frecuencias), añadir vías entre planos de masa junto a los cambios de plano de la señal evitará discontinuidad en las corrientes de retorno [12].

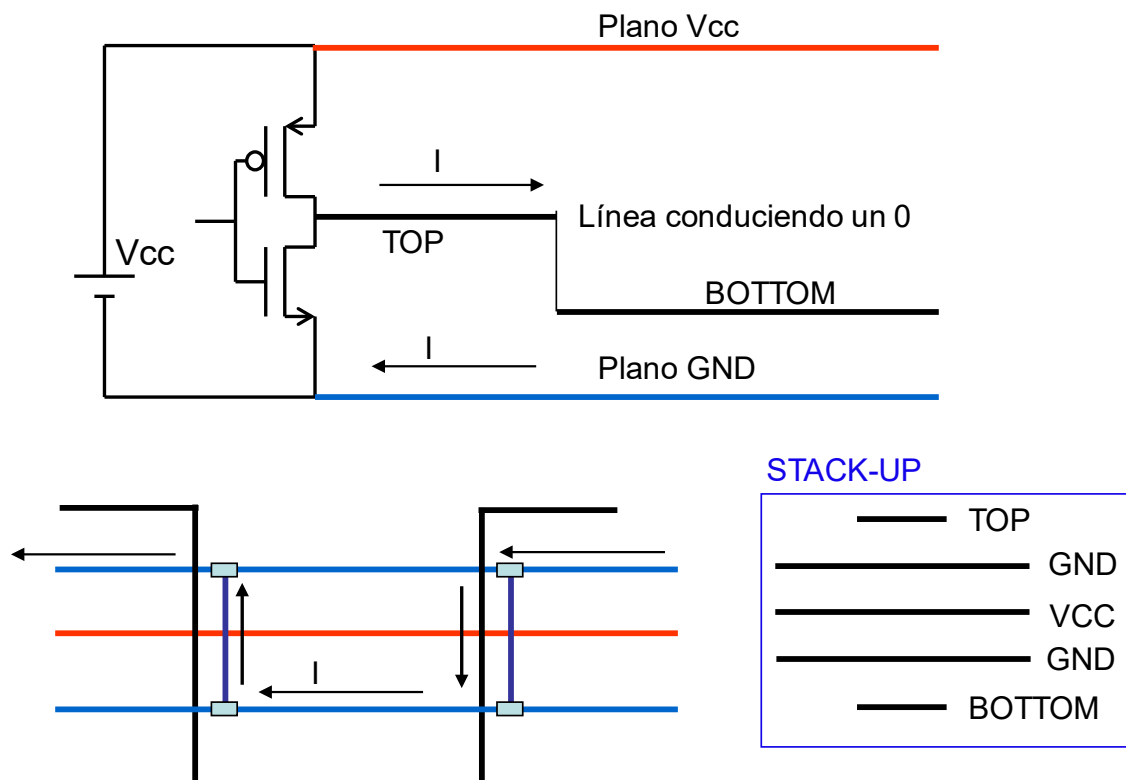


Figura 5.12. Recuerda cuidar el camino de las corrientes de retorno de señales críticas cuando cambien de capa. Fuente propia

Radiación por antenas parásitas

Una antena parásita (una antena no deseada o al menos no intencionada) es una estructura conductora que, si recibe energía en una banda de frecuencias que incluye alguna de aquellas a las que resuena, radiará. Habitualmente, las antenas parásitas son flotantes (su potencial eléctrico no está fijado) o se comportan como flotantes a partir de cierta frecuencia (si el punto o los puntos en los que fijamos su potencial no son pequeños respecto a la longitud de onda).

Dicho así, puede que no lo veas claro. Permite que te ponga algunos ejemplos de antenas parásitas:

- **Un área de cobre en capa externa no conectada a nada (flotante).** Esto sucede a veces cuando rellenamos de cobre las capas externas, bien para reducir la contaminación (evitamos eliminar al ambiente cobre sobrante), para mejorar la conductividad térmica o para homogeneizar la distribución de cobre por la superficie, cosa que el fabricante de PCBs nos agradecerá. En este caso, el potencial eléctrico respecto a masa (o respecto a cualquier otro nodo) es desconocido. Si acoplamos una señal, por ejemplo, de forma capacitiva desde una pista que pase cerca o por debajo, y excitamos alguna resonancia del área de cobre, habremos creado una *antena tipo parche microstrip*. **¿Cuál es la solución?** Generalmente, conectar el área de cobre a masa o a una tensión de alimentación. Pero no basta conectar a masa de cualquier forma, como veremos en el siguiente caso.



Figura 5.13. Ejemplo de potencial antena parásita, indicada por la flecha verde. Fuente:

<https://electronics.stackexchange.com/questions/257585/should-we-remove-unconnected-copper-island-among-connected-traces>

- **Un área de cobre en capa externa conectada a masa en un punto.** Si el área no es pequeña respecto a la longitud de onda de interés, la conexión a masa garantiza una tensión “estable” sólo en las inmediaciones de la conexión. Conforme nos alejemos, la distancia deja de ser eléctricamente pequeña y permitimos que un acoplamiento de energía (por ejemplo, de una pista cercana) induzca un potencial variable en el área de cobre. ¡Bingo! Otra antena tipo parche *microstrip*. **¿Cuál es la solución?** Conectar el área de cobre a masa en múltiples puntos, a una distancia inferior a $\lambda/10$. La Figura 5.14 muestra un ejemplo.
- **Un disipador sobre un circuito integrado no conectado a masa (o conectado en sólo un punto).** El disipador puede estar flotante bien porque el adhesivo o pasta térmicos son aislantes eléctricos, o porque el encapsulado es plástico. Si el propio circuito integrado o una pista cercana inyectan energía por acoplamiento capacitivo en el disipador, éste se puede comportar también como una antena tipo parche parásita y radiar. **¿Cuál es la solución?** Conectar el disipador a masa y hacerlo en más de un punto. También podemos usar un disipador cerámico, menos eficiente como disipador, pero una nulidad como antena, al no ser conductor.

- Una tarjeta pinchada sobre una placa base: la inductancia de las conexiones a masa en el conector hará que masa de tarjeta y masa de placa base sean distintas, formando una **antena monopolo** parásita que podrá radiar en sus frecuencias de resonancia. **¿Solución?** Minimizar la inductancia total de las conexiones entre masa de la tarjeta y de la placa base por el simple procedimiento de poner muchos pines de masa en el conector o conectores.

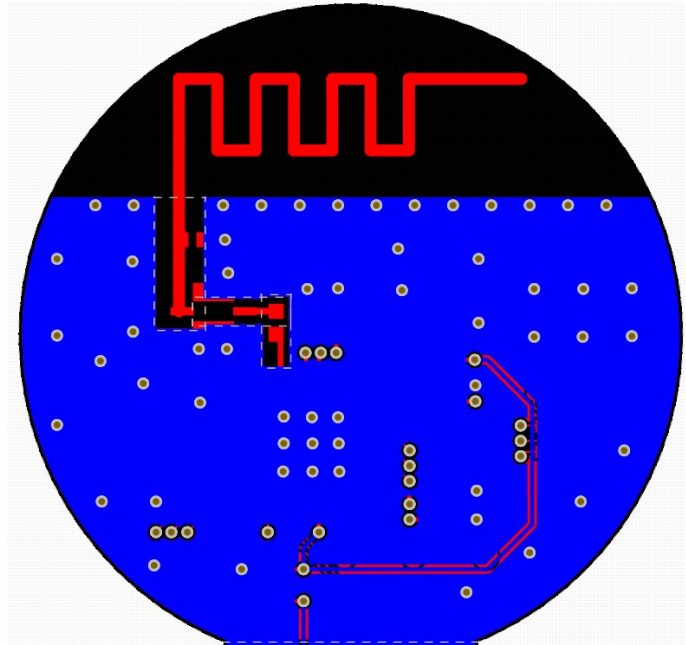


Figura 5.14. Plano de masa en capa *bottom* con múltiples vías al plano de masa en capa *top* para evitar efectos de antena parásita. Se trata del diseño *Bluetooth-Sentinel*, un ejemplo que viene con la distribución de Altium Designer.

- Un cable conectado a un PCB: es una variante de la anterior. Pero al ser la excitación en ruido en los planos de masa y alimentación, lo describiremos en una sección posterior.

Debemos describir, aunque sea sin entrar en detalles, cómo se excita y a qué frecuencias resuena una antena tipo parche y una antena tipo monopolo. En parte, por culturilla ingenieril. Pero fundamentalmente para que sepas reconocer una cuando la veas.

Antenas microstrip tipo parche

Comencemos por ver el aspecto que tiene una antena microstrip tipo parche intencionada.

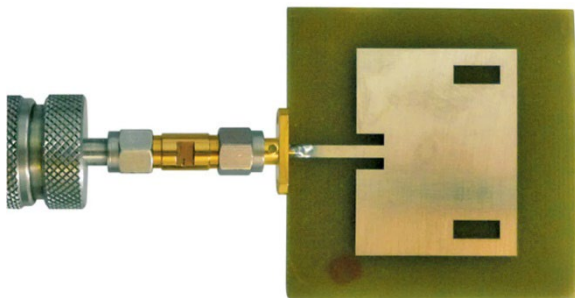
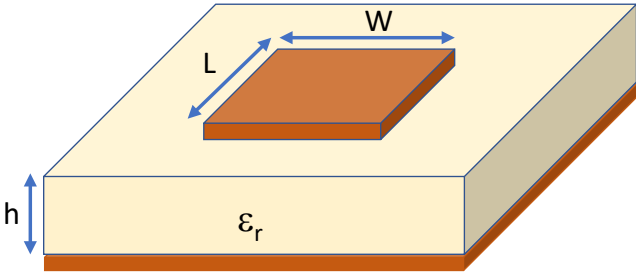


Figura 5.15. Antena para Wi-Fi/WiMAX.
Fuente: <https://www.mwrf.com/>

La Figura 5.15 muestra una antena real que se puede diferenciar de un área de cobre en uno de tus PCBs por las dimensiones y porque se alimenta de forma directa, no por acoplamiento por una pista cercana. Pero en esencia, es lo mismo.

En una primera aproximación, las frecuencias a las que resuena un rectángulo de cobre se calculan a partir de sus dos dimensiones, así como por la altura del parche sobre el plano de masa y la constante dieléctrica efectiva que es función del sustrato y de las dimensiones de la antena. La Figura 5.16 proporciona las expresiones necesarias para calcular las frecuencias de resonancia. Debes hacer m, n, p igual a 0, 1, 2, etc. Para calcular las frecuencias.

Una vez la antena básica está diseñada (dimensiones), hay que acoplar señal. La Figura 5.17 muestra dos formas de hacerlo: por ranura y por proximidad. En el primer caso, una pista cruza bajo una discontinuidad (ranura) en un plano, haciendo que radie y que dicha radiación se acople a la antena. En el segundo caso, una pista cercana o bajo la antena acopla energía por proximidad. Si la energía acoplada contiene alguna de las frecuencias de resonancia de la antena, radiará.



$$\epsilon_a = \frac{\epsilon_r + 1}{2} + \frac{\epsilon_r - 1}{2} \cdot \left(1 + 12 \frac{h}{W}\right)^{-0.5}$$

$$f_{mnp} = \frac{c_0}{2\sqrt{\epsilon_a}} \cdot \sqrt{\left(\frac{m}{h}\right)^2 + \left(\frac{n}{L}\right)^2 + \left(\frac{p}{W}\right)^2}$$

Figura 5.16. Antena tipo parche parásita (no intencionada) consistente en un área de cobre flotante en capa externa. Fuente propia

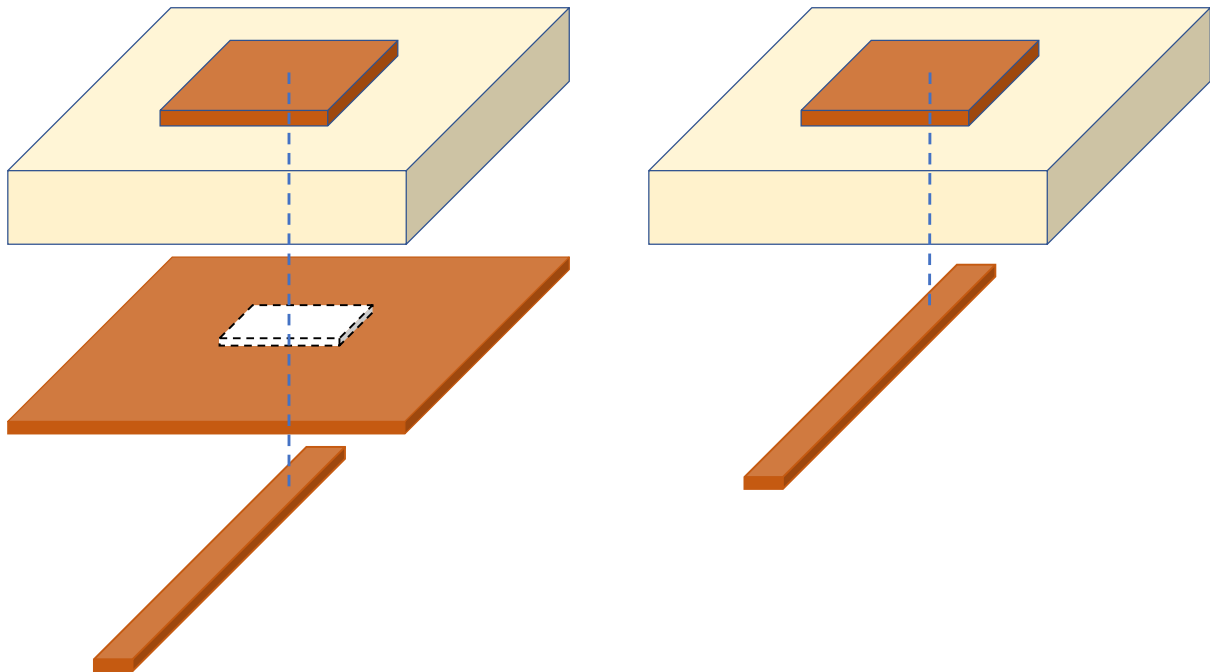


Figura 5.17. Dos formas de inyectar energía en una antena tipo parche. Por ranura (izquierda) y por proximidad (derecha). En el primer caso, aunque un plano separa la antena de la pista agresora, una ranura en el plano permite el acoplamiento de energía. En el segundo caso, el acoplamiento es directo entre la antena y una pista en una capa inferior sin planos entre las dos capas. Fuente propia

Radiación de una tarjeta pinchada en una placa base

Un módulo insertado en un *backplane* o una *motherboard* se comporta como un monopolo sobre un plano de masa; formando una antena muy eficiente a la frecuencia de resonancia del módulo. Las pequeñas diferencias de tensión originadas por la inductancia de las conexiones de masa en los conectores (en azul en la Figura 5.18) sirven de excitación para la antena. De este modo, **los planos de masa y alimentación del módulo hacen de antena**. La intensidad del campo radiado depende de la tensión de excitación de antena, que depende de la distancia entre pines de señal y masa (área del bucle, que afecta a la inductancia) y del número de pines de masa que rodean a un pin de señal.

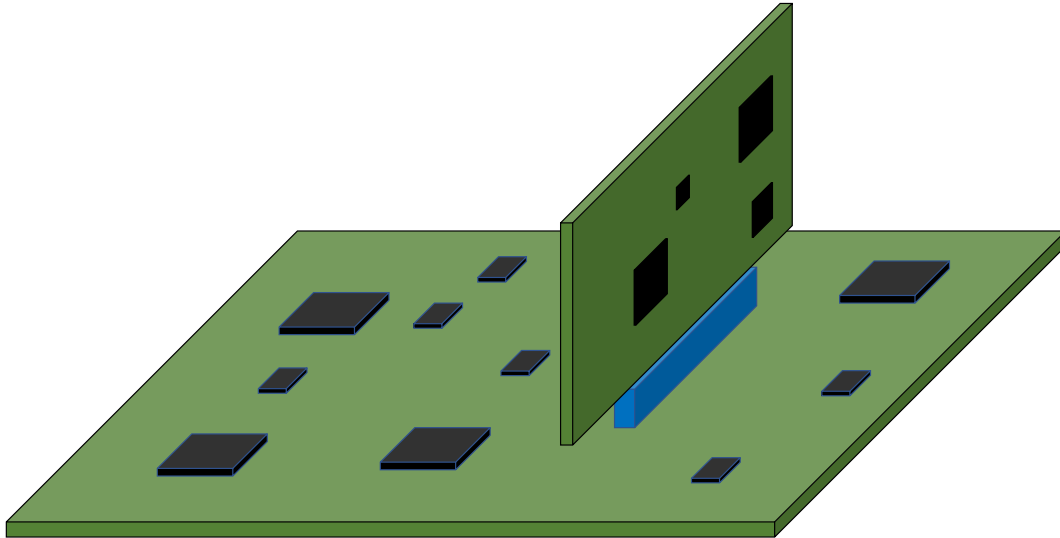


Figura 5.18. La tarjeta pinchada sobre una placa base radia como un monopolo. Es la impedancia (inductancia) de las conexiones a masa en el conector la que produce la excitación del monopolo. Fuente propia

Te recomiendo la lectura de la referencia [13] si quieres profundizar en la comprensión de esta subsección. El *test setup* descrito en el artículo consiste en un oscilador de 40 MHz seguido de un buffer en la placa base, que es alimentado a una placa hija pinchada verticalmente sobre ésta a través de un conector con un solo pin de masa). La radiación del monopolo (plano de masa en la tarjeta hija) produce una radiación que supera los límites que marca la normativa EMC definidos para equipos domésticos por EN55022 (ya obsoleta, ha sido sustituida por EN55032).

Este problema no es fácil de evaluar antes de producir un prototipo, y solucionarlo requiere mejorar las conexiones de masa y alimentación en los conectores para reducir su inductancia. Por lo general basta con aumentar el número de pines de masa y alimentación en los conectores.

Propagación de ruido en planos de masa y de alimentación

Este es un concepto que te puede costar. Porque hemos crecido (como ingenieros) con una fe en las propiedades mágicas de las masas: aquello a lo que llamemos masa, adquiere la increíble propiedad de estar a cero voltios constantes, en todos sus puntos. Pero aquello a lo que llamamos masa no es más que un conductor, generalmente de área elevada porque le dedicamos algún que otro plano entero en el PCB, pero que no se libra de las leyes de la física. Por tanto, entre dos puntos de ese conductor que llamamos masa no deja de haber una resistencia y una inductancia.

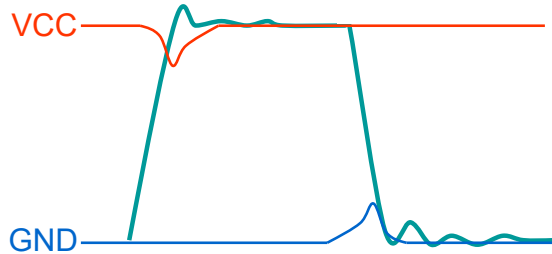


Figura 5.19. Ground bounce

Cuando un circuito integrado digital consume un pulso de corriente, esta corriente proviene de los planos de masa y de alimentación. Y, como hemos dicho, estamos sometidos al dictado de la expresión $V = R \cdot I + L \cdot dI/dt$. Es decir, se produce una caída de tensión en la línea o plano de alimentación y un rebote en la línea o plano de masa (*ground bounce*).

Nos hemos referido a circuitos digitales, y no analógicos: los primeros general estrechos pulsos y flancos que extienden su espectro en las frecuencias muy elevadas, donde mayor es el efecto del *ground bounce*.

Estos picos de corriente y estas fluctuaciones de tensión no se quedan como un efecto local, sino que encuentran una estructura por la que propagarse por todo el PCB: la *guía biplaca*. Una forma sencilla de visualizar este efecto es la onda que se propaga por un estanque cuando dejas caer una piedra.

Te presento de nuevo a la guía biplaca

Dos planos metálicos paralelos forman una *guía biplaca* que permite la propagación de ruido en masas y alimentaciones. Si la distancia entre planos es pequeña comparada con la longitud de onda ($h \ll \lambda/2$, lo que a efectos prácticos en electrónica ocurre siempre), se pueden propagar ondas electromagnéticas a partir de cualquier frecuencia.

Puedes verlo como un equivalente a una línea de transmisión *stripline*. Las perturbaciones se propagarán por todo el PCB y llagarán a los bordes, donde podrá producirse radiación.

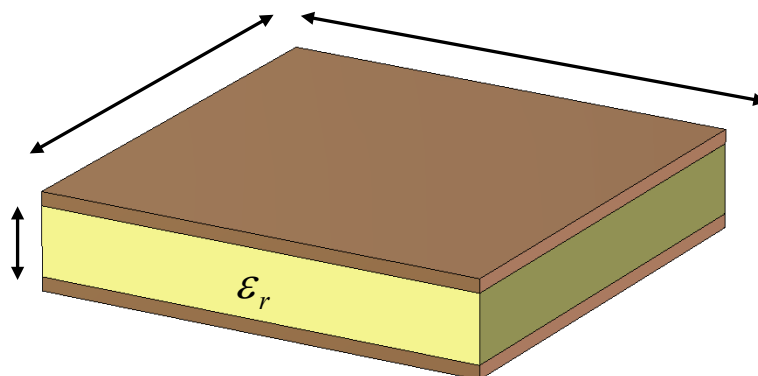


Figura 5.20. La guía biplaca permite la propagación de las perturbaciones por todo el PCB. Al llegar a los bordes, parte de la energía se refleja y parte se radia al exterior. La onda reflejada se propaga a su vez por todo el PCB, dando lugar a una interferencia entre múltiples ondas hasta que la atenuación va reduciendo la amplitud. Es lo mismo que ocurre en un pequeño estanque o en una bañera cuando dejas caer una piedra. Fuente propia

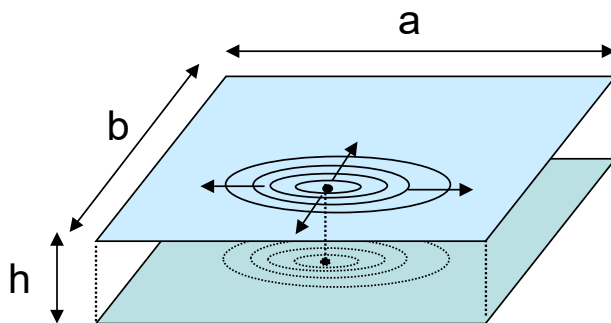
En [14] se describen expresiones para estimar la radiación por los bordes del PCB, aquí sólo daremos un límite superior:

$$|E_{max}| = V_{max} \cdot f \cdot \frac{\sqrt{\mu\epsilon_0 \cdot (W^2 + L^2)}}{r}$$

Vemos que la radiación de bordes aumenta linealmente con la frecuencia y es proporcional a la amplitud de la diferencia de tensión entre placas (V_{max}).

Cuidado con las resonancias

Las dimensiones de la guía biplaca determinan sus frecuencias de resonancia, donde la perturbación que se propaga por la estructura y la radiación por los bordes será mucho mayor a la que da la expresión anterior. Podemos calcular las frecuencias a las que resuena una guía biplaca a partir de la siguiente expresión:



$$f_{mn} = \frac{c}{2\sqrt{\epsilon_r}} \sqrt{\left(\frac{m}{a}\right)^2 + \left(\frac{n}{b}\right)^2}$$

Por ejemplo, un PCB de 10x10 cm² tendrá sus primeras tres frecuencias de resonancia a 231 MHz, 327 MHz, y 463 MHz.

Vías como generadores de perturbaciones

Las guías biplaca no se excitan sólo por ruido en los planos. En un PCB hay cientos de pequeñas antenas que también inyectan señal en estas estructuras: se trata de las vías de señal (Figura 5.21). Este efecto se produce con mayor eficiencia a frecuencias altas, y poco podemos hacer por evitarlo.

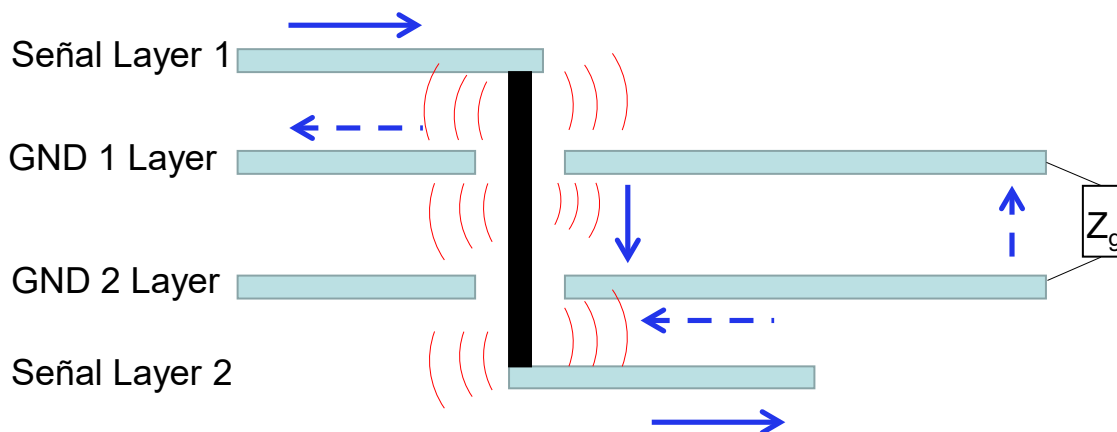


Figura 5.21. Una vía pasante entre capas *top* y *bottom* actúa como una pequeña antena que excita la guía biplaca (en este caso formada por las capas GND1 y GND2). Fuente propia

Reducir la generación de ground bounce: condensadores de desacoplo

Una forma de reducir el ruido producido por la conmutación de los circuitos digitales en los planos de masa y alimentación es **crear cortocircuitos en alta frecuencia mediante la red de condensadores de desacoplo**.

Dedicaremos un día a aprender a diseñar la red de condensadores de desacoplo. Por el momento bastará con mencionar algunas de sus características y propiedades:

- Un condensador de desacoplo, entre Vcc y masa es un almacén de carga cerca del circuito integrado. Las demandas de corriente de media-alta frecuencia (por debajo de 200 MHz) son proporcionadas por condensadores cercanos al integrado con menor inductancia que los planos (Figura 5.22 izquierda). El resultado es una menor fluctuación local en los planos de alimentación.
- Se establece una jerarquía de condensadores, que cubren bandas de frecuencias bajas, medias y medias-altas, en función de su inductancia parásita (Figura 5.22 derecha).
- En la banda de frecuencias en las que el condensador presenta baja impedancia (miliohmios), se produce un cortocircuito entre Vcc y masa que atenúa enormemente la propagación en la guía biplaca. Las vías entre planos de masa tienen el mismo efecto.

Guarda este concepto en la cabeza: lo desarrollaremos con detalle el Día 7.

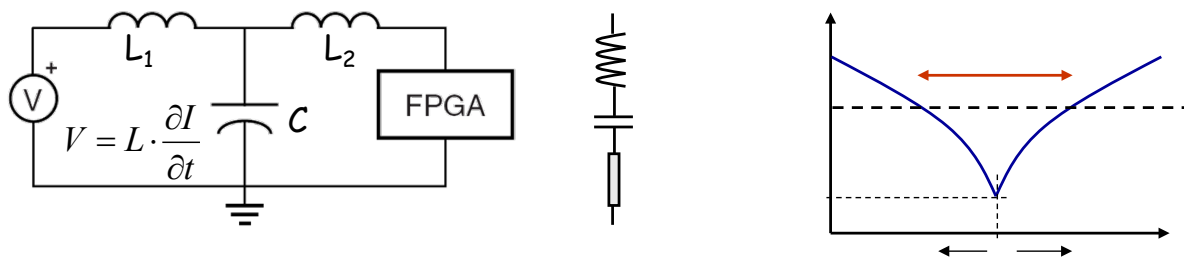


Figura 5.22. Condensadores de desacoplo: concepto (izquierda) y comportamiento en frecuencia de un condensador real (derecha). Fuente propia

Reducir la propagación del ground bounce: aislamiento entre secciones del diseño

La Figura 5.23 representa un diseño con el plano de masa partido (*split plane*) en tres secciones: analógica, digital y entrada/salida. Igualmente podría haber otra sección: potencia. Representa también un aspecto en el que hay menos consenso entre diseñadores. ¿Debe haber un plano único de masa? ¿O debo separar las masas analógicas, digitales y de potencia? ¿Por qué una masa de entrada/salida? ¿Y qué hay de chasis, dónde se conecta? ¿Cómo debemos interconectar las diferentes masas? ¿Con un pequeño puente de cobre? ¿Con una ferrita? ¿O todas las masas unidas a un único punto? Las decisiones que tomemos a este respecto van a afectar, a veces gravemente, a la funcionalidad del diseño y al resultado de ensayos EMC.

No vamos a abordar hoy esta cuestión en profundidad: vamos a centrarnos en la limitación de la propagación del *ground bounce*. En días posteriores volveremos al tema de la separación de las masas del diseño bajo puntos de vista distintos. Poco a poco.

Las áreas sucias del diseño serán la sección digital y la de potencia. Pero desde el punto de vista de la radiación del PCB, la relevante será la digital. La sección digital, con sus picos de demanda estrechos, inyecta perturbaciones de alta frecuencia en las guías biplaca del PCB. En la página 117 vimos que la radiación que escapa por los bordes del PCB es proporcional a la frecuencia. La radiación en modo diferencial es proporcional al cuadrado de la frecuencia. Las antenas parásitas en un PCB resuenan a altas frecuencias. En cambio, las secciones de potencia producen picos de elevada amplitud, pero de frecuencias mucho menores, que se radiarán con mayor dificultad. Un punto más a favor de la sección digital como fuente más peligrosa del ruido: en los ensayos de emisiones radiadas se mide sólo la energía radiada por encima de 30 MHz. Los circuitos de potencia no son rival. Será diferente en los ensayos de emisiones conducidas, a menos limitados a 30 MHz.

Volviendo a la Figura 5.23, observamos que la sección digital está separada de las demás partiendo en plano de masa en varias áreas. Con esto buscamos frenar la propagación del *ground bounce* generado en la sección

digital para que no ensucie la sección analógica. La unión entre diferentes masas puede hacerse mediante ferritas, que presentan una impedancia elevada a frecuencias altas y garantizan una equipotencialidad en continua.

En ocasiones resulta inevitable mantener la unión de los planos, si necesitamos dar continuidad a los caminos de retorno de algunas señales críticas. No hay una solución válida no mucho óptima para todos los diseños. Conocimientos y experiencia aconsejarán la mejor forma de abordar el esquema de masas en un diseño.

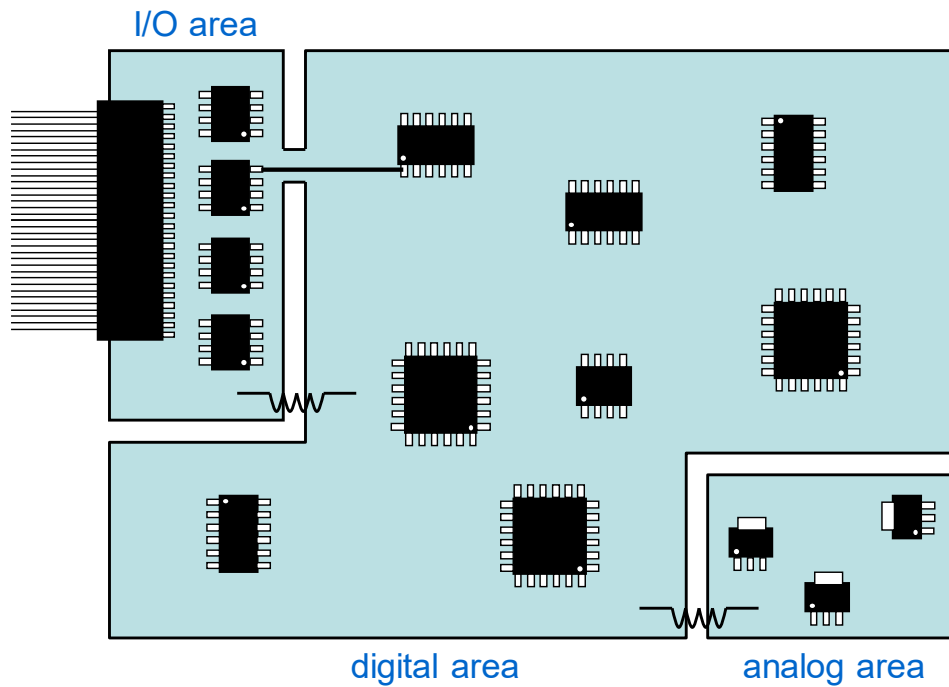


Figura 5.23. Separación de masa analógica, digital y masa de E/S para reducir el acoplamiento de ruido en secciones sensibles.
Fuente propia

Radiación por bordes del PCB

En las páginas 116 y 117 hemos hablado de cómo las señales acopladas a una guía biplaca se radiaban al exterior al alcanzar el borde del PCB. Aun no siendo la principal contribución de la radiación del PCB, vamos a comentar dos técnicas de reducción de la radiación.

Regla 20-H

Consiste en recortar el plano de alimentación por sus bordes una distancia igual a $20 \cdot H$, siendo H la separación entre planos de la guía biplaca, de modo que las líneas de campo se cierren entre VCC y GND reduciendo la radiación. La separación entre VCC y GND no afecta a la reducción del campo radiado [15], de modo que no es importante que los planos estén más o menos separados.

En la práctica, puedes encontrar complicado recortar $20 \cdot H$ el plano de alimentación. Debes saber que con $5 \cdot H$ ya se obtiene una importante reducción en la radiación (de forma orientativa, un factor 5), de modo que por poco que puedas hacer, habrá sin duda un beneficio. Mi consejo: no te vuelvas loco con esta regla. ¡Hay cosas cien veces más importantes que puedes hacer para reducir la radiación de un PCB!

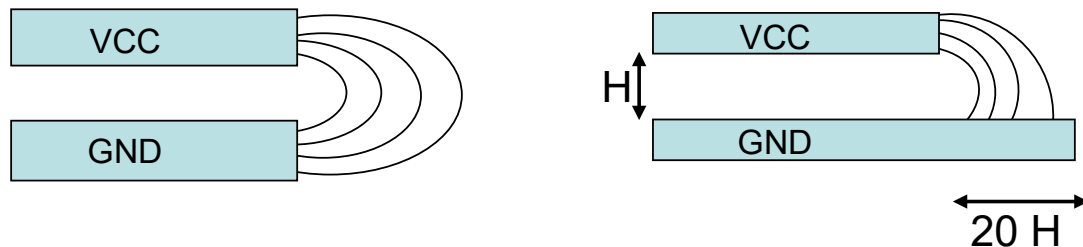


Figura 5.24. Regla 20-H para reducir la radiación por los bordes del PCB. Fuente propia

PCB stitching

Consiste en unir mediante vías varios planos de masa a lo largo del borde del PCB. Si la separación entre vías es suficientemente pequeña (no mayor de $\lambda/10$ a la frecuencia de interés), el campo queda confinado en el interior del PCB y es atenuado entre 10 y 30 dB. Eso sí, el campo confinado puede acoplarse a pistas y pasar a las capas externas a través de vías. Desde estas capas externas puede radiarse, reduciendo la efectividad del *stitching* [16]. El *stitching* produce un aumento de las resonancias y un aumento del *crosstalk* [17], lo que puede ser relevante en determinados diseños.

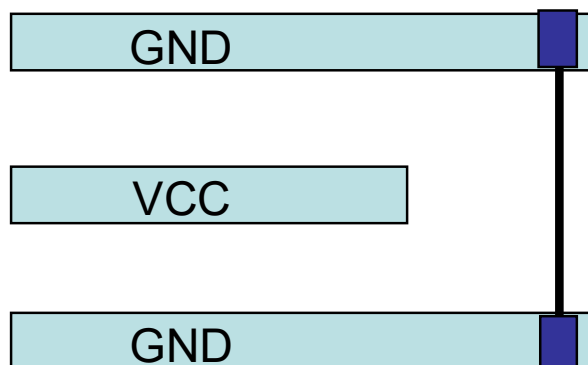


Figura 5.25. Concepto de PCB stitching para reducir la radiación por bordes del PCB. Fuente propia

Resumen después de la lección

Conocer las razones por las que radia un PCB nos ayudará a mejorar el rutado para sacar mejor nota en los ensayos de emisiones radiadas. Aunque asimilar los contenidos puede no ser sencillo, la buena noticia es que no será necesario, por lo general, realizar complicadas simulaciones EMC, sino que bastará con aplicar una serie de buenas prácticas.

Las pistas en capa externa junto a sus caminos de retorno, aún en condiciones ideales, radian por **modo diferencial**. Hemos comentado que la radiación es proporcional al área del bucle y al cuadrado de la frecuencia, de donde podemos derivar las siguientes buenas prácticas para rutado en capa externa:

- Hacer las pistas tan cortas como sea posible.
- Situar el plano de referencia tan cerca como sea posible de la capa de señal.
- Reducir el ancho de banda de la señal, bien eligiendo drivers tan lentos (es decir, tiempo de flanco elevado) como permita la aplicación, bien usando terminaciones de línea.
- Rutaremos por capa interna las señales críticas.

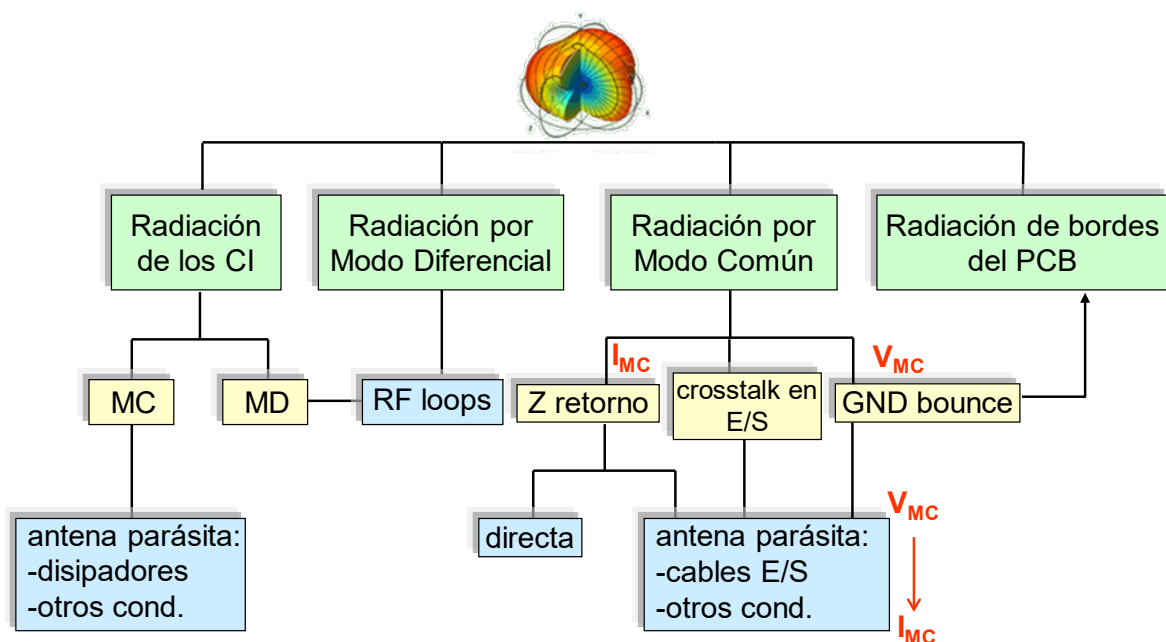


Figura 5.26. Diagrama resumen de las causas y tipos de radiación en un PCB

Cuando introducimos una discontinuidad en el camino de retorno de la señal, producimos conversión de modo diferencial a modo común, que se radia mucho más eficientemente. Esto es lógico, puesto que en modo diferencial la radiación de los caminos de señal y de retorno las corrientes circulan en sentidos opuestos, produciendo cancelación de los campos radiados. En cambio, en la **radiación por modo común**, ambos caminos se refuerzan. La radiación en modo común es proporcional a la frecuencia y la longitud de la pista, de modo que las recomendaciones anteriores son también válidas. Pero lo que debemos hacer, por encima de todo, es:

- Cuidar los caminos de retornos de corrientes. En especial, evitar que las pistas críticas crucen discontinuidades (ojo con *split planes*, ranuras y antipads) y garantizar que los cambios de capa de señal no supongan una discontinuidad.

Debemos evitar las estructuras conductoras flotantes o insuficientemente conectadas a una tensión de referencia, lo que denominamos **antenas parásitas**. Disipadores y áreas de cobre en capas externas no conectadas son estructuras candidatas a ser excitadas por acoplamiento de energía desde una pista cercana, por ejemplo. Si esta energía acoplada incluye alguna de las frecuencias de resonancia de la estructura, radiará eficientemente. Ahora que lo sabemos, tendremos cuidado con ir dejando antenas parásitas en nuestros diseños.

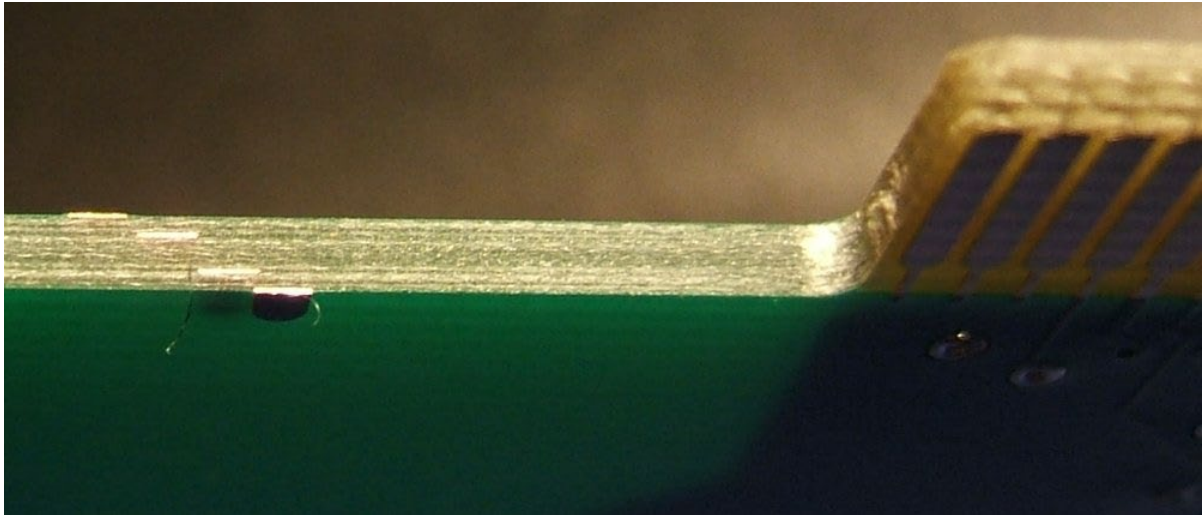
El [ruido en los planos de masa y alimentación](#) se propagará por todo el PCB (o por todo el plano, si usamos *split planes*). Al llegar a los bordes del PCB, parte de la energía se radiará. La regla 20-H nos da una forma razonablemente sencilla de reducir este efecto, que, como ya hemos comentado, no es de los más importantes. Pero todo suma. Finalmente, el ruido en los planos de masa y alimentación puede acoplarse a cables, dando lugar a una peligrosa radiación en modo común.

Recursos para aprender más

TEC: IEEE Transactions on Electromagnetic Compatibility

- [1] Mark I. Montrose, “EMC and the printed circuit board: design, theory and layout made simple”, IEEE Press, ISBN:0-7803-4703-X
- [2] TEC vol.46, no.4, p.580, “Electromagnetic Interference (EMI) Reduction From Printed Circuit Boards (PCB) Using Electromagnetic Bandgap Structures”
- [3] TEC vol.45, no.3, p.386, “The Radiation of a Rectangular Power-Bus Structure at Multiple Cavity-Mode Resonances”
- [4] Microwave and Wireless Components Letters vol.13, no.1, p.21, “A Novel Power Plane With Integrated Simultaneous Switching Noise Mitigation Capability Using High Impedance Surface”
- [5] TEC vol.40, no.2, p.111, “Study of the Ground Bounce Caused by Power Plane Resonances”
- [6] TEC vol.47, no.3, p.479, “Design and Modeling of High-Impedance Electromagnetic Surfaces for Switching Noise Suppression in Power Planes”
- [7] IEEE Transactions on Advanced Packaging, vol.22, no.3, p.274, “Reducing Simultaneous Switching Noise and EMI on Ground/Power Planes by Dissipative Edge Termination”
- [8] TEC vol.46, no.4, p.580, “Analytical Model for the Rectangular Power-Ground Structure Including Radiation Loss”
- [9] TEC vol.47, no.2, p.227, “Analysis on the Effectiveness of the 20-H Rule for Printed-Circuit-Board Layout to Reduce Edge-Radiated Coupling”
- [10] TEC vol.47, no.2, p.219, “On the Electromagnetic Radiation of Printed-Circuit-Board Interconnections”
- [11] TEC vol.43, no.4, p.671, “Power-Supply Decoupling on Fully Populated High-Speed Digital PCBs”
- [12] TEC vol.45, no.1, p.22, “Power-Bus Decoupling With Embedded Capacitance in Printed Circuit Board Design”
- [13] TEC vol.43, no.4, p.549, “Electrical Behavior of Decoupling Capacitors Embedded in Multilayered PCBs”
- [14] TEC vol.43, no.4, p.538, “EMI Mitigation With Multilayer Power-Bus Stacks and via Stitching of Reference Planes”
- [15] IEEE International Symposium on Electromagnetic Compatibility, vol.2, p.793, “Impact on Radiated Emissions of Printed Circuit Board Stitching”, 1999
- [16] TEC vol.46, no.1, p.33, “Numerical and Experimental Investigation of Radiation Caused by the Switching Noise on the Partitioned DC Reference Planes of High Speed Digital PCB”
- [17] TEC vol.44, no.1, p.165, “High-Performance Inter-PCB Connectors: Analysis of EMI Characteristics”

Día 6. Diseño de estructuras de PCB multicapa



[Creative Commons Attribution-Share Alike 3.0 Unported](#). Author: [Sieobserver](#) Christoph Kappel

Llegados a este punto, sabemos qué es una línea de transmisión y sus realizaciones en el PCB: líneas microstrip y stripline. Sabemos que la pista es la mitad de la historia y que el camino de las corrientes de retorno es la otra mitad. Sabemos que hemos de mimar el retorno. Sabemos lo que son las reflexiones, qué efectos producen y cómo reducirlas mediante diferentes técnicas de terminación. Sabemos cómo radian las pistas y otras estructuras en el PCB.

Es momento de aplicar los criterios y consideraciones que hemos aprendido a la definición de la estructura multicapa del PCB. No se trata simplemente de escoger una estructura estándar de un fabricante de PCBs. Hemos de planificar cuidadosamente cuántas capas de rutado vamos a necesitar, por qué capas va a ir cada grupo de señales, cuáles van a ser los caminos para las corrientes de retorno y cómo vamos a garantizar su continuidad. Hemos de decidir cuántos planos de alimentación necesitamos, si éstos serán de un solo nodo o si los partiremos en varias áreas separadas para diferentes alimentaciones (split planes). Y hemos de definir y asegurar las impedancias de línea que necesitamos en cada capa.

Diseñar una estructura de PCB multicapa consiste en responder a las preguntas anteriores y a otras más avanzadas.

Nuestros objetivos para esta lección son:

- Comprender los elementos constituyentes y el proceso de fabricación de un PCB
- Conocer las limitaciones que imponen los fabricantes
- Comprender las metas que buscamos al diseñar la estructura de un PCB (stack-up)
- Saber diseñar y especificar stack-ups

Estructura de un PCB multicapa

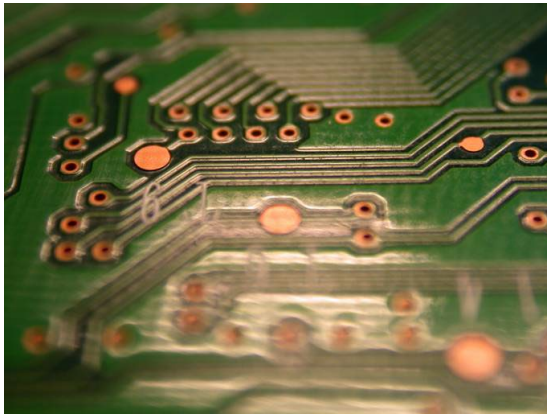


Figura 6.1. Imagen libre. Fuente: <http://www.imageafter.com>

Un PCB es un sándwich de varios pisos formado por capas de metal y capas de material aislante que dotan de rigidez al conjunto. Como metal, se elige cobre por su alta conductividad eléctrica, lo que reduce las pérdidas (atenuación) a bajas y medias frecuencias. Como material aislante se emplean generalmente capas de fibra de vidrio y resina epoxi, un material resistente al fuego (FR-4 o FR4, *flame retardant*).

Este material base, FR4 (aunque en determinadas aplicaciones se usan materiales cerámicos o de otro tipo) viene en forma de sándwich de tres capas: cobre-FR4-cobre. Con estos sándwiches se fabrican las capas internas del PCB. Estos sándwiches reciben el nombre de **núcleos (cores)**.

Para adherir entre sí los núcleos se usa **pre-preg** (pre-impregnado), que no es más que FR4 sometido a un proceso de curado (presión y calor) parcial, de modo que es todavía blando y adhesivo. Para evitar que se complete el polimerizado, el pre-preg ha de conservarse refrigerado hasta que se añada a la pila de capas que formarán el PCB.

Finalmente, como tercer material, junto a los núcleos y las capas de pre-preg, encontramos láminas de cobre, que se usan generalmente para las capas externas del PCB. Estas láminas se adhieren al núcleo adyacente mediante pre-preg. El resultado es un apilado de capas, que no por casualidad se denomina *stack-up* en inglés y *empilage* en francés.

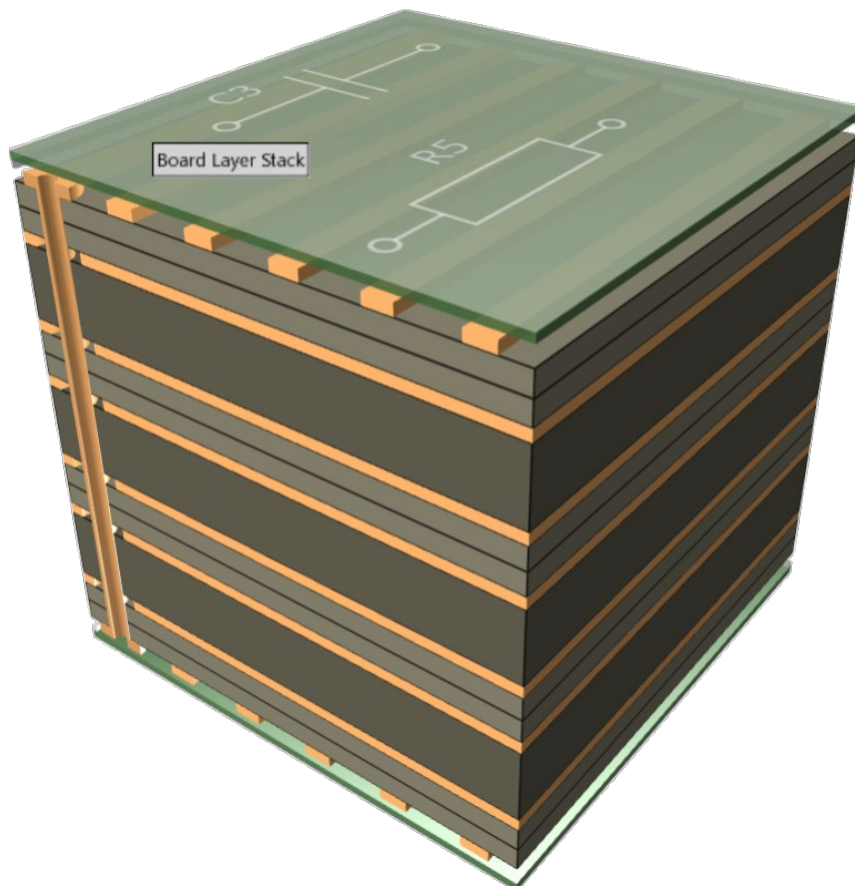


Figura 6.2. Esquema de una estructura PCB de 8 capas: tres núcleos definen las capas internas. Dos láminas delgadas de cobre, las externas. Cada uno de estos cinco elementos está adherido al adyacente por una o más capas de pre-preg. Fuente propia

En la Figura 6.2, un PCB de 8 capas metálicas, podemos observar los tres núcleos internos que forman las 6 capas metálicas internas. Estos tres núcleos están unidos entre sí mediante capas de *pre-preg*. Las capas metálicas *top* y *bottom* están adheridas a los núcleos también mediante *pre-preg*. Las capas de color verde son las capas de pasivación, máscara de soldaduras o *solder mask*, como prefieras nombrarla.

Una fotografía de un corte transversal de un PCB real (Figura 6.3) no permite diferenciar fácilmente entre *pre-preg* y el FR4 de los núcleos. Esto es lógico, puesto que el *pre-preg* no es más que FR4 parcialmente curado. Una vez se ensambla toda la estructura y se somete a presión y calor, se completa la polimerización del *pre-preg*, se elimina humedad y el PCB queda firmemente unido. Como las capas de *pre-preg* siempre están en contacto con cobre, éste no puede ser totalmente liso, sino que ha de tener cierta rugosidad para favorecer la adherencia.

Recuerda una regla sencilla pero importante: las capas internas se fabrican con núcleos. Las capas externas son láminas de cobre sobre pre-preg. Las reglas están para ser rotas y a veces puedes hacer las cosas de forma distinta, pero en lo sucesivo asumamos que las capas internas se fabrican con núcleos y que las capas externas son láminas de cobre sobre pre-preg.

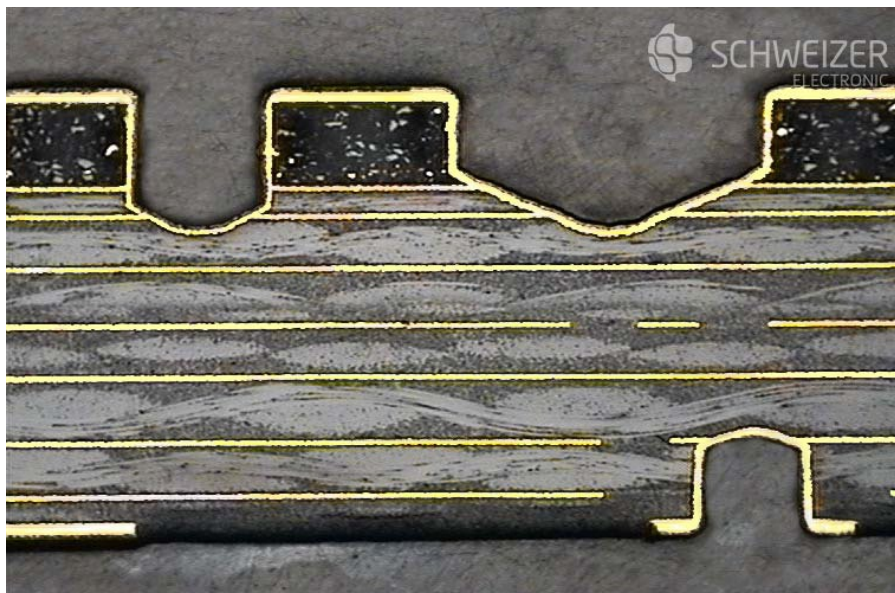


Figura 6.3. Corte transversal de un PCB. Fuente¹: <https://www.schweizer.ag/en/products-solutions/cross-section-of-the-month.html>.

FR4 y *pre-preg*

El material base está compuesto a base de resina epoxi reforzada con tela de fibra de vidrio (principalmente SiO₂, si bien están presentes otros óxidos, Figura 6.4 izquierda) y completamente curado (polimerizado por calor que produce endurecimiento). Cuando decimos “tela de fibra de vidrio” nos referimos realmente a un tramado de tela (Figura 6.4 derecha).

¹ Permiso solicitado a Schweizer en dos ocasiones para usar la imagen. Pendiente de respuesta desde el 4 de agosto de 2020.

Componente	Proporción
Dióxido de silicio	52-56 %
Óxido de calcio	16-25 %
Óxido de aluminio	12-16 %
Óxido de boro	5-10 %
Óxido de magnesio	0-5 %
Óxido de Na, K, Fe y Ti	0-3 %
Compuestos de flúor	0-1 %

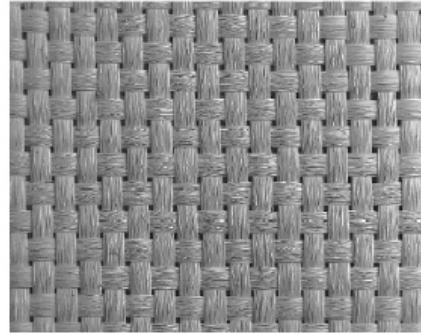


Figura 6.4. Composición de la fibra de vidrio (izquierda) y trama de la tela, en las que los filamentos tienen un grosor entre 2,5 y 13 micras (derecha).

Hay diferentes espesores de pre-preg más o menos estandarizados como el 106 (50-60 μm), 1080 (70-80 μm), 2116 (120-130 μm) y 7628 (180-200 μm). Los tipos 106 y 1080 en monocapa deben ser evitados (dentro de lo posible) para líneas diferenciales. La razón es que, si una pista del par cae sobre fibra y la otra sobre resina, hay diferencia apreciable de constante dieléctrica y por tanto de impedancia y velocidad de propagación.

Cualquier combinación lineal de los espesores de pre-preg disponibles (has de consultar al fabricante) es válida, de modo que tendrás una razonable flexibilidad para definir el valor de cada capa de pre-preg.

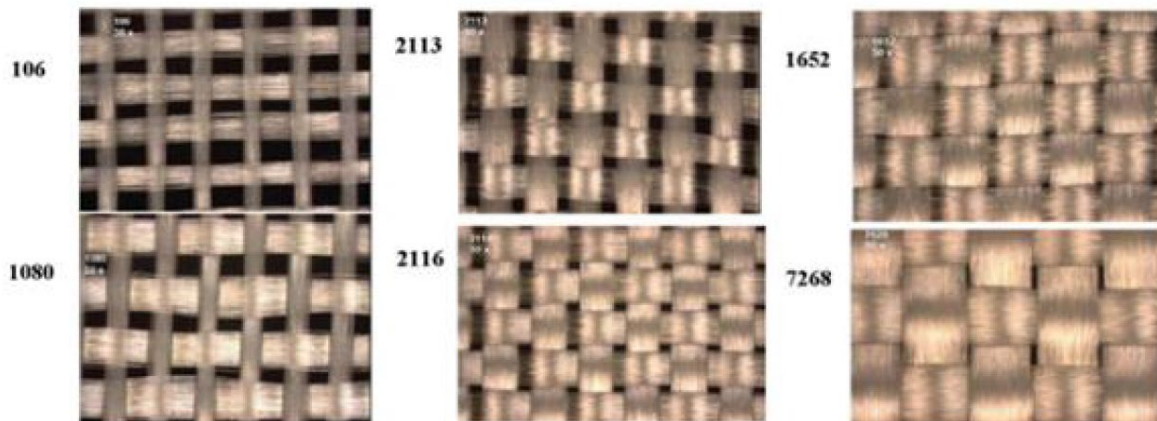


Figura 6.5. Seis de los tipos de trama de pre-preg más habituales. Fuente de la imagen: “PCB Stackup Design Considerations for Intel FPGAs”, AN-613, Intel, 28 de junio de 2017

Parámetros de fabricación de un PCB

Los fabricantes del PCB disponen de núcleos, capas de pre-preg y de láminas de cobre de distintos espesores. Tomamos como ejemplo las especificaciones de Lab-Circuits (www.lab-circuits.com), uno de los mejores fabricantes de PCBs en España:

- Espesores de cobre base: 17, 35 y 70 micras. A veces expresado en onzas (1 oz equivale a 34,3 micras)
- Grosor de FR4 en núcleos: 130, 180, 240, 350, 500, 510 y 730 micras
- Grosos de pre-preg: 106 (48mm), 1080 (65mm), 2113 (88mm), 7628 (175mm). Ya sabes que puedes usar combinaciones lineales de estos valores

Es importante que no tomes estos espesores como valores exactos. Una vez definida estructura, te recomiendo que contactes con el fabricante de PCBs, quien revisará tu estructura y te comentará que cierta capa de pre-preg será algo menor del valor nominal, o que el proceso de metalización de vías aumenta el espesor de las capas de cobre externas. Si es importante para el diseño tener impedancias bien definidas, este paso es imprescindible.

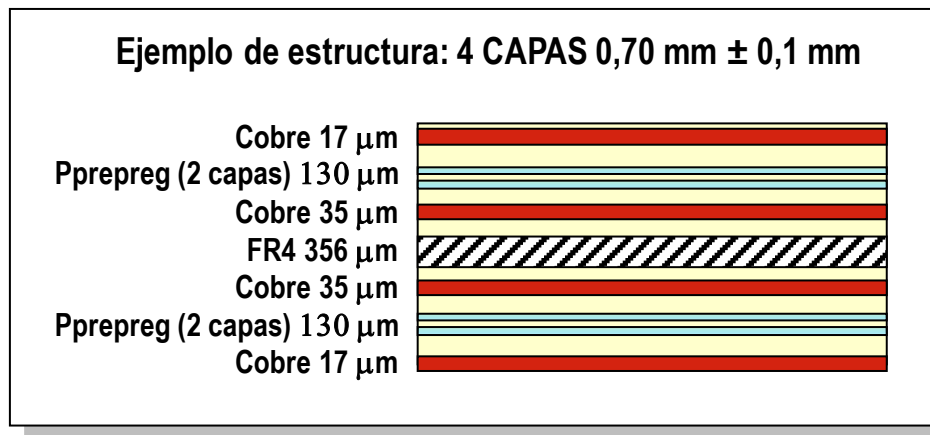
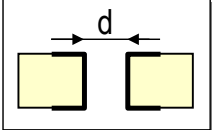
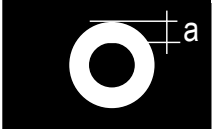
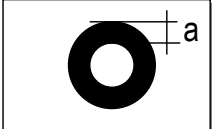


Figura 6.6. Ejemplo de especificación de un PCB. Fuente: www.lab-circuits.com

Los valores de ancho mínimo de pista y diámetro de vía, entre otros, determinan la clase del PCB. Cuanta mayor es la clase, más cara es la fabricación. Introduciremos estos y otros parámetros para la clase elegida como reglas en nuestra herramienta de rutado. Pasaremos a una clase superior cuando no consigamos rutar el diseño con los parámetros de la clase actual. La

Tabla 6.1 recoge algunos de los parámetros que determinan la clase.

Tabla 6.1. Algunos de los parámetros que definen la clase de fabricación de un PCB. Fuente: www.lab-circuits.com

	Parámetro	Clase 4	Clase 5	Clase 6	Clase 7
	diámetro mín. metalizado	0.3	0.3	0.2	0.15
	ancho/espacio mín. conduct. extern. (17 μ m Cu) conduct. intern. (17 μ m Cu)	0.2 0.15	0.15 0.125	0.125 0.1	0.1 0.075
	Aislamiento capas intern. masa alimentación	0.4	0.3	0.25	0.2
	corona mín. capa externa corona mín. capa interna	0.17 0.22	0.13 0.19	0.10 0.15	0.075 0.125

Rara vez he podido trabajar en clase 6, casi siempre he tenido que hacerlo en clase 7. Pero es sólo por el tipo de diseños que me ha tocado abordar. Recuerda: un valor de clase más baja implica también menor coste.

Vías

Otro parámetro importante y que debes verificar (acaba habiendo sorpresas a este respecto) es el *via aspect ratio*, es decir, el cociente entre longitud y diámetro interno de la vía. Hay un valor máximo que especifica el fabricante. Por ejemplo, para LabCircuits es de 13 en clase 7. Eso quiere decir que una vía pasante de 0,15 mm de diámetro no puede usarse en un PCB con un grosor superior a 1,95 mm. Como referencia, la mayoría de los PCBs tienen un espesor cercano a 1,65 mm (65 mils).

Las **vías pasantes** son aquellas que atraviesan completamente el PCB (Figura 6.7). Una vez ensamblada toda la estructura del PCB y eliminado el cobre no necesario de las capas *top* y *bottom*, dejando pistas y pads como especifican los ficheros del diseño, se taladran los orificios de las vías pasantes. Después, se somete el PCB a un proceso de metalización, en el que se deposita cobre para formar las paredes internas de las vías. En este proceso aumenta el espesor de cobre en capas externas alrededor de 20 micras. No tener esto en cuenta provocará una variación en la impedancia de las capas externas en torno a un ohmio. No es mucho, pero sale gratis hacerlo bien, ¿cierto?

Las **vías ciegas** se extienden desde una capa externa hasta una interna, sin atravesar todo el PCB. Sólo hay dos razones para usar este tipo de vía: (1) cuando el diseño es muy denso y queremos liberar espacio para rutado en las capas no afectadas y (2) cuando trabajamos con señales multi-gigabit y queremos mejorar la integridad de señal. ¿Cómo? Si una señal SATA ha de saltar de capa *top* a capa 3, el resto de la vía hasta *bottom* no sólo es inútil, sino que es una pequeña línea de transmisión terminada en circuito abierto (denominada **stub**) que puede provocar reflexiones. Así que, mejor eliminar el tramo innecesario. Se puede hacer de dos maneras: con una vía ciega o mediante *back-drilling* (un taladro, introducido desde la capa *bottom*, elimina el cobre que no hace falta hasta la capa 3). ¿Se puede crear una vía ciega hasta cualquier capa? No. Una vía ciega se crea en un proceso intermedio del ensamblado del PCB, haciendo un taladro pasante cuando hemos “montado” sólo una parte de los núcleos. Por tanto, es posible hacer un taladro de capa *top* a capa 3 (atravesamos un núcleo), o hasta capa 5 (atravesamos dos núcleos), pero no hasta capa 4. Por cierto, en la Figura 6.7, la vía ciega de *top* a capa 2 no tiene mucho sentido.

Una **vía enterrada** une dos capas internas. Y lo dicho para vías ciegas es de aplicación aquí también. Ha de atravesar un número entero de núcleos, de modo que es posible hacer vías enterradas entre capas 2 y 3, entre 2 y 5, pero no entre 2 y 4.

Como alternativa al taladrado mecánico está la **microvía láser**. Es posible atravesar un espesor de unas 100 micras de dieléctrico, no mucho más. De modo que si queremos usar esta solución habrá que tenerlo en cuenta al diseñar la estructura del PCB. El diámetro del taladro, unas 100 micras, es inferior al que se consigue con taladrado mecánico (unas 150 micras). Pero el pad de la microvía ha de ser mayor, de unas 300 micras. Es posible apilar dos o más microvías (*stacked microvias*), teniendo a un deficiente relleno de la microvía (Figura 6.7, derecha) como uno de los principales modos de fallo por estrés termo-mecánico.

En los años que llevo como diseñador, he usado casi siempre vías pasantes, y sólo en un par de diseños me he visto forzado a usar vías ciegas. Ten en cuenta que cada paso extra en el diseño cuesta dinero. El coste del PCB aumenta tal vez un 10-20% entre usar sólo vías pasantes y usar también otros tipos de vías.

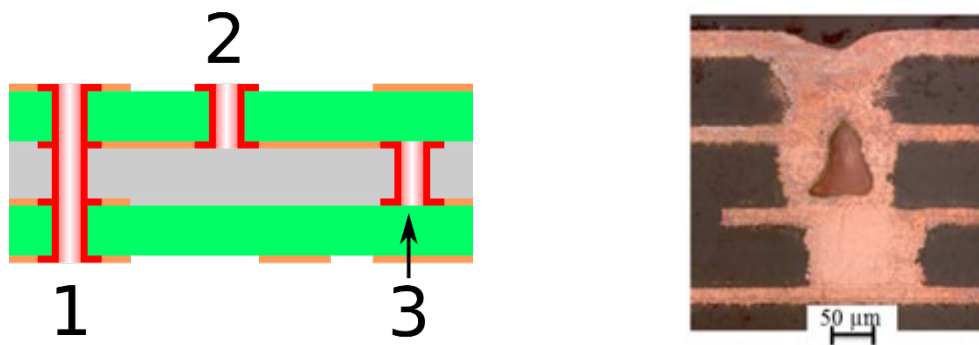


Figura 6.7. Tipos de vías: enterradas (3), ciegas (2), pasantes (1) y microvías (imagen de la derecha). Fuente de la imagen izquierda: Wikipedia, licencia CC-BY 3.0, autor [M adler](#). Fuente de la imagen derecha: Wikipedia, licencia CC-BY 3.0, autor [Yan Ning](#).

Impedancia controlada

Por “impedancia controlada” se entiende la verificación de que las imprecisiones y tolerancias del proceso de fabricación no han introducido variaciones de impedancia y que los valores de impedancia reales coinciden con los esperados. Son causa de variaciones de impedancia:

- Espesor de cobre y anchura de pista no uniformes
- Espesor y propiedades del dieléctrico no uniformes

Para verificar la ausencia de desviaciones importantes durante el proceso de fabricación respecto a los valores calculados, se utilizan placas de prueba denominadas cupones de test (*test coupons*). Se suele colocar un par de placas de prueba, una en cada extremo del panel (Figura 6.8). El diseñador del PCB debe especificar al fabricante qué capas contienen señales de impedancia controlada, el valor de impedancia(s) que queremos obtener en cada capa y la tolerancia.

Los cupones de test contienen pistas de unos 15 cm de longitud en cada capa de señal. Todos los planos de masa y alimentación se cortocircuitan mediante vías. Los extremos de cada pista se llevan a conectores en capa externa, a fin de poder medir posteriormente los parámetros más relevantes (impedancia de las pistas y posibles variaciones locales).

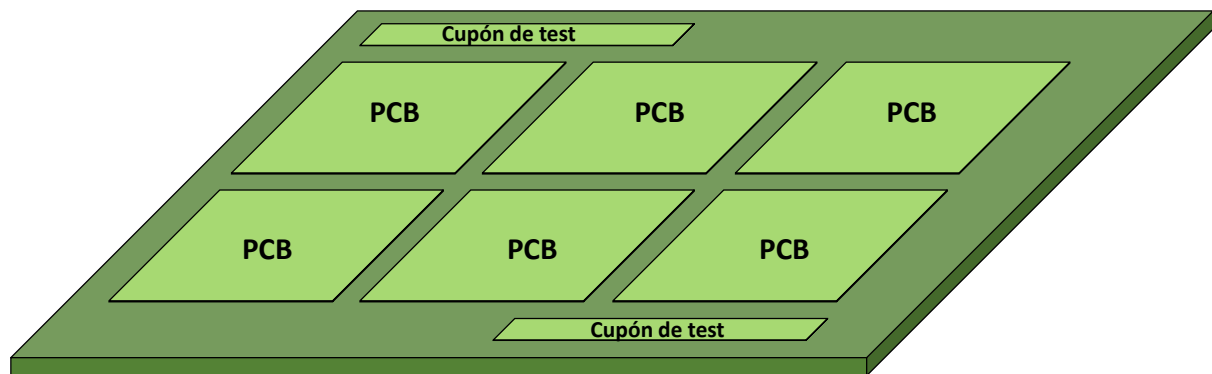


Figura 6.8. Cupones de test en un panel. Fuente propia

Una vez fabricado un panel, se analizan las pistas de los cupones de test mediante reflectometría en el dominio temporal (un instrumento que excita la señal con un pulso y recibe las reflexiones, lo que nos hablará tanto de la impedancia de línea como de la presencia de variaciones locales de impedancia). Si el cupón no pasa el test, el panel se desecha y se fabrica otro.

Obviamente, este proceso encarece el producto, pues requiere mano de obra adicional (el elemento más caro del proceso de producción) y obliga al fabricante a producir paneles extra si hay variaciones en el proceso.

Consideraciones para el diseño del stack-up

El stack-up define la estructura del PCB (número de capas, espesores de cobre y de dieléctrico y tipo de dieléctrico), si bien el proceso de diseño del stack-up es inseparable de la definición de la anchura de pistas para obtener las impedancias deseadas, tamaño de las vías, asignación de los planos de alimentación y masa y función que cumple cada capa.

¿Quién diseña el stack-up?

El número final de capas lo determinará el experto en layout, pero el diseño inicial del stack-up es (en mi opinión) responsabilidad del diseñador del sistema, que es quien conoce el diseño en profundidad, sabe qué líneas y buses son críticos, qué componentes son ruidosos y cuáles son sensibles. El experto en layout verá un nodo llamado, por ejemplo, “STR1” y eso no le dirá nada: lo mismo puede tratarse de una señal cuasi-estática y robusta que de un delicado reloj de 150 MHz.

Alguien que conozca el diseño debería no sólo proponer una primera versión del stack-up, sino del *floorplan* (ubicación de los componentes principales en el PCB) y del esquema de masas. También debe proporcionar información al experto en layout sobre el rutado preferido de determinados componentes y buses. Debe revisar periódicamente el progreso del rutado, proponiendo cambios allí donde haya habido una comunicación insuficiente o un problema de comunicación. En definitiva, quien diseña los esquemas adquiere (siempre en mi opinión) una responsabilidad que va mucho más lejos.

A medida que una señal pasa de una capa a otra, no debe notar un cambio de impedancia

Debes tener en cuenta que una señal, en todas las capas por donde pase, debe tener la misma impedancia: una diferencia de impedancias del 20% dará lugar a una reflexión del 10%. Normalmente definirás una impedancia cercana a 50 Ω para señales single-ended, y de 90 o 100 Ω para los pares diferenciales, según se trata de USB 2.0, PCI express, Ethernet, etc.

Una consecuencia de lo anterior es que la anchura de pista para lograr 50 Ω o la impedancia que desees puede ser distinta en cada capa. Debes introducir esta información en las reglas de diseño del programa de diseño de PCBs, de modo que la anchura de pista se ajuste automáticamente sin que necesites recordar los valores.

Señales que no se ven afectadas por la impedancia de línea (un pulsador, por ejemplo) no necesitan seguir estas reglas si hay un buen motivo, como, por ejemplo, asegurar suficiente sección de cobre para evitar sobrecalentamiento si la pista ha de llevar mucha corriente.

Simetría

Para lograr las impedancias de línea deseadas tendrás que ir jugando con los espesores de pre-preg y de núcleo, así como de anchura de pista en cada capa. Pero debes hacerlo asegurando que el stack-up sea simétrico respecto a su eje central: de lo contrario la estructura puede combarse (*warping*) debido a la diferencia de tensiones mecánicas.

Si puedes permitirte, usa sólo planos de masa para los caminos de retorno de las señales

Voy a pedirte que revises cuatro fragmentos breves de días anteriores: las secciones “Corriente de retorno” en la página 30, “(dos conductores)... que forma un bucle con una determinada inductancia” en la página 39, “Los caminos de retorno en líneas diferenciales”, en la página 92 y “Reducir la generación de señales en modo común” en la página 110.

Si, a medida que una señal cambia de capa sólo encuentra planos de masa como planos de referencia (caminos de las corrientes de retorno), dar continuidad a estos caminos sólo requiere situar vías que cortocircuiten los planos de masa cerca de las vías de señal. De esto ya hemos hablado. Veamos un ejemplo.

En la Figura 6.9 podemos ver a la izquierda un conector HDMI (J6) montado en capa bottom (azul) y los cuatro pares diferenciales asociados. El par superior pasa a una capa interna (Signal4, en amarillo, mira el stack-up en la Figura 6.10) y vuelve a capa bottom antes de llegar al buffer U47.

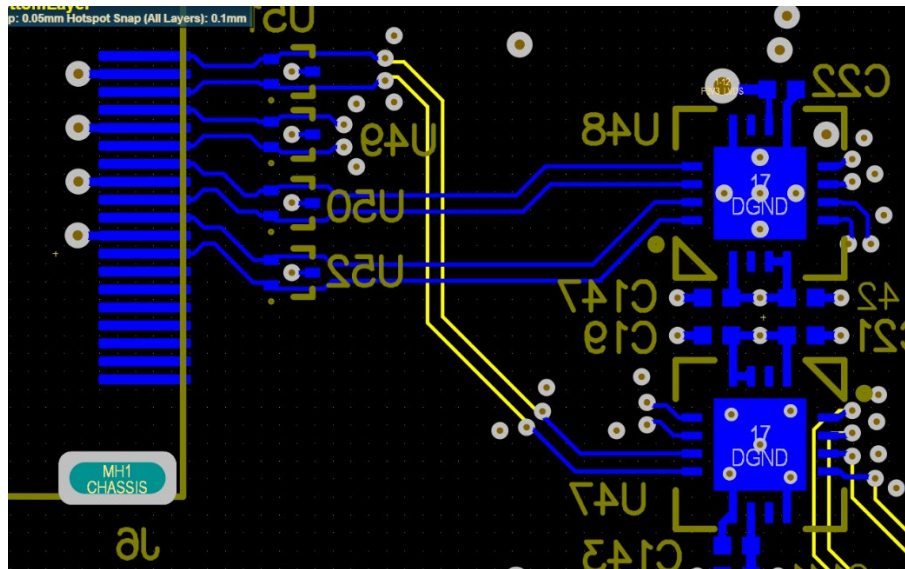


Figura 6.9. Damos continuidad a los caminos de retorno de corrientes con vías a masa. Fuente propia

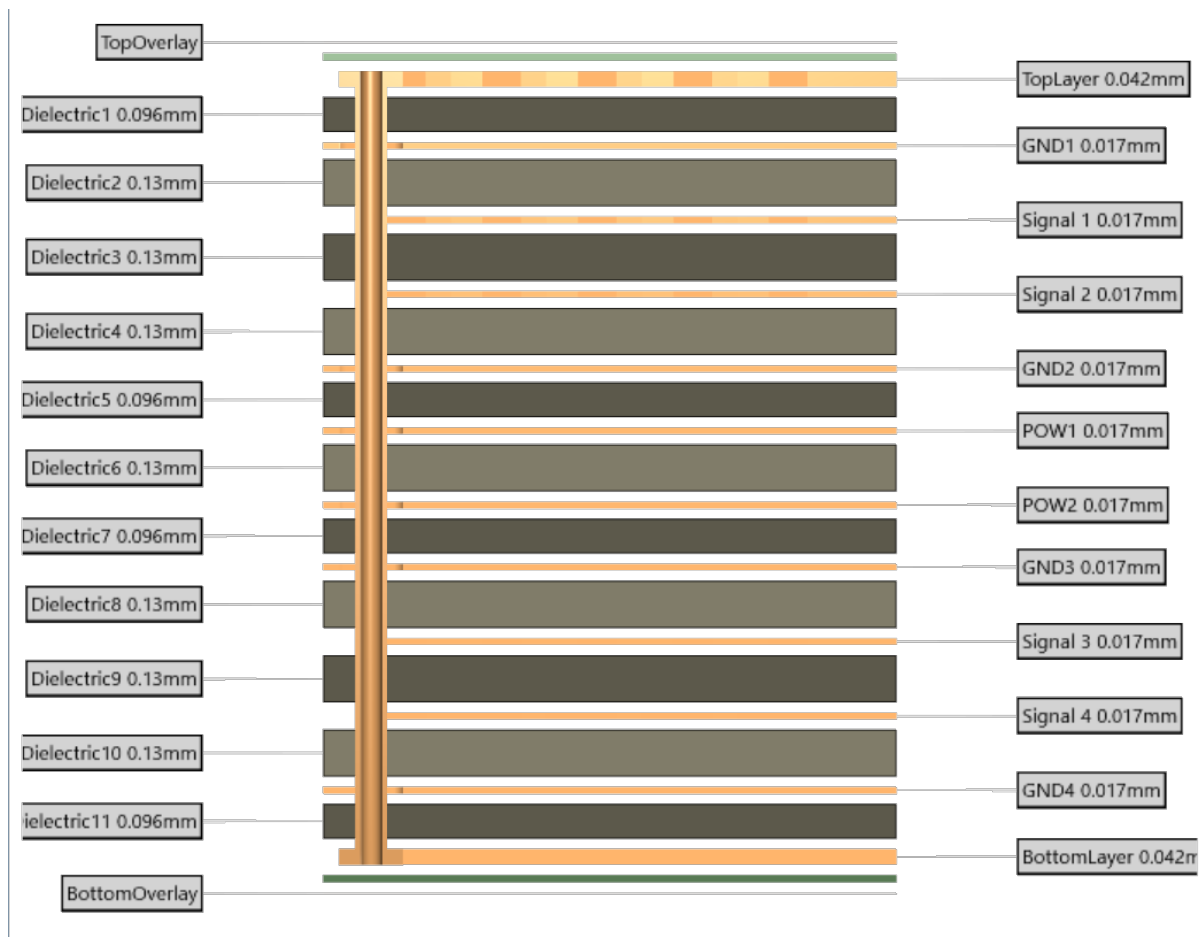


Figura 6.10. Stack-up del ejemplo de la Figura 6.9. Fuente propia

Podemos observar que, en esta estructura de 12 capas, las seis capas de señal (incluyendo *top* y *bottom*) tienen únicamente planos de masa como adyacentes. En la Figura 6.9, las vías junto a los cambios de capa del par diferencial son vías a masa y su objetivo es proporcionar continuidad a las corrientes de retorno. Fácil.

Imagina que intercambiamos GND4 y POW2, que es un plano de alimentación. ¿Cómo darías ahora continuidad a las corrientes de retorno?

No siempre podrás evitar tener planos de alimentación como planos de referencia. Por ejemplo, en PCIe (PCI express) el módulo tiene por especificación un espesor máximo de 62 mils (milésimas de pulgada), 1,57 mm, limitado por el conector en la placa base. Esto limita el número de capas internas. Disponer de tantos de planos de masa como en el stack-up de 1,8 mm del ejemplo anterior puede ser difícil, lo que obliga a usar planos de alimentación como referencia, con todos los problemas que conlleva.

Conviene (dentro de lo posible) alternar capas de señal y planos de masa para reducir crosstalk entre capas de señal

En la estructura de la Figura 6.10, las capas Signal1 y Signal2 son adyacentes, del mismo modo que Signal3 y Signal4. Signal1 tiene a 130 micras por encima un plano de masa (GND1), y 130 micras por debajo la capa Signal2. A su vez, ésta tiene 130 micras por debajo otro plano de masa (GND2). Como resultado, Signal1 tendrá más corriente de retorno por GND1 que por GND2, es decir, estará más acoplada a GND1 que a GND2.

La proporción entre corrientes de retorno por cada plano se calcula como presentamos en la sección “Corrientes de retorno en una stripline” en la página 53, con $h_1=130$ micras, $h_2=260$ micras. Resulta que una pista rutada en Signal1 tiene el 67% de la corriente de retorno en GND1 y el 33% en GND2. Si uno de los dos planos de referencia fuera de alimentación, tendríamos un serio problema con las corrientes de retorno cuando cambiáramos de capa.

Pero volviendo al punto que nos ocupa, la pequeña separación entre capas de señal (130 micras) nos obliga a evitar tener pistas superpuestas o muy cercanas en Signal1 y Signal2. Podemos optar por llevar cuidado o por un rutado a 90°. De haber podido insertar un plano de masa entre ambas capas no tendríamos esta restricción.

Conviene que un plano de alimentación sea adyacente a otro de masa para mejorar el desacoplo

Volviendo a la estructura de la Figura 6.10, ambos planos de alimentación (POW1, POW2) tienen adyacentes y a 96 micras de distancia un plano de masa. La capacidad entre cada plano y masa, asumiendo una constante dieléctrica de 4,2 es de:

$$C = \epsilon_0 \cdot \epsilon \frac{S}{d} = 8,85 \cdot \frac{10^{-12} F}{m} \cdot 4,2 \cdot \frac{1 \text{ cm}^2}{96 \text{ micras}} \approx 39 \text{ pF/cm}^2$$

Lo que no está nada mal, teniendo en cuenta que es una capacidad con muy baja inductancia y que por tanto será muy útil como parte de la red de condensadores de desacoplo en frecuencias altas (ya veremos qué es esto en la siguiente lección, Día 7).

Fíjate también en que los planos de alimentación son adyacentes y están cercanos, lo que provocará inyección de ruido por acoplamiento capacitivo entre ellos: ¡no hay soluciones perfectas!

Ejemplo de diseño

Queremos diseñar un stack-up para una placa de 8 capas, de las que 4 serán de señal. El diseño utiliza dos alimentaciones: +3.3V y +2.5V. El espesor final de la placa debe estar comprendido entre 1,6 y 1,8 mm. La impedancia de las líneas debe ser de 50 ohmios aproximadamente. Para los espesores disponibles de cobre, núcleos y pre-preg usa los indicados en la página 127. Recuerda que la estructura debe ser simétrica. Asumen un valor de 4,2 para la constante dieléctrica.

Bien, ¿por dónde comenzamos? Lo mejor para superar el miedo al papel en blanco es dibujar algo. Comencemos por dibujar la estructura de un PCB de 8 capas.

Paso 1: Dibujar la estructura

Sabemos del principio de la lección de hoy que las capas internas se fabrican con núcleos, esos sándwiches de cobre-FR4-cobre. Por tanto, dibujaremos tres núcleos. Entre núcleo y núcleo, pre-preg para adherirlos. Las capas externas sin láminas de cobre adheridas a los núcleos adyacentes mediante pre-preg. El resultado tiene el aspecto de la Figura 6.11.

¿Por qué es importante este paso? Porque por lo general los espesores que puedes conseguir con pre-preg y con núcleos son distintos. En una capa interna de rutado, estos espesores determinan las distancias a los planos de referencia y por tanto influyen mucho en la impedancia de línea. Saber si lo que separa L2 de L3, por ejemplo, es pre-preg o un núcleo es vital para el diseño. Te sorprenderías de cuántos errores se pueden cometer por no dar este primer paso: [un sencillo dibujo](#).

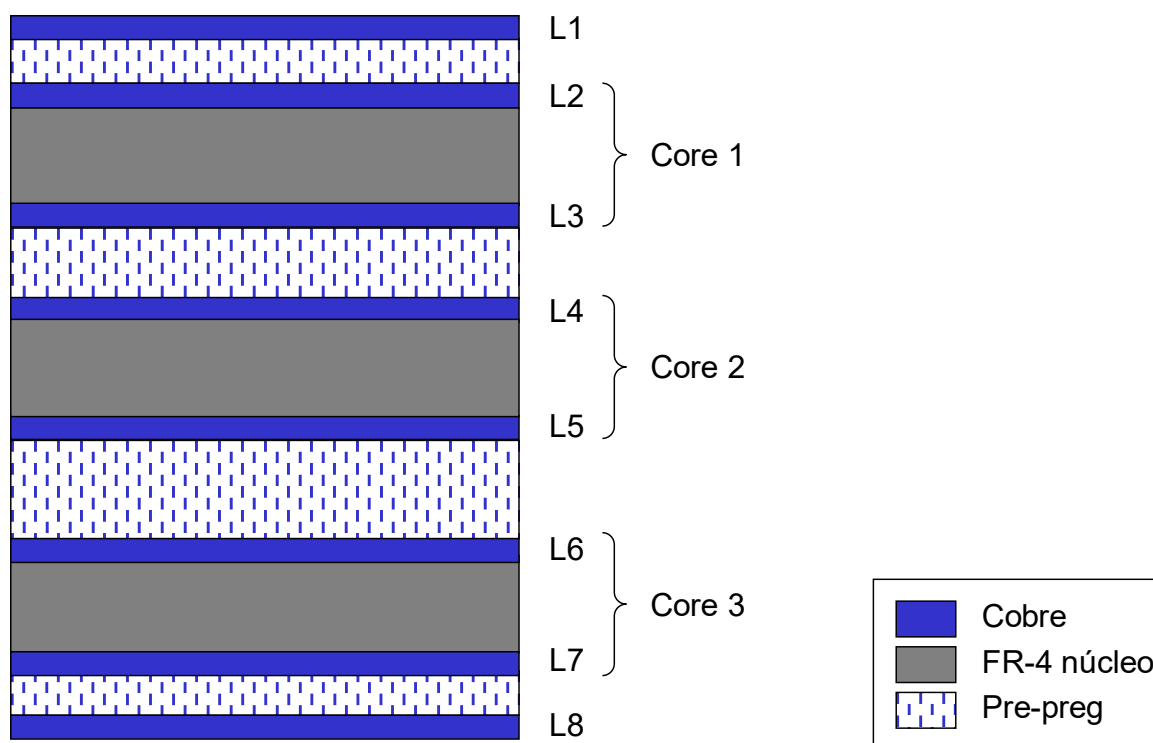


Figura 6.11. Un PCB de ocho capas consta de tres núcleos (*cores*), dos láminas de cobre para *top* (L1) y *bottom* (L8) y pre-preg para unir los diferentes elementos entre sí

Paso 2: Asignar funciones a las capas

Tenemos que asignar 4 capas de señal. Respecto a las alimentaciones, podríamos suponer que el diseño aceptará un *split-plane*, es decir, partir una capa en área de 2,5V y 3,3V. Pero vamos a asumir por razones didácticas que no es así y por tanto asignamos a cada alimentación una capa. Eso nos deja dos capas libres para masa. ¿Te parecen muchas? Ya verás que no.

Este paso no tiene solución única: habrá varias buenas, varias malas y al menos una excelente.

Aspiramos a buscar una solución de este último grupo. Las particularidades de cada diseño marcarán la diferencia y harán que una solución buena para un diseño sea excelente para otro. O al revés. Como el enunciado no nos dice nada sobre el diseño, vamos a hacer una discusión genérica que podrás usar en el futuro para tus propios diseños.

Capas externas

Vamos a comenzar suponiendo que las capas externas son de señal. Si tenemos en cuenta que una buena parte del área del PCB en capa *top* está ocupada por circuitos integrados y conectores, amén de área de cobre de disipación para reguladores, áreas de masa locales para osciladores y áreas de cobre de masa y de alimentación, resulta evidente que la capacidad de rutado en esta capa es limitada, aunque hay grandes variaciones de un diseño a otro. En cuanto a *bottom*, suele tener mayor capacidad de rutado, excepto en diseños que usen esta cara inferior para transferir calor a una plancha de aluminio o para planos de alimentación.

Però vam a suponer que L1 i L8 son capes de rutado. Ya sabes que senyals crítiques (que radien fuertemente o que sean muy susceptibles a ruidos externos) deberían ir rutadas por capes internas. Estas capes externas son ideales para líneas y buses lentos, así como para pares diferenciales de alta velocidad si necesitas reducir al máximo el número de vías por los que pasa la senyal.

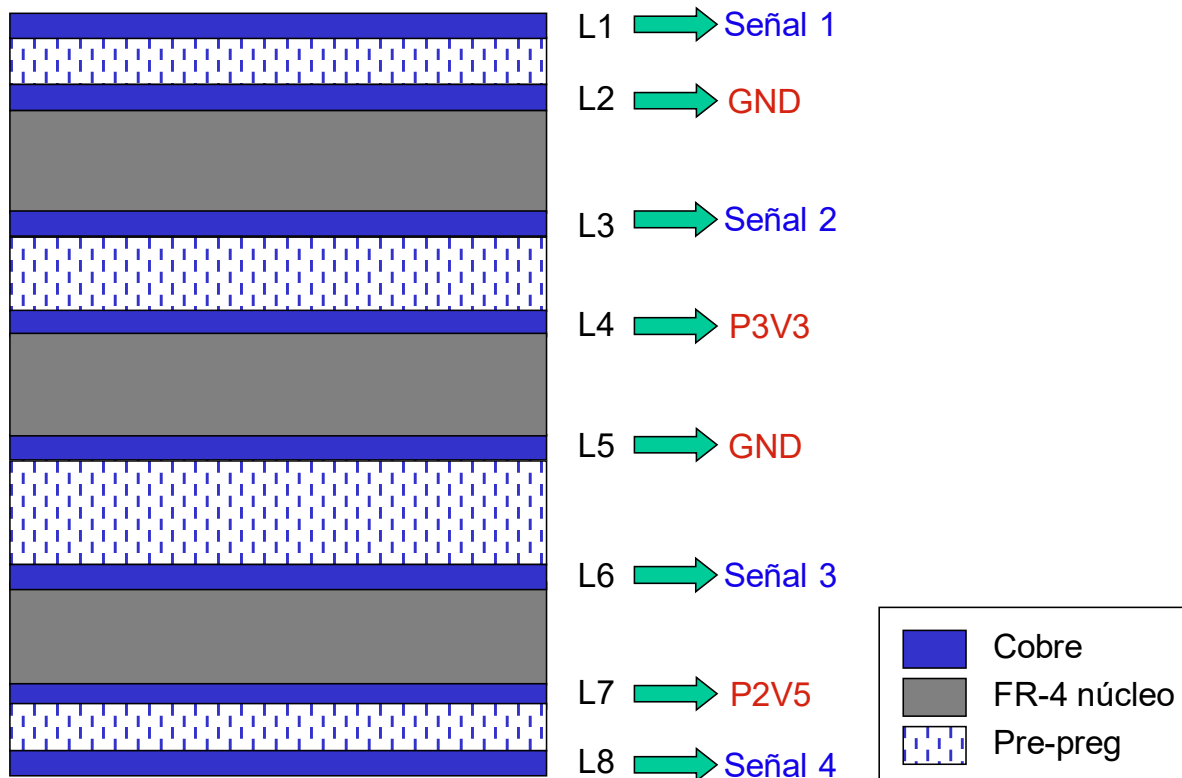


Figura 6.12. Ejemplo de asignación de función a las capas. No hay una solución única, sino un conjunto de soluciones malas, buenas y una o dos excelentes

Capas internas

Vamos a seguir la estrategia de evitar que haya dos capas de señal adyacentes. Veremos más adelante otro ejemplo donde tomamos la estrategia contraria, pero ahora mismo será más fácil asumir la primera. Evitar dos capas de señal contiguas elimina el problema de acoplamiento entre capas. Así que, bajo la capa de señal L1, ponemos L2 como plano de masa. L3 será capa de señal. Siguiendo este patrón (aunque no es necesario) L8 es señal, L7 será un plano de masa o alimentación, en este caso de 2,5V, y L6 será capa de señal. Nos queda por asignar L4 y L5. En la solución propuesta, L4 es P3V3 y L5 es masa.

Una aclaración sobre nomenclatura: posiblemente por compatibilidad con programas CAD, que no permitían comenzar el nombre de un nodo por una cifra ni contener caracteres especiales (tales como comas o puntos), es común hablar de “P2V5” para referirse a una alimentación de 2,5V. Por tanto, “P3V3” es 3,3V y podrás deducir qué es “P1V8”, “P1V”, etc.

¿Cómo sabemos si esta es una buena o una mala solución?

Debemos mirar las dos mitades de las señales: pistas y caminos de retorno. Como no hay capas de señal adyacentes, no hay acoplamiento entre ellas. Bien. Mirar los caminos de retorno marca la diferencia. Vamos allá:

L1 tiene como plano de referencia masa (L2). Si saltamos a L3, habrá que asegurar continuidad a la cara opuesta de L2 (lo que se puede conseguir simplemente a través del propio taladro de la vía de la señal). Pero parte de la corriente de retorno viajará también en L4 por el plano P3V3. ¿Cómo le damos continuidad de L4 a L2 cerca (muy cerca) de la vía por la que la señal pasa de L1 a L3? A bajas y medias frecuencias (hasta no más de 100 o 150 MHz) sólo podemos hacerlo mediante condensadores. En la Figura 6.13, un condensador (con sus vías a los planos de masa y P3V3) obliga a la corriente retorno a dar un pequeño rodeo, provocando una pequeña discontinuidad, pero tal vez aceptable si el condensador está muy cerca de la vía de señal. Por tanto, cerca de los circuitos integrados, donde hay condensadores, tenemos la posibilidad de dar continuidad a los retornos. Pero a mitad de placa, lejos de condensadores, esto no es posible.

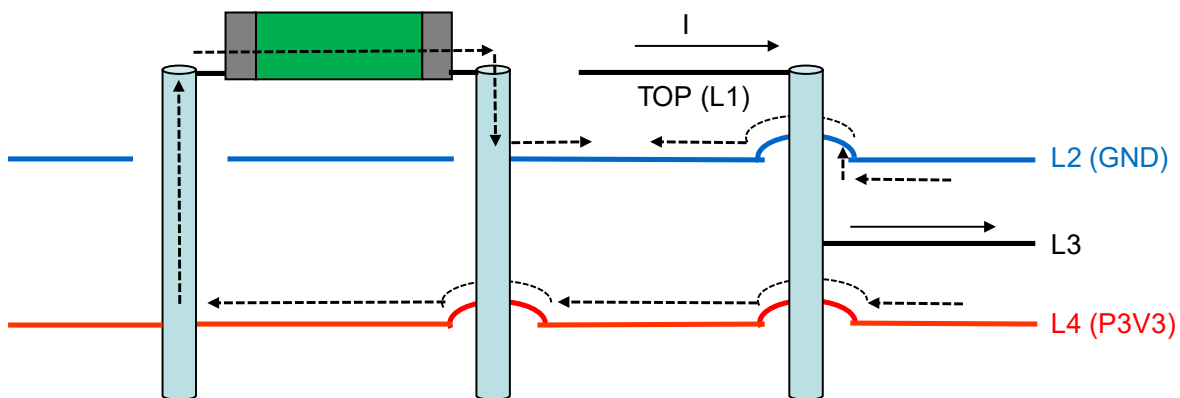


Figura 6.13. Camino de retorno de corrientes por los planos de referencia en el paso de señal entre L1 y L3. Fuente propia

Por lo tanto, a menos que nos limitemos a saltos entre L1 y L3 muy cerca de los integrados, tendremos problemas con señales en las que no podamos permitirnos una degradación significativa de la integridad de señal. Una entrada de un pulsador, una salida a una alarma, una línea de un bus de datos lento, pueden permitirse esta degradación. Una señal de reloj, una línea de un bus rápido, no.

El salto entre dos capas, las que sean, de la estructura de la Figura 6.12, será malo para señales críticas. **Pero...** siempre hay un pero. Hay esperanza: si el espesor del dieléctrico entre L2 y L3 es tres o cuatro veces menos que el espesor del dieléctrico entre L3 y L4, la corriente de retorno por el plano de masa será cuatro o cinco veces mayor que por P3V3 (ya hemos hablado de esto con anterioridad, recuerda), reduciendo mucho el problema que hemos mencionado en los saltos de señal entre L1 y L3. Pero el problema del salto entre L1 y L6 o L8 persiste.

De hecho, esto es importante, el espesor del dieléctrico entre L2 y L3 debe ser el mismo que entre L6 y L7. Porque debe haber simetría en la sección del PCB: de lo contrario, el PCB se combaría.

En el caso de espesor L2-L3 pequeño comparado con L3-L4, y por tanto L6-L7 pequeño comparado con L5-L6, nos encontramos con que entre L1 y L3 se puede saltar. Y entre L6 y L8 también. Pero las señales críticas no podrían salir de uno de estos pares.

Mi conclusión: este stack-up no es una buena solución. ¿Una alternativa? Mira la Figura 6.14. Cualquier salto entre L1, L3 y L4 no requiere más que añadir una vía de masa cerca de la vía de señal para dar continuidad a los caminos de retorno de corrientes (Figura 6.15).

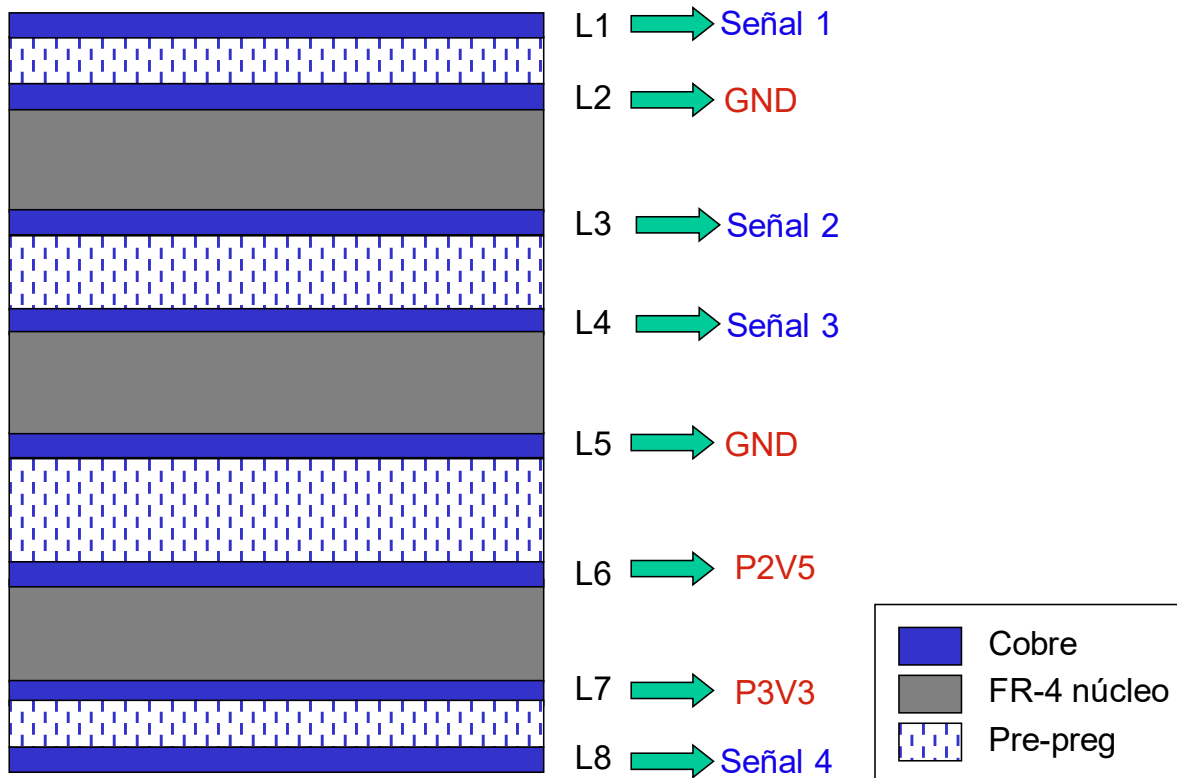


Figura 6.14. Stack-up alternativo

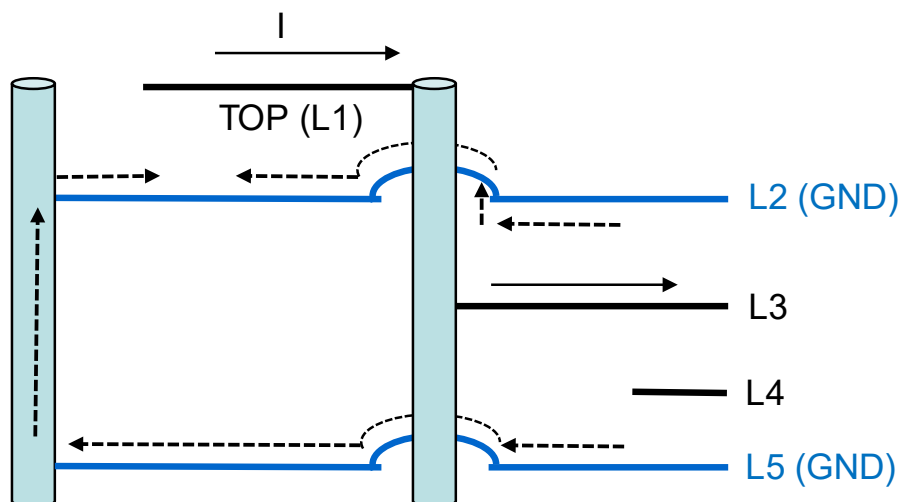


Figura 6.15. Retorno de corrientes a través de una vía a masa cerca de la vía de señal. Fuente propia

Si además consigo que el espesor L3-L4 sea al menos 3 veces mayor que el L2-L3 y L4-L5, minimizo el acoplamiento entre las capas de señal L3 y L4 y podré rutar libremente por estas capas. En caso contrario, tendré que evitar solapar pistas en L3 y L4. Por ejemplo, rutando en direcciones perpendiculares en cada capa. Dispongo así de tres capas para rutar señales en las que no quiero perjudicar la integridad de señal. La capa L4 queda para rutar señales lentas donde puedo permitirme degradación. Es un mejor stack-up.

¿Hay más alternativas? Claro. Me gusta más el stack-up de la Figura 6.16. Hemos dividido L6 en zonas de P2V5 y P3V3 (lo que se denomina *split plane*). En general, es posible hacerlo, y nos aporta la ventaja de que L8 también tiene como único plano de referencia masa y por tanto podemos realizar saltos entre las cuatro capas de rutado.

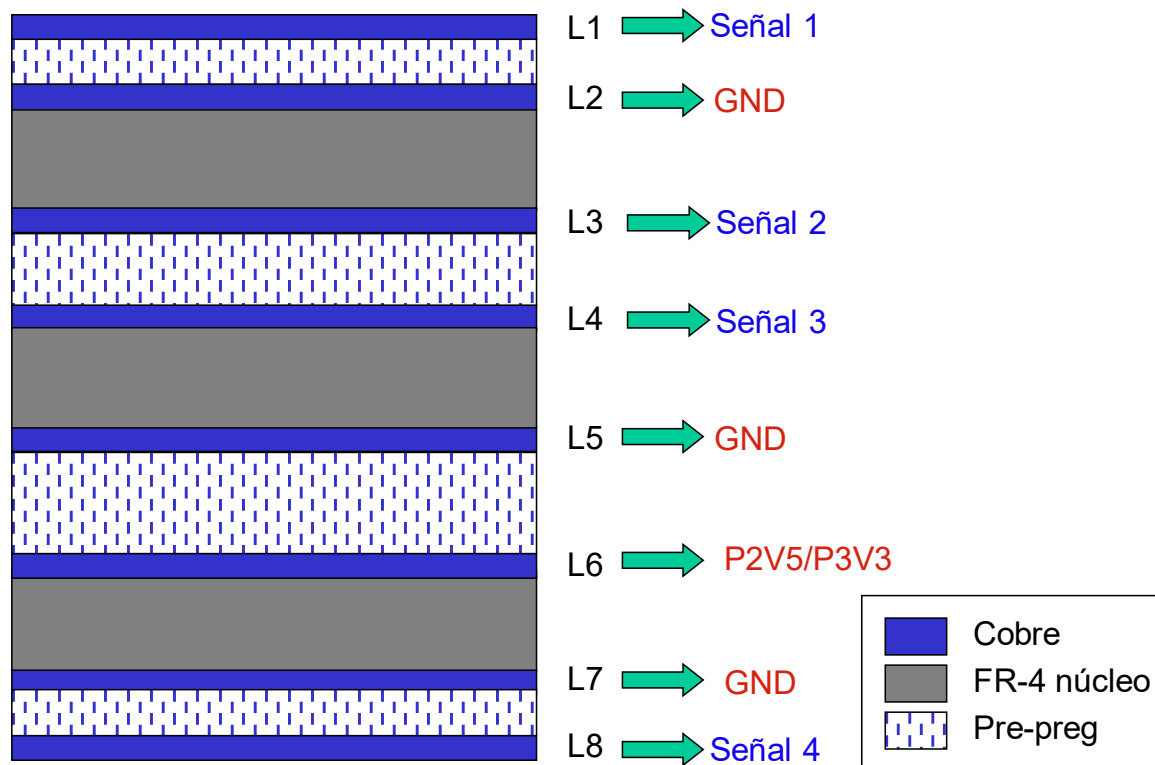


Figura 6.16. Un stack-up mejorado, empleando la técnica de split-plane

Añadimos un par de criterios más para evaluar la bondad de un stack-up

Vamos a mirar un aspecto más: **la proximidad entre planos de masa y de alimentación**. Esto es importante, porque los condensadores de desacoplo son efectivos sólo hasta tal vez 150 MHz. Por encima de esta frecuencia, es la capacidad entre planos la que aporta baja impedancias entre alimentación y masa. Ya sabes, $C = \epsilon \cdot \frac{S}{d}$, de modo que, a menor distancia entre planos, mayor capacidad. Es posible alcanzar 40-50 pF/cm² de capacidad con muy baja inductancia parásita.

También es relevante **la cercanía de los planos de masa y de alimentación a la(s) capa(s) donde se monten los componentes**. Porque minimizará la inductancia de las conexiones (por la menor distancia recorrida a través de las vías) y resultará en una impedancia menor a frecuencias altas.

Puedes evaluar la bondad de las tres propuestas anteriores respecto a estos dos criterios. Y llegarás a la conclusión de que darle la vuelta al último stack-up mejora las cosas (Figura 6.17).

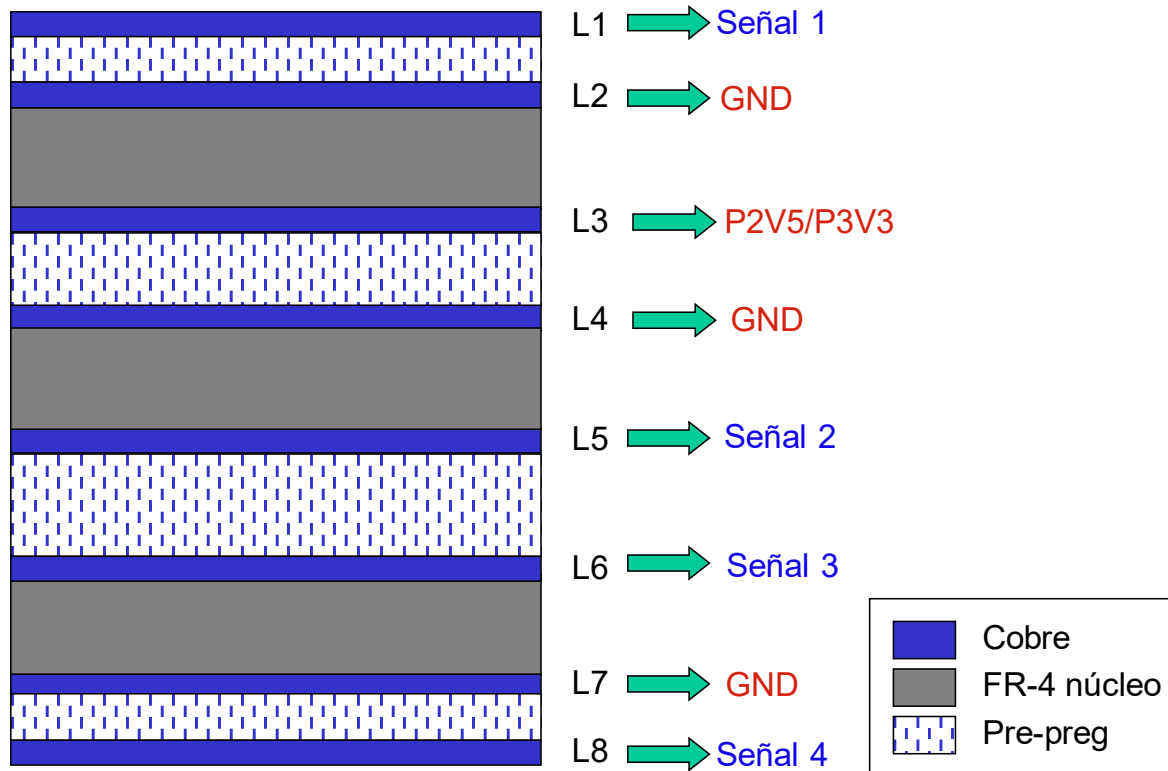


Figura 6.17. Stack-up mejorado

Paso 3: Diseñar las capas externas

Vamos a asumir el stack-up de la Figura 6.17. Comenzamos por diseñar las capas externas. **¿Por qué?** Porque tienes sólo dos grados de libertad para buscar 50 ohmios: la anchura de pista y la altura sobre el plano de referencia. En las capas internas tienes tres grados de libertad (se suma una distancia adicional al segundo plano de referencia).

En este momento te pediré buscar “**Saturn PCB toolkit**” en un navegador web. Regístrate, descarga e instala el programa. Y diseña conmigo.

Bien, estos son los materiales con los que vamos a trabajar en este ejercicio, pero debes preguntar a tu fabricante de PCBs qué espesores tiene disponibles:

- Espesores de cobre base: 17, 35 y 70 micras. A veces expresado en onzas (1 oz equivale a 34,3 micras)
- Grosor de FR4 en núcleos: 130, 180, 240, 350, 500, 510 y 730 micras
- Grosos de pre-preg: 106 (48mm), 1080 (65mm), 2113 (88mm), 7628 (175mm). Ya sabes que puedes usar combinaciones lineales de estos valores

Vamos a jugar un poco con la pestaña “Conductor Impedance”, eligiendo “Microstrip” como tipo de línea. Elegimos cobre base de 17 micras, al que habrá que sumar otras tantas por metalización de vías (*plating*). Elige 4,2 para el valor de la constante dieléctrica (es un valor estándar, pero debes usar el del material base que elijas a las frecuencias altas de la señal). Con dos capas de pre-preg 106 (96 micras de espesor) obtengo aproximadamente 50 ohmios con una pista de 170 micras de anchura (Figura 6.18).

El programa no contempla la máscara de soldaduras, lo que puede suponer un error en el orden de un ohmio. Pero vamos a aceptar este resultado.

Figura 6.18. Diseño de las capas externas: con dos pre-preg tipo 106 (96 micras) y una anchura de pista de 170 micras logramos los 50 ohmios. Captura de pantalla

Paso 4: Diseñar las capas internas

Elegimos “Stripline Asym” (stripline asimétrica, es decir, con distancias diferentes a los dos planos de referencia) como línea de transmisión. Señal 2 y Señal 3 tienen el espesor de un núcleo hasta el plano de masa. Y ambas capas están separada por pre-preg. Nos interesa un acoplamiento fuerte al plano adyacente de masa, y un acoplamiento menor (mayor distancia) hasta la otra capa de señal.

Fíjate en que la distancia hasta el segundo plano de referencia es igual al espesor de un núcleo más un pre-preg.

Con un núcleo fino de 130 micras y un pre-preg formado por dos capas 7628 (350 micras en total), obtenemos, para una anchura de pista de 150 micras, los 50 ohmios buscados (Figura 6.19).

Paso 5: Comprobaciones finales

Ahora hay que sumar el espesor de todas las capas, asumiendo 20 micras de máscara de soldaduras tanto en top como en bottom. El resultado (Figura 6.20) es de 1,5 mm, inferior al requerido en el enunciado. ¿Cómo lo aumentamos?

Si añadimos dos capas 2113 a las dos capas 1080 en los pre-preg internos, el espesor del PCB aumenta en 260 micras, subiendo hasta unos reconfortantes 1,76 mm.

El impacto de este cambio en la impedancia de las capas internas está en torno a 0,6 ohmios y se queda en 50,6 ohmios. No haría falta ni cambiar la anchura de pista. Y obtenemos un beneficio adicional: como la distancia de cada capa interna de señal a masa es de 130 micras, y la distancia entre capas de rutado es de 480 micras, tendré un acoplamiento fuerte a masa y un acoplamiento débil a la capa de señal adyacente, reduciendo el crosstalk entre ellas.

Conductor Impedance

Conductor Width (W)
0,15 mm

Conductor Height (H1)
0,130 mm

Conductor Height (H2)
0,48 mm

Narrow calculation mode

Diagrama de la estructura de capas:

Z_0
49.949 Ohms
 L_0
3.4155 nH/cm
 C_0
1.3690 pF/cm
 T_{pd}
68.3802 ps/cm

Options

Base Copper Weight

- ☐ 9um
- ☒ 18um
- ☐ 35um
- ☐ 53um
- ☐ 70um
- ☐ 88um
- ☐ 106um
- ☐ 142um
- ☐ 178um

Plating Thickness

- ☐ Bare PCB
- ☒ 18um
- ☐ 35um
- ☐ 53um
- ☐ 70um
- ☐ 88um
- ☐ 106um

Passive Circuits

- ☐ Microstrip
- ☐ Microstrip Embed
- ☐ Stripline
- ☒ Stripline Asym
- ☐ Dual Stripline
- ☐ Coplanar Wave

Units

- ☐ Imperial
- ☒ Metric

Substrate Options

Material Selection
Custom

Er
4,2

Tg (°C)
130

Temp Rise (°C)
20

Temp in (°F) = 36.0

Ambient Temp (°C)
22

Temp in (°F) = 71.6

Information

Total Copper Thickness
18 um

Via Thermal Resistance
179.3 °C/W

Via Count: **10**

Conductor Temperature
Temp in (°C) = N/A

Via Voltage Drop
17.9 °C/W per via

Temp in (°F) = N/A

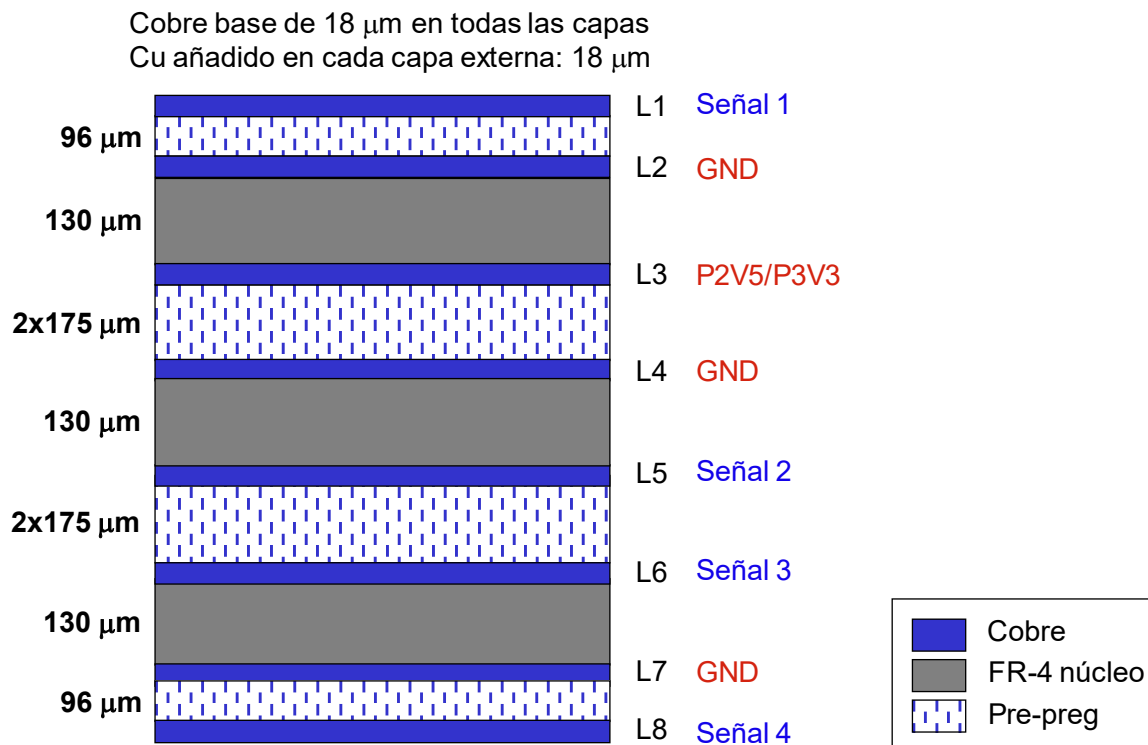
N/A

SATURN PCB DESIGN, INC.
Turnkey Electronic Engineering Solutions

Follow Us

f t in y

Figura 6.19. Diseño de las capas internas. Un núcleo de 130 micras (H1), dos capas de pre-preg 7628 (350 micras) y un núcleo de 130 micras (H2, 480 micras) y una anchura de pista de 140 micras nos dan los 50 ohmios que buscábamos. Captura de pantalla



Espesor total del PCB (con 2x20 μm de solder mask): 1,5 mm

Figura 6.20. El espesor final es de 1,5 mm. Es inferior al requerido (1,6-1,8 mm). ¿Cómo lo aumentamos?

Differential Pairs

Conductor Width (W) mm

Conductor Spacing (S) mm

Conductor Height (H) mm

W/H = 1.462
S/H = 2.692

Target Zdiff Ohms

Formula Restrictions:
0.1 < W/H < 3.0
0.1 < S/H < 3.0

Zdiff Differential Ohms

Zo Ohms

+/- Tolerance = 10%

Ohms

Ohms

Options

Base Copper Weight
☐ 9um
☒ 18um
☐ 35um
☐ 53um
☐ 70um
☐ 88um
☐ 106um
☐ 142um
☐ 178um

Units
☐ Imperial
☒ Metric

Substrate Options
 Material Selection

Er Tg (°C)

Temp Rise (°C)

Temp in (°F) = 36.0

Ambient Temp (°C)

Temp in (°F) = 71.6

Plating Thickness
☐ Bare PCB
☒ 18um
☐ 35um
☐ 53um
☐ 70um
☐ 88um
☐ 106um

Differential Layer
☒ Edge CpId Ext
☐ Edge CpId Int Sym
☐ Edge CpId Int Asym
☐ Edge CpId Embed
☐ Broad CpId Shld
☐ Broad CpId NShld

Information
 Total Copper Thickness 36 um
 Via Thermal Resistance N/A
 Via Count:

Conductor Temperature
 Temp in (°C) = N/A
 Temp in (°F) = N/A

Via Voltage Drop
 N/A

SATURN
PCB DESIGN, INC
Turnkey Electronic Engineering Solutions

Follow Us

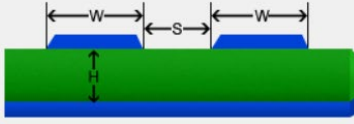


Figura 6.22. Diseño de pares diferenciales en capa externa, 100 ohmios. Captura de pantalla

Para terminar, unas lecturas recomendadas

No te hará ningún daño echar un vistazo a los siguientes artículos que puedes encontrar en Internet: [18] [19] [1] [20]. Las ideas son las que hemos trabajado hoy, pero leer a un autor distinto puede ayudarte a entender algún punto que no te haya quedado claro.

Día 7. Diseño de redes de desacoplo

Los términos “red de condensadores de desacoplo” o “red de desacoplo” son desafortunados. Porque no aclaran nada. Y porque son varias las funciones que cumplen estas redes en un PCB.

El diseño de una red de desacoplo consiste en determinar cuántos condensadores, de qué valor, tipo y tamaño debemos añadir al diseño, así como dónde colocarlos, para lograr una baja impedancia (típicamente del orden de unos miliohmios a cientos de miliohmios) entre una tensión de alimentación y masa.

¿Por qué queremos baja impedancia entre una alimentación y masa? Por los siguientes motivos: para mantener una alimentación estable pese a las demandas pulsadas y abruptas de los circuitos integrados digitales, para cortocircuitar y por tanto atenuar el ruido que se propaga por las guías biplaca y para proporcionar caminos para las corrientes de retorno entre masa y alimentación.

Todos los circuitos integrados del PCB deben “ver” una baja impedancia entre alimentación y masa. Esto requiere en unos casos un simple condensador de 100 nF cerca del integrado (realmente el mayor valor que puedas alcanzar en el encapsulado más pequeño que puedas permitirte), decenas de condensadores en otros casos.

Los condensadores reales presentarán baja impedancia sólo en un margen de frecuencias. Para cubrir todo el espectro en frecuencia de los pulsos de corriente, debemos establecer un diseño jerárquico con condensadores que cubran los rangos de bajas, medias y altas frecuencias. No obstante, por encima de 150 MHz será prácticamente imposible obtener baja impedancia con condensadores. La capacidad entre planos en el PCB, si está diseñada adecuadamente, complementará la red de desacoplo a partir de aproximadamente 200 MHz. No obstante, habrá que hacer un estudio de hasta qué frecuencia debemos proporcionar baja impedancia: los encapsulados de los circuitos integrados introducirán una barrera o límite por encima de la cual nuestros esfuerzos no serán rentables.

Todas las alimentaciones en el PCB deben tener una red de desacoplo asociada, de la que forman parte los condensadores a la entrada y a la salida de los reguladores lineales y conmutados.

El diseño de la red de desacoplo se puede realizar en primera aproximación con herramientas sencillas y de bajo o nulo coste. Lo importante es comprender los qué, cómo y por qué y derivar una metodología y buenas prácticas para obtener redes de desacoplo buenas. Una vez más, lo óptimo será enemigo de lo bueno.

Comenzaremos la lección de hoy exponiendo las funciones que cumple una red de desacoplo. A continuación, estudiaremos cómo se comporta un condensador real y qué tipos de condensadores debemos usar. En este momento presentaremos una metodología y estudiaremos un ejemplo.

En un diseño real partiremos tal vez de la red de desacoplo que incorpora un diseño de referencia, que no tiene por qué ser adecuada para nuestra aplicación. Copiarla ciegamente puede resultar en una red insuficiente (no funcional) o sobredimensionada (costosa). Cotejar esta red con la que hayamos diseñado nos ayudará a afinar la solución.

Por último, un aviso: el diseño de redes de desacoplo presenta mucha incertidumbre, lo que nos obligará a manejar márgenes de seguridad amplios y te dejará posiblemente con una sensación de insatisfacción. Hablaremos sobre esto más adelante.

Funciones de un condensador de desacoplo

Primera función: eliminar fluctuaciones de baja frecuencia en la alimentación

Hablemos de esos condensadores de elevado tamaño, cerca de la entrada de alimentación, a veces electrolíticos de aluminio, en diseños de menor tamaño cerámicos multicapa. Esos condensadores conforman la denominada **capacidad de bulk**. ¿Por qué crees que están ahí?

Colocamos los condensadores lo más cerca posible de la entrada de alimentación a la placa para independizar al diseño de la inductancia de los cables de alimentación e impedir la penetración de ruido de baja frecuencia en el diseño. Explicaremos esto con calma más adelante.

Estos condensadores de elevado valor y por tanto de encapsulados grandes (alta inductancia) sólo sirven para filtrar las perturbaciones de baja frecuencia, porque su inductancia es elevada y por tanto su impedancia es alta a alta frecuencia. Dicho de otro modo, proporcionan una fuente de carga con baja impedancia a frecuencias a las que la inductancia de los cables provocaría demasiada caída de tensión, anulando así el efecto negativo de la inductancia de los cables.

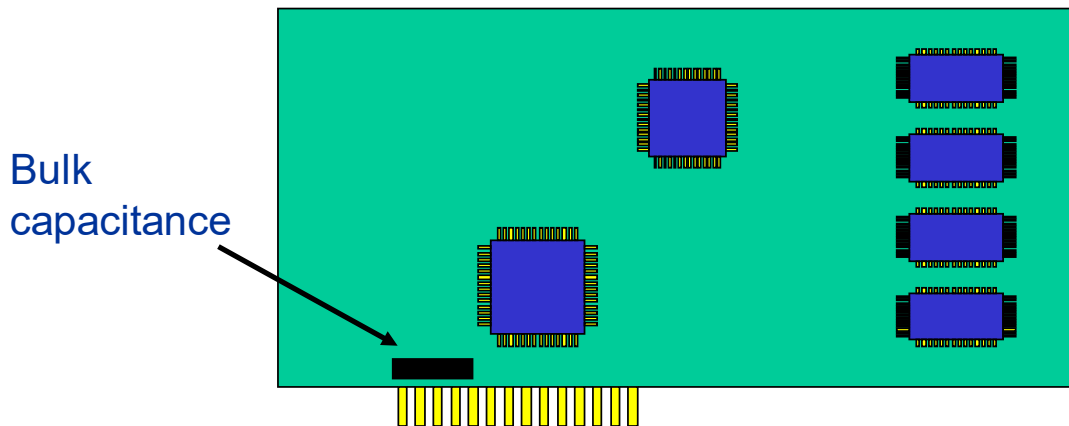


Figura 7.1. La capacidad de *bulk* (que se puede traducir como abultada) se coloca a la entrada de la alimentación, junto al conector. Fuente propia

Los cables de alimentación provocan una caída de tensión según la expresión $V \propto L \cdot \frac{\partial I}{\partial t}$, lo que puede ser inaceptable en muchos casos. La estrategia que se persigue con la **capacidad de bulk** (gruesa, de bulto) es proporcionar una fuente de carga con una inductancia mucho menor. En la Figura 7.2, L_2 (inductancia entre la capacidad de bulk y los circuitos integrados) es mucho menor que L_1 (inductancia de las conexiones de alimentación), dando así lugar a menores caídas de tensión de alimentación.

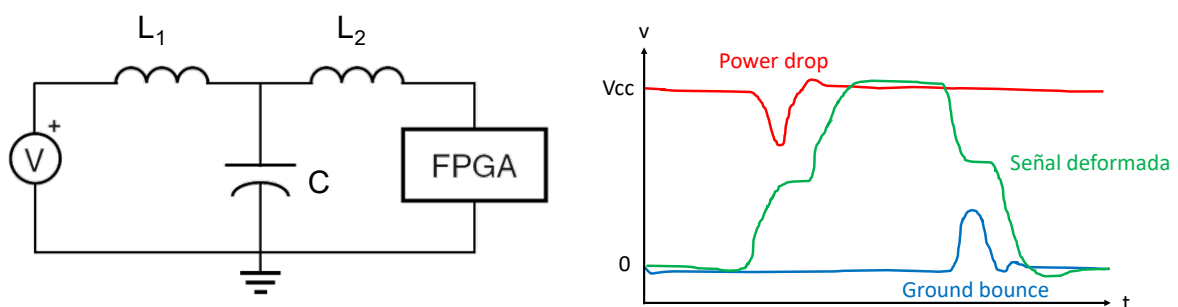


Figura 7.2. Concepto de la capacidad de bulk (izquierda), donde buscamos reducir la inductancia entre los circuitos integrados y la fuente de carga, dando así lugar a caídas de tensión y rebotes en las masas menores (derecha). Fuente propia

Los rebotes en la alimentación y en masa tiene su efecto en las señales generadas por los integrados, deformándolas y afectando así a la integridad de señal.

Conociendo la máxima ΔV que tolera el sistema (dato que extraeremos de las hojas de datos de los componentes) y la máxima ΔI (una estimación del pico de la demanda de corriente), determinamos la máxima reactancia admisible entre masa y alimentación:

$$X_{\max} = \Delta V / \Delta I$$

Determinamos la frecuencia a la que el cableado de alimentación presenta esta reactancia:

$$f_{ind-cable} = \frac{X_{\max}}{2 \cdot \pi \cdot L_{cable}}$$

La inductancia del cableado (en nH) se puede estimar como:

$$L_{cable} \approx 4 \cdot l \cdot \ln \frac{2H}{D}$$

Donde l es la longitud del cable en cm, H es la separación media entre cables y D el diámetro de los cables. H y D deben estar en las mismas unidades. Para garantizar una baja impedancia entre masa y alimentación por encima de $f_{ind-cable}$ necesitamos un condensador que calculemos para tener una reactancia $X_{\max}$ a $f_{ind-cable}$:

$$C_{bulk} = 1 / (2\pi \cdot f_{ind-cable} \cdot X_{\max})$$

Ejemplo numérico

Supongamos los siguientes parámetros para el cable: L_{cable} ($l=0,5m$, $H=5mm$, $D=2mm$) = 322 nH

Para un ΔI de 4A, con un 5% de tolerancia a 3,3V obtenemos que la reactancia debe ser menor que: 40 mΩ. El cable alcanza esta reactancia a 20 kHz.

Calculamos el valor del condensador que presenta 40 mΩ a 20 kHz: 200 μF. Estos valores sugieren el empleo de un condensador electrolítico de muy baja ESR (ya que la impedancia total será la suma de la ESR y la reactancia del condensador). Aplicando un margen de seguridad, escogemos el valor de 330 μF. Este margen de seguridad se justifica porque como ya sabes, el valor nominal de la capacidad disminuirá al polarizar el condensador, con el envejecimiento y la temperatura.

Kemet, uno de los grandes fabricantes de condensadores, ofrece en su web un simulador online (<http://ksim.kemet.com/>) de impedancia de sus condensadores. La Figura 7.3 muestra una simulación de un componente adecuado para nuestro ejemplo.

El margen útil de frecuencias se extiende hasta aproximadamente 2 MHz. Si añadimos un par de condensadores cerámicos multicapa (MLCC, *multi-layer ceramic capacitor*) de 10 μF y tamaño 1206, el margen útil se extiende hasta los 20 MHz (Figura 7.4), ¡ganando así una década!

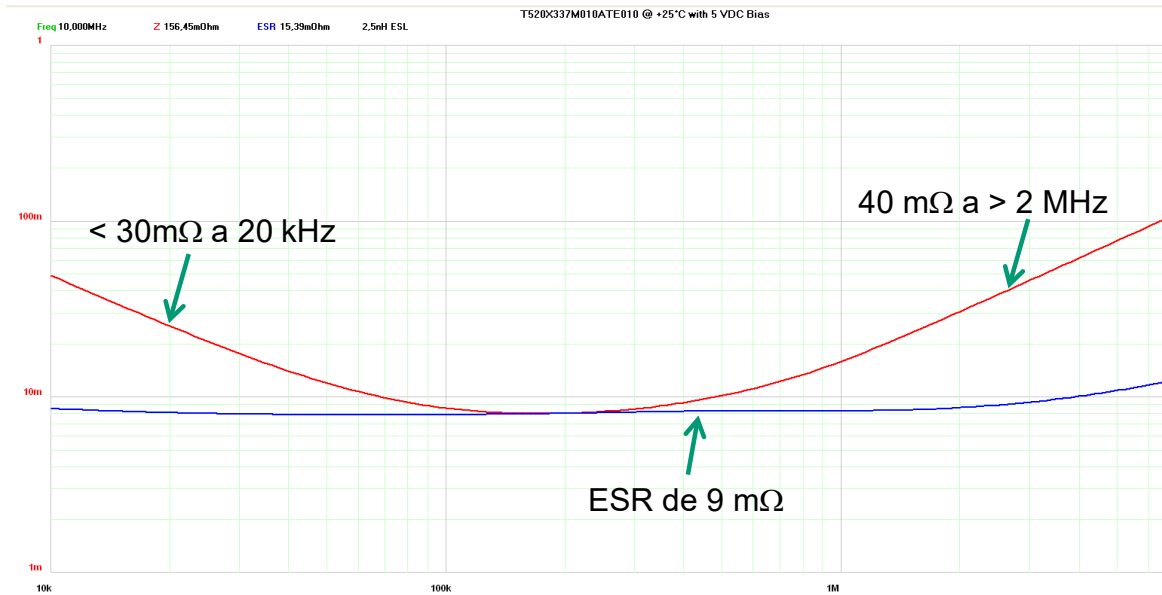


Figura 7.3. Simulación con Kemet Spice Software (330 µF SMD size X, conductive polymer). Fuente propia

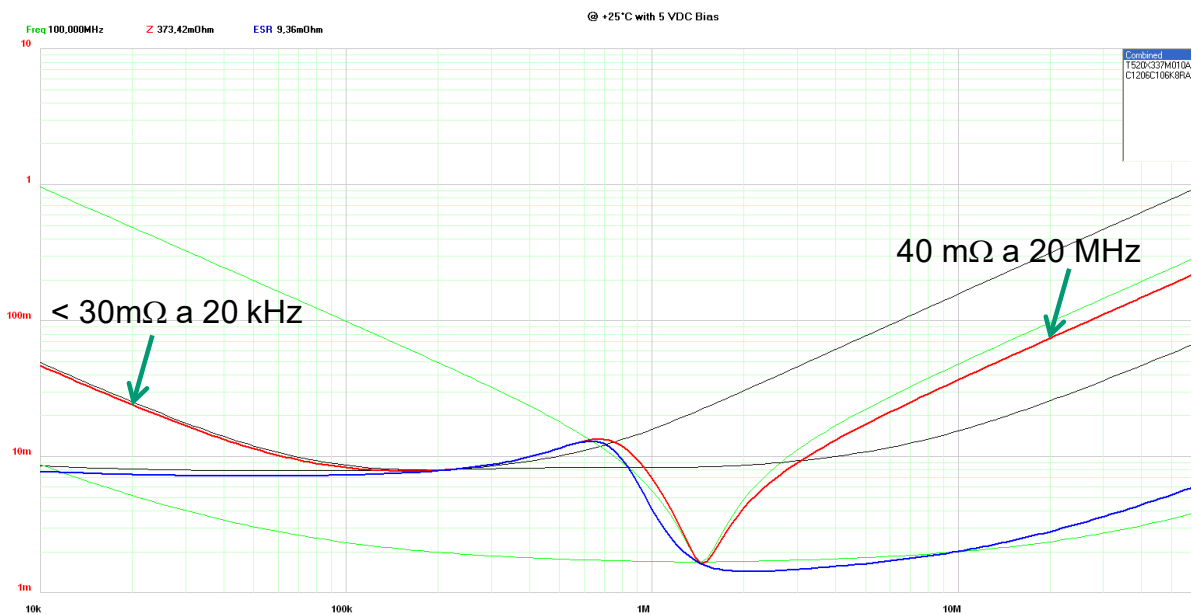


Figura 7.4. Extendemos el margen de frecuencias útil a 20 MHz añadiendo 2x 10 µF 1206 MLCC. Fuente propia

Segunda función: eliminar fluctuaciones de media frecuencia en la alimentación

La capacidad de *bulk*, por su tamaño y características, es lenta. Esto quiere decir que su elevada inductancia le impide entregar picos de carga rápidamente. Estos picos de demanda instantánea debe proporcionarlos un grupo de condensadores de menor tamaño (porque menor tamaño implica menor inductancia), colocados junto a los circuitos integrados (para reducir también la **inductancia total del lazo**) lo que constituye un segundo nivel jerárquico de la red de desacoplo.

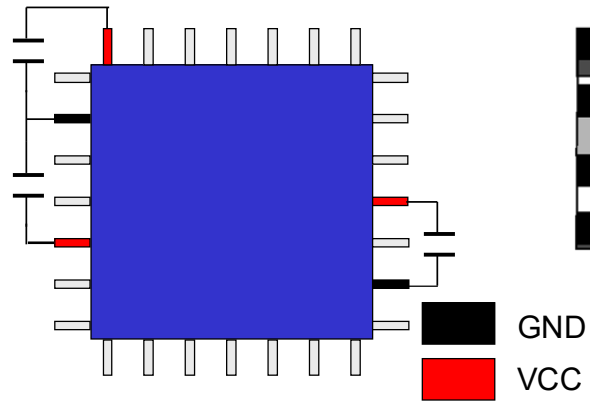


Figura 7.5. Las componentes de frecuencias medias (hasta tal vez 150 MHz) de los picos de corriente son proporcionadas por condensadores de pequeño tamaño junto a los circuitos integrados. Fuente propia

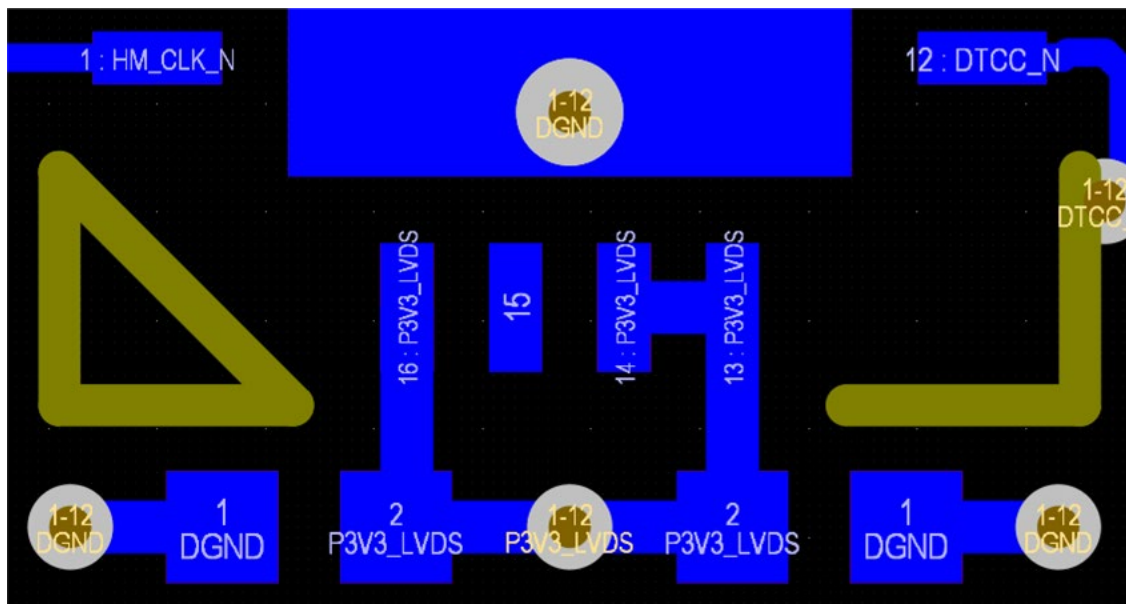


Figura 7.6. Detalle de la ubicación y conexión de dos condensadores de desacoplo tamaño 0402 entre P3V3_LVDS y DGND a los pines de alimentación de un circuito integrado. Sin haber una única forma de hacerlo, el criterio es claro: reducir la inductancia total del lazo tanto como sea posible. La conexión de los condensadores a los planos de masa y alimentación no es óptima y añade un poco de inductancia extra.

Es importante entender el concepto de **inductancia total del lazo**. Un condensador tiene una inductancia propia, que depende básicamente de sus características geométricas (sobre todo longitud y anchura). Cuando se monta en el PCB, la inductancia de las pequeñas pistas y de las vías hasta los planos de alimentación y masa (denominada **inductancia de montaje**) se suma a la inductancia de los planos en el bucle (*capacitor connection inductance loop*, bucle 1 en la Figura 7.7).

La distancia entre el condensador y el circuito integrado implica también una pequeña inductancia (*plane inductance loop*, bucle 2 en la figura), que puede ser del orden de 0,5-1 nH. Finalmente, el encapsulado y las conexiones a los planos de alimentación (*IC connection inductance loop*, bucle 3 en la figura) completa la lista de elementos que conforma la inductancia total del lazo.

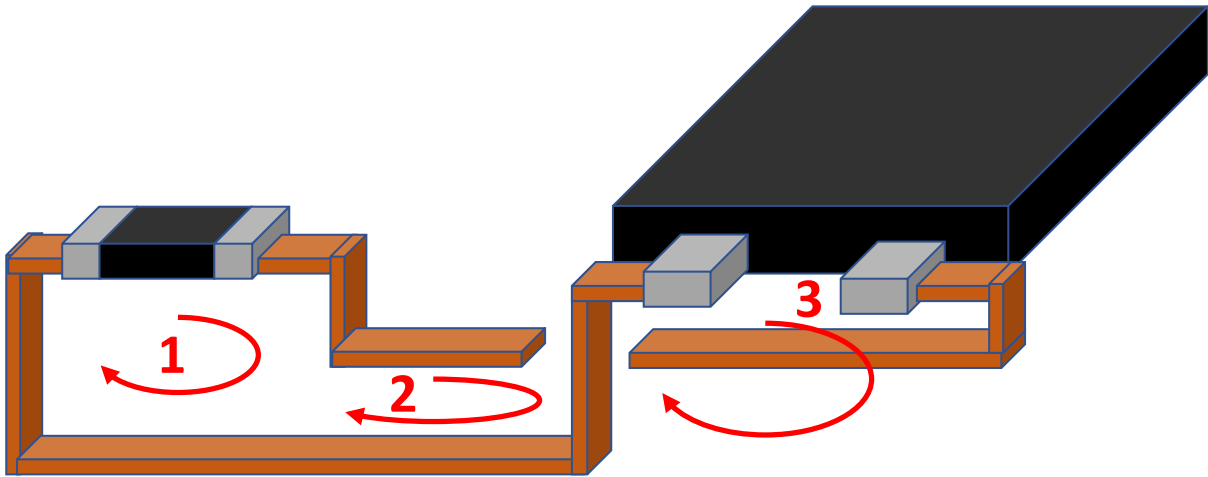


Figura 7.7. La inductancia total del lazo está formada por la del condensador de desacople y sus conexiones a los planos, la del circuito integrado y sus conexiones a los planos y la de los planos en la distancia que separa ambos componentes. Fuente propia

Tu objetivo es minimizar la suma de la inductancia de estos tres bucles. Ahí va una lista de cosas que puedes hacer:

- Escoge el encapsulado más pequeño que puedas permitirte. En función del valor de la capacidad y de su *rating* de tensión (te recomiendo escoger al menos un valor doble del nominal), así de la capacidad del montador de trabajar con encapsulados muy pequeños (como 0201), es posible que no puedas bajar de un 0603 o de un 0402 (mira la Figura 7.12).
- Intenta conectar los condensadores a los planos con baja inductancia (mira la Figura 7.14). Esto será especialmente importante en integrados que requieran desacople a alta frecuencia
- Intenta colocar el condensador tan cerca como sea posible del circuito integrado, reduciendo así la inductancia de los planos
- Intenta colocar planos de masa y de alimentación cerca de la capa top, con el fin de reducir la distancia a recorrer a lo largo de las vías y por tanto su inductancia

Decirlo es más fácil que hacerlo, porque son muchos los criterios que debes satisfacer y la red de desacople, como decía Neruda, “...es simplemente una ola alta sobre las olas”. Es decir, tendrás que alcanzar soluciones de compromiso entre muchos requisitos y, por ejemplo, ubicar el plano de alimentación cerca de la capa de componentes tal vez choque con otras necesidades.

Tercera función: proporcionar un camino de baja inductancia para las corrientes de retorno

Ya hemos hablado de esta función en días anteriores. Cuando la continuidad de las corrientes de retorno requiera un camino de baja impedancia entre una alimentación y masa, el uso de condensadores es una opción aceptable hasta tal vez 100-150 MHz. Por tanto, señales con muy alto ancho de banda deberán evitar esta problemática.

La red de condensadores de desacople, que agrupa estos componentes en las inmediaciones de los circuitos integrados, permite saltos de capa de señal que impliquen cambio de plano de referencia entre una alimentación y más siempre que se cumplan tres condiciones:

- Que el cambio de capa de señal ocurra en las inmediaciones de los circuitos integrados
- Que ambos circuitos dispongan de redes de desacople entre la alimentación en concreto y masa
- Que el ancho de banda de la red de desacople alcance la frecuencia de codo de las señales

De no cumplirse lo anterior, la red de desacople no estaría cumpliendo con la función y estaríamos ante un problema de integridad de señal y potencialmente de EMC.

Cuarta función: reducir la propagación de ruido

La Figura 7.8 muestra una estructura formada por un plano de alimentación sobre otro de masa al que se le ha dado una forma irregular (dos recortes en forma de rectángulo, en blanco, Figura 7.8 arriba a la izquierda). Una fuente de ruido (un par de pines de alimentación y masa de un circuito integrado digital que demandan corriente de alimentación de forma pulsada) excita la guía biplaca, lo que debería inundar de ruido todo el PCB. La presencia de un pequeño grupo de condensadores de desacoplo (Figura 7.8 arriba, derecha) crea varios cortocircuitos para el ruido, atenuando el nivel de perturbación que alcanza el extremo opuesto del PCB (Figura 7.8 abajo).

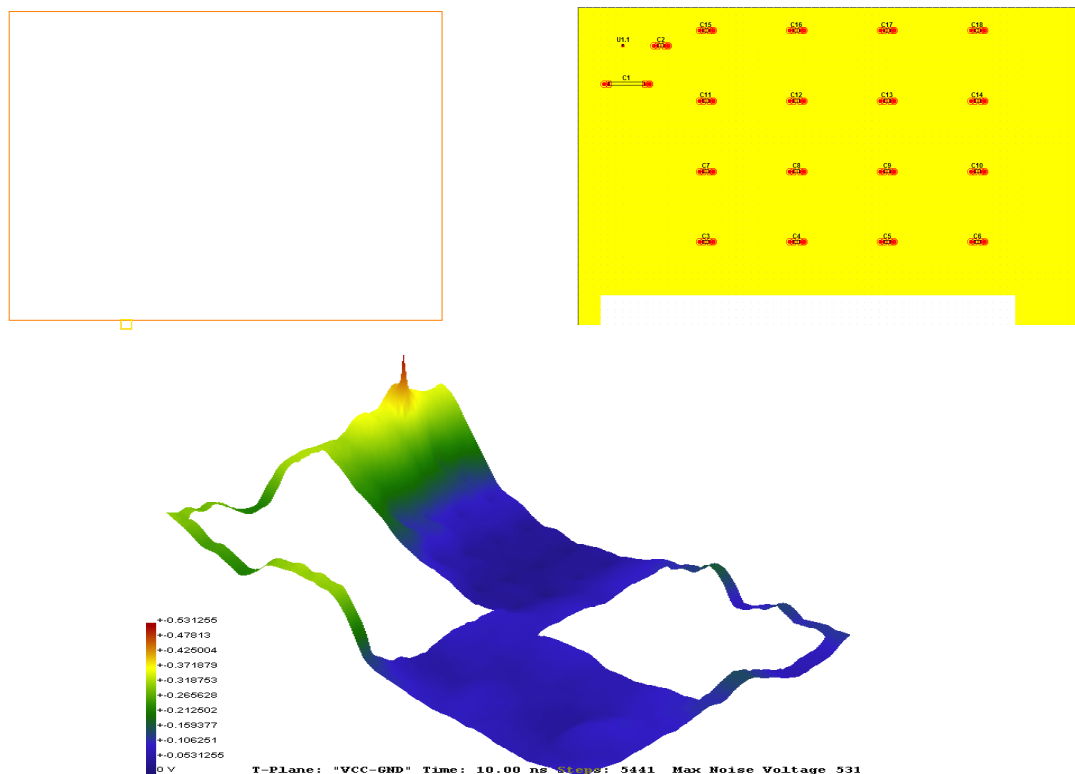


Figura 7.8. Reducción de la propagación de ruido en una guía biplaca mediante condensadores de desacoplo. Fuente propia

Comportamiento en frecuencia de un condensador real

Un condensador real presenta, además de una capacidad eléctrica, una pequeña resistencia serie equivalente (ESR) y una inductancia serie (ESL), conformando un circuito resonante serie. La ESR es del orden de miliohmios y la inductancia parásita de 1 a 5 nH. Una representación de la impedancia de este circuito resonante en frecuencia (Figura 7.9) nos permite entender por qué un condensador sólo nos sirve a efectos de desacoplo en el **margen útil de frecuencias** en el que su impedancia esté por debajo de un valor dado.

Este margen útil se extiende más allá de la resonancia, cuando el componente se comporta ya como una inductancia, porque lo que nos interesa es el valor de la impedancia y no el signo de la reactancia.

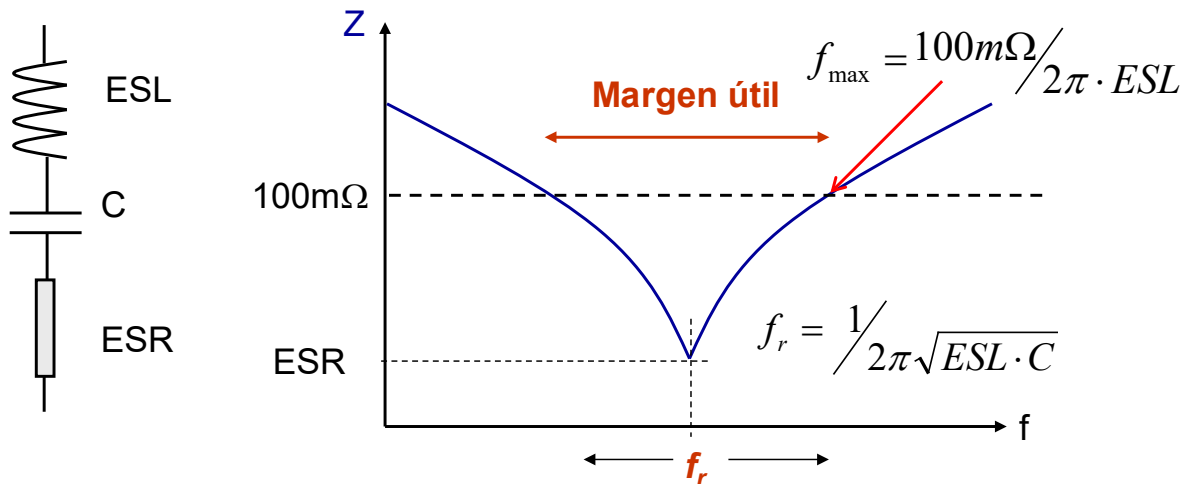


Figura 7.9. Comportamiento en frecuencia de un condensador real. Fuente propia

Extensión del ancho de banda útil mediante combinación de condensadores

Una primera estrategia consiste en poner dos condensadores de diferente valor en paralelo. Hay que elegir los condensadores para que la curva de impedancia equivalente no contenga zonas por encima de la impedancia objetivo y, sobre todo, evitando que haya **picos de anti-resonancia**.

En una **anti-resonancia** la impedancia equivalente es muy alta, y se debe a que una inductancia y una capacidad de dos circuitos serie RLC en paralelo resuenan. Por ejemplo, en la Figura 7.10, la inductancia del condensador de 22 nF (que domina por encima de la frecuencia de resonancia de la curva izquierda) resuena con la capacidad del condensador de 100 pF (curva derecha). **Si algún circuito integrado excita esta frecuencia, tendremos un ruido elevado.** Realmente el cálculo de la frecuencia de anti-resonancia es algo más complejo. Ocurre a $f = \frac{1}{2\pi\sqrt{LC}}$, donde L es la suma de las inductancias de los dos condensadores y C es el equivalente serie de los dos condensadores.

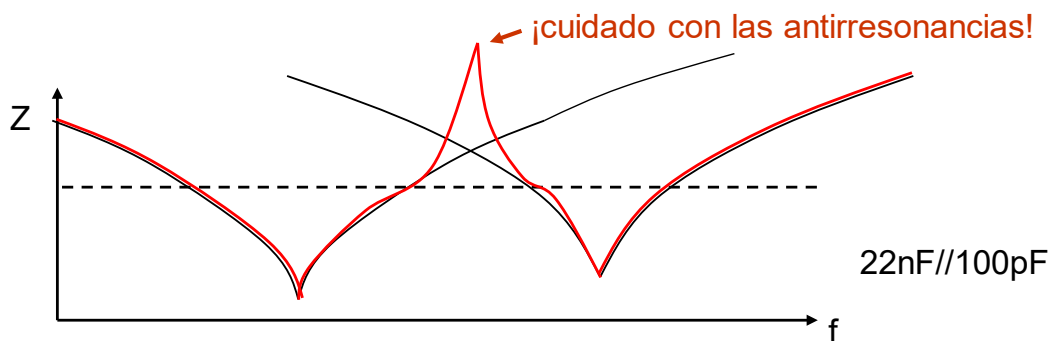


Figura 7.10. Extensión del margen útil de frecuencia poniendo dos condensadores distintos en paralelo. Fuente propia

Para evitar el problema anterior, también se puede optar por usar un único valor de capacidad, como es el caso de la Figura 7.11. Con dos condensadores idénticos en paralelo, baja la curva equivalente (la ESR equivalente es la mitad) y aumenta el margen útil.

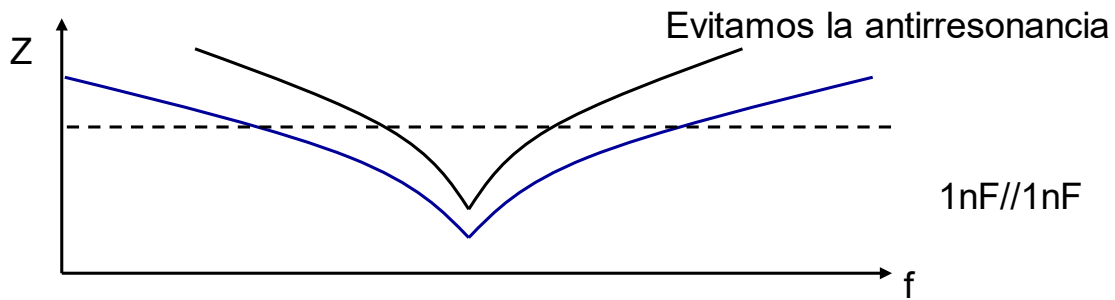


Figura 7.11. Extensión del margen útil de frecuencia poniendo dos condensadores idénticos en paralelo. Fuente propia

Hay que tener claro qué quiere decir “poner dos condensadores en paralelo”. Si la distancia entre ellos es muy pequeña, no hay que hacer ninguna consideración más allá de añadir a la ESL la inductancia del montaje (la de las pistas y vías hasta llegar a los planos de alimentación y de masa) a cada una de las dos ramas RLC que representan cada condensador y considerar ambas ramas en paralelo.

Pero a medida que los separamos, aparece también la inductancia de los planos, cuyo efecto es proporcional a la frecuencia. Por eso, si estamos cubriendo un rango de decenas de MHz, los condensadores deben estar muy juntos. A bajas frecuencias, incluso una separación de 10 cm podrá ignorarse.

Tipos de condensadores de desacoplo

Los condensadores utilizados en desacoplo son básicamente de tres tipos:

- **Cerámicos multicapa:** Presentan muy baja inductancia debido a su pequeño tamaño, así como muy baja ESR. Los usaremos para los condensadores de desacoplo de pequeño valor (generalmente no por encima de 47 μF), en el encapsulado más pequeño que encontremos. No son condensadores polarizados.
- **Electrolíticos de tántalo:** Presentan también baja inductancia, pero en cambio su ESR es alta. Están disponibles en valores de capacidad superiores a los electrolíticos. Son condensadores polarizados.
- **Electrolíticos de aluminio:** Los usaremos cuando necesitemos valores elevado de capacidad. Son también condensadores polarizados.

De estos tres tipos, por su baja ESR y ESL nos interesan los cerámicos multicapa. ¿Los conocías o los has usado alguna vez? Tus diseños estarán plagados de este tipo y sólo excepcionalmente usarás condensadores de otras clases. **Excepción:** algunos reguladores necesitan una ESR más elevada por razones de estabilidad del lazo de realimentación, y por eso recomiendan condensadores electrolíticos o de tántalo.

Condensadores cerámicos multicapa (MLCC)

Elegir qué MLCC (*multi-layer ceramic capacitor*) usar sin estar familiarizado con sus características sería un grave error. Lo primero que tenemos que saber es que están disponibles en distintos dieléctricos agrupados en dos clases: NPO (Clase I), X7R, X5R e Y5V (Clase II). Las dos clases difieren en estabilidad, en el margen de temperatura y en valores disponibles.

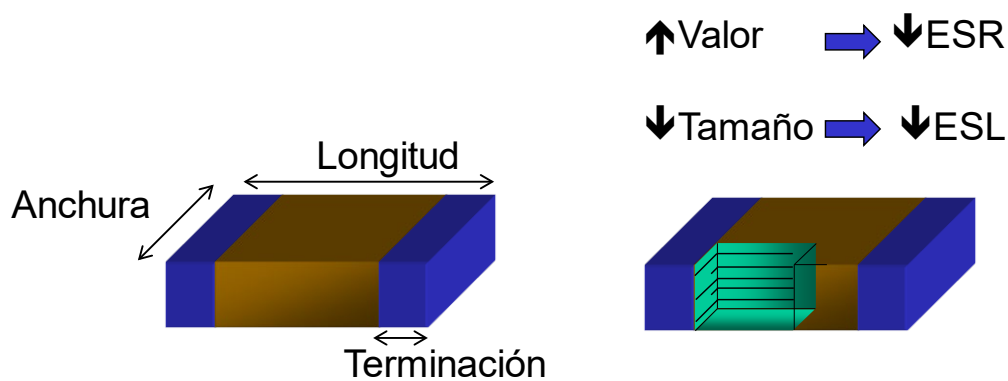
Los de tipo **NPO** son los más caros. Están disponibles sólo para valores pequeños (hasta 10 nF) y son bastante estables con la temperatura (del orden de 300 ppm/°C). Su uso suele restringirse a osciladores y circuitos sintonizados.

Un condensador de **clase II**, por ejemplo, con dieléctrico X7R está disponible hasta varios microfaradios y presenta una variación de $\pm 15\%$ en el margen de temperatura. Con dieléctrico Y5V encontramos condensadores de hasta aproximadamente 10 μF y con una variación entre +30 y -80% en temperatura. Por tanto, escogeremos valores de capacidad por encima de los nominales para hacer frente a estas variaciones.

Tabla 7.1. Comparación entre diferentes tipos de condensadores MLCC

Dieléctrico	NPO (COG)	X7R, X5R	Y5V
Características	Ultra estable	Estabilidad media	Baja estabilidad
Aplicaciones	Filtrado Temporización	Filtrado Temporización Desacoplo	Filtrado Desacoplo

Los condensadores que usaremos preferentemente para desacoplo serán, por precio, estabilidad y valores disponibles, los MLCC de clase II con dieléctricos X5R o X7R. Sólo si necesitamos valores muy elevados o si lo exige la estabilidad de los reguladores, usaremos condensadores de tantalio o de aluminio.



Tamaño	0402	0603	0805	1206	1210	1808	1812
Longitud (mm)	1	1,6	2	3,2	3,2	4,5	4,5
Anchura (mm)	0,5	0,8	1,25	1,6	2,5	2,03	3,2
Terminación (mm)	0,25	0,4	0,5	0,6	0,75	0,75	0,75

Figura 7.12. Tamaño de los encapsulados SMD habitualmente utilizados para condensadores MLCC. Fuente propia

Inductancia de montaje

Normalmente, los circuitos integrados no toleran más de un $\pm 5\%$ de variación respecto a la tensión de alimentación nominal. Teniendo en cuenta esta variación permisible de tensión (ΔV) y la demanda de corriente (ΔI), obtenemos un objetivo para la impedancia máxima (Z_{target}) del desacoplo entre continua y una frecuencia límite que definiremos más adelante:

$$Z_{\text{target}} = \frac{\Delta V}{\Delta I}$$

Debemos tener en cuenta la inductancia de las conexiones de los condensadores a los planos y de las vías. De este modo, la frecuencia de resonancia teórica del componente es mayor que cuando está montado en el sistema.

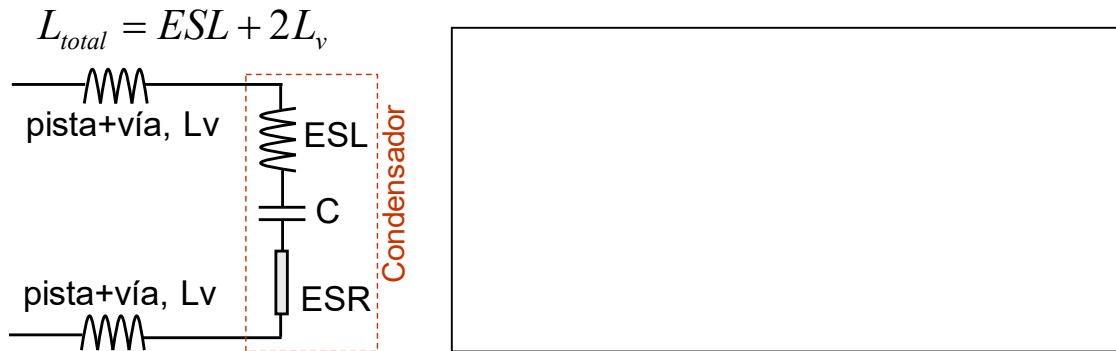


Figura 7.13. Disminución de la frecuencia de resonancia de un condensador al incluir el efecto del montaje en el PCB. Fuente propia

Las pistas y las vías de conexión a los planos de alimentación y masa presentan una inductancia que depende de cómo se realicen estas conexiones. No todos los montadores permiten *via-in-pad* (Figura 7.14, derecha), y todo aquel que haya rutado un PCB de cierta complejidad sabe que poner dos vías por pad ocupa demasiado espacio. Siendo realistas, lo mejor que podemos hacer la mayor parte de las veces es la configuración de la Figura 7.14, centro.

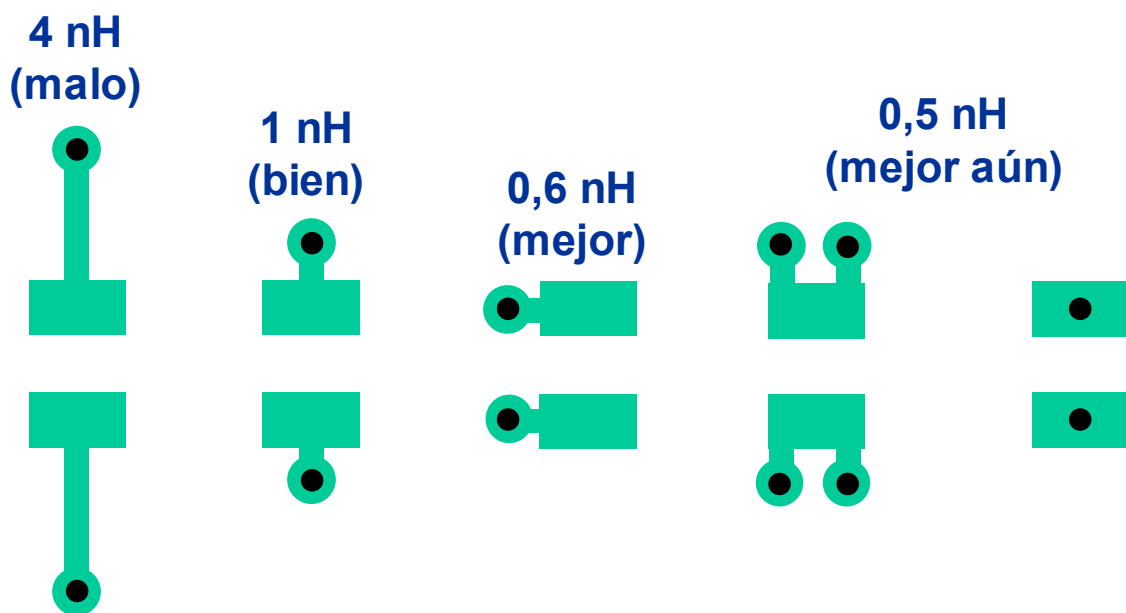


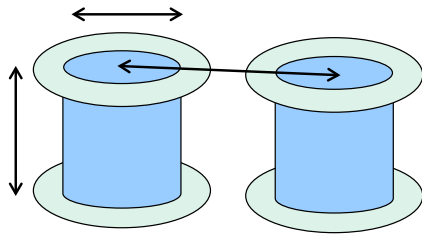
Figura 7.14. Valores orientativos para la inductancia de montaje del condensador en el PCB

Echa un vistazo a la Figura 7.6. ¿Cómo debería haber conectado los condensadores a masa y alimentación para reducir la inductancia de montaje? Es un buen ejercicio dibujarlo en papel.

Pero también sabe cualquiera que haya rutado un PCB medianamente complejo, que no pocas veces nos vemos obligado a rutar como se indica en las dos configuraciones de la izquierda. Como es habitual, choca lo académico con lo real.

Si queremos hacer estimaciones más realistas, la Figura 7.15 y la Figura 7.16 nos permite estimar la inductancia de las vías y de las pistas que las conectan al componente. Las estimaciones de los ejemplos nos permiten comprender que el componente, a menudo, es quien menos inductancia aporta al total. Un condensador de desacoplo será poco efectivo si presenta excesiva inductancia. Esta es una de las razones por las que colocamos varios en paralelo.

Inductancia de un par de vías (h en mils): $L(\text{pH}) = 10h \cdot \ln\left(\frac{2S}{d}\right)$



Por ejemplo, un par de vías de 1,6 mm de alto, de diámetro, separadas 2 mm presentan **1,9 nH**

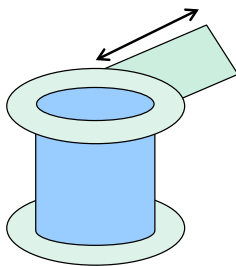
Al reducir el diámetro a 0,4 mm la inductancia baja a **1,45 nH**

Por lo tanto, la **impedancia** es demasiado alta:

$$Z \approx \frac{\pi L}{t_f} = \frac{\pi \cdot 1,9 \text{ nH}}{1 \text{ ns}} = 6 \Omega$$

Figura 7.15. Estimación de la inductancia de un par de pistas. Fuente propia

La **inductancia** de una de las pistas de longitud L se calcula como: $Z_0 \cdot L / c_m$, donde C_m es la velocidad de la luz en el medio.



Por ejemplo, una pista de 1 mm, con $Z_0=50\Omega$, $\epsilon_r=4$, presenta una inductancia de **333 pH**

El valor de dos pistas (666 pH) debe sumarse al de las vías en el ejemplo de la diapositiva anterior, alcanzando **2,57 nH**

Para una señal de 1 ns de tiempo de flanco, la **impedancia** es demasiado alta:

$$Z \approx \frac{\pi L}{t_f} = \frac{\pi \cdot 2,57 \text{ nH}}{1 \text{ ns}} = 8 \Omega$$

Figura 7.16. Estimación de la inductancia de una pista corta. Fuente propia

Inductancia de los planos de masa y de alimentación (plane loop inductance o spreading inductance)

En la Figura 7.7 se incluye la inductancia del camino entre el condensador y el circuito integrado como parte de la inductancia total a considerar. Sin tener que recurrir a simulaciones electromagnéticas, podemos usar una sencilla expresión para estimar el valor de la inductancia del camino formado por los planos de masa y de alimentación.

Si h es la separación entre planos de masa y de alimentación, L es la distancia entre condensador y circuito integrado y w es la anchura de los planos, siempre que $w > 10 \cdot h$ es válida la siguiente aproximación:

$$L_{\text{loop}} = \mu_0 \cdot h \cdot L / w, \text{ siendo } \mu_0 = 4 \cdot \pi \cdot 10^{-7} \text{ H/m}$$

Por ejemplo, con $h=200\text{ }\mu\text{m}$, un plano cuadrado de cualquier tamaño presenta una inductancia de 251 pH. Si el plano es rectangular, cuenta cuántos cuadrados contiene y multiplica por este valor. Este resultado es importante porque (1) a la hora de usar *split planes* puedes sentirte tentado a hacer planos de alimentación estrechos, lo que resultaría en una inductancia muy elevada y (2) se pone en evidencia la necesidad de ubicar los condensadores que cubren altas frecuencias muy cerca de los circuitos integrados.

Ubicación de los condensadores de desacoplo

Para ser efectivos, los condensadores deben ubicarse dentro de un “radio de acción” que depende del margen de frecuencias en el que presentan baja impedancia. En [21] se propone como criterio que este radio de acción es $1/40$ de la longitud de onda correspondiente a la frecuencia de resonancia del condensador:

$$\text{radio} = \frac{\lambda}{40} = \frac{c_0 / \sqrt{\epsilon_r}}{40 \cdot f_r}$$

De este modo, un condensador de bulk de 100 μF y una inductancia de 10 nH, que resuena a unos 160 kHz, debe colocarse a menos de 23,5 metros de los circuitos integrados a los que debe servir carga. Es decir, que su radio de acción abarca sin problemas todo el PCB.

En cambio, un condensador de 4,7 nF y una inductancia de 0,7 nH, que resuena a aproximadamente 88 MHz, debe colocarse a menos de 4,3 cm del circuito integrado a proteger.

Este método de cálculo no tiene en cuenta la inductancia de los planos de masa y de alimentación y mi consejo es que si te es posible reduzcas en al menos un factor 3 esta distancia.

La herramienta Altera PDN Design Tool

En las siguientes secciones vamos a desarrollar una metodología sencilla, pero muy empleada, para diseñar la red de desacoplo para un circuito integrado. En las siguientes páginas te explicaré la metodología, añadiré algunas explicaciones y pondré algunas capturas de pantalla de un ejemplo de diseño.

Estas capturas son de una hoja de cálculo para Excel desarrollada por Altera (fabricante de FPGAs) en 2009, que forma parte de Intel desde 2015. Esta herramienta usa un modelo compatible con el que vamos a exponer. Si buscas “**Altera PDN Tool**” en el navegador, encontrarás la página de descarga. El fichero que necesitas está bajo el nombre “*Power Delivery Network (PDN) Tool (ZIP) (device agnostic)*”. Y aquí va la primera captura, que resume el modelo definido en esta hoja de cálculo.

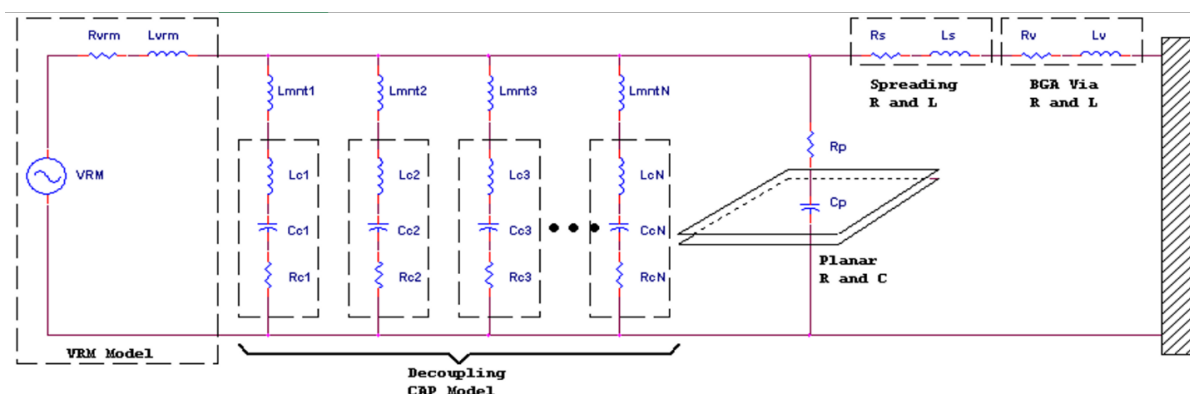


Figura 7.17. Modelo de la red desacoplo en la herramienta Altera PDN Design Tool. Captura de pantalla

El modelo (Figura 7.17, izquierda) comienza por considerar un regulador, cuya impedancia equivalente (modelo de primer orden) es un circuito RL serie formado por 1 m Ω y 10-30 nH. A continuación, encontramos la red de condensadores, cada uno con su circuito RLC serie, incluyendo la inductancia de montaje. Este último aspecto se modela en detalle, al permitir definir los parámetros geométricos para dos topologías habituales de montaje (vías a los lados del condensador y vías en los extremos, mira la Figura

7.18). La herramienta usa fórmulas (similares a las que hemos presentado en las páginas anteriores) para calcular la inductancia de montaje.

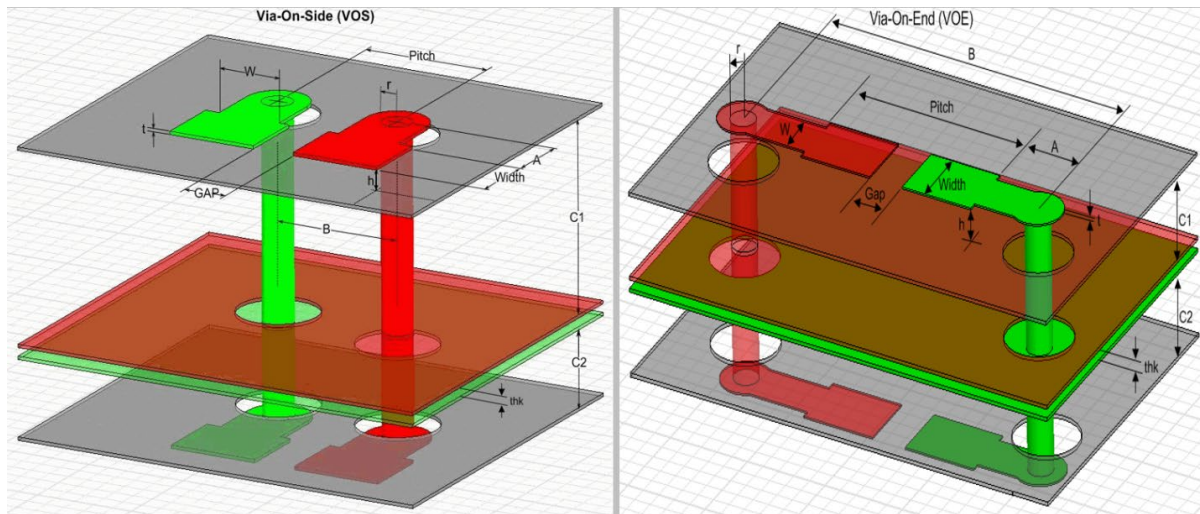


Figura 7.18. Modelado de la inductancia de montaje en Altera PDN Design Tool. Captura de pantalla

Altera PDN Design Tool tiene una biblioteca de componentes estándar, definiendo una inductancia típica para cada tamaño de encapsulado (por ejemplo, 0,4 nH para 0402 y 0,5 nH para 0603) y un valor de ESR que depende tanto del encapsulado como del valor de capacidad (la ESR disminuye al aumentar la capacidad).

Hemos comentado en páginas anteriores que los condensadores dejan de ser útiles (dejan de tener impedancias en el orden de miliohmios) a partir de 100-150 MHz. Por encima de esta frecuencia sólo podemos contar con la capacidad entre planos de masa y de alimentación. La herramienta permite, tras definir la geometría (largo y ancho de los planos, constante dieléctrica y distancias del plano de alimentación al plano o a los planos de masa) calcular la capacidad entre planos. También se calcula la resistencia óhmica del plano.

Un penúltimo elemento del modelo está formado por la resistencia e inductancia en el camino entre los condensadores y los pines (o bolas) del encapsulado. De estos dos términos, el relevante es la inductancia, que se conoce como *spreading inductance*. Pero en este caso la herramienta no realiza cálculos; simplemente permite elegir entre tres escenarios de 15, 30, 45 pH o definir un valor concreto.

Por último, se modela la resistencia e inductancia de las vías desde los pines o bolas del encapsulado del circuito integrado a los planos de alimentación.

En resumen, el modelo contempla todos los elementos relevantes, siendo detallado en unos y muy poco preciso en otros. En cualquier caso, vale para nuestros propósitos.

EJEMPLO 1: red de desacoplo para 64 ADCs

Un módulo de adquisición de datos de 64 canales que diseñé en 2017 contiene un amplificador y un ADC por canal (matriz de 8x8 celdas en la parte central del módulo, Figura 7.19). Los ADCs se alimentan a 3,3V a partir de un regulador lineal (rodeado en rojo en la Figura 7.19, el esquema se muestra en la Figura 7.20 y el rutado en la Figura 7.21). Aunque la eficiencia es muy pobre (rondando el 50%, pues regulamos de 6V a 3,3V) se usa un regulador lineal para reducir el ruido.

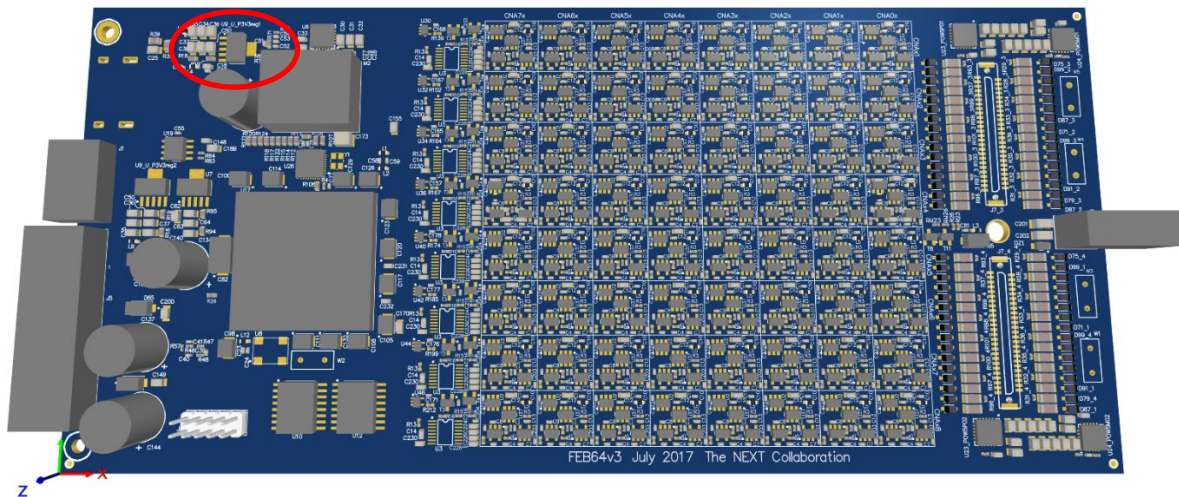


Figura 7.19. Módulo de adquisición de datos de 64 canales del ejemplo. Fuente propia

En el esquema del regulador, la ferrita L9 tiene la función de atenuar el ruido de alta frecuencia proveniente de un regulador conmutado que genera la tensión de entrada (6V). Los tres condensadores de entrada C34-C36 de tamaño 0805 deberían presentar baja impedancia a baja frecuencia. Lo comprobaremos con el simulador. La capacidad a la salida del regulador está formada también por tres condensadores 0805 de 10 μ F (C37, C38 y C50): esta será la capacidad de bulk de la red de desacoplo de esta alimentación.

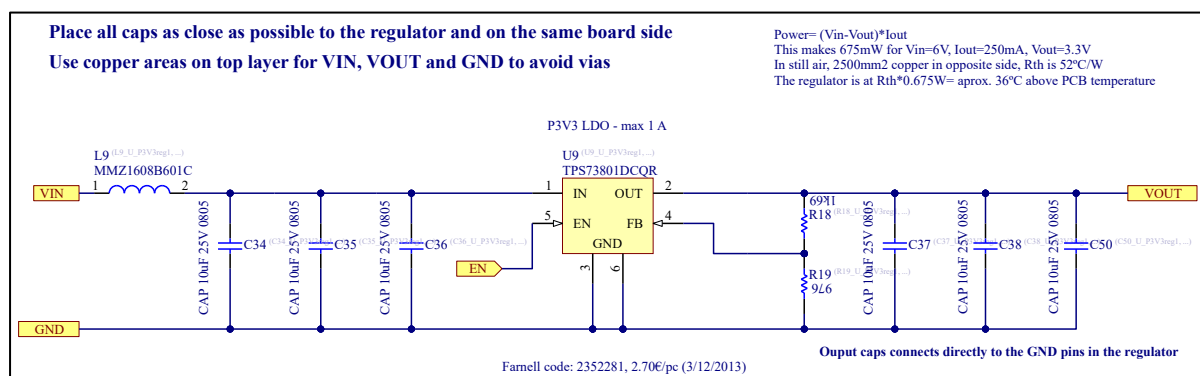


Figura 7.20. Esquema del regulador de 3,3V. Fuente propia

En la Figura 7.21, hemos usado áreas de cobre en capa *top* para las tensiones de entrada, salida y masa. Los condensadores de entrada y de salida están conectados directamente a estas áreas, reduciendo el acoplamiento de ruido desde otras partes del PCB. Estas tres áreas se conectan a planos internos mediante vías.

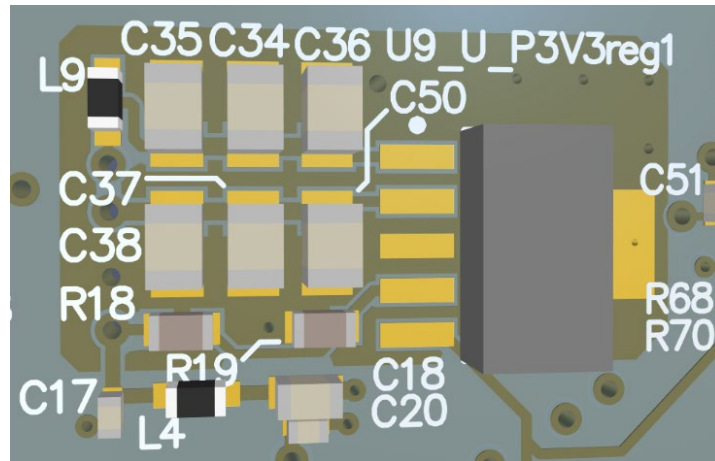


Figura 7.21. Rutado del regulador de 3,3V. Fuente propia

El regulador alimenta los ADCs de 64 canales. Cada ADC dispone de un condensador de 4,7 μF 0603 (montado en capa *bottom*) y otro de 100 nF 0402 como desacoplo local en capa *top* (Figura 7.22).

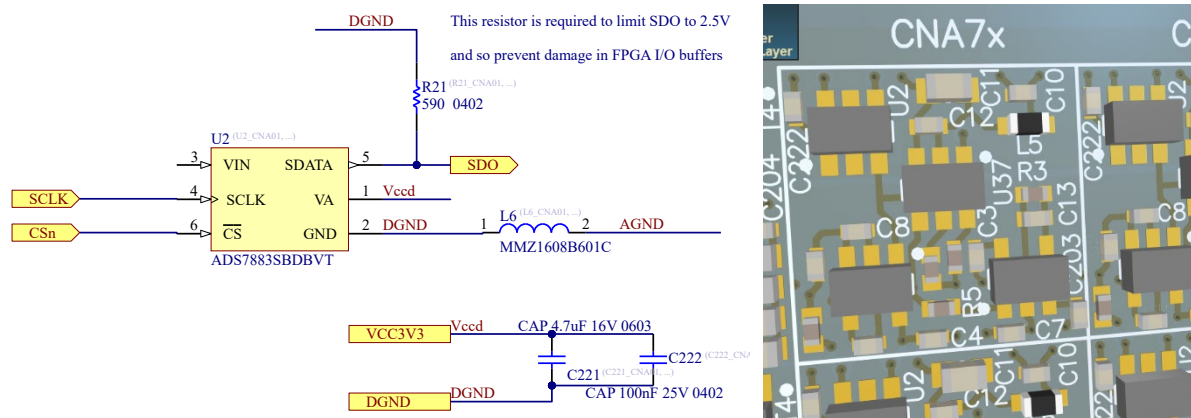


Figura 7.22. Esquema y rutado (capa *top*) de un canal, mostrando el ADC (U2) y un condensador de desacoplo de 100 nF (C222). El segundo condensador de desacoplo (C221) asociado al ADC está en capa *bottom*. Fuente propia

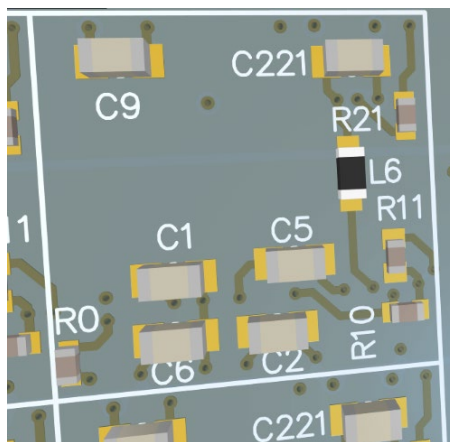


Figura 7.23. Rutado de un canal en capa *bottom*. El condensador de desacoplo es C221, de 4,7 μF . Fuente propia

Los ADCs son de 12 bit, digitalizan a 1 Msa/s (*megasamples per second*, 1 MHz) y consumen ligeramente por debajo de 2 mA cada uno, 128 mA en total para todos los ADCs. Para limitar el ruido, vamos a permitir

menos de 1 mV (un LSB *-least significant bit-* del ADC equivale a 805 μV) de variación en la alimentación de 3,3V. La red de desacoplo debe evitar, en todo el margen de frecuencias entre continua y una frecuencia a determinar, una impedancia superior a:

$$Z_{\text{target}} = \frac{\Delta V}{\Delta I} = \frac{1 \text{ mV}}{2 \text{ mA}} = 500 \text{ mohm}$$

La inductancia del pin y del *bonding* en el encapsulado presenta ya aproximadamente 1,1 nH [22]. Esta inductancia alcanza 500 m Ω a 72 MHz. Esto quiere decir que no merece la pena invertir esfuerzos en mantener la impedancia baja por encima de 72 MHz, ya que el propio encapsulado nos limita.

Una regla en el diseño de redes de desacoplo es determinar f_{target} , es decir, aquella frecuencia a la que la inductancia del encapsulado ya presenta tanta impedancia como la que podemos permitirnos (Z_{target}). Nuestros esfuerzos en el diseño de la red de desacoplo deben llegar hasta f_{target} , no más.

Ya tenemos nuestros dos parámetros de diseño de la red de desacoplo: impedancia máxima entre alimentación y masa (Z_{target}) y frecuencia. Ahora podemos usar Alterna PDN Design Tool para evaluar la solución escogida.

La Figura 7.24 muestra la gráfica de impedancia para tres condensadores 0805 de 10 μF (capacidad a la salida del regulador), un condensador 4,7 μF 0603 y otro de 100 nF 0402. Hemos fijado f_{target} a 70 MHz y Z_{target} a 500 m Ω , y observamos que, exceptuando un pico en torno a 270 kHz, la curva de impedancia está muy por debajo del límite que nos hemos impuesto. El pico resulta de la resonancia entre la inductancia del regulado lineal (que el modelo fija a 10 nH) y los condensadores de alto valor (4,7 y 10 μF). Como el fabricante del regulador recomienda en su hoja de datos precisamente una capacidad a la salida de 10 μF , podemos asumir que dicha resonancia no se producirá en la realidad (es decir, que el modelo del regulador no se corresponde con la realidad).

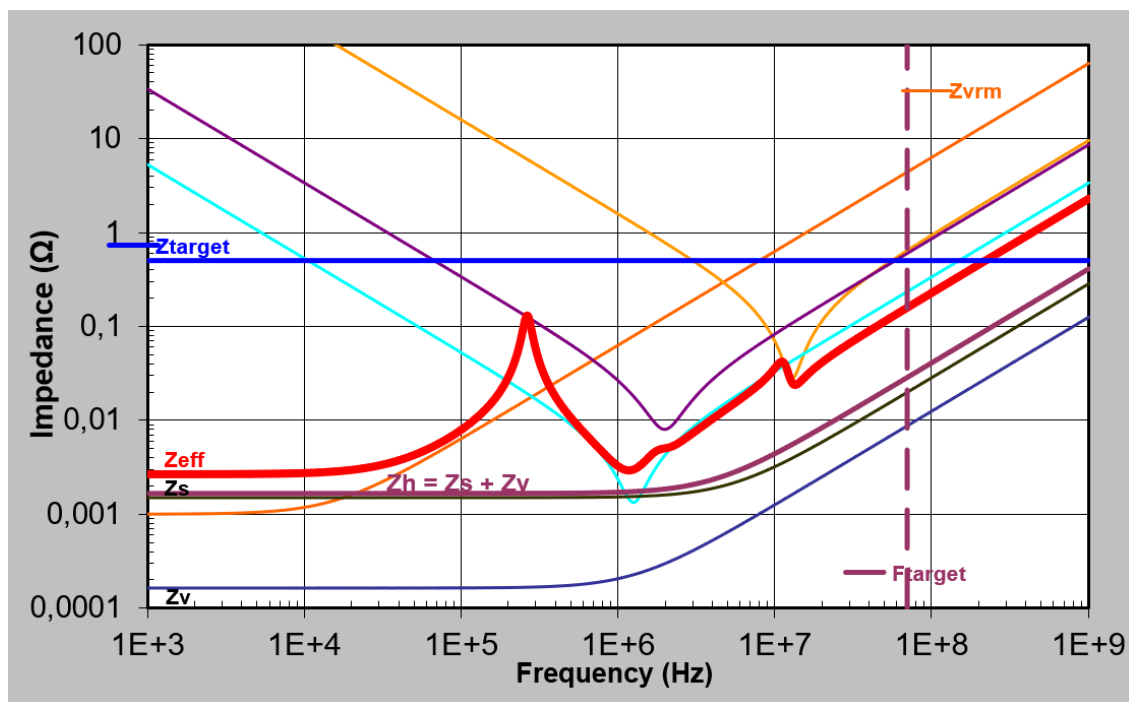


Figura 7.24. Curva de impedancia (rojo) para la capacidad a la salida del regulador y los dos condensadores asociados a un ADC. Captura de pantalla

Si simulamos ahora qué ocurre si sólo dejamos los tres condensadores de 10 μF a la salida del regulador, resulta que ahora la impedancia objetivo (Z_{target}) es 1mV/128mA (64 ADCs consumiendo 2 mA cada uno), unos 7,8 mohm. En este caso, sólo estamos por debajo de la impedancia objetivo hasta aproximadamente 5 MHz, lo que justifica que haya desacoplo local en las celdas (canales).

La Figura 7.25 muestra (izquierda) la simulación considerando un único ADC (Z_{target} de 500 m Ω) y un condensador de 4,6 μF 0603, sin los tres condensadores a la salida del regulador. A la derecha, añadiendo también un condensador de 100 nF 0402.

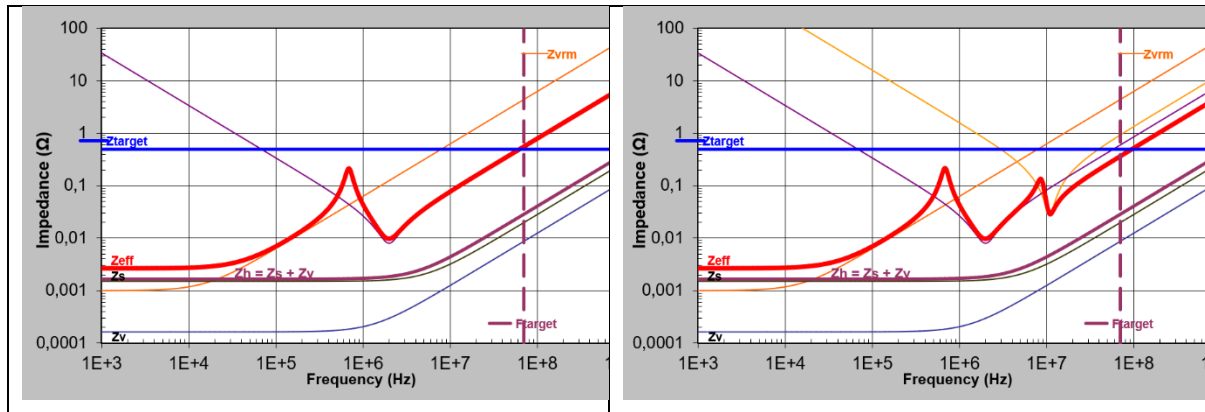


Figura 7.25. Simulación de la alimentación de un ADC sin considerar la capacidad de bulk. Capturas de pantalla

Lo que podemos deducir de este pequeño estudio es que el desacoplo de cada canal es adecuado, y que a la salida del regulador tal vez basta con un solo condensador de 10 μF y no tres.

EJEMPLO 2: red de desacoplo para la interfaz LVDS

El módulo de adquisición de datos dispone de una interfaz LVDS formada por cuatro pares diferenciales de 200 Mb/s. Las señales digitalizadas de los 64 canales se envían a un módulo concentrador (2 pares diferenciales), se recibe un reloj (tercer par) y señales de configuración y control (el cuarto par). Un cable HDMI conecta el módulo de adquisición con el concentrador. La Figura 7.26 muestra la cara *bottom* del módulo de adquisición, indicando en rojo la zona donde se ubican los buffers LVDS y el conector HDMI.

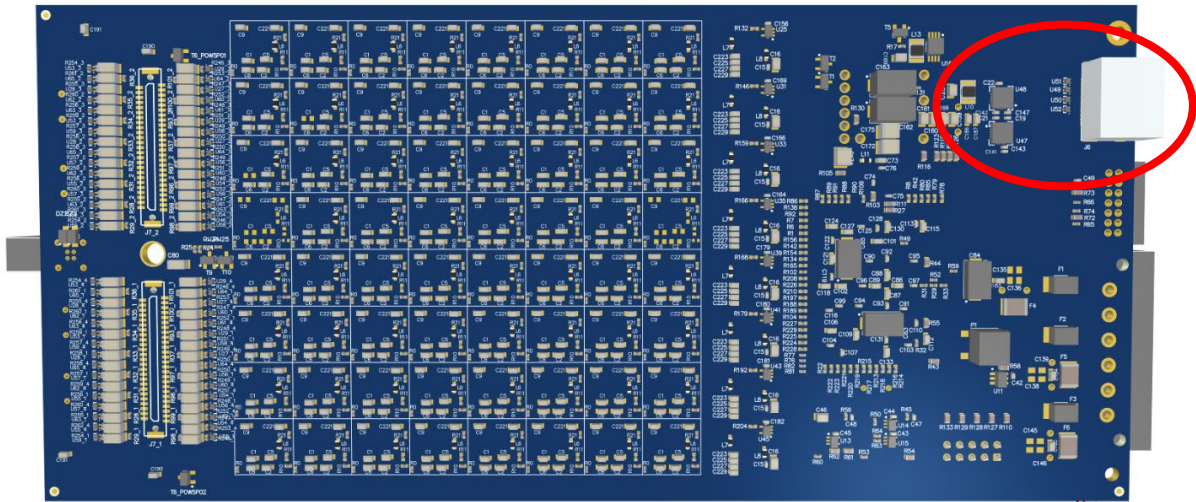
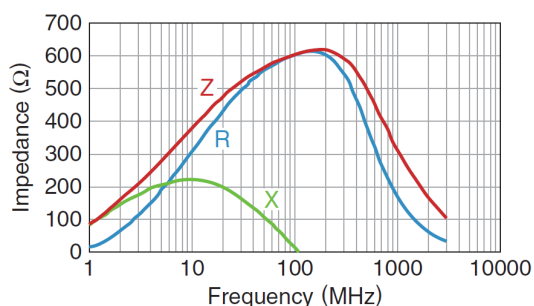


Figura 7.26. Zona (en rojo) donde se ubica la interfaz LVDS. Fuente propia

El conector HDMI (J6) tiene carcasa metálica y va conectada a chasis (pines MH1-MH4). Este nodo (chasis) está separado de la masa digital que queremos mantener libre de las perturbaciones que puedan introducirse por el conector (por ejemplo, descargas electrostáticas, o ESD). Nos aprovechamos del conector y cable HDMI, pero usamos sus señales con un protocolo y señalización diferentes: dos pares diferenciales para salida de datos (HM_DOUTx), una entrada diferencial de reloj (HM_CLK) y otra entrada para comandos (configuración y sincronización, HM_CMD).

Ya en el módulo, montamos una protección ESD para cada par diferencial, que consiste básicamente en diodos Zener en inversa capaces de soportar mucha corriente (denominados diodos TVS, *transient voltage suppressor*), cuyos ánodos están conectados a chasis, de modo que las corrientes de los transitorios quedan separadas de la masa digital.

En la Figura 7.27, abajo, podemos ver los dos buffers LVDS, que se interponen entre la FPGA del módulo de adquisición y el conector, protegiéndola de sobretensiones y sobrecorrientes. Estos dos *buffers* montados en capa *bottom*, DS25CP152TSQ, U47 y U48 en la Figura 7.28, son los dos circuitos integrados para los que queremos realizar un estudio de la red de desacoplo.



Los buffers se alimentan a 3,3V. Esta tensión se genera en un regulador lineal (U9_U_P3V3reg1 en la Figura 7.29). Con el fin de reducir el ruido que podría pasar entre el plano de P3V3 y la alimentación de los buffers LVDS, se introduce un filtro en pi formado por una ferrita y dos condensadores (C17, C18 y L4, ver Figura 7.29 y Figura 7.30). La ferrita L4 (MMZ1608B601C de TDK) tiene una resistencia en continua menor de 0,4 ohm, una impedancia de 600 ohm a 100 MHz y un rating de 500 mA.

Una ferrita no deja de ser (en primera aproximación) una red RLC paralelo con un factor de calidad bajo, lo que evita la aparición de resonancias.

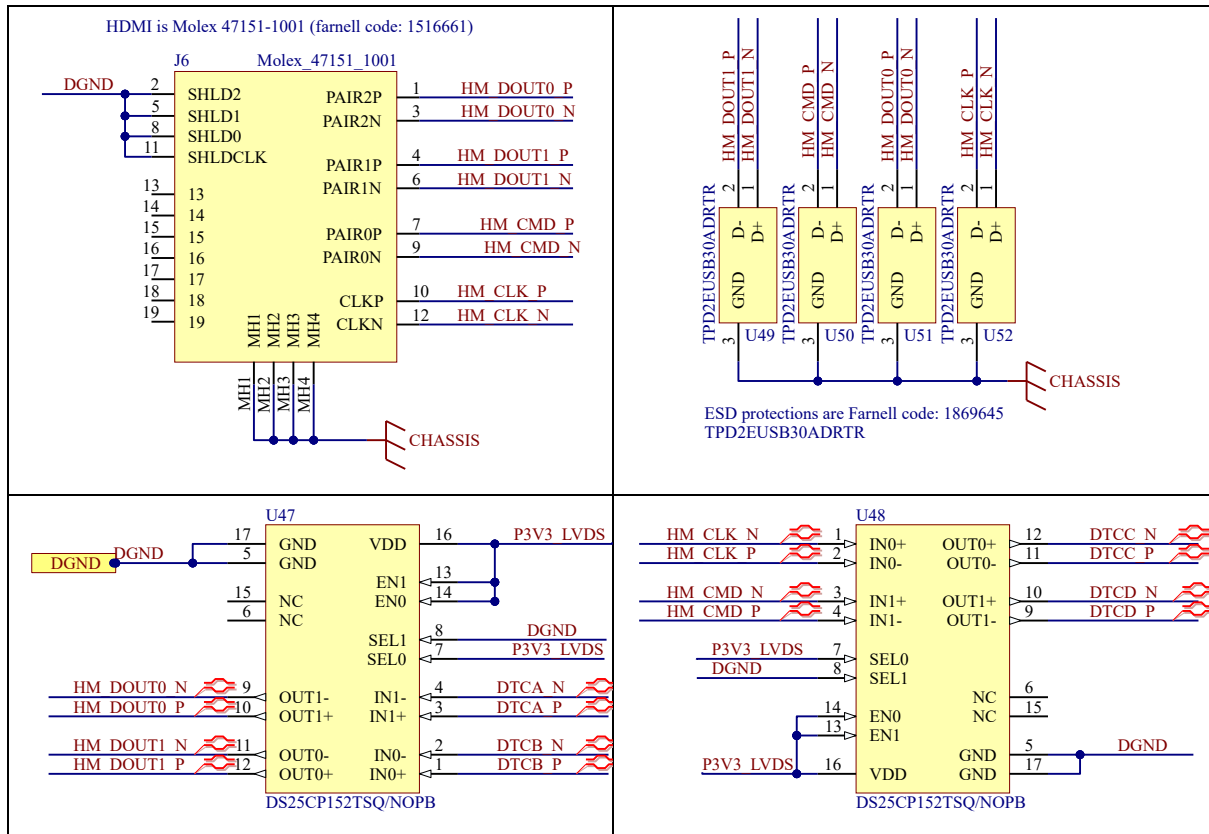


Figura 7.27. Detalla del diagrama esquemático del conector HDMI (arriba, izquierda), protecciones ESD (arriba, derecha) y buffers LVDS U47 y U48 (abajo). Fuente propia

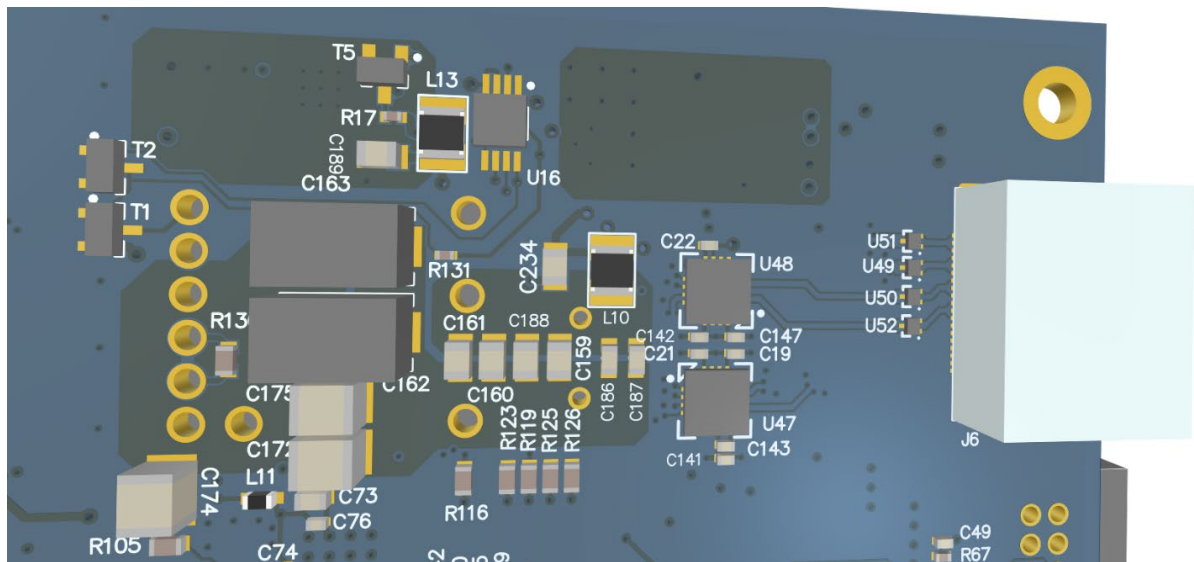


Figura 7.28. Detalle del rutado en *bottom*: buffers LVDS (U47, U48), conector HDMI (J6), protecciones ESD (U49-U52) y siete condensadores de desacoplo rodeando a los buffers. Fuente propia

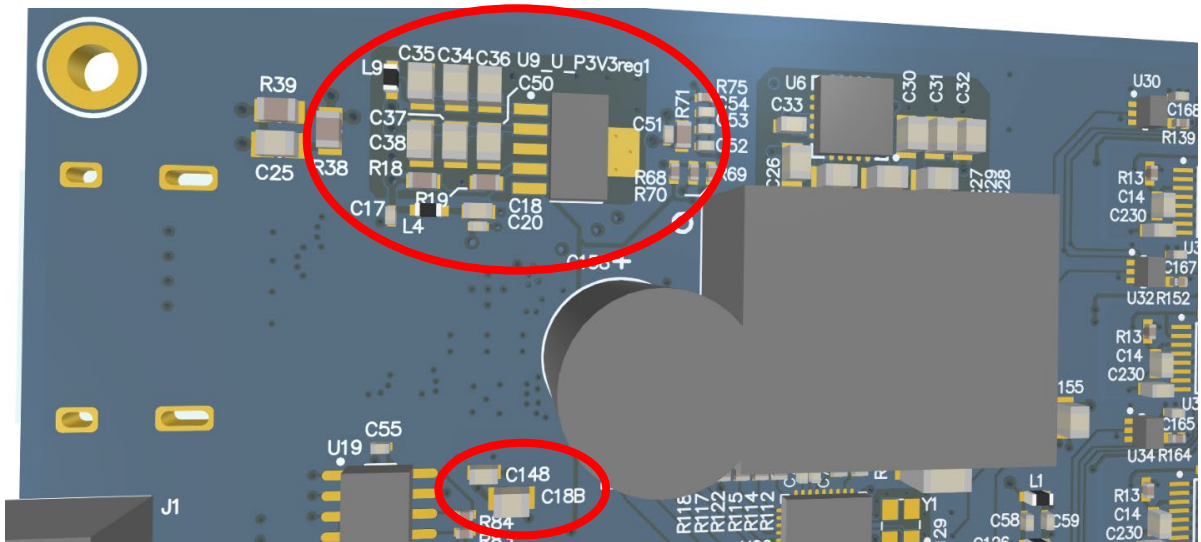


Figura 7.29. Detalle del *layout* mostrando el regulador lineal U9_U_P3V3reg1 que genera 3,3V, así como la red de filtrado formada por C17, L4 y C18. También en rojo, C18B y C148 son condensadores de 10 y 1 μ F que hacen la función de capacidad de bulk para la red de desacoplo de la alimentación de los buffers LVDS. Fuente propia

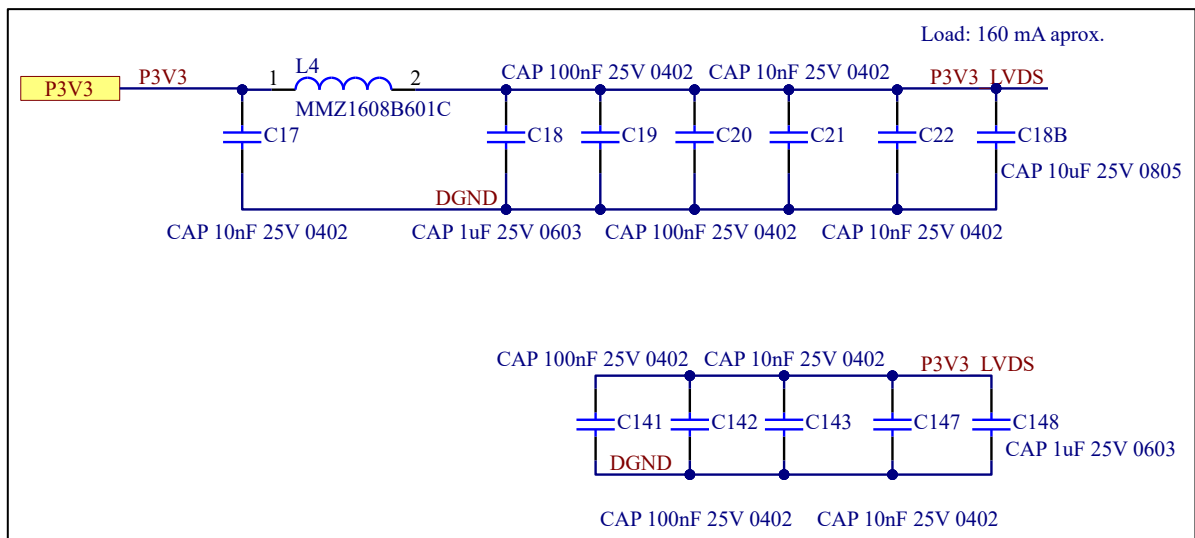


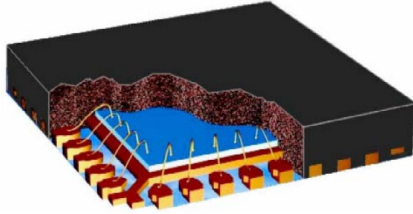
Figura 7.30. A partir de P3V3, una red de filtrado (C17, L4, C18) cumple el doble propósito de atenuar el ruido que otros componentes alimentados por P3V3 introducen en la alimentación de los buffers LVDS (P3V3_LVDS) y el que los buffers LVDS provocarían en otros componentes. El resto de los condensadores (10 en total) conforman la red de desacoplo para los buffers LVDS. Fuente propia

Tras el filtro en pi, la red de desacoplo está formada por un condensador 10 μ F 0805 (C18B) y otro de 1 μ F (C148), actuando como capacidad de bulk (Figura 7.29, abajo), así como cuatro condensadores (dos de 100 nF y dos de 10 nF) por cada buffer.

¿Cómo evaluamos la bondad de esta red de desacoplo? En primer lugar, hemos de determinar Z_{target} y f_{target} . Cada buffer (DS25CP152TSQ de Texas Instruments) tiene un consumo de corriente máximo de 77 mA. No sabemos cuánta de esta corriente es continua o estática (que no debe ser suministrada por la red de desacoplo), Y de la que es dinámica, tampoco conocemos su espectro en frecuencia. Hay herramientas que permiten definir el pico de corriente como una señal triangular o pulsada y estimar su espectro. Pero en una estimación a mano vamos a ser conservadores y, sin más datos, suponemos que todo el consumo es dinámico.

El buffer se alimenta a 3,3V y tiene un margen de ± 300 mV. Estos 300 mV se han de repartir entre el ruido que el consumo de los integrados provoca en la alimentación, así como al ruido acoplado desde el exterior y desde otras partes del módulo. Siendo conservadores, permitiremos una ΔV de 150 mV. De este modo:

$$Z_{target} = \frac{\Delta V}{\Delta I} = \frac{150 \text{ mV}}{77 \text{ mA}} \approx 2 \text{ ohm}$$



El encapsulado es un QFN de 21x18,5 mm² y 16 pines. Para determinar f_{target} (es decir, a qué frecuencia la inductancia parásita del encapsulado y del bonding alcanzan Z_{target} y por tanto hasta qué frecuencia tiene sentido que invirtamos esfuerzo en la red de desacoplo) podemos obtener información a partir de un modelo de simulación del componente. En concreto, una nota de aplicación de Texas Instruments [23] nos da una referencia para la inductancia de un encapsulado QFN de 16 pines, en torno a 0,9 nH. Resulta entonces:

$$f_{target} = \frac{Z_{target}}{2\pi L} = \frac{2 \text{ ohm}}{2\pi \cdot 0,9 \text{ nH}} \approx 350 \text{ MHz}$$

Este resultado, habida cuenta de que con condensadores difícilmente alcanzaremos un desacoplo por encima de 150 MHz, puede requerir considerar también la capacidad en entre el plano de alimentación y el (o los) planos de masa adyacentes. Veremos si en nuestro caso es necesario.

Si comparamos la inductancia parásita de varios encapsulados de 16 pines, vemos que puede haber diferencias de hasta un factor 4, y por tanto la f_{target} variará proporcionalmente.

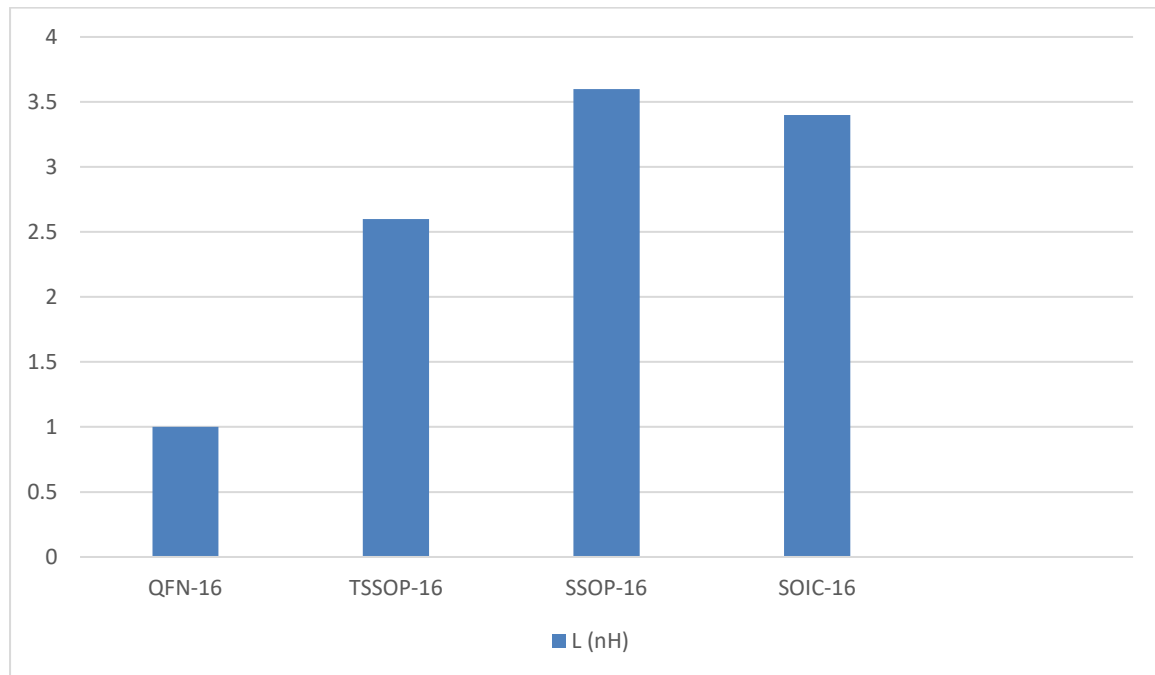


Figura 7.31. Comparación de parásitos típicos de varios encapsulado de 16 pines (datos extraídos de [23])

Reflexionemos un momento sobre el enorme margen de seguridad que supone emplear este método de diseño: ante el desconocimiento del espectro de los pulsos de corriente, asumimos que el valor nominal de consumo se da en TODAS las frecuencias, resultando en una Z_{target} con un margen de seguridad aplísimo. Estamos matando moscas a cañonazos.

Vamos a llevar Z_{target} , f_{target} y la red de desacoplo a Altera PDN Design Tool para evaluar la bondad del diseño. Ignorando la capacidad entre planos, resulta la curva de impedancia de la Figura 7.32.

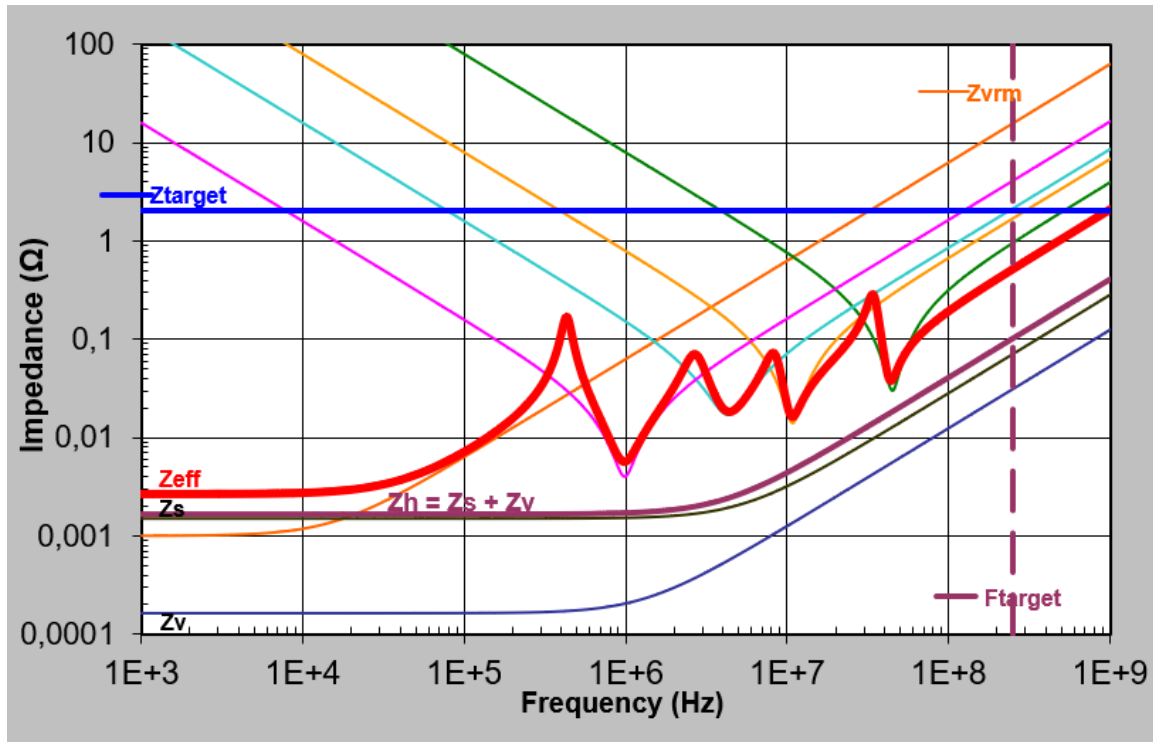
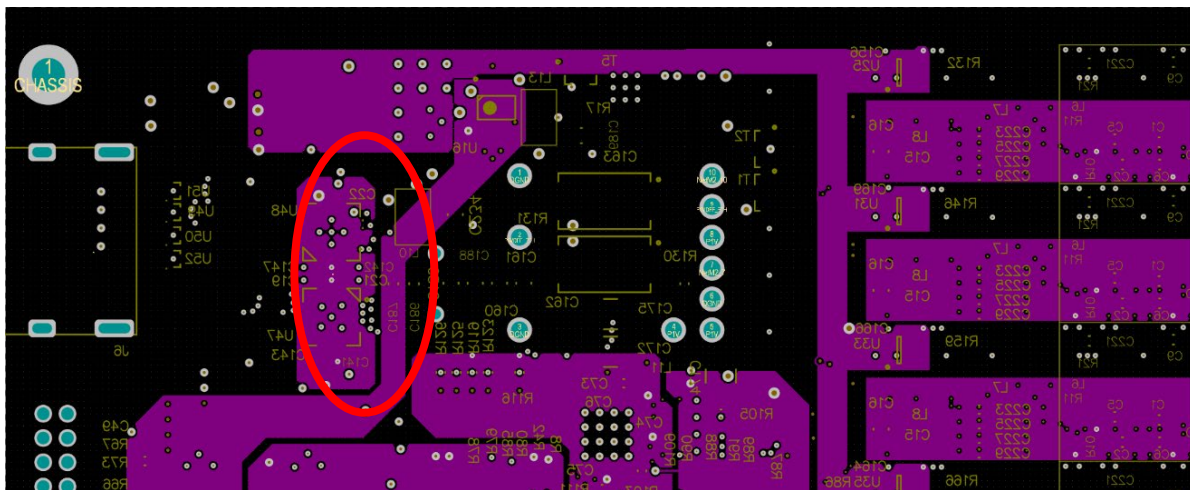


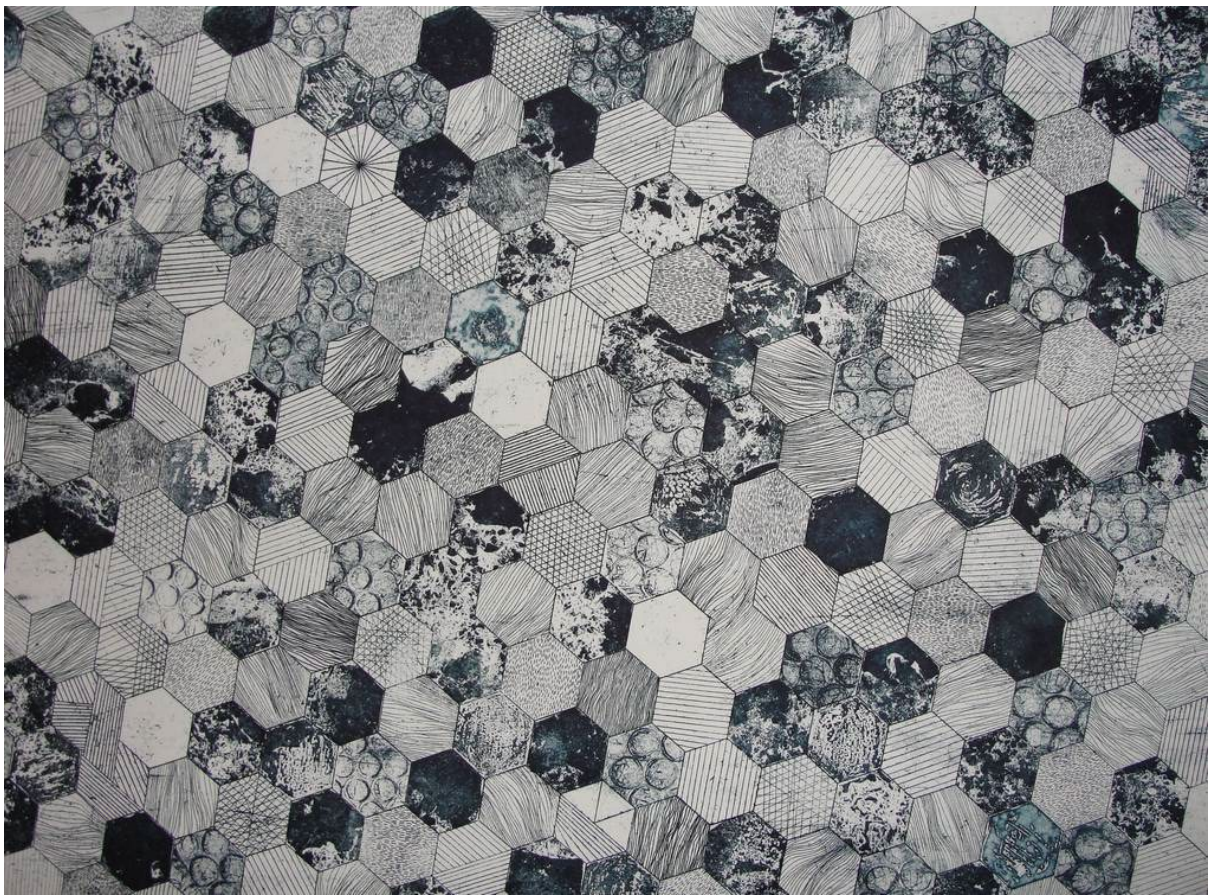
Figura 7.32. Impedancia de la red de desacoplo, asumiendo un regulador lineal y sin tener en cuenta la capacidad entre planos. La frecuencia objetiva está en 250 MHz, porque es el máximo valor que permite la herramienta, pero debemos fijarnos en el valor a 350 MHz. Captura de pantalla

La baja demanda de corriente nos permite, excepcionalmente, disponer de un margen útil que se extiende más allá de f_{target} . El diseño es adecuado, incluso se podría eliminar algún condensador de los cuatro de pequeño valor. Pero no suponen un problema para su rutado en este diseño, hay espacio suficiente. Puedes probar a jugar con la herramienta y probar otras posibilidades.

Cabe preguntarse sobre el efecto de la capacidad entre planos. La tensión P3V3_LVDS dispone de una pequeña área de 95 mm² en capa 7, siendo la capa 8 masa y estando a 130 micras de distancia. Introduciendo estos datos en la pestaña “Plane Cap” de la herramienta, observamos que la curva de impedancia sólo varía muy cerca del GHz.



Día 8. Metodología de diseño para integridad de señal



Fuente de la imagen: www.pexels.com

Ahora que has aprendido las bases de la integridad de señal (líneas de transmisión, impedancia de línea, caminos de retorno, reflexiones, terminaciones, diseño de estructuras multicapa, diseño de redes de desacoplo) y has asimilado un conjunto de buenas prácticas, es momento de integrar este conocimiento en la metodología de diseño de un producto electrónico.

Vamos, en primer lugar, a presentar los análisis de integridad en señal pre-layout (antes de rutar, incluso, diría yo, antes de completar los diagramas esquemáticos) como una valiosísima herramienta que te permitirá ahorrar tiempo y dinero en tu próximo producto que incluya al menos una sección digital. El

estudio de terminaciones y topologías de línea, de las distancias entre componentes en el diseño, la planificación de la estructura del stack-up y el diseño de las redes de desacoplo forman parte de esta fase.

Una vez rutado el diseño, verificaremos que las simulaciones aproximadas que realizamos antes de rutar siguen arrojando resultados positivos. Y añadiremos un estudio de crosstalk (acoplamiento de energía entre líneas) así como una verificación de la continuidad de los caminos de retorno y un estudio del power integrity (aunque sea en su versión más sencilla de estudio de caída de tensiones de alimentación en continua). Todo esto formará parte del análisis de integridad de señal post-layout.

La metodología de diseño es necesariamente iterativa: puede hacer falta volver a una fase previa (como cambiar a topología de un nodo o cambiar una terminación de línea) para corregir errores encontrados en una fase posterior del diseño.

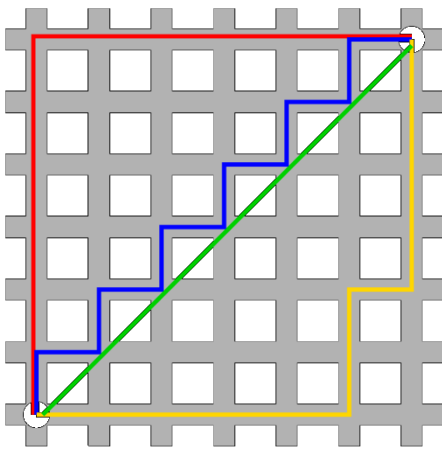
Todavía no hemos hablado de cómo realizaremos las simulaciones de integridad de señal. Presentaremos los modelos IBIS como un estándar de la industria que permite, sin que el fabricante desvele información propietaria, modelar el comportamiento de buffers de entrada y de salida que nos permitan simular las formas de onda en las líneas del PCB.

Análisis de integridad de señal pre-layout

¿Es posible tener una certeza razonable de que un bus va a funcionar correctamente sin haber construido un prototipo o ni tan siquiera haber diseñado el PCB? Este es el objetivo de las simulaciones de integridad de señal pre-layout (antes de rutar). En este tipo de análisis realizamos simulaciones analógicas de un diseño no rutado que incluyan el efecto de los buffers de E/S y de las líneas de transmisión para estimar reflexiones, *ground bounce* y retardos en una línea digital.

Todo lo que necesitas es saber la posición relativa de los circuitos integrados, asumir un valor de impedancia de línea y disponer de un modelo sencillo de cómo se comportan los buffers de entrada/salida. Bien, puede que haga falta alguna otra cosa si vamos a evaluar líneas gigabit o multi-gigabit. Pero comencemos por lo sencillo.

Para la primera condición, conocer la posición relativa entre circuitos integrados y la distancia entre ellos, suelo utilizar unas herramientas avanzadas de diseño electrónico denominadas PowerPoint y Presentaciones de Google. Basta con dibujar a escala 1:1 (suele caber en pantalla) el contorno del PCB y dentro de él rectángulos que representan a los circuitos integrados, también a escala 1:1. A este dibujo le llamamos *floorplan*. A partir de aquí es fácil estimar las distancias rectilíneas o **distancias Manhattan**.



Si conoces la posición de dos puntos que vas a conectar mediante una pista de PCB, pero no sabes cómo vas a rutar (líneas verde, azul o amarilla en la Figura 8.1) puedes hacer una aproximación mediante la suma de la distancia vertical y horizontal entre los dos puntos (línea roja). La denominamos distancia Manhattan, aludiendo al diseño en cuadrícula de la mayoría de las calles de la isla de Manhattan.

Debes identificar todas las líneas y buses críticos de tu diseño (aquellos con longitud mayor de la crítica, recuerda lo que estudiamos el Día 2).

En líneas punto a punto, colocaremos entre los dos buffers (uno de entrada, otro de salida) una línea de transmisión de la impedancia que hayas decidido y de la longitud Manhattan correspondiente.

Figura 8.1. Concepto de distancia Manhattan

En líneas multipunto, debes definir la topología (distancias Manhattan entre buffers) y simular múltiples escenarios, en el caso en el que haya más de un *driver* en la línea (un bus de datos con un microprocesador y varias memorias es un buen ejemplo).

La Figura 8.2 recoge la metodología de trabajo. Un editor gráfico te permite dibujar la topología, añadiendo líneas de transmisión (deberás definir como mínimo impedancia y retardo) y añadir *buffers* de entrada y de salida. En el simulador, asignaremos un modelo a cada *buffer*, ya sea uno genérico o (mucho mejor) un modelo de simulación proporcionado por el fabricante del circuito integrado. Por lo general serán modelos IBIS (los estudiaremos dentro de un rato), aunque muchos simuladores aceptan también modelos SPICE.

Con estas entradas, el simulador producirá unas formas de onda analógicas que deberás evaluar como aceptables o no con los criterios que estudiamos los Días 2 y 3. Si el resultado es satisfactorio, ya tienes unos criterios sobre cómo rutar esta pista o bus.

Si no es satisfactorio, debes introducir cambios (cambiar la topología, añadir terminaciones, acortar el bus o variar la distancia entre circuitos integrados, variar el buffer -si es configurable-) y en casos más drásticos simplemente decidir que la línea o bus no va a funcionar y que has de replantarte el diseño. Todas estas pruebas (denominadas **análisis “what if...?”**, “¿qué pasaría si...?”) llevan sólo unos pocos minutos y te pueden ahorrar meses de tiempo y miles de euros.

Para obtener resultados más exactos y para hacer estimaciones de *crosstalk* debemos rutar el PCB y hacer un análisis *post-layout*, pero en este momento ya obtienes unas reglas claras para el rutado y una alta probabilidad de que aplicando estas reglas el diseño funcione correctamente.

Definición de la topología en un editor gráfico

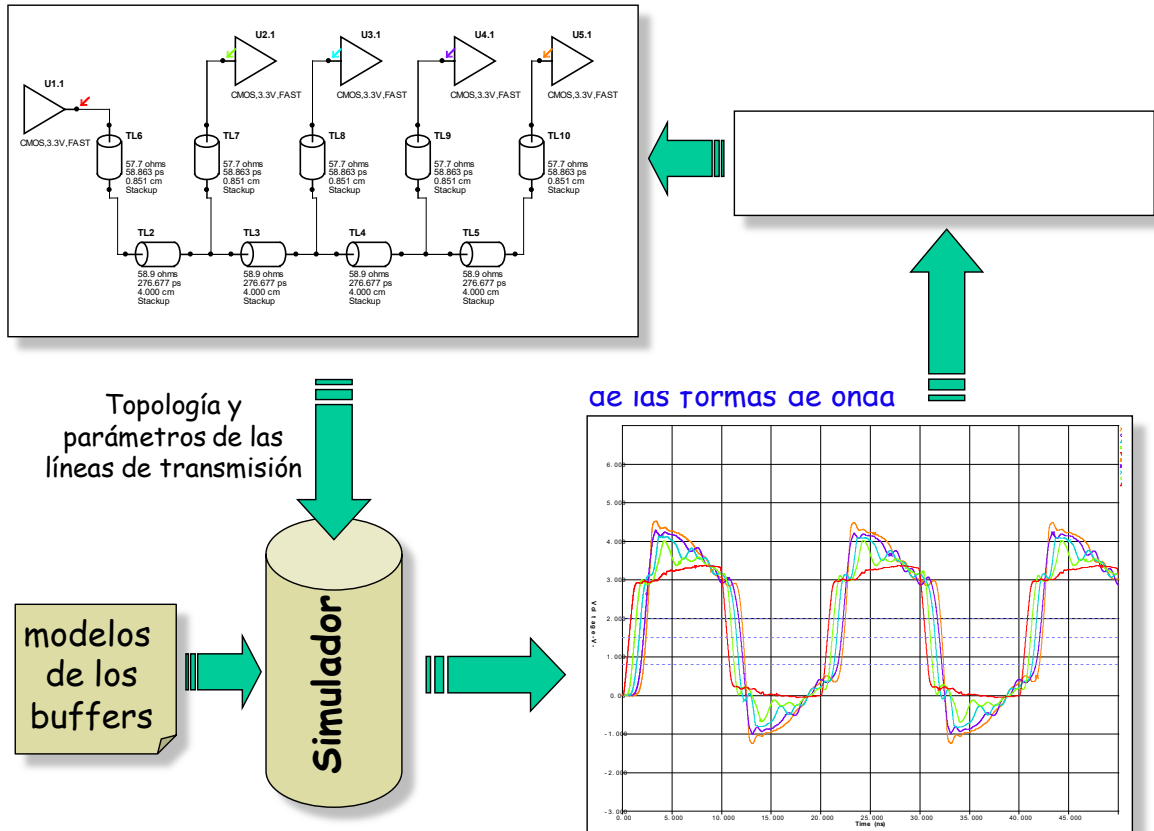


Figura 8.2. Análisis de integridad de señal pre-layout. Fuente propia

Análisis de integridad de señal post-layout

Una vez rutado el diseño, ya no es necesario estimar distancias Manhattan, dibujar la topología de los nodos ni definir la impedancia de las líneas: toda esta información, y usando distancias reales en lugar de estimadas, pueden ser extraídas automáticamente del diseño rutado.

En la Figura 8.3 se refleja este cambio: el diseño rutado es importado en la herramienta de simulación, el resto del proceso es similar al realizado en las simulaciones *pre-layout*. Pero trabajar con un diseño rutado ofrece una ventaja adicional: el conocimiento preciso de la geometría (por qué capa se ruta cada pista, distancia entre pistas y cercanía de los planos de alimentación y de masa -por donde circulan las corrientes de retorno-) permite realizar además un estudio de *crosstalk*, o diafonía.

Si los resultados de las simulaciones no son satisfactorios, deberemos introducir modificaciones en el diagrama esquemático y/o re-rutar la sección del PCB afectada y volver a simular las líneas de interés. El coste en tiempo de una modificación y una nueva simulación es mayor que en un análisis *pre-layout*, pero decididamente sigue implicando ahorros importantes en tiempo y costes en el desarrollo del producto electrónico.

Extracción de parámetros del PCB

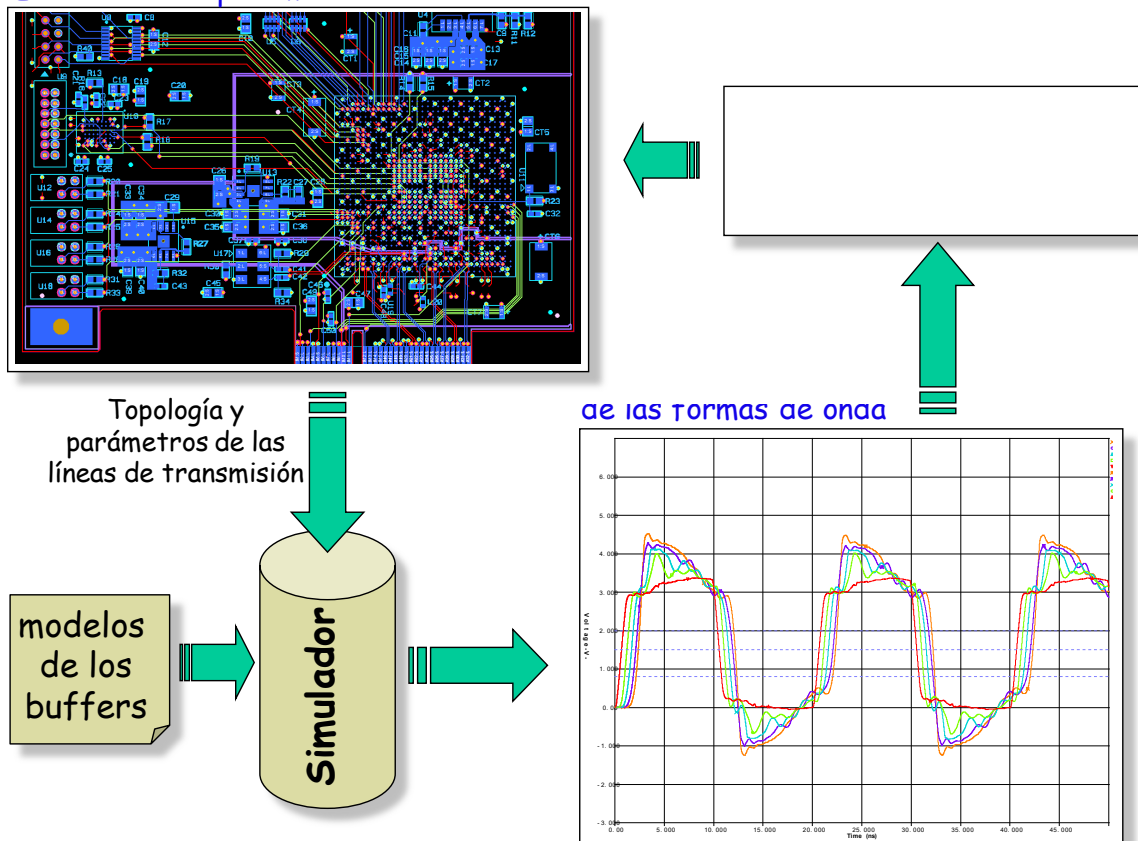


Figura 8.3. Análisis de integridad de señal post-layout. Fuente propia

Metodología de diseño

Lo que hemos estudiado en los días anteriores sugiere incorporar pasos en la metodología de diseño. **La captura del diagrama esquemático no está completa hasta que hayamos realizado también el diseño de la red de condensadores de desacoplo y hayamos propuesto un *floorplan*.**

En este momento abordaremos el diseño de la estructura del PCB multicapa, con todo lo que ello implica: número y función de cada capa, definición de anchuras de pista en cada capa para mantener la impedancia constante a medida que la señal salte entre capas, estudio de la continuidad de los caminos de las corrientes de retorno, generación de la documentación mínima de especificación y envío al fabricante del PCB elegido para su verificación.

Como tercer paso, realizaremos un análisis de integridad de señal *pre-layout*, y aunque esto se debería hacer con una herramienta de altas prestaciones como HyperLynx, también es posible realizar estudios sencillos con otras herramientas mucho menos potentes como son las simulaciones de integridad de señal de Altium (de hecho, Altium no permite simulaciones *pre-layout*, es necesario rutar antes, pero puede ser un rutado tentativo, no definitivo).

Todos los pasos descritos hasta el momento, que son iterativos, culminan en una validación de los diagramas esquemáticos.

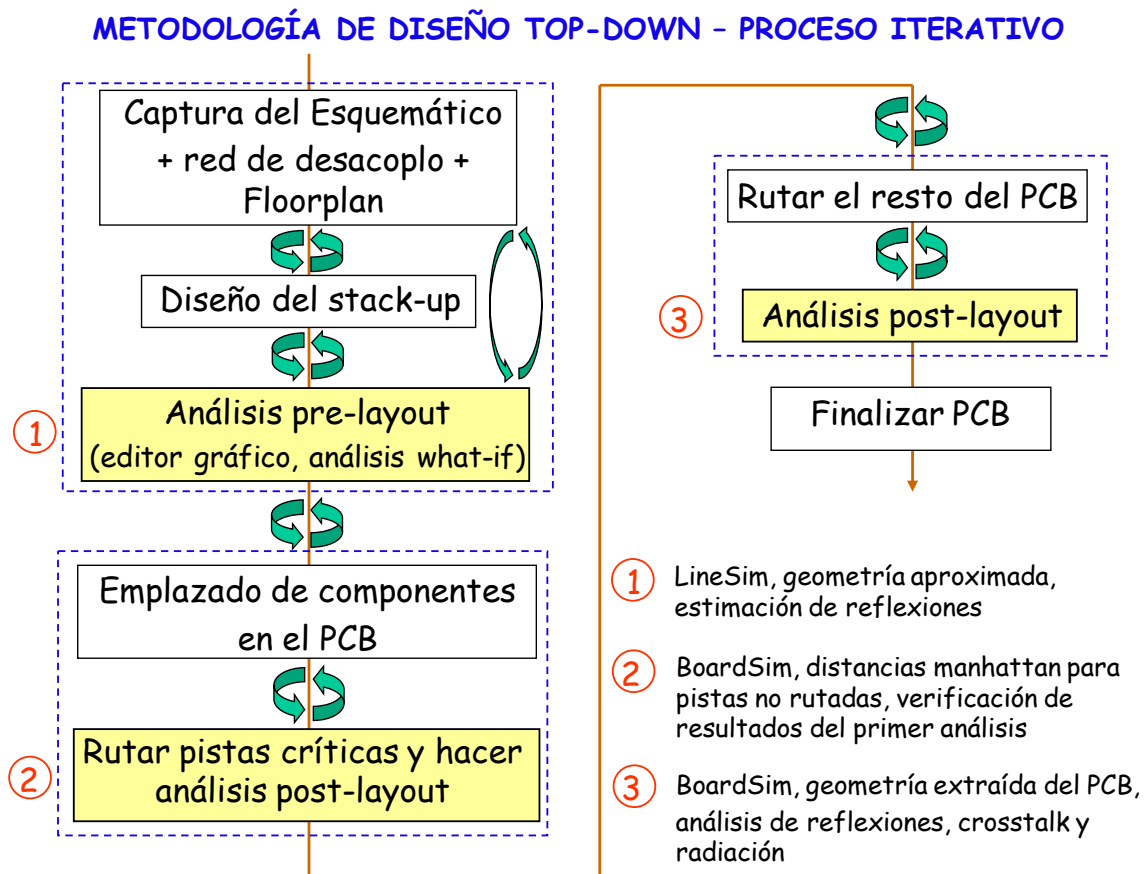


Figura 8.4. Metodología de diseño incluyendo la integridad de señal. Fuente propia

Es ahora cuando comienza el rutado: emplazamiento de los componentes en el PCB (nos serviremos del *floorplan* que ya tenemos hecho como guía) y rutado sólo de las pistas más críticas, aquellas que queramos que sean más cortas y con el menos salto de capas. Hablamos de líneas USB, SATA, PCI Express, todos los relojes, y en general señales que debemos mimar. Una simulación *post-layout* de estas líneas críticas nos permitirá validar el trabajo realizado. Sólo entonces debemos rutar el resto del PCB y, al finalizar, realizar una nueva simulación *post-layout* que incluya también un estudio del *crosstalk*.

Modelos IBIS

A principios de los años 90, Intel estaba desarrollando el bus PCI. Ahora el estándar para placas base de ordenador es PCI Express, basado en enlaces punto a punto diferenciales de alta velocidad, pero en los años 90 el bus paralelo PCI de 32 o 64 bits y 33/50/66 MHz, con sus 132-528 Mbyte/s (con PCI Express se superan 15 Gbyte/s) representaba el no va más en ancho de banda disponible para comunicar procesador, memoria y dispositivos de E/S.

Intel se encontró con que, sin realizar simulaciones analógicas de las formas de onda de las señales, los diseñadores de placas base de ordenador no iban a conseguir diseños exitosos. Es decir, popularizar PCI no sólo iba a requerir enseñar a los diseñadores una nueva disciplina (la integridad de señal), sino que iba a requerir nuevas herramientas de simulación.

SPICE permite simular líneas de transmisión. Si añadimos los modelos SPICE de los buffers de E/S ya tenemos una herramienta válida. El problema estriba en que un modelo SPICE es una representación física de los transistores, diodos y en general de la estructura que emplea el buffer. Y esto es información propietaria que el fabricante no siempre está dispuesto a proporcionar.

En 1993, Intel invita a otras empresas a formar el IBIS Open Forum para desarrollar un formato de especificación común. Aparece la versión 1.0 de la especificación IBIS (*I/O Buffer Information Specification*). La idea es desarrollar modelos basados en parámetros, tablas corriente-tensión y tensión-tiempo que representen el comportamiento de un buffer si necesidad de proporcionar información sobre cómo está construido. Como dificultad adicional, esto requiere simuladores no SPICE.

Las especificaciones de las diferentes versiones IBIS están disponibles en <https://ibis.org/specs/>

Evolución de la especificación IBIS

En las versiones 1.x (1993), los modelos de buffers CMOS y TTL de salida son ficheros ASCII que contienen tablas corriente-tensión (para los elementos 1, 2 y 3 en la Figura 8.5), parásitos del encapsulado (elemento 5) y velocidad de conmutación (elemento 4).

- Elementos 1 y 2: Transistores de pullup y pulldown (TTL, CMOS, BiCMOS) definidos por las tablas corriente-tensión.
- Elemento 3: los diodos de limitación también están definidos por tablas corriente-tensión
- Elemento 4: tiempos de subida y bajada expresados como un parámetro dV/dt
- Elemento 5: modela las capacidades, resistencia e inductancia parásitas en el pin y en el bonding

Los modelos de buffers de entrada sólo requieren los elementos 3 y 5.

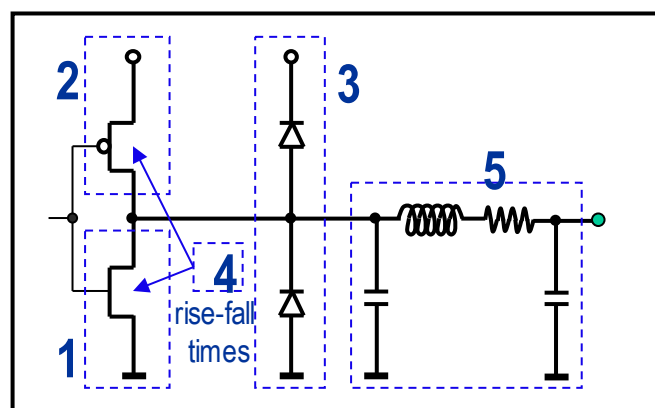


Figura 8.5. Secciones de un modelo IBIS

A partir de la versión 2.1, en 1995, IBIS se convierte en un estándar ANSI/EIA-656. Añade capacidad para modelar no sólo buffers CMOS y TTL, sino también ECL, PECL y líneas diferenciales. Los modelos IBIS que encontrarás habitualmente en la web de los fabricantes son por tanto de versiones 2.1 o posteriores.

Versiones posteriores añaden nuevas características, como modelado AMI y Touchstone, que quedan fuera del ámbito de esta introducción.

Ejemplo de fichero IBIS (.ibs): Buffer de reloj CY2305

CY2305, de Cypress Semiconductors, es un buffer de reloj con PLL con 5 salidas. Las salidas provienen directamente de la entrada o pasan a través de un PLL para eliminar el retardo interno (*zero-delay clock buffers*).

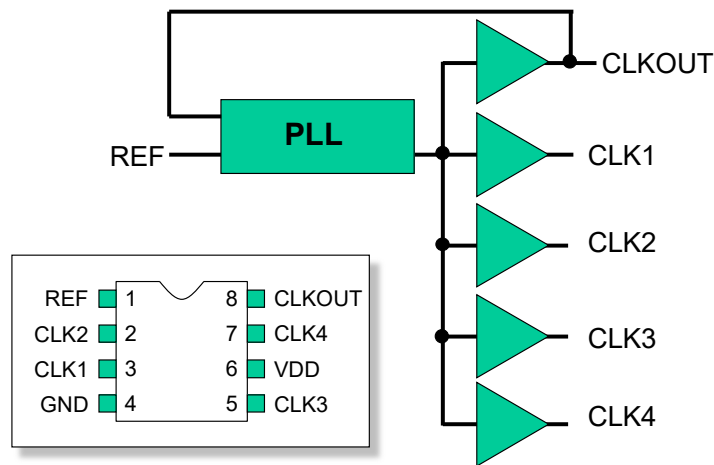


Figura 8.6. Diagrama funcional del buffer de reloj CY2305. Fuente propia

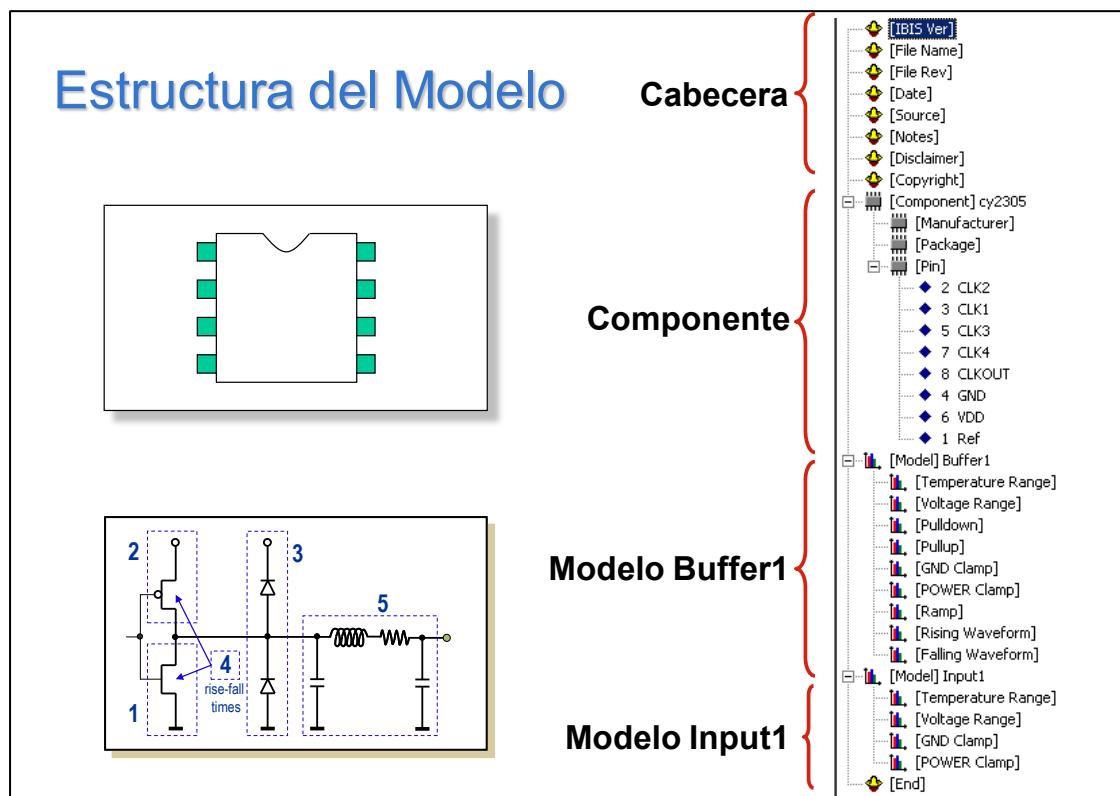


Figura 8.7. Estructura del fichero IBIS. Captura de pantalla

Un modelo IBIS no deja de ser un fichero de texto. En primer lugar, nos encontramos con una **cabecera**, que nos habla de la versión IBIS, nombre del modelo y otra información de identificación. Las palabras clave reservadas (*keywords*) van entre corchetes. Las líneas que comienzan por una barra vertical son comentarios.

```

*****
| IBS file cy2305.ibs created by s2ibis2 version 0.91BETA
| North Carolina State University Electronics Research Laboratory 1995
| *****
[IBIS ver]    2.1
[File name]   cy2305.ibs
[File Rev]    0
[Date]       Monday, October 26, 1998
[Source]      This file originated at Cypress Timing Technology
[Notes]       File created by JMO.
[Disclaimer]  This information is for modeling purposes only and is not guaranteed.
[Copyright]   Copyright Cypress Semiconductor, Inc. 1997

```

Después de la cabecera viene una sección de **declaración de componente y parásitos del encapsulado**, donde se indican los valores típicos, mínimo y máximo.

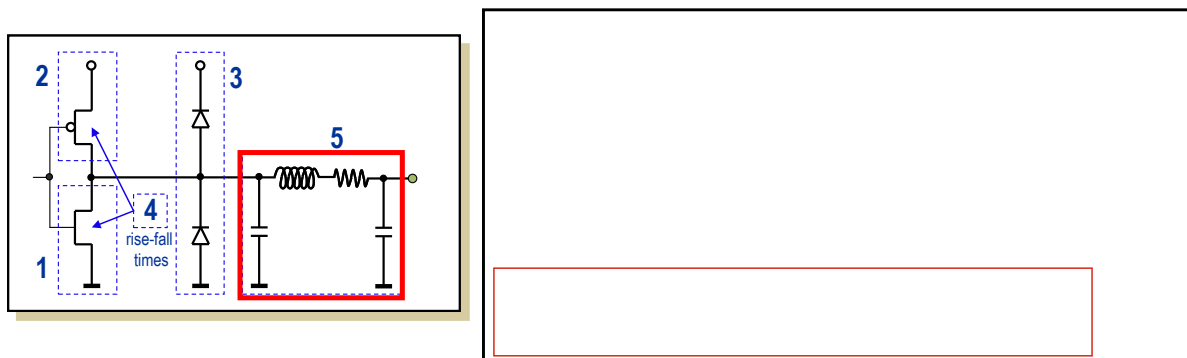


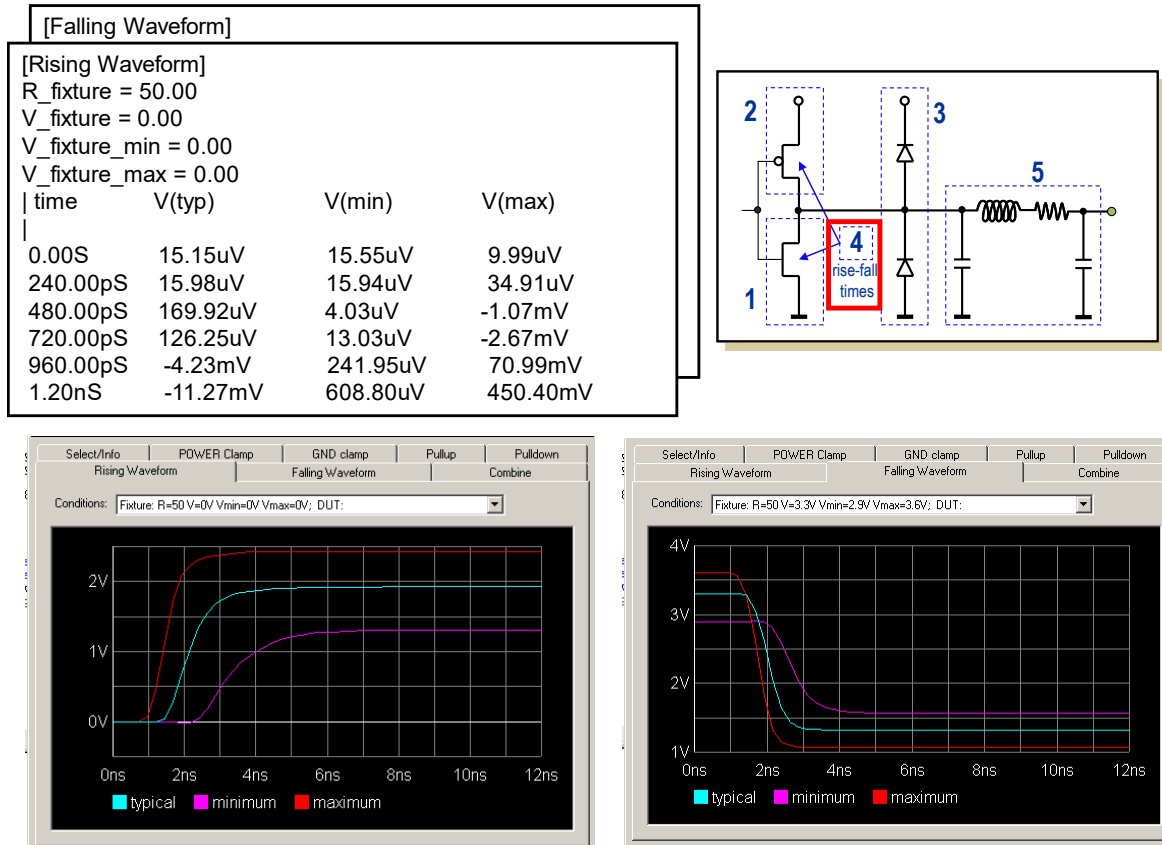
Figura 8.8. Sección de declaración del componente en el modelo IBIS

Tras palabra reservada [Pin] se define, para cada pin del encapsulado, su nombre, modelo y parásitos (si son distintos de los especificados tras la palabra reservada [Package]). Por cierto, es de aquí de donde podemos obtener el valor de la inductancia del pin y del bonding para estimar f_{target} en el diseño de la red de desacoplo.

[Pin]	signal_name	model_name	R_pin	L_pin	C_pin
2	CLK2	Buffer1			
3	CLK1	Buffer1			
5	CLK3	Buffer1			
7	CLK4	Buffer1			
8	CLKOUT	Buffer1			
4	GND	GND			
6	VDD	POWER			
1	Ref	Input1			

Se usan por defecto los valores de la sección anterior (encapsulado)

El fichero IBIS termina con la definición de los modelos de los buffers. Se especifica, para cada modelo, su nombre, tipo, umbrales lógicos, capacidad del buffer (sin incluir la parásita del encapsulado) y el rango de temperaturas y tensiones para simulación de casos típico, mínimo y máximo. Por ejemplo, el modelo "Buffer 1" declarado para las cinco salidas de reloj comienza así:



Un buffer de entrada requiere un modelo mucho más sencillo y se limita a la capacidad de entrada y a las tablas tensión-corriente de los diodos de protección ESD. Así, el modelo Input 1 para el componente CY2305. Captura de pantalla

PARTE 2

Fundamentos de Compatibilidad Electromagnética

Día 9. La Directiva Europea de EMC



Fuente de la imagen: <https://www.pexels.com/es-es/foto/calle-camino-carretera-de-una-sola-mano-536/>

Durante la mayor parte del Siglo XX (1933-1992) la regulación sobre emisiones radioeléctricas se refería a las interferencias sobre las transmisiones de radio producidas por otros equipos de radio y otras fuentes (como motores de encendido por chispa). No hacía falta regular emisiones ni susceptibilidad radiada ni conducida en productos electrónicos que no fueran radiadores intencionados. Simplemente porque la electrónica era, hasta los años 80, mayoritariamente analógica. Y eso, por lo general, quiere decir bajas emisiones.

No fue hasta la popularización de la electrónica digital a finales de los 70 y principios de los 80, con sus flancos de señal cada vez más abruptos y frecuencias de reloj cada vez más altas, cuando la capacidad de un producto de emitir interferencias de forma radiada y conducida comenzó a ser preocupante. También aumentó la susceptibilidad de los equipos, a medida que los microprocesadores y las máquinas de estados tomaban las riendas del comportamiento de los productos.

*El resultado fue la necesidad de regular y en 1992 aparece en **Europa** la directiva 89/336/EEC, conocida como Directiva EMC, en plena vigencia de 1996 y que supuso un shock para la industria y los supervisores europeos. La Directiva define dos caminos de conformidad: (1) una auto-certificación, en la que el fabricante o importador declaran (no demuestran) que cumplen los criterios según estándares armonizados y (2) el certificado por parte de un organismo competente de un expediente técnico de construcción.*

En **Estados Unidos**, la norma militar MIL-STD-461 fijaba ya en los años 60 no sólo límites a emisiones, sino que introdujo por primera vez en unas normas el concepto de susceptibilidad y de inmunidad. En los años 70, motivado por la proliferación de dispositivos digitales, la FCC (Federal Communications Commission) publica normas en materia de compatibilidad electromagnética.

En la actualidad hay alrededor de una decena de regulaciones nacionales o supranacionales en materia de EMC, si bien por el tamaño de su mercado las más influyentes son la europea y la FCC Estadounidense. Pero estas normas suelen estar basadas en estándares internacionales IEC o CISPR, por lo que en muchos casos los métodos de medida y los niveles exigibles son similares, no así los procedimientos burocráticos.

- Directiva EMC (Europa) 
- FCC (US) 
- EAC (Rusia) 
- VCCI (Japan) 
- BSMI (Taiwan) 
- Industry Canada 
- Australia and New Zealand 
- RRA Korea 
- CCC (China) 

Vocabulario básico en EMC

Conocer el vocabulario de una disciplina es condición *sine qua non*, requisito, para poder aprenderla. Vamos a comenzar con los términos más básicos e iremos introduciendo otros a lo largo de la lección de hoy.

Entorno electromagnético

Se define como la totalidad de los fenómenos electromagnéticos variables con el tiempo que existen en una región dada. Esto incluye señales electromagnéticas deseadas y no deseadas. Se puede describir por fuentes que pueden estar activas o mediante parámetros medibles como voltajes y corrientes eléctricas, e intensidades de campo eléctrico y magnético.

En otras palabras: tu producto será comprado por un cliente y puesto en servicio en un entorno específico. Allí recibirá perturbaciones conducidas por la línea de alimentación, como son el efecto de descargas de rayos cercanos, maniobras de conmutación en centrales eléctricas, conmutación de motores en el mismo edificio o descargas electrostáticas por parte de usuarios. También estará sometido a radiación electromagnética (como emisoras de radio cercanas o teléfonos móviles) y a transitorios que lleguen por conexiones de entrada/salida. Es cierto que las perturbaciones serán diferentes en cada instalación en particular, y eso constituye el entorno electromagnético en el que ha de operar tu producto.

Perturbación electromagnética

Se define como todo fenómeno electromagnético susceptible de afectar el funcionamiento de un dispositivo, de un aparato o de un sistema. Puede ser:

- un ruido electromagnético
- una señal no deseada
- una modificación del medio de propagación en sí mismo

Esta definición queda englobada dentro de la anterior: perturbación es todo fenómeno electromagnético que podría afectar a tu equipo. Si no afecta, se queda en perturbación. Si afecta, tendrá otro nombre: **interferencia**.

Interferencia electromagnética (electromagnetic interference, EMI)

Cualquier anomalía aportada al funcionamiento de un dispositivo, de un aparato o de un sistema por una perturbación electromagnética. Hay que definir qué es anomalía, y ahí el fabricante tiene mucho que decir. Por ejemplo, que la salida de audio de un receptor de radio incorpore un tono de 200 Hz 90 dB por debajo del nivel medio de la señal, no será considerado una anomalía. Que se interrumpa el audio durante medio segundo, sí. Las anomalías pueden dar lugar a una **degradación de funcionamiento**.

Degradación de funcionamiento

Desviación no deseada de las características de funcionamiento de un dispositivo, aparato o sistema en presencia de una interferencia electromagnética. Lo que puedes detectar en un ensayo de inmunidad es una degradación de funcionamiento. Podrá ser aceptable o no, lo que implicará que el resultado del ensayo es positivo o negativo. Como fabricante, puedes definir en el manual de usuario lo que consideras que es una degradación de funcionamiento aceptable. Otra cosa es lo que piense el usuario. Pero es lícito que escribas en el manual de usuario que la proximidad de una emisora de radio FM u otros radiadores similares pueda provocar una degradación de la calidad de audio en un *walkie-talkie* de juguete y que consideras que es aceptable.

Compatibilidad Electromagnética (EMC)

Según IEC 61000-1-1, se define como la capacidad de cualquier aparato, equipo o sistema para funcionar satisfactoriamente en su entorno electromagnético, sin producir perturbaciones electromagnéticas sobre cualquier cosa de ese entorno:

- **Funcionar satisfactoriamente**, implica que el equipo es tolerante con otros equipos, siendo poco susceptible (o suficientemente inmune) a señales electromagnéticas que otros equipos ponen en el ambiente
- **No producir perturbaciones electromagnéticas** intolerables, de manera que la emisión de señales electromagnéticas por parte del propio sistema no origine problemas de EMI con otros equipos

La compatibilidad electromagnética debe ocuparse de tres problemas diferentes, describiendo la capacidad de:

- Los sistemas eléctricos y electrónicos para funcionar sin interferir en otros equipos
- Que dichos sistemas deben funcionar como deben en un entorno electromagnético específico
- No provocar interferencias consigo mismo (por ejemplo, un regulador conmutado en una esquina del PCB inyectando ruido en un circuito de video en la esquina opuesta)

Una EMC eficaz requiere un sistema diseñado, fabricado y comprobado según el entorno electromagnético al que se destina.

Immunidad Electromagnética

El aparato, equipo o sistema debe ser capaz de operar adecuadamente en ese entorno sin ser interferido por otros o por él mismo. La inmunidad indica hasta qué punto puede contaminarse electromagnéticamente el entorno antes de que afecte de manera adversa a un equipo.

Susceptibilidad Electromagnética

Es complementario a la inmunidad. Indica la incapacidad de un dispositivo, equipo o sistema de funcionar sin degradación en presencia de una perturbación electromagnética.

Emisión Electromagnética

Fenómeno por el cual la energía electromagnética emana de una fuente. La energía electromagnética puede alcanzar a un sistema:

- Por radiación (interferencias radiadas)
- Por conducción (interferencias conducidas)

Estas dos formas de interacción o acoplamiento pueden estar presentes simultáneamente y pueden dar lugar a efectos no deseados dentro y fuera del sistema que contiene las fuentes de perturbación.

Emisión Radiada

Es la componente de energía de radiofrecuencia transmitida a través de un medio en forma de campo electromagnético.

Emisión Conducida

Es la componente transmitida a través de un medio físico como un cable o hilo.

Descarga electrostática (Electrostatic Discharge, ESD)

Transferencia de electricidad en el que intervienen dos cuerpos con diferente nivel electrostático. En otras palabras, el desagradable pero breve chispazo que recibes en ocasiones cuando bajas del coche, descienes por un tobogán de plástico o das la mano a otra persona. Este fenómeno tiene capacidad destructiva sobre los circuitos integrados y someter los equipos a descargas electrostáticas es una de las pruebas que hay que superar para cumplir con la normativa EMC Europea.

Solución a los problemas EMC

No hay otra solución que una de estas tres: reducir la interferencia en origen, reducir el acoplamiento o aumentar la inmunidad del equipo víctima.

Cuando ves un cable de alimentación con un engrosamiento en un extremo, es que se ha añadido una ferrita para atenuar ruido de frecuencias medias y/o altas: se ha resuelto el problema reduciendo el acoplamiento.

Cuando tardas tres días más en entregar un diseño rutado de lo que esperaba tu jefe, posiblemente sea porque has buscado reducir la interferencia en origen, evitando emisiones por antenas parásitas, añadiendo terminaciones de línea y eliminando discontinuidades en los caminos de retorno de la señal.

Cuando añades diodos TVS para proteger tus entradas/salidas de descargas ESD, has buscado aumentar la inmunidad de tu equipo.

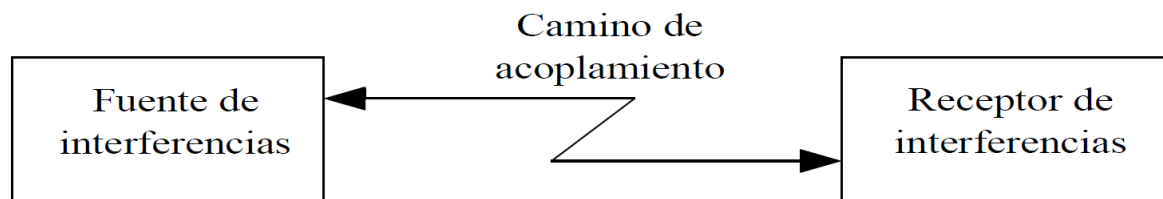


Figura 9.1. Planteamiento clásico de la solución a un problema EMC: reducir la interferencia en origen, reducir el acoplamiento o aumentar la inmunidad del equipo víctima.

Ante cualquier problema de EMC debes identificar estos tres actores: fuente, camino y víctima. El siguiente paso es decidir en cuál será más rentable invertir esfuerzos y coste. Los especialistas en EMC van equipados con cosas tales como papel de aluminio, cinta de cobre, un surtido de ferritas para cables, malla metálica, antenas de campo cercano, un analizador de espectros de mano y una sonda de corriente. Con todo esto evalúan la causa del problema y la posible solución. La experiencia es su activo más valioso.

Si estás interesado en toda esta parafernalia, en [24] puedes encontrar algo interesante que leer. Ármate de este equipo (o tanto como te permita tu presupuesto) y aplica unos pocos principios básicos para entender el problema y ensayar una solución: a veces basta con apantallar (rodear el equipo de papel de aluminio), reducir imperfecciones en el apantallamiento (tapar ranuras y juntas con cinta de cobre), atenuar ruido en los cables (añadir ferritas), sustituir un componente muy ruidoso (que detectas con sondas de campo cercano), conectar a chasis una pantalla de un cable mal conectada (lo que podrás descubrir con tus propios ojos o con un multímetro) o descubrir puntos sensibles a descargas ESD empleando un encendedor piezoeléctrico de cocina.

Diseñando para EMC

El diseño electrónico para EMC puede enfocarse desde dos puntos de vista conocidos comúnmente como:

- **Crisis approach:** el diseñador se despreocupa de la EMC hasta que ha finalizado la fase de pruebas funcionales del producto, o peor aún, cuando el equipo falla en las pruebas EMC, o peor todavía, cuando llegan las primeras quejas de clientes. Lamentablemente, es la forma habitual de trabajar en demasiadas empresas.
- **Systems approach:** este enfoque considera la EMC desde la fase de desarrollo conceptual del producto. El diseñador se anticipa a los problemas de EMC antes de empezar el diseño, implementa soluciones a dichos problemas y testea los prototipos tan a fondo como sea posible.

Con este segundo enfoque la EMC llega a ser una parte integral de la ingeniería electrónica y eléctrica, la ingeniería mecánica y el desarrollo de *firmware/software*, pues generalmente las soluciones requieren el concurso de varias disciplinas. La EMC forma parte del diseño y no es un costoso añadido al final de proceso de desarrollo del producto.

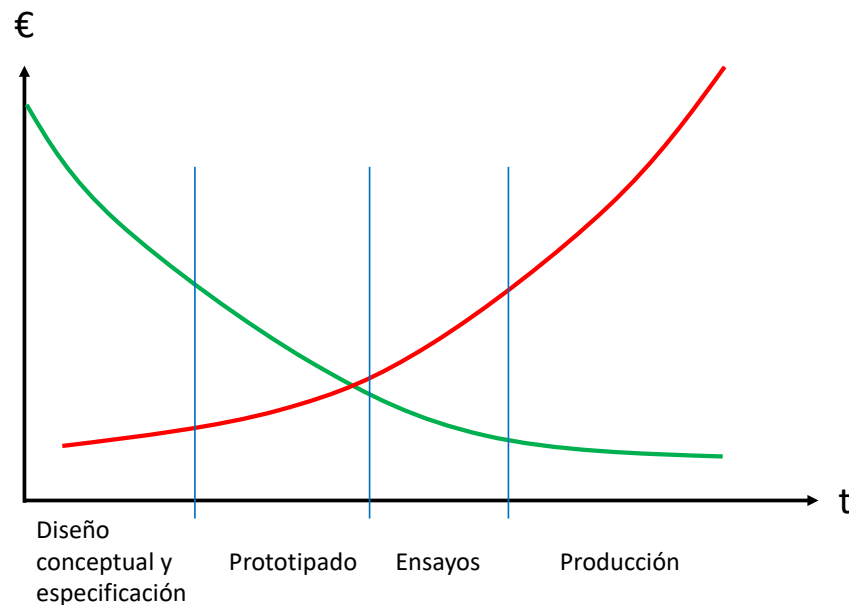


Figura 9.2. Crisis approach (curva roja) frente a systems approach (curva verde)

El problema no es sólo de voluntad. Para integrar la EMC en el proceso de diseño la empresa necesita ingenieros y técnicos con la formación adecuada. En los últimos años ha crecido la demanda de perfiles con conocimientos en EMC en España (y presumo que será igual en otros países). Una posible razón es la internacionalización de las empresas de diseño y producción electrónica, que se ven forzadas a cumplir con exigencias y metodologías de trabajo que llevan tiempo considerando la EMC como una parte más del diseño y no una ruleta rusa a la hora de superar los ensayos EMC. Por experiencia propia, colaborando con varias empresas en el desarrollo de nuevos productos, a menudo basta con aplicar buenas prácticas durante la fase de diseño para superar los ensayos EMC sin problemas. También es cierto que hay sectores más duros (como el ferroviario y aeronáutico) en los que esquivar los problemas sólo con buenas prácticas en fase de diseño es más complicado. ¡Para qué vamos a engañarnos!

Marco normativo

Lo siento. Esto será al menos tan aburrido de leer para ti como lo es de escribir para mí. Consolémonos pensando que tenemos que pasar juntos por este trance. Vamos allá...

La [Directiva 89/336/EEC](#), que como hemos comentado tuvo plena vigencia desde 1996, no es el marco normativo actual. Desde esa fecha se han producido los siguientes cambios:

- En el año 2000 se publica la R&TTED (*Radio and Telecomm Terminal Equipment Directive*) 1999/5/EC. Reemplaza (en la práctica amplía) a la directiva EMC en equipos de radio y equipos terminales de telecomunicación (definido como producto o componente de este destinado a su conexión directa o indirecta por medio de interfaces de redes públicas de telecomunicación). En esencia, mantiene los requisitos de la directiva EMC y añade otros específicos.
- En 2004 se publica la [segunda directiva EMC](#) (2004/108/EC). Con el cambio, ya no es necesaria la aprobación del documento técnico de construcción por parte de un organismo notificado si optamos por esta vía (es opcional). Se clarifica la diferencia entre aparatos e instalaciones fijas, definiendo los requisitos que se impone a estas últimas.
- En 2014 nace la [tercera directiva EMC](#) (2014/30/EU). Ahora los equipos terminales de telecomunicación vuelven a estar bajo la directiva EMC. Los receptores y cualquier otro equipo con antena (incluso un receptor GPS) ahora están bajo la RED.
- También en 2014, se publica la [RED](#) (*Radio Equipment Directive*, 2014/53/EU), de obligado cumplimiento desde junio de 2017. Sustituye a la R&TTED del año 2000. Afecta a cualquier equipo que de forma intencionada o no intencionada emita o reciba RF a través de una antena (radar, RFID, Wifi, BT, NFC, receptores GPS, etc.) ¿Qué requisitos deben cumplir los equipos bajo el ámbito de la RED?
 - Cumplir la Directiva de Baja Tensión -*Low Voltage Directive*, LVD-, eliminando las excepciones por baja tensión (50V AC, 75V DC). Estas excepciones permitían, por ejemplo, no tener que someter a ensayo de LVD a equipos alimentados por baterías de 12V
 - Cumplir la Directiva EMC (a través de la familia de normas [EN 301 489](#))
 - Incorporar los requisitos de seguridad de las Directivas LVD y EMC
 - Cumplir las normas de uso eficiente del espectro (es decir, no radiar de forma apreciable en otros canales)
 - En casos especiales, otros requisitos

En este [enlace](#) [25] encontrarás un resumen de la adaptación de la Directiva 2014/53/EU a la legislación española.

Aclaremos un poco lo anterior

Cualquier producto electrónico que tenga una antena (WiFi, BLE, lo que sea), aunque sólo sea para recepción (como en el caso de GPS) es capaz de radiar, aunque sea de forma no intencionada, y por tanto está sujeto a la Directiva de Equipos de Radio (RED), que va más allá de la Directiva EMC (EMCD). En el ámbito TIC (tecnología de la información y de las comunicaciones) gran parte de los equipos tienen interfaces inalámbricos, por lo que será necesario conocer bien no sólo la EMCD, sino la RED.

Sólo aquellos equipos desprovistos de antena quedan sujetos a la Directiva EMC.

Directiva europea de EMC

Aparato

Equipo terminado o combinación de equipos comercialmente disponible como una única unidad funcional, destinada a un usuario funcional. Son excepciones: equipos de radioaficionado, equipos de radio, equipos para aeronaves, equipos para demostración en ferias, y otros tipos de productos con directivas propias (equipos de radio, médicos, automoción, marino, tractores, etc.)

Componente (no afectado por la Directiva) vs aparato

La distinción está en si va destinado a un usuario final o no. En el caso de un componente, el integrador es el responsable de que el aparato final cumpla la directiva EMC. Cabe la duda de qué pasa con un componente que pueda ser ambas cosas, como una tarjeta gráfica para PC. Si se vende a un cliente final en una tienda, la tarjeta debe haber superado los ensayos EMC (forma más habitual de demostrar cumplimiento de la directiva).

Instalación fija vs aparato

Una instalación fija es un sistema que contiene varios tipos de aparatos, que es fabricado, instalado y operado de forma definitiva en una ubicación predefinida. Como la instalación no se puede mover a un laboratorio de ensayo EMC, la instalación debe realizarse con buenas prácticas de ingeniería EMC que deben ser documentadas y mantenidas por una persona responsable. Lo normal es seguir una de estas dos estrategias: (1) test de EMC limitados y/o (2) usar aparatos con marcado CE e instalarlos siguiendo las indicaciones de los fabricantes. La instalación fija no necesita llevar marcado CE ni es obligatorio un certificado de conformidad.

Sistema vs aparato

Al contrario que una instalación fija, un sistema, que contiene varios aparatos, sí está concebido para cambiar de ubicación. Y, por tanto, es susceptible de ser sometido a ensayo en un laboratorio de EMC. Por ejemplo, un PC compuesto de monitor, teclado, ratón y torre es un sistema. Debe ser sometido a test y certificado como un conjunto.

La Directiva EMC Europea se suma a la estrategia que ha encontrado la Unión para evitar fracasar en la aprobación de sus normas por parte de los Estados miembros, que requiere unanimidad: recoger en el texto sólo los requisitos esenciales y dejar para otros reglamentos el detalle fino. De este modo, quien lea las pocas páginas de la Directiva quedará sorprendido al ver que ésta se limita a enunciar que los requisitos esenciales de la directiva EMC para un aparato son:

- Las perturbaciones electromagnéticas que genera no deben exceder un nivel que produzcan degradación de prestaciones en otros aparatos
- Las perturbaciones recibidas en un uso normal (entorno electromagnético) no deben degradar inaceptablemente las prestaciones
- Las condiciones anteriores deben demostrarse mediante una evaluación de EMC del aparato

No se explicita cómo y con qué criterios se hace esta evaluación. La Directiva EMC es aplicable a:

- Nuevos productos europeos (y se refiere a cada unidad individual) que vayan a ser puestos en el mercado europeo (venta del fabricante a un distribuidor)
- Productos nuevos y existentes de terceros países (y se refiere a cada unidad individual) que vayan a ser puestos en el mercado europeo. El importador es el responsable de que el producto cumpla con la Directiva

La Directiva no es aplicable a productos reacondicionados, pero sí a los mejorados. Por lo tanto, una vez puesto en el mercado o en servicio, la regulación europea sobre EMC ya no es aplicable al aparato: [la regulación sólo busca evitar la entrada en el mercado de productos que no cumplan la directiva.](#)

Marcado CE

Todo producto comercializado o puesto en servicio dentro de la Unión Europea debe llevar marcado CE, siempre que una Directiva Europea le sea de aplicación. Un producto ha de cumplir con los requisitos esenciales de todas las Directivas que le sean de aplicación para poder llevar el marcado CE (Figura 9.3 izquierda).

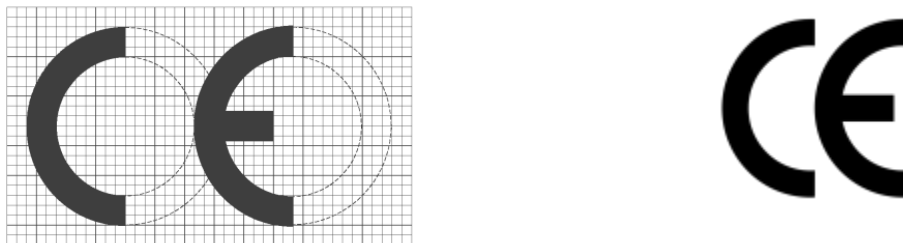


Figura 9.3. ¡Cuidado! No confundas el diseño de la marca CE (derecha) con el símbolo “China Export” (izquierda)

El marcado CE de un producto constituye una declaración por parte de la persona física o jurídica que lo ha colocado de que el producto se ajusta a las disposiciones comunitarias y que se han realizado los procedimientos pertinentes de evaluación de conformidad.

El marcado CE debe ser colocado por el fabricante o su representante autorizado y debe tener al menos 5 mm de alto. Debe colocarse en un lugar visible del producto o placa de características (marcado de forma visible, legible e indeleble). Como excepción, si no es posible hacerlo en el producto, se hace en el embalaje o en la documentación que lo acompaña. Así que cuando veas un producto con marcado CE ya sabes lo que significa: “este producto cumple con todas las directivas europeas que le son de aplicación”.

Declaración de conformidad

Además del marcado CE, el fabricante o importador debe emitir una **Declaración de Conformidad (DdC)**, que contiene (Figura 9.4):

1. Una referencia a las Directivas aplicables
2. La identificación del aparato
3. El nombre y la dirección del fabricante y/o representante autorizado en la Comunidad Europea
4. Una referencia fechada a las especificaciones respecto a las cuales se declara la conformidad
5. Fecha de la declaración
6. Identidad y firma de la persona facultada para comprometer al fabricante o su representante autorizado

Entonces, si un producto lleva marcado CE y viene acompañado de una DdC, ¿podemos tener la seguridad de que realmente cumple con la normativa? Un estudio de hace unos años concluía que hasta un 40% de los convertidores de potencia que se vendían en el mercado europeo no cumplían con los requisitos. Ten en cuenta que ningún regulador u organismo nacional o europeo te piden que demuestres nada. Si careces de suficientes escrúpulos, mientes en la DdC y pones el marcado CE. No serías el primero. Pero te ruego que no lo hagas.

Los reguladores nacionales, con sus escasos recursos, comprueban productos comerciales al azar, o actúan bajo sospecha o denuncia. Si un regulador nacional descubre que has mentido, puede prohibir la comercialización en ese país y advertirá al resto de países, que podrán hacer lo mismo. Si eres una pequeña empresa, te enfrentas a una catástrofe económica. Si eres más grande, puedes sufrir daños en tu reputación. En cualquier caso, te enfrentarás a una multa. En este escenario, la tentación de recortar costes y plazos puede ser grande. Pero el fabricante/importador debe ser responsable y no engañar a sus clientes; a largo plazo no es una buena estrategia.

Este es el modelo propuesto en la EN 17050 mencionada anteriormente.

MODELO RECOMENDADO DE DECLARACIÓN DE CONFORMIDAD EN 17050	
DECLARACIÓN DE CONFORMIDAD	
Nombre del emisor:	☐
Dirección:	
Declaramos bajo nuestra exclusiva responsabilidad la conformidad del producto:	☐
(nombre del aparato, marca, modelo, fabricante)	
al que se refiere esta declaración, con la(s) norma(s) u otro(s) documento(s) normativo(s)	☐
(título y/o número y fecha de la(s) norma(s) u otro(s) documento(s) normativo(s))	
Información adicional:	☐
de acuerdo con las disposiciones de la Directiva 99/05/CE, del Parlamento Europeo y del Consejo de 9 de marzo de 1999, (transpuesta a la legislación española mediante el Real Decreto 1890/2000, de 20 de noviembre de 2000).	
Lugar y fecha de emisión.	☐
Firmado por:	☐
Firma.	

Figura 9.4. Elementos que deben aparecer obligatoriamente en la Declaración de Conformidad

Declaración de conformidad CE	
<i>Fabricante: Brother Industries Ltd., 15-1, Naeshiro-cho, Mizuho-ku, Nagoya 467-8561, Japón</i>	
<i>Planta: Brother Technology (Shenzhen) Ltd., NO6 Gold Garden Ind. Nanling Buji, Longgang, Shenzhen, China</i>	
<i>Declaramos que:</i>	
<i>Descripción de los productos:</i>	<i>Impresora láser</i>
<i>Nombre de producto:</i>	<i>HL-2035</i>
<i>Número de modelo:</i>	<i>HL-20</i>
<i>cumple las disposiciones de las directivas aplicadas: Directiva de baja tensión 2006/95/EEC y directiva de compatibilidad electromagnética 2004/108/EC.</i>	

Estándares armonizados que se aplican:

Seguridad *EN60950-1:2001*

EMC EN55022: 1998 + A1: 2000 +A2: 2003 Clase B

EN55024: 1998 + A1: 2001 +A2: 2003

EN61000-3-2: 2006

EN61000-3-3: 1995 + A1: 2001 +A2: 2005

Año en el que se aplicó por primera vez la marca CE: 2008

Expedido por: Brother Industries, Ltd.

Fecha: 6 de febrero de 2008

Lugar: Nagoya, Japón

Firma: _____

Junji Shiota

Administrador general

Quality Management Dept.

Printing & Solutions Company

Figura 9.5. Ejemplo de DdC. (http://support.brother.com/g/s/id/html/doc/prINTER/hl2035/spa/html/ug/appendix12_2_7.html)

La DdC de la Figura 9.5 es sólo un ejemplo que he encontrado en internet. Fíjate en que al final hay un alto directivo de la empresa que firma y se hace responsable. Declara, entre otros, que cumple con los estándares EN55022:1998 con adendas de 2000 y 2003, y EN55024:1998, con adendas de 2001 y 2003, que son aplicables a equipos de tecnologías de la información, lo que es claramente aplicable a una impresora. También declara conformidad con EN61000-3-2 y EN61000-3-3 (emisiones de armónicos y otras perturbaciones en la red AC).

Por supuesto, estas normas ya están obsoletas:

- EN55022 ha sido reemplazada por EN55032, que en la actual versión de esta obra fue revisada por última vez en 2016.
- EN55024 fue revisada en 2011, con una adenda en 2015 (es la versión más reciente hasta el momento)

No hace falta decir que debes usar siempre la versión más reciente de las normas. Por cierto, las normas “EN” son estándares armonizados. Vamos a ver qué es esto.

Directiva de Equipos de Radio (RED)

Si vas a comercializar un equipo sometido a la RED, lo primero que debes hacer es leerte el Real Decreto 188/2016, de 6 de mayo, por el que se aprueba el Reglamento por el que se establecen los requisitos para la comercialización, puesta en servicio y uso de equipos radioeléctricos, y se regula el procedimiento para la evaluación de la conformidad, la vigilancia del mercado y el régimen sancionador de los equipos de telecomunicación. Ánimo. Lo que sigue es un breve (y por supuesto incompleto) resumen.

Equipos sometidos a la RED

El ámbito de aplicación son los **equipos radioeléctricos**, exceptuando equipos y kits de radioaficionado, equipos marinos y para aeronaves, kits para profesionales para I+D y equipos radioeléctricos para defensa, seguridad pública y del Estado.

Por “equipo radioeléctrico” entendemos aquel que emite o recibe intencionadamente, o que necesite como accesorio una antena para emitir o recibir. ¿Por qué cubre esta directiva a receptores, como puede ser un receptor **GPS**? Porque si acoplas (por supuesto no intencionadamente) ruido a la antena, estás emitiendo. Productos con directiva propia, como los juguetes, si disponen de radios, están también sujetos a la RED.

¿Cómo demostrar la conformidad con la RED?

Al igual que ocurre con la Directiva EMC, la conformidad con las normas armonizadas supondrá la presunción de conformidad con la RED. Lo habitual es seguir la vía definida como **control interno de la producción** (Módulo A en la Figura 9.6), en la que el fabricante (el importador es responsable de que el fabricante lo haya hecho, si el equipo proviene de fuera de la UE) pasa, sobre una o varias unidades de muestra, los ensayos que marcan las normas en un laboratorio externo y obtiene un informe con los resultados. Emite una **declaración responsable del fabricante**, quien debe custodiar toda la documentación relevante (esquemas, planos, resultados de los ensayos, software utilizado en los ensayos, etc.) durante un periodo de diez años. Periódicamente, el fabricante debería someter un número de unidades a ensayo para verificar que se mantiene la calidad de producción.

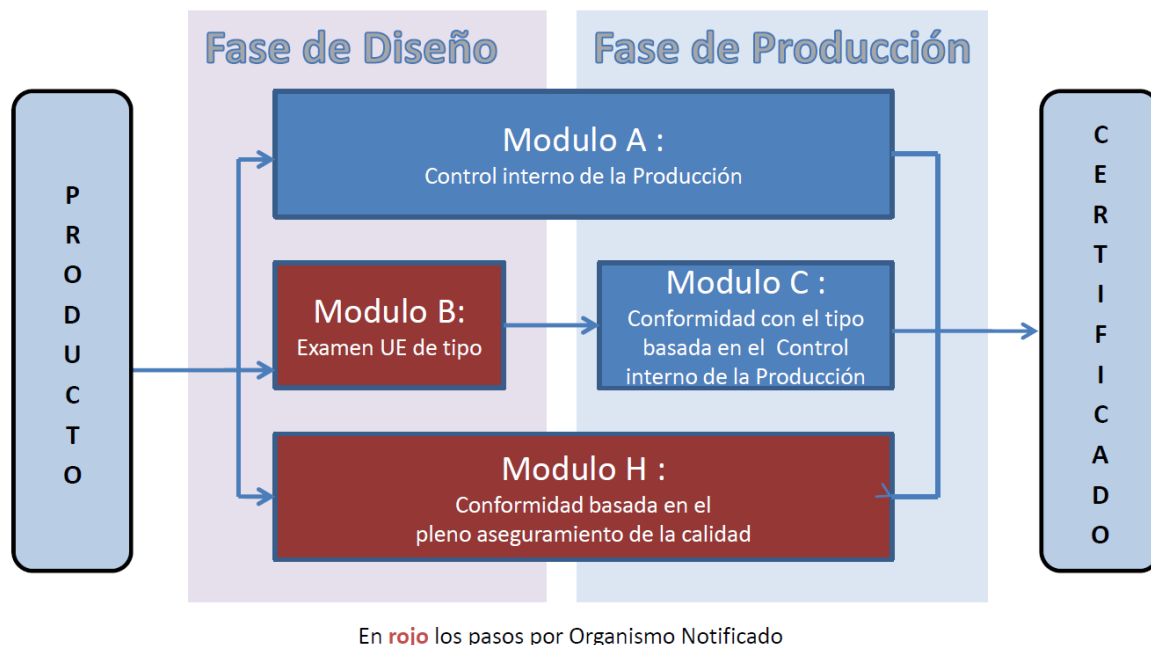


Figura 9.6. Evaluación de la conformidad según la RED. Figura extraída de [25]

Para que un equipo sometido a la RED pueda comercializarse en España, debe cumplir con el CNAF (cuadro nacional de atribución de frecuencias) y con la regulación española en cuanto a potencia, uso, etc. Puede ocurrir que un equipo cumpla con estos requisitos especiales de un país, pero no de otro: en ese caso debe indicarse de forma visible y clara en el embalaje y en las instrucciones. La forma de hacerlo está recogida en el BOE ([enlace](#)).

Hay un registro obligatorio para ciertas categorías de equipos. Citando al BOE: “*Antes de la introducción en el mercado, los fabricantes registrarán en el sistema central de registro que la Comisión Europea pondrá a disposición de los fabricantes los tipos de determinadas categorías de equipos radioeléctricos que presenten un bajo nivel de conformidad con los requisitos esenciales. La obligación de registro será de aplicación a partir del día 12 de junio de 2018.*”

Declaración de conformidad

Al igual que con la EMCD, el fabricante emite una DdC cuyo modelo recomendado se recoge en la Figura 9.7.

DECLARACION DE CONFORMIDAD	
Nombre del fabricante:	
Dirección: [Identificación del declarante (nombre, domicilio, grado de representación del fabricante)]	
Declaramos bajo nuestra exclusiva responsabilidad la conformidad del producto: [Nombre del equipo radioeléctrico, marca, modelo, fabricante, país de fabricación, número de lote o de serie, en su caso, procedencia y número de ejemplares] [fotografía en color del equipo radioeléctrico]	
al que se refiere esta declaración, con la(s) norma(s) u otros documento(s) normativo(s) [Título y/o número y fecha de la(s) norma(s) u otro(s) documento(s) normativo(s)]	
de acuerdo con las disposiciones de la Directiva 2014/53/UE del Parlamento Europeo y del Consejo de 16 de abril de 2014, así como de las disposiciones (otra legislación de armonización de la Unión Europea).	
(Si procede): El organismo notificado (nombre, número) ha efectuado (descripción de la intervención) y expedido el Certificado de examen UE de tipo:	
(Si procede): Accesorios (incluida la referencia al software) que permiten que el equipo declarado funcione como está previsto y sin alterar la evaluación de los requisitos esenciales aplicables al mismo.	
Información adicional (si es preciso):	
Firmado en nombre de:	
(nombre, cargo, firma)	Lugar, fecha de expedición

Figura 9.7. Modelo recomendado para la declaración de conformidad de la RED. Fuente: [BOE](#)

Enlaces para más información sobre la RED

DIRECTIVA: <http://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex:32014L0053>

REAL DECRETO 188/2016: <https://www.boe.es/boe/dias/2016/05/10/pdfs/BOE-A-2016-4444.pdf>

REGLAMENTO DE EJECUCIÓN 2017/1354: <http://www.boe.es/doue/2017/190/L00007-00010.pdf>

LISTADO DE NORMAS ARMONIZADAS: https://ec.europa.eu/growth/single-market/european-standards/harmonised-standards/rte_es

GUIA DE APLICACIÓN DE LA RED: <http://ec.europa.eu/docsroom/documents/23321>

FAQ de la RED: <http://ec.europa.eu/DocsRoom/documents/24921>

USO DEL ESPECTRO EN PAÍSES EUROPEOS: <http://www.efis.dk/>

Estándares armonizados

La Directiva EMC sólo fija los requisitos esenciales:

- Las perturbaciones electromagnéticas que genera no deben exceder un nivel que produzcan degradación de prestaciones en otros aparatos
- Las perturbaciones recibidas en un uso normal (entorno electromagnético) no deben degradar inaceptablemente las prestaciones
- Las condiciones anteriores deben demostrarse mediante una evaluación de EMC del aparato

Un fabricante o importador debe demostrar que cumple con estos requisitos. Una forma de hacerlo, de hecho, la más difundida, aunque no la única, es **demostrar el cumplimiento de estándares armonizados derivados de organizaciones europeas (CEN, CENELEC, ETSI)**. Por ejemplo, cumplir con la norma genérica *UNE-EN 61000-6-1:2007 Compatibilidad electromagnética (CEM). Parte 6-1: Normas genéricas. Inmunidad en entornos residenciales, comerciales y de industria ligera. (IEC 61000-6-1:2005)* servirá para demostrar que el producto cumple con el segundo requisito, siempre que se trate de un producto para uso residencial, comercial o de industria ligera y que no haya una norma específica de familia de producto.

Vamos a aclarar qué es esto de las normas genéricas, de familia de producto y qué son las normas básicas.

- **Normas Básicas:** especifican métodos de medida, instrumentación a utilizar, límites máximos de medición. Es decir, se trata de estandarizar el entorno en el que se realiza el ensayo y la forma de aplicar estímulos o de realizar medidas de forma que en España y en Australia el mismo ensayo proporcione los mismos resultados. No dice qué nivel de emisiones máximo puede tener un producto. Pero dice qué antena receptora utilizar, a qué distancia y cómo debe ser la sala donde se realice el ensayo. Las normas básicas proceden de IEC (International Electrotechnical Commission) o de CISPR (International Special Committee on Radio Interference). Este es un listado de las normas más básicas más habituales IEC armonizadas como normas europeas (*european norm*, EN):

EN 61000-3-2 Armónicos en la red de alimentación AC

EN 61000-3-3 Fluctuaciones de tensión y *flicker* en la red de alimentación AC

EN 61000-4-1 Visión general ensayos inmunidad

EN 61000-4-2 Descargas electrostáticas (ESD)

EN 61000-4-3 Campos electromagnéticos de alta frecuencia radiados

EN 61000-4-4 Transitorios eléctricos rápidos (EFT)

EN 61000-4-5 Impulsos de alta energía (ondas de choque)

EN 61000-4-6 Campos electromagnéticos de alta frecuencia conducidos

EN 61000-4-7 Armónicos e inter-armónicos

EN 61000-4-8 Campos magnéticos a frecuencia de red

EN 61000-4-9 Campos magnéticos pulsados

EN 61000-4-10 Campos magnéticos amortiguados

EN 61000-4-11 Fallos, fluctuaciones, cortes y micro cortes en la alimentación AC

EN 61000-4-12 Ondas amortiguadas en la alimentación

- **Normas de Producto:** Describen requisitos específicos de una familia de productos. Hay muchos tipos de producto que tienen normas específicas, como equipos de tecnología de la información, equipos de iluminación, electrodomésticos, equipos médicos, etc. Por ejemplo, los equipos de tecnología de la información están sujetos a las normas:

EN-55024: Equipos de tecnología de la información. Características de inmunidad

EN-55032: Compatibilidad electromagnética de equipos multimedia. Requisitos de emisión. Se trata de la norma CISPR-32 armonizada en Europa.

- **Normas Genéricas:** Sólo se usan cuando no existe la norma de producto del equipo bajo ensayo. Las cuatro normas siguientes han sido armonizadas en Europa a partir de normas IEC 61000-6-x:

EN 61000-6-1:2002. Inmunidad en entornos Residencial, comercial e industria ligera.

EN 61000-6-2: 2002. Inmunidad en entornos Industriales.

EN 61000-6-3:2002. Emisión en entornos Residencial, comercial e industria Ligera.

EN 61000-6-4: 2002. Emisión en entornos Industriales.

La gran pregunta es, ¿qué normas EMC debe cumplir mi equipo? Si no tienes experiencia previa, tendrás que bucear en el catálogo de normas UNE y es fácil que acabes equivocándote. No hay una página web institucional o privada donde escribas “alarma para bicicletas” y te devuelva el listado de normas EMC a cumplir. De modo que tienes varias opciones:

- Contratas a un experto (que para eso están)
- Consultas al laboratorio EMC donde harás las medidas
- Miras la DdC de un par de productos de la competencia y obtienes gratis el listado de normas a cumplir. Después lo compruebas con el catálogo de normas UNE, las compras (son de pago) y las estudias

Mi recomendación, sobre todo si tu empresa está comenzado a dar sus primeros pasos en EMC, es seguir estas tres vías a la vez. En cualquier caso, es importante que alguna persona de tu empresa (dos sería mejor, para evitar un escenario de “*bus factor* = 1”) vayan desarrollando experiencia en normativa y ensayos EMC, y que participen en el desarrollo del producto en una fase temprana. Te voy a dar un ejemplo: si tu producto ha de ser ensayado a descargas ESD de 8 kV, estaría bien saberlo de antemano y añadir ya en fase de diseño las protecciones necesarias.

*(En informática, **bus factor** o factor autobús es un término usado en proyectos de desarrollo de software, que alude a una gran cantidad de información vital de un proyecto de software limitada solamente a uno o unos pocos desarrolladores, impidiendo la continuación del proyecto en el hipotético caso de que estos desarrolladores clave sean atropellados por un autobús) – Cita de Wikipedia*

Aclaración sobre alta, baja y media frecuencia en EMC

Independientemente de la familia de producto, hay un conjunto de ensayos que son comunes a todas las normas. En primer lugar, suele haber normas de emisiones y normas de inmunidad. Las primeras especifican los ensayos y límites de emisión, pero han de diferenciar entre emisiones conducidas y emisiones radiadas. Vale la pena dedicar un par de párrafos a hablar un poco sobre esto.

Para medir las emisiones radiadas debes colocar una antena, preferentemente en capo lejano ($d > \lambda/2\pi$), lo que para 10 MHz equivale a casi 5 metros. Esto implica que, a frecuencias bajas y medias, el tamaño de una cámara anecoica (con superficies absorbentes para evitar reflexiones) sea prohibitivo. Además, la eficiencia de un producto electrónico para radiar a frecuencias bajas y medias es muy pobre (un radiador de dimensiones inferiores a $\lambda/10$ será un mal radiador). Ambos factores conspiran para definir **HF** (*high frequency*, alta frecuencia) como aquella superior a 30 MHz ($\lambda=10$ m, por lo que un PCB de 10-20 cm de largo, equivalente a $\lambda/50$ o $\lambda/100$, será un pobre radiador), límite inferior para medir la energía radiada por un equipo con una antena.

A frecuencias medias (**MF**, *medium frequency*), entre 150 kHz y 30 MHz (longitudes de onda entre 2 km y 10 m), los radiadores más eficientes son los cables conectados a los equipos y no los equipos, y resulta más barato medir la corriente que circula por los cables y extrapolar a partir de esta medida la radiación equivalente. Hablamos entonces de emisiones conducidas (por los cables), pero no son más que emisiones radiadas medidas de forma indirecta y de forma más barata.

Hablamos de baja frecuencia (*low frequency*, **LF**) a la banda comprendida entre la frecuencia de red y 2 kHz (su armónico número 40). Esta banda no nos interesa desde el punto de vista de la radiación, sino desde el ruido propagado (como distorsión armónica y otros transitorios) en las líneas de alimentación AC. Por cierto, en español hablamos de armónicos de la frecuencia de red, sin hache, mientras que en inglés se habla de *mains harmonics*, con hache.

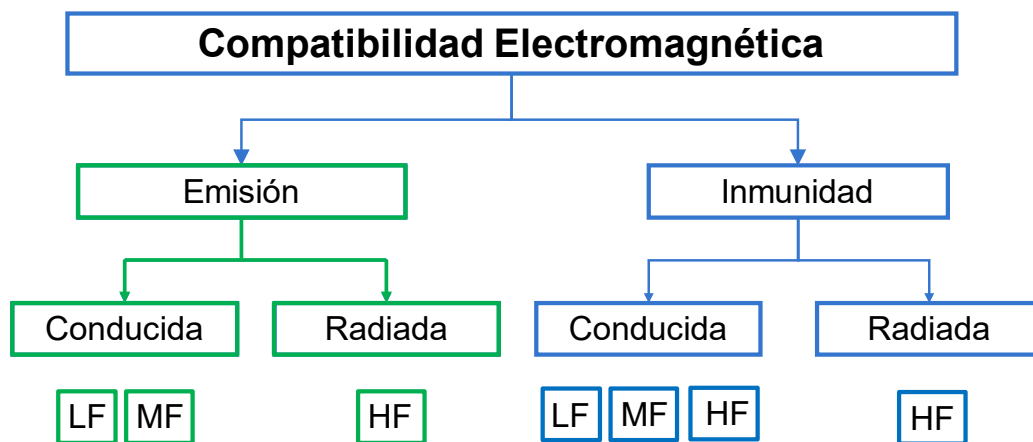


Figura 9.8. Es importante distinguir entre tres bandas de frecuencia (LF hasta 2 kHz, MF desde 150 kHz a 30 MHz, y HF a partir de 30 MHz), los ensayos radiados y conducidos y los ensayos de emisión y de inmunidad.

Es importante insistir en la relación entre el tamaño de un objeto y su capacidad de radiar. En el caso de un hilo eléctricamente corto, la potencia radiada será de $P=R_R \cdot I^2$, donde R_R es la resistencia de radiación y se calcula como:

$$R_R = 160 \cdot \pi^2 \cdot (l/\lambda)^2$$

La resistencia de radiación aumenta hasta aproximadamente $l=\lambda/2$ y luego comienza a presentar un patrón periódico. Una antena en $\lambda/4$ tiene una resistencia de radiación de casi 100 ohmios. Por ejemplo, a 30 MHz, un cable de 1 m de longitud tiene una resistencia de radiación de casi 16 ohmios, mientras que la longitud equivalente de un PCB (15 cm) tendría una resistencia de radiación de 0,35 ohmios, casi 45 veces inferior. Por eso, a frecuencias inferiores a 30 MHz, consideramos la contribución de los cables en la radiación y no consideramos la de los PCBs.

Hemos comentado la división en tres bandas de frecuencia principales en emisión. Pero ¿qué hay de los ensayos de inmunidad? **Debes tener presente el principio de reciprocidad: una estructura es tan buena antenna emitiendo como recibiendo.** Por lo tanto, si un PCB de 10 cm de largo es una aceptable antenna recibiendo a 300 MHz (su longitud eléctrica es cercana a $\lambda/10$) y podemos someterlo a un campo eléctrico desde una antenna en el interior de una cámara anecoica de dimensiones razonables, las cosas cambian si queremos someter el equipo a una interferencia de 3 MHz: el PCB será una muy mala antenna y sólo la energía acoplada en los cables conectados al equipo, que tendrán una longitud bastante superior al PCB, será relevante. De nuevo, es mucho más rentable inyectar una corriente eléctrica de 3 MHz en los cables, con un factor de conversión entre μA y $\text{dB}_{\mu\text{V/m}}$, y usar una instalación de medida de dimensiones razonables.

El resultado es la distinción entre ensayos de inmunidad radiada en HF, ensayos de inmunidad conducida en MF (que son realmente ensayos de inmunidad radiada indirectos) y ensayos de inmunidad conducida de LF en líneas de alimentación AC. Cabe hacer la puntualización de que, en inmunidad, en función del tipo de producto, también se ensaya inmunidad conducida en HF (Figura 9.6).

Ensayos de inmunidad

La Figura 9.7 resume los ensayos de inmunidad a los que se somete a un equipo (*device under test*, DUT, en la jerga EMC).

La envolvente y los puertos (conectores) del equipo son sometidos a **descargas electrostáticas (ESD)** aplicadas con una pistola ESD según IEC-61000-4-2. Este tipo de perturbaciones, de miles de voltios de amplitud, aunque de muy corta duración y por tanto de baja energía, son destructivas para la mayor parte de los circuitos integrados. Nuestro objetivo, como diseñadores, será añadir las protecciones necesarias para evitar la destrucción de hardware. Pero no podremos evitar que una entrada digital conmute o que una entrada analógica sufra un pulso corto (*glitch*), de modo que lo más habitual es que debemos incorporar protecciones a nivel de software para rechazar estos falsos pulsos.

Un tipo de perturbación similar en duración y energía, aunque de origen muy diferente, son las **ráfagas o EFT (*electrical fast transient*)**. Con esta perturbación (definida en IEC-61000-4-4) se pretende simular el efecto de conexiones y desconexiones de cargas inductivas (motores) en la línea de alimentación AC, por tanto, se ensaya en la entrada de alimentación AC del equipo. Como no es infrecuente llevar por la misma bandeja de cables líneas AC, DC y de señal, también debe introducirse esta perturbación en las líneas de señal y de alimentación DC, a menudo sólo si los cables superan los 3 metros de longitud (depende del tipo de producto). Las mismas protecciones que usamos para ESD sirven para EFT, lo que simplifica el diseño.

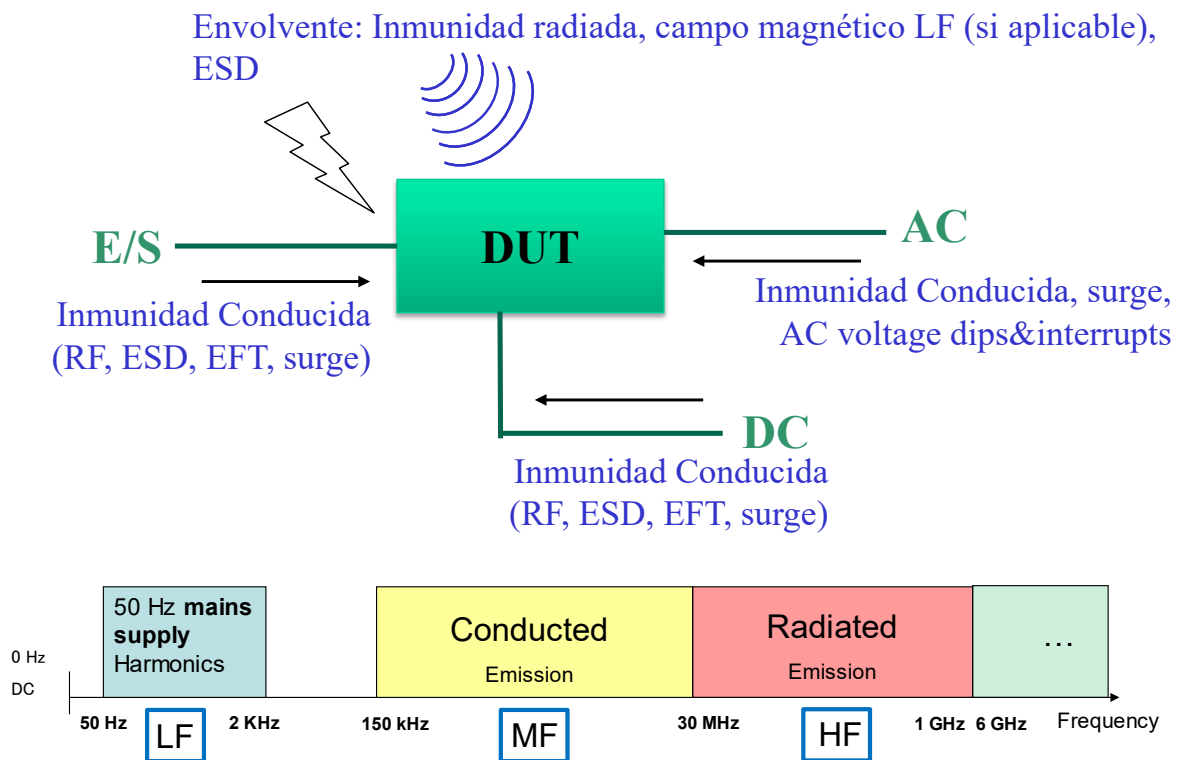


Figura 9.9. Resumen de los principales tipos de ensayo a los que el laboratorio de EMC someterá a tu equipo

ESD y EFT son perturbaciones pulsadas y de origen bien definido. No son continuas. Esto implica (de nuevo, depende de la familia de producto) que la norma pueda permitir que durante la perturbación sea aceptable una cierta degradación de funcionamiento. Por ejemplo, un monitor de ordenador podría parpadear durante una descarga ESD, recuperándose inmediatamente después. Esto sería, por lo general, aceptable por la norma.

Las **ondas de choque (surge)** emulan el efecto de la descarga de un rayo en las inmediaciones. Son perturbaciones de alta energía, con una alta capacidad destructiva, y requieren protecciones especiales (tubos de descargas de gas y varistores). La forma de onda de las descargas y su aplicación al DUT está definida en la norma IEC-61000-4-5. Las ondas de choque se aplican a las líneas de alimentación AC, así como a

líneas de señal y de alimentación DC externas. Por ejemplo, un bus RS-485 tendido entre una piscifactoría y una caseta de control en la playa es una línea externa, sujeta al efecto directo de la caída de un rayo. Una conexión Ethernet dentro de un edificio no está sujeta a este riesgo y por tanto no necesita ser sometida a este ensayo.

Los ensayos frente a **inmunidad conducida** en MF (generalmente de 150 kHz a 30 MHz, aunque depende de la familia de producto) implican inyectar en los cables conectados al equipo un tono de frecuencia variable, modulado en amplitud a 1 kHz, que supone una perturbación continua. ¿Por qué modular el tono a 1 kHz?: para facilitar la detección del efecto en el DUT. La perturbación penetra en el equipo, se rectifica en la primera unión PN que encuentra y es filtrada paso bajo por las capacidades que encuentra a su paso. De este modo, se produce una demodulación AM que se puede ver o escuchar como un tono de 1 kHz.

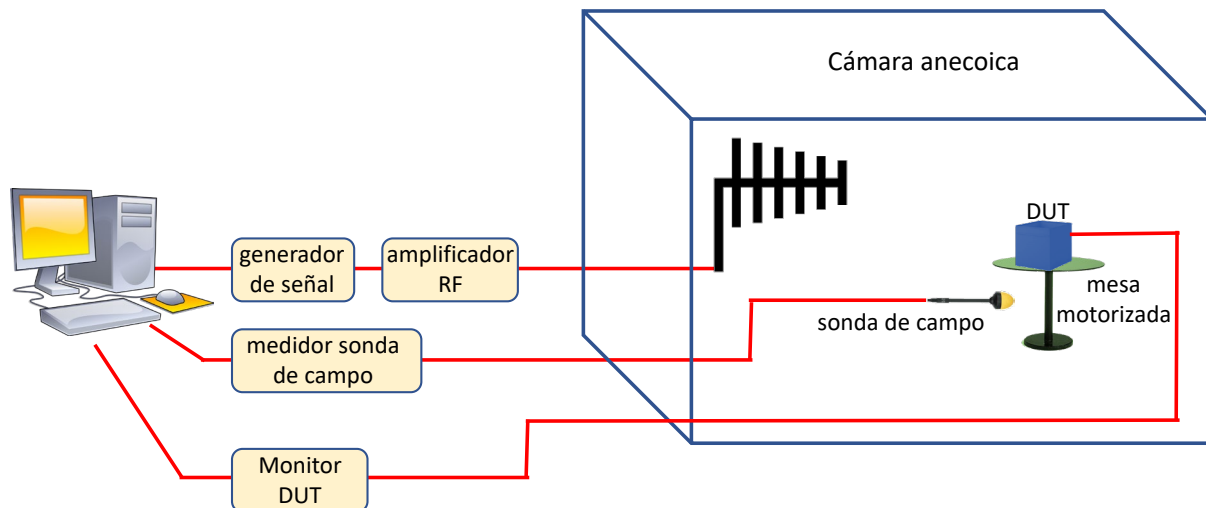


Figura 9.10. Ensayo de inmunidad radiada según IEC-61000-4-3, esquema simplificado

Equivalentemente, los ensayos de **inmunidad radiada** en HF implican el acoplamiento en los cables y en el equipo de una señal radiada, también modulada en amplitud con un tono de 1 kHz. Por ejemplo, IEC 61000-4-3 especifica la prueba de 80-1000 MHz, con señal modulada en AM al 80% 1 kHz. Se pueden dar incrementos de frecuencia en pasos del 1%, con una duración de al menos 0,5 s por paso, si bien el barrido puede ser más lento por las características del DUT y su software de prueba.

Para el ensayo es necesario un generador y amplificador de señal, así como un par antenas (para cubrir todo el margen de frecuencias), una sonda de campo para verificar la intensidad del campo en el DUT, una mesa motorizada para rotar 360° el DUT y un equipo (monitor del DUT) que verifica el correcto funcionamiento del DUT en cada frecuencia del ensayo.

Algunos equipos deben ensayarse también frente a **campos magnéticos de baja frecuencia (LF)**. Los equipos para uso en redes de distribución o instalaciones eléctricas deben ser ensayados en campo magnético (según norma IEC 61000-2-7, Figura 9.11). El resto sólo deben ensayarse en campo magnético a 50 Hz si contienen partes sensibles a campos magnéticos, según IEC-61000-4-8. En el ensayo, se rodea el equipo de un bucle que crea un campo uniforme (dentro de 3 dB) de 1-100 A/m (continuo) o 300-1000 A/m (corta duración).

Adicionalmente a las anteriores (perturbaciones pulsadas ESD, EFT y onda de choque), radiofrecuencia conducida y radiada y campo magnético de baja frecuencia), los equipos alimentados desde la red AC deben ser ensayados también frente a **bajadas de tensión e interrupciones en la línea AC**, definidos por IEC-61000-4-11 (Figura 9.12).

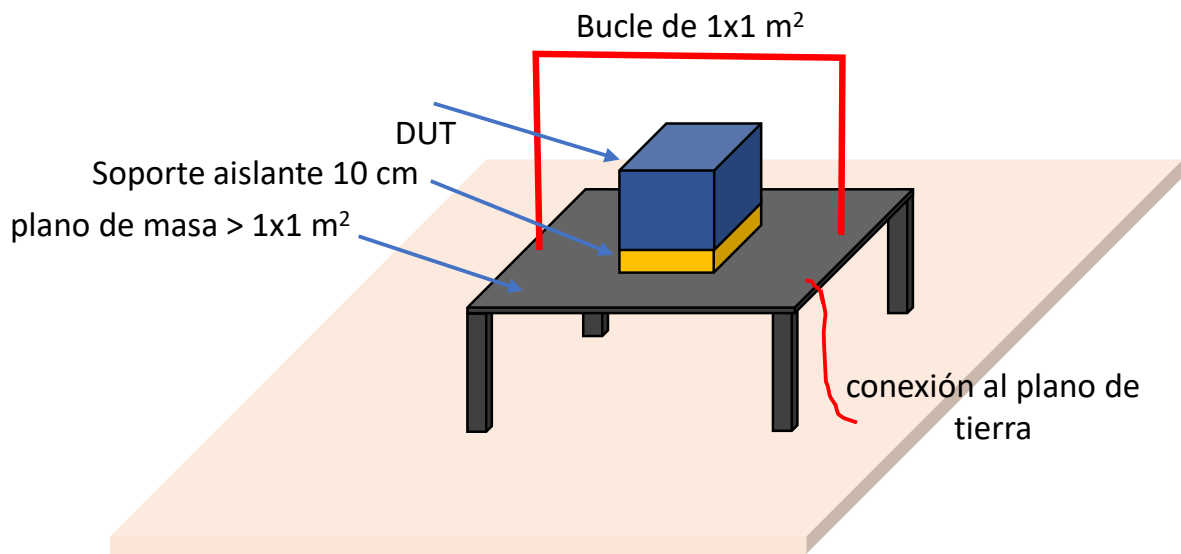


Figura 9.11. Ensayo de inmunidad a campos magnéticos LF. El bucle puede adoptar tres posiciones ortogonales (planos xy, xx, yz)

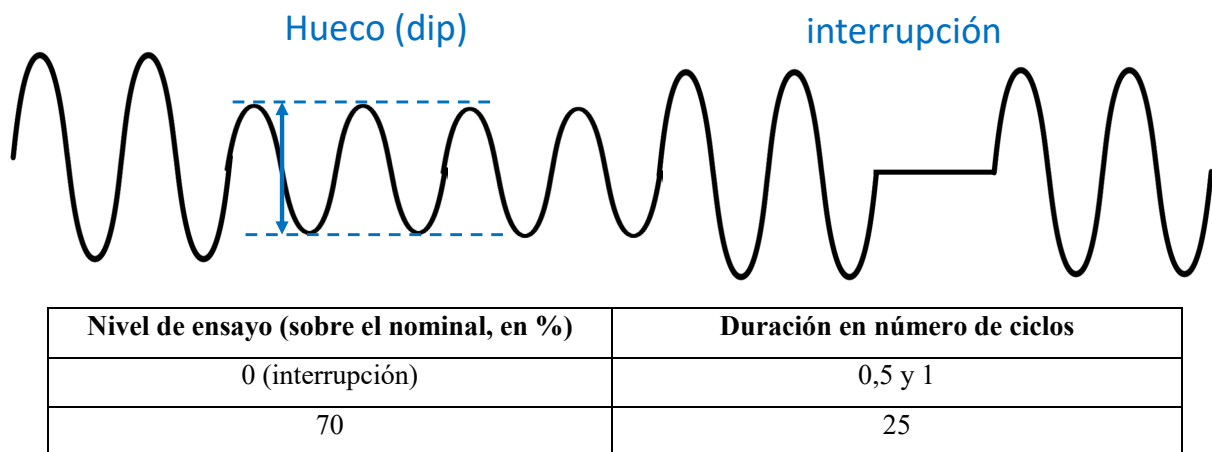


Figura 9.12. Caídas en la alimentación AC definidas en IEC-61000-4-11, Clase 2 (aplicaciones no industriales y no especialmente sensibles). El transitorio puede iniciarse y terminar en cualquier ángulo de la fase de la señal. La norma también define que se debe ensayar el equipo frente a interrupciones breves (de 250 ciclos para dispositivos de clase 2)

Criterios de aptitud en los ensayos de inmunidad

Superar o fallar un ensayo de emisiones depende de quedarse por encima o por debajo de una línea, un límite que marca la norma para tu producto. Es binario: pasa o no pasa y no hay discusión. Como mucho, podemos hablar de la incertidumbre de medida y cómo interpretarla (lo que haremos en la página 202).

En inmunidad las cosas son más complejas, pues, ¿qué quiere decir que un DUT supera un ensayo de inmunidad? Pensemos, por ejemplo, en un monitor de ordenador sometido a un ensayo de onda de choque. ¿Será aceptable que la pantalla se apague y se reinicie? ¿O que parpadee? ¿O no debe haber degradación apreciable de prestaciones? Está claro que la destrucción, el daño físico implica no superar el test. Pero ya no se trata de algo tan simple como superar o no un límite. De hecho, gran parte de la tarea de detectar un fallo o mal comportamiento recae en el software de prueba que debes entregar junto al DUT al laboratorio de ensayo EMC.

EN 55024 define 3 criterios de aptitud en test de inmunidad, que se aplican de forma general a los productos electrónicos:

- **Criterio de aptitud A: el equipo debe operar correctamente antes, durante y después del test.** No se permite degradación de funcionamiento o pérdida de función. Pero “El nivel de aptitud puede sustituirse por una pérdida de aptitud admisible” definida por el fabricante. Es decir, que en principio las funcionalidades del producto no deben verse afectadas, pero si en el manual de usuario dices que ante una descarga ESD la pantalla puede parpadear y consideras como fabricante que eso es aceptable, es válido.
- **Criterio de aptitud B: el equipo debe operar correctamente tras el test.** Durante el test se permite degradación de funcionamiento o pérdida de función, pero no cambio en datos almacenados o en el estado de operación. De nuevo “El nivel de aptitud puede sustituirse por una pérdida de aptitud admisible” definida por el fabricante. Es decir, el equipo no puede reiniciarse o perder el estado.
- **Criterio de aptitud C: Se permite una pérdida temporal de función, siempre que el sistema pueda recuperarse por sí mismo o mediante un control externo.** Es decir, con tal que el equipo pueda recuperar la funcionalidad tras el test mediante un reset (automático o manual), el equipo supera el ensayo. Lo que no está permitido son fallos *hard* (daños o destrucción)

Habitualmente, para RF radiadas o conducidas se ha de cumplir el criterio A. Esto es lógico, ya que se trata de perturbaciones continuas y el equipo no puede estar permanente con sus funcionalidades degradadas.

Para transitorios ESD y EFT se debe cumplir generalmente el criterio B. Son transitorios que ocurren de forma más o menos esporádica (con una alta dependencia del entorno en el que trabaja el equipo) y puede ser admisible una degradación leve y transitoria de funcionalidades.

Para transitorios tipo onda de choque se debe cumplir generalmente el criterio C (en líneas de señal) o B (en líneas de alimentación). Son transitorios de alta energía y baja ocurrencia y es por tanto aceptable una afectación más severa a las funcionalidades.

Cabe llamar la atención sobre el papel que juega el *software del equipo*. Ante una perturbación externa, los filtros y protecciones que añadamos al equipo pueden evitar la destrucción y atenuar la perturbación, pero con demasiada frecuencia no podrán evitar cambios de estados lógicos en entradas, incluso corrupción en relojes, señales de *reset* y de control. Se hace necesario rechazar *glitches* (pulsos estrechos fruto de perturbaciones externas) por software. Un ejemplo sencillo: ante una entrada de control que proviene de la botonera de un ascensor, el microprocesador debe rechazar pulsos más cortos de 100 ms (coge un pulsador y un osciloscopio e intenta conseguir un pulso más estrecho, te reto). Ciertamente en este ejemplo un condensador en la línea junto al pin del microprocesador podrá hacer esta función, pero no está de más la redundancia y hay muchos otros ejemplos que requieren un firme control de errores por *software*: desde *watchdog timers* hasta comprobaciones lógicas que detecten combinaciones o secuencias de entradas imposibles.

Ensayos de emisión

Hemos comentado anteriormente (página 195) que las emisiones en HF (por encima de 30 MHz) se miden en cámara anecoica con una antena, pero que en medias frecuencias es más viable medir la corriente en los cables como aproximación al campo radiado. Por este motivo distinguimos entre emisiones radiadas y conducidas.

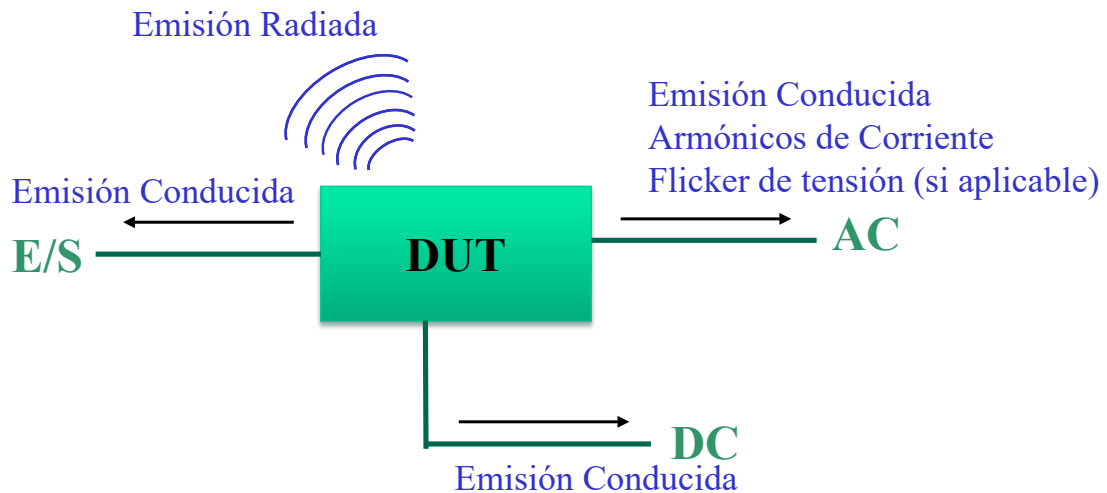


Figura 9.13. Ensayos de emisión

Emisiones radiadas

Las emisiones radiadas se miden normalmente entre 30 MHz y 1 GHz (el margen depende de la familia de productos), siguiendo el método de medida descrito en CISPR 16-1. Actualmente, según CISPR 32 el margen superior de frecuencia es de 1 GHz para equipos con frecuencias internas de hasta 108 MHz. Si hay frecuencias internas de hasta 500 MHz, el límite de medida sube a 2 GHz. Con frecuencias internas de hasta 1 GHz hay que extender el rango de medida a 5 GHz. Para frecuencias internas por encima de 1 GHz hay que medir hasta 6 GHz, que es el caso de equipos que incorporen microprocesadores de altas prestaciones.

Se realiza un barrido en frecuencia, realizando mediadas de cuasi-pico con un ancho de banda de resolución de 120 kHz hasta 1 GHz. Entre 1 y 6 GHz se miden valores de pico o de promedio con un ancho de banda de resolución de 1 MHz. Un detector de cuasi-pico se comporta como un detector de pico seguido de un integrador, de modo que suaviza la respuesta y requiere un tiempo de estabilización de la medida antes de saltar a la siguiente frecuencia de medida.

Un escaneo de todo el margen de frecuencias puede alcanzar varias horas, teniendo en cuenta que en cada frecuencia hay que girar el equipo sometido a ensayo 360°, variar la altura de la antena entre 1-4 m para buscar el máximo (y eliminar nulos) y medir con antena en polarización horizontal y vertical. La situación empeora si las máximas emisiones o el ciclo de operación cambia cada N segundos. En ese caso, habrá que esperar dicho tiempo entre pasos de variación de la altura de la antena o de giro de la mesa: es muy importante optimizar el software de la configuración de test si quieres evitar perder días haciendo las medidas. El *test setup* está definido por las normas y define distancias entre antena y DUT de 3 o 10 m. El DUT se ubica sobre una mesa no conductora motorizada a 0,8 m de altura sobre el suelo.

Una vez realizadas las medidas, se representan en una gráfica con el eje vertical expresado en $\text{dB}_{\mu\text{V/m}}$ y con el eje horizontal en frecuencia en escala logarítmica. Se comprueba en esta gráfica si en alguna frecuencia se superan los límites que marca la norma. En la Figura 9.14 se indican los límites que marca la norma CISPR-32 hasta 1 GHz.

Hay que hablar, aunque sea brevemente, sobre la **incertidumbre de medida**. ¿Qué pasa si en una frecuencia excedo el límite por 1 dB? ¿No podría ser que la incertidumbre de medida en el ensayo juegue en mi contra y refleje un valor superior al real? Bien, CISPR-16-4-2 recoge esta problemática.

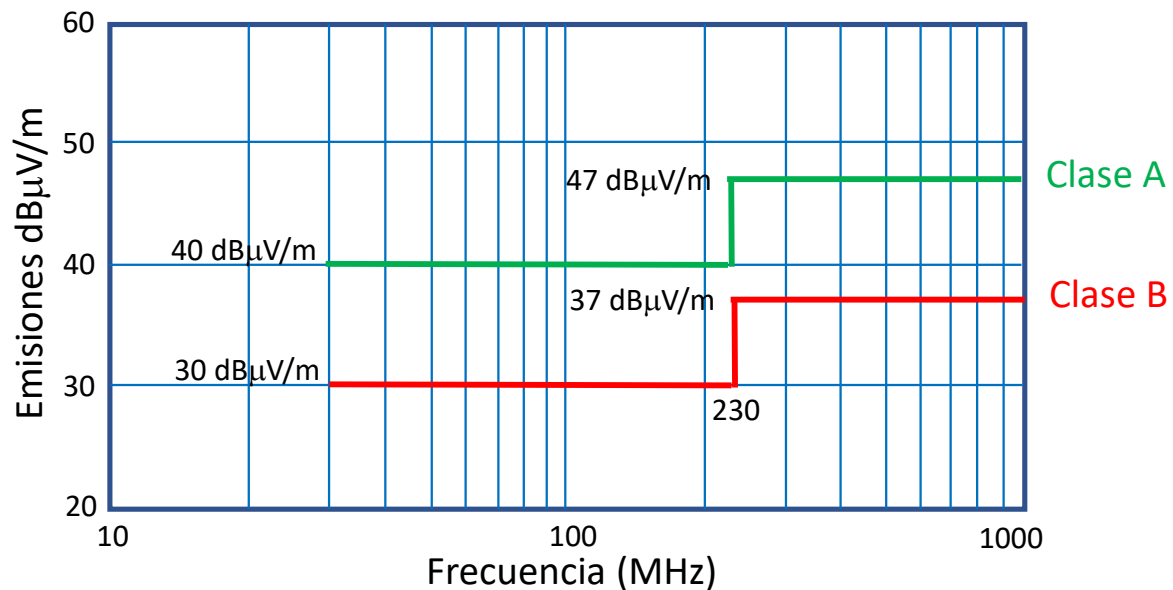


Figura 9.14. Límites de emisiones radiadas según la norma CISPR-32 (EN55032) hasta 1 GHz. La clase B se refiere a productos para entornos residenciales y comerciales. La clase A aplica a productos industriales

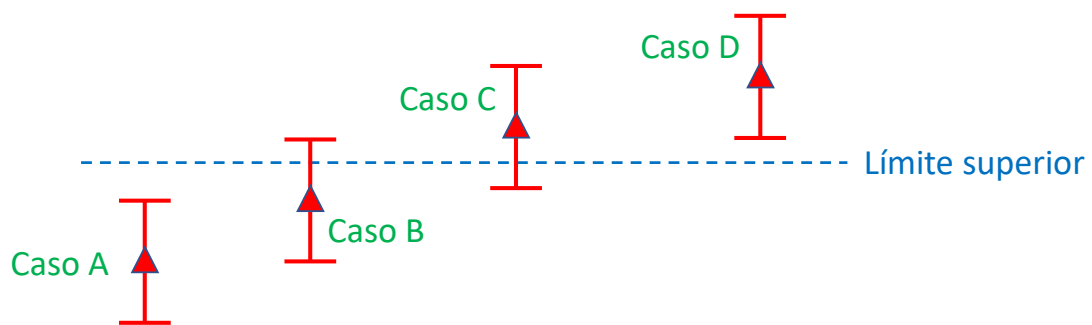


Figura 9.15. CISPR 16-4-2 y la incertidumbre de medida

En la Figura 9.15 se recogen los cuatro casos que se pueden dar. En cada uno, el triángulo muestra el valor medido, el intervalo de incertidumbre y el límite que marca la norma (línea horizontal).

- Caso A: cumple, ya que medida $\pm$ incertidumbre está por debajo del límite.
- Caso D: no cumple, ya que medida $\pm$ incertidumbre está por encima del límite.
- Caso C: se considera que no cumple, ya que la medida está por encima del límite
- Caso B: si incertidumbre del laboratorio es baja ($\leq 5,1$ dB), se considera que cumple. Si la incertidumbre del laboratorio es mayor que 5,1 dB, hay que sumar a la medida la diferencia entre la incertidumbre y 5,1 dB y determinar entonces si el valor obtenido es superior al límite.

¿Dónde se realizan los ensayos de emisiones radiadas?

En principio pueden realizarse al aire libre, sin más que asegurarnos de que no hay elementos metálicos cercanos que puedan producir reflexiones y falsear las medidas. Las recomendaciones CISPR definen que, para medidas con antena a 10 m del DUT, no puede haber objetos que reflejen las ondas (coches, farolas, vallas metálicas) en una elipse de eje mayor de 20 m y eje menor de 17,3 m. Se ha de instalar un plano de tierra, generalmente mediante una rejilla metálica, de unas dimensiones mínimas (típicamente extendiéndose 1 m alrededor de DUT y antena y cubriendo la distancia entre ambos). El problema de usar este tipo de instalaciones, que son muy económicas, reside en complicaciones asociadas a la meteorología (lluvia, viento y efecto de las variaciones de temperatura y humedad sobre los equipos).

La alternativa a estos sitios al aire libre (OATS, *open-area test site*) es una costosa cámara anecoica (generalmente semi-anecoica, con suelo reflectante), con paredes absorbentes de radiofrecuencia, al alcance sólo de unas pocas empresas y de los laboratorios que se dedican profesionalmente a hacer ensayos EMC.

Emisiones conducidas

Continuemos con el ejemplo de la norma CISPR-32, aplicable a equipos multimedia y de tecnología de la información. El margen de frecuencias comprende el rango 150 kHz-30 MHz. El método de ensayo se detalla en CISPR 16-1-2 y varía según el tipo de cable (alimentación o puerto de comunicaciones). **El objetivo de este ensayo es comprobar que el nivel de señal en los cables en cada frecuencia no supera el límite que marca la norma.**

La Figura 9.16 muestra la disposición de elementos en el ensayo de emisiones conducidas. Requiere un plano de masa de $2 \times 2 \text{ m}^2$, una mesa no conductora, una red de estabilización de impedancia -ahora veremos por qué-, y (ahora viene lo caro) un receptor. Como elemento adicional, se suele incluir un limitador a la entrada del receptor para no dañarlo por sobrecarga. El acoplamiento entre el cable y el plano de tierra influye en la medida, de modo que la norma da instrucciones sobre cómo disponer el cable.

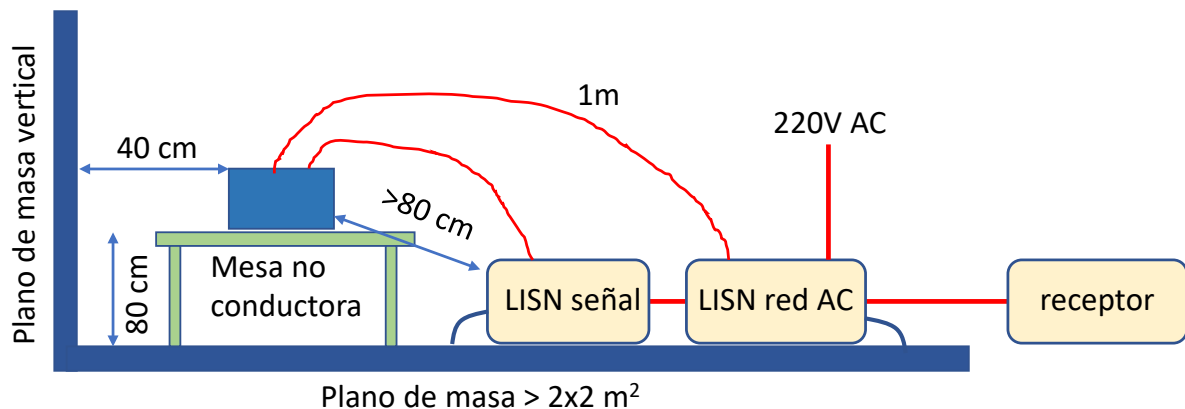


Figura 9.16. Disposición de los equipos en el ensayo de emisiones conducidas para puertos de comunicaciones y de alimentación AC

En todos los casos es necesario introducir entre el DUT y el receptor (medidor) un equipo especial denominado **ISN** (*impedance stabilization network*, *red de estabilización de impedancia*). Las razones son varias y pueden resumirse en los siguientes puntos:

- Para garantizar la repetibilidad de las medidas, el DUT debe ver siempre una misma impedancia en el cable.
- El receptor tiene una entrada apantallada no diferencial y con impedancia de entrada de 50 ohmios. El cable que viene del DUT y en el que queremos medir el nivel de emisiones puede ser diferencial o no, apantallado o no y con una impedancia de línea muy distinta de 50 ohmios: necesitamos una red que extraiga la señal de interés y la adapte al receptor.
- Queremos medir sólo la perturbación que el DUT pone en el cable, por tanto, necesitamos desacoplar de la medida las perturbaciones que añaden otros elementos de la red.

Se suele denominar **LISN** (*line impedance stabilization network*) a la red de adaptación empleada para medir emisiones en líneas de alimentación y que cumple con las funciones anteriores.

CISPR 16-1-2 define un circuito para LISN (Figura 9.17). La inductancia de 50 μH y la resistencia de 5 ohm definen la impedancia, que es razonablemente estable (45-48 ohm) en el rango 150 kHz-30 MHz. El resto de componentes sirven como desacoplo y filtro adicional. El resultado es un filtro que limita el paso de interferencias entre la red AC o DC y equipos auxiliares, por un lado, y el DUT y el receptor de medida por otro. Esta red se repite para cada uno de los cables de alimentación (fase y neutro en AC, o positivo y negativo en DC). **Es decir, no se mide el ruido en modo común sino en modo normal.** Esto es importante a la hora de

diseñar el filtro de alimentación: hemos de añadir filtrado para ambas componentes y no pensar que como el modo normal radia mucho menos que el modo común, debemos filtrar sobre todo este último.

Hay un video en Youtube que te recomiendo ver para mejorar tu comprensión de los que es un LISN (si bien explica el LISN definido en CISPR-25): <https://www.youtube.com/watch?v=QPJzp66Yvzs>.

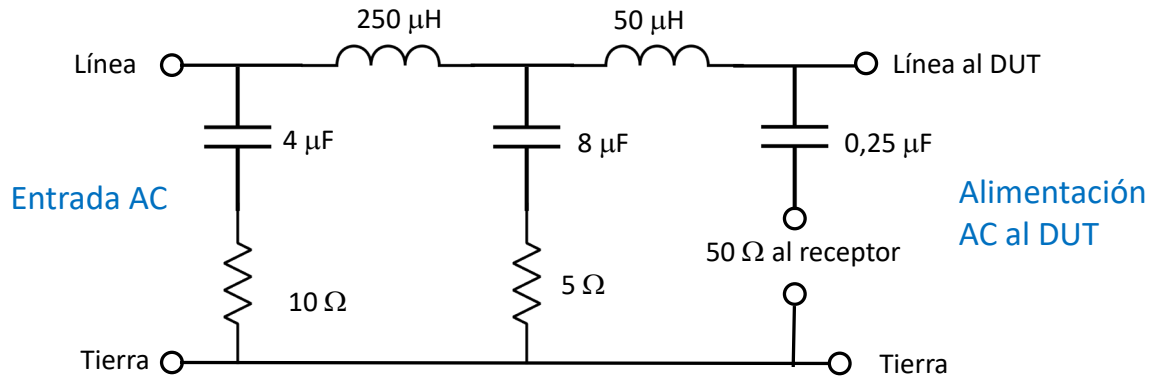


Figura 9.17. LISN definido por CISPR 16-1-2

Para realizar medidas en cables que no sean de alimentación, se utilizan diversas redes (ISN) que siguen cumpliendo las funciones de estabilización de la impedancia, desacoplamiento de ruido entre el lado del DUT y el resto del *setup* de medida, y la extracción de la señal de medida para el receptor. Estas redes están adaptadas al tipo de cable y de señal, pero el principio básico es el mismo que el que hemos estudiado para el LISN.

Por último, hay que comentar que diferentes tipos de cables/señales requieren medidas de tensión, de corriente o ambas. Pero esto queda al cuidado del laboratorio de ensayo. Para ti, como diseñador, y siendo la medida de emisiones conducidas en las líneas de alimentación las más comunes, debes conocer el circuito del LISN (o su curva de impedancia) como ayuda para el diseño del filtro.

Emisiones en baja frecuencia

Además de medida de emisiones radiadas a partir de 30 MHz y conducidas entre 150 kHz y 30 MHz, también se realizan medidas de baja frecuencia en puertos de alimentación AC. En concreto, las normas introducen límites a los armónicos inyectados en la línea AC (hasta 2 kHz) y al parpadeo (*flicker*) que el consumo discontinuo del equipo puede producir en una luminaria y ser percibido por una persona.

Día 10. Diseño de protecciones frente a transitorios ESD



Los ensayos de inmunidad frente a RF radiada y conducida no tienen (en la mayoría de los casos) capacidad destructiva, pero los ensayos de inmunidad frente a transitorios de alta tensión (ESD, EFT y onda de choque) sí. Esto obliga a añadir protecciones en las conexiones de entrada/salida y de alimentación de nuestro equipo.

En el caso de la ESD (descarga electrostática), unos diseñadores se limitan a añadir un diodo TVS (a veces llamado transil), lo que es un error como veremos. Otros diseñadores copian una red de protección probada anteriormente en otro producto de su empresa, sin saber si es una buena opción para el nuevo diseño. Por último, hay diseñadores que simplemente no añaden protecciones y fían el éxito en el ensayo ESD (y a la robustez del producto una vez en uso) a la suerte.

No me he encontrado todavía con ningún ingeniero que tuviera una comprensión adecuada del diseño de protecciones frente a ESD. De modo que dedicaremos una lección completa a entender los conceptos necesarios y a aprender una metodología que, tras años de enseñarla a decenas de ingenieros que trabajan en prácticamente todos los sectores, considero suficientemente probada.

Una visión general de los transitorios de alta tensión

Comencemos por la normativa. IEC 61000-4-1 define tres tipos de transitorios en test de inmunidad:

- **Descargas electrostáticas (ESD)**, con pulsos de típicamente 4-8 kV y una decena de amperios, anchura de pulso de unas decenas de nanosegundos y por tanto energía moderada, unas decenas de milijulios (mJ). Esta energía es más que suficiente para destruir circuitos integrados desprotegidos.
- **Ráfagas de transitorios rápidos (*electrical fast transient*, EFT)**, con tensiones y corrientes, anchuras de pulso y energías del mismo orden de magnitud que las descargas ESD. Como consecuencia, los mismos dispositivos y técnicas que protegen los circuitos de descargas ESD valdrán también como protecciones frente a EFT, pero con la consideración de que al tratarse de un ensayo que implica miles de pulsos, y no sólo unas decenas como en el caso de ESD, las protecciones deben ser capaces de sobrevivir a un mayor número de transitorios. Este transitorio emula el efecto de encendido y apagado de motores y otras cargas inductivas en la línea de alimentación AC, que pueden acoplarse capacitivamente a otros cables en la misma canaleta o mazo de cables.
- **Ondas de choque (*surge*)**, resultado de descargas de rayos que se acoplan a cables y conductores. Son transitorios mucho más lentos y de mayor duración que ESD y EFT, con corrientes de cientos de amperios y tensiones del orden del kilovoltio. Con una energía de varias decenas de julios, son altamente destructivos y requieren dispositivos de protección distintos.

Tabla 10.1. Características de los transitorios de alta tensión

	V (kV)	I (A)	t _r (ns)	width	energy
ESD	4-8	una decena	1	60 ns	1-10s mJ
EFT (pulse)	0.5-2	decenas	5	50 ns	4 mJ
EFT (burst)	0.5-2	decenas	N.A.	15 ms	100s mJ
Lightning surge	0.5-2	cientos	1250	50 µs	10-80 J

Ante un transitorio no continuo, las normas no exigen (por lo general) un criterio de aptitud A, sino B y en ciertos casos C. Como siempre, hay variaciones en función del tipo de producto. **Habitualmente, para ESD y EFT se pide criterio de aptitud B, y para onda de choque B o C (en función de si hablamos de un puerto de alimentación o de señal).**

Tener esto en mente es muy importante. Cuando lleves tu producto a un ensayo de, por ejemplo, ESD, estará permitido (criterio de aptitud B) que durante la descarga el equipo sufra una degradación de prestaciones, pero tras el transitorio debe recuperarse sin haber sufrido cambio de estado ni de memoria. Es decir, no sólo basta con que las protecciones impidan un daño físico al equipo, también debemos evitar (y aquí el software debe jugar un papel importante) que el equipo se reinicie, se quede “colgado” o que se produzcan errores *soft* relevantes. Recuerda lo que dijimos al respecto en la página 200.

¿Con qué elementos podemos proteger nuestro diseño?

Hay dos estrategias para evitar que el transitorio llegue a los elementos sensibles de nuestro producto, e implica usar protecciones de tipo **crowbar** o de tipo **clamping**.

Un dispositivo tipo **crowbar** basa su efectividad en presentar una impedancia muy baja frente al transitorio, conduciéndolo lejos de los circuitos sensibles, generalmente a chasis o a un área metálica elevada. La potencia disipada es baja, por lo que pueden hacer frente a transitorios muy agresivos como son las ondas de choque. Un tubo de descarga de gas (ver sección Los modelos SPICE para varistores se hacen en base a un ajuste a un polinomio. Con frecuencia dan problemas de sintaxis en LTspice. Si no quieres ir retocando el modelo para corregir los errores, puedes recurrir a otro simulador, como TINA de Texas Instruments.

Tubos de descarga de gas (GDTs) en la página 248) es un ejemplo de dispositivo crowbar. Su desventaja radica en que su transición de alta a baja impedancia es lenta, dejando expuesta la circuitería sensible al transitorio durante un tiempo suficiente para su destrucción.

Un dispositivo tipo **clamping** usa una estrategia diferente: limitan (recortan) la tensión que supere un determinado nivel, como es el caso de diodos TVS y de varistores. Su acción puede ser muy rápida (en la escala del nanosegundo en el caso de diodos y de varistores de montaje superficial) y son por tanto preferidos para proteger frente a transitorios rápidos (ESD, EFT).

El indiscutible rey de las protecciones frente a transitorios de alta tensión es una variante del diodo Zener con un área de la unión PN elevada para soportar picos de corriente elevados sin sufrir daños. A este tipo de diodos les llamamos **diodos TVS (transient voltage suppressors)** o diodos transil. Sus ventajas son su rápida respuesta (inferior al nanosegundo, lo que los hace aptos para transitorios rápidos como ESD y EFT), la existencia de versiones con capacidad eléctrica muy pequeña (lo que los hace adecuados incluso para puertos digitales rápidos como Ethernet o USB) y el hecho de que no se degradan sometidos a múltiples transitorios. Existen versiones de potencia aptas para ser empleadas incluso con transitorios de alta energía, como puede ser una onda de choque de baja amplitud.

En aplicaciones de muy alta frecuencia, la capacidad eléctrica de un TVS puede resultar excesiva. Existe un dispositivo, un tipo de **resistencia variable con la tensión (VVR)** con capacidad eléctrica ultra-reducida, del orden de femtofaradios (fF). Pero soporta niveles de energía reducidos y su capacidad de recortar (*clamp*) o limitar tensiones es también reducida.

Por encima de los diodos TVS, en cuanto a capacidad de sobrevivir a pulsos de mayor energía, encontramos diferentes versiones de **varistores**. Se trata de dispositivos bidireccionales, con capacidades eléctricas mucho mayores a los diodos TVS (y por tanto no aptos para interfaces digitales rápidos), tiempos de respuesta rápidos en el caso de varistores multicapa SMD y lentos en varistores estándar. Se degradan al ser sometidos a pulsos repetidos, reduciendo paulatinamente la tensión de recorte (*clamping voltage*). Las versiones de potencia pueden soportar ondas de choque reducidas, si bien los transitorios de mayor energía requieren otras soluciones como son los **tubos de descarga de gas** (una cápsula vidrio con dos o tres terminales metálicos llena de un gas que, sometido a una alta tensión, se ioniza creando un camino de baja impedancia y limitando así la tensión entre terminales).

Tabla 10.2. Tipos de dispositivos de protección frente a transitorios de alta tensión

dispositivo	$t_{\text{respuesta}}$	capacidad	Adecuado para	Comentarios
Diodo TVS	< 1ns	Varios pF	ESD, EFT, surge	Líneas de señal, alim. DC y líneas de señal rápidas
VVR	< 1ns	< 1pF		Líneas de señal rápidas
Varistor multicapa SMD	< 1ns	10pF-2,5nF		En PCBs, líneas de señal y de alimentación
Filtro LC	-	alta		Líneas de señal y de alimentación
Varistor estándar	100s ns	10pF-10nF	surge	Líneas de alim. AC
Gas Tube	μs	< 1pF		Telecom. y líneas de alimentación

Normativa IEC para ESD

IEC 61000-4-2 es la norma básica de aplicación a productos que puedan sufrir descargas electrostáticas por parte de usuarios/operadores (es decir, casi todos los productos). La norma define:

- Forma de onda típica de corriente de descarga
- Niveles de ensayo
- Equipo, instalación y procedimiento de ensayo. Calibración e incertidumbre de medida
- Dos métodos de descarga: por contacto (preferente) y en el aire (cuando no se pueda hacer por contacto)
- Dos tipos de aplicación: directa (sobre el dispositivo) e indirecta (sobre un plano conductor adyacente –simula ESD de un operador sobre objetos próximos al DUT)

Recuerda que una norma básica dice cómo hacer el ensayo. La severidad de la perturbación que debe soportar el DUT y el criterio de aptitud exigible, así como otras condiciones particulares, quedan definidos en las normas específicas de familias de productos o en las normas genéricas, si no hay norma específica de aplicación.

Tabla 10.3. Niveles de descarga por contacto o al aire definidos en IEC 61000-4-2

Descarga por contacto		Descarga en el aire	
Nivel	Tensión de ensayo (kV)	Nivel	Tensión de ensayo (kV)
1	2	1	2
2	4	2	4
3	6	3	8
4	8	4	15

La norma IEC-61000-4-2 define la forma de onda (Figura 10.2) que como puedes comprobar se compone de dos pulsos: un primer pico estrecho y de alta corriente, debido a la carga acumulada en la punta de la pistola ESD, y un segundo pulso, con mucha mayor energía y anchura, que se corresponde con la descarga de la capacidad interna de la pistola. He añadido unas líneas rojas discontinuas para destacar la superposición de estos dos pulsos.

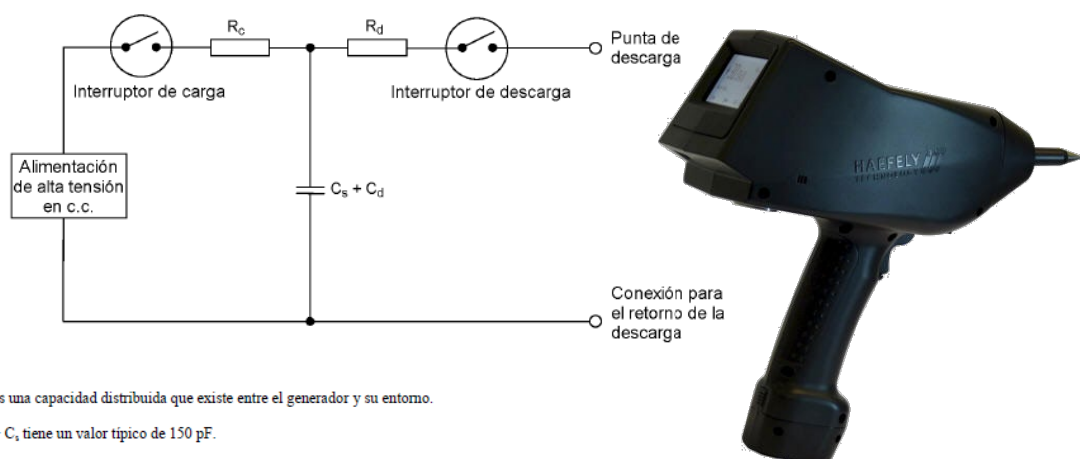
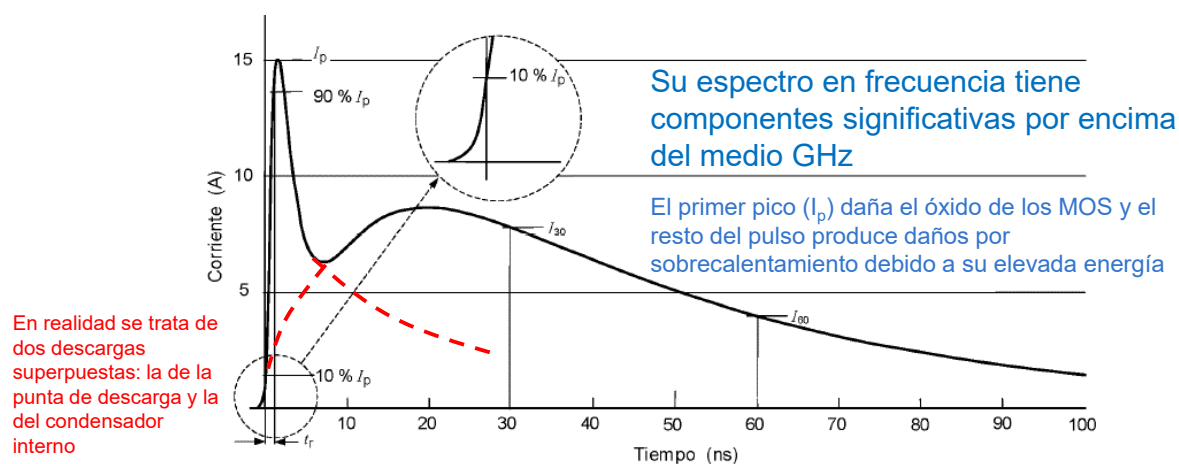


Figura 10.1. Modelo del generador de descargas (conocido como pistola ESD) y ejemplo de generador

El tiempo de subida del pulso ESD es de orden del nanosegundo, lo que implica que tiene componentes espectrales hasta cientos de MHz, con todo lo que ello implica.



Nivel	Tensión de pico (kV)	Pico de corriente (A) $\pm 15\%$	Tiempo de subida (ns) $\pm 25\%$	Corriente a 30 ns (A) $\pm 30\%$	Corriente a 60 ns (A) $\pm 30\%$
1	2	7,5	0,8	4	2
2	4	15	0,8	8	4
3	6	22,5	0,8	12	6
4	8	30	0,8	16	8

Figura 10.2. Forma de onda del pulso ESD

El generador ESD está diseñado para emular la descarga producida por una usuario u operador humano. Por eso el circuito incluye una capacidad de 150 pF (simular a la del cuerpo humano) y una resistencia de descarga de 330 ohmios (Figura 10.1). El circuito incluye un interruptor para cargar la capacidad d 150 pF y otro para aplicar la descarga.

Las descargas se aplican siempre respecto a tierra y la norma específica que la conexión de la pistola a tierra es a través de un cable de 2 metros. La forma de la punta metálica de la descarga afecta a la forma de onda, de modo que incluso este detalle está detallado en la norma.

Inciso: simulación del pulso ESD

Antes de seguir describiendo cómo se llevan a cabo los ensayos ESD vamos a introducir la idea en la que se basa la metodología de diseño de protecciones que te propongo: prototipado virtual mediante simulaciones SPICE.

Para poder hacerlo necesitamos cinco elementos:

- Modelos SPICE para las fuentes de perturbaciones
- Modelos SPICE para las protecciones (TVS, ferritas, varistores, etc.)
- Modelos SPICE para los circuitos integrados que queremos proteger
- Un simulador SPICE
- Confianza en que la metodología sea válida

Afortunadamente, todo esto está disponible y sin coste. Vamos a limitarnos a introducir el modelo de una pistola ESD. Dejaremos para otras secciones la presentación del resto de elementos, el estudio de un par de ejemplos y una discusión sobre las luces y alguna sombra de esta metodología (nadie es perfecto).

Bien, para hacer un modelo SPICE de una fuente de pulsos ESD compatible con la norma IEC 61000-4-2, vamos a plantear dos ramas paralelas de circuitos RLC serie: uno representará la descarga de la punta de la sonda y otro el condensador de descarga (Figura 10.3). En [26] podrás encontrar una propuesta de valores. La mía parte de $C_1=50$ pF, $C_2=100$ pF, $R_2=330$ Ω , $R_1=1$ Ω , $L_1=100$ nH, $L_2=1.5$ μ H. A partir de ahí fui

afinando el modelo para ajustarlo a la curva IEC y se obtiene una forma de onda que cumple con I_{pico} , $I_{30\text{ns}}$, $I_{60\text{ns}}$. Puedes usar en los ejercicios que haremos más adelante los valores de la Figura 10.3 o los de [26].

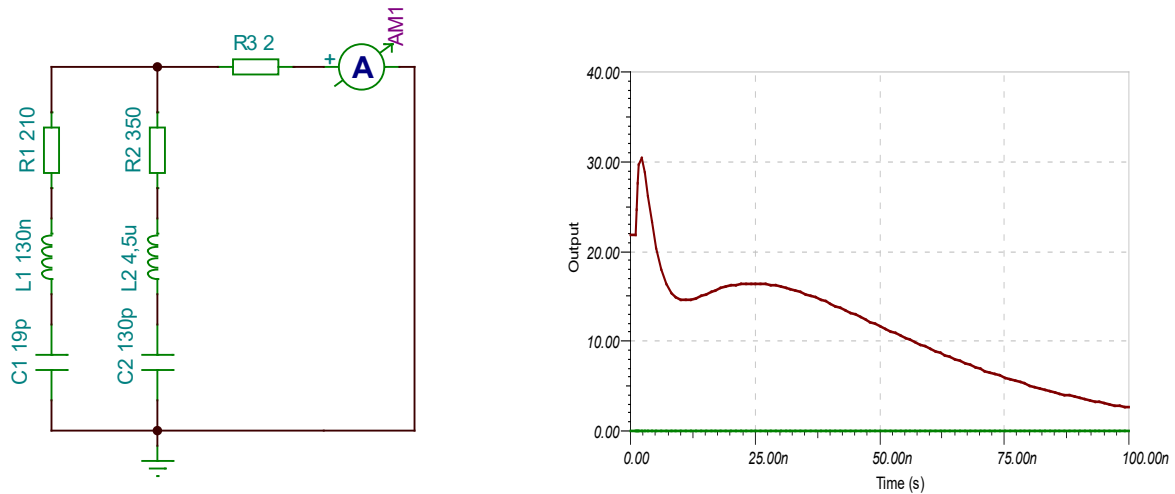


Figura 10.3. Modelo SPICE para un generador ESD. Nos permitirá ensayar virtualmente protecciones ESD sin necesidad de realizar test destructivos

¡Ahora dispones de una pistola ESD virtual! Las condiciones iniciales de C1 y C2 se fijan para la tensión de ensayo (4 kV, por ejemplo, o la que necesites ensayar en tu equipo). Una vez modelos protecciones y circuito a proteger, podrás estimar tensiones, corrientes, potencias disipadas y anchuras de pulsos sin tener que destruir nada.

Procedimiento de ensayo ESD

El DUT (*Device Under Test*) debe funcionar en su modo más sensible, el software de prueba comprobará el correcto funcionamiento de todas las funcionalidades, testeará la integridad de los datos en memoria, de los datos enviados y recibidos por los puertos y todo aquello que sea relevante.

La norma genérica, o de producto nos da el nivel de las descargas a aplicar (IEC61000-4-2 es sólo norma básica que nos dice cómo hacer el ensayo) en un *setup* como el de la Figura 10.4. Es fácil reproducirlo en un rincón de un laboratorio. Requiere una mesa no conductora (¿madera?) sobre un plano de tierra (y efectivamente puesto a tierra). El DUT se coloca sobre un soporte aislante, que a su vez se apoya sobre un plano de acoplamiento horizontal (PAH) de cobre 1,6x0,8 m² (que bien puede ocupar toda la superficie de la mesa) y junto a un plano de acoplamiento vertical (PAV) de 0m5x0,5 m² de cobre, también sobre el soporte aislante. Ambos planos de acoplamiento van unidos a tierra mediante resistencias de 1 MΩ. El DUT está separado 10 cm tanto del borde del PAH como del PAV. Todo este protocolo permite que los ensayos sean repetibles.

Bueno, son repetibles los **disparos por contacto** (la punta de la pistola ESD tocando el DUT o los planos PAH, PAV. Menos repetibles son los **disparos por aire** (aquellos a los que no podamos acceder con la punta de la pistola, también en los contactos de un conector de plástico -como Ethernet-), con el generador cargado y acercándola al punto sensible hasta que salta el arco, lo que depende principalmente la humedad del aire. Como es un parámetro menos controlable, la prueba es también menos repetible.

Eligiendo los puntos de test

Con la pistola configurada a 20 disparos por segundo, nos acercamos al DUT e identificamos los puntos sensibles, tomando nota de su ubicación y del tipo de disparo que será necesario (por contacto o al aire).

¿Dónde no aplicaremos descargas?: donde, durante un uso normal, un usuario no accede. Es decir, quedan excluidos superficies y puntos de mantenimiento, tales como contactos de baterías reemplazables o superficies no accesibles tras la instalación. Tampoco se ensaya sobre contactos donde vayan conectores blindados (se harán descargas sobre el blindaje). No se aplican descargas en conectores sensibles

funcionalmente a ESD y que estén adecuadamente etiquetados como tal (por ejemplo, en entradas de antena GPS o GPRS).

Realizando el test

Aplicaremos al menos diez descargas simples de la polaridad más sensible en los puntos seleccionados. Las descargas van separadas al menos 1 segundo (más si es necesario para determinar si ha habido un fallo: el software de test puede requerir su tiempo). En equipos no conectados a tierra se descargará en DUT entre disparo y disparo.

Hay que realizar tanto descargas directas en los puntos sensibles como descargas indirectas sobre los planos PAH, PAV. **¿Qué sentido tiene hacer descargas indirectas?**: emulan la situación en la que un humano se descarga sobre una superficie o chasis metálico y se produce un acoplamiento capacitivo desde esta superficie hasta el DUT. Este acoplamiento se realiza sin resistencia serie y por tanto puede resultar más abrupto y agresivo.

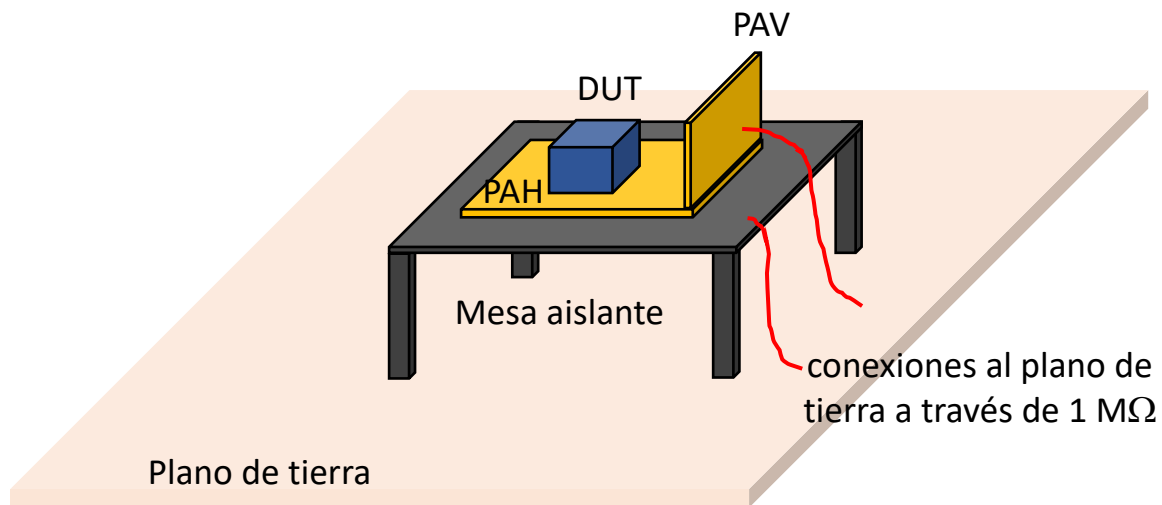


Figura 10.4. Test setup para ensayos ESD

Primera línea de protección: la caja

Cajas de plástico

Un espesor de 1 mm de PVC presenta una tensión de ruptura dieléctrica de aproximadamente 30 kV. Si utilizamos ABS, la tensión de ruptura sería de 16 kV. Y si fabricamos la caja en PLA con una impresora 3D, podemos estar entre 30 y 60 kV. Por lo tanto, en ensayos ESD típicos con tensiones entre 4 y 8 kV, podemos decir que cualquier caja servirá para prevenir la ESD. ¿Completamente? No. La caja presenta juntas y un número variable de orificios para el paso de cables y conectores.

En la Figura 10.5 se ilustra este punto. No sabemos qué hay en la junta o en el orificio. Tal vez sólo plástico perfectamente limpio y aire. Tal vez grasa de las manos de un montador. Tal vez suciedad conductora. Lo más conservador es suponer que una junta o un orificio son conductores y que la descarga alcanzará el lado interior.

Asumiendo una tensión de ruptura dieléctrica del air entre 1 y 3 kV/mm (en función de la humedad), y para ensayos ESD de hasta 8 kV, podemos decir que 1 cm es una distancia segura entre cualquier componente o elemento conductos del diseño y una junta u orificio.

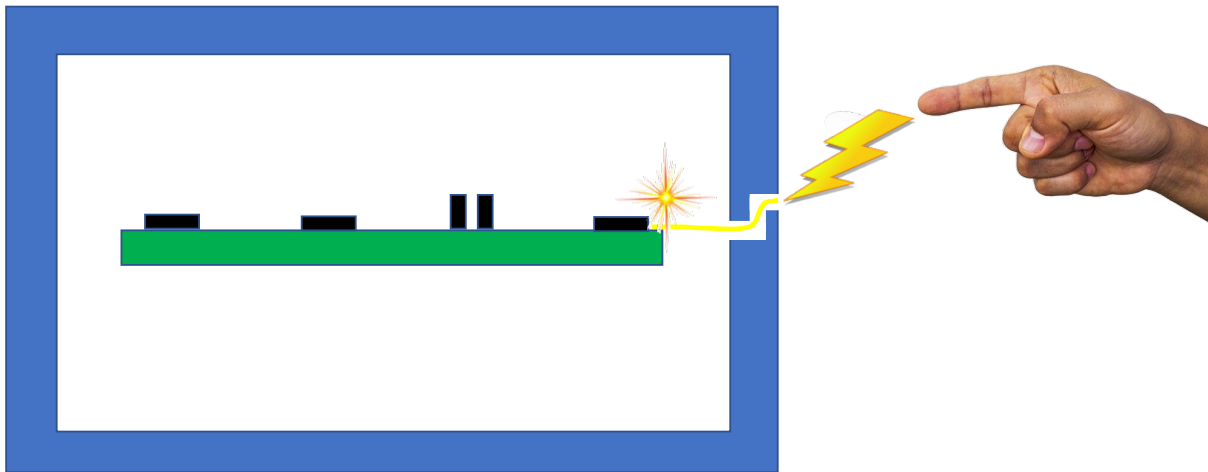


Figura 10.5. ESD a través de una junta en una caja de plástico. Fuente propia

¿Y si la caja es de plástico con pintura conductora en su lado interior?

Teniendo en cuenta que la pintura conductora permitirá que la descarga alcance toda la superficie interior de la caja, modificamos la afirmación anterior: asumiendo una tensión de ruptura dieléctrica del air entre 1 y 3 kV/mm (en función de la humedad), y para ensayos ESD de hasta 8 kV, podemos decir que 1 cm es una distancia segura entre cualquier componente o elemento conductos del diseño y cualquier superficie interior de la caja.

Placa de descarga ESD

En el caso en que no podamos permitirnos una separación (*gap*) de 1 cm entre la caja el PCB, podemos hacer lo mismo que una calculadora electrónica y tantos otros dispositivos de mano sin chasis metálico: usar una placa metálica para capturar la mayor parte de la descarga y alejarla de la electrónica (Figura 10.6). El concepto se resume en la Figura 10.7, aunque hay que explicarlo desde más de un punto de vista.

Interponiendo una placa metálica entre las juntas y el PCB, derivamos la descarga a una superficie metálica no conectada tierra. Superficie metálica que presenta una capacidad a tierra que se cargará a miles de voltios por efecto de una descarga ESD. Lentamente, a través de iones en el aire, a través del contacto con otra superficie o con el operador, la placa se descargará lentamente a tierra. Pero eso ocurre mucho después del que el pulso ESD, que no se extiende más allá de un par de cientos de nanosegundos, no sea más que un recuerdo lejano.



Figura 10.6. La calculadora que me acompaña desde primer curso de carrera, en 1989, una Casio fx-4500P y su placa de descarga ESD (derecha)

Esta placa interna puede ir conectada a tierra o no. Que esté conectada a tierra sólo acelera su descarga tras el pulso ESD; pero no tiene influencia en el par de cientos de nanosegundos del transitorio. ¿Por qué?: porque la inductancia de la conexión a tierra, generalmente un cable de longitud superior a un metro, supone una impedancia muy elevada a las frecuencias donde el transitorio tiene su energía.

La placa interna metálica hace uso de su capacidad de espacio libre y capacidad a tierra para acumular la carga y evitar que vaya a parar a la electrónica.

Pero ahora hemos creado otro problema: tenemos una placa metálica cargada a una tensión elevada junto a la electrónica. ¿No podrá producirse una descarga entre ambos? Sí. Y si no hay descarga, ¿no podrá acoplarse una fuerte interferencia por acoplamiento capacitivo? También. Para evitar lo primero y reducir lo segundo, conectamos placa interna y electrónica con un hilo (forma de expresar una conducción eléctrica de no baja impedancia a alta frecuencia). La impedancia del hilo bastará para limitar la carga que se transfiere a la electrónica, dejando la mayor parte acumulada en la placa interna, o esa es la intención.

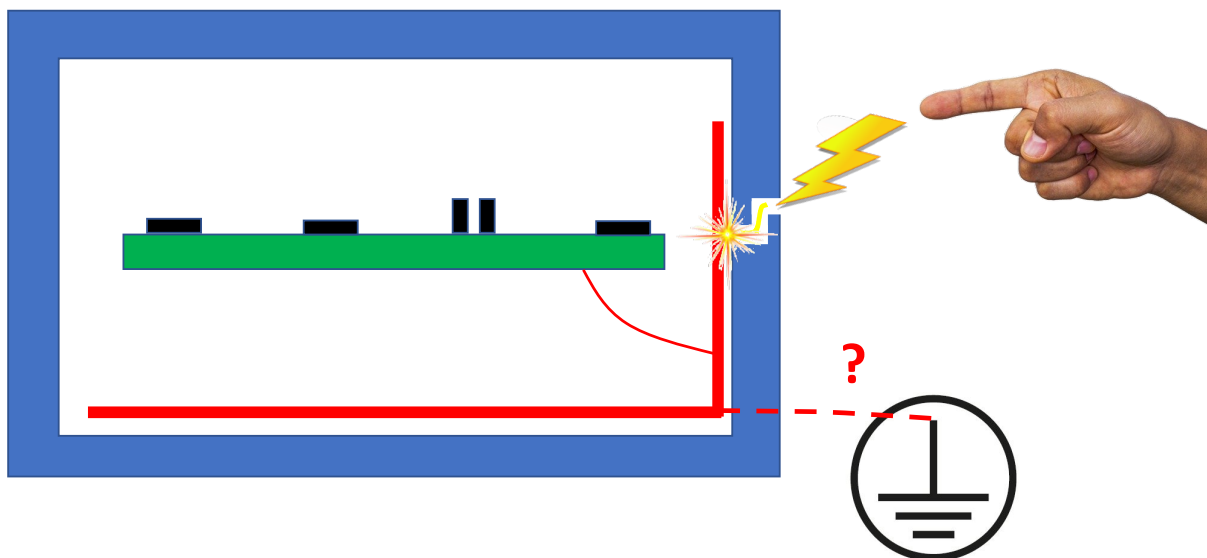


Figura 10.7. Placa metálica interna para capturar las ESD

¿Y si no podemos permitirnos una placa metálica interna?

Por falta de espacio, coste o cualquier otro condicionante podemos tener que abandonar la solución anterior. En este caso, podemos crear en el PCB un área de masa de E/S (entrada/salida), tan grande como sea posible que abarque conectores y protecciones ESD pero que no se solape con otros planos. Esta área de masa sucia estará conectada a la masa de circuito en un punto, típicamente mediante una ferrita. Siempre que sea posible, la mas de E/S, masa sucia, estará unida a tierra con una conexión de baja inductancia. La capacidad de esta área podrá acumular cierta cantidad de carga ($C=Q/V$). Una segunda (e importante función) consiste en

derivar a tierra las elevadas corrientes que circularán por los diodos TVS y por los varistores, evitando que creen perturbaciones en la masa de circuito.

Si tampoco disponemos de una masa de E/S, tendremos que usar la masa del circuito como nodo al que derivar los transitorios ESD, lo que es una pobre solución, al provocar transitorios elevados en masa que tendrán consecuencias en el funcionamiento del producto.

Si ni siquiera tenemos un plano de masa, ¡tendremos un gran problema!

Hace unos meses, un antiguo alumno me planteó un problema de protección ESD. Su producto (un actuador lineal controlado por una pequeña electrónica) tenía la electrónica a escasos milímetros del actuador lineal, que era metálico, y estaba expuesto al ambiente y por tanto a descargas ESD. Las descargas sobre el actuador metálico se acoplaban a la electrónica por ruptura dieléctrica. Y no sabía cómo evitarlo. Conectar el actuador a tierra no era una solución posible o no resolvía el problema, no recuerdo ahora los detalles.

Le sugerí interponer una delgada capa de kapton (poliimida) entre la electrónica y el actuador. Esto consiguió evitar la ruptura dieléctrica y atenuar la carga acoplada a la electrónica lo suficiente como para que el producto funcionara correctamente.

En electrónica, una vez construido el prototipo o el producto, el 90% de los problemas son mecánicos.

Cajas metálicas

Usar cajas metálicas facilita la captura de la descarga y su deriva a tierra, pero también provoca problemas (Figura 10.8) similares a los que comentamos con la caja de plástico recubierta internamente de pintura conductora. Vamos a explicarlo de nuevo. Durante la descarga ESD la caja alcanzará una elevada tensión respecto a tierra, menor cuanto mayor sea el área de la caja (recuerda, $C=Q/V$). Esto puede producir:

- Descargas secundarias entre la caja y el circuito: podemos evitarlas dejando suficiente separación (*air gap*) entre ambos
- Inyección de corriente en los nodos del circuito a través de las capacidades parásitas. Podemos reducirlas conectando la caja a la masa del circuito en un punto, con una conexión que presente algo de impedancia (por inductancia) a las elevadas frecuencias implicadas en la descarga.

Aunque tomemos las medidas anteriores, pueden seguir produciéndose inyecciones de corriente (por acoplamiento capacitivo). Si queremos un mayor nivel de atenuación, podemos añadir una segunda caja interna.

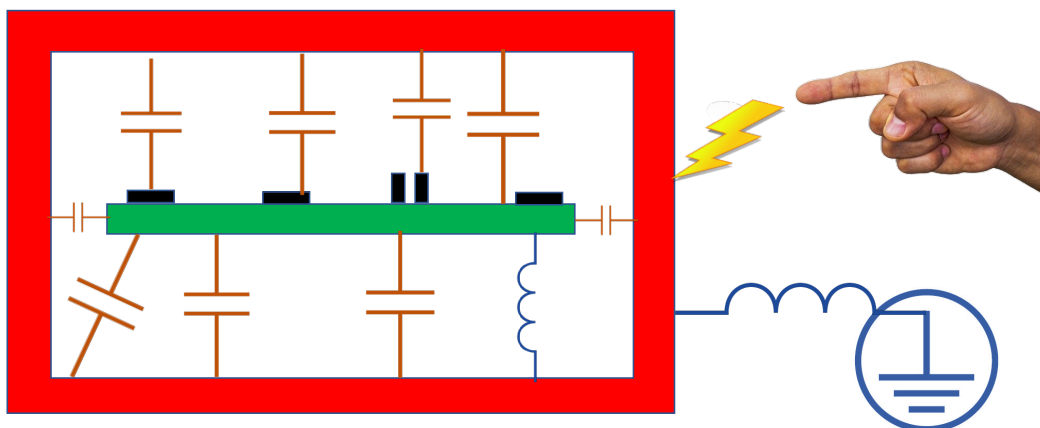


Figura 10.8. Efecto de la ESD en un equipo con envoltorio metálica. Fuente de la imagen: http://www.compliance-club.com/archive/old_archive/991215.htm#_Toc7417944

¿Qué pasa con equipos aislados operados por baterías?

La diferencia es que no hay conexión de la caja a tierra. El efecto de la descarga es el mismo (la impedancia de la conexión caja-tierra es por lo general tan elevada a alta frecuencia que podemos ignorarla), pero la descarga de la caja será mucho más lenta.

Conclusión: la caja no protege completamente

Las aberturas (paso de cables, ranuras de ventilación, juntas de montaje) proporcionan un camino de entrada de la ESD a la electrónica. Las cajas conductoras facilitan la inyección de corrientes por acoplamiento capacitivo entre la caja y el circuito. Pero lo que una caja no puede evitar es la descarga ESD que entre en nuestro equipo por un cable. La única solución a este problema es capturar la descarga a la entrada de nuestro producto y derivarla sin que afecte a su funcionamiento. Y esto implica el uso de protecciones ESD y su correcta conexión.

Los fabricantes de circuitos integrados dedican una parte de la valiosa área del dado de silicio a diodos de protección: generalmente, un diodo de la E/S a VCC, otro diodo a masa. De esta forma, los diodos evitan que la tensión a la entrada del integrado se aleje más de 0,7 V de estos límites, siguiendo la curva V/I típica del diodo. Pero esto puede provocar corrientes elevadas por los diodos de protección, destruyéndolos y dejando al circuito integrado desprotegido.

A mayor área del diodo de protección, mayor capacidad de conducir corriente. Sólo los *transceivers*, circuitos pensados para ser conectados al exterior, son diseñados con el requisito de elevadas áreas de diodos de protección. La inmensa mayoría de los integrados, expuestos a una descarga ESD, serán destruidos.

De modo que tenemos que añadir protecciones junto a los conectores de E/S para aumentar el nivel de protección. Vamos a estudiar diferentes dispositivos, comenzando por el diodo TVS. Pero te adelanto un par de conclusiones:

- El diodo TVS es sólo la mitad de la solución. Es necesario añadir una impedancia entre el diodo TVS y el circuito integrado que reduzca la corriente del transitorio por los diodos internos de protección. Muchos diseñadores ignoran este segundo elemento, lo que da lugar a protecciones insuficientes.
- Los transitorios ESD son perturbaciones de alta frecuencia. No tener esto en cuenta al conectar las protecciones en el PCB puede hacer que sean completamente inútiles.

Diodos TVS: características, consideraciones y ejemplos de uso

Un diodo TVS (*Transient Voltage Suppressor*) es un diodo Zener con una unión PN de gran área para soportar corrientes instantáneas elevadas. Como contrapartida, también aumenta la capacidad de la unión PN y con ello la carga al circuito, sobre todo en el caso de señales digitales rápidas. ¿Por qué es negativo para una conexión 100 Mb Ethernet que coloquemos un TVS con, digamos, 1 nF de capacidad? Porque eso implica cargar o descargar la capacidad cada 10 ns. Y ya sabes, $I = C \cdot \Delta V / \Delta T = 1 \text{ nF} \cdot 2,5 \text{ V} / 10 \text{ ns} = 250 \text{ mA}$. Creo que es pedirle demasiado al *transceiver*, que tendrá una corriente máxima mucho menor, y como resultado la señal en un par diferencial será un pequeño diente de sierra.

La Figura 10.9 recoge las configuraciones habituales de un diodo TVS. A la izquierda, aquella con la que te sentirás más cómodo. Se emplea cuando no necesitamos una capacidad muy baja y cuando la línea a proteger sea unipolar y positiva. A su derecha, la versión bipolar: con dos TVS enfrentados permitimos que la tensión de trabajo pueda ser positiva o negativa respecto a masa. Cualquier exceso de tensión será recortado por los diodos.

Las dos configuraciones de la derecha, que emplean diodos de conducción, te plantearán dudas. En ambas, el pulso ESD pasa primero por un diodo en directa y entonces alcanza el TVS. La capacidad de los dos diodos queda en serie, y ya sabes que eso implica que la capacidad resultante será menor. Empleando este truco, estas configuraciones se usan para líneas digitales rápidas. La configuración de la derecha va un paso más lejos: si polarizamos el Zener en inversa, hacemos la zona de deplexión mayor: alejamos las cargas y bajamos la capacidad. Pero no se te ocurra conectar directamente a V_{CC} : estarías exponiendo la alimentación al pulso ESD: una resistencia de pull-up aislará alimentación y pulso ESD y mantendrá el TVS polarizado.

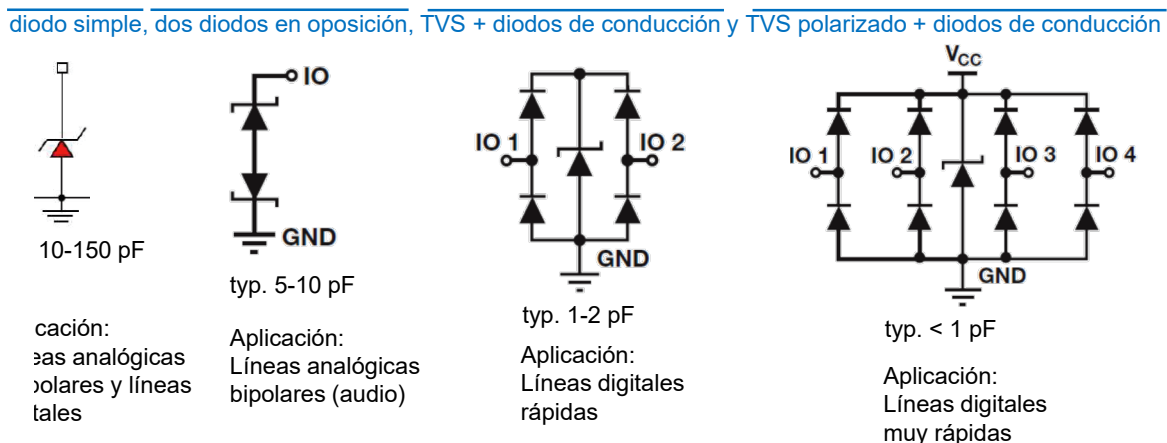


Figura 10.9. Configuraciones habituales de diodos TVS

Ahora que ya tenemos criterios para escoger la configuración, vamos a hablar de cómo debemos rutar el TVS en el circuito impreso. En la Figura 10.10, vemos una pista horizontal. El conector queda a la izquierda, el circuito a proteger a la derecha. El diodo TVS, cuya huella son dos pads SMD, debe ser rutado con longitud cero entre la pista de señal y el cátodo. Es decir, la pista va al cátodo y de aquí se ruta al resto del circuito.

¿Es realmente necesario esto? Vamos a hacer unos números rápidos. Una pista típica tiene 9 nH/cm. Ante un transitorio de 20A/ns, ¿qué tensión cae en la conexión entre la pista a proteger y el TVS si ésta tiene una longitud de 1 cm? Respuesta: $V = L \cdot \Delta I / \Delta T = 9 \text{ nH} \cdot 20 \text{ A} = 180 \text{ V}$. Es decir, si tu TVS debía recortar, por ejemplo, a 20V, ahora lo hará a 200V.

La conexión entre el ánodo y la masa de E/S o chasis debe ser, por los mismos motivos, una conexión de muy baja inductancia.

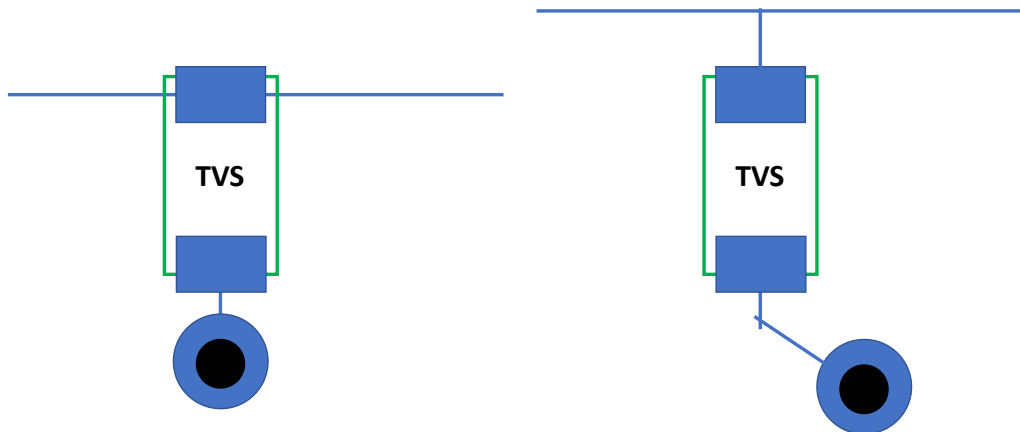


Figura 10.10. El cátodo del diodo TVS debe situarse sobre la línea a proteger. El ánodo debe ir conectado a masa de E/S o chasis con la mínima longitud posible. ¿Por qué? Cada nanohenrio de inductancia de las conexiones eleva significativamente la tensión a la que el TVS recorta la amplitud del pulso ESD. La figura de la izquierda representa un rutado correcto. La figura de la derecha hace un uso incorrecto de la protección

Ejemplo de diodo TVS: On Semiconductor ESD5Z2.5T1G

Vamos a estudiar las características de un TVS a partir de las especificaciones que el fabricante da en su hoja de datos. Sin ninguna razón en particular, escogemos el ESD5Z2.5T1G de On Semiconductor. Lo primero que vemos es su hoja de datos es esto:





n
; **SOD-523**

IEC-61000-4-2 (ESD): $\pm 30\text{kV}$ (contacto y aire)

IEC-61000-4-4 (EFT): 40 A

Surge: (pulso 8x20ms): 240 W

¿Quiere decir el fabricante que protege a mi equipo de una descarga ESD de hasta 30 kV? Para nada. Quiere decir que el TVS soporta este nivel de descarga sin ser destruido. Pero no dice nada sobre el pulso residual que deja pasar hacia tu electrónica ni sus efectos. Así que no te confíes por leer “30 kV” en la especificación ESD.

Además de su resistencia a descargas ESD, se nos dice también su resistencia a los otros dos tipos usuales de transitorios de alta tensión: ráfagas (*electrical fast transient*, EFT) y onda de choque (surge).

Este TVS de pequeño tamaño está disponible en versiones con tensiones de trabajo entre 2,4V y 12V.

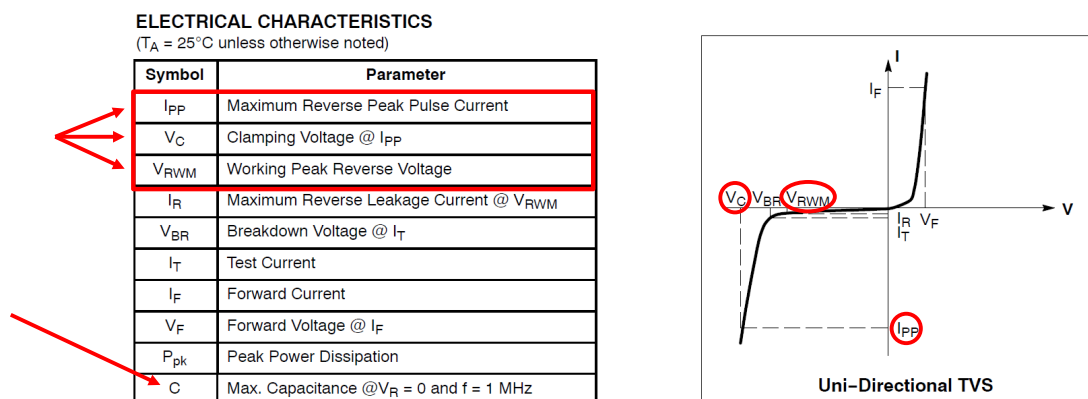


Figura 10.11. Parámetros principales de un diodo TVS

A la hora de elegir un TVS has de comprobar que la **tensión de trabajo en inversa** es compatible con tu aplicación. Por ejemplo, en la Figura 10.12, a la tensión de trabajo de 2,5V la corriente de fugas es de 6 μA . Si la señal útil tiene una amplitud de 2,5V o menos y puedes permitirte esta corriente de fugas, estupendo. Si no es así, puedes elegir una versión con una tensión de trabajo ligeramente superior para asegurarte una corriente de fugas más reducida.

Debes verificar también que la **capacidad del TVS** no va a cargar en exceso la línea. En la Figura 10.12 se indica una capacidad típica de 145 pF: no será adecuado para líneas ni buses de frecuencia de más de una decena de megahercios.

Otro parámetro en el que fijarte será la tensión de recorte (V_C , *clamping*). El producto de esta tensión y la corriente que circula por el diodo en esa condición suele coincidir con la potencia de pico que puede disipar el diodo. En este caso, los valores son de 10,9 V, 11 A y 120 W. Pero no debes pensar que la tensión máxima entre terminales del diodo será tan baja. La Figura 10.12, en su parte inferior, muestra las formas de onda entre terminales cuando aplicas un pulso ESD de 8 kV de polaridad positiva y negativa. El pico alcanza casi 40V, pero es muy estrecho. Le sigue una meseta de unos 240 ns de anchura y una tensión cercana a 10V.

ELECTRICAL CHARACTERISTICS ($T_A = 25^\circ\text{C}$ unless otherwise noted, $V_F = 1.1\text{ V Max.}$ @ $I_F = 10\text{ mA}$ for all types)

Device*	Device Marking	V_{RWM} (V)	I_R (μA) @ V_{RWM}	V_{BR} (V) @ I_T (Note 2)	I_T	V_C (V) @ $I_{PP} = 5.0\text{ A}^\dagger$	V_C (V) @ Max $I_{PP}^\dagger$	I_{PP} (A) [†]	P_{pk} (W) [†]	C (pF)
		Max	Max	Min	mA	Typ	Max	Max	Max	Typ
ESD5Z2.5T1G	ZD	2.5	6.0	4.0	1.0	6.5	10.9	11.0	120	145

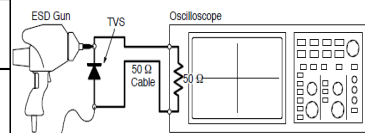


Figure 4. Diagram of ESD Test Setup



Figure 1. ESD Clamping Voltage Screenshot Positive 8 kV contact per IEC 61000-4-2

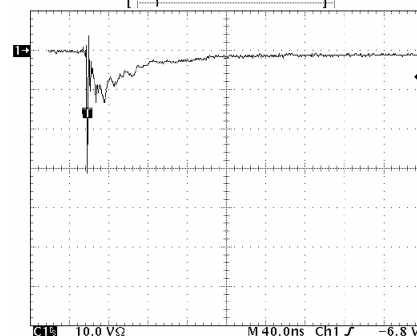


Figure 2. ESD Clamping Voltage Screenshot Negative 8 kV contact per IEC 61000-4-2

Figura 10.12. Tensión de trabajo y tensión de recorte (*clamping*)

El pico tiene poca energía, por ser tan estrecho. La parte peligrosa del pulso residual es la meseta. Volveremos a esto más tarde, vamos primero a echar un vistazo a algunos ejemplos de aplicación.

Aplicaciones: protección para SD/ μSD

Una solución para esta interfaz, que añade algo de filtrado EMI, es el dispositivo Infineon BGF117. Esta protección, con 8 pF de carga capacitiva, es adecuada para un bus SD que por especificación debe ser cargado con un máximo de 40 pF.



El filtro en pi (Figura 10.13) ayuda a reducir ruido conducido desde otras partes del sistema. Los diodos bidireccionales limitan la amplitud de la ESD. Las resistencias de pull-up ahorran espacio en placa.

Las resistencias serie de 40 ohmios complementan a los diodos, limitando la corriente que circula hacia el circuito de interfaz.

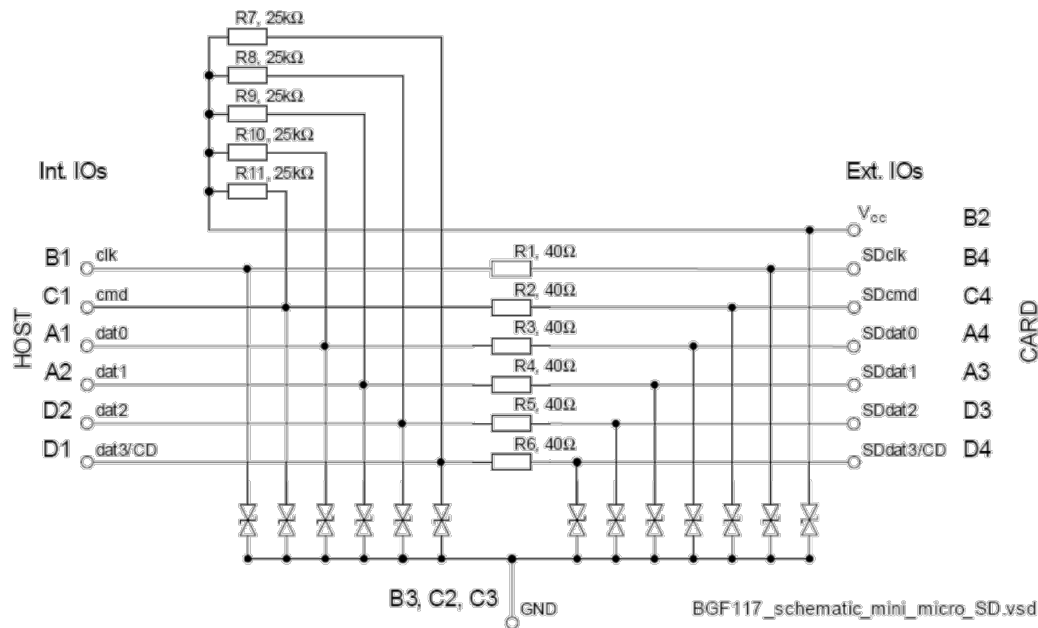


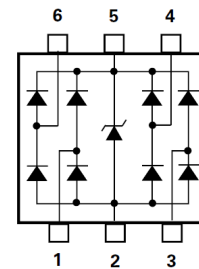
Figura 10.13. Estructura interna de la protección Infineon BGF117

Aplicaciones: protección para SD/μSD (un segundo ejemplo)

Ejemplo de diseño real basado en SP3050 de Littelfuse

Objetivo: protección ESD

- Capacidad: 2pF
- ESD, IEC61000-4-2, ±20kV contact, ±30kV air
- EFT, IEC61000-4-4, 40A (5/50ns)
- Lightning, IEC61000-4-5, 10A (8/20μs)



Esta protección para cuatro líneas es de baja capacidad (2 pF) y puede usarse con el pin 5 (cátodo del Zener) al aire o conectado a alimentación con una resistencia de pull-up, en función de si necesitamos o no una capacidad eléctrica muy reducida.

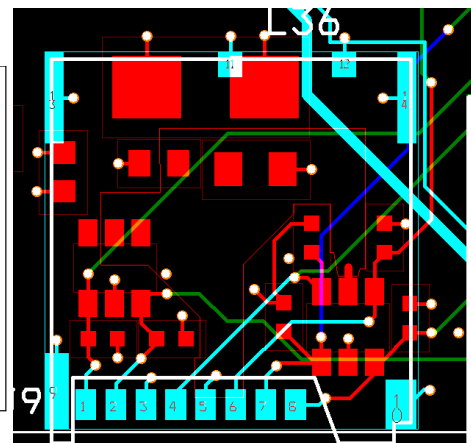
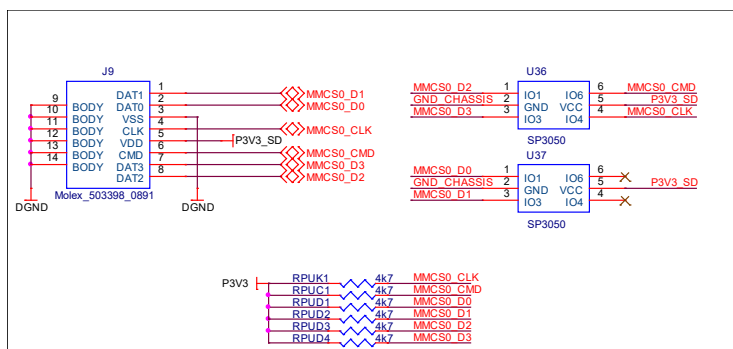


Figura 10.14. Diagrama esquemático y layout de un conector de tarjeta de memoria micro-SD y de sus protecciones

En la Figura 10.14 puedes ver un ejemplo de diseño real con esta protección donde se comenten algunos errores. Las protecciones ESD son los dos integrado de 6 pines en capa *bottom* (rojo). El producto superó los ensayos ESD en laboratorio, pero el diseño era mejorable a coste cero. **¿Puedes encontrar un par de errores?** Después sigue leyendo.

En primer lugar, el pin 5 de las protecciones está conectado directamente a la alimentación (P3V3_SD) y no a través de un pull-up. En este diseño en concreto no era tan grave porque esta alimentación está separada por una ferrita de P3V3, que alimenta muchos otros integrados del producto. La ferrita proporcionaba algo de aislamiento, aunque, insisto, mejor haber usado una resistencia de pull-up.

En segundo lugar, la línea que parte del pin 4 del conector viola la regla que enunciamos al final de la página 217: la línea debe ir hasta la protección, y del *pad* de la protección parte la pista que va al interior de nuestro diseño. Es cierto que la violación es muy pequeña, un par de milímetros tal vez, pero degrada las prestaciones de la protección.

Observa también cómo el ánodo del diodo TVS está conectado a chasis y no a masa de circuito. Se trataba de un producto embarcado atornillado a una estructura metálica (chasis). Otra mejora en el diseño sería llevar también a chasis los pines 9-14 (cuerpo) del conector de memoria micro-SD. Así, cualquier descarga ESD iría a chasis y no a la masa del circuito. Otra mejora más: el pin 3 del conector (masa) también debería llevar una protección TVS.

Aplicaciones: Ethernet

La misma protección SP3050 de Littlefuse se puede utilizar también en una interfaz Ethernet. En la Figura 10.15, izquierda, observamos los choques en modo común (CMC1, CMC2) así como los transformadores que suele incluir internamente un conector Ethernet. Estos elementos atenúan interferencia y perturbaciones, lo que nos proporciona una ayuda. Curiosamente, en Ethernet sólo usamos dos pares diferenciales (transmisión y recepción) de los cuatro que hay en una conexión RJ-45. Los dos pares no utilizados no se dejan flotando, ya que el ruido acoplado en estas líneas podría propagarse e ir rebotando de un extremo a otro de la línea. En su lugar, se ponen terminaciones para el modo común (R3, R4, 75 ohmios). También se termina el modo común para los pares diferenciales utilizados (R1, R2, toma intermedia de los transformadores).

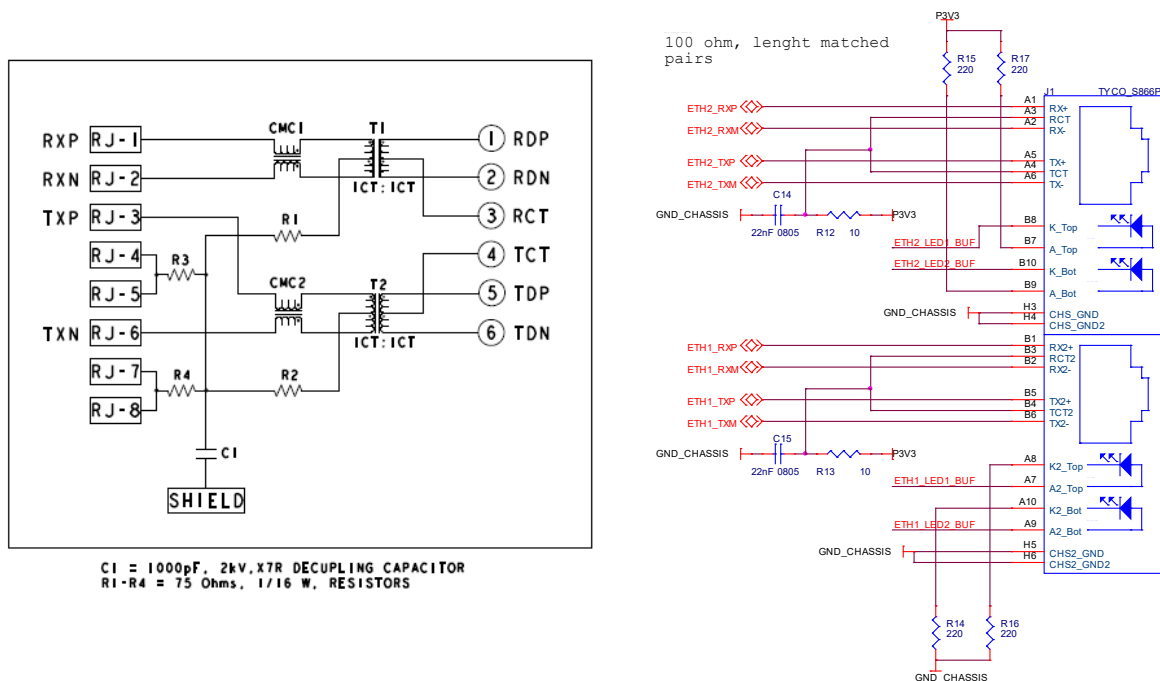


Figura 10.15. Esquema interno típico de un conector Ethernet con transformadores internos (izquierda). Diagrama esquemático del ejemplo de diseño (derecha)

En la Figura 10.15, derecha, tienes el diagrama esquemático del diseño rutado en la Figura 10.16, un conector Ethernet doble. La pregunta que te hago es: ¿están adecuadamente posicionadas y rutadas las protecciones ESD (las dos huellas de 6 pines en capa bottom, en rojo)?

La respuesta es un clamoroso no. Los cuatro pares diferenciales parten de los pines A1-A2, A5-A6, B1-B2 y B5-B6. Deberían partir de las protecciones ESD. El camino debería ser conector-protecciones-circuitos. Los 5-8 mm, *grosso modo*, que hay entre los pines y los SP3050 degrada enormemente las prestaciones de las protecciones.

Y, sin embargo, ese diseño también superó con éxito los ensayos ESD. A veces la electrónica es muy agradecida, aunque la tratemos a patadas. Otras veces no hay forma de que las cosas vayan bien, aunque la tratemos con mimo. Pero es una cuestión de probabilidades. En este caso hubo suerte. Y ayuda: los *transceivers* Ethernet suelen sobrevivir a descargas ESD de 2 kV.

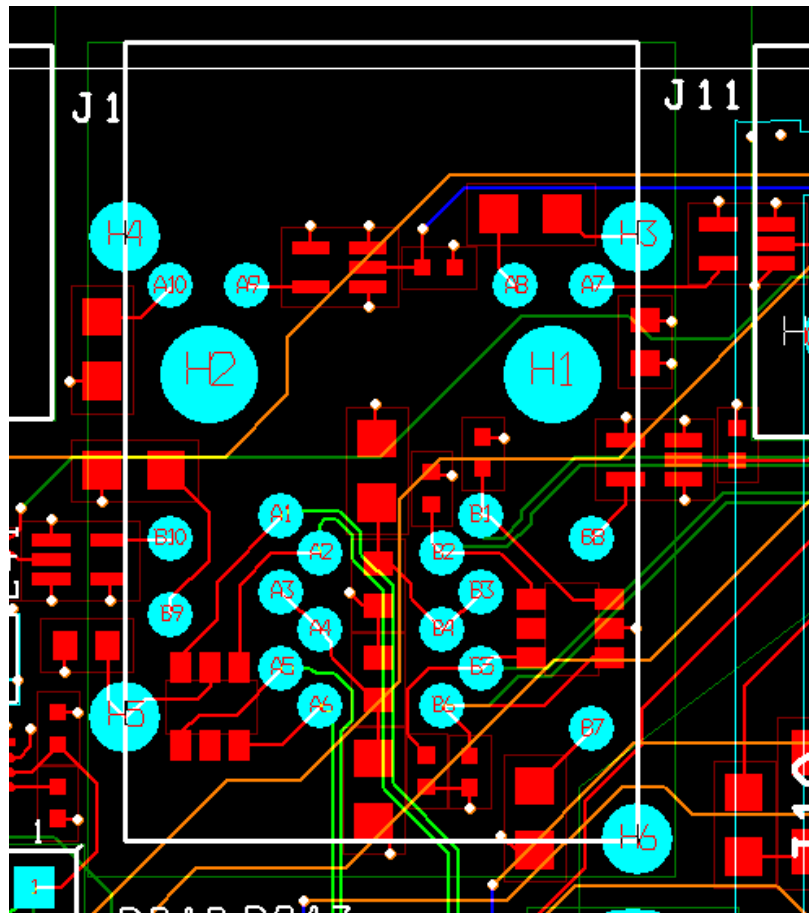


Figura 10.16. Ejemplo de diseño mal rutado: conector RJ-45 doble con protecciones ESD

No hemos hablado sobre si es buena o mala práctica extender el plano de masa de circuito bajo el conector Ethernet, o si debemos tener chasis y no masa bajo el conector, o si no debemos tener nada. Equilibrar los requisitos de EMC (como derivar las perturbaciones ESD fuera de la masa limpia) y de integridad de señal (como es la continuidad de las corrientes de retorno) no tiene solución única, a menudo distintas notas de aplicación dan recomendaciones en aparente conflicto y lo mejor que podemos hacer es leer y reflexionar. Te dejo dos referencias para leer si te asaltan las dudas: [27] y [28].

Aplicaciones: USB 2.0

En este caso usaremos una protección que no es realmente un diodo TVS, sino una resistencia variable con la tensión (VVR), y por tanto bidireccional. Se trata del dispositivo PGB2010402 de Littlefuse. Según su

hoja de datos, la tensión de *clamping* es 40V, la capacidad es de tan sólo 70 fF. Con una huella 0402, es realmente pequeño y cómodo de usar.

Si estudias la Figura 10.17, podrás identificar la disposición de cada elemento de la protección y cómo se ha realizado el rutado. Hay una pequeña modificación que supondría una mejora. Te dejo descubrirla, ¡ya me comentas por correo electrónico!

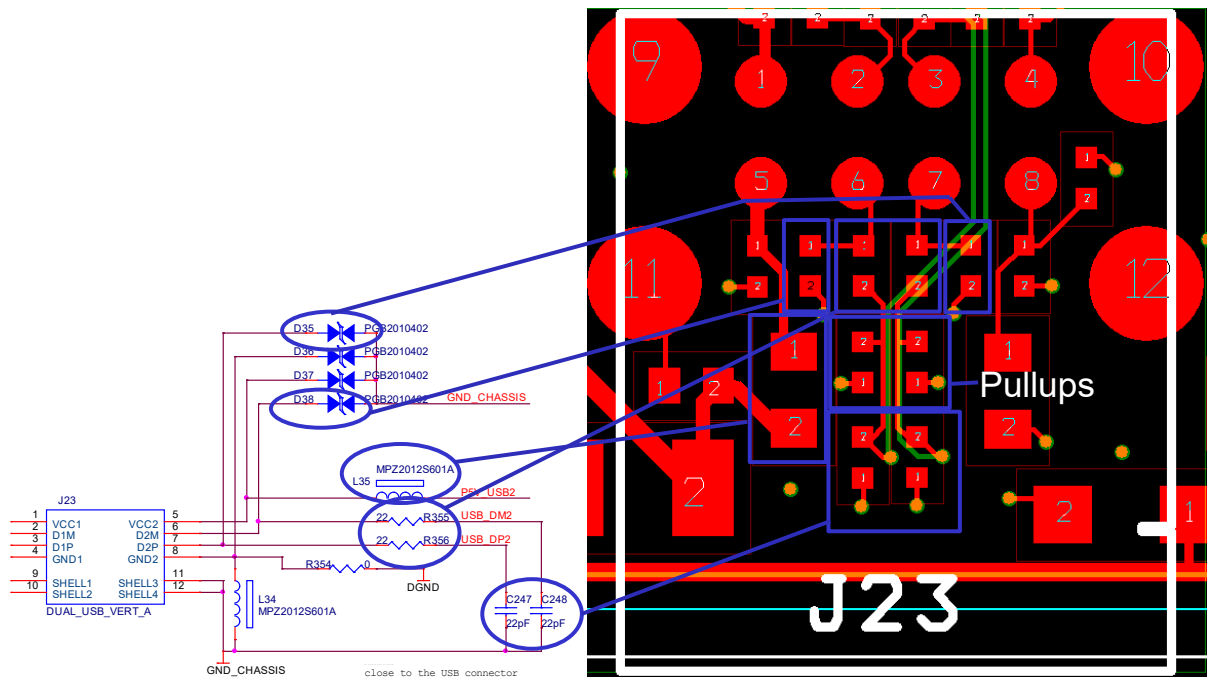


Figura 10.17. Diagrama esquemático y layout de un conector USB doble y sus protecciones

Igual te extraña la conexión de GND2 del conector a masa del circuito (GND) mediante una resistencia de 0 ohmios y a chasis (GND_CHASSIS) mediante una ferrita, a la vez que hay una protección ESD. Esta configuración te da la flexibilidad para conectar GND2 directamente a GND, o mediante una ferrita. O para llevar a chasis con una resistencia de 0 ohmios o con una ferrita. Cada instalación final tiene sus características. Si conoces el tipo de instalación que harán tus clientes y el tipo de ruido que puede estar presente, decidirás montar por defecto una solución u otra.

El TVS es sólo la mitad de la red de protección: añadiendo una impedancia en serie

En la página 219 vimos que un diodo TVS no recorta perfectamente todo lo que exceda de la tensión de trabajo, ni siquiera de la tensión de *clamping*. El diodo deja pasar una onda residual que se caracteriza por un pico alto y estrecho de varias decenas de voltios y una meseta mucho más ancha de menos amplitud. En la Figura 10.18 se reproduce esta forma de onda residual, cuya parte más peligrosa no es el pico, sino la meseta.

El equivalente es una fuente de tensión de varios voltios de amplitud y una resistencia del orden del ohmio. La corriente a través de las protecciones en el circuito integrado puede ser muy alta y es necesario limitarla. ¿Cómo? Mediante una impedancia serie (resistencia o ferrita típicamente) tras el diodo TVS.

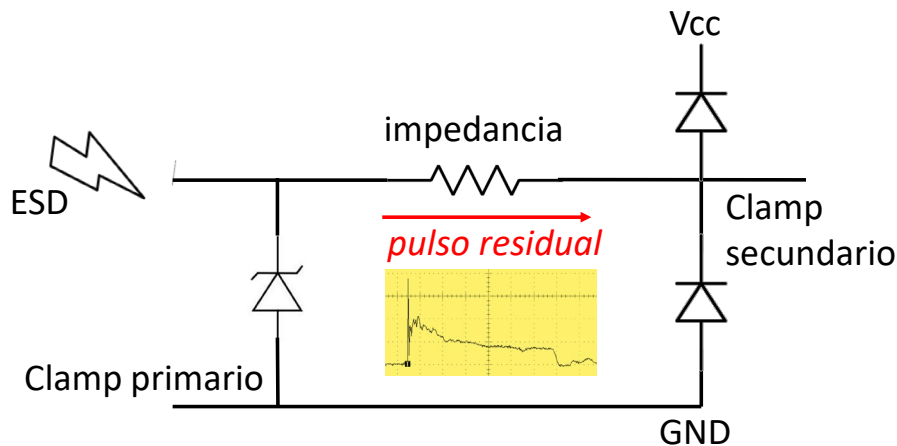


Figura 10.18. La onda residual que deja pasar el diodo TVS debe ser atenuada por una impedancia en serie que reduzca la corriente que pasa por los diodos de protección del circuito integrado (*clamp secundario*). Aunque usemos un TVS de 5V, la meseta de esta onda residual puede tener 7-25 V de amplitud.

Vamos a hacer un pequeño experimento virtual mediante una simulación SPICE. Te cuento cómo montamos el modelo. En primer lugar, tomamos el modelo de la fuente ESD que presentamos en la página 210. A continuación, montamos el diodo TVS (cuyo modelo SPICE hemos de descargar de la página web del fabricante).

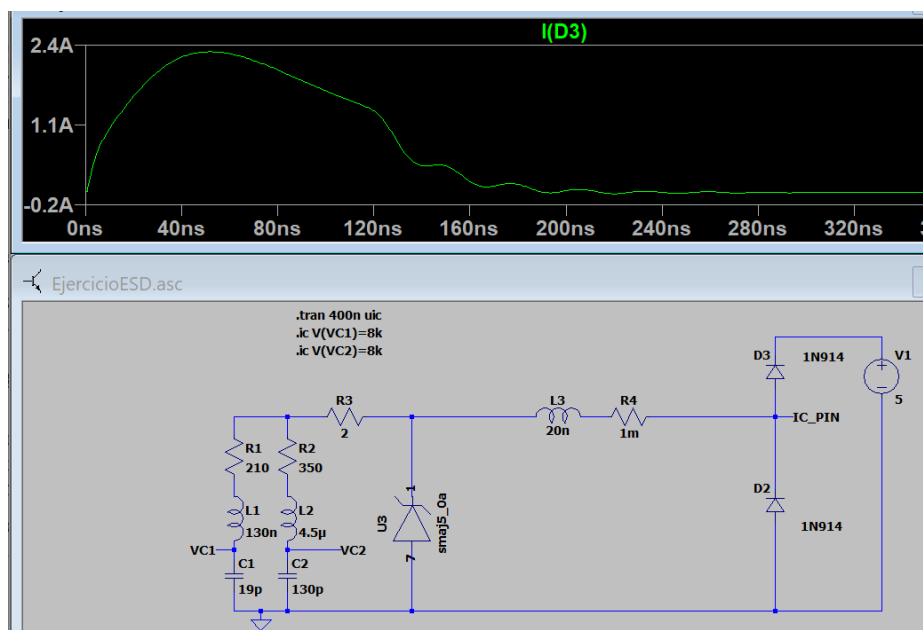


Figura 10.19. Ejemplo de simulación de protecciones ESD con SPICE. Captura de pantalla

Tras el TVS, hemos considerado la inductancia de 2-3 cm de pista hasta el circuito integrado a proteger. ¿Cómo lo modelamos? Muy fácil: un diodo a Vcc (cualquier diodo normal nos sirve, sólo queremos evaluar la corriente que lo atraviesa) y otro diodo a masa.

Bien, la corriente por el diodo de protección del circuito integrado alcanza 2,4A para un pulso de polaridad positiva, 1,5A para polaridad negativa (hay que simular ambos), lo que es excesivo y puede provocar la destrucción del integrado. Añadiendo una resistencia de 56 ohmios, limitamos la corriente unos 62 mA: una reducción de dos órdenes de magnitud en la corriente y la viabilidad de la protección al precio de un céntimo.

Si eres de los que suelen colocar sólo el diodo TVS, sin impedancia serie tras él, y has tenido problemas con los ensayos ESD, apreciarás la potencia e inmediatez de esta metodología de diseño de protecciones basada en simulación SPICE.

¿Cómo sabemos cuánta corriente pueden soportar los diodos de protección del IC?

No esperes encontrar este dato en la hoja de datos de un circuito integrado, es algo muy infrecuente. Si no nos da esa información (lo más habitual) podemos obtenerla destruyendo unos cuantos integrados y midiendo a qué corriente de pico se producen los daños.

Una tercera posibilidad es **mantener, para circuitos integrados digitales, la corriente de pico en el transitorio por debajo de 100 mA**. Tras consultar a algunos diseñadores microelectrónicos, este valor de consenso nos ofrece una referencia para las simulaciones SPICE.

En OPAMPs y otros circuitos integrados analógicos, si no hay información sobre el nivel de tolerancia a ESD, es una buena práctica limitar la corriente transitoria a 3-5 mA.

Otra posibilidad es que en la hoja de datos venga especificado un nivel de tolerancia a descarga ESD. En ese caso, podemos hacer un poco de ingeniería inversa: simulamos la descarga de la tensión especificada en los diodos de protección y medimos la corriente de pico que han de soportar: este valor de dará el orden de magnitud de la corriente que puede soportar el integrado.

¿Qué tipo y valor de impedancia serie podemos añadir?

Con carácter general, para líneas digitales rápidas montamos en el PCB el diodo TVS y una resistencia de un valor (22-100 ohm) tan alto como permita la aplicación. En líneas digitales lentas podemos aumentar el valor de la resistencia, incluso estudiar la viabilidad de un circuito RC, que puede llegar a hacer innecesario (en líneas muy lentas o cuasi-estáticas) el diodo TVS.

En líneas de alimentación no podemos usar resistencias (ya sabes, por aquello de la ley de Ohm), pero puedes usar ferritas.

Simulaciones SPICE

¿Qué simulador?

En principio, cualquier simulador SPICE es válido. Te recomiendo usar un simulador SPICE gratuito como LTspice. Puedes descargar el simulador desde la web de varias empresas, como Analog Devices o Würth Elektronik. Aquí tienes los enlaces para la descarga:

- <https://www.analog.com/en/design-center/design-tools-and-calculators/ltspice-simulator.html#>
- https://www.wuerth-elektronik.com/web/en/electronic_components/produkte_pb/bauteilebibliotheken/ltspice/ltspice.php

Tal vez el aspecto que menos me gusta es el editor gráfico, pero otros como la importación de *netlists* y creación de nuevos símbolos están bastante bien logrados.

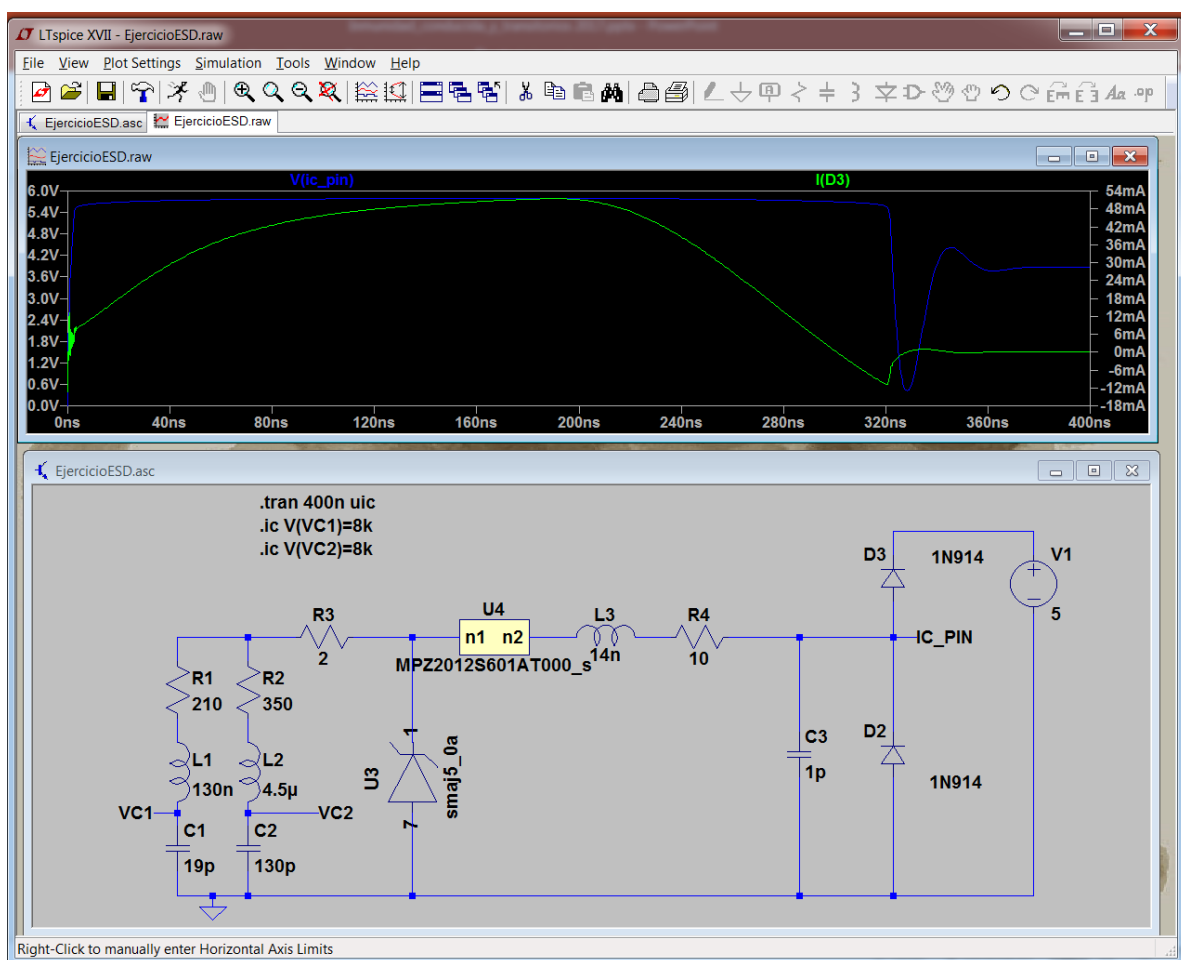


Figura 10.20. LTspice XVII

Importar modelos y crear símbolos en LTspice

Una vez has descargado en un fichero de texto el modelo SPICE del diodo TVS, ferrita o de otro tipo de componente, basta con abrirlo desde LTspice, pinchar en la línea en la comienza la descripción del subcircuito (comando .SUBCKT), hacer clic con el botón derecho del ratón y seleccionar la opción “Create Symbol” del menú desplegable.

Una ventana emergente te preguntará si quieres crear un símbolo a partir de este subcircuito. Acepta y LTspice creará un símbolo sencillo que podrás instanciar.

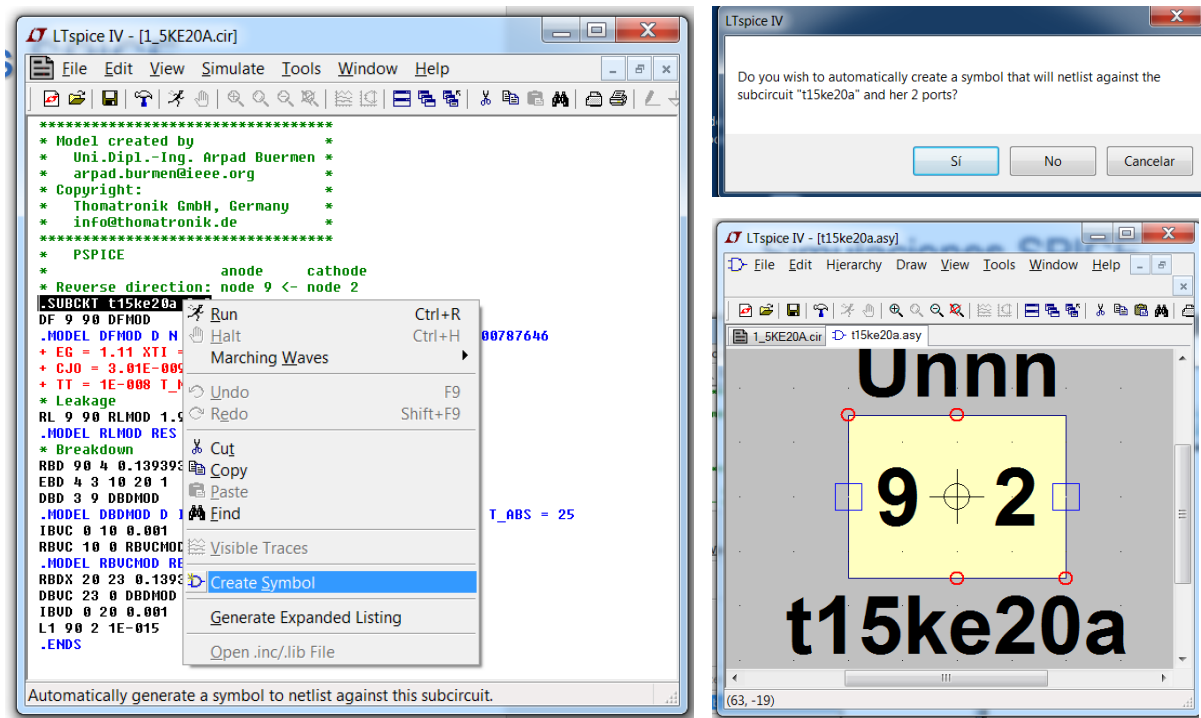


Figura 10.21. Pasos para importar un modelo SPICE y crear un nuevo símbolo. Captura de pantalla

Modelo de una ferrita

Un circuito RLC paralelo con una pequeña resistencia serie (del orden de decenas de miliohmios) es una aproximación aceptable al comportamiento de una ferrita. Un modelo más preciso añade un segundo elemento RLC paralelo en serie con los anteriores.

Si observas la Figura 10.22, los dos modelos para una ferrita MPZ1608S601A no muestran una diferencia muy acusada, de modo que aceptaremos tanto modelos sencillos como complejos para nuestras simulaciones.

Por debajo de la frecuencia de resonancia, la ferrita se comporta como una inductancia, pero las resistencias evitan que entre en resonancia con otras capacidades cercanas. Esta es una de las grandes ventajas de una ferrita frente a una inductancia: da lugar a circuitos resonantes de bajo factor de calidad, lo que evita *ringing* y otros efectos normalmente indeseados. En resonancia, la impedancia es elevada (típicamente hasta unos cientos de ohmios). Por encima de la resonancia su comportamiento es capacitivo.

La ferrita del *netlist* (MPZ1608S601A de TDK) está pensada para líneas de alimentación, si bien es útil también en líneas de señal. Su capacidad de corriente es de 1A (hasta 85°C) y presente una impedancia de 600 ohmios a 100 MHz. Debes elegir con cuidado las ferritas, en función de qué margen de frecuencias quieres atenuar preferentemente y la corriente en continua que deba soportar.

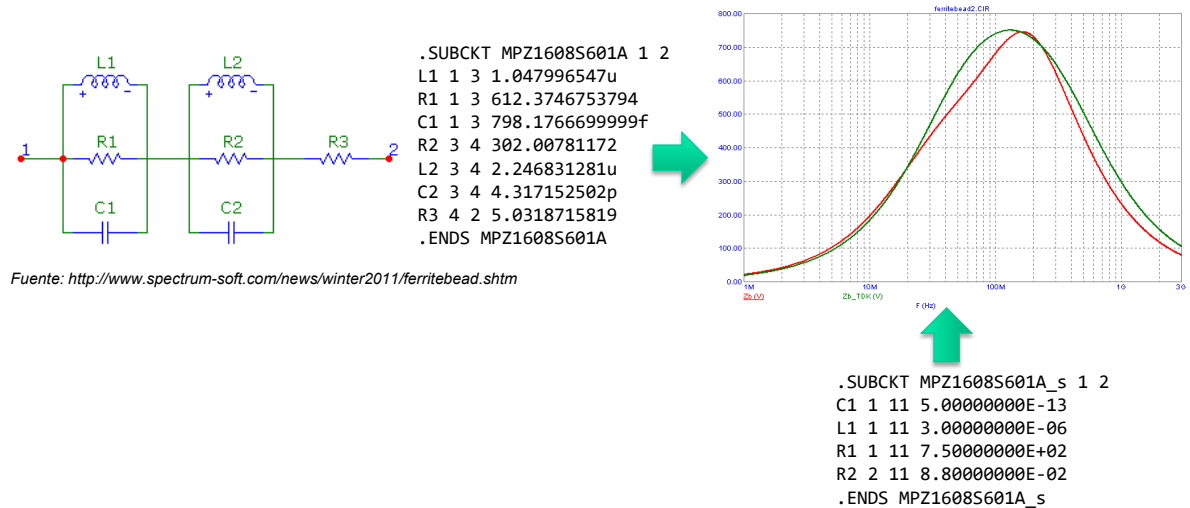


Figura 10.22. Modelo sencillo de una ferrita (circuito RLC paralelo con resistencia serie)

Debes comenzar por elegir el material correcto para la ferrita. Normalmente hay tres tipos de materiales:

- **Polvo de hierro:** Sólo útil para frecuencias bajas y medio-bajas (20 MHz)
- **Manganeso-zinc (MnZn):** Útil hasta frecuencias en torno a 80 MHz, son utilizadas para filtrar interferencias conducidas (150 kHz - 30 MHz)
- **Níquel-zinc (NiZn):** permite atenuar señales hasta 1 GHz, depende del fabricante y del modelo, y son por tanto las más usadas en EMC

Modelo de un diodo TVS – ejemplo: SMAJ5.0A

Como el diodo es un componente soportado por cualquier simulador SPICE, el modelo de un diodo TVS es muy sencillo. Ahí va un ejemplo:

```
.SUBCKT SMAJ5.0A 1 2
* TERMINALS: A K
Done 1 2 DtvS
Rleak 1 2 0.013meg
.MODEL DtvS D (IS=1.0e-5 RS=0.0391 N=1.5 IBV=10m BV=6.46 CJO=1400p)
.ENDS
```

Part Number (Uni)	Part Number (Bi)	Marking		Reverse Stand off Voltage V_R (Volts)	Breakdown Voltage V_{BR} (Volts) @ I_T		Test Current I_T (mA)	Maximum Clamping Voltage V_C @ I_{pp} (V)	Maximum Peak Pulse Current I_{pp} (A)	Maximum Reverse Leakage I_R @ V_R (μA)	Agency Approval
		UNI	BI		MIN	MAX					
SMAJ5.0A	SMAJ5.0CA	AE	WE	5.0	6.40	7.00	10	9.2	43.5	800	X
Parameter		Symbol	Value	Unit							
Peak Pulse Power Dissipation at $T_A=25^\circ C$ by 10x1000 μs waveform (Fig.1)(Note 1), (Note 2)		P_{PPM}	400	W							
Power Dissipation on infinite heat sink at $T_A=50^\circ C$		$P_{M(AV)}$	3.3	W							
Peak Forward Surge Current, 8.3ms Single Half Sine Wave (Note 3)		I_{FSM}	40	A							
Maximum Instantaneous Forward Voltage at 25A for Unidirectional only (Note 4)		V_F	3.5V/6.5	V							

Figura 10.23. Características principales del ampliamente utilizado SMAJ5.0A, producido por varios fabricantes

Ejercicio

Determina, mediante simulación, si la ferrita es necesaria para protegerlo frente a transitorios ESD.

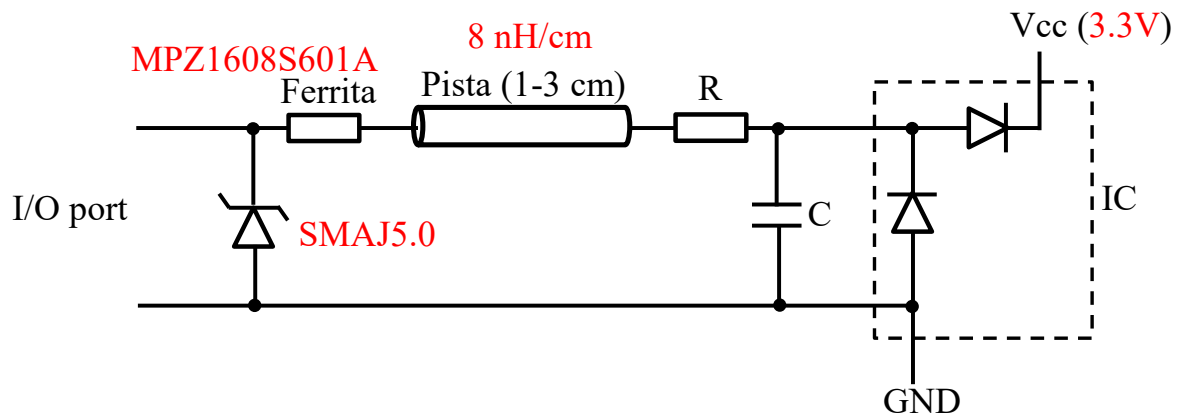


Figura 10.24. Esquema para el ejercicio de simulaciones ESD con LTspice

Solución

En primer lugar, debemos descargar los modelos SPICE de la ferrita y del diodo TVS. Si buscas por “MPZ1608 TDK SPICE” llegas fácilmente a la siguiente página, donde puedes descargar un fichero .zip para ferritas MPZS601: <https://product.tdk.com/info/en/technicalsupport/tvcl/general/beads.html>.

Al extraer el .zip, encuentras el fichero “MPZ1608S601ATA00_s.mod”. Se trata de un fichero de texto y lo puedes abrir desde LTspice. Crea un nuevo símbolo asociado a este modelo (mira cómo en la página 227) y guárdalo en el directorio de trabajo. Puedes editarlo antes de guardarlo. Sigue los mismos pasos para importar el modelo del SMAJ5.0A.

Ahora puedes crear el esquema (si es tu primera vez con LTspice tendrás que explorar un poco y buscar algún tutorial en internet, que es lo que tuve que hacer yo en su momento, no te quejes). El esquema debe quedarte más o menos así:

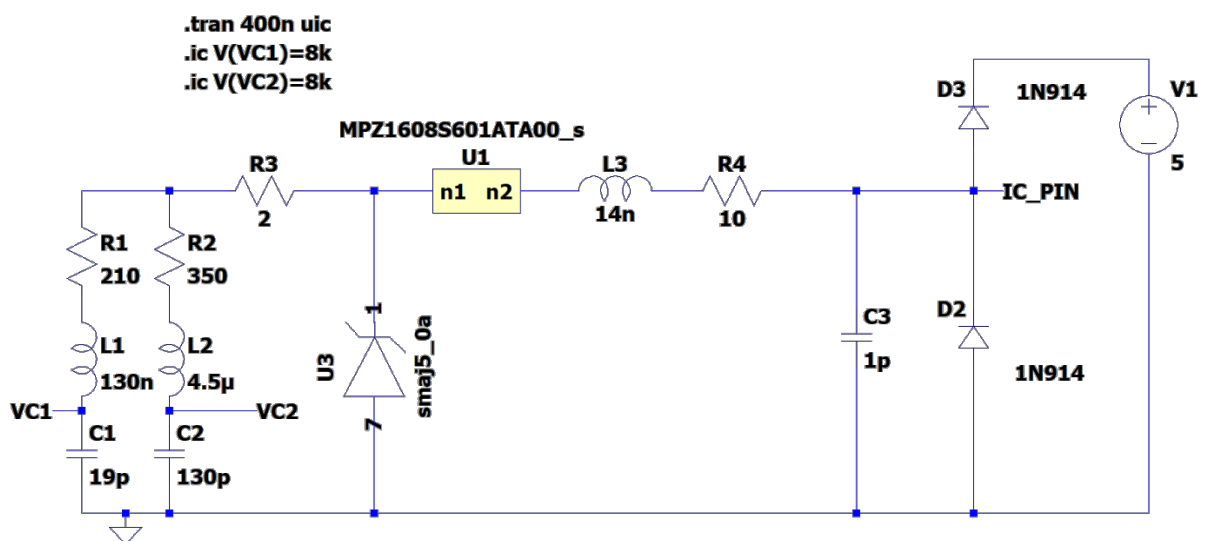


Figura 10.25. Esquema del modelo

Hay dos directivas del simulador que debes conocer:

`.tran 400n uic`, dice al simulador que realice un análisis transitorio de 400 ns usando las condiciones iniciales (*uic, use initial conditions*).

`.ic V(VC1)=8k`, dice al simulador que la condición inicial (*ic, initial condition*) para el nodo VC1 (condensador C1) es de 8 kV. Variando las condiciones iniciales de VC1 y VC2 cargas tu pistola ESD virtual a la tensión deseada. Por ejemplo, “-8k” para un pulso de polaridad negativa.

Ahora ya puedes simular el circuito y visualizar la corriente por el diodo D3, que apenas supera los 50 mA, lo que sería aceptable.

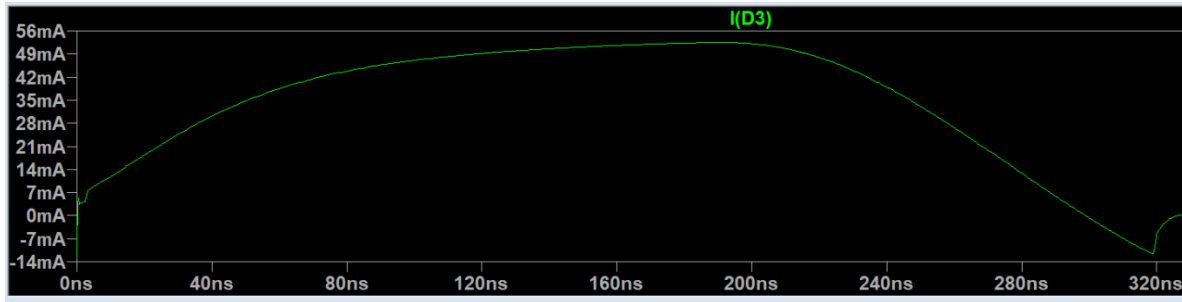


Figura 10.26. Simulación del esquema de la Figura 10.25

Si eliminamos la ferrita, el pico de corriente supera los 275 mA. Excesivo (recuerda que habíamos dado como referencia no superar 100 mA en circuitos integrados digitales). Si aumentamos R4 a 56 ohm, la corriente de pico en D3 no supera 65 mA. De modo que la respuesta es que sí, se puede eliminar la ferrita.

Protecciones basadas en aislamiento

Para terminar este día de estudio, vamos a echar un vistazo a otro tipo de soluciones basadas en aislamiento para proteger nuestro equipo frente a ESD y otros transitorios de alta tensión. Tomemos como ejemplo el transceiver RS-485 aislado de Analog Devices ADM2461E ([enlace](#) a la hoja de datos).



Proporciona un aislamiento de 5,7 kV (hasta 1 minuto, bajando la tensión a poco más de 1 kV si se aplica de manera continua). Para ello debemos garantizar al menos 8,3 mm de separación entre elementos conductores (pines, pistas) de los lados aislados, ya sea a través del aire o del cuerpo del dispositivo. Puede trabajar hasta 125°C, incorpora protección a ESD de ± 12 kV (descarga por contacto) y ± 15 kV (descarga al aire). Según el fabricante, supera los ensayos de emisiones radiadas (no dice nada sobre conducidas) EN 55032 Clase B (entorno residencial

más restrictivo, recuerda) en un PCB de dos capas. Eso sí, reduciendo la velocidad de transmisión a 500 kbps.

El diagrama funcional de este *transceiver* (Figura 10.27) muestra el bloque de protecciones ESD y el de aislamiento. También muestra alimentaciones (VDD1, VDD2) para cada uno de los lados aislados, lo que requiere una alimentación aislada: una complicación adicional.

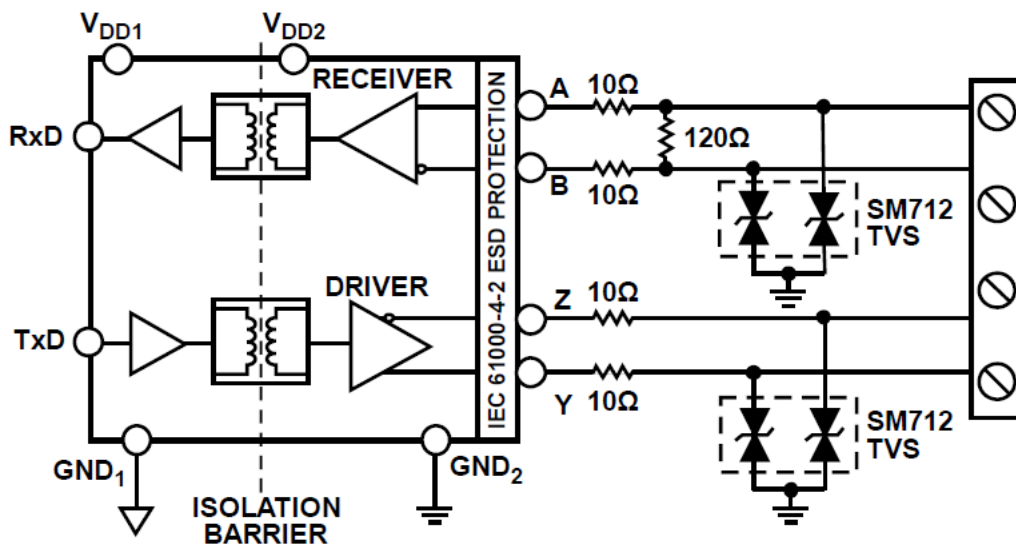


Figura 10.27. Diagrama funcional del *transceiver* aislado RD-485 ADM2461E, con protecciones adicionales para transitorios EFT y onda de choque

Para proporcionar protección frente a transitorios de mayor energía (los estudiaremos mañana) el fabricante propone añadir diodos TVS bipolares, en este caso del tipo SM712. Se trata de protecciones específicas para RS-485 disponibles a través de varios fabricantes (te dejo en [enlace](#) a su hoja de datos).

Día 11. Diseño de protecciones para EFT y onda de choque



Esta imagen de dominio público (versión original de Mircea Madau), tomada en las afueras de Oradea, Rumanía, el 17 de agosto de 2005, recoge la descarga de rayos durante una tormenta que causó fuertes inundaciones en el sur de país.

Normativa para EFT

Norma básica

IEC 61000-4-4, EFT: norma básica de aplicación a productos que puedan sufrir transitorios eléctricos rápidos en ráfagas en conexiones de alimentación, señal y tierra.

EFT (acrónimo de *electrical fast transient*) es uno de los tres transitorios de alta tensión definidos en las normas de inmunidad para una amplia gama de productos. Este transitorio emula el efecto de conmutación de cargas inductivas, relés y contactores en las líneas de alimentación AC. Los transitorios se propagan por la red de alimentación, acoplándose al equipo víctima por conducción. Esta perturbación puede acoplarse capacitivamente a cables de señal y de alimentación DC si los cables van junto a los de alimentación (por ejemplo, en la misma canaleta).

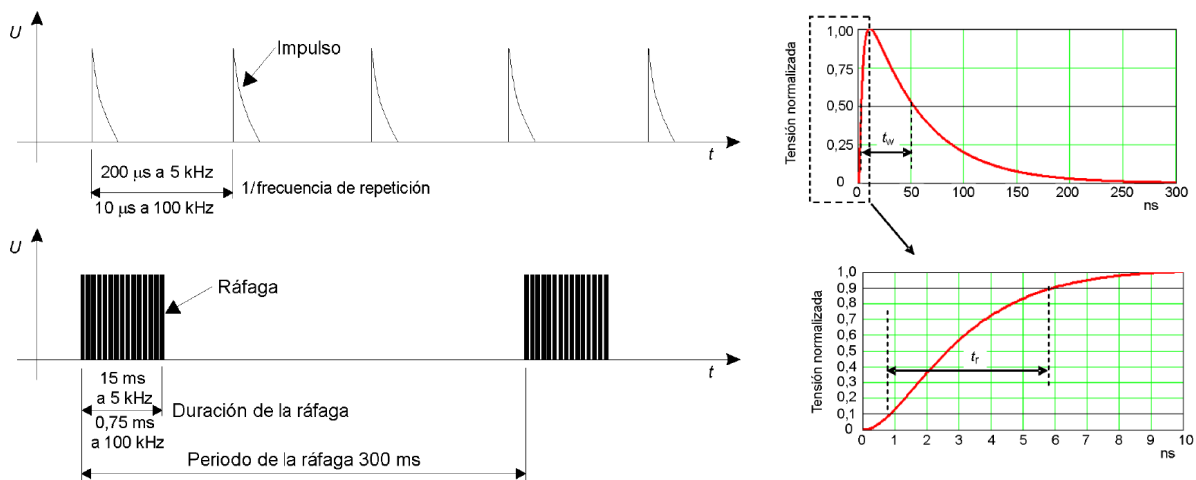


Figura 11.1. Características de la forma de onda de un pulso, de una ráfaga y de un tren de ráfagas EFT

Un pulso EFT se define con características que lo hacen muy similar a un pulso ESD: un tiempo de subida breve (cerca de 1 ns en ESD, cerca de 5 ns en EFT), anchura de pulso similar (50 ns hasta caer al 50% de amplitud en tensión) y tensión de pico típicamente comprendida entre medio y dos kilovoltios (4 u 8 en ESD). Por tanto, una protección frente a ESD servirá necesariamente frente a EFT.

A menudo se critica que esta forma de onda acoplada se aleja de la realidad, donde tras varios metros de cableado, el filtrado paso bajo y las resonancias la cambian por completo. Pero es lo que dicta la norma y por tanto lo que debe aplicarse en el ensayo.

Una diferencia relevante entre ambos tipos de transitorios reside en que si bien el ensayo ESD implica aplicar pocas (típicamente diez) descargas espaciadas al menos un segundo, en el caso de EFT deberemos aplicar un tren de ráfagas de en torno a un minuto de duración (para cada polaridad), estando las ráfagas separadas 300 ms entre sí, constando cada ráfaga de típicamente 75 pulsos separados 200 μs, aunque estas cifras pueden cambiar en función del tipo de producto.

Si bien un pulso EFT tiene una energía moderada, unos 4 mJ sobre una carga de 50 Ω, debemos tener en cuenta que en 300 ms se producen 75 pulsos y la potencia equivalente (J/s) es de 1 W. Este es un dato importante a la hora de diseñar las protecciones, pues componentes SMD en encapsulados de pequeño tamaño como puede ser un tipo 1206 (frecuente en varistores o incluso en diodos TVS) no son capaces de disipar esta potencia al ambiente sin elevar excesivamente su temperatura y sufrir daños.

IEC 61000-4-4 define la forma de onda del pulso, las características de la ráfaga, del equipo, instalación y procedimiento de ensayo. La Tabla 11.1 recoge los niveles de tensión de pico aplicables en función del tipo de línea.

La norma define **una red de acoplamiento/desacoplamiento para líneas de alimentación AC y DC** (Figura 11.2). Vemos que la salida del generador se acopla a las líneas de alimentación a través de un condensador en serie de 33 nF. Las ferritas, inductancia y condensadores forman la sección de desacoplamiento. Su función es evitar que la perturbación afecte al equipo auxiliar en el ensayo. Por ejemplo, si estamos ensayando un televisor, tendremos conectada una fuente de señal en sus entradas HDMI. Esta fuente de señal es un equipo auxiliar, no debe ser sometida a la perturbación y por tanto se conecta a la red AC protegida por la sección de desacoplamiento.

El acoplamiento en líneas de señal es distinto (Figura 11.3) y consiste en un triángulo metálico de 1 m de longitud y lados de 7 cm. La capacidad equivalente entre esta estructura (denominada pinza de acoplamiento capacitiva) y el cable de señal es de aproximadamente 100 pF.

Tabla 11.1. Niveles de tensión de pico aplicables en función del tipo de línea

E/S alimentación y tierra		E/S señal	
Nivel	Tensión de cresta (kV)	Nivel	Tensión de cresta (kV)
1	0,5	1	0,25
2	1	2	0,5
3	2	3	1
4	4	4	2

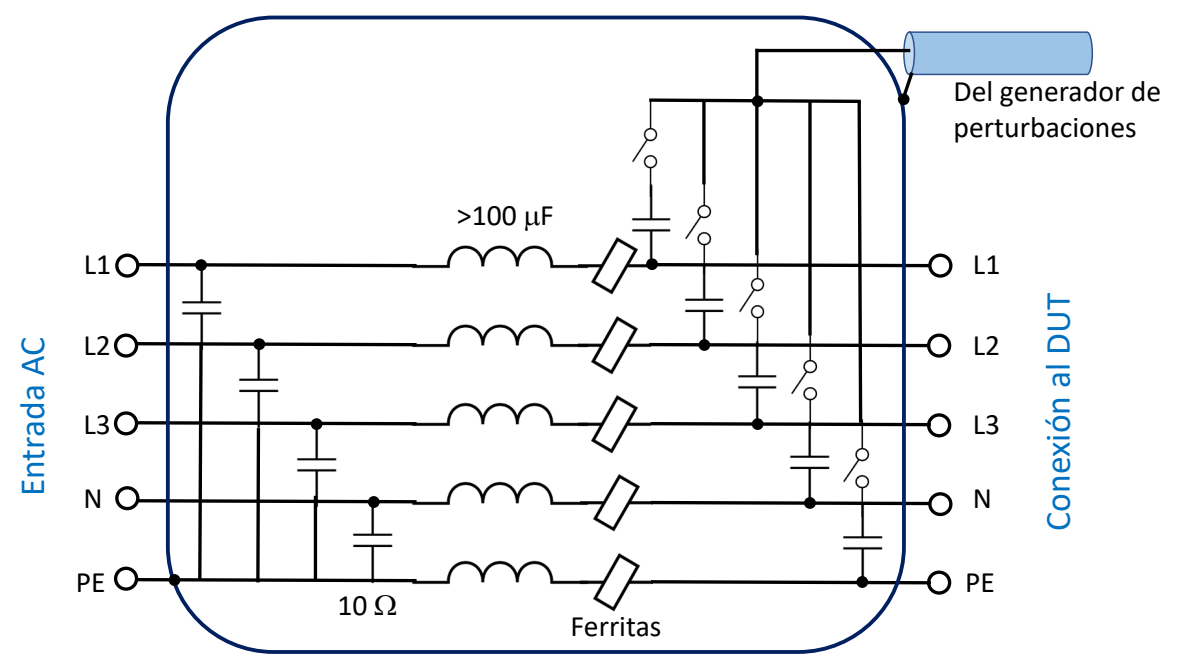


Figura 11.2. Red de acoplamiento/desacoplamiento para líneas de alimentación AC y DC



Figura 11.3. Acoplamiento de perturbaciones EFT a líneas de E/S mediante una estructura normalizada. La imagen muestra un modelo comercial ofrecido en <https://www.theemcshop.com/>

Debes tener presente que la perturbación se aplica simultáneamente y en modo común respecto a tierra a todas las líneas de alimentación y de señal (donde la norma de familia de producto así lo requiera). Esto tiene sentido, ya que una EFT emula un transitorio que induce corriente en todas las líneas.

Normas de familia de productos

Las normas genéricas o de familia de producto establecen el nivel de la perturbación en el ensayo de inmunidad, así como otras condiciones particulares. Por ejemplo, [EN 55024](#), aplicable a equipos de tecnología de la información:

- Sólo se aplica el ensayo a cables de señal si su longitud es mayor a 3 m (se entiende de longitudes menores no darán lugar a un acoplamiento importante entre las líneas de AC y las de señal).
- El criterio de aptitud establecido por esta norma es el B, al igual que para los ensayos de inmunidad a ESD.
- La frecuencia de repetición de los pulsos es de 5 kHz, excepto para aparatos xDSL que será de 100 kHz. El nivel de ensayo en líneas de señal, alimentación DC (tengan o no convertidor AC/DC separado) y alimentación AC es 0,5 kV.

Simulando pulsos EFT en SPICE

En [27] podemos encontrar un circuito para simular un generador de pulsos EFT, que hemos reproducido en la Figura 11.4 junto a una simulación de la forma de onda con una carga de 50 ohmios (IEC 61000-4-4 define dos cargas -50 ohmios y 1 k Ω - para las que el generador debe cumplir con las especificaciones de tiempo de subida, anchura de pulso y amplitud de pico).

Este no es el único circuito posible. Por ejemplo, en [28] se propone un circuito distinto para el generador.

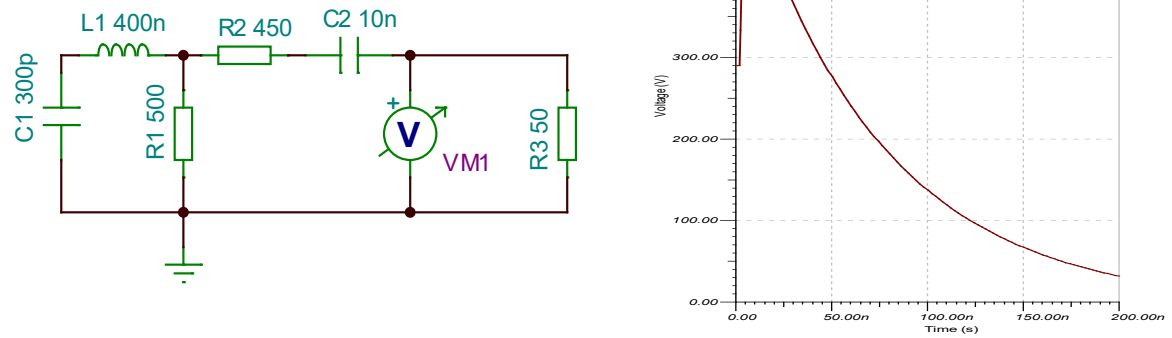


Figura 11.4. Circuito para simular pulsos ESD [27]

La Figura 11.5 muestra un ejemplo de simulación sobre una línea de señal analógica que presenta una carga de 50 ohmios. Observa que, tras el circuito generador, hemos añadido en serie una capacidad de 100 pF que representa a la pinza capacitiva. De haber simulado una línea de alimentación hubiéramos usado una capacidad de 33 nF.

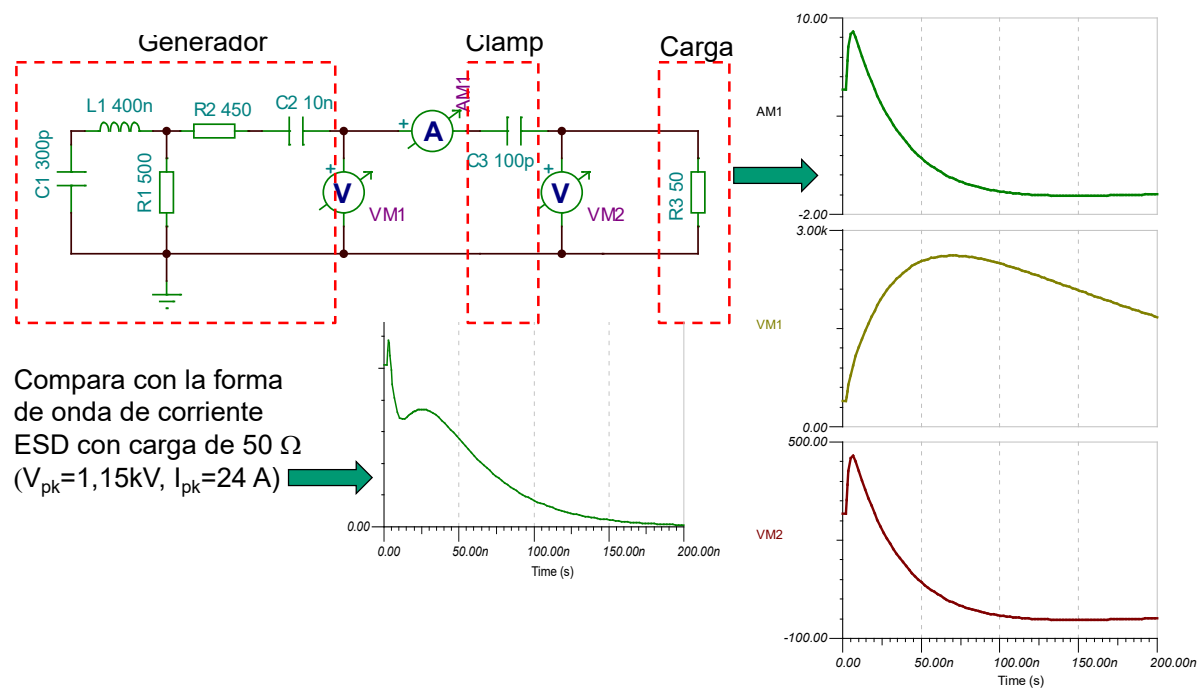


Figura 11.5. Simulación de un pulso EFT en una línea de señal analógica que presenta una carga de 50 ohmios

Ejercicio de descarga EFT

Determina, mediante simulación, si es necesaria una ferrita para proteger al circuito integrado frente a pulsos EFT de 500 V de pico o si por el contrario basta con una pequeña resistencia.

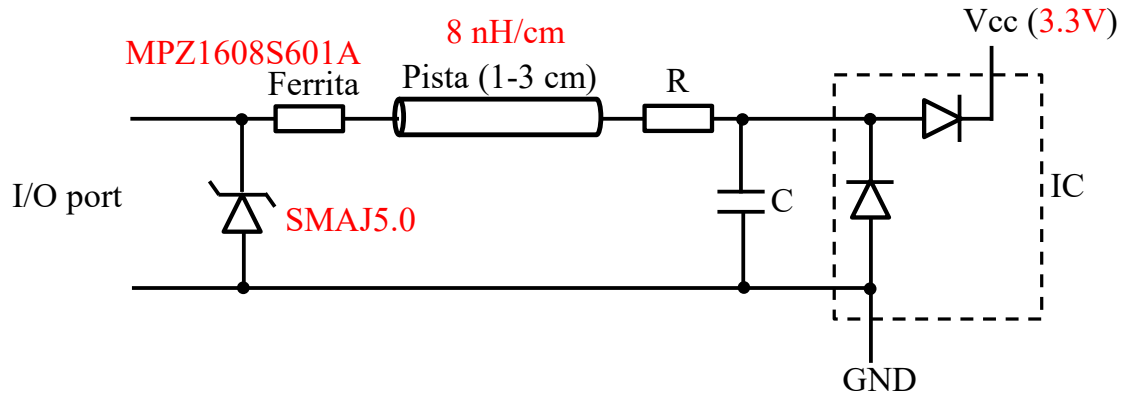


Figura 11.6. Ejercicio de simulación SPICE para pulsos EFT

El esquema en LTspice está recogido en la Figura 11.7. Sólo con la ferrita, la corriente de pico es de unos 140 mA, excesivo según los criterios que fijamos ayer. Por tanto, hay que aumentar la impedancia en serie añadiendo una resistencia. En su lugar, vamos a reemplazar la ferrita por una resistencia. No nos debe extrañar que con una resistencia de 56 ohmios en lugar de la ferrita la corriente de pico baje a unos aceptables 60 mA, un resultado similar al que obtuvimos ayer en el ejercicio de simulación ESD.

Para conseguir un pulso de 500 V sobre una resistencia de 50 ohmios, como marca la norma, debemos cargar inicialmente C1 a una tensión de 5.2 kV.

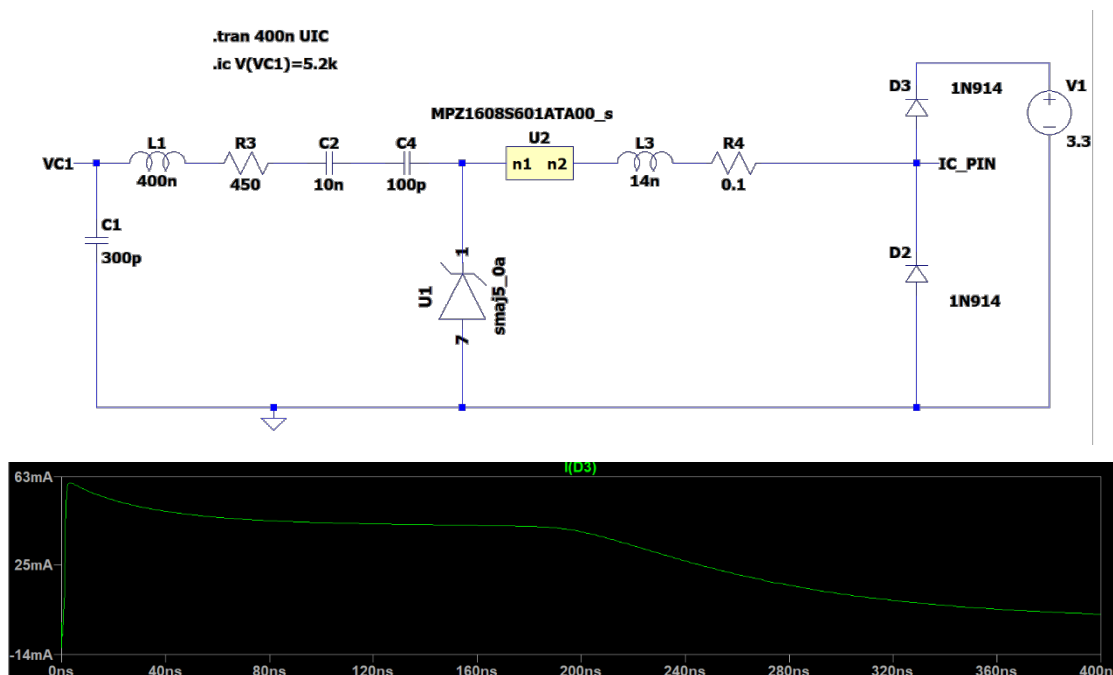


Figura 11.7. Diagrama del circuito en LTspice (arriba) y simulación con una resistencia de 56 ohmios en lugar de la ferrita (abajo)

Normativa para onda de choque (Surge)

Norma básica

La norma básica **IEC 61000-4-5** define los ensayos de inmunidad a las ondas de choque. Esta perturbación de alta tensión emula sobretensiones en líneas de alimentación AC provocadas por descargas atmosféricas (rayos) cercanas a los cables, maniobras en sistemas de gran potencia, fallos en la red eléctrica, etc. También pueden afectar a las alimentaciones DC y señales externas (como líneas R485, ADSL y telefonía). La energía del pulso de prueba puede ser miles de veces mayor que la de ESD y EFT y por tanto es necesario utilizar dispositivos de protección diferentes.

La forma de onda de tensión en circuito abierto se conoce como pulso 1,2/50 μs (haciendo referencia al tiempo de subida y a la anchura de pulso -reducción al 50% de amplitud-) y se muestra en la Figura 11.8, izquierda.

La forma de onda de corriente en cortocircuito se conoce como pulso 8/20 μs (haciendo referencia al tiempo de subida y a la anchura de pulso -reducción al 50% de amplitud-) y se muestra en la Figura 11.8, derecha. Su amplitud de pico en tensión, con una carga de 2 ohmios, debe ser la mitad de la que alcanza la forma de onda 1,2/50 μs . Ambas formas de onda hacen referencia al mismo ensayo.

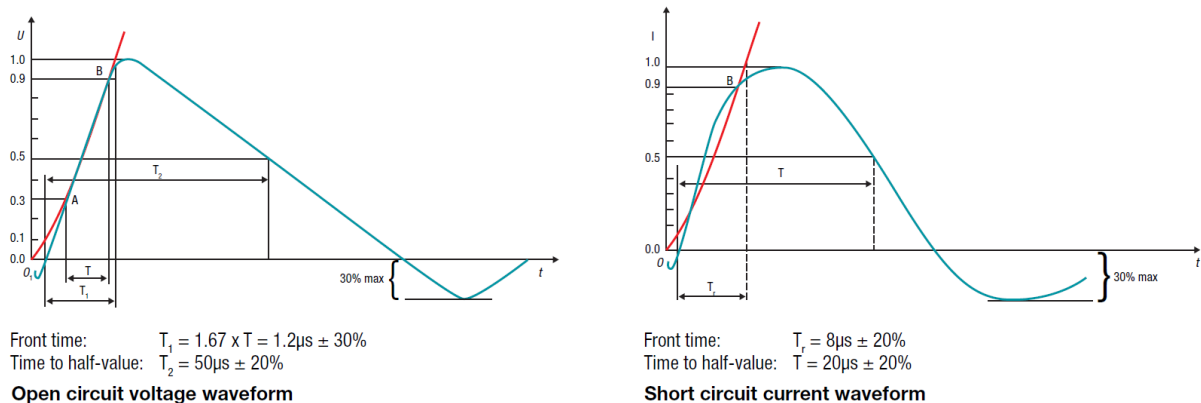


Figura 11.8. Formas de onda de la descarga (onda de choque): tensión en circuito abierto (izquierda) y corriente en cortocircuito(derecha), normalizadas. Fuente [29]

El pulso producido por el generador se acopla a la línea a ensayar por medio de una de las redes de acoplamiento definidas en la norma:

- Para líneas de alimentación AC y DC, cuando el pulso se aplica en modo normal entre dos líneas, el terminal de referencia se conecta directamente a una línea, el otro terminal a la segunda línea mediante un condensador de 18 μF .
- Para líneas de alimentación AC y DC, cuando el pulso se aplica entre una línea y tierra, el terminal de referencia se conecta a tierra, el otro terminal a la línea mediante un condensador de 9 μF y una resistencia de 10 ohm.
- En líneas diferenciales no apantalladas, la perturbación se aplica en modo común con una resistencia de 40 ohm a cada línea del par.

La norma define otras redes de acoplamiento para otras situaciones y debes conocerlas si quieres simular correctamente el ensayo, pues la red de acoplamiento afecta al pulso realmente aplicado.

Normas de familia de productos

EN 55024 define los niveles que han de superar los **equipos de tecnología de la información**. La tensión de pico en circuito abierto es de 500 V para puertos de alimentación DC, 1 kV (ambas polaridades) para líneas de señal y L-N en red AC y de 2 kV para L-E y N-E en líneas de red AC. El criterio de aptitud es B, excepto en puertos de señal, que es C.

Simulación de ondas de choque

EN 61000-4-5 ofrece un modelo del circuito equivalente del generador (Figura 11.9), pero no da el valor de los componentes.

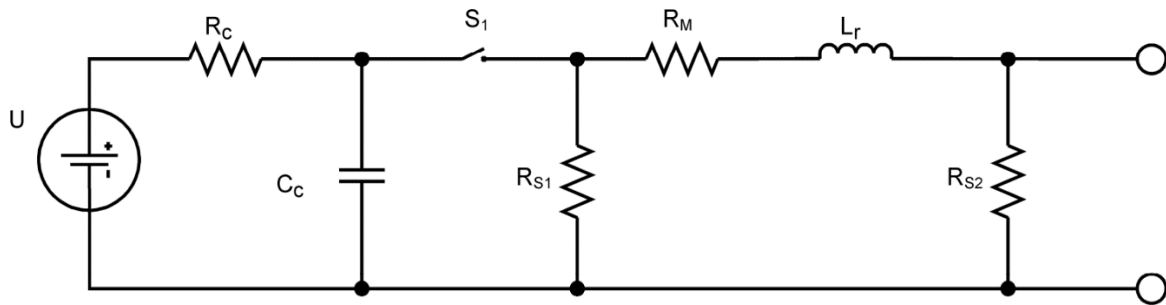


Figura 11.9. Modelo equivalente del generador de ondas de choque. Fuente: <https://verimod.com> ([enlace](#))

Afortunadamente, hay quien se ha tomado la molestia en estimar los valores de los componentes [30]:

$$\begin{array}{lll}
 C_C = 6,038 \mu\text{F} & L_r = 10,37 \text{ mH} & R_{S1} = 25,105 \Omega \\
 R_{S2} = 19,80 \Omega & R_M = 0,941 \Omega & V_{IC} = 1082 \text{ V (para una tensión de pico de 1 kV)}
 \end{array}$$

Con estos valores, en la Figura 11.10 comprobamos que las formas de onda en circuito abierto y en cortocircuito se corresponden con la especificación de la norma.

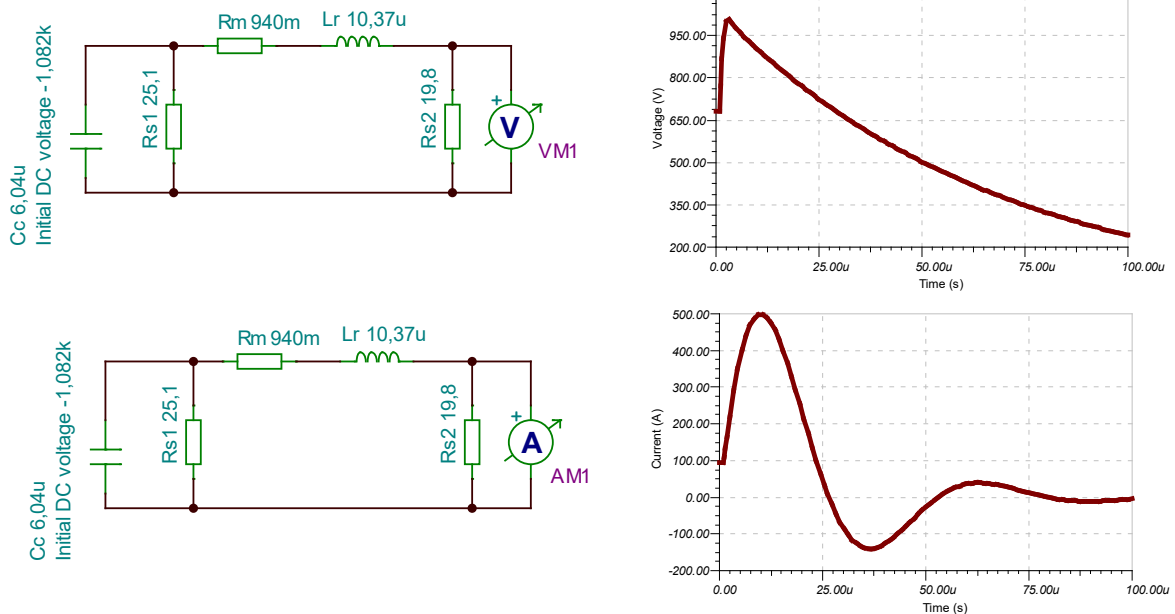


Figura 11.10. Simulación SPICE del generador de ondas de choque

Dispositivos de protección para ondas de choque

La energía de las ondas de choque es miles de veces superior a la de ESD o EFT. Por tanto, son necesarios dispositivos de protección distintos:

- Varistores de óxidos metálicos (MOVs) de alta potencia
- Varistores de óxidos metálicos (MOVs) multicapa de alta velocidad
- Diodos TVS de alta potencia
- Tubos de descarga de gas

Pensando en instalaciones en el interior de edificios, la protección primaria frente a ondas de choque (que, recuerda, emulan descargas de rayos en las inmediaciones de la instalación) la aportan las compañías eléctricas y de telecomunicación a la entrada del edificio. Suele tratarse de protecciones que actúan sobre el modo común: transformadores (red AC) o tubos de descarga de gas de cada línea a tierra (en el caso de líneas de comunicaciones). Si se trata de una instalación propia (por ejemplo, una red externa RS-485 en un campo de cultivo) hemos de proporcionarla nosotros.

Si queremos proteger frente a ESD, EFT y ondas de choque necesitamos un **esquema de protección de dos niveles**. Mira la Figura 11.11. Funciona así: al inicio del transitorio, cuando la tensión es de unos pocos voltios, comienza a conducir el diodo TVS. El varistor y el tubo de descarga de gas comienzan a conducir corriente a tensiones bastante más elevadas. La corriente a través del diodo TVS hace que caiga tensión en las resistencias (o resistencia, si es mejor no tener elevación en la línea de masa). Esta resistencia (o ferrita, o resistencia variable con la tensión, que también podría servir) se calcula (y se comprueba mediante simulación) para que el tubo de gas o el varistor comiencen a conducir antes de que se alcance la corriente máxima en el diodo TVS.

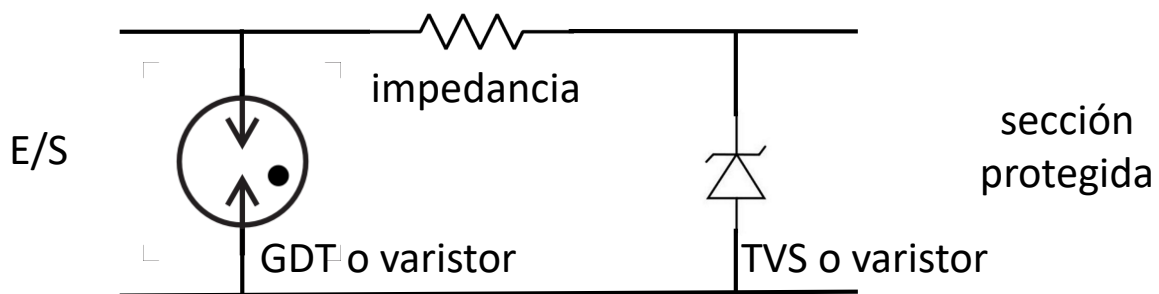


Figura 11.11. Esquema de protección a dos niveles

Diodos TVS de alta potencia

Lo primero que nos llama la atención es el encapsulado: axial en lugar de SMD. Esto es así para aumentar su volumen y por tanto para poder soportar mayor corriente. También ayuda este mayor volumen a evacuar el calor de manera más eficaz al ambiente.

Tomemos como ejemplo el diodo TVS de potencia 1.5KE33A. Los primeros caracteres (1.5K) nos indican que la potencia de pico que puede disipar es de 1,5 kW. Esto coincide con el producto de la tensión máxima de *clamping* (recorte) y corriente a esta tensión ($45,7\text{V} \cdot 33,3\text{ A} = 1,52\text{ kW}$). Los dígitos “33” indican que a esa tensión en voltios la corriente en inversa es de 1 mA.

La capacidad eléctrica de este diodo es de 1 nF: olvídate de usarlo con líneas analógicas de alta frecuencia o digitales de elevada frecuencia de reloj. Por ejemplo, a 10 MHz, su impedancia (para pequeña señal) sería de 16 ohmios aproximadamente. Pero ya su tensión de trabajo, cerca de los 30V, te indica que este dispositivo será utilizado en líneas de alimentación.

1.5KE series de Littelfuse

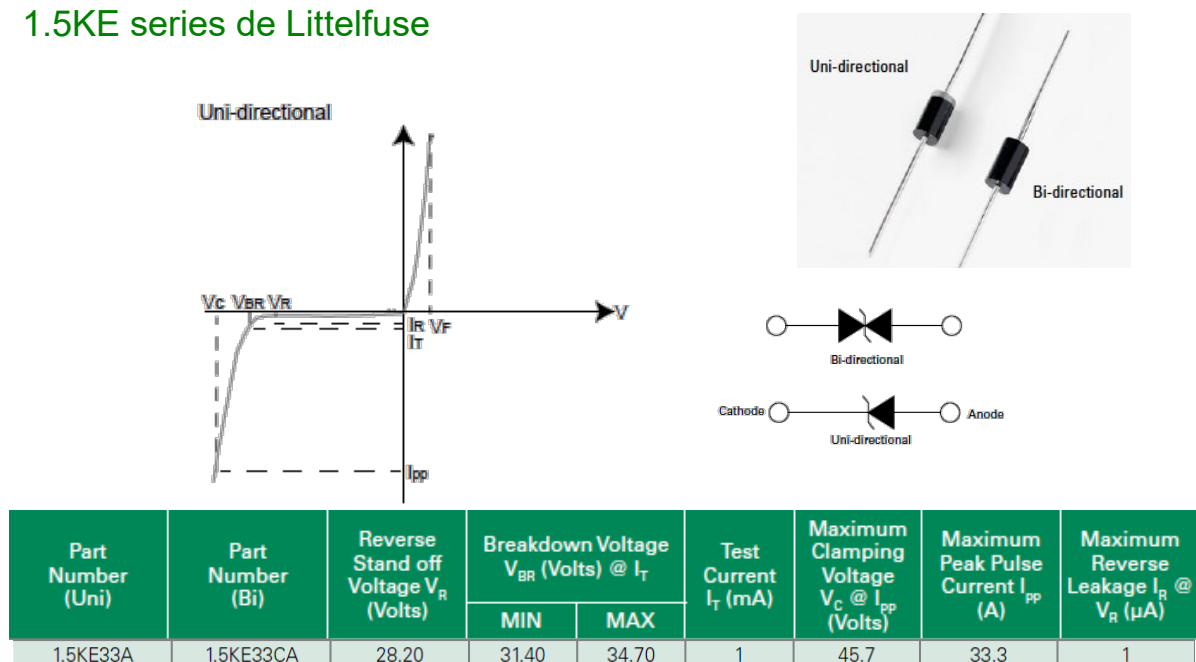


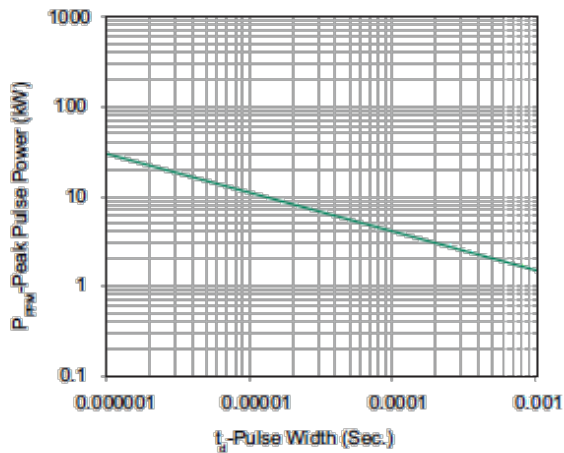
Figura 11.12. Características del diodo TVS 1.5KE33A

La máxima potencia de pico que soporta este diodo (1,5 kW) está especificada para un pulso de 1 ms de anchura (Figura 11.13). Si el pulso es más estrecho es de esperar que pueda soportar mayores potencias de pico. Por ejemplo, soportará hasta 4 kW de pico para pulsos de 100 μs de anchura. Esto juega normalmente a nuestro favor. Una simulación SPICE puede darte una estimación de la anchura de pulso, de la potencia de pico disipada y por tanto del margen de seguridad que hay en el diseño.

Un segundo ajuste que hay que realizar es el *derating con la temperatura*: los 1,5 kW nominales pueden ser disipados sólo hasta 25°C. A temperaturas superiores hay que ajustar este valor conforme a la Figura 11.14.

Combinando el efecto de la anchura de pulso y de la máxima temperatura de funcionamiento podemos estimar la máxima potencia de pico admisible y compararla con el resultado de las simulaciones.

1.5KE33A de Littelfuse



Maximum Ratings and Thermal Characteristics ($T_A=25^{\circ}\text{C}$ unless otherwise noted)			
Parameter	Symbol	Value	Unit
Peak Pulse Power Dissipation by 10x1000 μs Test Waveform (Fig.2) (Note 1)	P_{PPM}	1500	W
Steady State Power Dissipation on Infinite Heat Sink at $T_L=75^{\circ}\text{C}$ (Fig. 6)	P_D	6.5	W
Peak Forward Surge Current, 8.3ms Single Half Sine Wave Unidirectional Only (Note 2)	I_{FSM}	200	A
Maximum Instantaneous Forward Voltage at 100A for Unidirectional Only (Note 3)	V_F	3.5/5.0	V
Operating Junction and Storage Temperature Range	T_J, T_{STG}	-55 to 175	$^{\circ}\text{C}$
Typical Thermal Resistance Junction to Lead	$R_{\theta JL}$	15	$^{\circ}\text{C/W}$
Typical Thermal Resistance Junction to Ambient	$R_{\theta JA}$	75	$^{\circ}\text{C/W}$

Figura 11.13. Potencia de pico en función de la anchura de pulso (diodo TVS 1.5KE33A)

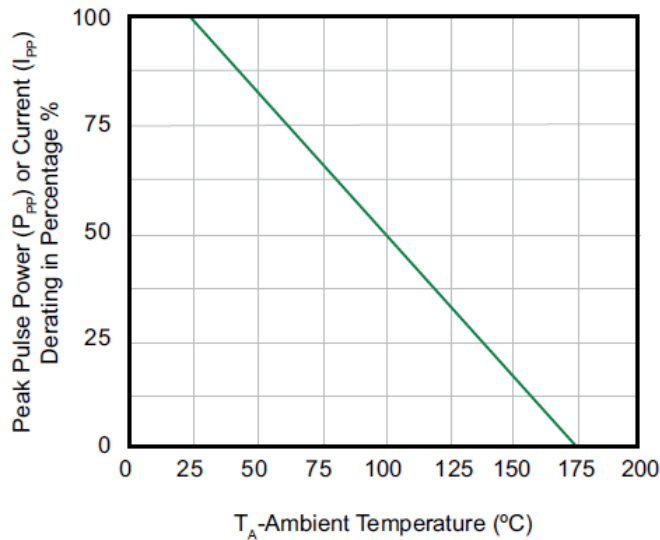


Figura 11.14. Derating con la temperatura (diodo TVS 1.5KE33A)

Ejemplo

Estudia un circuito de protección de una línea DC de 18V basado en un diodo TVS 1.5KE20A y un filtro en “T” con dos inductancias de 1 mH (Figura 11.15). La máxima temperatura ambiente para el producto será de 50°C. Modelamos la carga mediante el condensador de entrada del regulador (100 μ F) y una resistencia de valor igual a V/I , siendo $V=18$ V e $I=225$ mA el consumo del circuito. Como se trata de una línea de alimentación, la red de acoplamiento según la norma es un condensador en serie de 18 μ F.

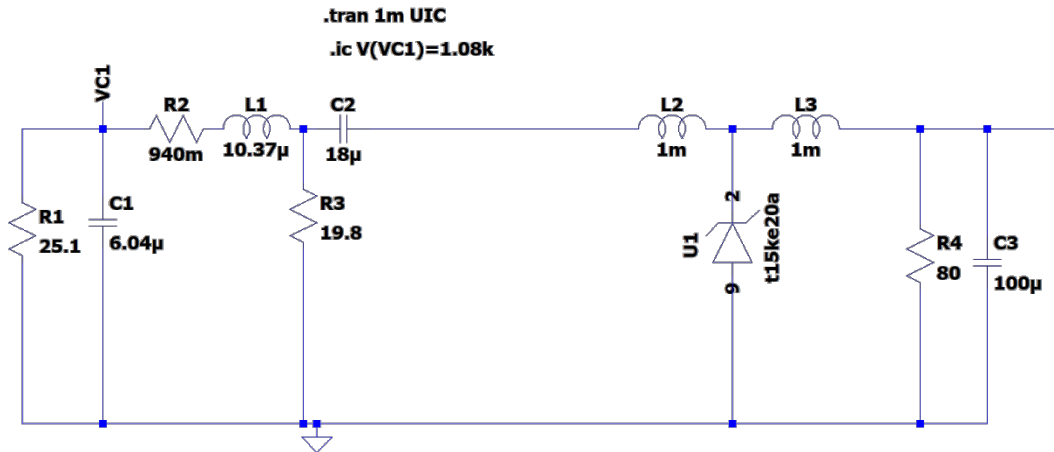


Figura 11.15. Diagrama del circuito en el simulador

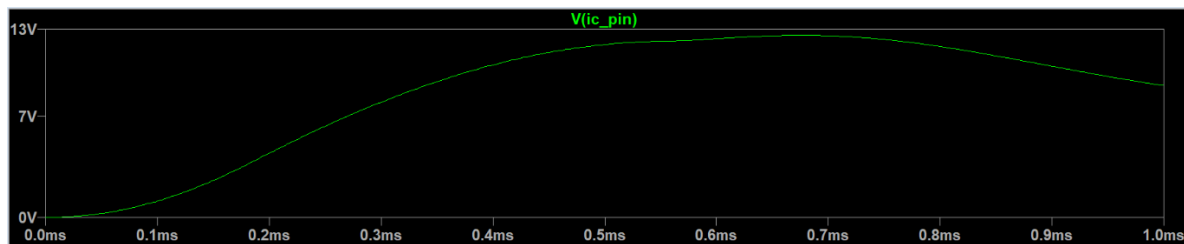


Figura 11.16. Tensión en la carga. El valor de pico es de 12,6 V

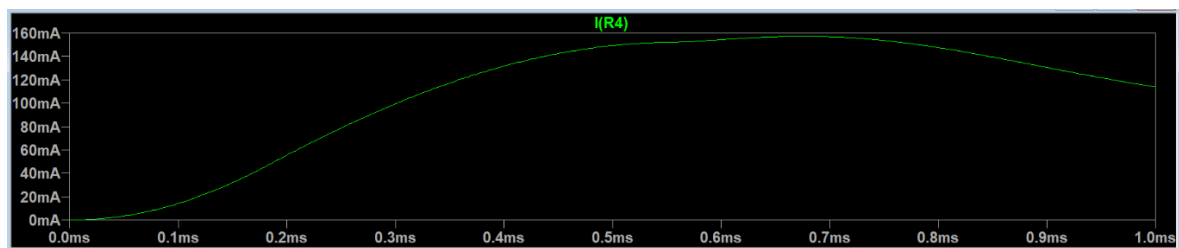


Figura 11.17. Corriente en la carga de 80 ohmios. El valor de pico es de 160 mA

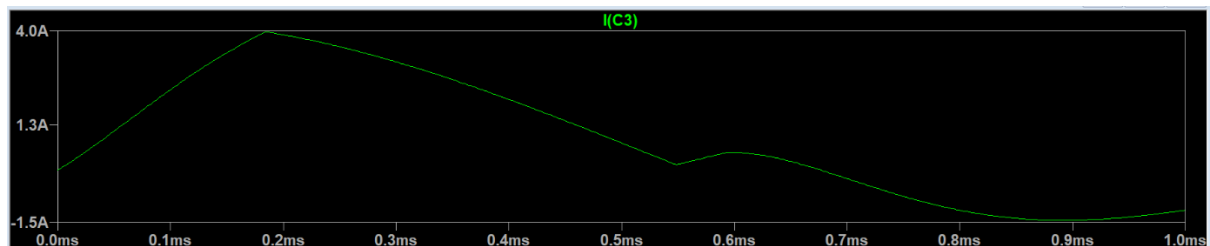


Figura 11.18. Corriente a través del condensador de 100 μ F en la carga, con un pico de 4 A

Las figuras anteriores muestran el transitorio en la carga: una sobretensión de 12,6V y una corriente AC en el condensador de entrada del regulador de 4 A. Hay que escoger un condensador de entrada capaz de hacer frente a estos valores. Del mismo modo, hay que escoger un regulador capaz de soportar la sobretensión (12V nominales más 12,6V).

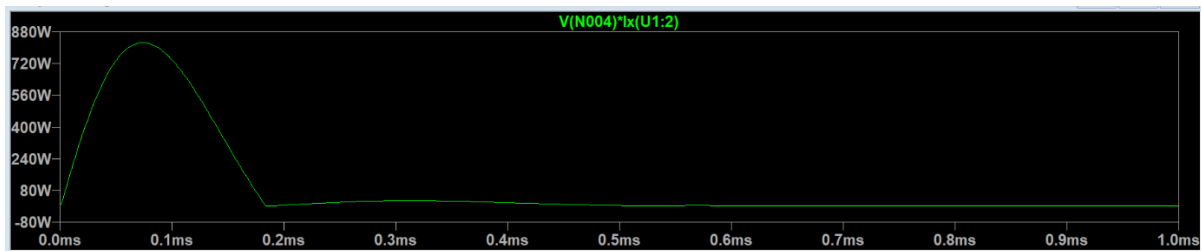


Figura 11.19. Potencia instantánea en el diodo TVS, con un valor de pico de 825W y una anchura de pulso inferior a 200 μ s

Hay que comprobar también la potencia disipada en el diodo TVS (Figura 11.19). En LTspice, debes mantener la tecla Alt pulsada para seleccionar la opción de potencia instantánea. Echando un vistazo a la hoja de datos del 1.5KE20A (Figura 11.20) observamos que a 20°C la potencia de pico para un pulso de 200 μ s es de 3 kW. A 50°C debemos considerar que se alcanza sólo el 90% de este valor, 2,7 kW. Con 825 W de pico, el factor de seguridad es de 3. Más que suficiente.

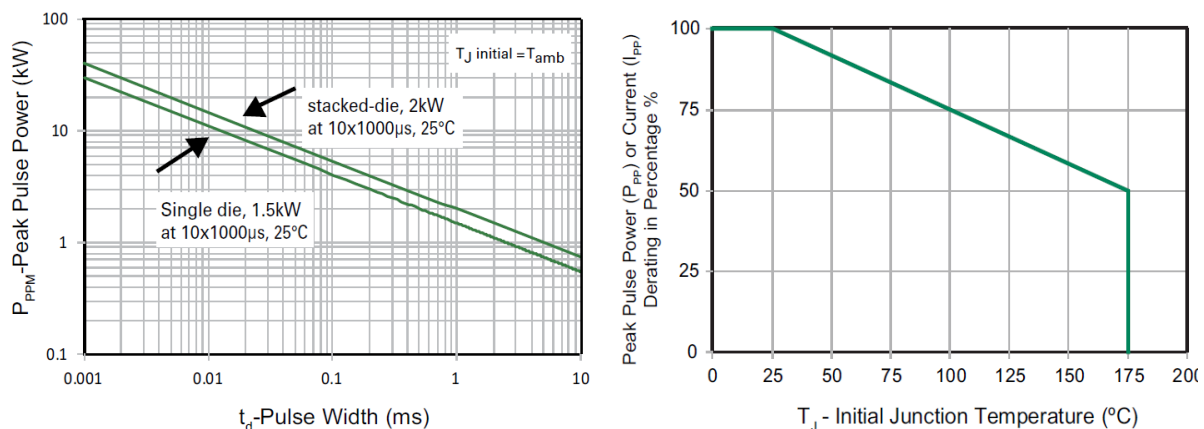
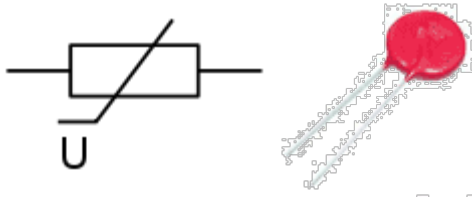


Figura 11.20. Características del TVS 1.5KE20A

Dejo como ejercicio que evalúes el comportamiento del circuito ante un pulso de polaridad negativa.

Varistores

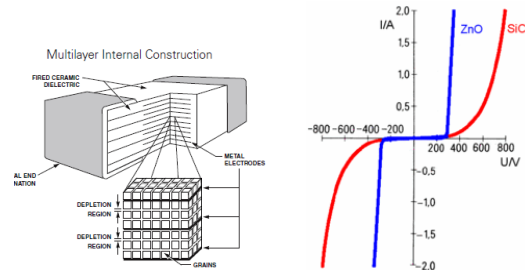
El término varistor es una contracción de “variable resistor”. El tipo más usual (denominado MOV, *metal-oxide varistor*) está formado por granos de óxido de zinc y otros metales que forman multitud de diodos dispuestos en todas las direcciones. Así es, ¡es bidireccional y simétrico! Por debajo de una diferencia de tensión de ruptura, presenta alta impedancia. Una vez superada esta tensión crítica, tenemos una baja impedancia.



Se degradan tras cada transitorio absorbido, en función de la corriente conducida. Así, corrientes pequeñas no degradan apreciablemente el dispositivo incluso después de millones de transitorios. Corrientes medias dan lugar a una degradación de después de varios miles de transitorios. Corrientes elevadas limitan la vida útil del varistor a unas pocas decenas de pulsos. Conocer los límites es complicado, pues los fabricantes sólo proporcionan una información limitada.

V14MLA1206N de Littelfuse

– 1.4 nF, respuesta < 1ns, SMD 1206



Part Number	Maximum Ratings (125°C)				Specifications (25°C)			
	Maximum Continuous Working Voltage		Maximum Non-repetitive Surge Current (8/20µs)	Maximum Non-repetitive Surge Energy (10/1000µs)	Maximum Clamping Voltage at 1A (or as Noted) (8/20µs)	Nominal Voltage at 1mA DC Test Current		Typical Capacitance at f = 1MHz
	V_{MDC} (V)	V_{MAO} (V)	I_{TM} (A)	W_{TM} (J)	V_C (V)	V_{NDC} Min (V)	V_{NDC} Max (V)	C (pF)
V14MLA1206N	14.0	10.0	150	0.400	32.0	15.9	20.3	1400

¿Cómo falla un varistor?

Generalmente, cuando fallan debido a un impulso moderado, lo hacen en **cortocircuito**. Esto permite detectar el fallo y reemplazar el varistor. Si el impulso absorbido es de mucha energía, pueden estallar, quedando en **circuito abierto** (y dejan de proteger al equipo). En este segundo caso, no es fácil detectar el fallo sin una inspección visual.

Para proteger a los varistores es frecuente verlos asociados a fusibles o PTCs, que evitan sobre corrientes y prolongan su vida útil.

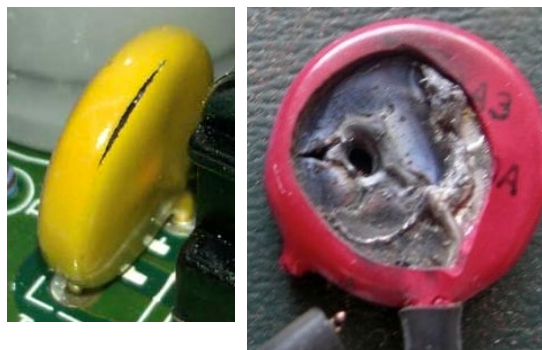


Figura 11.21. Varistores de lenteja que han estallado, lo que es fácil de apreciar a ojo desnudo

¿Puede sobrevivir un varistor a un ensayo EFT?

Os traslado una pregunta que me hace hoy mismo (mientras escribo esta sección) un ingeniero de una empresa con la que he colaborado en ocasiones. La pregunta es buena, y podemos aprender algo de ella: los varistores se degradan con el uso, y es algo que debes tener muy en cuenta en el diseño.

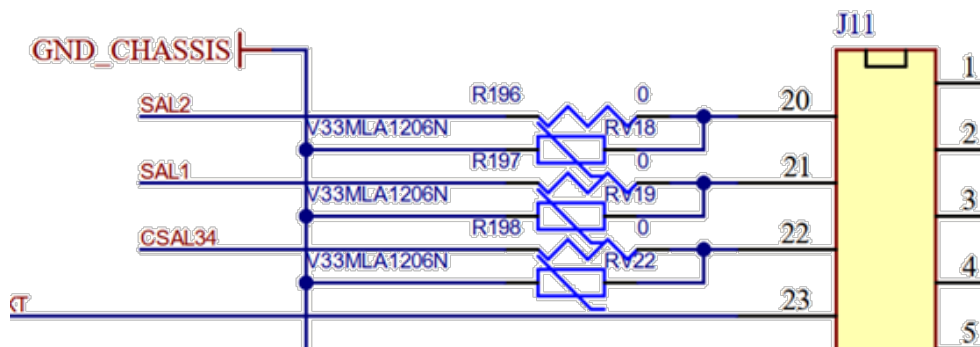
“Hola compañero,

En el diseño (...) la mayoría de los puertos/conectores están protegidos con varistores. Para la protección del nuevo equipo que estamos desarrollando, se nos plantea una duda con las pruebas EFT, y es si la repetitividad de los pulsos incluidos en la misma puede afectar al varistor por envejecimiento.

¿Sabes de alguna nota o documento que haga referencia al asunto, o experiencia personal?

Gracias.”

El texto viene acompañado de este esquema:



Esta ha sido mi respuesta:

“Hola!

Un varistor como el que montáis está diseñado para soportar un único transitorio a su energía máxima (0.8 J en el modelo del esquema, conduciendo un corriente de 180 A para un pulso ancho), duración y energía que es bastante mayor de la de un pulso EFT (unos 4 mJ y tal vez 200 ns de transitorio). El datasheet dice que ante un transitorio que provoque una corriente de 3A en el dispositivo con una anchura de pulso de 20 μ s, el dispositivo no muestra degradación tras 10.000 pulsos. Pararon la prueba aquí, no dicen cuanto más aguanta.

*Un ensayo EFT de 1 minuto supone 15.000 pulsos (75 pulsos/burst, 1 burst cada 300 ms, 60 segundos) de 4 mJ cada uno. **Puedes simular la corriente de pico y el ancho de pulso en el varistor ante un pulso EFT** y a partir de aquí sacar algunas conclusiones.*

*La potencia media durante un 1 minuto de prueba EFT es de 1 W, lo que no es poca cosa. Si miras la letra pequeña en el datasheet, casi no se ve, dice que "Average power dissipation of transients for 0402, 0603, 0805, **1206** and 1210 sizes not to exceed 0.03W, 0.05W, 0.1W, **0.1W** and 0.15W respectively". **Según esto, una prueba EFT debería degradar este varistor de encapsulado 1206**, tal y como temes.*

Una nota de aplicación que trata sobre estos temas es:

https://m.littelfuse.com/~media/electronics_technical/application_notes/varistors/littelfuse_selecting_a_littelfuse_varistor_application_note.pdf

Así que, sin duda, una prueba EFT va a envejecer al varistor, pero vas a testear de 1 a 3 unidades y el resto de las que produzcas van a ir a la instalación del cliente sin sufrir el test, y por tanto sin degradación.

Ante transitorios moderados, los varistores suelen morir en cortocircuito, pero antes han ido degradándose, bajando poco a poco su tensión de recorte (clamp). Así que una baja impedancia en el varistor o una tensión entre terminales menor de la esperada es un indicador de que hay que reemplazarlo.

No acabo de entender si tu pregunta es que un varistor puede degradarse durante la prueba EFT y hacer que el equipo no pase el test. O que tras el ensayo EFT el varistor haya quedado tocado y eso afecte a otras pruebas.”

Modelo SPICE de un varistor

Tomemos como ejemplo el V14MLA1206N, un MOV para automoción. Se trata de un varistor multicapa SMD, por tanto, rápido y válido para ESD, EFT y otros transitorios. Se usa en lugar de un diodo TVS cuando, a igualdad de volumen, se necesita absorber mayor potencia. Como desventaja, limita a tensiones muy por encima de la nominal: no es tan bueno como un TVS a la hora de recortar (*clamping*).

Continuous	ML Series	Units
Steady State Applied Voltage:		
DC Voltage Range (V_{MIDC})	3.5 to 120	V
AC Voltage Range ($V_{WACIRMS}$)	2.5 to 107	V
Transient:		
Non-Repetitive Surge Current, 8/20 μ s Waveform, (I_{ns})	4 to 500	A
Non-Repetitive Surge Energ. 10/1000 μ s Waveform. (W_{ns})	0.02 to 2.5	J



Part Number	Maximum Ratings (125° C)					Specifications (25°C)		
	Maximum Continuous Working Voltage		Maximum Non-repetitive Surge Current (8/20µs)	Maximum Non-repetitive Surge Energy (10/1000µs)	Maximum Clamping Voltage at 1A (or as Noted) (8/20µs)	Nominal Voltage at 1mA DC Test Current		Typical Capacitance at f = 1MHz
	V _{M(DC)}	V _{M(AC)}	I _{TM}	W _{TM}	V _C	V _{N(DC) Min}	V _{N(DC) Max}	C
	(V)	(V)	(A)	(J)	(V)	(V)	(V)	(pF)
V14MLA1206N	14.0	10.0	150	0.400	32.0	15.9	20.3	1400

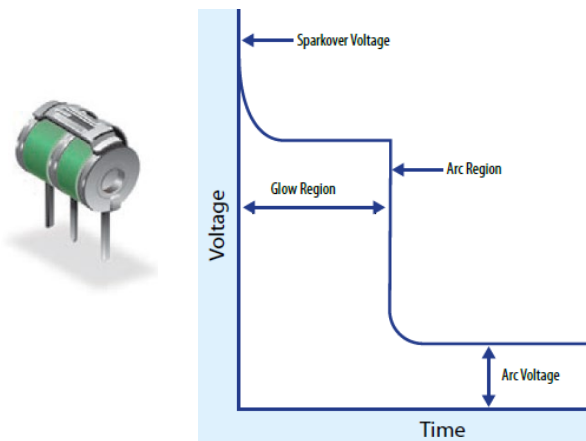
[illegible]

Figura 11.22. Ejemplo de modelo SPICE de un varistor

Los modelos SPICE para varistores se hacen en base a un ajuste a un polinomio. Con frecuencia dan problemas de sintaxis en LTspice. Si no quieres ir retocando el modelo para corregir los errores, puedes recurrir a otro simulador, como TINA de Texas Instruments.

Tubos de descarga de gas (GDTs)

Su uso principal es la protección de líneas externas (RS-485, ADSL, Ethernet) a la entrada de edificios o instalaciones. Por su baja capacidad (aprox. 1 pF) no cargan a líneas de datos rápidas. Se trata de dispositivos tipo *crowbar*, lo que permite manejar transitorios elevados con una mínima disipación de potencia en el dispositivo. Las configuraciones de tres terminales (el central va conectado a tierra o chasis) son ideales para protecciones línea a línea y línea a tierra.



if a Gas Discharge Tube (GDT) Primary Protector “, rev2, www.bourns.com, 2008

¿Cómo funciona?

1. Por debajo de la tensión umbral hay alta Z entre terminales
2. Cuando la tensión aumenta debido al transitorio, comienza la ionización del gas (*glow region*), donde comienza a circular corriente (mA) por el GDT
3. La ionización da lugar a una avalancha y crea un camino de baja impedancia
4. Cuando el transitorio termina, el GDT vuelve a su estado de alta impedancia

Bourns 2038 series (SMD 7.5x5 mm, 3 terminales, <1pF)



Characteristic	2038-15-SM
DC Sparkover $\pm 25\%$ @ 100 V/s L1/L2 to Gnd (NOTE 1)	150 V
Typical Impulse Sparkover L1/L2 to Gnd 100V/ μ s 1000V/ μ s	350 V 500 V

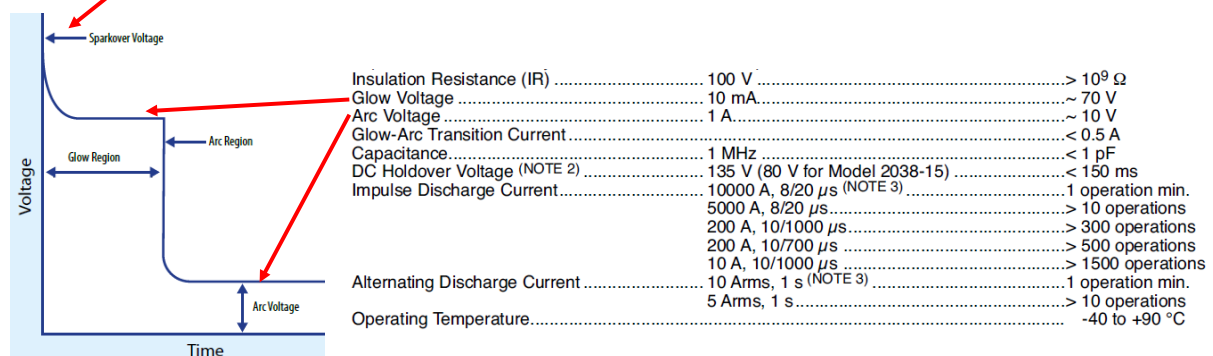


Figura 11.23. Ejemplo de GDT: Bourns 2038-15-SM

Consideraciones prácticas

Mantener al GDT en la región de “glow” durante mucho tiempo acorta la vida del dispositivo por una elevada disipación térmica. El GDT es un dispositivo lento y por tanto es necesario una protección secundaria para absorber la energía inicial del transitorio: esto es algo que ya habíamos comentado al comprar dispositivos tipo *crowbar* y tipo *clamping*.

Los electrodos del GDT se degradan tras cada disparo. El modo de fallo es un cortocircuito debido a la deposición de metal en las paredes de la cápsula, creando un camino de baja impedancia. De este modo, no sólo los varistores sino también los tubos de descarga de gas se degradan. La Figura 11.24 muestra la vida útil (en descargas) de un dispositivo.

10000 A, 8/20 μ s (NOTE 3)	1 operation min.
5000 A, 8/20 μ s	> 10 operations
200 A, 10/1000 μ s	> 300 operations
200 A, 10/700 μ s	> 500 operations
10 A, 10/1000 μ s	> 1500 operations

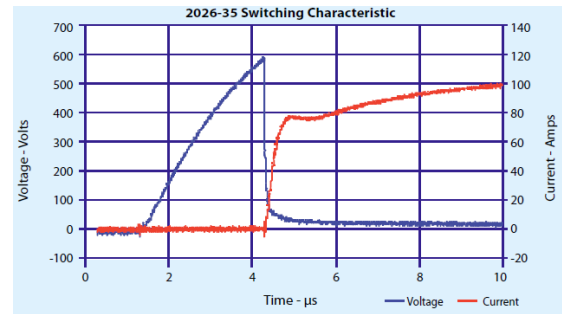
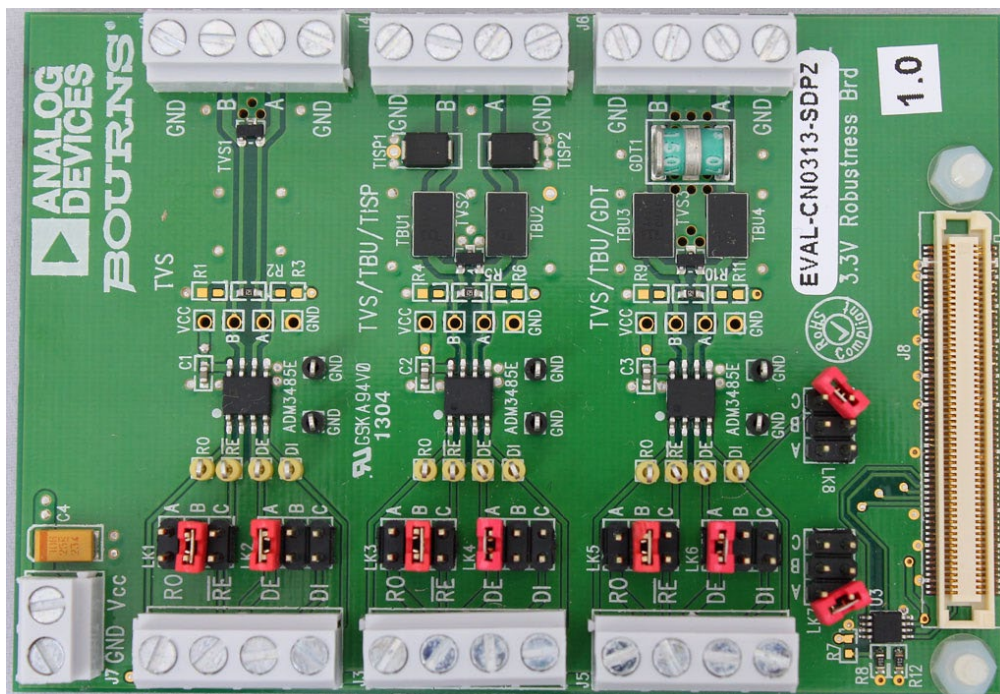


Figura 11.24. Ejemplo de vida útil (en descargas) del Bourns 2038-15-SM en función de la corriente de pico del transitorio

La Figura 11.25, que procede de una nota de aplicación de Analog Devices, recoge tres circuitos de protección. El de la derecha incorpora un tubo de descarga de gas junto a la regleta de entrada. Como ves, su volumen es aceptable y se puede incorporar a un PCB sin problemas.



Fuente: “EMC Compliant RS-485 Transceiver Protection Circuits”, Analog Devices, Circuit Note CN-0313

Figura 11.25. Un transceiver RS-485 protegido mediante tres circuitos distintos, en una placa de evaluación de Analog Devices [31]

Día 12. Estudio de un caso de aplicación de normativa EMC: productos IoT



Fuente de la imagen: <https://www.pexels.com>. Imagen libre

Lo que sigue es un caso real de consultoría en EMC. Una empresa valenciana (que llamaré Cute-4 S.L.) tenía desarrollado e introducido en el mercado un equipo #Producto_antiguo#. Recientemente ha desarrollado tres nuevos equipos que son una evolución del anterior: #Producto1#, #Producto2# y #Producto3#. Cute-4 S.L. desea comercializarlos tanto en la UE como el EE. UU. Los equipos incluyen interfaces inalámbricos (WiFi y/o Bluetooth), un microcontrolador, sensores, entrada de alimentación DC y puerto de comunicaciones.

Acuden a mi preguntando qué normativa deben cumplir los nuevos productos, con la duda de si la certificación del producto antiguo es válida, y me piden también una revisión de los diseños desde el punto de vista de la EMC.

Ya desde el primer contacto la empresa plantea preguntas directas, como la siguiente: “Si tenemos una misma PCB con mismo rutado, y certificamos una versión que lleva todos los componentes soldados, y luego en producción no se montan ciertos componentes, ¿hay que pasar otra certificación? Un ejemplo: Tenemos una placa con un sensor de temperatura y un sensor de movimiento. Certificamos la PCB con los dos sensores y luego se hacen variantes no montando el de temperatura, o no montando el de movimiento para que sean dos productos distintos, pero con misma PCB y puede que misma carcasa.”

No puedo reproducir aquí el informe completo, obviamente, por motivos de privacidad, pero recojo información suficiente para que comprendas el proceso de estudio del caso:

- Identificar la normativa aplicable*
- Descargar y leer la versión más reciente de la normativa*
- Identificar los ensayos y condiciones específicas de aplicación al producto*
- Razonar si las distintas variantes del producto requieren una sola certificación o varias, para cada ensayo*
- Hacer propuestas al diseño para mejorar el desempeño del producto en los ensayos*

Obviamente, hacer todo esto tiene más sentido antes de comenzar el diseño, no después: la empresa hubiera ahorrado tiempo, dinero y sustos...

Extracto del informe (con comentarios añadidos) sobre los productos a certificar por Cute-4 S.L.

Introducción

En el informe “Productos a certificar” de Cute-4 S.L. se recogen las variantes de los productos #Producto1#, #Producto2# y #Producto3# que desea certificar en materia de compatibilidad electromagnética en Europa y/o en Estados Unidos.

Se trata de módulos sensores con interfaz inalámbrica (WiFi y/o Bluetooth) en el interior de edificios, para aplicaciones comerciales.

Lo primero que debemos tener claro es que **todos estos productos son emisores y/o receptores de radiofrecuencia, y como tales están sujetos en Europa a la Directiva de Equipos de Radio (RED, *Radio Equipment Directive* 2014/53/UE)** que ha reemplazado (2015) a la antigua Directiva R&TTE (Directiva de radio y equipos terminales de telecomunicación) bajo la que se certificó el equipo #Producto_antiguo#. Uno de los cambios más relevantes en la nueva normativa es que se considera como producto que cumple la directiva a una combinación específica de hardware y software. Por tanto, un mismo hardware con distintas versiones de firmware requeriría, en principio, más de un proceso de evaluación de cumplimiento.

El campo de aplicación de los equipos de Cute-4 S.L. no está recogido en ninguna de las excepciones de la Directiva (como son los equipos y kits para radioaficionados, equipos marinos o aeronáuticos y equipos de evaluación para profesionales en aplicaciones de I+D).

La RED establece tres requisitos esenciales (si bien ciertas categorías de productos pueden estar sujetas a requisitos adicionales):

1. **Cumplimiento de la Directiva de baja Tensión (LVD, *Low Voltage Directive* 2014/35/UE**, protección de la salud y seguridad para personas y animales domésticos y la protección de los bienes), pero **sin aplicar el límite inferior**. Es decir, aunque los voltajes que maneja el equipo sean inferiores a 75 VDC o 50 VAC, el equipo no está exento.
2. **Cumplimiento de lo dispuesto en la Directiva EMC (2014/30/UE).**
3. **Demostrar el uso eficiente del espectro** con el fin de evitar interferencias con otros equipos en la misma banda. Esto incluye aspectos tales como potencia radiada, ancho de banda ocupado, emisiones fuera de banda y emisiones espurias en recepción. **Este requisito ha de ser necesariamente verificado en un laboratorio especializado**, a diferencia de punto 2 anterior. Será de aplicación EN 300-328 o EN 300-440 en función de la interfaz de radio que equie cada producto Cute-4 S.L.

A partir del 12 de junio de 2018, será necesario la inclusión en un **registro central** de aquellos productos pertenecientes a ciertas categorías, no determinadas en la actual redacción de la RED. Es conveniente leer el artículo 10 de la RED, en especial en lo referente a la información y etiquetado a incluir en el equipo.

La actual declaración de conformidad del equipo #Producto_antiguo# de 5 de mayo de 2015 hace referencia al cumplimiento de la Directiva R&TTE (1995/7/EC) que ya no está en vigor. No obstante, esta declaración de conformidad es válida para las unidades ya comercializadas, no así para las unidades que hayan sido introducidas en el mercado con posterioridad. La disposición transitoria de 2014/53/UE reza:

“Los Estados miembros no impedirán, para los aspectos contemplados en la presente Directiva, la comercialización ni la puesta en servicio de equipos radioeléctricos regulados por la misma que sean conformes con la correspondiente legislación de armonización de la Unión aplicable antes del 13 de junio de 2016 y que hayan sido introducidos en el mercado antes del 13 de junio de 2017.”

Directiva EMC europea para equipos de radio

El segundo requisito fijado por la RED, el cumplimiento de lo dispuesto en la Directiva EMC está ligado por lo general al cumplimiento de los estándares de familia de producto.

Es de aplicación para los equipos a certificar de Cute-4 S.L. la norma [UNE EN 301 489-1](#) v2.1.1 de 2016 (*Common technical requirements*), que define los ensayos de emisión e inmunidad que debe superar el equipo. En función de la interfaz RF, la norma anterior es matizada por [EN 301 489-17](#) v3.1.1 de 2017 (*Specific conditions for Broadband Data Transmission Systems*) en el caso de WiFi y Bluetooth.

La actual declaración de conformidad del equipo #Producto_antiguo# de 5 de mayo de 2015 hace referencia al cumplimiento de EN 301 489-1 y -17 versiones v1.9.2 y 2.2.1 respectivamente.

Ha habido algunos cambios en la normativa desde entonces. Por ejemplo, el ensayo de inmunidad a RF radiada ha de hacerse ahora hasta 6 GHz, mientras que el equipo se testeó por la norma anterior hasta 2,7 GHz. Los ensayos de emisiones se hacen ahora conforme a EN-55032, mientras que el equipo se midió conforme a EN-55022.

EN 301 489-1 establece tres clases de equipos: fijos, portátiles y embarcados (uso en vehículos). La clase de producto determina los ensayos EMC a superar. A la vista del documento “Productos a certificar” entiendo que se trata de productos fijos. Este estándar determina los test de emisiones y de inmunidad que deben superar los equipos, en concreto:

Test de emisiones

- Emisión radiada según EN-55032, para clase B, de 30 MHz a 6 GHz
- Si existen, emisión conducida por puertos de red entre 150 kHz y 30 MHz, según EN-55032, clase B
- Emisión conducida de 150 kHz a 30 MHz por puertos de alimentación AC o DC (en DC sólo si el cable de alimentación es mayor de 3m, en AC es obligatorio), según EN-55032, medida promediada y cuasi-pico.

Nota: Sobre la prueba de emisiones conducidas en puertos DC, para el resto del estudio del caso: Si la alimentación DC (por conector USB) viene de un convertidor AC/DC (equipo auxiliar), hay que medir las emisiones en el lado AC del convertidor. Si viene de un convertidor DC/DC, en el lado de entrada de alimentación del convertidor. En estos productos, se usa un convertidor AC/DC que se vende con el equipo, de modo que se medirán emisiones en el lado AC del convertidor.

- Emisión conducida en puertos de red cableada (como Ethernet), lo que no aplica a estos equipos.
- Emisión de armónicos en red AC (según EN-61000-3-2/A1, clase A)
- Test de fluctuaciones y parpadeo en red AC (según EN-61000-3-3)

Por ejemplo, un equipo no conectado a alimentación AC ni con E/S de señal sólo deberá pasar test de emisiones radiadas (envolvente) y conducidas en el puerto de alimentación DC.

Test de inmunidad

- Descargas electrostáticas (ESD) según IEC 61000-4-2. Test a ± 4 kV por contacto, ± 8 kV por descarga en el aire.
- Transitorios rápidos (EFT) en modo común en puertos (con cable > 3 m) de E/S, señal, control y alimentación DC según IEC 61000-4-4, con nivel de 500 V pico. También en puertos AC (1 kV pico).

- Ondas de choque según IEC 61000-4-5. En productos fijos, solo en puertos de alimentación AC (2 kV línea a tierra, 1 kV entre líneas) y en puertos de red con cable interno de > 30 m (500 V pico, línea a tierra y pantalla a tierra)
- Inmunidad RF radiada según IEC 61000-4-3, de 80 MHz a 6 GHz, campo de 3 V/m modulado en AM 1 kHz al 80%
- Si procede, inmunidad a interrupciones y caídas en la alimentación AC (según IEC 61000-4-11) con caídas al 0% (0.5, 1 y 250 ciclos) y a 70% (25 ciclos)
- Inmunidad a corrientes inyectadas en modo común (IEC 61000-4-6) en puertos de señal y de alimentación AC y DC, de 150 kHz a 80 MHz, 3Vrms, modulación AM 1kHz 80%. Para puertos no AC, aplicar sólo si el cable es > 3m.

EN 301 489-17 define una banda de exclusión en inmunidad para equipos operando en la banda de 2,4 GHz comprendida entre 2.28 GHz y 2.6035 GHz.

Uso eficiente del espectro

Para dispositivos que incluyen WiFi, Bluetooth y ZigBee, es de aplicación la norma **EN 300 328** v2.1.1 de 1/1/2017 (Emisiones de sistemas de transmisión de datos de banda ancha en 2.4 GHz).

La actual declaración de conformidad del equipo #Producto_antiguo# de 5 de mayo de 2015 hace referencia al cumplimiento de EN 300 440.2 v1.4.1 (*Short Range Devices; Radio equipment to be used in the 1 GHz to 40 GHz frequency range*). Es aplicable a productos como juguetes y RFID. En mi opinión, debemos aplicar la EN 300 328, no la EN 300 340.

Normativa FCC para equipos de radio

Los productos que equipan un radiotransmisor ("radiadores intencionados") deben cumplir con los establecido en FCC Parte 15 sub-parte C (abreviado como FCC Part 15C"). Recomiendo el siguiente [enlace](#) para obtener una visión de conjunto de los requisitos. Los requisitos generales (antena, emisiones radiadas y emisiones conducidas) están recogidas en las siguientes tres secciones FCC Part 15C:

- FCC Part 15, Subpart C, **15.203**, establece como requisito que el equipo sólo debe usarse con la antena especificada por el fabricante. Básicamente, dice que si la antena no es reemplazable (como es el caso de antenas impresas en PCB) el cumplimiento del requisito es trivial. Si la antena es reemplazable, debe asegurarse que el usuario no puede instalar una antena distinta a la especificada. Por tanto, el uso de conectores estándar plantea un problema.
- FCC Part 15, Subpart C, **15.207**, establece los límites de emisiones conducidas (150 kHz - 30 MHz) en líneas de alimentación para equipo alimentados red la red AC directamente o a través de cargadores/adaptadores o indirectamente a través de otros equipos (PoE, USB, etc.) Es un ensayo muy parecido al europeo (define los mismos niveles).
- FCC Part 15, Subpart C, **15.209**, establece los límites de emisión radiada a partir de 9 kHz para radiadores intencionados, hasta el décimo armónico de la máxima frecuencia del equipo (ya sea de transmisión o de la parte digital).

Para equipos transmitiendo en la banda de 2.4 GHz se establecen las siguientes provisiones adicionales que pueden afectar a Bluetooth, BLE, WiFi y ZigBee:

- FCC Part 15, Subpart C, **15.247**, establece ciertas condiciones a sistemas que emplean *frequency hopping* (como Bluetooth) y modulación digital (como WiFi, BLE y ZigBee), tales como el ancho de banda mínimo a 6 dB, densidad espectral de potencia, emisiones fuera de banda y emisiones conducidas por el puerto de antena, entre otras. Límite de potencia mucho mayor que en 15.249. Aplicable a radiadores de banda ancha.
- FCC Part 15, Subpart C, **15.249**, establece también condiciones de emisión para radiadores intencionados en la banda de 2.4 GHz. Límite de potencia en torno a 1 dBm, radiadores de banda estrecha. No parece aplicable a los interfaces de radio equipados en los productos de Cute-4, S.L. que estamos considerando.

La actual declaración de conformidad del equipo #Producto_antiguo# es según 15.247.

Certificación FCC

Para radiadores intencionados es necesario obtener certificación a través de un *Telecommunication Certification Body* (TCB), que revisa los resultados obtenidos por un laboratorio certificado. Gracias a un *acuerdo de reconocimiento mutuo* del que participan seis países europeos, en España hay (según búsqueda en el sitio web de FCC) dos laboratorios acreditados para hacer medidas en FCC Part 15 C en:

- Málaga (DEKRA, antiguo AT4)
- Barcelona (LGA Technology Center, APPLUS+)

De estos dos, en el momento de elaborar el informe en 2018, *Dekra era también TCB reconocido por FCC*.

Los equipos deben ser etiquetados de la forma usual (FCC Subpart A, section 15.19). Adicionalmente, hay que incluir, por ejemplo, para equipos Bluetooth, una etiqueta declarando que “Contains FCC ID: ‘FCC ID of the Bluetooth module’”.

En la *base de datos de FCC* queda recogida información del producto. Por ejemplo, para la actual certificación de #Producto_antiguo# aparecen una serie de documentos descargables como resultado de una búsqueda trivial “FCC ID search” en internet. Estos documentos pdf incluyen, entre otros, manual de usuario, fotografías internas y externas, informe de los ensayos EMC y certificación del TCB.

Recomendaciones de cara a la certificación de los cuatro productos basados en la plataforma #Producto1#

Para el **mercado CE**, esta familia de productos sólo deberá pasar tres ensayos EMC (test de emisión radiada, inmunidad ESD e inmunidad RF radiada) además de EN 300 328 (uso eficiente del espectro).

Un cambio en el firmware o en el hardware requiere, en principio, una certificación para cada variante. La opción de montar o no una segunda interfaz RF es causa suficiente para requerir certificación por separado.

En el caso de los dos productos cuya única diferencia es la envolvente de plástico, bastará con repetir el ensayo de inmunidad ESD para la segunda variante, siendo el resto de test compatibles con los realizados sobre la primera variante.

En el caso de la variante que incluye un imán opcional y un tornillo metálico adicional, todos los test serán equivalentes, siempre y cuando un experto informe de que estos elementos adicionales no afectan a la antena del equipo.

Para el **mercado FCC** no se exigen test de inmunidad, pero sí de emisiones radiadas, emisiones conducidas por puertos de alimentación y uso del espectro. Equipos con las diferencias indicadas en el hardware o en firmware requieren, en mi opinión, certificados diferentes.

En el caso de los dos productos cuya única diferencia es la envolvente de plástico, al usar mismo hardware y software, entiendo que pueden tener un mismo certificado.

Nota sobre el rutado del módulo #Producto1#

- Se identifica un par diferencial USB que no parece haber sido rutado correctamente como un par diferencial (la separación entre pistas varía). No parece haberse tenido en cuenta que la impedancia diferencial ha de ser de 90 ohm. Se ruta en capa bottom sobre tres discontinuidades del plano partido de alimentación. Se ruta sin intentar evitar los *antipads* en su ya dañado plano de referencia de retorno de corrientes. Se ruta adyacente y muy cerca de una señal de reloj. Se recomienda resolver todos estos aspectos.
- Se identifica un segundo par diferencial USB que tampoco parece rutado como línea diferencia del 90 ohm

Entiendo que el bus USB es funcional (es decir, que se ha probado y funciona correctamente), pero tal y como está rutado, por lo expuesto anteriormente, aumentará la radiación del producto y podrá poner en peligro superar el ensayo de emisiones radiadas.

- El bus USB no tiene protecciones ESD/EFT junto al conector. El microprocesador utilizado tiene muy poca tolerancia a ESD. Recomiendo añadir, entre los conectores y el filtro RC, protecciones ESD de muy baja capacidad adecuadas para USB.

Recomendaciones de cara a la certificación de los dos productos basados en la plataforma #Producto2#

Para el **mercado CE**, esta familia de productos sólo deberá pasar tres ensayos EMC (test de emisión radiada, inmunidad ESD e inmunidad RF radiada) además de EN 300 328 (uso eficiente del espectro).

Las dos variantes de esta plataforma (...) emplean Bluetooth. Las principales diferencias entre las variantes que puedan afectar a los ensayos son:

- (a) La presencia de la conexión USB en la variante #A#, por la prueba ESD
- (b) La presencia del sensor de presencia en la variante #P# y la correspondiente apertura en la envolvente, por la prueba ESD

Antes estas diferencias mi consejo es pasar test ESD por separado a las dos variantes.

Asumiendo cable USB menor de 3 metros en la variante #A#, los equipos deben ser sometidos, además de a la prueba de inmunidad a ESD (a ± 4 kV por contacto, ± 8 kV por descarga en el aire), a los test de:

- Emisión radiada según EN-55032, para clase B, de 30 MHz a 6 GHz
- Inmunidad RF radiada según IEC 61000-4-3, de 80 MHz a 6 GHz, campo de 3 V/m modulado en AM 1 kHz al 80%

Las dos variantes pueden tener comportamiento diferente debido no sólo a que usan diferentes versiones de firmware, sino a la presencia del cable USB que actúa como antena.

Por tanto, mi consejo es pasar también ensayos de emisión radiada e inmunidad a RF radiada por separado a las dos variantes. De modo que la certificación va a ser independiente para cada variante.

Para el **mercado FCC** no se exigen test de inmunidad, pero sí de emisiones radiadas, emisiones conducidas por puertos de alimentación y uso del espectro. Las dos variantes, en mi opinión, requieren certificados diferentes.

Nota sobre el rutado del módulo #Producto2#

- El bus I2C en el conector USB no tiene protección ESD. Podría dañar los sensores y MCUs conectados. Recomendando añadir protecciones ESD para las líneas de datos, alimentación y masa. Recomendando crear un área de masa sucia (o masa de E/S) donde derivar los transitorios (ánodos de los TVS de protección ESD). Si no es posible, se puede usar la misma masa de circuito (con menor eficiencia).
- Hay áreas de masa en las que se podrían poner vías a masa top-bottom para evitar propagación de RF
- La línea RX es innecesariamente larga
- Que el PCB sea de dos capas lo hace más susceptible a RF externa y aumenta la radiación del PCB. Si hay problemas al pasar los ensayos, cabría plantearse pasar a 4 capas y añadir planos continuos de masa y de alimentación.

Recomendaciones de cara a la certificación de los dos productos basados en la plataforma #Producto3b#

Para el **mercado CE**, el producto PoE está certificado y su comercialización puede seguir adelante, pese a que la normativa EMC de aplicación ha cambiado. Para la variante WiFi, el cambio a alimentación AC y el nuevo módulo de comunicaciones hacen necesario pasar de nuevo los test.

Esta nueva variante deberá pasar test de emisiones radiadas, inmunidad ESD e inmunidad a RF radiada. En caso de querer re-certificar en producto #Producto_antiguo#, cada una de las dos variantes deberá pasar los test completos por separado, ya que: (1) tiene distinto firmware, (2) distinto hardware, (3) distinta envolvente.

Adicionalmente a los ensayos recogidos para los productos #Producto1# y #Producto2#, la variante con alimentación AC debe superar:

- Emisión conducida de 150 kHz a 30 MHz por el puerto de alimentación AC, según EN-55032, medida promediada y cuasi-pico.
- Emisión de armónicos en red AC (según EN-61000-3-2/A1, clase A)
- Test de fluctuaciones y parpadeo en red AC (según EN-61000-3-3)
- Inmunidad a transitorios rápidos (EFT) en modo común en puerto AC según IEC 61000-4-4, con nivel de 1 kV pico.
- Inmunidad a ondas de choque según IEC 61000-4-5 en puerto de alimentación AC (2 kV línea a tierra, 1 kV entre líneas)
- Inmunidad a interrupciones y caídas en la alimentación AC (según IEC 61000-4-11) con caídas al 0% (0.5, 1 y 250 ciclos) y a 70% (25 ciclos)
- Inmunidad a corrientes inyectadas en modo común (IEC 61000-4-6) en el puerto de alimentación AC, de 150 kHz a 80 MHz, 3V rms, modulación AM 1kHz 80%.

Para el **mercado FCC** no se exigen test de inmunidad, pero sí de emisiones radiadas, emisiones conducidas por puertos de alimentación y uso del espectro. Equipos con las diferencias indicadas en el hardware y en firmware requieren, en mi opinión, certificados diferentes.

Nota sobre el rutado del módulo #Producto3b#

- En este módulo sí hay protecciones en la entrada USB. Si es idéntico al diseño #Producto_antiguo# que ya superó pruebas, recomiendo usar este esquema de protección para el resto de equipos con entrada USB.
- Al tratarse de un PCB de 2 capas en el que las líneas USB, y el resto en general, no tiene un camino de retorno de corrientes bien definido, el módulo es más propenso a radiar y a ser susceptible a RF radiadas.

Día 13. Blindajes: atenuando interferencias radiadas



Fuente de la imagen: [Wikimedia](#), bajo licencia Creative Commons Attribution-Share Alike 2.0. Autor: [mattbuck](#)

*La niebla en un muelle londinense nos proporciona una metáfora para entender cómo actúa un blindaje. La luz del sol se refleja en la niebla. La luz que consigue esquivar la reflexión es absorbida a medida que se adentra en la niebla. Ambos efectos, reflexión y absorción, atenúan la luz del sol. **Un blindaje metálico actúa como una niebla de electrones, reflejando y atenuando las ondas que lo penetran.***

En el simplificado lenguaje de la ingeniería, la onda incidente hace vibrar a los electrones libres, transfiriéndoles energía, y es esta vibración sincronizada de cargas la que provoca la aparición de la onda reflejada. Pero no basta con una delgada capa de 1, 10 o 100 átomos de espesor: una pantalla demasiado delgada será prácticamente transparente. Si queremos dar una explicación al fenómeno de absorción, debemos buscarla en las corrientes de Foucault (o corrientes eddy) que originan el efecto pelicular.

¿Qué es un blindaje o pantalla?

Un blindaje o pantalla es una superficie metálica conductora que rodea un sistema electrónico y cuya función es impedir la propagación de campos eléctricos, magnéticos o electromagnéticos desde el exterior hacia el interior del dispositivo o desde el interior del dispositivo hacia el exterior (Figura 13.1). El blindaje puede tomar la forma de una caja metálica, pintura conductora o láminas metálicas o suficientemente conductoras (como ciertas formas del carbono).

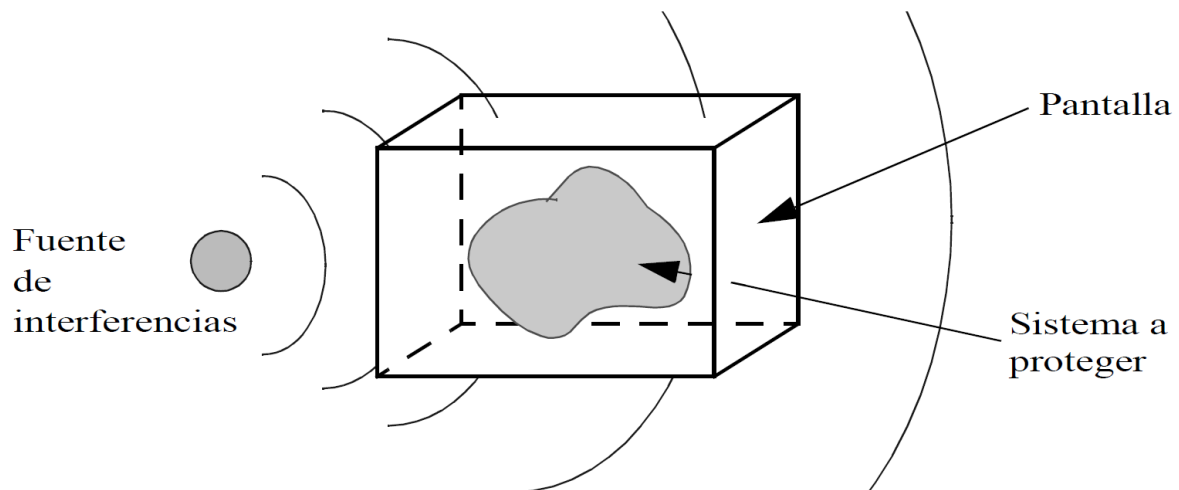


Figura 13.1. Blindaje. Fuente propia

Cuando una onda electromagnética alcanza una lámina conductora (metal o pintura conductora, por ejemplo), el **blindaje o pantalla**, se producen dos efectos que resultan en una disminución de la energía que logra atravesarlo. Estos efectos son la **reflexión** y la **absorción**. Cuando el espesor de la pantalla es menor que la profundidad de penetración (concepto que definiremos más adelante), hay que tener en cuenta también las múltiples reflexiones que tienen lugar dentro del blindaje.

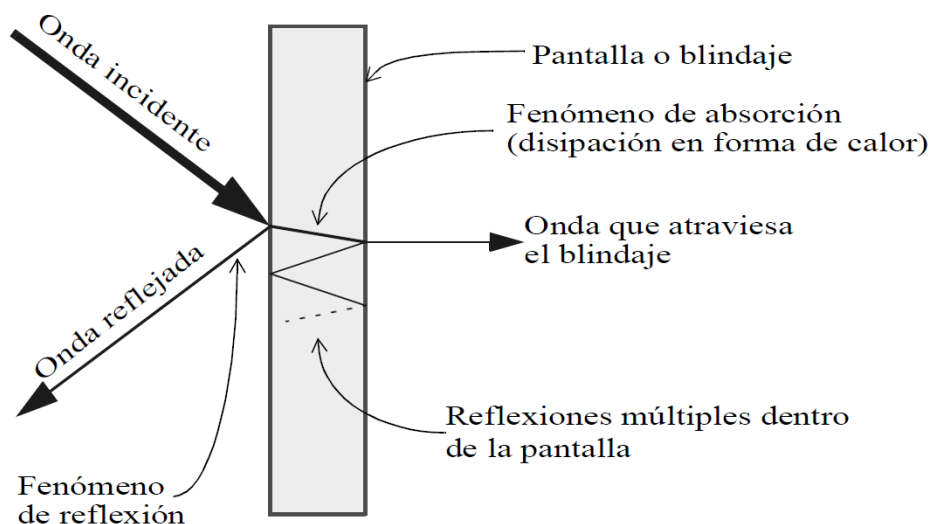


Figura 13.2. Reflexión y absorción en un blindaje. La simplificación de que la reflexión se produce en la interfaz aire-blindaje (útil para hacer cálculos) no debe ocultar que en realidad en este fenómeno participa un espesor no despreciable de la pantalla. Fuente propia

El blindaje puede tomar la forma de una carcasa metálica cerrada que rodea el equipo completo (como en un ordenador), sólo una sección (para proteger, por ejemplo, un chip GPS o un ADC de alta velocidad) o un cable (piensa en un cable coaxial).

La interacción entre el blindaje y la onda electromagnética depende del espesor, conductividad y permeabilidad magnética de la pantalla, así como de la frecuencia y de las características de la onda: ¿estamos en campo lejano, en campo cercano predominantemente eléctrico o magnético? (si lo necesitas, repasa la sección Campo cercano y campo lejano en la página 32).

Los parámetros anteriores resultan en una figura de mérito, de nominada [efectividad del blindaje](#) a una frecuencia dada, que expresa en decibelios la atenuación en potencia -ya sabes, $20 \cdot \log(E_i/E_o)$ o $20 \cdot \log(H_i/H_o)$ - que sufre de la onda al atravesar la pantalla.

Esta efectividad teórica se ve muy degradada por [ranuras, juntas y aberturas en el blindaje](#). Como regla sencilla, si mantienes la dimensión mayor de los orificios pequeña comparada con la longitud de onda (al menos $\lambda/10$) la degradación es pequeña. Por eso verás por lo general orificios de ventilación circulares y pequeños en lugar de ranuras alargadas (si bien la normativa de seguridad eléctrica es otro motivo para hacer esto: evitar que puedas llegar con los dedos a elementos a tensión potencialmente peligrosa).

¿Qué necesitas aprender de blindajes y pantallas?

En primer lugar, será necesario que comprendas [cómo funciona una pantalla electrostática o jaula de Faraday](#). En baja frecuencia, los electrones pueden moverse la suficiente distancia entre pico y valle de la onda como para permitir una redistribución de carga: aparece así un campo eléctrico que se opone al externo, logrando un campo eléctrico nulo en la región rodeada por la pantalla. A medida que aumenta la frecuencia del campo los electrones apenas tienen tiempo de moverse para dar lugar a una adecuada redistribución de cargas, atenuando el efecto de jaula de Faraday.

En cuanto al campo magnético continuo o de baja frecuencia, una pantalla electrostática no tiene efecto apreciable sobre el campo. Sólo si la pantalla está fabricada con un material ferromagnético lograremos canalizar las líneas de campo en su interior, evitando su aparición en la región interior a proteger, pero esto no es una atenuación.

[A frecuencias mayores de unos pocos kilohercios, dependemos de los fenómenos de reflexión y de absorción](#), así como un tercer efecto consecuencia de los anteriores: las reflexiones múltiples dentro del espesor de una pantalla delgada. Comprendiendo estos efectos podrás elegir el material de la pantalla y su espesor para lograr una determinada atenuación a una frecuencia determinada. Para esto necesitas tener expresiones para [calcular la efectividad del blindaje](#), y definir así el espesor de la pantalla.

A continuación, tener unas nociones de [cómo degrada la efectividad del blindaje una ranura](#) te permitirá fijar requisitos mecánicos, así como la conveniencia de usar soluciones comerciales para evitar fugas en las juntas.

Como ves, hay que asimilar varios conceptos (lo que lleva tiempo y esfuerzo) y aplicar algunas expresiones sencillas para hacer cálculos (lo que no requiere una cosa ni la otra). Vamos a comenzar en la siguiente sección por presentar el fenómeno de reflexión.

Lo que ya sabías: apantallamiento electrostático

En primer curso de carrera, en la asignatura de física, y sin duda también en bachillerato, habrás estudiado electrostática, es decir, los fenómenos eléctricos que se producen debido a distribuciones de cargas en equilibrio. Y entre estos fenómenos te hablaron de la ley de Gauss para el campo eléctrico y del **apantallamiento electrostático**.

Una pantalla electrostática consiste en una superficie metálica que rodea a un equipo o circuito a proteger. Un campo electrostático externo produce una redistribución de cargas en la pantalla (mira la Figura 13.3), creando un campo eléctrico de sentido contrario (más carga negativa cerca de la positiva externa, más carga positiva cerca de la negativa externa), de modo que el campo en el volumen interior es nulo. No es necesario conectar la pantalla a tierra.

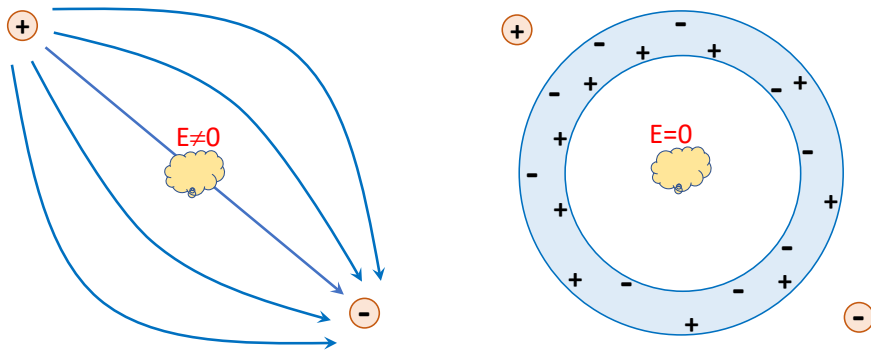


Figura 13.3. La redistribución de cargas en la pantalla compensa el campo electrostático externo, resultando en un campo nulo en el volumen interior. Fuente propia

¿Ocurre lo mismo cuando queremos proteger al exterior del efecto de cargas en el interior? Se produce igualmente una redistribución de cargas en la pantalla, pero debido a que la carga neta del conjunto no es nula, sigue produciéndose un campo eléctrico en el exterior (Figura 13.4 izquierda): sólo si conectamos la pantalla a tierra (o a un conductor grande) puede neutralizarse la carga total del sistema y permitir que el campo eléctrico en el exterior sea nulo (Figura 13.4 derecha).

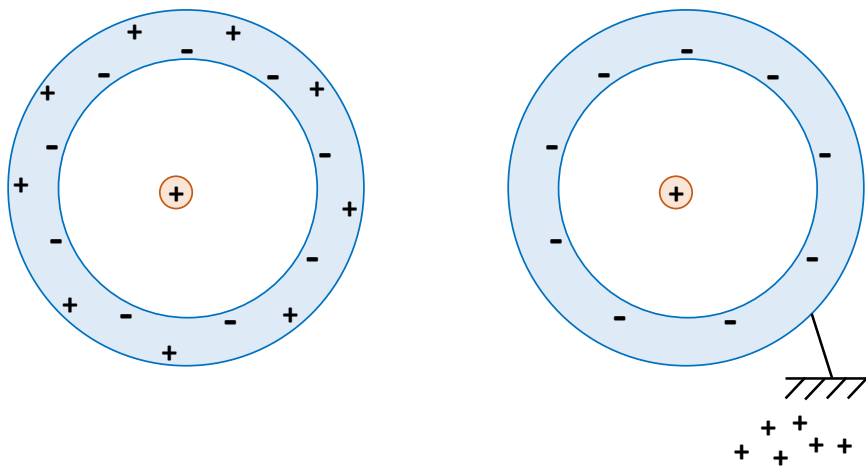


Figura 13.4. Para proteger al exterior de campos eléctricos generados en el interior de la pantalla electrostática, debemos conectarla a tierra para cargarla (Nota: si ahora eliminamos la conexión a tierra, la pantalla queda cargada). Fuente propia

Lo anterior es válido para campos estáticos o de una variación tan lenta que permita tiempo la redistribución de cargas. De hecho, es posible construir una pantalla electrostática de madera: sólo has de esperar el tiempo suficiente (menos de 30 segundos en el experimento que encontrarás en este [enlace](#), y que te recomiendo encarecidamente estudiar) a que se produzca la redistribución de cargas.

¿Qué pasa con los campos magnéticos estáticos? Si no hay variación ($\partial B/\partial t=0$) no se inducen corrientes ni efecto alguno en la pantalla, de modo que una pantalla metálica de cobre o de aluminio será transparente al campo: no habrá forma de apantallarlos. Sólo si la pantalla es ferromagnética encauzará el campo a través suyo y evitará que el volumen interior de la pantalla esté sometido a la acción del campo.

Resumiendo: una superficie conductora que rodea un circuito electrónico no tiene efecto alguno frente a campos magnéticos estáticos o de muy lenta variación. En el caso de campos eléctricos estáticos o de lenta variación, el circuito queda protegido de la acción de campos externos. El exterior sólo queda protegido de la acción de nuestro circuito si está conectado a tierra.

*Pero en EMC nuestro problema se extiende a frecuencias mucho más elevadas, hasta varios GHz. Una vez abandonamos las muy bajas frecuencias una pantalla deja de funcionar como hemos descrito y depende de los fenómenos de **reflexión y de absorción** para atenuar los campos que inciden sobre su superficie. Ambos fenómenos se basan en la interacción fotón-electrón (llamada difusión), o si prefieres decirlo de otra manera, en la vibración inducida por los campos en los electrones libres del metal.*

*Veremos en las siguientes secciones cómo reflexión y absorción tienen un efecto variable con la frecuencia (disminuye el primer efecto, aumenta en segundo). Pero llega una frecuencia a la que los electrones ya no pueden moverse lo bastante rápido y la efectividad de la pantalla decrece: se trata de la **frecuencia de plasma**. Esto ocurre en el rango de los petahercios (10^{15} Hz), en la luz ultravioleta, así que no debe preocuparnos desde el punto de vista de la EMC.*

Reflexión en una superficie conductora

Hay varias interpretaciones o modelos que explican el fenómeno de la reflexión. Puedes leer la excelente explicación que da Richard Feynmann en su magnífica obra “Electrodinámica cuántica” [32] sobre la reflexión de la luz en un espejo en el capítulo 1 y sobre la interacción electrón-fotón en el capítulo 3. Pero has de tener en cuenta que, en radiofrecuencia, las dimensiones de átomos y moléculas son menores que la longitud de onda, por lo que puedes optar por una descripción macroscópica (ondas) en vez de microscópica (fotones).

Por tanto, puedes preferir quedarte con la idea más sencilla de que el campo incidente hace moverse a los electrones libres de la pantalla, lo que crea una onda que se radia hacia el exterior (es decir, se emplea la energía incidente para radiar una onda hacia el exterior).

O puedes hablar de impedancias de onda y de blindaje e interpretar la reflexión tal y como hablamos en el Día 2, que es una útil herramienta de cálculo. Vamos a optar por esta última, por ser la más sencilla de comprender y de calcular, si bien la más alejada de lo que realmente está ocurriendo.

Al llegar a la pared exterior del blindaje, el campo (eléctrico, magnético o electromagnético) se encuentra con un cambio brusco de impedancia y parte de él es reflejado y no atraviesa el blindaje. El coeficiente de reflexión, que mide la fracción de campo que es reflejado, se calcula como:

$$|\rho| = \left| \frac{Z_m - Z_\omega}{Z_m + Z_\omega} \right|$$

donde Z_m es la **impedancia del medio** que forma la pantalla y Z_ω es la **impedancia del medio** por el que se propaga la onda antes de llegar a la pantalla (también llamada impedancia de onda). De forma general, la impedancia de un medio, **en campo lejano** (cuando estamos más allá de $\lambda/2\pi$ de la fuente), se calcula como:

$$Z_m = \sqrt{\frac{j\omega\mu}{\sigma + j\omega\epsilon}}$$

donde se pueden hacer dos aproximaciones para medios no conductores y medios conductores, respectivamente. **Para un medio no conductor como el aire o el vacío**, si despreciamos la conductividad:

$$\sqrt{\mu_0/\epsilon_0} = 120\pi \approx 377\Omega$$

valor ya conocido por el lector como impedancia característica (cociente E/H) de una onda en campo lejano en el aire o en el vacío. **Para un medio conductor**,

$$Z_m = (1 + j) \cdot \sqrt{\frac{\pi\mu f}{\sigma}}$$

Pero la impedancia del medio en campo cercano no sigue estas expresiones. Depende de si se trata de un campo predominantemente magnético (E/H pequeño) o eléctrico (E/H elevado), según la Figura 13.5.

A medida que nos alejamos de la fuente, el cociente E/H cambia y tiende a 377Ω . Si tienes en cuenta que $\lambda/2\pi$ es de 2 cm a 2,4 GHz y de 2 m para 24 MHz, podemos considerar que casi siempre nos encontraremos en campo lejano y asumir que la impedancia de onda es de aproximadamente 377Ω .

Una pantalla, siendo conductora, presentará una impedancia pequeña a baja frecuencia, que aumentará con $\sqrt{f}$ aproximadamente.

Para campo lejano o campo eléctrico cercano, $Z_\omega \gg Z_m$ a frecuencias bajas y medias, y el coeficiente de reflexión será alto (cercano a la unidad). Como consecuencia, cualquier pantalla metálica, por delgada que sea, proporcionará una buena protección.

A medida que aumenta la frecuencia, el fenómeno de reflexión es cada vez menos importante (los electrones pueden moverse una distancia menor en cada ciclo de la onda, disminuyendo el efecto de reflexión) y habrá que fiar la efectividad del blindaje al fenómeno de absorción.

Para campo magnético cercano el coeficiente de reflexión es pequeño y una parte significativa de la energía de la onda penetra en el blindaje.

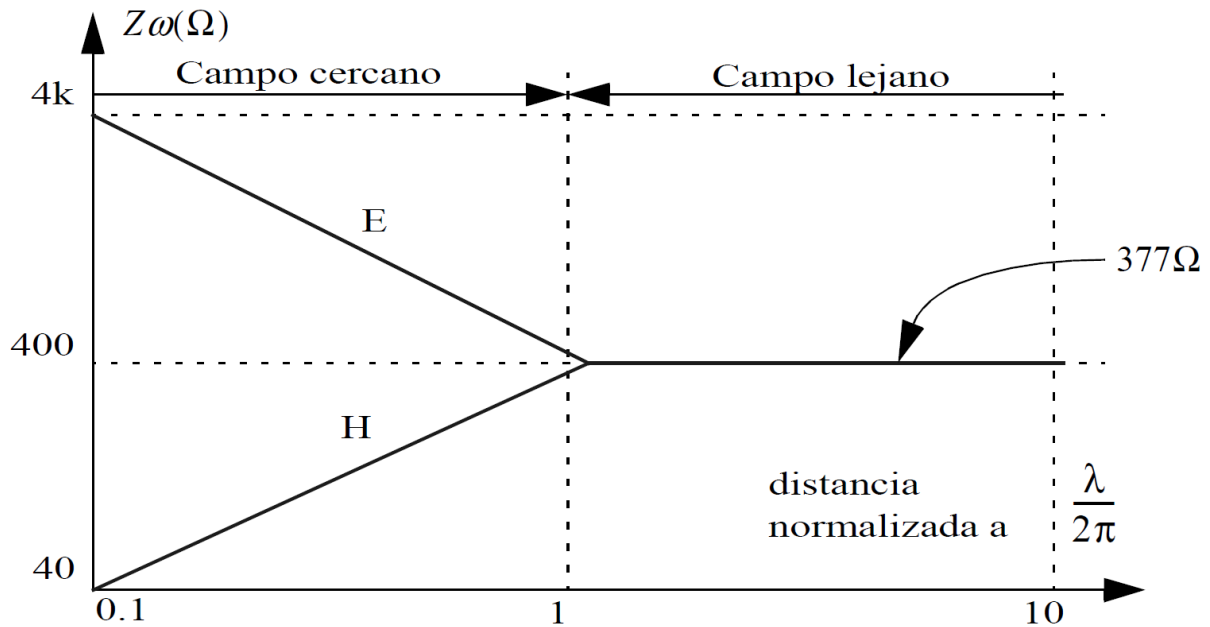


Figura 13.5. Impedancia de un medio no conductor (en este caso, aire) en campo cercano y en campo lejano. En campo eléctrico (magnético) cercano, la intensidad disminuye como $1/d^3$, mientras que el magnético (eléctrico) lo hace como $1/d^2$. En campo lejano, ambos se atenúan como $1/d$. Fuente propia

Absorción en un medio conductor

Cuando una onda electromagnética penetra en el blindaje, parte de la energía es disipada en forma de calor por efecto Joule sobre las corrientes inducidas en el material del blindaje. ¿Recuerdas lo que estudiaste sobre corrientes eddy (o de Foucault)?

El efecto pelicular como fundamento de la absorción en un blindaje

La mayoría de los textos sobre blindajes renuncian a proporcionar una explicación física que fundamente la acción de un blindaje y comienzan directamente proponiendo expresiones de cálculo. No me parece mal, de hecho, a menudo es conveniente. Simplemente creo que tu curiosidad (y la mía) quedará mejor satisfecha si entendemos qué está ocurriendo exactamente cuando decimos que una onda se atenúa al penetrar en un blindaje. Como ingenieros, podemos cometer el atrevimiento de intentar atisbar el mecanismo oculto dentro de la caja negra: ya hay demasiadas cajas negras en la ingeniería.

La ley de Faraday de la inducción

Si consideras la superficie delimitada por un hilo conductor cerrado (espira, bucle), y si hay una variación de flujo magnético a través de esta superficie, los electrones libres se verán empujados a lo largo del conductor, dando lugar a una corriente eléctrica (carga en movimiento). El sentido de la corriente es tal que el flujo magnético creado se opone a la variación que da lugar a la corriente. Es decir, un campo magnético variable da lugar a un campo eléctrico variable:

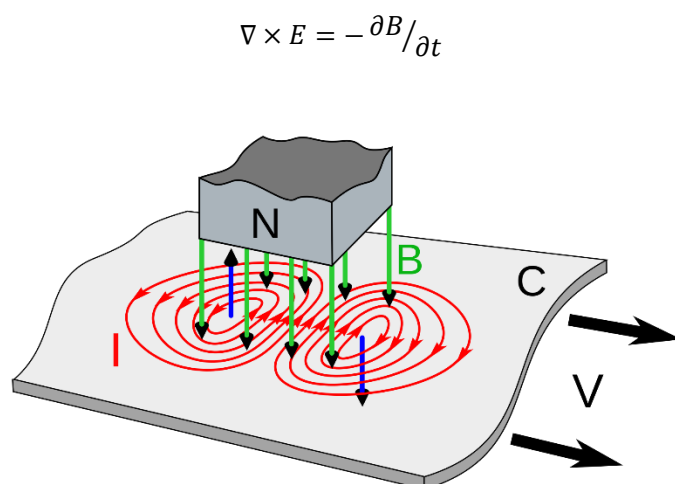


Figura 13.6. Generación de corrientes eddy en un metal bajo la acción de un campo magnético variable. Créditos de la imagen: By Chetvorno - Own work, CC0, <https://commons.wikimedia.org/w/index.php?curid=34157558>

Corrientes de Foucault o corrientes eddy

Acabamos de decir que un campo magnético variable en el interior de un conductor da lugar a líneas de campo eléctrico perpendiculares a la dirección del campo que se cierran, forzando a los electrones a formar corrientes en círculos que crean un campo magnético que se opone al que las creó.

Como resultado, el campo se reduce dentro del conductor, y la energía del campo se disipa como energía térmica a medida que los electrones chocan entre sí. A estas corrientes eléctricas en círculos, inducidas por un campo magnético variable, se les denomina corriente de Foucault o corrientes eddy.

Este fenómeno aumenta con la frecuencia del campo (onda) incidente, ya que el término $\partial B / \partial t$ es proporcional a la frecuencia.

También aumenta con la conductividad del metal, ya que, a mayor conductividad, mayor es la corriente inducida.

El tercer factor que afecta a la generación de las corrientes eddy es la permeabilidad magnética del conductor, ya que, a mayor permeabilidad, mayor es el flujo magnético.

Todo esto combinado da lugar al...

Efecto pelicular

Si un campo magnético variable se atenúa a medida que penetra en un medio conductor, porque las corrientes eddy inducidas van disipando en forma de calor la energía del campo, y si este efecto aumenta con la frecuencia de la onda y con la conductividad y permeabilidad magnética del medio conductor, te resultará al menos creíble que te diga que hay una ley exponencial decreciente que expresa la amplitud del campo a medida que penetra una profundidad x en el metal:

$$E(x) = E_0 \cdot e^{-x\sqrt{\sigma\pi\mu f}}$$

$$H(x) = H_0 \cdot e^{-x\sqrt{\sigma\pi\mu f}}$$

Donde definimos la profundidad de penetración, δ , como aquella a la que el campo (y las corrientes inducidas) se atenúa un factor $1/e$, lo que representa $20 \cdot \log(e) \approx 8,7$ dB. El 98% del campo es atenuado al penetrar cuatro veces δ . La profundidad de penetración se calcula como:

$$\delta = \frac{1}{\sqrt{\sigma\pi\mu f}}$$

Y las expresiones anteriores pueden reescribirse como:

$$E(x) = E_0 \cdot e^{-x/\delta}$$

$$H(x) = H_0 \cdot e^{-x/\delta}$$

Recuerda que $\mu_0 = 4\pi \cdot 10^{-7}$ y sus unidades son T·m/A (tesla·metro/amperio). En el cobre, a 100 MHz, $\delta = 6,61$ μm . De modo que, en alta frecuencia, un campo que penetre en un blindaje se atenúa fuertemente tras unas cuantas decenas de micras. Un campo de baja frecuencia penetrará y posiblemente atravesará el blindaje.

Las pérdidas por absorción son iguales para campos cercanos y lejanos y, al contrario que en el fenómeno de reflexión, aumentan con la frecuencia.

Cálculo de blindajes

La efectividad de un blindaje puede expresarse según las expresiones siguientes para campos eléctricos y magnéticos. E_i , H_i son las intensidades de campo incidentes y E_o , H_o las intensidades de campo que atraviesan el blindaje.

$$S_E = 20 \cdot \log \frac{E_i}{E_o}$$

$$S_H = 20 \cdot \log \frac{H_i}{H_o}$$

Para ondas electromagnéticas en campo lejano se cumple que $S_E = S_H$.

La efectividad, expresada en decibelios, es la suma de tres componentes: pérdidas por reflexión (R), pérdidas por absorción (A) y un factor de corrección por reflexiones múltiples (B). Esto se expresa como:

$$S = R + A + B$$

Las expresiones clásicas que permiten calcular los distintos componentes (R, A, B) se muestran en la Figura 13.7, que deben usarse según se indica en la Figura 13.8.

A, R, B en dB

f : frecuencia en MHz

t : espesor del blindaje en cm

d : distancia desde la fuente en metros

δ : profundidad de penetración en cm

$$\delta = 0,0066 \cdot \sqrt{\sigma_r \mu_r} f$$

- 1) $A = 1314,3 \cdot t \cdot \sqrt{\mu_r \sigma_r} f$
- 2) $R = 108,1 - 10 \cdot \log(\mu_r f / \sigma_r)$
- 3) $R = 141,7 - 10 \cdot \log(\mu_r d^2 f^3 / \sigma_r)$
- 4) $R = 74,6 - 10 \cdot \log\left(\frac{\mu_r}{d^2 f \sigma_r}\right)$
- 5) $B = 20 \cdot \log(1 - e^{-2t/\delta})$

Figura 13.7. Expresiones para el cálculo de blindajes

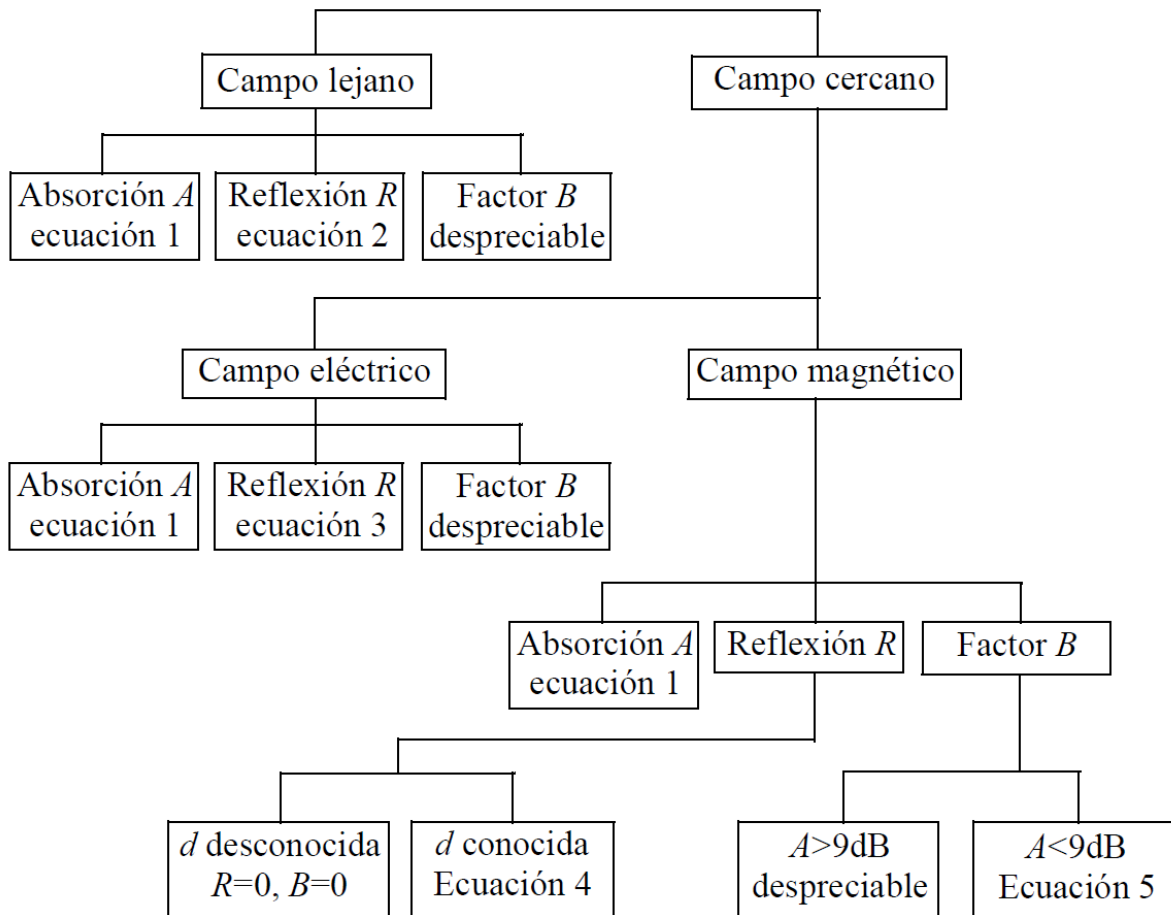


Figura 13.8. Algoritmo para cálculo de blindajes

Un blindaje, **para ser efectivo contra campo eléctrico cercano**, debe tener una alta conductividad (σ) para que las pérdidas por reflexión sean altas (cobre, cromo o aluminio. La plata y el oro son demasiado caros).

Para ser efectivo contra campos magnético cercano, debe tener una alta permeabilidad (μ) para que las pérdidas por absorción sean altas (por ejemplo, mediante materiales ferromagnéticos como el hierro o el acero).

Un **campo lejano** se atenuará adecuadamente por reflexión (y queremos una superficie con buena conductividad), y si es necesario, se hará uso del espesor del material para aportar pérdidas por absorción.

Efecto de las ranuras

En la práctica, las pantallas suelen tener orificios de ventilación, están atravesadas por cables no apantallados y las juntas no son estancas (por ejemplo, la carcasa de un PC). De hecho, estas imperfecciones degradan tanto la efectividad del blindaje que 60 dB está ya considerado como una buena efectividad.

En general, **es preferible tener varios orificios pequeños que uno sólo pero grande**. La degradación de la efectividad es función de la mayor dimensión del orificio y no de su área.

- Aumentar en un factor L la longitud del orificio degrada la efectividad del blindaje en aproximadamente $20 \cdot \log(L)$
- Tener n orificios pequeños degrada la efectividad del blindaje de forma proporcional a la raíz cuadrada de n , es decir como $10 \cdot \log(n)$
- Combinando los dos puntos anteriores, para un área dada de ranuras (generalmente determinada por criterios térmicos) **degradamos la efectividad del blindaje la mitad de decibelios si usamos un grupo de pequeños orificios que uno solo de la misma área total**.

Recuerda que son las corrientes inducidas en el blindaje las que originan tanto la reflexión como la absorción de la onda incidente, y allí donde hay ranura no hay corrientes inducidas. Cuanto menor sea la distorsión en las corrientes inducidas (lo que se logra con orificios pequeños comparados con la longitud de onda), menor la reducción de efectividad del blindaje.

Materiales empleados en los blindajes para RF

Lo que sigue es un resumen de un artículo publicado en diciembre de 1993 en la revista EMC Test&Design [32], con algunos añadidos de cosecha propia.

Hablamos de RF, radiofrecuencia, para referirnos a ondas en campo lejano. En este caso, para ahorrar material, la mejor opción es usar un buen conductor para tener una elevada reflexión.

El **cobre** es uno de los materiales más utilizados, junto a una de sus aleaciones, el latón. Sus buenas características (maleable, no se oxida rápidamente) ayudan a trabajarlo y darle formas complejas si es necesario. Su alta conductividad, el hecho de que el óxido se elimina fácilmente (y sigue siendo un razonable conductor) y que es poco susceptible a la corrosión son factores que juegan a su favor. Su elevado coste y el hecho de no ser ferromagnético juegan en su contra.

Lo encontramos habitualmente en pantallas RF con un espesor comprendido entre una décima de un milímetro (lo más habitual) y un milímetro (si necesitamos mayor resistencia mecánica o pérdidas adicionales por absorción a bajas frecuencias).

El **aluminio** es también un material muy usado. Su conductividad es aproximadamente la mitad que la del cobre, por lo que requiere un mayor espesor de material para alcanzar la misma efectividad del blindaje que otro de cobre. Es ligero pero resistente y puede ser extruido, lo que lo hace adecuado para estructuras de mayor tamaño. En contra, su óxido presenta una muy baja conductividad (de hecho, es un material cerámico) y es propenso a la corrosión. El aluminio se oxida en pocas horas de exposición a la atmósfera, de modo que las juntas y uniones que requieran continuidad eléctrica deben ser limpiadas de óxido (por ejemplo, por abrasión) antes de realizarse.

El **acero**, al ser ferromagnético, es usado para aplicaciones que requieran apantallar campo magnético de baja frecuencia. El cobre y el aluminio no sirven para este propósito, siendo superiores en todas las demás aplicaciones. El óxido es un problema grave, y aún con un tratamiento de galvanización, sigue siendo sensible a la corrosión. Es un material difícil de trabajar, de modo que su uso suele estar restringido a campo magnético en baja frecuencia. Para hacerlo más complicado, hay muchas variedades de acero y sus propiedades magnéticas cambian en función de la variedad.

Cobre, aluminio y acero son los tres materiales más utilizados. Hay que prestar especial atención a las uniones entre placas de metal, eliminando óxido, corrosión y suciedad antes de realizar las uniones. Para garantizar la continuidad eléctrica se usan materiales de interfaz.

Una cinta de cobre con adhesivo es un buen material de interfaz hasta 1 GHz. También se suelen usar tiras de cobre, de cobre estañado o de estaño. En este último caso, resulta que es compatible (para evitar la corrosión) con el cobre, el aluminio y el acero.

Tabla 13.1. Profundidad de penetración en el metal (μm) y atenuación ideal en dB en campo lejano para un espesor de 1 mm

	1 kHz	1 MHz	100 MHz	1 GHz
Cu	2100 / 108	66 / 210	6,6 / 1372	2,1 / 4202
Al	2600 / 104	82 / 187	8,2 / 1074	2,6 / 3265
Acero inox 304	13000 / 75	415 / 82	41 / 246	13 / 677

Tal y como puedes ver en la Tabla 13.1, más importante que la elección del material es evitar orificios y ranuras en el blindaje. El acero inoxidable 304 es de tipo austenítico (es razonablemente maleable, soldable e inoxidable) pero pierde casi todas sus propiedades magnéticas: no debes suponer que un acero tiene una elevada permeabilidad, hay muchos tipos de aceros.

Corrosión galvánica: cuándo no juntar dos metales

Diferentes metales tienen diferentes conductividades, y por tanto se produce una diferencia de potencial entre ellos al entrar en contacto. Esto da lugar a desplazamiento de electrones desde el metal anódico (por ser el más propenso a cederlos, que por tanto se oxida) al catódico (que gana electrones). En presencia de un electrolito (por ejemplo, humedad), el metal anódico es el que se corroe, es decir, sus iones positivos se disuelven en el electrolito, eliminando así metal del ánodo, picómetro a picómetro, hasta que se satura el electrolito o (lo más habitual) hasta que desaparece el metal que hace de ánodo.

Para minimizar la corrosión galvánica, nunca hay que poner en contacto directo dos metales que estén alejados en la [serie galvánica](#), una lista ordenada de metales. Del más catódico al más anódico, podemos citar la siguiente serie resumida:

Oro, plata, titanio, acero inoxidable, latón, cobre, acero, estaño, aluminio, zinc y magnesio

Por tanto, poner en contacto directo cobre y aluminio es una mala idea, ya que el aluminio tenderá a sufrir corrosión. El estaño, en el centro de la tabla, es compatible (como ya hemos comentado) con cobre, acero y aluminio.

Qué metal usar como material de interfaz tiene su importancia desde el punto de vista de la corrosión y el estaño se presenta como una buena opción, junto a otras buenas propiedades: fácil de soldar y buena conductividad eléctrica, tanto en su forma pura como oxidada.

Ejemplo de cálculo de blindajes

Hace tres o cuatro años, un físico con quien suelo colaborar me pidió calcular la atenuación que una lámina de 2-5 mm de espesor de un tipo especial de acero inoxidable (aleado con cromo y níquel, para hacerlo resistente a la corrosión, pero con muy baja permeabilidad magnética) tendría sobre un campo magnético de 1 kHz en campo cercano.

Por supuesto, los físicos saben mil veces más que un ingeniero sobre los fundamentos, mecanismos y efectos de primer y segundo orden que rigen la propagación de campos electromagnéticos a través de metales. Y son capaces de realizar complicadas simulaciones para obtener resultados muy exactos.

Pero nosotros, los ingenieros, sabemos echar mano de expresiones simplificadas y sin rubor alguno (fruto de nuestra ignorancia sobre la complejidad subyacente) devolver un resultado numérico. A continuación, mi respuesta a su solicitud (en inglés, pues la consulta se hizo en esa lengua), que reutilizo aquí como ejemplo de cálculo de blindajes.

Dear ...,

Austenitic stainless steels, especially grades 310 and 316, have very low magnetic permeability, especially when in annealed condition. Nevertheless, as cold working increases the magnetic permeability, the calculations below are to be taken as optimistic.

In general, the higher the nickel to chromium ratio the more stable is the austenitic structure and the less magnetic response that will be induced. AK steel 316 has a permeability of 1.02 (when in the annealed condition) and a resistivity of 74 mΩ·cm, equivalent to a conductivity of 1.351e+6 Siemens/m.

The sheet thickness will be in the range 2-5 mm and the field frequency will be below 1 kHz. With these parameters we can calculate the attenuation of the steel sheet.

Reflection losses

We start by calculating the impedance of the steel sheet at 1 kHz:

$$Z_{\text{sheet}} = \sqrt{\frac{2\pi\mu f}{\sigma}} = \sqrt{\frac{6.28 \cdot 10^3 \cdot 4\pi \cdot 1.02 \cdot 10^{-7}}{1.352 \cdot 10^6}} = 77.2 \mu\Omega$$

*The wave impedance is E/H, and it will be very low in the near-field region. We can calculate it at a distance d=0.2 m from the magnetic field source as (I have assumed we are in near field, d<wavelength/(2*pi)):*

$$Z_{\text{wave}} = 2\pi f \mu_0 d = 6.28 \cdot 10^3 \cdot 4\pi \cdot 10^{-7} \cdot 0.2 = 1.58 \text{ m}\Omega$$

Reflection losses are $Z_{\text{wave}}/(4 \cdot Z_{\text{sheet}}) = 4.86$, or 13.7 dB

Absorption losses

This is the main attenuation mechanism for low-frequency magnetic fields. We will start by calculating the skin depth at several frequencies:

$$\delta_{1\text{kHz}} = \frac{1}{\sqrt{\sigma\pi\mu f}} = \frac{1}{\sqrt{1.35 \cdot 10^6 \cdot 3.14 \cdot 1.02 \cdot 4\pi \cdot 10^{-7} \cdot 10^3}} = 13.55 \text{ mm}$$

One delta thickness produces an attenuation factor of 1/e. So, a 5-mm sheet will produce an attenuation factor of $e^{5/13.57} = 1.44$, or 3.2 dB

Correction factor due to multiple reflections inside the sheet

As the sheet is thin compared to the skin depth, absorption is low and thus we must consider the effect of multiple reflections inside the metal sheet (thus reducing the effectiveness of the shield). This can be considered as:

$$1 - e^{-2t/\delta} = 0.52, \text{ or } -5.7\text{dB}$$

This subtracts to the attenuation due to absorption and reflection.

Overall loss

$$\text{Reflection} \cdot \text{Absorption} \cdot \text{Correction} = 4.86 \cdot 1.44 \cdot 0.52 = 3.64, \text{ or } 11.2 \text{ dB}$$

This means that the H field passing through the stainless-steel sheet is 1/3.64 of the incident field.

Conclusions

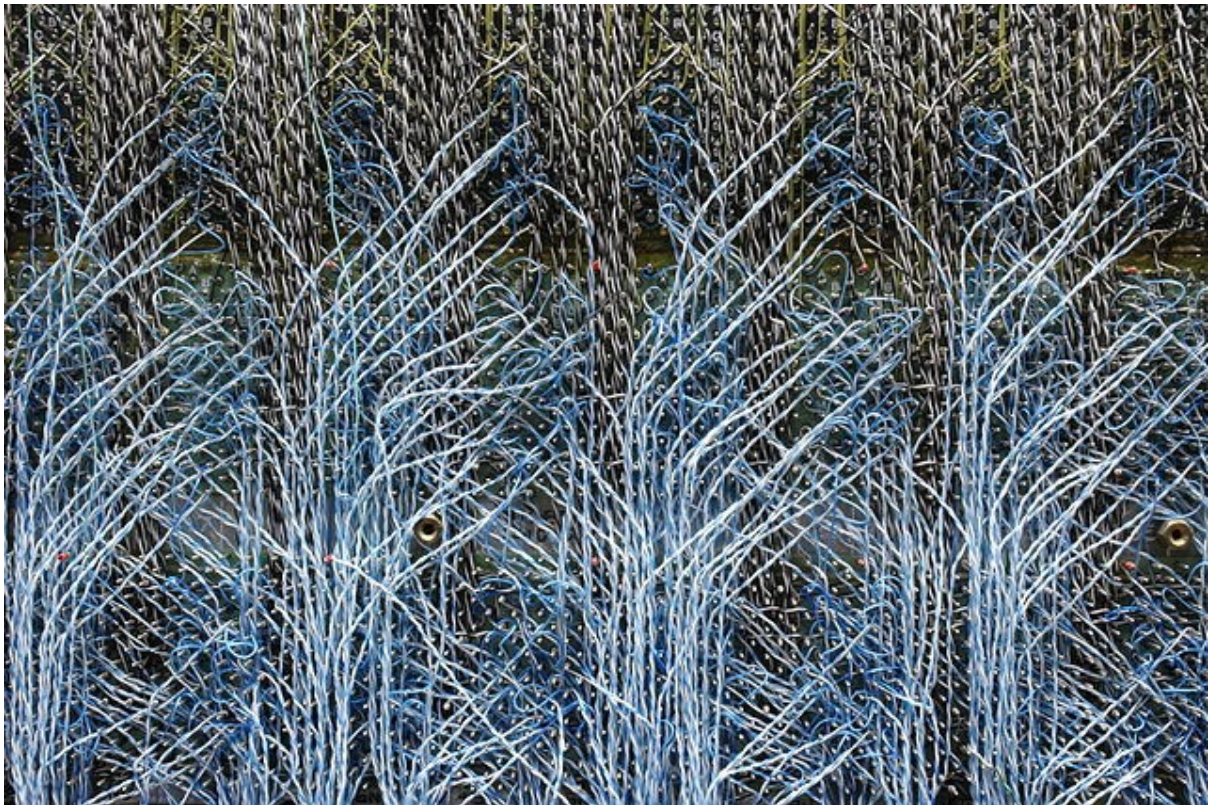
The loss is higher than I expected. Put in the right numbers (like frequency and d) and get the real number.

Vuelta a 1993...

En ese año, Acorn Computers Ltd. publicó un interesante y fácil de leer documento, [EMC Design Guidelines](#) [33], que dedica buena parte de su extensión al cálculo de blindajes y al efecto de las ranuras. Por aquel entonces ni siquiera era de obligado cumplimiento la EMCD, pero los diseñadores ya necesitaban reglas y principios para reducir emisiones e incrementar la inmunidad de sus equipos.

También habla de particionado de funciones en el PCB y otras buenas ideas que pueden ayudarte a reflexionar. Si sientes nostalgia por el pasado, echa un vistazo al documento en este [enlace](#).

Día 14. Interferencias en líneas de PCB, cables planos, coaxiales y pares trenzados



Una compleja instalación de pares trenzados. [Imagen](#) con licencia Creative Commons Attribution-Share Alike 3.0 Unported. Autor: [Shieldforyoureyes](#)

Desde el punto de vista de la EMC, los cables presentan problemas asociados al acoplamiento de energía (diafonía o crosstalk) entre líneas dentro del cable, entre cables y entre el entorno y el cable. Hay que distinguir entre acoplamiento en campo cercano y campo lejano, entre baja, media y alta frecuencia. ¿Nuestro objetivo? Aprender a hacer un mejor uso de cables planos, pares trenzados, cables coaxiales y apantallados para reducir problemas de inmunidad y de emisión.

Dejaremos para otro día el empleo de filtros y ferritas, con frecuencia asociados a cables que entran o salen de nuestro equipo, centrándonos hoy exclusivamente en las interconexiones. Va a ser un día inevitablemente confuso que va a requerir lectura, dudas, nueva lectura, iluminación, más dudas, una nueva lectura...

Hoy iremos saltando entre dominios, intentando comprender fenómenos y efectos en relación con los cables y la EMC: qué ocurre en alta frecuencia y qué ocurre en baja frecuencia, qué ocurre en campo cercano y en campo lejano, qué sucede con el campo eléctrico y con el campo magnético, autoinmunidad e inmunidad a interferencia externas, ... No hay una respuesta categórica del tipo "conecta siempre el blindaje a del cable a tierra por ambos extremos". La respuesta correcta es un compromiso que surge de analizar las distintas fuentes de interferencia que afectan a una situación concreta y sopesar cuánto ganamos y cuánto perdemos con cada posible solución.

Diafonía capacitiva entre dos cables eléctricamente cortos

Comencemos por algo básico. Cuando existe una diferencia de potencial entre dos conductores separados una cierta distancia, se crea un condensador, una capacidad eléctrica. Por ejemplo, dos cables cercanos a diferente potencial eléctrico crean una capacidad eléctrica. Un plano de masa y la caja metálica de un equipo crean una capacidad eléctrica, si existe (que existirá) diferencia de potencial eléctrico entre ambos.

Este condensador no es un elemento concentrado (como un componente electrónico) sino que se trata de una capacidad distribuida. Esta capacidad distribuida, generalmente no deseada, recibe en nombre de **capacidad parásita**. Si el circuito es eléctricamente corto (recuerda, esto quiere decir que las longitudes involucradas son pequeñas comparadas con la longitud de onda λ a la frecuencia que estemos considerando), representamos este efecto como un condensador ideal entre los dos conductores (C en la Figura 14.1). Esta capacidad forma parte del “**esquemático oculto**”: el conjunto de elementos parásitos no recogidos en un diagrama esquemático y que alejan el comportamiento de circuito real del ideal.

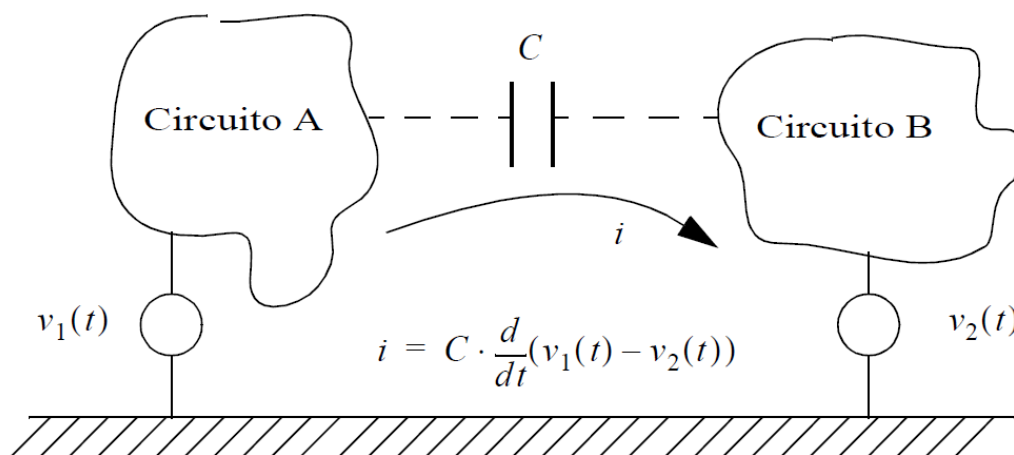


Figura 14.1. Interferencia por acoplamiento capacitivo. Fuente propia

Si la diferencia de potencial entre los dos conductores es variable, se induce una corriente, denominada **corriente de desplazamiento**, fruto de la atracción y repulsión de electrones en cada electrodo debidas al campo eléctrico, sin que haya electrones que salten realmente entre ambos conductores.

Esta corriente de desplazamiento es directamente proporcional al valor de la capacidad y a la velocidad de variación (derivada) de la diferencia de tensión con el tiempo (Figura 14.1). Podemos distinguir dos casos prácticos de acoplamiento capacitivo:

1. Acoplamiento entre un circuito y un chasis o plano de masa. La importancia del chasis o el plano de masa estriba en que representa un conductor de área elevada y, lo deseemos o no, todos los nodos del circuito van a presentar una capacidad de acoplamiento (parásita) a este conductor extenso.
2. Acoplamiento entre cables o pistas de circuito impreso, fenómeno denominado **diafonía capacitiva**, también *capacitive crosstalk* en la literatura en lengua inglesa.

Para simplificar el problema podemos suponer que un conductor va a actuar como fuente de interferencias (existe una diferencia de potencial variable entre éste y un plano de masa o tierra) y el otro conductor va a ser el receptor de la interferencia (Figura 14.2). Si los cables son eléctricamente cortos, la corriente interferente circulará por la resistencia R , dando lugar a una tensión acoplada en el conductor receptor de interferencias. Si los cables se comportan como una línea de transmisión a las frecuencias de la interferencia, el frente de ondas inyectará una corriente en el cable víctima, corriente que se propagará tanto hacia la derecha como hacia la izquierda, y dará lugar a distintas formas de onda en función de las impedancias que encuentre la interferencia en los extremos del cable víctima.

En el caso de baja frecuencia (no consideramos líneas de transmisión) la tensión inducida sobre la resistencia es, expresada en el dominio temporal, el producto de la resistencia R y la corriente de desplazamiento en la capacidad parásita:

$$v_R(t) \approx R \cdot C \cdot \frac{d}{dt}v(t)$$

Esta aproximación es válida siempre que $R \ll 1/\omega C$. Cuanto mayor sea R mayor será la tensión acoplada, por lo que la diafonía capacitiva es muy importante en circuitos víctima con alta impedancia de entrada.

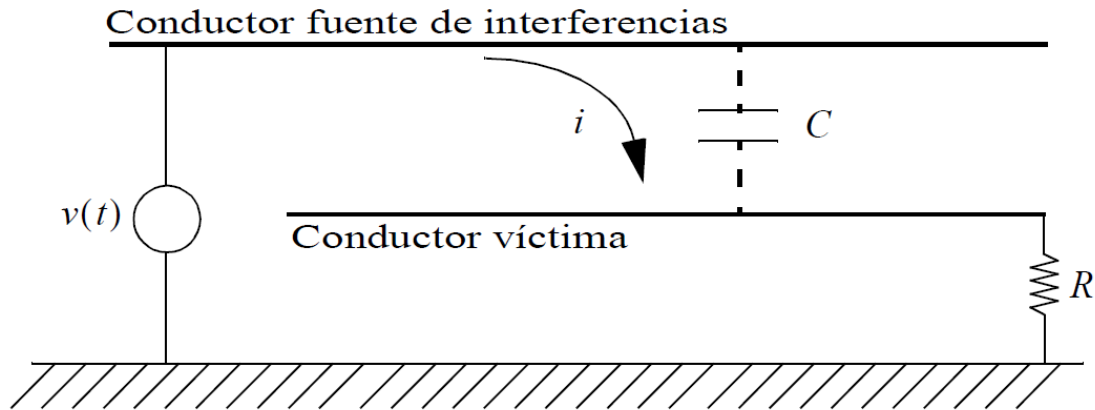


Figura 14.2. Diafonía capacitiva. Fuente propia

Ejemplo: Calcular la tensión inducida por diafonía capacitiva entre la red eléctrica y el circuito de la Figura 14.3. Suponemos una carga R_L de 1 kΩ.

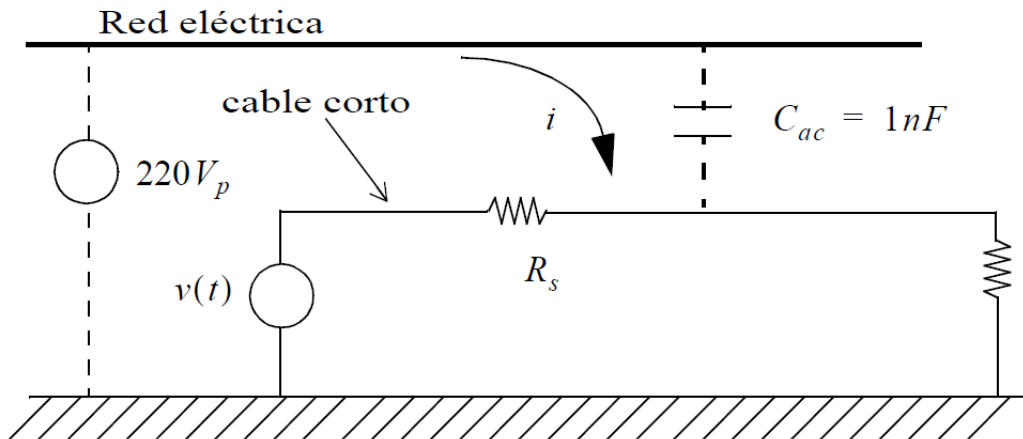


Figura 14.3. Ejemplo de acoplamiento de la red eléctrica. Fuente propia

Podemos resolver el circuito aplicando superposición, por lo que la tensión en la carga será igual a la suma de las contribuciones de las dos fuentes (la de señal y la tensión de red). Si llamamos R al paralelo de R_s y R_L , entonces:

$$v_r(t) \approx R \cdot C_{ac} \cdot \frac{\delta}{\delta t}v(t)$$

El valor máximo de la derivada es $220 \cdot 100 \cdot \pi$, lo que multiplicado por C_{ac} da un valor de $69 \mu A$. El valor de pico de la tensión inducida en R_L será entonces de $69 mV$. Para comprobar que la aproximación es válida basta con evaluar la expresión $R \ll 1/\omega C_{ac}$, que es cierta en este ejemplo.

Ejemplo: Diafonía capacitiva entre dos pistas de un circuito impreso (Figura 14.4).

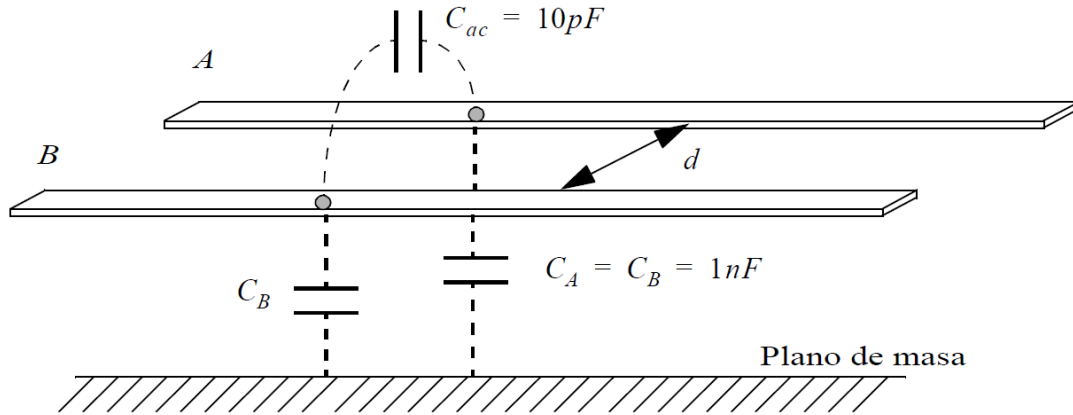


Figura 14.4. Diafonía capacitiva entre pistas de un circuito impreso. Fuente propia

Por la pista A se propaga una señal digital de $3 V$ de pico y un tiempo de subida de $5 ns$. La pista B tiene una resistencia equivalente a masa de 50Ω . Si no tenemos en cuenta C_A ni C_B y asumiendo una línea eléctricamente corta, la tensión inducida será:

$$V_{ind} = R_L \cdot C_{ac} \cdot \frac{3V}{5ns} = 50 \cdot 10^{-11} \cdot \frac{3}{5 \cdot 10^{-9}} = 300mV$$

La simplificación es válida mientras se cumpla que $R_L = 50 \Omega \ll 1/\omega \cdot C_{ac}$, lo que deja de ser cierto a partir de unos $30 MHz$. El modelo equivalente del circuito, teniendo en cuenta C_A y C_B se muestra en la Figura 14.5.

Observamos que C_A no influye en el análisis. La capacidad C_B en paralelo con la resistencia de carga forma un filtro paso bajo con una frecuencia de corte definida por $f_c = 1/(2\pi R_L C_{ac})$, que en nuestro ejemplo es de aproximadamente $3,2 MHz$.

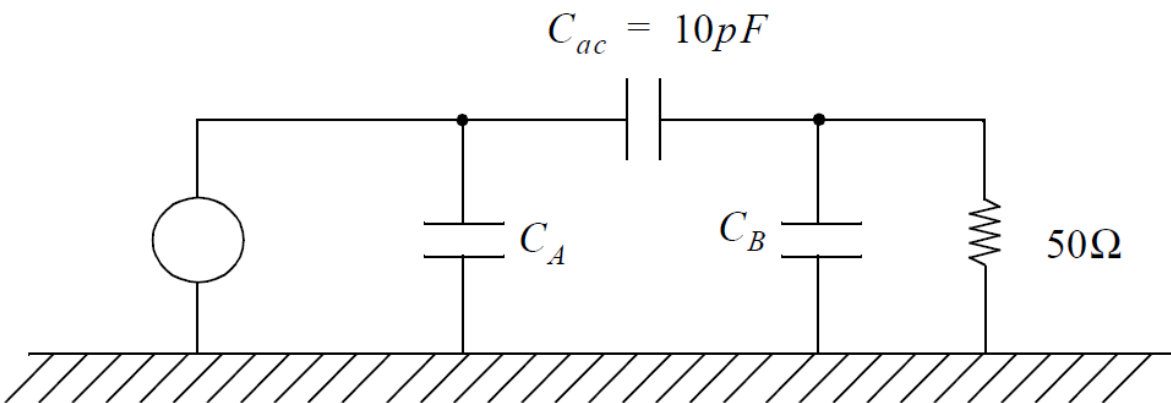


Figura 14.5. Circuito equivalente de la Figura 14.4. Fuente propia

Si tenemos en cuenta que una señal digital de tiempo de subida t_r tiene componentes espectrales significativas hasta una frecuencia determinada por la expresión $0,5/t_r$, que en nuestro ejemplo es de 100 MHz, resulta que la aproximación es válida par pistas de hasta 7,5 cm (equivalente a $\lambda/20$). A partir de ahí, tendremos que estudiar el problema con otro método.

En una placa de circuito impreso, el mejor modo de reducir la diafonía capacitiva es aumentar la separación entre pistas y evitar trazar tramos largos de pistas paralelas.

Resumiendo: el acoplamiento capacitivo se manifiesta mediante la aparición de una corriente acoplada. Aumentar la separación entre los conductores, disminuir la velocidad de variación de las tensiones implicadas, disminuir la separación pista-plano (aumentar C_B en la Figura 14.5) son medidas para reducir el acoplamiento.

Diafonía inductiva sobre bucles eléctricamente cortos

Este fenómeno recibe también el nombre de **crosstalk inductivo**. Consiste en la aparición de una diferencia de potencial en bucles cercanos a un conductor por el que circula una corriente variable (Figura 14.6). La diafonía es nula en corriente continua, como ya sabes por la [Ley de Faraday](#).

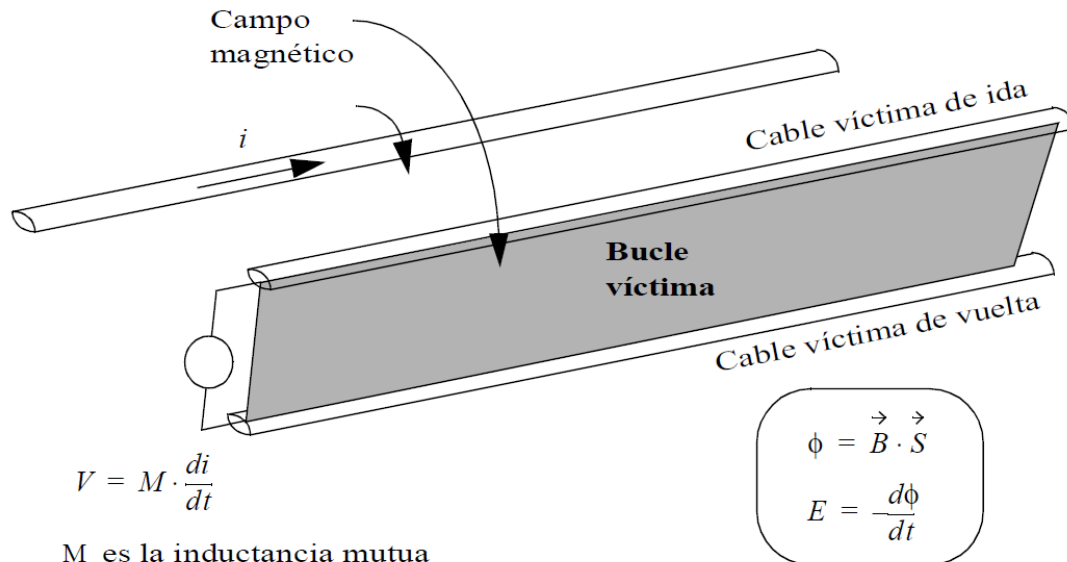


Figura 14.6. Fenómeno de la diafonía inductiva. Fuente propia

La tensión generada es menor cuanto más pequeña sea la variación de la corriente que produce la interferencia y cuanto más pequeño sea el área del bucle víctima. La diafonía inductiva es más grave cuando afecta a circuitos de baja impedancia (al contrario de lo que ocurría con la diafonía capacitiva). **Un cable plano supone el caso más desfavorable de diafonía inductiva** (Figura 14.7), debido a:

- a. Hay varios conductores paralelos de gran longitud
- b. Hay generalmente una única línea de masa

Esto implica la existencia de grandes bucles y coeficientes de inductancia mutua elevados. La diafonía en un cable plano puede reducirse de tres maneras:

1. Disminuyendo el área del bucle (por ejemplo, utilizando cables cortos o trenzando líneas de señal con líneas de masa)
2. Disminuyendo el coeficiente de acoplamiento entre hilos (intercalando líneas de señal y de masa en el cable, apantallando el cable o situando un plano de masa bajo el cable)
3. Reduciendo la velocidad de variación de la señal interferente

Consideremos la Figura 14.7 e imaginemos que el hilo 8 es de masa y por el hilo 4 introducimos una señal. Está claro que en los bucles formados por cada uno de los otros hilos y masa se inducirá un fuerte acoplamiento inductivo. Este tipo de esquema (una única línea de masa) se usa cuando queremos reducir la anchura del cable, las señales agresoras son poco abruptas (producen poca interferencia) y las víctimas son robustas (pueden soportar interferencia sin degradar las prestaciones del equipo).

Un esquema más robusto supone disponer de un mayor número de líneas de masa, lo que reduce el área de los bucles agresores y víctima y reduce el ruido en masa. En un caso casi ideal, cada línea de señal tiene una línea de masa y ambos hilos van trenzados. Como veremos en la siguiente sección, el ruido acoplado es menor si empleamos pares trenzados.

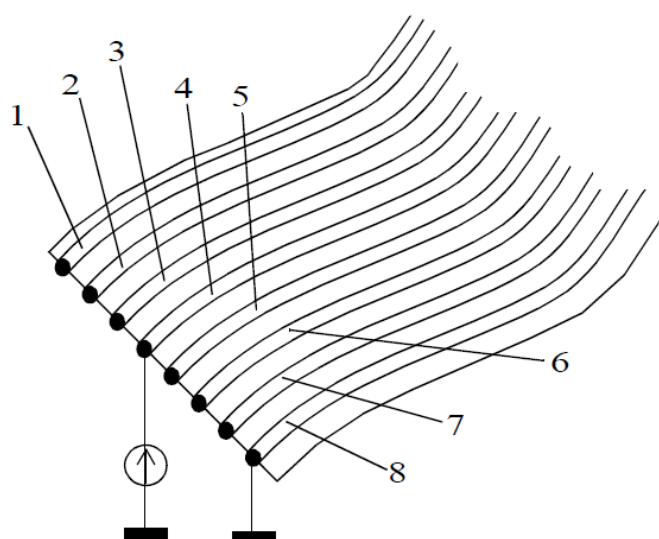


Figura 14.7. Cable plano. Fuente propia

Diafonía entre pistas de un PCB (líneas no eléctricamente cortas)

Cuando las líneas no son eléctricamente cortas tenemos que considerarlas como líneas de transmisión, cada una con sus dos conductores. **Recuerda: la pista es sólo la mitad del camino de la señal. El (o los) planos de referencia son la otra mitad de la línea de transmisión.** Resulta intuitivo que a mayor distancia pista-plano, mayor *crosstalk* de tipo inductivo (los bucles son de mayor área) y capacitivo (por comparación de la capacidad pista-plano frente a pista-pista, es lo mismo que estudiamos al hablar de la impedancia diferencial el Día 4, ¿te acuerdas?) También, la proximidad entre líneas aumenta la energía acoplada entre la agresora y la víctima. **Por tanto, la diafonía o *crosstalk* en un PCB es un problema dominado por la geometría.**

Pero hay un segundo ingrediente: lo abruptos (rápidos) que sean los flancos en la línea agresora. Al aumentar $\partial V/\partial t$ y $\partial i/\partial t$ (también $\partial E/\partial t$ y $\partial H/\partial t$) aumenta el acoplamiento de energía entre líneas de transmisión: recuerda aquello de $i=C \cdot \partial V/\partial t$ y $V=L \cdot \partial i/\partial t$. **Dicho de otro modo, sólo los flancos en la línea agresora generan un acoplamiento en la víctima**, de forma que la interferencia en la línea víctima se parece a la derivada de la señal en la línea agresora (Figura 14.9).

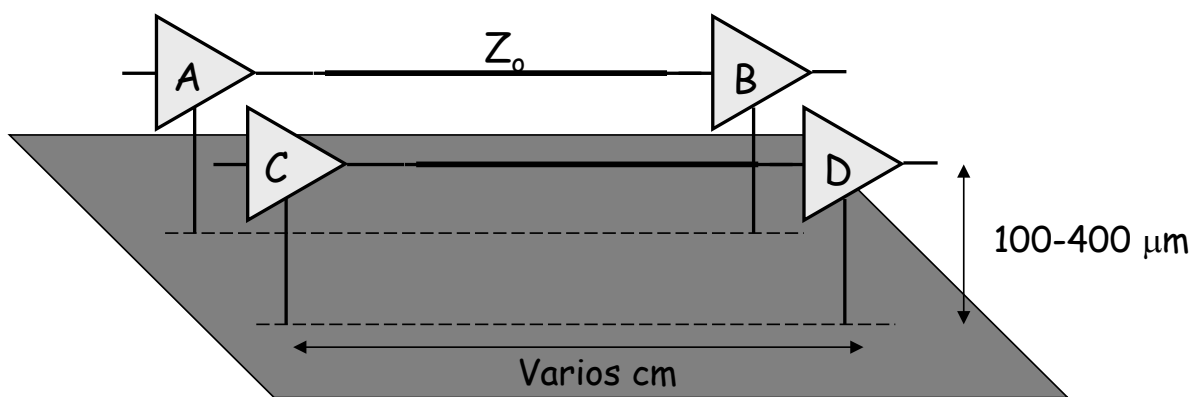


Figura 14.8. El acoplamiento de energía (*crosstalk*) entre dos pistas en un PCB está determinado por la geometría y por lo abruptos que sean los flancos en la línea agresora. Fuente propia

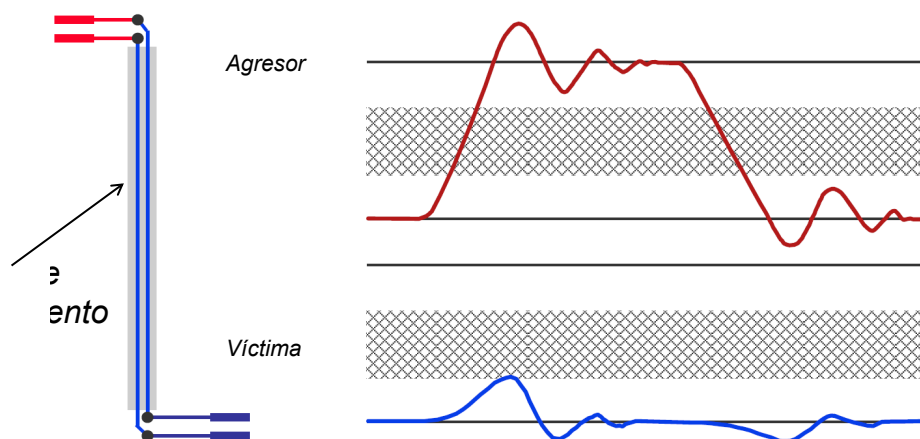


Figura 14.9. Sólo los flancos en la línea agresora generan un acoplamiento en la víctima

¿Sabes cómo identificar si el *crosstalk* es inductivo o capacitivo?

Al cambiar la polaridad de la corriente en el circuito agresor, la polaridad de la interferencia cambia sólo si el *crosstalk* es de origen inductivo. La realidad es que siempre vas a observar una mezcla de ambos, pero en ciertos casos domina uno de los dos y es importante saberlo para poder reducirlo con las medidas adecuadas.

Vocabulario básico: NEXT, FEXT, forward y backward

Tenemos que pensar en un frente de ondas propagándose por la línea agresora, y cómo este frente de ondas induce a medida que se propaga una onda electromagnética en la línea víctima (Figura 14.10). Y, aquí viene lo interesante, esta onda acoplada se propaga en ambas direcciones por la línea víctima: la que viaja en el mismo sentido que el frente de ondas agresor se denomina *forward crosstalk*, mientras que el que lo hace en el sentido contrario recibe en nombre de *backward crosstalk*. Si lo prefieres, son las componentes directa e inversa de la onda acoplada.

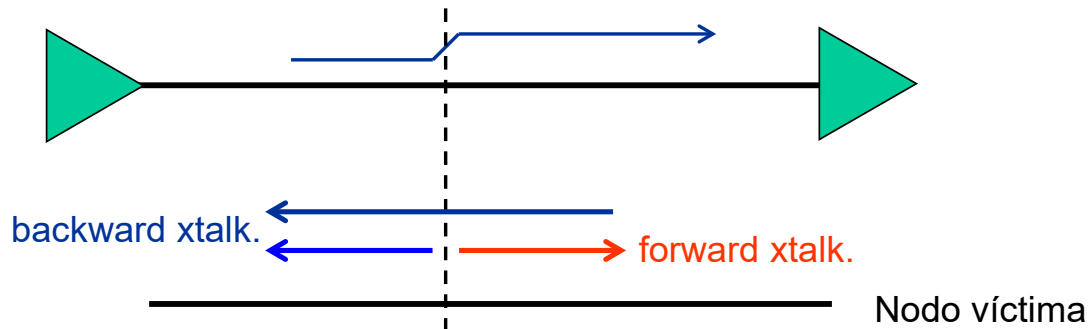


Figura 14.10. Crosstalk directo (*forward*) e inverso (*backward*). Fuente propia

El *crosstalk capacitivo* induce una señal que se propaga en el nodo víctima con la misma polaridad en cada dirección (componentes directa e inversa).

El *crosstalk inductivo* crea una componente que se propaga en el nodo víctima con polaridad opuesta en cada dirección.

Por tanto, en el nodo víctima aparece una componente *de crosstalk* en el extremo opuesto, **Far-End Crosstalk o FEXT** (capacitivo restado del inductivo, que tienden a anularse, sobre todo en líneas *stripline*) y otra componente en el extremo cercano **Near-End Crosstalk o NEXT** (capacitivo sumada a inductivo, se refuerzan). En líneas *microstrip* la componente directa es muy pequeña.

Todo esto es confuso, no te preocupes. Vamos a explicarlo algo mejor, pero sin dedicar más tiempo del necesario: nuestro objetivo es aprender a controlar la amplitud de *crosstalk*, no aprender a dibujar su forma de onda.

Un caso simplificado: estudio de crosstalk en líneas sin reflexiones

Supongamos que no hay reflexiones en la línea agresora ni en la línea víctima. Y supongamos que las líneas son eléctricamente largas. En la Figura 14.11 hemos definido un experimento con estas condiciones: dos líneas acopladas de 50 Ω y 15 cm de longitud en capa externa.

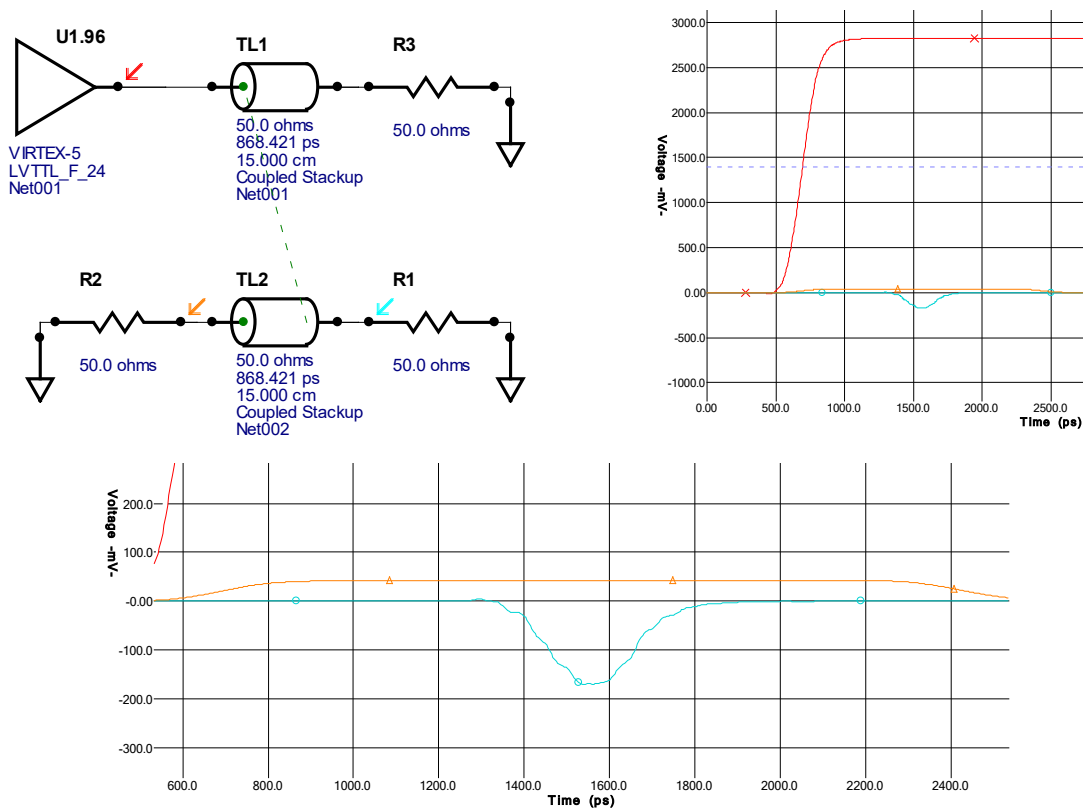


Figura 14.11. Experimento para estudiar el *crosstalk* directo e inverso. Fuente propia

El **crosstalk directo** (en azul claro en la Figura 14.11) crea un pulso de anchura constante igual al tiempo de duración del flanco del pulso interferente (460 ps en este ejemplo) y su amplitud aumenta con la longitud de la línea, hasta llegar a un máximo. Es como si en una bañera empujaras el agua con la mano en una dirección: la amplitud de la ola aumenta con la distancia, se va acumulando la energía. Pero la ola no se hace más ancha.

El **crosstalk inverso** (en naranja) crea un pulso de anchura igual al doble del tiempo de propagación de la línea y de amplitud constante. Sus aproximadamente 1700 ps de anchura son el doble que el tiempo de propagación en la línea (868 ps). Siguiendo con el símil de la bañera, al desplazar la mano creas una onda también en sentido opuesto al del movimiento de la mano. La creas desde el primer instante y durante todo el tiempo que mueves la mano (tiempo de propagación de la línea), y cuando la detienes, la última energía creada todavía se tiene que propagar hasta el otro extremo de la bañera: de ahí que la duración de este pulso sea el doble del tiempo de propagación de la línea.

Observa que ambos tipos de *crosstalk* tienen amplitudes de signo contrario. En este experimento simplificado, el *forward crosstalk* coincide con el FEXT y el *backward crosstalk* con el NEXT. En el caso general, se suman dando lugar a formas de onda más complejas. Vamos a estudiarlo en el experimento de la siguiente sección.

Estudio de crosstalk en un caso más real

Modificamos el experimento anterior para permitir reflexiones en ambos extremos, tanto de la línea agresora como

de la víctima. En la

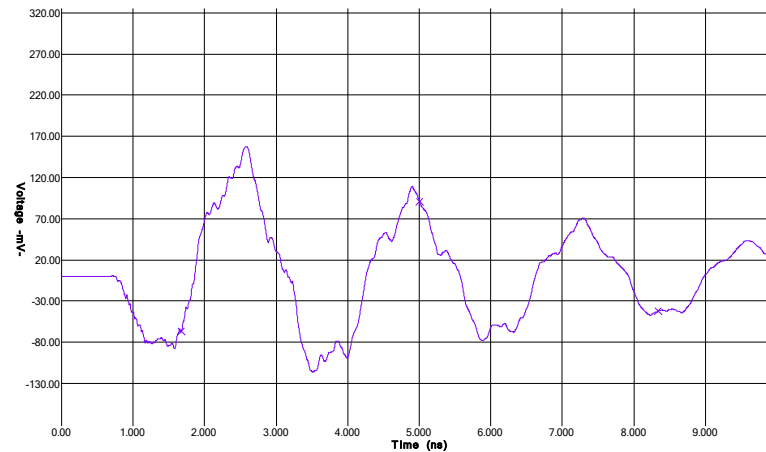


Figura 14.12, U3 está definido como “stuck low”, es decir, que genera un nivel bajo constante. Podemos observar las formas de onda NEXT (en U3) y FEXT (en U4). Fuente propia

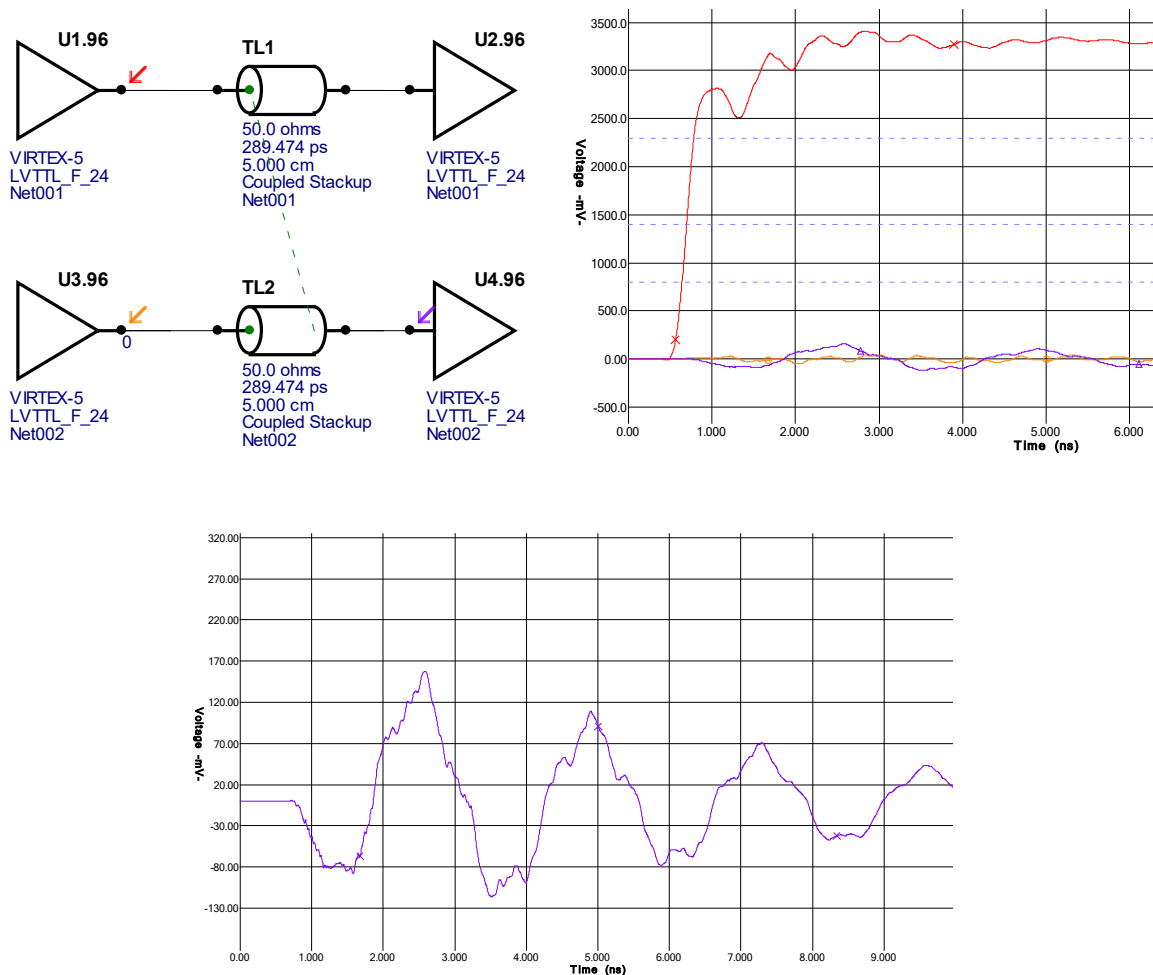


Figura 14.12. Experimento de *crosstalk* en una situación realista. El FEXT (abajo) tiene una amplitud de pico de casi 170 mV. Fuente propia

Para reducir el *crosstalk* puedes:

- Separar las pistas: normalmente, una separación entre pistas igual a tres veces su anchura es una regla sencilla de diseño que suele dar buenos resultados
- Escoger una familia lógica o un *driver* configurable con tiempos de transición tan grandes como puedas permitirte
- Aumentar la proximidad de la pista al plano (o planos) de referencia empleando dieléctricos delgados
- Evitar que las pistas discurren paralelas a lo largo de tramos largos
- Tener especial cuidado con el rutado de las señales más peligrosas (relojes) y más sensibles (reset, por ejemplo)

Cables planos

Ribbon cable en inglés, es posiblemente el tipo de conexión más empleado para interconectar dos módulos electrónicos a distancias cortas. La Figura 14.13 recoge cuatro ejemplos de cable plano.

Arriba a la izquierda, el más habitual: un conjunto de hilos metálicos paralelos envueltos en un material aislante. No hay blindaje o pantalla, ni plano de masa de referencia, y el área de los bucles formados por los hilos de señal y de retorno son mayores cuanto menor es el número de líneas de masa. Arriba a la derecha, la misma configuración, pero más compacta: pistas de cobre sobre un sustrato flexible (típicamente de kapton, el sustrato sobre el que se fabrican los PCBs flexibles). Abajo, dos variantes que mejoran las prestaciones: usar pares trenzados señal-masa (izquierda) y disponer de un plano de masa continuo bajo las señales (derecha).

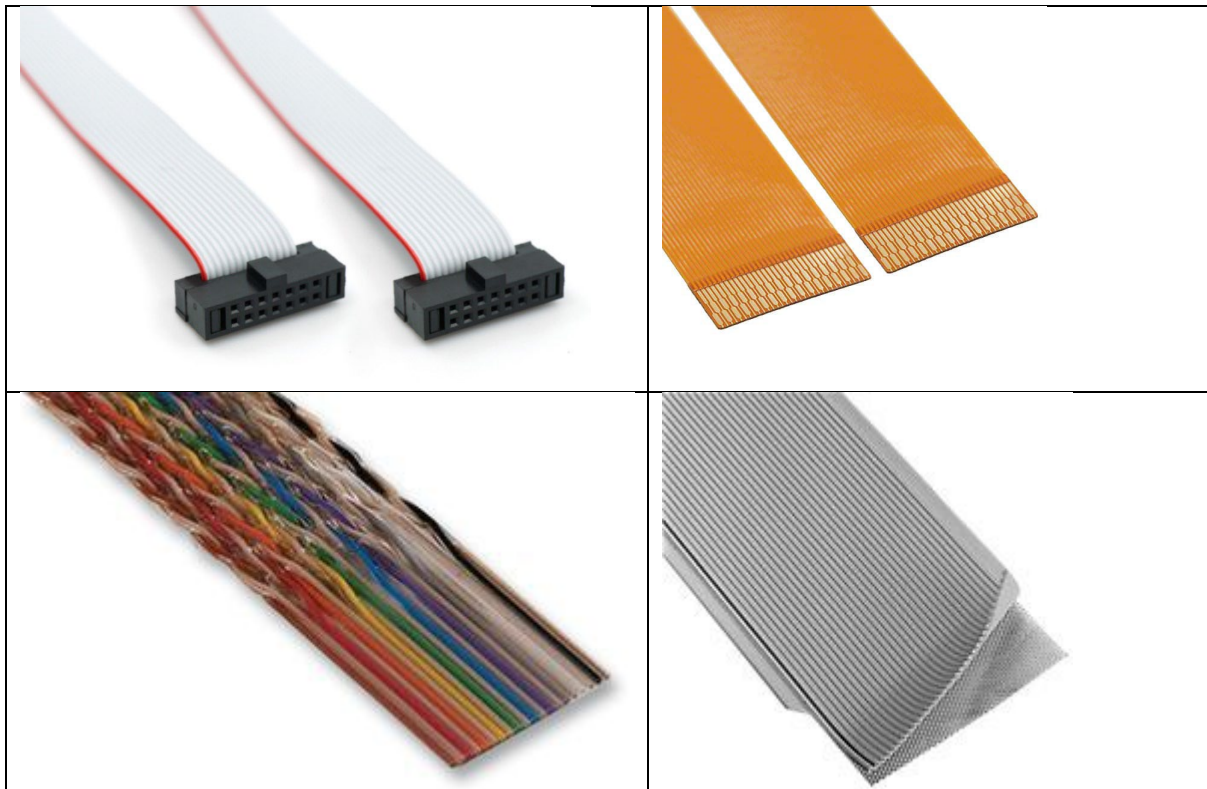


Figura 14.13. Cuatro tipos de cables planos.

Hace tres años, diseñando una interconexión de 4 metros para leer 64 fotodiodos, de los que teníamos que llevar hasta la electrónica de lectura tanto el ánodo como el cátodo de cada sensor, decidimos utilizar cuatro cables planos de 51 hilos. El patrón en cada cable plano era “masa-ánodo-cátodo” repetido 16 veces, más una línea adicional de masa. Los dos hilos restantes se usaban para excitar un LED para probar los fotodiodos.

Ánodo y cátodo de cada fotodiodo se comportan como un par diferencial, lo que reduce el ruido acoplado en pares vecinos. El hilo de masa entre señales de distintos fotodiodos contribuye a reducir aún más la interferencia entre fotodiodos. Como descubrimos en las pruebas, habíamos hecho un buen trabajo controlando el *crosstalk* en el cable.

El problema fue otro, y resultó toda una sorpresa [34]: extendimos junto a una ventana un cable plano de 4 metros, con fotodiodos en un extremo y la electrónica de lectura (un integrador de la corriente recibida en cada microsegundo) en el otro. Sin señal que excitara los fotodiodos, llegábamos a leer hasta 500 mV_{pp} a la salida del cable, sin amplificar la señal, y hasta 8 V_{pp} a la salida de la electrónica.

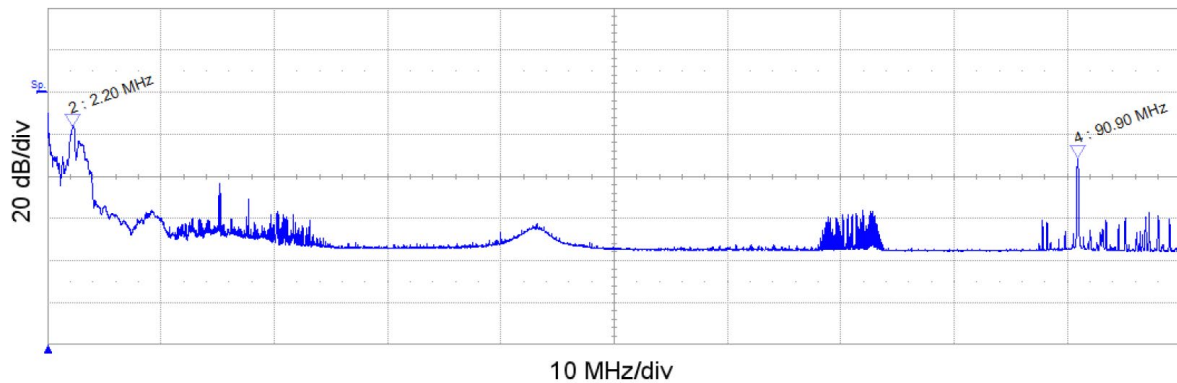


Figura 14.14. Espectro del ruido captado en un cable plano de 4 metros de longitud, extendido junto a una ventana. Fuente: [34]

Usando un analizador de espectros pudimos averiguar el origen de esta interferencia. La señal captada por el cable se podía clasificar en tres bandas de frecuencia:

- 2-3 MHz: Ráfagas cortas y periódicas de 2 a 3 ms de duración, separadas varios microsegundos entre sí. Tocando con la punta de la sonda en una conexión a tierra de la red eléctrica observábamos esta misma interferencia. Por lo tanto, era originada por algún equipo (o equipos) en otra parte del edificio, y era propagada (y radiada) por la instalación eléctrica del edificio.
- 9-75 MHz: Ruido de baja amplitud debido a emisiones de TV, HF y VHF.
- 90 MHz: Un pico de gran amplitud procedente de una emisora de FM a menos de un kilómetro de distancia.

De estas tres bandas, la primera era la más relevante para nuestro problema, ya que:

1. El ancho de banda de la etapa de entrada de nuestra electrónica era de 20 MHz
2. La electrónica tenía un integrador (tiempo de integración de 1 μ s), de modo que las altas frecuencias quedaban fuertemente atenuadas

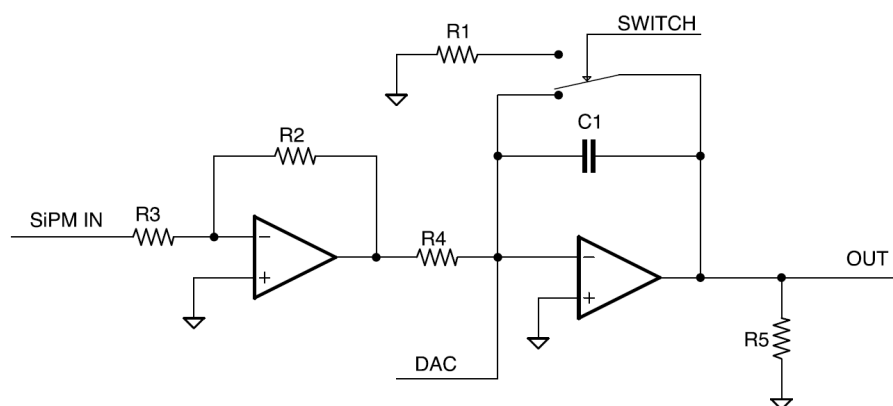


Figura 14.15. Primera versión de la electrónica de lectura, basada en un *switched (gated) integrator*. Fuente: [34]

Esta primera versión de la electrónica sólo amplificaba e integraba la señal del ánodo de cada fotodiodo. Como el ruido se acopla en el cable plano principalmente en modo común, decidimos transformar la etapa de entrada no diferencial en pseudo-diferencial (Figura 14.16), consiguiendo una importante reducción del ruido. Esto lo hicimos en dos iteraciones. En la segunda añadimos R28 para mejorar la simetría de las líneas. Aunque el CMRR final no era muy elevado (tan sólo 37 dB) conseguimos un buen factor de atenuación.

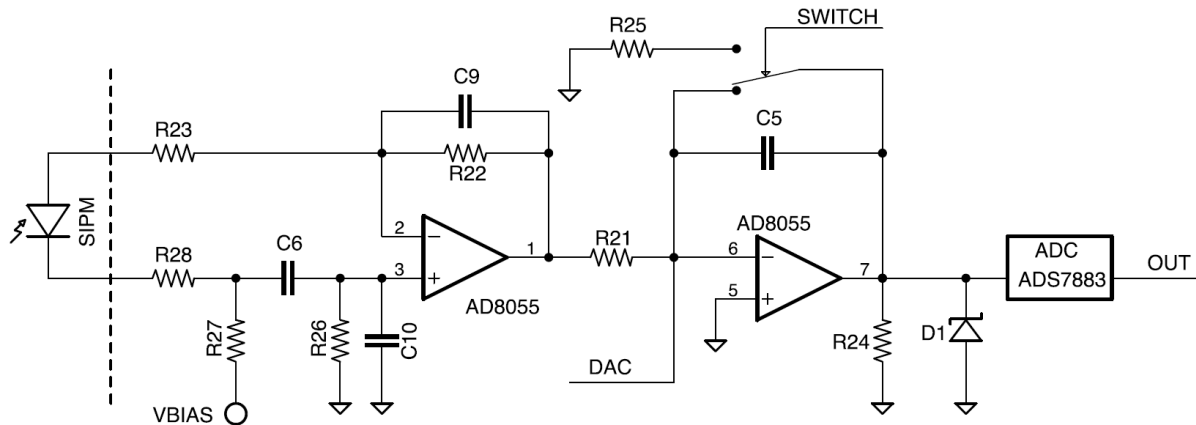


Figura 14.16. Etapa de entrada modificada para reducir el ruido acoplado en modo común en el cable plano. Fuente: [34]

La segunda mejora consistió en rodear cada grupo de 4 cables planos con una malla metálica en forma de tubo con una luz de 1 mm. Esta pantalla presentaba una atenuación máxima al ruido a 1 MHz según el fabricante, justo en la banda que más nos interesaba atenuar. Sin conectar la pantalla a masa en ninguno de los dos extremos, conseguimos una reducción moderada de ruido.

Con el fin de evitar ruido de baja frecuencia debido a bucles de masa, y sospechando que el principal mecanismo de acoplamiento era capacitivo, conectamos la pantalla a masa sólo en el lado de la electrónica. De este modo, conseguimos una reducción de ruido adicional, dejándolo en tan sólo 12 mV_{pp}, justo por debajo de la amplitud correspondiente a la detección de un fotón.

De esta experiencia podemos extraer varias conclusiones importantes:

- Un cable plano y largo es un fantástico receptor de interferencias (antena).
- Un problema de interferencias rara vez se soluciona modificando un único aspecto del sistema. Cuando obtenemos mejoras mediante tres o cuatro modificaciones, sus efectos suelen ser multiplicativos, lo que permite alcanzar mejoras que en un principio nos podían parecer imposibles.
- El primer paso para resolver un problema es identificar la causa. En nuestro caso, un acoplamiento de baja frecuencia (ráfagas de 2-3 MHz) desde la red eléctrica del edificio, que identificamos con un analizador de espectros. Buscar el origen de la interferencia hubiera sido complicado y, aun así, lo normal es que el “dueño” del equipo no estuviera entusiasmado por tener que añadir filtros o tomar otras medidas. Hacer nuestro diseño más robusto era la opción sobre la que teníamos más influencia.
- Las interferencias en un cable plano pueden reducirse mediante tres estrategias principales: apantallar el cable (lo que hicimos con una malla de tubo), reducir el área de los bucles (ánodo, cátodo y masa iban en hilos contiguos) y eliminar el ruido acoplado en modo común (simetrizamos la impedancia de cada hilo y modificamos la primera etapa de la electrónica para hacerla diferencial). Tuvimos que hacer uso de las tres para alcanzar nuestro objetivo (ruido inferior a la señal de un fotoelectrón). Todavía podríamos haber reducido más el ruido usando un cable plano de pares trenzados, lo que hubiera sido necesario si el ruido acoplado hubiera tenido una fuerte componente de acoplamiento inductivo, pero no fue el caso.

Cables de pares trenzados

Cuando el problema es una interferencia acoplada capacitivamente, la mejor solución es apantallar el cable (vimos un buen ejemplo en la sección anterior). Cuando la interferencia se acopla inductivamente, apantallar es poco eficiente (en baja frecuencia, un blindaje es casi transparente para el campo magnético) y hay dos estrategias principales:

- Reducir el área del bucle
- Buscar una cancelación parcial de la energía inducida

Ambas estrategias se consiguen trenzando los dos conductores de la línea, tal y como se muestra en la Figura 14.17. Al trenzar los hilos nos aseguramos de que la separación entre ellos es mínima. Pero lo mismo se podría decir de dos hilos contiguos en un cable plano. Comprenderemos la ventaja de trenzar los hilos si consideramos un campo magnético variable que se propaga en dirección perpendicular a la pantalla (o papel) y la (lo) atraviesa. La polaridad de la tensión inducida es opuesta cuando el cable azul está sobre el rojo que cuando es el rojo el que está arriba. De este modo se consigue una razonable cancelación de las tensiones inducidas por acoplamiento inductivo. Para reducir el capacitivo, ya conocemos algunas estrategias a añadir a este cable trenzado (apantallando y/o simetrizando los hilos y el receptor, que será diferencial).

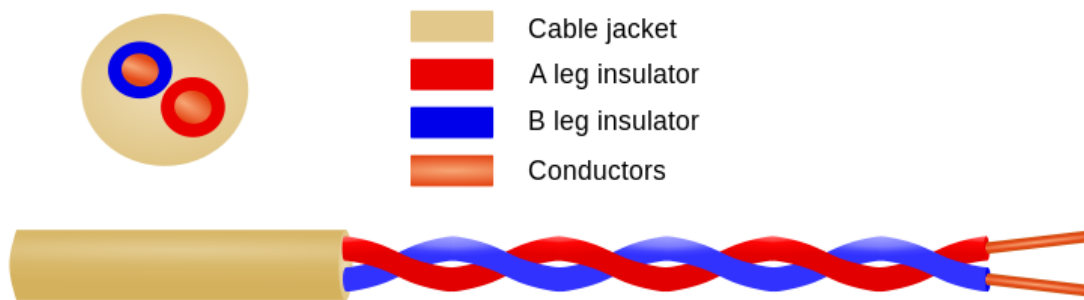


Figura 14.17. Par trenzado sin pantalla. [Imagen](#) con licencia CC-SA 3.0

Por ejemplo, un cable de red como el de la Figura 14.18 (*fully shielded twisted pair*) lleva una hoja de aluminio rodeando cada par, así como una pantalla adicional rodeando el conjunto de los cuatro pares.



Figura 14.18. Cable con cuatro pares trenzados y apantallados individualmente. [Imagen](#) con licencia CC-BY-SA 3.0.

Un artículo de 1979 con el sugerente título de “*Unscrambling the misteries about twisted wire*” [35], intenta responder entre otras a estas dos preguntas:

1. ¿Cuánto reducimos la interferencia acoplada al trenzar los hilos?
2. ¿Cuántas vueltas por metro debe tener el trenzado?

En un estilo claro y directo (nada que ver con los actuales artículos de investigación, en los que se busca hacerlo todo tan críptico como sea posible para vestir una aportación poco relevante con una aparente complejidad), el autor concluye mediante experimentación que no es relevante el número de vueltas por metro (prueba con 13, 26 y 52 vueltas por metro obteniendo resultados similares). La interferencia acoplada se reduce en más de 50 dB hasta 40 kHz si lo comparamos con un par de hilos paralelos separados 0,9 mm. A partir de esta frecuencia, la reducción se degrada 20 dB/década hasta 4 MHz. Por encima de 4 MHz, no hay ventaja apreciable en trenzar los hilos.

Este experimento se hace en campo cercano, con el par trenzado o par de hilos paralelos a 2,54 cm (una pulgada) de distancia de la fuente de campo magnético. Es decir, las conclusiones sólo serían válidas para estudiar el acoplamiento entre cables dentro de un mismo mazo, no para entender qué ocurre ante radiación electromagnética en campo lejano.

No obstante, lo que ocurre si estudiamos un escenario a frecuencias más altas y con víctima y agresor separados una mayor distancia, no cambia:

Un par trenzado es mucho más efectivo que dos hilos paralelos y juntos para reducir emisión y susceptibilidad a campos magnéticos hasta varis decenas de kHz, reduciéndose la diferencia a partir de ese punto y desapareciendo la ventaja a partir de aproximadamente 1 MHz.

Cables coaxiales

Este tipo de cable consta de dos conductores concéntricos separados por un material dieléctrico en el que el exterior hace a la vez de camino de retorno para la señal y de pantalla. El conductor interior suele ser sólido, mientras que la pantalla está formada por varias capas de malla o de hoja metálica. El conjunto va recubierto por una capa aislante y protectora.

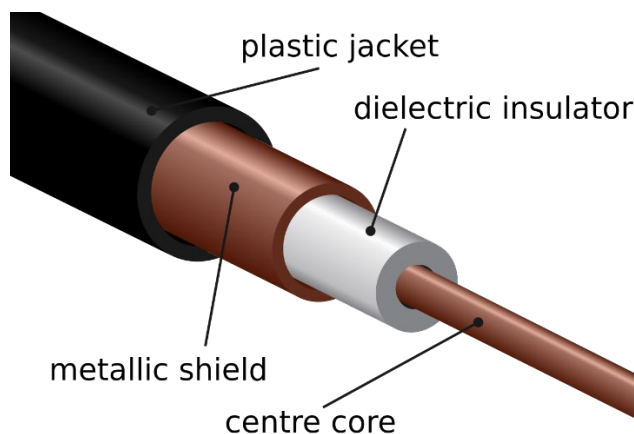


Figura 14.19. Construcción de un cable coaxial (fuente: [Wikipedia](#), imagen con licencia CC-BY 3.0)

Aunque lo cierto es que no siempre la pantalla será el único camino de retorno y no siempre actuará como pantalla, vamos a considerar que es así (caso ideal) y a partir de ese escenario iremos describiendo las no idealidades y limitaciones del cable coaxial.

Cable coaxial en un escenario ideal

Conductor interno y pantalla son los caminos de ida y de retorno de la señal. En consecuencia, los campos eléctrico y magnético quedan contenidos dentro de la estructura y no hay radiación al exterior. **Suena bien, ¿verdad?**

Si la pantalla está conectada a un conductor de área elevada (léase masa o tierra) que pueda mantener equipotencial la pantalla, un **campo eléctrico** exterior no podrá afectar al conductor interior, que sólo “ve” en su entorno un conductor equipotencial que lo rodea. La equipotencialidad sólo será posible en cables eléctricamente cortos a la frecuencia de interés.

Un **campo magnético** externo inducirá una corriente (vale, induce una diferencia de potencial que da lugar a una corriente) en ambos conductores. Pero si la pantalla está conectada a masa/tierra por ambos extremos, puede haber circulación de corriente por la pantalla, que inducirá a su vez (hay un coeficiente de acoplamiento mutuo) una corriente en el conductor interno, de signo opuesto al que induce el campo externo. E, idealmente, se anulan a partir de unos pocos kilohercios en adelante.

Fin del escenario ideal. Que dista mucho de ser el real. Pero es un bonito sueño.

Cable coaxial eléctricamente largo

Conforme subimos en frecuencia pasan dos cosas:

- Respecto a ese campo externo, entraremos en zona de campo lejano (recuerda que la frontera estaba en torno a $\lambda/2\pi$). Así que hay que considerarlo como un campo electromagnético.
- El cable se comporta como una línea de transmisión y deja de ser equipotencial. De hecho, la pantalla se comporta como una antena.

Si la conecto la pantalla a masa/tierra sólo por un extremo, el efecto de antena aumenta hasta que su longitud es de $\lambda/4$. A partir de esta longitud se produce un mínimo para una longitud de $\lambda/2$, otro máximo en $3\lambda/4$, y así sucesivamente.

Este es uno de los motivos por los que en alta frecuencia es recomendable conectar la pantalla por ambos extremos: reducimos la longitud eléctrica de la antena por la mitad, pero esto no evita la radiación, sólo doble la frecuencia a la que estamos en el escenario de $\lambda/4$.

Un ejemplo: a 30 MHz, la longitud de onda es de 10 metros. Un cable más largo de un metro es claramente una línea de transmisión.

Conclusión: sólo cuando el cable es eléctricamente corto, el cable coaxial es (razonablemente) inmune a campos eléctricos externos.

Cable coaxial y bucles de masa en baja frecuencia

Si la pantalla está conectada a masa/tierra por ambos extremos, ¿qué porcentaje de la corriente de retorno circula por la pantalla y qué porcentaje por masa/tierra? Depende de la impedancia de cada camino y puede crear un ruido de baja frecuencia, tal y como comentaremos en la próxima sección.

Conexión del blindaje en cables apantallados

En electrónica, un sistema de transmisión de señales requiere un medio y un modo de propagación, normalmente un par de cables (*pair*). Este par de cables puede estar trenzado (*twisted pair*) y/o apantallado (*shielded pair*), según el grado deseado de protección frente a interferencias acopladas y radiadas.

La protección frente a interferencias conducidas se llevará a cabo mediante filtrado y una correcta planificación de las conexiones a masa, a fin de minimizar en lo posible la problemática producida por los denominados bucles de masa. En este apartado vamos a estudiar la conexión a masa de los cables apantallados. **La pregunta es: ¿qué extremos del cable es conveniente conectar a masa y por qué?**

Si la fuente es flotante (es decir, no conectada a tierra directamente o a través de una impedancia definida) y el cable es coaxial (un único conductor de señal interno), la pantalla debe utilizarse como retorno de corriente y debe ser conectada a masa en el receptor, tal y como hemos explicado en secciones anteriores y se muestra en la Figura 14.20.

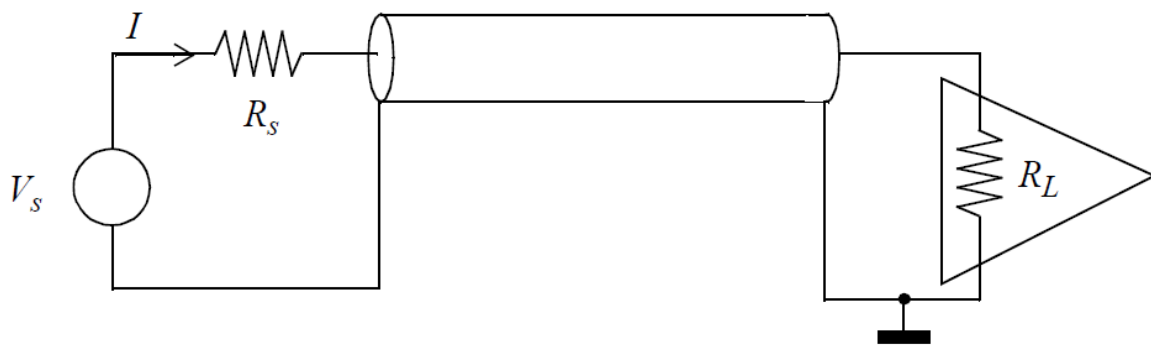


Figura 14.20. Conexión del blindaje cuando la fuente es flotante. Fuente propia

La Figura 14.21 muestra una fuente y una carga conectadas mediante un cable coaxial en el que los dos extremos están conectados a tierra (la fuente ya no es flotante, se dice que está **referida a masa**, aunque lo esté a tierra, a veces la terminología es confusa). Hay dos caminos de retorno para la corriente I : la pantalla del cable coaxial (I_p) y el sistema de masas o tierra (I_t).

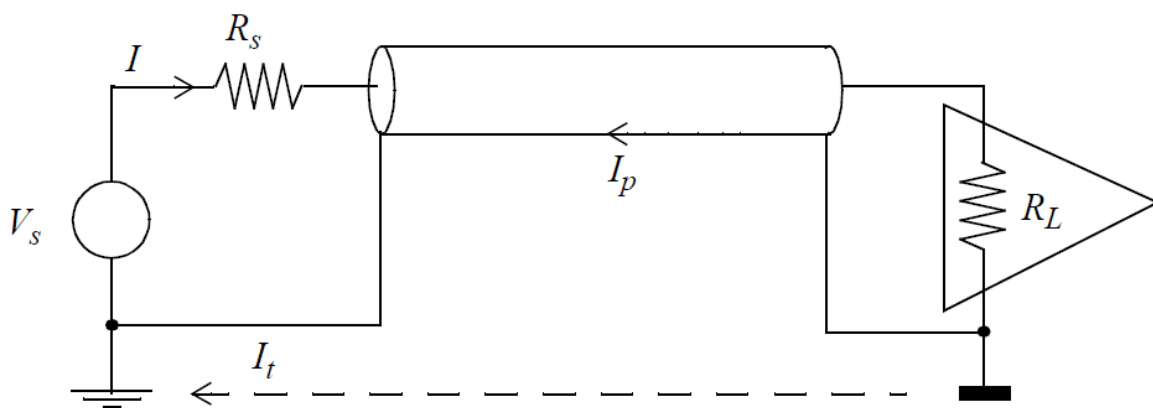


Figura 14.21. Cable coaxial con los dos extremos a masa. Fuente propia

En general, la impedancia del circuito a través del sistema de masas o de tierra no está definida, por lo que a priori no podemos decir si la corriente retornará por la pantalla o por masa/tierra.

A altas frecuencias, la impedancia del camino a través de la pantalla es menor (una menor área del bucle implica una menor inductancia y por lo tanto una menor impedancia) y como consecuencia la corriente

retornará por la pantalla. Como regla práctica, la corriente retornará por la pantalla si la frecuencia de la señal es al menos cinco veces mayor que la **frecuencia de corte del cable** (cociente entre la resistencia y dos pi veces la inductancia del conductor, $R/2\pi L$). Así, para un cable del tipo RG-58C, este límite se sitúa en 10 kHz.

A baja frecuencia, si el cable de conexión no es muy largo, conectar la pantalla a masa en los dos extremos crea un circuito de baja impedancia cerrado por masa/tierra en el que pueden acoplarse campos magnéticos externos, dando lugar a la aparición de señales de modo común indeseadas. Conectar la pantalla por un único extremo abre este bucle y es por lo tanto una mejor conexión (Figura 14.22), siempre que el modo común (diferencia de potencial entre masas) no sea elevado.

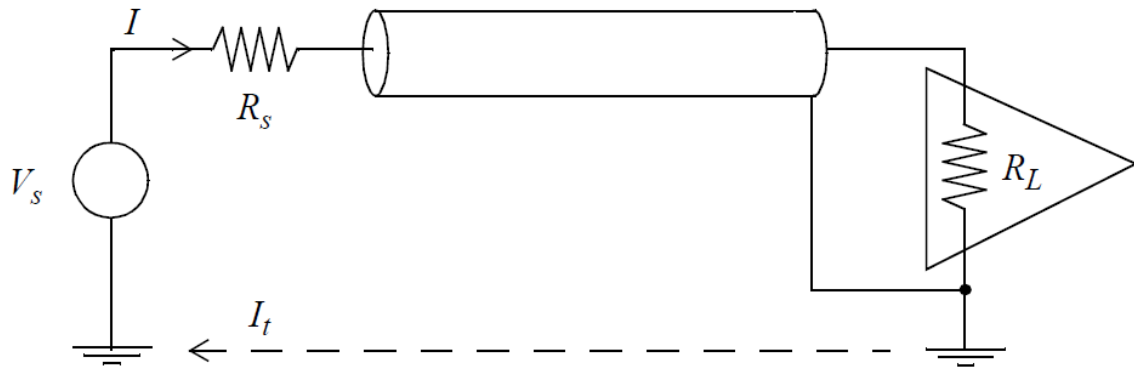


Figura 14.22. Cable coaxial con único extremo a masa. Fuente propia

En alta frecuencia, conectaremos la pantalla a masa por ambos extremos para reducir la radiación de interferencias (por la ley de Ampère) y para garantizar que el blindaje sea equipotencial.

...

En resumen: conectar el blindaje sólo por un extremo es buena idea en baja frecuencia: obtenemos apantallamiento frente a campos eléctricos y evitamos acoplamientos magnéticos de baja frecuencia en modo común. En alta frecuencia solemos optar por conectar el blindaje por ambos extremos: obtenemos apantallamiento frente a campos magnéticos y mejoramos la equipotencialidad del blindaje. Dejar el blindaje sin conectar por ambos extremos no sólo no proporciona protección frente a ningún tipo de campo, sino que puede aumentar el ruido. En cualquier caso, hay que evaluar qué solución (conectar por un extremo o por los dos) es más adecuada en cada aplicación concreta.

Bibliografía

- [1] B. Olney, «Stackup Planning (part 3),» 2015. [En línea]. Available: https://www.icd.com.au/articles/Stackup_Planning_Pt3_PCBD-Aug2015.pdf.
- [2] F. M. a. H. M. J. Toledo, «Past, present and future of data acquisition systems in high energy physics experiments,» *Microprocessors and Microsystems*, p. 27 (2003) 353–358, 2003.
- [3] Cypress Semiconductors, «AN1202 - A Comparison of CML and LVDS for High-Speed Serial Links,» 20 Diciembre 2002. [En línea]. Available: <https://www.cypress.com/file/74011/download>.
- [4] Texas Instruments, «High-Speed Layout Guidelines for Signal Conditioners and USB Hubs,» Texas Instruments, <http://www.ti.com/lit/an/slla414/slla414.pdf>, 2018.
- [5] L. W. Ritchey, «Differential Signal Design, Part 1,» *The PCB Design Magazine*, Aug. 2013, pp. 8-23, 2013.
- [6] B. Olney, «Differential pair Routing,» *The PCB Magazine*, p. 48, 2011.
- [7] Altium, «PCB Routing Angle Myths: 45-degree Angle Versus 90-degree Angle,» 4 Mayo 2018. [En línea]. Available: <https://resources.altium.com/p/pcb-routing-angle-myths-45-degree-angle-versus-90-degree-angle>.
- [8] L. W. Ritchey, «Differential Signal Design, Part 2,» *The PCB Design magazine*, Septiembre 2013.
- [9] C. Paul, Introduction to Electromagnetic Compatibility, 2nd ed., John Wiley & Sons, 2006.
- [10] C. Z. a. D. T. Hubing, «Maximum Radiated Emission Calculator: Differential-Mode EMI Algorithm,» 2013.
- [11] EMRG, «Michigan State University. Electromagnetics Research Group,» [En línea]. Available: https://www.egr.msu.edu/emrg/sites/default/files/content/module9_radiated_emissions.pdf.
- [12] M. I. Montrose, EMC and Printed Circuit Board: Design, Theory, and Layout Made Simple, New York: John Wiley & Sons, 1998.
- [13] M. Leone, «On the Electromagnetic Radiation of Printed-Circuit-Board Interconnections,» *IEEE Transactions on Electromagnetic Compatibility*, pp. VOL. 47, Issue 2, 2005.
- [14] C. Z. a. D. Todd, «Maximum Radiated Emission Calculator: Power Bus EMI Algorithm,» 2013.
- [15] M. Montrose, «Analysis on the Effectiveness of the 20-H Rule for Printed-Circuit-Board Layout to Reduce Edge-Radiated Coupling,» *IEEE Transactions on Electromagnetic Compatibility*, pp. vol.47, no.2, p.227, 2005.
- [16] X. Ye, «EMI Mitigation With Multilayer Power-Bus Stacks and via Stitching of Reference Planes,» *IEEE Transactions on Electromagnetic Compatibility*, pp. vol.43, no.4, p.538, 2001.
- [17] G. Haussmann, «Impact on Radiated Emissions of Printed Circuit Board Stitching,» de *IEEE International Symposium on Electromagnetic Compatibilit*, 1999.
- [18] B. Olney, «Stackup Planning, Part 1,» Junio 2015. [En línea]. Available: https://www.icd.com.au/articles/Stackup_Planning_Pt1_PCBD-June2015.pdf.
- [19] B. Olney, «Stackup Planning, Part 2,» Julio 2015. [En línea]. Available: https://www.icd.com.au/articles/Stackup_Planning_Pt2_PCBD-July2015.pdf.

- [20] B. Olney, «The Perfect Stackup (for high speed designs),» Noviembre 2011. [En línea]. Available: https://www.icd.com.au/articles/Perfect_Stackup_PCB-Nov2011.pdf.
- [21] G. Li, «Optimization design method of decoupling capacitor in PCB hardware of electric vehicle controller,» *Energy Procedia. The 8th International Conference on Applied Energy – ICAE2016*, vol. 105, p. 3201 – 3206, 2017.
- [22] Avago Technologies, «Linear Models for Diode Surface Mount Packages - Application Note 1124».
- [23] Texas Instruments, «Application Report SCBA017D - Quad Flatpack No-Lead Logic Packages,» 2004.
- [24] K. Wyatt, «An EMC troubleshooting kit,» 1 junio 2012. [En línea]. Available: <https://www.edn.com/an-emc-troubleshooting-kit-part-1a-emissions/>.
- [25] Ministerio de Energía, Turismo y Agenda Digital, «DIRECTIVA RED 2014/53/EU TRANSPUESTA A LA LEGISLACIÓN ESPAÑOLA POR REAL DECRETO 188/2016,» [En línea]. Available: <https://avancedigital.gob.es/equipos-telecomunicacion/Documents/Directiva-RED.pdf>. [Último acceso: Julio 2020].
- [26] U. P. d. M. A. I. Luca Buccolini, «Simple SPICE ESD Generator Circuit based on IEC61000-4-2 Standard,» [En línea]. Available: <https://www.youspice.com/simple-spice-esd-generator-circuit-based-on-iec61000-4-2-standard/>. [Último acceso: 10 mayo 2020].
- [27] D. Q. Zhu Wu, «The Device Parameters Simulation of Electrical Fast Transient Generator,» de *ICSGCE*, Chengdu, China, 2011.
- [28] A. K. V. Wallash, «Measurement, Simulation and Reduction of EOS Damage by Electrical Fast Transients on AC Power,» de *Electrical Overstress/Electrostatic Discharge Symposium*, Reno, NV, USA, 2010.
- [29] Texas Instruments, «SLYY152 - Demystifying surge protection,» 2018.
- [30] D. a. H. B. Powell, «Introduction to Voltage Surge Immunity Testing,» IEEE Power Electronics Society, Denver Chapter Meeting, Denver, September 2007.
- [31] Analog Devices, «CN0313 - EMC Compliant RS-485 Transceiver Protection Circuits,» [En línea]. Available: <https://www.analog.com/en/design-center/reference-designs/circuits-from-the-lab/cn0313.html>. [Último acceso: 2020].
- [32] J. Weibler, «Properties for metals used for RF shielding,» *EMC Test&Design*, 1993.
- [33] Acorn Computers Ltd., «EMC Design Guidelines,» 1993. [En línea]. Available: http://chrisacorns.computinghistory.org.uk/docs/Acorn/DN/Acorn_Dev_EMCDesignGuidelines.pdf. [Último acceso: julio 2020].
- [34] J. R. Samaniego, «Study and design of the front-end and readout electronics for the tracking plane in the NEXT experiment,» 2017. [En línea]. Available: <https://riunet.upv.es/handle/10251/86285>.
- [35] R. B. Cowdell, «Unscrambling the misteries about twisted wire,» de *IEEE International Symposium on Electromagnetic Compatibility*, San Diego, 1979.
- [36] C. C. C. Keith Armstrong, «Design Techniques for EMC – Part 6,» Cherry Clough Consultants, [En línea]. Available: http://www.compliance-club.com/archive/old_archive/991215.htm#_Toc7417944. [Último acceso: 10 mayo 2020].
- [37] S. Marum, «Protecting Circuits from the Transient Voltage Suppressor's Residual Pulse during IEC 61000-4-2 Stress,» de *EOS/ESD Symposium*, 2009.