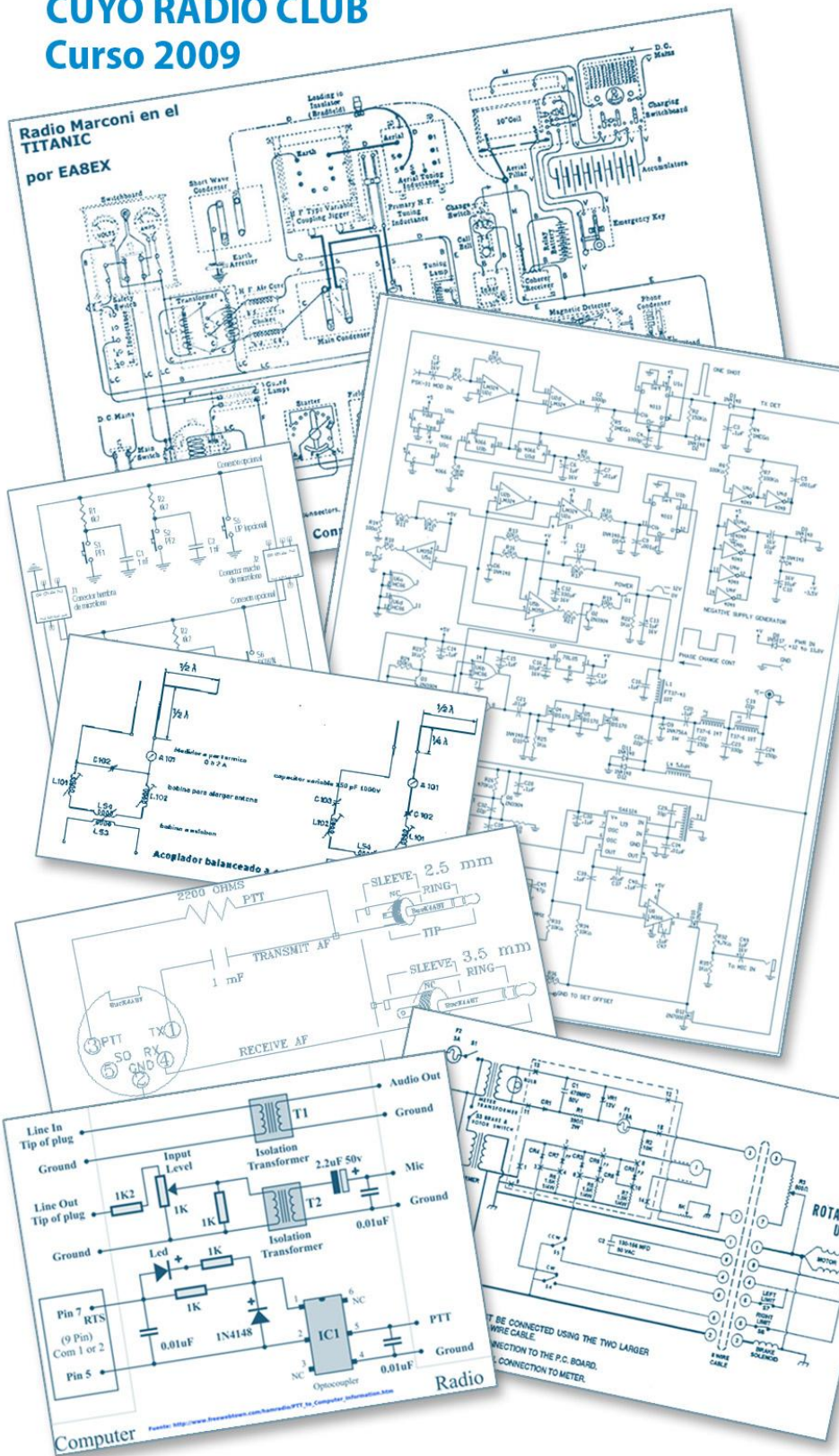


CUYO RADIO CLUB Curso 2009

TÉCNICA



LU1MA

TÉCNICA

INTRODUCCIÓN

Estos apuntes pretenden ayudar al aspirante a radioaficionado a dar sus primeros pasos en el camino de la radiotécnica. Están orientados a aquellas personas que sus conocimientos técnicos son muy escasos y por ello, se tratará de explicar los fenómenos radioeléctricos de una manera simple y didáctica.

La materia es muy amplia y para aquellos aspirantes que deseen conocer “**algo más**”, diré que hay literatura muy completa que tratan los temas de manera más profunda.

Al final de cada capítulo se da un listado con algunos libros que se pueden consultar.

OSVALDO D. PERALTA
LU3MAM (L33M)

CUYO
RADIO
CLUB

TEMA: 1

TEORÍA ATÓMICA

Si bien en el banco de preguntas técnicas que hoy existe, no se exige este tema, considero que el aspirante debe tener una clara idea de esta teoría puesto que allí nacen los fenómenos eléctricos, magnéticos y electromagnéticos.

Comenzaremos diciendo que todo lo que nos rodea está formado por pequeñas partículas llamadas **ÁTOMOS**. Para entender esto, supondremos que tomamos un trozo de hierro y lo dividimos sucesivamente hasta obtener una última partícula tal que si la divido deja de ser hierro, es decir se altera su naturaleza. Esa última partícula que no ha perdido sus características originales se llama **ÁTOMO**.

Ahora bien, supongamos que tenemos un grano de sal, cuyo nombre en química es **CLORURO de SODIO** y su fórmula es **CINa** y hacemos la misma experiencia anterior. La última partícula que obtendremos de sal no será un átomo de sal sino que, por tratarse de una sustancia compuesta, será una **MOLÉCULA** de CLORURO de SODIO. Si ahora por algún medio químico separamos esta molécula obtendremos un átomo de CLORO y un átomo de SODIO. Vemos entonces que las moléculas son partículas formadas por dos o más átomos de distintos elementos.

Otro ejemplo es la molécula de agua, que está formada por dos átomos de hidrógeno y uno de oxígeno. Existen en la naturaleza **107 elementos químicos** que conforman este mundo maravilloso de los cuales 90 se encuentran en la naturaleza, mientras que el resto han sido obtenidos en forma artificial en laboratorios. Sus distintas y variadas combinaciones han permitido distinguir alrededor de medio millón de sustancias compuestas. A modo de ejemplo citamos algunos elementos: oxígeno, carbono, mercurio, oro, plata, calcio, yodo, azufre, cloro, etc.

Ya desde la antigüedad se conocía la existencia del átomo. Su nombre proviene del griego y significa "**indivisible**" puesto que se creía que era imposible dividirlo. El sabio **Tales de Mileto** (600 a.c.) sabía que algunos materiales tratados de cierta manera tenían propiedades extrañas que no podían ser explicadas en esos tiempos. Es el caso del ámbar, una resina de color amarillo que se extraía de las coníferas. Si una barra de ámbar se frotaba con una piel de gato, la barra de ámbar luego atraía objetos livianos como por ejemplo cabellos. Se decía que la barra había adquirido "algo" que le daba la propiedad de atraer objetos.

Esta propiedad del ámbar le dio origen a la palabra electricidad ya que ámbar, en griego se dice "Elektron". Se decía que la barra estaba "electrificada". Fue el físico danés **Neils Bohr** quién en 1913 publicó su teoría atómica y es la que consideraremos para nuestro estudio.

En ella decía que todos los cuerpos están formados por átomos. Que los átomos están constituidos por un **NÚCLEO** central en el que reside casi toda la masa del átomo y su carga positiva, rodeado de los **ELECTRONES** que poseen carga negativa y están situados a distancia relativamente grandes del núcleo. Los electrones no son atraídos por el núcleo puesto que giran a gran velocidad alrededor del núcleo, de tal manera que la fuerza centrífuga, producida por éste giro, contrarresta la fuerza de atracción ejercida por el núcleo. Se podría comparar éste modelo atómico con un sistema solar en el que el sol sería el núcleo y los planetas representarían a los electrones. El camino que siguen los electrones alrededor del núcleo se denomina **ÓRBITA**.

Debemos decir también que en el núcleo se hallan dos partículas, una denominada **NEUTRÓN** que no posee carga y otra denominada **PROTÓN** que es la que le da carga positiva al núcleo. Todos los electrones tienen la misma carga negativa; todos los protones tienen exactamente la misma carga positiva. Además cualquier átomo normal, que contiene igual número de protones que de electrones, no presenta carga en exceso. Esto conduce a la conclusión de que las cargas de un protón y un electrón, aunque de signo opuesto, son de igual valor. El átomo más simple que se conoce es el de Hidrógeno que posee un protón y un electrón.

El hierro posee 26 protones y 26 electrones, el cobre posee 29 protones y 29 electrones. Los electrones están agrupados en distintas órbitas, así el átomo de cobre tiene 4 órbitas al igual que el átomo de hierro, lo que varía es la cantidad de electrones que hay en la última órbita.

Electricidad

Los átomos de una sustancia cualquiera, como ya se ha dicho, contiene un número igual de protones y electrones. Por consiguiente, la materia no presenta efectos eléctricos y se dice que es eléctricamente **NEUTRA** o que está descargada. Si alteramos por algún procedimiento el equilibrio que existe entre protones y electrones, esto es, que un cuerpo tenga exceso de electrones o defecto de electrones se dice que dicho cuerpo está **CARGADO**, o dicho más técnicamente: se ha **IONIZADO**. Hay muchas maneras de alterar dicho equilibrio.

El más antiguo es el fenómeno de electrificación por frotamiento. En realidad el nombre adecuado sería electrificación por contacto puesto que lo que se hace es poner en contacto íntimo los cuerpos. Si se frota una barra de ebonita con una piel de gato y se aproxima a una bolita de médula de saco suspendida por un hilo de seda, la bolita es atraída en un principio por la ebonita y luego del contacto es repelida.

Dos bolitas de médula de saco que se han tratado de la misma manera se repelen mutuamente. Por otra parte, cada bolita es atraída por la piel de gato. ¿Qué ha ocurrido? Al frotar la ebonita con la piel de gato hay una transmisión

espontánea de electrones de la piel de gato a la ebonita. Es decir que la ebonita tendrá un exceso de electrones o dicho de otra forma se ha cargado negativamente y la piel de gato ha perdido electrones es decir que ha quedado cargada positivamente.

Al acercar la barra de ebonita con carga negativa a la bolita que tiene todos sus átomos equilibrados, se produce una atracción debido a que los electrones de la bolita son repelidos y quedan sobre la superficie enfrentada con la ebonita los protones de los átomos de la bolita. Por ello, al tratarse de cargas de distinto signo, se produce la atracción. Cuando entran en contacto, algunos electrones en exceso de la barra de ebonita pasan a la bolita y quedan sobre la superficie de la bolita. Ahora hay electrones sobre las dos superficies y se produce la separación. De los experimentos realizados con la ebonita al frotarla con la piel de gato se deduce que la carga adquirida por uno de ellos es exactamente igual a la perdida por el otro, es decir que no hay creación de carga eléctrica sino una transmisión de carga de un cuerpo a otro.

Supongamos ahora que tenemos dos cuerpos, uno ionizado positivamente es decir con defecto de electrones y otro ionizado negativamente es decir con exceso de electrones. Si unimos ambos cuerpos con un alambre de cobre, ¿qué ocurrirá? La respuesta es que el cuerpo cargado positivamente tomará electrones del alambre de cobre hasta quedar neutralizado. Pero el alambre de cobre ha quedado ahora con un defecto de electrones, es decir ha quedado con carga positiva. El alambre tomará electrones del cuerpo cargado negativamente hasta neutralizarse. Vemos que ha habido un movimiento de electrones desde el cuerpo cargado negativamente hacia el cuerpo cargado positivamente. Este movimiento es, como dijimos, exclusivamente de electrones y de allí deriva la palabra **ELECTRICIDAD**.

Otra forma de expresar lo sucedido es diciendo que hubo una **CORRIENTE** de electrones o una **CORRIENTE ELÉCTRICA**. Podemos concluir también que por el hecho que se produzcan movimientos de electrones, éstos desarrollan una fuerza determinada que es la Fuerza electrónica en movimiento y que se la denomina universalmente con el nombre de **FUERZA ELECTROMOTRIZ** y que se debe precisamente a la fuerza desarrollada por los electrones, de manera que cuanta más cantidad de ellos exista en un cuerpo, mayor fuerza electromotriz podrá desarrollar al entregar electrones a otro cuerpo que posea muy pocos.

Resistencia

Imaginemos el mismo caso anterior pero en lugar de unir los cuerpos cargados con un alambre de cobre, utilizamos un hilo de seda o una cinta de goma. Ocurrirá que no habrá circulación de electrones y se dice que estos elementos son malos conductores, **DIELÉCTRICOS** o **AISLADORES**. A los elementos que permiten la circulación fácilmente de electrones se los llama **CONDUCTORES**.

Que un material sea conductor o aislador depende de la cantidad de electrones que posea en la última órbita de sus átomos.

Si en el ejemplo anterior unimos los cuerpos cargados con un alambre de aluminio parecería que no hay diferencia con el alambre de cobre, pero en realidad el restablecimiento de cargas "**tarda más tiempo**". Esto se debe a que la cantidad de electrones que posee en la última órbita el cobre es distinta a la del aluminio.

Ésta es la causa por la que un material es más conductor que otro y a esta característica se la llama **RESISTENCIA ELÉCTRICA**.

Hubo un notable físico alemán llamado **GEORGE SIMON OHM** (1789-1854) que estudió la relación que hay entre la Fuerza Electromotriz (**FEM**), la Corriente eléctrica y la Resistencia eléctrica de los materiales. La ley que rige esta relación se conoce como **LEY DE OHM**.

TEMA: 2

LEY DE OHM

Como habíamos dicho, la relación que existe entre la Corriente Eléctrica, la Fuerza Electromotriz (o también llamada Diferencia de Potencial o Tensión) y la Resistencia Eléctrica, fue hallada por **G.S.OHM** y dicha relación se conoce con el nombre de **LEY DE OHM**.

La demostración de ésta ley la haremos de una forma sencilla y haciendo analogía con fenómenos que son vistos por nosotros diariamente. Vale decir también que esta ley tiene una demostración matemática que puede ser consultada en cualquier libro de Física o electricidad.

Consideremos un tanque con agua que está a cierta altura del nivel del piso y que de él sale una cañería de un diámetro determinado. Mientras más elevado esté el tanque con más fuerza saldrá el chorro de agua por la cañería y viceversa, cuanto más cerca esté del piso el tanque, tendremos menos caudal de agua o dicho de otra manera habrá menor corriente de agua. Hay una relación directa entre la altura del tanque y el caudal de agua. Son directamente proporcionales.

Veamos ahora un circuito eléctrico formado por un generador de corriente continua (o **pila**, o **batería**) que llamaremos **V**, unido por un conductor a una resistencia que llamaremos **R**. Por el circuito, debido a la diferencia de potencial que existe entre los extremos de la pila, circulará una corriente eléctrica desde el polo negativo de la pila hacia el polo positivo, que llamaremos **I**. Haciendo la analogía con el tanque de agua, comparamos la altura del tanque con la diferencia de potencial entre los bornes de la pila.

A mayor diferencia de potencial, tendremos mayor corriente eléctrica en el circuito y a menor diferencia de potencial tendremos menor circulación de corriente en el circuito de nuestro ejemplo. Es decir que: **LA INTENSIDAD DE CORRIENTE ES DIRECTAMENTE PROPORCIONAL A LA DIFERENCIA DE POTENCIAL**. Que sea directamente proporcional significa que si un valor aumenta el otro también aumenta en forma proporcional y si uno disminuye, el otro también disminuye proporcionalmente.

Ahora volvamos al ejemplo del tanque. Dejaremos la altura invariable y, de alguna manera, imaginemos que estrangulamos un poco la cañería de agua. ¿Qué ocurre con la corriente de agua? Disminuye. Y si lo estrangulamos aún más, menor es la corriente de agua que sale por el extremo del caño.

Lo que hacemos al estrangular la cañería, es impedir que circule agua o dicho de otra manera le ofrecemos resistencia al paso del agua. Podemos decir entonces que a mayor resistencia tenemos menor corriente de agua. Son inversamente proporcionales. Traslademos este fenómeno a nuestro circuito. Si aumentamos el valor de la resistencia eléctrica, disminuirá la intensidad de corriente eléctrica que circula por el circuito. Si disminuimos la resistencia eléctrica, la intensidad de corriente aumentará. Decimos entonces que: **LA INTENSIDAD DE CORRIENTE ELÉCTRICA ES INVERSAMENTE PROPORCIONAL A LA RESISTENCIA ELÉCTRICA DEL CIRCUITO**.

Las conclusiones halladas anteriormente se pueden expresar matemáticamente por la siguiente fórmula:

$$I = \frac{V}{R}$$

Que es la expresión de la **LEY de OHM** y dice que:

LA INTENSIDAD DE UNA CORRIENTE ELÉCTRICA ES DIRECTAMENTE PROPORCIONAL A LA DIFERENCIA DE POTENCIAL ENTRE LOS EXTREMOS DEL CONDUCTOR E INVERSAMENTE PROPORCIONAL A SU RESISTENCIA ELÉCTRICA.

Unidades

Así como la distancia se mide en metros, la masa en gramos y el tiempo en segundos, hay un conjunto de unidades que identifican a la diferencia de potencial, la intensidad de corriente y la resistencia eléctrica.

DIFERENCIA DE POTENCIAL: también llamada fuerza electromotriz o tensión, se la representa generalmente con la letra **V** y su unidad de medida es el **VOLT** ó **VOLTIO** en honor al físico italiano **ALESSANDRO VOLTA** que fue el inventor de la pila voltaica.

INTENSIDAD DE CORRIENTE: se la representa con letra **I**, siendo su unidad de medida el **AMPERE** ó **AMPERIO**, en homenaje a **ANDRÉ MARIE AMPERE**, notable físico francés que introdujo muchos de los conceptos de electricidad y magnetismo.

RESISTENCIA ELÉCTRICA: normalmente representada por la letra **R** y se mide en **OHM** en honor a **G. S. OHM** como hemos dicho anteriormente. A esta unidad se la suele simplificar usando la letra griega omega (Ω).

Múltiplos y Submúltiplos

Cuando hablamos de grandes distancias, generalmente no usamos el metro, sino que las expresamos en kilómetros, al igual que si hablamos de un cuerpo muy pesado lo hacemos en kilogramos o en toneladas. Si la distancia es pequeña hablamos de centímetros o milímetros. Todo esto se hace por una cuestión de práctica y conveniencia. Lo mismo ocurre con las unidades de tensión, intensidad de corriente y resistencia eléctrica.

A continuación presentaremos una tabla con los múltiplos y submúltiplos más utilizados en la práctica:

AMPERE (A)	MÚLTIPLO	Kilo Ampere	(KA)	10^3 A
	SUBMÚLTIPLO	miliampere	(mA)	10^{-3} A
		Microampere	(μ A)	10^{-6} A
		Nanoampere	(nA)	10^{-9} A
		Picoampere	(pA)	10^{-12} A
VOLT (V)	MÚLTIPLO	KiloVolt	(KV)	10^3 V
	SUBMÚLTIPLO	milivolt	(mV)	10^{-3} V
		Microvolt	(μ V)	10^{-6} V
		Nanovolt	(nV)	10^{-9} V
		Picovolt	(pV)	10^{-12} V
OHM (Ω)	MÚLTIPLO	Kilohm	(k Ω)	10^3 Ω
	SUBMÚLTIPLO	Megaohm	(M Ω)	10^6 Ω
		miliohm	(m Ω)	10^{-3} Ω
		Microhm	(μ Ω)	10^{-6} Ω

Diremos también que la tensión se mide con un instrumento denominado **VOLTÍMETRO** y que se conecta en **PARALELO** con el elemento a medir. La corriente eléctrica se mide con un instrumento denominado **AMPERÍMETRO** y se conecta en **SERIE** con el circuito donde se realizará la medición y la resistencia se mide con un **OHMÍMETRO**. Estos tres instrumentos pueden conseguirse por separado o vienen los tres en un solo instrumento denominado: **MULTÍMETRO**, **TESTER** o **POLÍMETRO**.

Características físicas de una resistencia

En un circuito electrónico podemos identificar rápidamente las resistencias ya que se tratan de pequeños cilindros de color marrón en donde se han pintado tres o cuatro franjas de color que nos indican su valor en ohms y la tolerancia. Tienen, además, terminales de alambre en sus extremos.

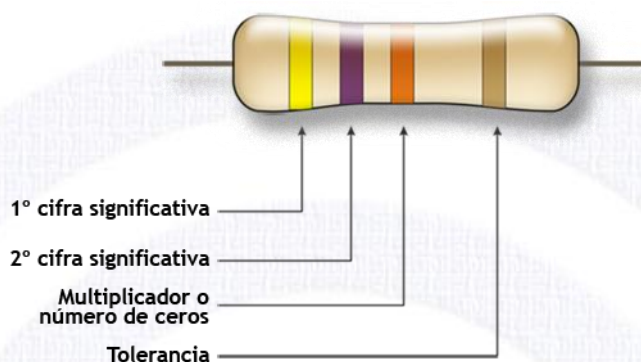
Es un componente imprescindible en la construcción de cualquier equipo electrónico, ya que permite distribuir adecuadamente la tensión y la corriente en todos los puntos necesarios.

Se utilizan materiales con resistividades altas para la fabricación de resistencias de los cuales mencionaremos los más utilizados: **Carbón**, **Alambre de Nicrón** y **Metal depositado**.

En cuanto a la tolerancia de una resistencia diremos que este término aparece como consecuencia de la imposibilidad de obtener un valor óhmico totalmente exacto en la fabricación de la misma. Se hace necesario, entonces establecer los extremos máximos y mínimos entre los que estará comprendida la resistencia. Estos valores, normalmente se expresan como un porcentaje del valor en ohms asignado teóricamente. Existen resistencias de gran precisión en su valor, lo que implica valores de tolerancias muy bajos, pero habrá que tener en cuenta que su precio aumentará considerablemente y sólo serán necesarias en aplicaciones muy especiales. Normalmente las tolerancias estandarizadas son las de $\pm 5\%$, $\pm 10\%$, $\pm 20\%$, aunque esta última está desapareciendo del mercado debido a su poca utilización y a que los procesos de fabricación han mejorado progresivamente.

Código de colores

Para identificar el valor de una resistencia se utiliza un sistema por medio de colores que permite obtener toda la gama de valores comerciales. A este sistema se lo llama código de colores y consiste en pintar alrededor de la resistencia y en un extremo, cuatro anillos de colores determinados, correspondiendo los dos primeros a los números indicativos del valor de la tabla, el tercero es el número de ceros (o multiplicador) que es necesario agregar y el cuarto, la tolerancia.



Tabla

La tabla de los colores se da a continuación:

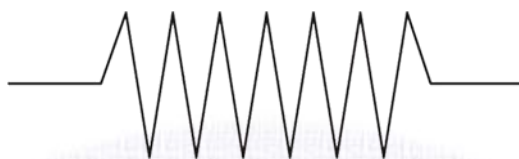
	COLOR	1º CIFRA	2º CIFRA	MULTIPLICADOR	TOLERANCIA
	Negro	0	0	x 1	-
	Marrón	1	1	x 10	±1 %
	Rojo	2	2	x 100	±2 %
	Naranja	3	3	x 1.000	---
	Amarillo	4	4	x 10.000	---
	Verde	5	5	x 100.000	---
	Azul	6	6	x 1.000.000	---
	Violeta	7	7	-----	---
	Gris	8	8	-----	---
	Blanco	9	9	-----	---
	Dorado	--	-	x 0,1	±5 %
	Plateado	-	-	x 0,01	±10 %

Para el dibujo en donde la primera franja es **AMARILLO**, la segunda franja es **VIOLETA** y la tercera es **ROJO**, el valor será de 4 7 0 0 OHMS y es de una tolerancia de ± 5% ya que la cuarta franja es **DORADO**.

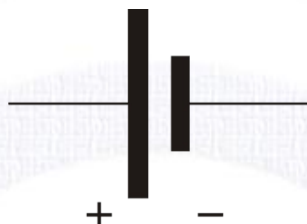
El aspirante se irá familiarizando con la identificación de los valores de resistencia a través de la práctica y llegará a determinar su valor por simple observación.

Símbolo de la resistencia

Como todos los componentes reales de un circuito, las **resistencias** se representan mediante un símbolo. El símbolo para la resistencia eléctrica es:



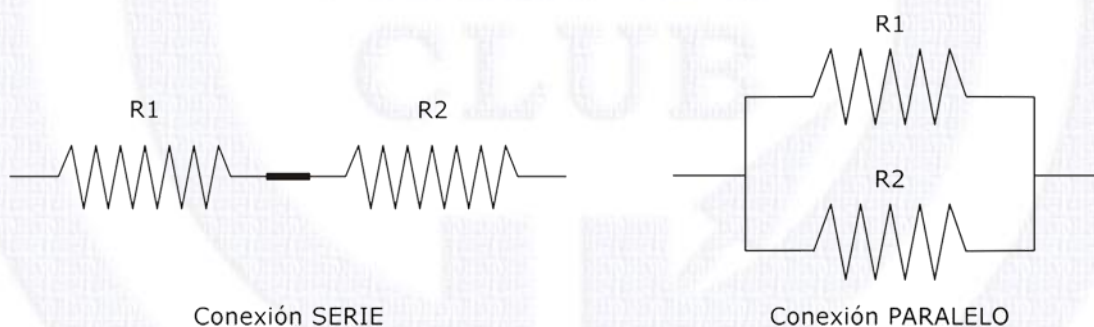
También el **generador de corriente continua, pila o batería** tiene su símbolo, y es:



En el que la barra más larga representa el polo o borne positivo. En la medida que se estudien otros componentes veremos su simbología.

Asociación de resistencias

Si observamos un circuito electrónico veremos que hay más de una resistencia interconectada en el mismo. En la gran mayoría de los casos, estas resistencias eléctricas se encuentran conectadas de una de las siguientes dos formas:



En ambos casos, siempre es posible encontrar un único valor de resistencia que sea equivalente al conjunto de resistencias que forman el circuito.

Esa resistencia se llama precisamente **RESISTENCIA EQUIVALENTE** y se puede demostrar matemáticamente que la resistencia equivalente en un circuito con resistencias en serie es igual a la suma de las resistencias que forman dicho circuito. Expresado en forma matemática diremos que:

$$R_e = R_1 + R_2 + \dots + R_n$$

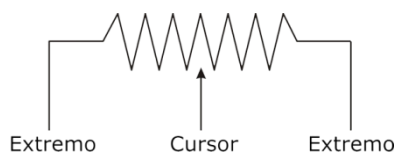
En el caso de tener resistencias en paralelo, la inversa de la resistencia equivalente es igual a la suma de las inversas de las resistencias que forman el circuito. Esto es:

$$\frac{1}{R_e} = \frac{1}{R_1} + \frac{1}{R_2} + \dots + \frac{1}{R_n}$$

El valor $1/R_e$ se le denomina **CONDUCTANCIA ELÉCTRICA**, ya que es la inversa de la resistencia eléctrica y su unidad es el **MHO** o **SIEMENS**.

Con frecuencia es necesario disponer de una resistencia cuyo valor óhmico pueda variarse mediante un control mecánico. Esta función es la que realizan los **POTENCIÓMETROS** (o resistencias variables) y se trata simplemente de

una resistencia que posee un tercer terminal conectado a un punto intermedio entre sus bornes, la posición de este terminal se ajusta por la rotación de un eje. La figura muestra el símbolo eléctrico del potenciómetro.



En un potenciómetro la resistencia medida entre sus extremos no varía con el desplazamiento del cursor y este valor es el dato de adquisición del dispositivo. Como ejemplo diremos que cuando compramos un potenciómetro de 10 kohm, la resistencia medida entre sus extremos será de 10 kohm (10.000 ohms), independientemente de la posición que se encuentre el eje que gobierna el cursor. Pero la resistencia medida entre cualquiera de los extremos y el cursor, se podrá variar a voluntad girando el eje. Obteniéndose valores que van desde 0 (cero) ohm hasta 10.000 ohms.

Potencia eléctrica

Cuando una resistencia eléctrica es atravesada por una corriente eléctrica, la energía eléctrica suministrada a la misma se transforma en energía térmica, es decir se genera calor. Este efecto se utiliza en la construcción de calentadores, estufas, planchas, etc. La energía térmica así generada se denomina **POTENCIA ELÉCTRICA**, se representa normalmente con la letra **P** y su unidad de medida es el **WATT** o **VATIO (W)** y queda determinada numéricamente por el producto entre la corriente eléctrica que atraviesa el elemento a considerar y la diferencia de potencial que aparece entre sus bornes.

$$P = I \times V \text{ (Watt)}$$

Si bien el proceso de calentamiento es utilizable en muchos casos (estufa, etc.) cuando se trata de componentes eléctricos, que han de colocarse en un circuito, una elevación de temperatura resulta perjudicial y hasta puede llegar a destruir al mismo. Mientras más grande sea físicamente el elemento, mayor será la capacidad que tenga para disipar el calor generado y podrá soportar mayores corrientes.

Es importante entonces tener en cuenta qué potencia va a disipar una resistencia porque hay que dar ese dato al comerciante a la hora de adquirir el elemento.

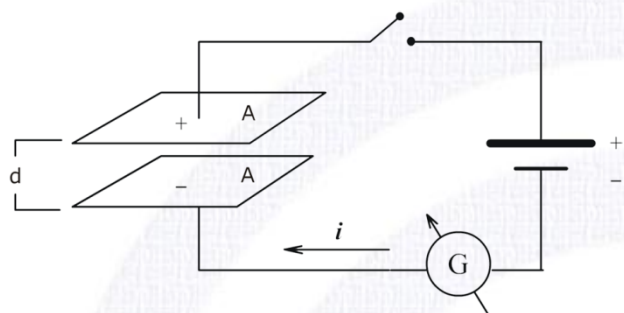
La potencia eléctrica se mide con un instrumento llamado **WATÍMETRO** o **VATÍMETRO**.

TEMA: 3

EL CAPACITOR

Otro de los componentes más utilizados en la construcción de equipos electrónicos, son los condensadores o también llamados **CAPACITORES**. Comenzaremos explicando la filosofía de funcionamiento de estos componentes de una manera sencilla para el aspirante.

Supongamos que tenemos dos placas conductoras de área A , paralelas, enfrentadas y separadas una distancia d . Conectadas a los extremos de un generador de corriente continua, como indica la figura.



Las dos placas no se tocan, están separadas por aire. El aire es un mal conductor de la corriente eléctrica es decir es un **DIELÉCTRICO** o **AISLANTE**.

En el instante de cerrar el interruptor vemos que la aguja del instrumento se ha movido rápidamente y ha vuelto, luego, a su posición de equilibrio. El movimiento de la aguja nos dice que ha habido una circulación de corriente en el circuito.

Pero ¿cómo es esto posible, si el circuito está interrumpido por el aire que hay entre las placas metálicas? Recordemos rápidamente la teoría electrónica que vimos en el **TEMA 1**.

Pues bien, el aire está formado por átomos, que a su vez están formados por protones y electrones. Los electrones se mueven a gran velocidad alrededor del núcleo. Vimos además que la diferencia entre un buen conductor y un mal conductor es la cantidad de electrones que hay en la última órbita. Es decir que un buen conductor tiene más electrones disponibles que un mal conductor. Cuando cerramos el interruptor, los electrones disponibles en la placa metálica, son atraídos por el borne positivo de la batería, y llegan a la otra placa metálica pasando por el instrumento y la batería.



Pero, al llegar aquí no pueden seguir su camino puesto que se encuentran con el circuito abierto (está el dieléctrico) y por ello la corriente eléctrica cesa inmediatamente y el instrumento deja de indicar. La placa conectada al borne positivo de la batería ha quedado con un defecto de electrones es decir con carga positiva y la placa conectada al borne negativo de la batería ha quedado con un exceso de electrones es decir que ha quedado con carga negativa.

Las dos placas cargadas producen un campo eléctrico y el aire está sometido a ese campo eléctrico.

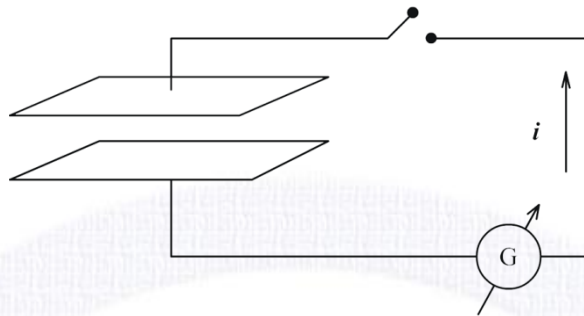
Cuando un cuerpo descargado (conductor o dieléctrico) se coloca dentro de un campo eléctrico como el formado por las dos placas metálicas se produce una redistribución de las cargas del cuerpo.

Estudiaremos especialmente lo que ocurre con un dieléctrico dentro de un campo eléctrico.

Los electrones y los protones de las moléculas del aire que separa las dos placas (dieléctrico), se desplazan por la acción del campo eléctrico formado por las placas metálicas cargadas, y ciertas regiones adquieren un exceso de carga positiva o negativa.

Cuando desconectamos la batería, el campo eléctrico cesa, pero el dieléctrico queda polarizado como indica la figura.

Si ahora conectamos este dispositivo “cargado” en un circuito como el de la figura, veremos que cuando cerramos el interruptor, la aguja del instrumento indica por un instante pasaje de corriente pero en sentido contrario al anterior, es decir cuando estaba la batería conectada.



Podemos concluir que el dispositivo formado por las dos placas metálicas (o conductoras) paralelas enfrentadas y separadas por un material dieléctrico, que en este caso es aire, ha sido capaz de retener carga eléctrica (**electrones**), o dicho de otra manera, es capaz de almacenar energía eléctrica. Tal dispositivo recibe el nombre de **CAPACITOR PLANO** o simplemente **CAPACITOR**.

La carga o cantidad de electricidad **Q** que puede almacenar un capacitor es proporcional a la tensión aplicada **V** y a la **CAPACIDAD** o **CAPACITANCIA C** del capacitor.

$$Q = V C$$

Por lo que podemos decir que la **CAPACIDAD** es:

$$C = \frac{Q}{V}$$

Donde **Q** es la carga almacenada medida en **COULOMB** y **V** es la tensión aplicada medida en **VOLT**.

Como todos los factores que intervienen en la electricidad (resistencia, inductancia, etc.) se miden de acuerdo a unidades de comparación bien establecidas; la capacidad tiene también su unidad de medida que se llama unidad de capacidad o **CAPACITANCIA** y recibe el nombre de **FARADIO** en memoria al físico **MICHAEL FARADAY**, que estudió los fenómenos inherentes a los capacitores.

El Faradio se define de la siguiente manera: se dice que un capacitor tiene una capacidad de un Faradio cuando al aplicar una diferencia de potencial de un Volt, se almacena una carga eléctrica de un Coulomb. El Faradio es una unidad muy grande que en la práctica no tiene mayor aplicación, por esta razón se utilizan submúltiplos del Faradio.

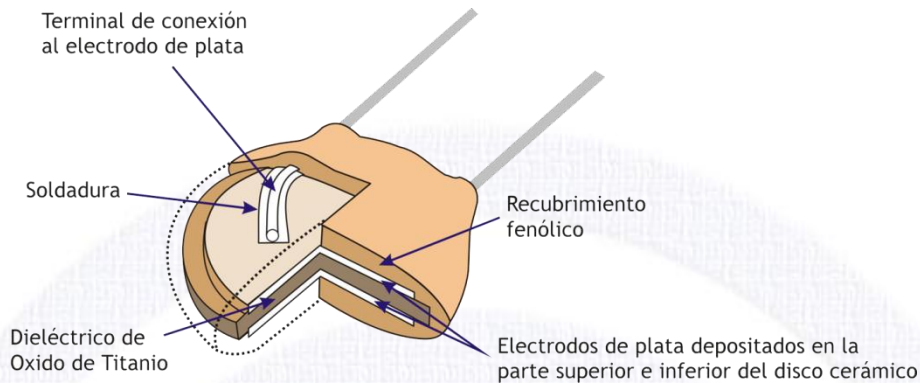
	MICROFARADIOS	(μF)	10^{-6} FARADIOS
SUBMÚLTIPLOS	NANOFARADIOS	(nF)	10^{-9} FARADIOS
	PICOFARADIOS	(pF)	10^{-12} FARADIOS

La capacidad “**C**” de un capacitor o condensador dependerá del tamaño del área “**A**” de las placas enfrentadas, de la distancia “**d**” que separa las placas y del tipo de aislante utilizado en la construcción. Lo dicho podemos expresarlo matemáticamente mediante la siguiente fórmula:

$$C = \frac{K \times \epsilon_0 \times A}{d}$$

Donde ϵ_0 es una constante y **K** es un coeficiente que nos permite comparar la calidad del aislante utilizado en la fabricación del capacitor, con respecto al vacío y se llama coeficiente dieléctrico. Para el vacío **K** es igual a 1 (uno). A continuación se listan algunos materiales usados en la fabricación de capacitores y su correspondiente coeficiente dieléctrico.

VACÍO: 1,0
 POLIESTIRENO: 2,5
 MICA: 6,8
 CERÁMICOS: 6 a 20



De acuerdo con la fórmula, la capacidad de un capacitor será mayor, cuanto mayor sea su área, menor sea su distancia entre placas y mejor la calidad del dieléctrico utilizado. Así, por ejemplo, dos capacitores que tengan igual área e igual distancia entre placas, pero están contruidos con distintos materiales dieléctricos, tendrán capacidades distintas.

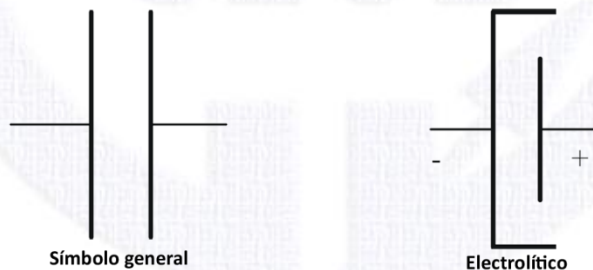


El tipo de dieléctrico utilizado en la construcción de capacitores sirve para clasificarlos. Es por eso que encontramos capacitores de polyester, cerámicos, mica, de papel, etc.

Otro tipo de capacitores que encontramos son los capacitores electrolíticos, cuya característica más sobresaliente es el elevado valor de capacidad que pueden tener en un reducido volumen. Están formados por dos electrodos metálicos sumergidos en una solución conductora o electrolítica. Deben conectarse con la polaridad indicada en el envase para que se forme sobre los electrodos una película aislante que es el dieléctrico.



El símbolo eléctrico de un capacitor es:



Hasta ahora hemos hablado de capacitores fijos de aplicación general, pero a veces es necesario poder variar la capacidad y es por ello que existen los **CAPACITORES VARIABLES**. Ellos permiten variar el valor de la capacidad dentro de ciertos límites.

Están contruidos por dos conjuntos de placas metálicas entrelazadas, uno de los conjuntos es fijo llamado vulgarmente “chapas fijas” y el otro móvil, llamado “chapas móviles” solidario a un eje alrededor del cual pueden girar. Al girar el eje varía el área enfrentada de las chapas y por lo tanto varía la capacidad.

La mayoría de estos tipos de capacitores son de dieléctrico de aire.

El control de sintonía de un receptor de radio, es un capacitor de este tipo.

Otro tipo de capacitor ajustable son los llamados **TRIMMER**, pero se los utiliza en circuitos donde se los ha de ajustar una sola vez para poner a punto el circuito. Son de tamaño más reducido y el dieléctrico generalmente es mica o polyester.

Asociación de capacitores

Los capacitores pueden conectarse, al igual que las resistencias, en **SERIE** y en **PARALELO**. Cuando asociamos dos o más condensadores en paralelo, podemos encontrar el valor de la capacidad equivalente de acuerdo con la siguiente fórmula:

$$C_p = C_1 + C_2 + \dots + C_n$$

Al poner dos o más capacitores en paralelo lo que hacemos es incrementar el área “A” enfrentada y por ello la capacidad aumenta.

Cuando asociamos capacitores en serie la capacidad equivalente se encuentra de acuerdo con:

$$\frac{1}{C_s} = \frac{1}{C_1} + \frac{1}{C_2} + \dots + \frac{1}{C_n}$$

Al poner dos o más capacitores en serie lo que hacemos es incrementar la distancia “d” entre placas y por lo tanto disminuye la capacidad.

La asociación de capacitores es exactamente a la inversa de la asociación de resistencias.

Si tenemos dos capacitores de igual valor conectados en serie, la capacidad equivalente será igual a la mitad del valor de uno de ellos. Si tenemos dos capacitores de igual valor conectados en paralelo, la capacidad equivalente será igual al doble de la capacidad de uno de ellos.

Tensión de ruptura

Cuando se aplica tensión a las placas de un capacitor, se ejerce una fuerza considerable sobre los electrones y protones de los átomos del dieléctrico. Como se trata de un aislador, los electrones no se desprenden de los átomos como lo hacen en los materiales conductores. Sin embargo, si la tensión aplicada es demasiado grande la fuerza sobre los electrones también lo será y el dieléctrico será “roto”, se perforará y permitirá el paso de corriente. Ese valor de tensión se lo conoce como **TENSIÓN DE RUPTURA** y depende del tipo de dieléctrico y del espesor del mismo.

Es un dato que debe tenerse en cuenta y normalmente está indicado por el fabricante sobre el componente. Por ejemplo, si un condensador lleva impreso: 22 μ F - 35 Volt, indica que es un capacitor de 22 μ F y que soporta como máximo una tensión de 35 Volt. Superado este valor, se perforará el dieléctrico y el componente quedará inutilizado. Vulgarmente se dice que se “pincha”.

Aplicaciones

Las aplicaciones de los capacitores son muchas y las iremos viendo a medida que avance el curso, pero ya estamos en condiciones de decir que los capacitores bloquean el paso de una corriente continua entre dos puntos de un circuito.

TEMA: 4

MAGNETISMO

Introducción

El estudio del magnetismo tiene especial importancia en la radiotécnica, puesto que el conocimiento de sus fenómenos nos permite entender las leyes del electromagnetismo.

Todos los equipos de radio tienen que trabajar bajo la influencia de los campos magnéticos. Sin ir más lejos las ondas de radio que se emiten en una antena, son ondas electromagnéticas. Los generadores de energía eléctrica que nos brindan la electricidad para nuestras casas y para la industria tienen su origen en el magnetismo.

Como sucede con la electricidad, no podemos ver el magnetismo, pero podemos estudiar los numerosos fenómenos magnéticos que se presentan en la vida diaria, medirlos y por lo tanto compararlos.

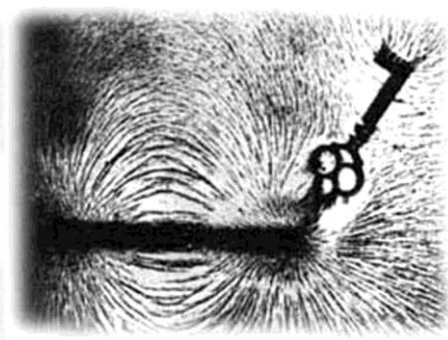
El fenómeno del magnetismo fue conocido por la humanidad en épocas muy remotas, en distintos lugares y en diferentes épocas. Hace muchos siglos el hombre descubrió que suspendiendo de un hilo un trozo de cierta clase de mineral de hierro, de manera que pudiera moverse libremente, éstos señalaban invariablemente la estrella polar (que daba la dirección del polo norte). Este dispositivo (llamado **brújula**) fue usado por los navegantes escandinavos y por los chinos en el año 218 de nuestra era.

Este mineral de hierro es **óxido de hierro** (Fe_3O_4) se lo llama **magnetita** y es el denominado "**IMÁN NATURAL**". El nombre de magnetita se debe a que las mejores muestras de óxido de hierro se encontraron originariamente en la ciudad de Magnesia en el Asia Menor.

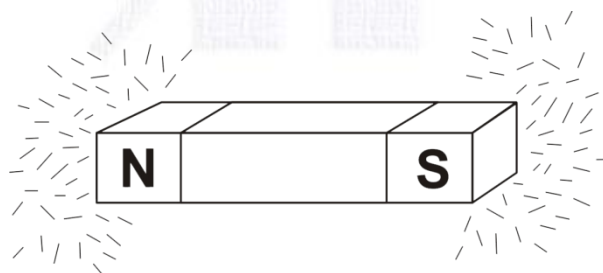
Se denominan imanes naturales porque se los encuentran en la tierra ya imantados. Veremos más adelante que es posible fabricar nuestros propios imanes.

Todos hemos visto la acción de un imán sobre un trozo de hierro. Un clavo, una aguja, cualquier trozo pequeño de hierro son atraídos por un imán cuando lo acercamos a ellos. Como dijimos anteriormente no podemos ver al magnetismo pero sí podemos estudiar los fenómenos que causan.

Tomaremos una barra de imán y le ataremos un hilo de manera que gire libremente. Acercaremos la barra a un recipiente con limaduras de hierro y observaremos que la mayoría de las limaduras han sido atraídas por los extremos de la barra imantada y muy pocas limaduras hay en el centro de la barra. Podemos decir con toda certeza que hay zonas en el imán que atraen a los elementos con mayor o menor facilidad. En la fotografía se puede apreciar este fenómeno.



Dichas zonas se las llama **POLOS** del imán y se les dio, arbitrariamente, **POLO NORTE** y **POLO SUR**. La fuerza de atracción del imán es mayor en los polos.



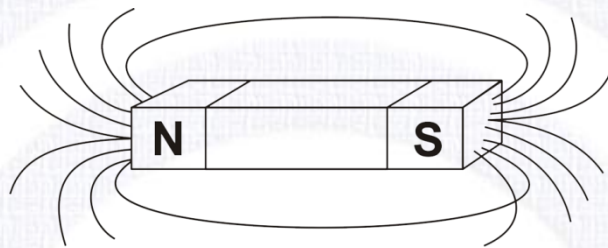
Es también conocido que los polos de igual nombre se repelen y polos de distinto nombre se atraen. Esto se puede comprobar fácilmente con dos barras imantadas enfrentadas.

Campo magnético

Todo el espacio que rodea a la barra de imán se lo denomina **CAMPO** y cualquier elemento de hierro que esté en ese campo estará bajo la influencia del imán. Para especificar aún mejor las cosas se dice que está bajo la influencia de un **CAMPO MAGNÉTICO**.

Al campo magnético no lo vemos, pero para poder estudiarlo lo representamos con líneas que rodean al imán. Estas líneas imaginarias llamadas **LÍNEAS DE CAMPO** están más concentradas donde el campo magnético es más intenso, es decir en la cercanía de los polos y menos concentradas donde el campo es menos intenso.

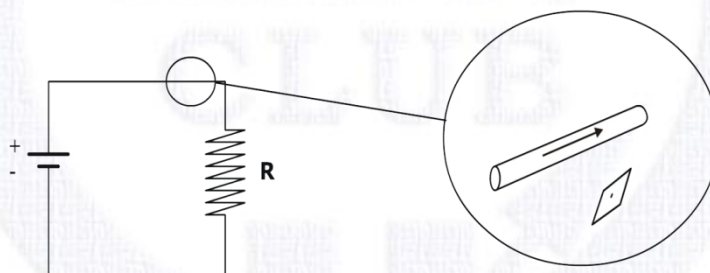
También se les da un sentido que es desde el polo norte al polo sur por afuera del imán y del polo sur al polo norte por dentro del imán, es decir que son líneas cerradas.



Si sobre una barra imantada colocamos un papel o cartulina delgada, y esparcimos limaduras de hierro sobre el papel, veremos lo que se denomina espectro o fantasma magnético. Es decir que las limaduras de hierro se agrupan de acuerdo a la intensidad del campo magnético. Lo que vemos no es el campo magnético sino una imagen de éste.

Ahora que ya se ha dado una introducción sobre el tema, avanzaremos un poco más.

Consideremos un circuito eléctrico, como se indica en la figura:

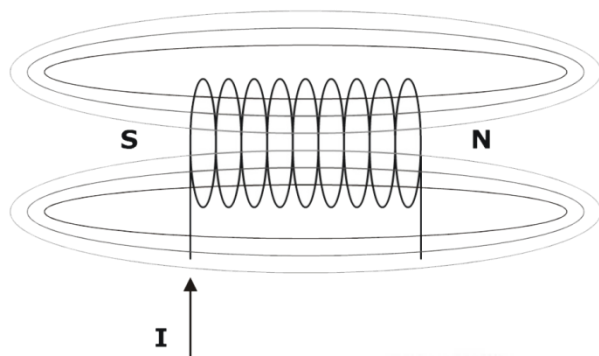


Si a este circuito le acercamos una aguja imantada (o una brújula), veremos que cuando circula corriente la aguja se desvía de su posición, es decir que ocurre un fenómeno de magnetismo.

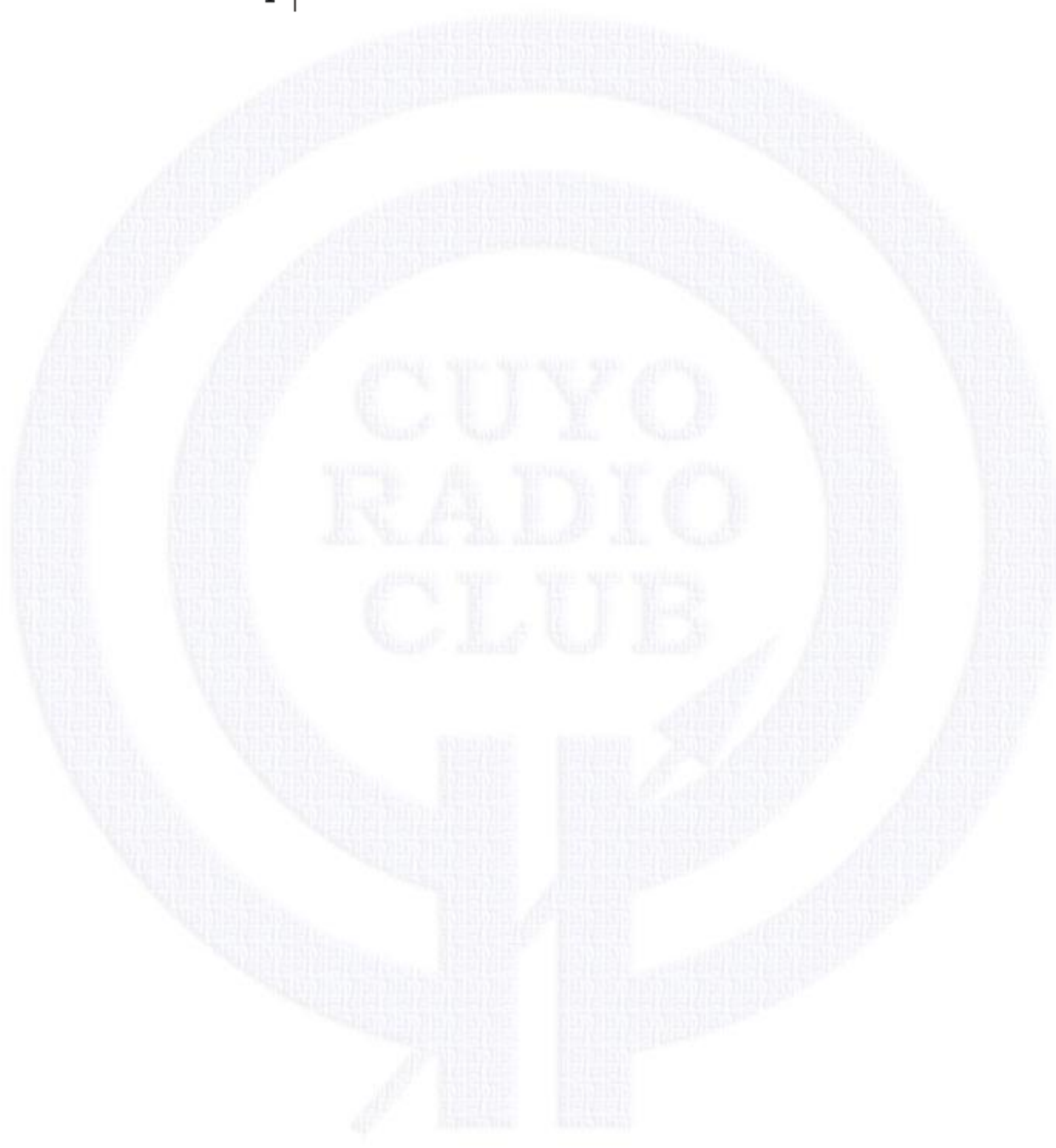
Este experimento fue realizado por el físico **Danés Hans Oersted en 1819** y demostró que cada vez que circula corriente por un conductor, se genera alrededor de mismo un campo magnético. Si pudiéramos colocar alrededor del conductor un papel o cartulina y esparciéramos limaduras de hierro, veríamos que las limaduras se acomodan en anillos concéntricos alrededor del conductor.

Ésta es la distribución que tiene el campo magnético que se forma alrededor de un conductor cuando circula por él una corriente eléctrica.

Ahora bien, si en lugar de tener un conductor rectilíneo, hiciéramos un arrollamiento con el alambre y le hiciéramos circular una corriente, veríamos que el efecto sobre la aguja imantada se acentúa y si colocáramos un papel y esparciéramos limaduras de hierro, el espectro del campo magnético formado alrededor del arrollamiento sería igual al formado por una barra de imán.



F →
 Polos de igual nombre
 se rechazan



Si hiciéramos circular corriente por el arrollamiento en el sentido indicado y acercáramos un imán a un extremo como indica la figura, veríamos que el imán es repelido con una fuerza F .

Lo que me indicaría que el polo formado en el extremo del arrollamiento es norte, pues polos de igual nombre se repelen y de distinto nombre se atraen. Si cambiamos el sentido de circulación de la corriente en el arrollamiento el imán será atraído lo que significa que el polo formado en el extremo del arrollamiento ahora es sur.

El físico Ampere llamó al alambre arrollado SOLENOIDE, también se lo conoce con el nombre de bobina.

Inducción

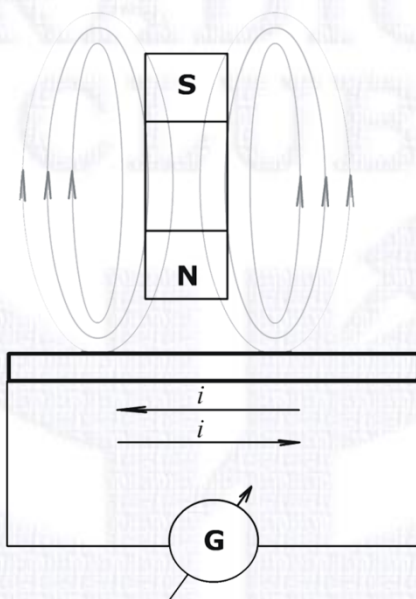
En 1831 el físico inglés Miguel Faraday pensó que el efecto inverso al descubrimiento de Oersted era posible por lo que realizó el siguiente experimento: Colocó un conductor dentro de un campo magnético y de alguna manera hizo variar el campo magnético.

Tomó un imán y lo acercó a un conductor, que tenía conectado en sus extremos un instrumento que detecta el pasaje de corrientes muy pequeñas (**galvanómetro**). Observó que si el imán está quieto, no ocurría nada, pero si movía el imán en dirección perpendicular al conductor, de manera que el conductor cortara a las líneas de campo, el instrumento acusaba pasaje de corriente.

Si lo movía en forma paralela al conductor, tampoco ocurría nada. La explicación de este fenómeno es que los electrones del conductor son desviados por el efecto del campo magnético cuando el imán se acerca al conductor, ese movimiento de electrones no es otra cosa que una corriente eléctrica que es detectada por el instrumento.

Cuando alejo el imán los electrones vuelven a su posición original causando una corriente eléctrica pero en sentido contrario. Podemos notar que se ha conseguido una corriente eléctrica en el conductor sin la intervención de pilas, generadores, etc. Esta corriente se la llama **CORRIENTE INDUCIDA** y a la fuerza eléctrica que da origen a esa corriente se la llama **FUERZA ELECTROMOTRIZ INDUCIDA (F.E.M. inducida)**.

Este importantísimo descubrimiento dio origen a los generadores de corriente, dínamos, alternadores, etc., pues su funcionamiento está basado en el principio que acabamos de explicar.



Si en lugar de usar un conductor rectilíneo usáramos un solenoide observaríamos que la F.E.M. inducida es mayor. Además mientras más vueltas tenga el solenoide o si movemos el imán cada vez más rápido (aumento el número de vaivenes por segundos), o si ponemos un imán más grande (mayor intensidad de campo magnético), más grande será la F.E.M. inducida. Esto lo puedo expresar matemáticamente como:

$$\text{F.E.M. inducida} = N^{\circ} \frac{\Delta\phi}{\Delta t}$$

Ley de lenz

Hasta ahora sabemos que si hacemos circular corriente por un conductor se genera alrededor de él un campo magnético, si la corriente es constante (corriente continua) el campo magnético también lo será, pero si la corriente es variable, es decir que aumenta y disminuye en el tiempo, el campo magnético variará de la misma forma. También sabemos que si movemos un conductor dentro de un campo magnético se induce una corriente en dicho conductor.

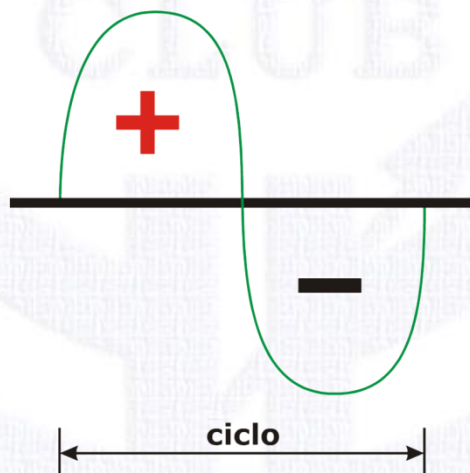
El físico alemán **H.F. LENZ** experimentalmente observó que cuando introducimos un imán dentro del solenoide, se genera una corriente inducida. Pero esa corriente inducida al circular por el solenoide crea un campo magnético y la polaridad de ese campo magnético es tal que se opone a que el imán siga introduciéndose dentro de la bobina. Si ahora trato de sacar el imán, cambia el sentido de circulación de la corriente inducida y esa corriente crea un campo magnético cuya polaridad es tal que se opone a que retiremos el imán.

Esto es lo que expresa la ley de Lenz:

El sentido de una F.E.M. inducida es tal que se opone a la causa que la produce.

Autoinducción

Estudiaremos ahora lo que ocurre cuando el solenoide es atravesado por una corriente que varía su amplitud en el tiempo. Dicha corriente se dice que es alterna ya que toma valores positivos y negativos en forma alternada. Podemos apreciar en la figura que la corriente toma el valor cero, crece hasta llegar a un máximo positivo, luego disminuye hasta llegar a cero nuevamente, continúa disminuyendo hasta llegar a un máximo negativo y luego crece hasta hacerse cero nuevamente. Este ciclo se cumple indefinidamente y por ello se dice que es corriente **ALTERNADA**.



Si se hace pasar una corriente alterna por una bobina de alambre, se forma un campo magnético alrededor de la bobina. Conforme aumenta la corriente (semiciclo positivo) se expande el campo magnético que cruza transversalmente los conductores de la misma bobina, induciendo en ellos una segunda corriente. Esa corriente tiene una dirección tal que se opone a la corriente original (**de acuerdo a la ley de LENZ**). En otras palabras la dirección de la corriente inducida es tal que tenderá a reducir la corriente original y de tal modo tenderá a oponerse a la expansión del campo magnético.

Cuando la corriente original comienza a disminuir el campo magnético que rodea a la bobina comienza a retraerse, igual que el caso anterior, cruza transversalmente las espiras de la bobina y otra vez se induce una segunda corriente en la bobina. La dirección de esta corriente inducida es tal que se opone a la corriente original que está disminuyendo. La corriente inducida tenderá a mantener una circulación de corriente en la bobina durante cierto tiempo, después que la corriente original haya cesado. La corriente inducida, por lo tanto tiende a oponerse a la desaparición del campo magnético.

Lo mismo ocurre en el semiciclo negativo.

Inductancia

La propiedad de un circuito de oponerse a todo cambio de la circulación de corriente que pasa por él, se llama **INDUCTANCIA**, debido a que esta oposición está motivada por tensiones inducidas en el propio circuito por el campo magnético variable, cualquier causa que afecte la cantidad de flujo magnético afectará también a la inductancia. La unidad empleada para medir la inductancia de un circuito se llama **HENRY (Hy)** en honor al físico norteamericano **JOSEPH HENRY**. Se puede definir al Henry como la inductancia que existe en un circuito cuando una variación de corriente de un Ampere por segundo produce una F.E.M. inducida de un Voltio. El símbolo con que se representa la inductancia es **L**. También se dice que es la propiedad de un circuito eléctrico en el cual los cambios en el flujo de corriente producen cambios en el campo magnético. Es considerada como una inercia eléctrica. En radio resulta conveniente emplear el miliHenry (mHy) que es 10^{-3} Hy o el microHenry (μHy) que es 10^{-6} Hy

Inductores en serie y en paralelo

Los inductores, como las resistencias se pueden conectar en serie, en paralelo o en circuitos combinados **serie-paralelo**. La inductancia total de varios inductores conectados en serie (siempre que el campo magnético de un inductor no pueda actuar sobre las espiras del otro), es igual a la suma de las inductancias de cada inductor. La fórmula es:

$$L_{\text{total}} = L_1 + L_2 + L_3 + \dots \text{etc}$$

Si dos o más inductores se conectan en paralelo (siempre que no haya acoplamiento entre sus campos magnéticos) su inductancia total puede hallarse aplicando la siguiente fórmula:

$$\frac{1}{L_{\text{total}}} = \frac{1}{L_1} + \frac{1}{L_2} + \frac{1}{L_3} + \dots \text{etc}$$

Obsérvese que esta relación es similar a la que existe entre resistencias en paralelo.

Sobre estos temas hay sólo una pregunta en el banco (**pregunta N°5**) pero como dije al principio es necesario que el aspirante comprenda muy bien los fenómenos del magnetismo, electromagnetismo y todo lo relacionado con ellos.

TEMA: 5

GENERADORES DE CORRIENTE

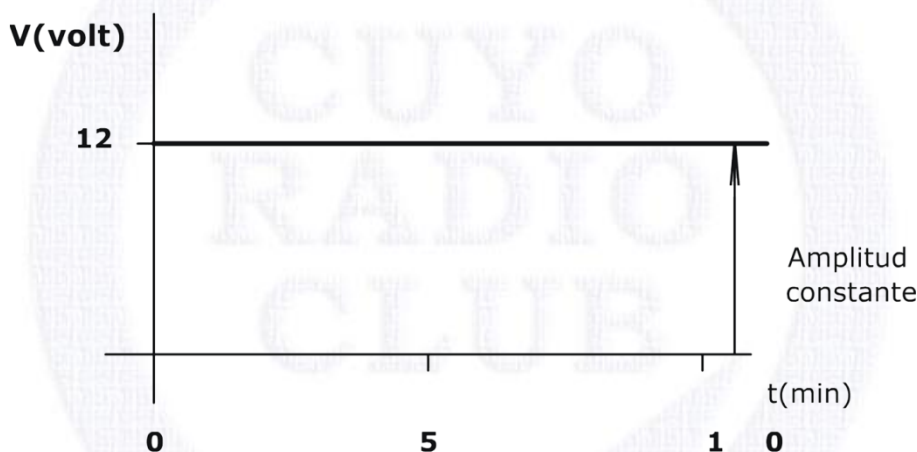
Introducción

Todos los circuitos eléctricos o electrónicos necesitan una fuente de energía eléctrica para su funcionamiento. Estos dispositivos reciben el nombre de **GENERADOR DE Tensión, FUENTE DE ALIMENTACIÓN** o simplemente **FUENTE**. De acuerdo como entregan energía eléctrica a través del tiempo, clasificaremos a los generadores en: generadores de tensión de corriente continua y de corriente alterna.

Generadores de tensión de corriente continua

Las pilas que adquirimos en el comercio, la batería del automóvil, los paneles solares de los satélites, el pack de baterías del handie son generadores de corriente continua.

Esto significa que si la batería es de, por ejemplo de 12 volts, nos entregará ese valor de tensión en forma constante (en realidad se descarga lentamente). Si representamos lo dicho anteriormente en un gráfico, vemos que el valor de tensión se mantiene constante en el tiempo es decir que a los 5 minutos tenemos 12 v. y a los 10 minutos también y así sucesivamente. Decimos que la **AMPLITUD** de la tensión se mantiene constante. La tensión o la corriente que se comporta de ésta manera se llama **CONTINUA**, pues su amplitud continuamente tiene el mismo valor.



Los dispositivos que trabajan con tensiones de corriente continua vienen indicados por el fabricante con las letras **DCV** o **DC** que significan **Direct Current Volt (Volt de Corriente Continua)** y **Direct Current (Corriente Continua)** respectivamente.

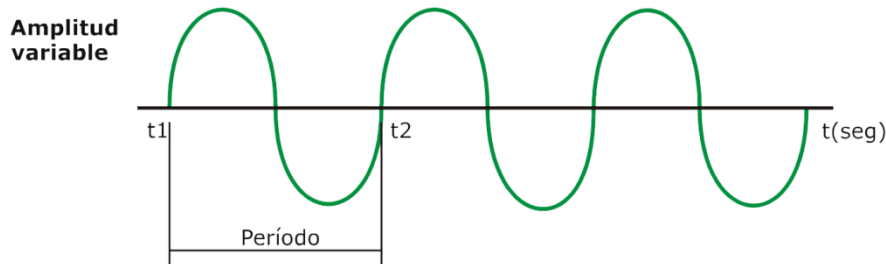
Una de las características más importantes de estos generadores es que poseen polaridad es decir que un borne tiene polaridad positiva y viene indicado con **“+”** y otro borne tiene polaridad negativa y viene indicado con **“-”**. **Nunca debe invertirse la polaridad en los dispositivos que trabajan con DCV pues se corre el riesgo de causarle severos daños.**

Otra característica de los generadores de corriente continua es la capacidad para entregar corriente, este valor viene indicado en la batería y es un dato de fábrica. Si el generador dice 12 volt - 1 Amper significa que podrá entregar como máximo 1 Amper de corriente. Si le conectamos una carga que le obligue a entregar más de 1 Amper, la tensión caerá y se dice que la fuente se **“aplata”**.

Generadores de tensión de corriente alterna

El toma corriente que tenemos en nuestro domicilio es un generador de corriente puesto que de él podemos tomar energía eléctrica para alimentar los electrodomésticos (heladera, T.V. radio, horno de microondas, etc.). El valor de la tensión de ese generador es de **220 volt** pero difiere de un generador de corriente continua en que su amplitud no es constante en el tiempo, sino que varía y va tomando alternadamente valores positivos y negativos. Este tipo de tensión (o corriente) recibe el nombre de **ALTERNA**.

En particular, si la variación de la amplitud sigue las variaciones de la función trigonométrica seno, se dice que es alterna senoidal. La representación gráfica es:



Como se ve el valor instantáneo cambia permanentemente pasando por cero (t_1) y creciendo hasta un valor máximo positivo para luego comenzar a disminuir y hacerse nuevamente cero. Ahora su polaridad ha cambiado y vemos que aumenta pero en sentido negativo hasta llegar a un máximo (negativo) luego comienza a disminuir hasta hacerse cero (t_2) y así sucesivamente. A diferencia de una tensión de corriente continua en la que solo es necesario conocer el valor de tensión que entrega (que es constante en el tiempo) y su polaridad, en una tensión de corriente alterna senoidal será necesario conocer más detalles como ser: El tiempo que tarda en pasar de t_1 a t_2 . La cantidad de veces que la onda se repite por segundo, etc.

Así podemos definir:

CICLO: Es el intervalo en el cual la onda no se repite. En el gráfico tenemos un ciclo entre los tiempos t_1 y t_2

PERÍODO: Se define como el tiempo que tarda la onda en completar un ciclo, generalmente se representa con la letra T o P y su unidad de medida es el segundo. En la figura el tiempo transcurrido entre los puntos t_1 y t_2 resulta ser el período de la onda dibujada.

FRECUENCIA: Se define como la cantidad de ciclos que tiene una onda alterna durante el intervalo de tiempo de 1 segundo. Dicho de otra manera es la cantidad de ciclos que hay por segundos. Matemáticamente la frecuencia resulta ser la inversa del período T de una onda, se la representa generalmente con la letra f y su unidad de medida es el ciclo sobre segundos o **HERTZ**, en honor al físico **HEINRICH HERTZ** y se abrevia **Hz**.

A veces es necesario usar múltiplos de esta unidad como ser:

$$\begin{aligned}\text{KILOHERTZ (KHz)} &= 10^3 \text{ Hz} \\ \text{MEGAHERTZ (MHz)} &= 10^6 \text{ Hz} \\ \text{GIGAHERTZ (GHz)} &= 10^9 \text{ Hz}\end{aligned}$$

Matemáticamente la frecuencia es:

$$f = \frac{1}{T} \text{ [Hz]}$$

Así, la onda senoidal suministrada por la empresa eléctrica local es de 50 ciclos por segundos o dicho de otra forma es de 50 Hz. Lo que quiere decir esto es que en un tiempo de un segundo hay 50 ciclos. En Europa la frecuencia industrial también es de 50 Hz, pero en Estados Unidos y Canadá es de 60 Hz. En base a la frecuencia las ondas se clasifican mediante un convenio internacional de frecuencias en función del empleo a que están destinadas:

- a) **SONORAS:** Ondas comprendidas entre 20 Hz y 15 KHz.
- b) **ULTRASONIDOS:** Ondas comprendidas entre 15 KHz y 30 KHz.
- c) **ELECTROMAGNÉTICAS:** Ondas con frecuencias superiores a 30 KHz

Las ondas electromagnéticas se las denominan también de Radiofrecuencia (R.F.) porque son utilizadas en radio comunicaciones y dentro de esta clasificación se las subdivide en bandas de frecuencias de la siguiente manera:

FRECUENCIAS	CLASIFICACIÓN	DESIGNACIÓN		USO
30 a 300 KHz	Baja Frecuencia	BF	(LF)	Radio-Navegación
300 a 3000 KHz	Media Frecuencia	MF	(MF)	Radio - O. Media
3 a 30 MHz	Frecuencia	AF	(HF)	Radio - O. Corta
30 a 300 MHz	Muy Alta Frecuencia	MAF	(VHF)	T.V. - F. Modulada
300 a 3000 MHz	Ultra Alta Frecuencia	UAF	(UHF)	T.V. - Radar - Radio
3 a 30 GHz	Súper Alta Frecuencia	SAF	(SHF)	Radar- Enlace Radio
30 a 300 GHz	Extremadamente Alta Frec.	EAF	(EHF)	Radar - Enlace Radio

Si la frecuencia de la red domiciliar es de 50 Hz, podemos calcular el período de la siguiente manera:

Es decir que el tiempo que tarda en cumplirse un ciclo es de 20 milisegundos.

$$f = \frac{1}{T} [\text{Hz}] \quad ; \quad f = \frac{1}{T} [\text{seg}]$$

$$T = \frac{1}{50} = [0,02 \text{ seg.}]$$

Continuando con las definiciones, debemos también agregar las siguientes:

VALOR PICO: Es el valor máximo que alcanza con respecto al nivel de cero volt (positivo o negativo). También se lo llama valor máximo o de cresta y se lo suele representar con las letras vp.

VALOR PICO A PICO: Resulta ser el máximo positivo más el máximo negativo que puede alcanzar. Para una onda senoidal simétrica como la estudiada es dos veces el valor pico. Se representa con las letras

$$V_{pp} = 2 \times V_p$$

VALOR EFICAZ: Resulta un tanto dificultoso determinar la cantidad de energía eléctrica que suministra un generador de corriente alterna senoidal. Esto se debe a que su magnitud varía permanentemente y debido a que la onda es simétrica, su valor promedio resulta ser cero.

Si la empresa de energía eléctrica nos cobrase en base al valor promedio, las facturas de “luz”, serían gratuitas.

Tampoco nos puede cobrar en base al valor pico debido a que no es cierto que en todo momento nos esté suministrando ese valor. Se llega entonces a un valor de compromiso que, desde el punto de vista de la potencia eléctrica suministrada, es equivalente a un generador de tensión de corriente continua imaginario que desarrolla dicha potencia.

Ese valor de tensión recibe el nombre de **VALOR EFICAZ** y lo podemos definir de la siguiente manera: *El valor eficaz de una onda senoidal de corriente alterna es el valor que debería tener un generador de corriente continua para generar igual cantidad de calor sobre una resistencia de carga, que la tensión de corriente alterna considerada.*

Se puede demostrar matemáticamente que el valor eficaz de una onda senoidal queda determinado como el valor pico dividido la raíz cuadrada de 2, y se lo representa con las letras: v_{ef}

$$v_{ef} = \frac{v_p}{\sqrt{2}} = \frac{v_p}{1,4142} = 0,707v_p$$

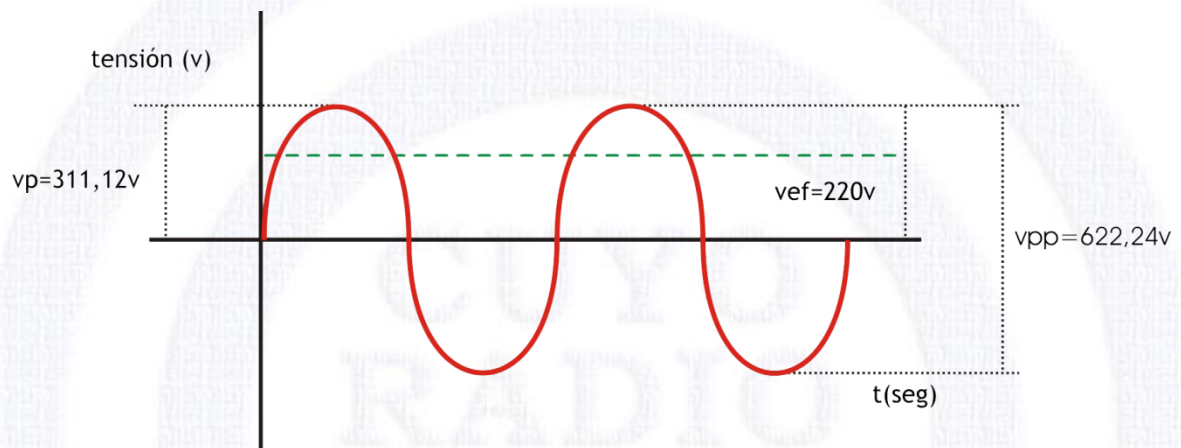
El valor eficaz es el parámetro más importante para determinar la magnitud de una tensión de corriente alterna senoidal. Tal es así que cuando en un artefacto eléctrico leemos el valor de la tensión de trabajo recomendada (220v, 110v, etc.) se está haciendo referencia al valor eficaz de dicha tensión, aunque no se lo mencione expresamente.

Si el valor eficaz de la tensión domiciliar es 220v, podemos hallar el valor pico despejando de la fórmula por lo que tendremos:

$$v_p = \sqrt{2} \cdot v_{ef} = 1,4142 \times 220v = 311,12v$$

Y el valor pico a pico de la tensión domiciliar será:

$$v_{pp} = 2 \times v_p = 2 \times 311,12v = 622,24v$$



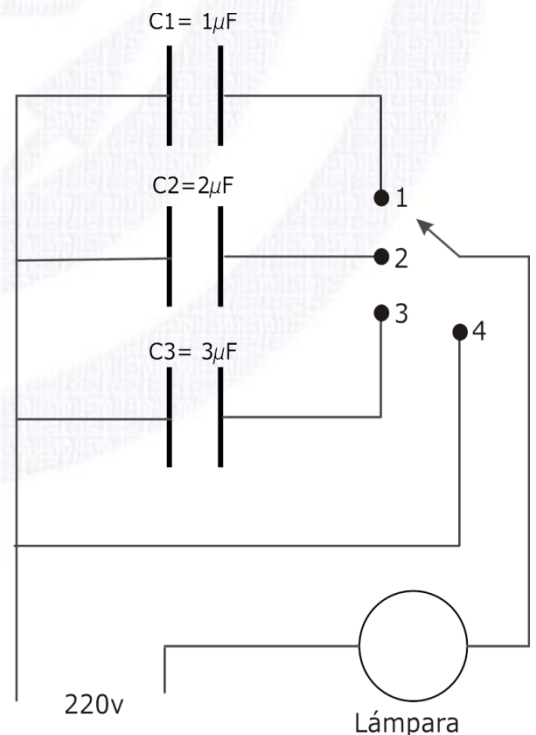
Reactancia capacitiva

Veremos ahora cómo se comporta un capacitor y una bobina cuando lo alimentamos con tensión alterna. Comenzaremos con el capacitor y para entender lo que ocurre partiremos de un experimento sencillo: Supongamos que tenemos tres capacitores cuyos valores son: $1\mu F$, $2\mu F$ y $4\mu F$ respectivamente, conectados en paralelo entre sí y alimentados por una tensión alterna de 220v.

Además tenemos una llave selectora que permite ir conectando de a un capacitor a la vez y una lámpara incandescente.

Si ponemos la llave en la posición 1 notaremos que la lámpara se enciende, tomando su filamento un color rojizo. **¿Cómo puede ser esto posible si el circuito está interrumpido por el dieléctrico del capacitor?** Lo que ha sucedido es que por tratarse de corriente alterna, el capacitor comienza a cargarse hasta un valor máximo que corresponde al valor máximo que alcanza el voltaje de la fuente, a partir de ese punto el valor del voltaje decrece hasta llegar a cero. Como el condensador se halla cargado cuando el voltaje tiende a disminuir, resulta que la diferencia de potencial que adquirió el condensador es mayor que el voltaje aplicado, puesto que éste, después de llegar a un valor máximo, disminuye hasta cero.

Resultando que si el potencial del condensador es mayor que el del circuito, entonces el condensador se descargará dando origen a una corriente en sentido contrario al de la carga y que atravesará el filamento de la lámpara. Así llegamos a un momento en que el voltaje de la fuente eléctrica es cero y a partir de ese punto, se invierte la polaridad produciéndose una corriente que cargará al condensador en el sentido contrario al anterior, hasta un valor máximo cuando el voltaje aplicado alcance al valor máximo negativo a partir del cual irá disminuyendo hasta llegar al valor cero. En este momento el condensador que se había cargado, se descargará en el circuito dando origen a una corriente en sentido contrario al de la carga. Esta carga y descarga se repite 50 veces por segundo por lo que la rapidez con que la corriente atraviesa el filamento en ambos sentidos lo mantiene encendido. Si la llave la colocamos en la posición 2, conectaremos ahora el C_2 de $2\mu F$ (el doble



del anterior). Veremos que el filamento se enciende más que con el capacitor de 1μF. La explicación de este fenómeno es que por ser C2 de mayor capacidad que C1, la carga que admite será mayor y entregará mayor corriente durante la descarga. Si ahora colocamos la llave en la posición 3, la lámpara se encenderá aún más, debido a que el capacitor es de 4μF. En la posición 4 no hay capacitor y por lo tanto la lámpara enciende normalmente. Se puede demostrar matemáticamente que en un circuito capacitivo puro la intensidad de corriente adelanta 90° con respecto a la tensión que la produce. Vemos que el capacitor en corriente alterna se comporta como una “resistencia”. El término correcto es **REACTANCIA**, que por tratarse de un capacitor se llama **REACTANCIA CAPACITIVA**. La reactancia se mide en Ohm (Ω) si la capacidad está expresada en Faradios y la frecuencia en Hz y se la representa con la letra **X_C**.

Matemáticamente la reactancia capacitiva está dada por:

$$X_C = \frac{1}{2 \times \pi \times f \times C} \quad [\Omega]$$

Es decir que **la reactancia es inversamente proporcional a la frecuencia y a la capacidad**. A mayor capacidad, menor reactancia. Si la frecuencia es cero (el caso de corriente continua) la reactancia capacitiva es un valor infinitamente grande, que equivale a decir que el circuito está abierto y no circula corriente, y es lo que ocurre cuando colocamos un capacitor en corriente continua.

Reactancia inductiva

Recordemos cuando estudiamos inductancia lo que ocurría cuando una corriente variable atravesaba la bobina. Si se hace pasar una corriente alternada por una bobina de alambre, se forma un campo magnético alrededor de la bobina. Conforme aumenta la corriente (semiciclo positivo) se expande el campo magnético que cruza transversalmente los conductores de la misma bobina, se induce una F.E.M. que da origen a una segunda corriente. Esa corriente tiene una dirección tal que se opone a la causa que la produce (ley de Lenz).

En otras palabras, la dirección de la corriente inducida es tal que tenderá a reducir la corriente original y tenderá a oponerse a la expansión del campo magnético. Cuando la corriente original comienza a disminuir, el campo magnético que rodea la bobina comienza a retraerse, igual que en el caso anterior, cruza transversalmente las espiras de la bobina y otra vez se induce una segunda corriente cuya dirección es tal que tenderá a mantener la circulación de la corriente original durante un cierto tiempo, después que la corriente original haya cesado. Lo mismo ocurre para el semiciclo negativo. Vemos que hay una oposición a que la corriente aumente o disminuya según sea el caso. La bobina tiene la propiedad de retardar el cambio de la corriente que circula por ella, es decir retarda la variación de intensidad. A esa propiedad se la llama **INDUCTANCIA**.

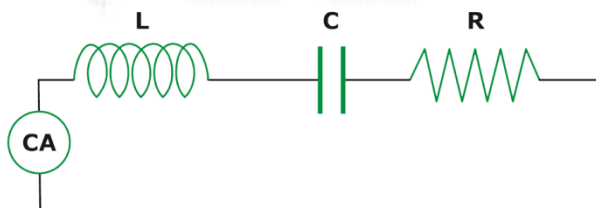
Se puede demostrar que en un circuito inductivo puro la F.E.M. inducida está retrasada 90° respecto a la intensidad de corriente que la produce. Esa oposición a la variación de la corriente recibe el nombre de **REACTANCIA**, y por tratarse de un fenómeno que ocurre en la inductancia, recibe el nombre de **REACTANCIA INDUCTIVA**. Se mide en Ohm (Ω) si la inductancia está en Henry y la frecuencia en Hz, se la representa con las letras **X_L** y matemáticamente la reactancia inductiva está dada por:

$$X_L = 2 \times \pi \times f \times L \quad [\Omega]$$

Es decir que es directamente proporcional a la frecuencia y a la inductancia. Si la frecuencia es cero (es el caso de la corriente continua), la reactancia inductiva es cero, es decir no hay ninguna oposición al paso de la corriente. Es un corto circuito.

Impedancia

Supongamos que tenemos un circuito formado por una bobina, un capacitor y una resistencia conectados en serie y alimentados con una tensión alterna senoidal como indica la figura.



Como el circuito está alimentado con corriente alterna, la bobina tiene una reactancia **X_L (Ω)**, y el capacitor tiene una reactancia **X_C (Ω)**, además de la resistencia **R (Ω)** que hay en el circuito. Se puede demostrar matemáticamente

que el circuito presenta una “**resistencia total**” a la circulación de la corriente eléctrica y su nombre correcto es de **IMPEDANCIA**, puesto a que “*impide*” el paso de la corriente eléctrica. Se mide en **Ohm** y se la representa con la letra **Z**.

En el caso de los tres elementos en serie la impedancia está dada por:

$$Z = \sqrt{R^2 + X_L^2 - X_C^2} \quad [\Omega]$$



TEMA: 6

GENERADORES DE RADIOFRECUENCIA

Introducción

En el capítulo anterior vimos que las ondas de corriente alterna se podían clasificar en base a su frecuencia. En particular estudiaremos las ondas usadas en radiocomunicaciones (ondas electromagnéticas o de radiofrecuencia) que van desde los 30 KHz a 300GHz.

Generación de ondas de radiofrecuencia

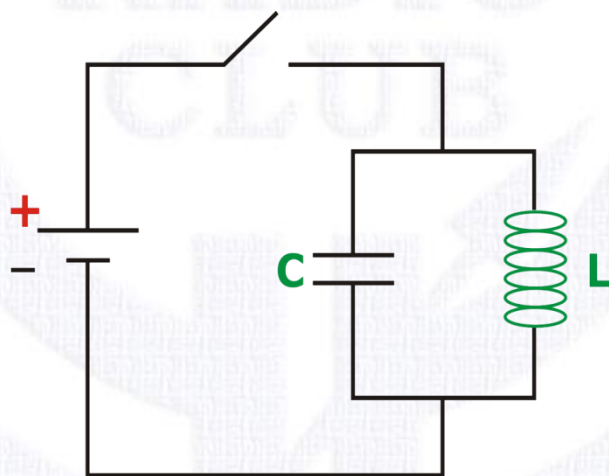
Una onda de corriente alterna se puede producir de varias maneras. La corriente alterna usada en nuestros domicilios y en la industria, se produce utilizando grandes generadores (llamados alternadores) instalados en usinas eléctricas. Este procedimiento se usa cuando tales corrientes son de baja frecuencia.

Cuando necesitamos generar una corriente de frecuencia elevada (**radio frecuencia**) debemos utilizar un “**OSCILADOR**”. Existen muchos tipos de osciladores, veremos solamente dos tipos de ellos: El oscilador L - C y el oscilador a cristal. El oscilador es el corazón de los transmisores, frecuentemente el transmisor mismo.

Oscilador L – C

Se trata de osciladores que trabajan con una bobina (con inductancia L) y un capacitor (con capacidad C).

El oscilador más elemental que podemos fabricar es:



En donde tenemos una fuente de corriente continua, un capacitor conectado en paralelo a una bobina y un interruptor. Trataremos de explicar cómo funciona este circuito, con los conocimientos adquiridos hasta el momento. Lo primero que hacemos es cerrar el interruptor e inmediatamente lo abrimos.

Lo que ocurre en el circuito se detalla a continuación: al cerrar el interruptor, el capacitor se carga a su valor máximo, al abrir el interruptor, el capacitor que se halla cargado tiene conectado entre sus extremos una bobina y por lo tanto comenzará a descargarse a través de ella.

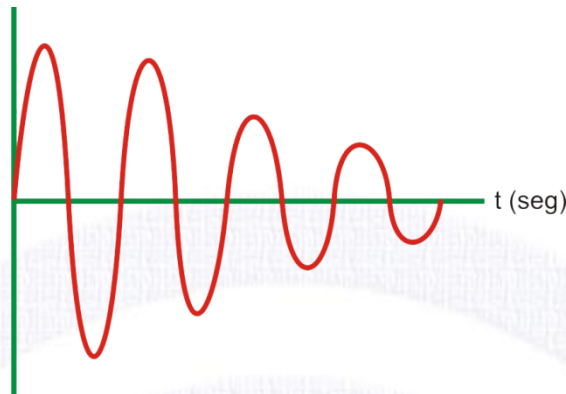
La corriente de descarga es una corriente variable en el tiempo (va disminuyendo a través del tiempo) por lo tanto al atravesar la bobina, se produce una fuerza electromotriz inducida que, de acuerdo con la ley de Lenz, se opone a la causa que lo produce.

Esta F.E.M. inducida carga nuevamente el capacitor pero ahora en sentido opuesto al anterior. Nuevamente comienza a descargarse a través de la bobina y se repite el proceso.

Se puede decir que hay un intercambio de energía electrocinética del capacitor en energía electromagnética en la bobina. Vemos que se producen oscilaciones de energía entre L y C .

El circuito formado por una bobina y un capacitor en paralelo recibe el nombre de “**CIRCUITO TANQUE**”.

Esto ocurre en la teoría, pero en la práctica las cosas son un poco diferentes. Sabemos que en todo circuito está presente una resistencia (aunque sea la resistencia del conductor), por eso la energía se va disipando en el tiempo hasta hacerse cero. Una onda de éste tipo recibe el nombre de “ONDA AMORTIGUADA” y es como la muestra la figura.



Nótese que la frecuencia no varía, lo único que varía es la amplitud. Para evitar que la onda se amortigüe, debemos entregar al circuito una energía de mantenimiento. Podemos decir que es lo mismo que ocurre cuando a un niño que juega en una hamaca, debemos darle un nuevo empujón para que siga hamacándose, caso contrario transcurrido un cierto tiempo se detendría. Hay dispositivos electrónicos que son capaces de hacer este trabajo como lo son las válvulas electrónicas y los transistores.

Frecuencia de la señal

Podemos decir que para cada par de bobina - capacitor existe una única frecuencia de oscilación que es la frecuencia de resonancia y se puede calcular con la fórmula vista anteriormente.

$$f = \frac{1}{2 \times \pi \times \sqrt{L \times C}}$$

A modo de ejemplo calcularemos la frecuencia de un circuito tanque que posee una bobina de $50 \mu\text{H}$ y un capacitor de 40 pF .

$$f = \frac{1}{2 \times 3,14 \times \sqrt{50 \times 10^{-6} \text{Hy} \times 40 \times 10^{-12} \text{F}}}$$

$$f = 3.558.812 \text{Hz} \approx 3,55 \text{MHz}$$

Nótese que la capacidad está en faradios y la inductancia en Henrios por lo que la frecuencia nos da en Hz.

En los libros de radio, existen tablas y ábacos que nos permiten encontrar los valores de un circuito oscilante.

Oscilador a cristal

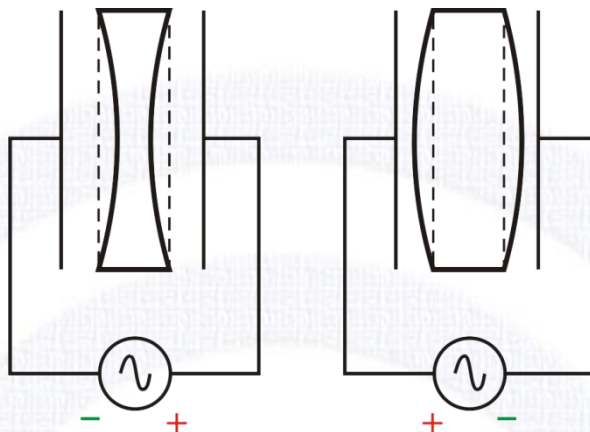
Los cristales de Cuarzo

El componente denominado cristal tiene externamente el aspecto de una caja metálica cerrada de la que asoman por su base un par de terminales de conexión. En el interior de este encapsulado se encuentra una lámina de cuarzo en forma circular o rectangular, que presenta sobre sus dos superficies unas metalizaciones unidas eléctricamente a los terminales de conexión, mediante dos hilos conductores.

Esta lámina de cuarzo es la encargada de realizar la función principal de este componente, merced a sus propiedades piezoeléctricas.

Explicaremos el fenómeno natural llamado “**PIEZOELECTRICIDAD**”. La palabra “piezo” proviene del griego y significa **presión**, esto da una idea del fenómeno. La piezo electricidad significa “**electricidad por presión**”. Se trata de una particularidad que presentan los cristales de ciertas sales minerales. Muchos son los cristales que poseen efectos piezoeléctricos, pero tres son los que han dado mejores resultados: sal de la **ROCHELLE**, **TURMALINA** y **CUARZO**.

Cuando se comprime uno de estos cristales, se desarrolla una diferencia de potencial entre sus dos caras. El aumento de presión da como resultado un aumento de la diferencia de potencial



Además si se aplica una tensión eléctrica entre sus dos caras paralelas se origina en éste una deformación mecánica.

Al eliminar esta tensión la lámina recupera su forma original, pero para llegar a ella pasará por una serie de estados intermedios semejantes a una oscilación ya que en la primera aproximación sobrepasará la forma primitiva, debido a la inercia mecánica, deformándose en sentido contrario y volviendo hacia atrás hasta que al cabo de un cierto tiempo se detendrá.

La frecuencia a la que se produce este fenómeno es fija y depende exclusivamente del cristal, pudiendo ser considerada como su frecuencia natural de oscilación.

Estructura cristalina

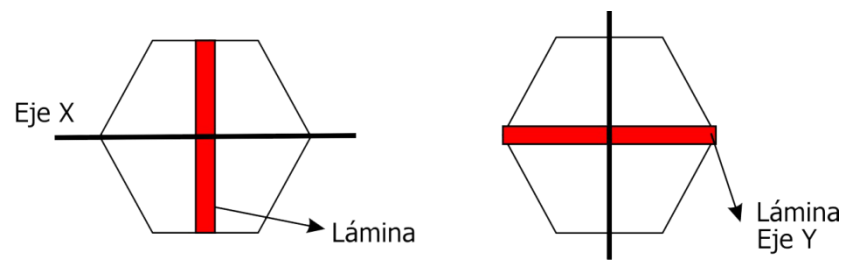
El cuarzo es un mineral formado por anhídrido de silicio cuya fórmula es:



se encuentra en la naturaleza en diferentes variedades, aunque para la aplicación que nos interesa, únicamente se emplea la formada por cristales prismáticos hexagonales, acabados en pirámides por sus caras extremas. Sobre este cristal se trazan unos ejes imaginarios que son los siguientes:

- Eje óptico, designado comúnmente por la letra z, que pasa por los vértices de las pirámides de los extremos, atravesando el cuerpo del cristal.
- Ejes eléctricos, designados por la letra x, que pasa por los centros de las aristas laterales del prisma, atravesando el centro geométrico del cristal. Existen en número de tres.
- Ejes mecánicos, designados con la letra y, que se trazan por los centros de las caras del prisma, atravesando el centro geométrico del cristal. También éstos forman un total de tres.

La lámina que debe formar el cristal oscilador se talla del cuerpo cristalino original en forma perpendicular a uno de los ejes eléctricos (x) o a uno de los mecánicos (y).



Las láminas así obtenidas presentan el efecto piezoeléctrico descrito anteriormente.

Resonancia

Si en lugar de una tensión continua, se aplique otra que varíe con una frecuencia igual a la propia de la lámina, de forma que se encuentren en resonancia, se reforzarán notablemente las vibraciones propias del cristal, produciéndose así una oscilación mantenida por éste y estabilizada, al ser su propia de resonancia. Esta resonancia desaparece en cuanto la frecuencia de la tensión de excitación se aparta un cierto número de Hertz de la propia de la lámina, deteniéndose la oscilación del cristal.

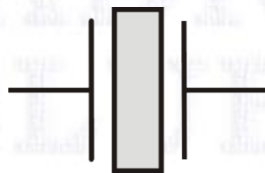
Formas de operación

El cristal puede trabajar según dos modos diferentes de operación: resonancia serie y resonancia paralelo. Las frecuencias resultantes de estas formas de trabajo son diferentes y en el momento del diseño del cristal se elige una de las dos, con el objeto de favorecer el modo elegido frente a otro.

En algunos modelos de cristal, sobre todo en aquellos preparados para funcionar con altas frecuencias, no se emplea la oscilación mecánica fundamental de la lámina sino que se les hace trabajar a unas frecuencias superiores múltiplos de aquélla, encontrándose así diversos tipos de cristales que funcionan en el 3°, 5° y hasta el 7° armónico (3, 5 ó 7 veces la frecuencia natural de oscilación), alcanzando frecuencias de centenares de Mhz.

Una característica importante que siempre habrá de ser tenida en cuenta es la variación de la frecuencia con la temperatura que cada modelo de cristal presenta y que se encuentra cuantificada en el catálogo. Suele expresarse en partes por millón, por grado centígrado (ppm/°C) y también en forma de porcentaje. Así, por ejemplo, si se tomó un cristal de 1Mhz con una variación con la temperatura de más menos 20 ppm/°C, significará que la frecuencia se desplazará en mas menos 20 Hz por cada grado de cambio de temperatura, por lo tanto si se va a encontrar en un ambiente en el que puede tenerse una oscilación de 10°C, la frecuencia va a moverse en mas menos 200 Hz, es decir entre 999.800 Hz y 1.000.200 Hz.

El símbolo del cristal es:



Si quisiéramos un oscilador a cristal de una frecuencia de 3,55 MHz deberíamos conseguir un cristal de esa frecuencia en el comercio o bien encargar su fabricación.

La virtud del cristal es la perfecta estabilidad, es decir que el valor de la frecuencia de oscilación es bastante estable, por ello se lo utiliza en osciladores patrones.

Oscilador de frecuencia variable (o.f.v.)

Hasta ahora hemos visto osciladores de una sola frecuencia. Es más útil tener un oscilador que pueda cubrir varias frecuencias. Se puede tener un oscilador con varios cristales y con una llave selectora ir seleccionando la frecuencia deseada. Otra forma y tal vez la más usada es la de tener un oscilador L - C en el cual el capacitor no es un capacitor fijo sino que es un capacitor variable. Con esto se logra que cada vez que movemos el eje del capacitor, variamos la frecuencia del oscilador. Así, por ejemplo, podemos generar frecuencias que van desde 3.5 MHz hasta 3,75 MHz con una sola bobina y un capacitor variable del valor adecuado.

Podríamos dejar fijo el capacitor y variar la inductancia (L) de la bobina y obtendríamos el mismo resultado. En la práctica es mucho más fácil variar la capacidad que la inductancia.

Las ondas electromagnéticas que genera un oscilador de radiofrecuencia, se propagan en forma de un campo magnético y eléctrico que oscilan perpendicularmente entre si. Reciben el nombre de ondas electromagnéticas.

Velocidad de la luz

Todas las radiaciones electromagnéticas se propagan en el espacio libre a la misma velocidad, conocida con el nombre de velocidad de la luz. A esta velocidad se la suele representar con la letra c, es una de las constantes universales y su valor resulta ser:

$$c = 300.000 \text{ km. / seg.}$$

Longitud de onda

Cuando las frecuencias de las ondas electromagnéticas son muy elevadas, se suele utilizar este concepto para caracterizarlas. La longitud de onda se define como la distancia que recorre una onda electromagnética de frecuencia f , durante un tiempo igual a su período T . Se representa la longitud de onda con la letra griega λ (lambda). Teniendo en cuenta que la señal viaja a la velocidad de la luz podemos calcularla haciendo el cociente entre el espacio recorrido (distancia) y el tiempo empleado en recorrerla, es decir:

$$\text{Velocidad} = \text{Distancia} / \text{tiempo}$$

O si lo ponemos en función de la distancia (*que es el espacio recorrido*) tenemos:

$$\text{Distancia} = \text{velocidad} \times \text{tiempo}$$

Donde: Distancia es la longitud de onda

Velocidad es la **velocidad de la luz**

Tiempo es el **período T**

$$\lambda = c \times T$$

Recordando que $T = 1/f$ y reemplazando tenemos:

$$\lambda = \frac{c}{f}$$

Si la velocidad se mide en m/seg. y la frecuencia en hertz, la longitud de onda se mide en metros. De acuerdo con la ecuación anterior, existe una relación perfectamente definida y muy importante entre la frecuencia de una radiación determinada y su longitud de onda. Por ejemplo, ondas de audio frecuencias tienen longitudes de onda del orden de los kilómetros, mientras que sistemas de comunicaciones radiales de alta frecuencia operan con ondas que tienen longitudes de onda del orden de los milímetros. Por acuerdos internacionales se han asignado determinadas bandas de longitudes de onda, posibles de ser empleadas en determinadas aplicaciones, especialmente en el área de las comunicaciones: radiodifusión, televisión, usos militares, radar, etc.

Se ilustra a continuación un gráfico que muestra el espectro de las ondas electromagnéticas.

10 Km	1 Km	100 m	10 m	1 m	1 mm	1 nm	10 nm	
LF	MF	HF	VHF	RADAR	IR	LUZ VISIBLE	UV	RAYOS X y GAMMA

La luz visible corresponde a solo una estrecha banda del espectro electromagnético que incluyen muchos tipos de onda. Los rayos gamma tienen una frecuencia extremadamente elevada y una longitud de onda baja. Las ondas de radio tienen una frecuencia mucho menor y una gran longitud de onda.

Cada color de la luz visible corresponde a una frecuencia particular del espectro. La luz violeta tiene la menor longitud de onda que puede detectar el ojo humano y el color rojo la mayor.

El resto de los conceptos referidos a este tema, se detallarán en los apuntes de propagación y antenas.

TEMA: 7

TRANSFORMADORES

Introducción

El transformador es una máquina estática, esto es, que no tiene piezas en movimiento. Se emplea en las unidades de alimentación de receptores, emisores, amplificadores, etc. Así, transfiriendo energía de un circuito a otro, por medio del flujo magnético que lo rodea, permite obtener tensiones de valores diferentes o iguales a las de la fuente de alimentación, por medio de arrollamientos secundarios, aislados entre sí, de acuerdo con las necesidades de trabajo.

El funcionamiento del transformador se basa en los fenómenos de inducción, por variación del flujo magnético, y por eso no pueden trabajar nunca con corriente continua. En la mayoría de sus aplicaciones el transformador trabaja con corriente de forma sinusoidal, pudiendo tener ésta, sin embargo, otras formas de onda (rectangular, compleja, etc.). De la forma de onda depende el rendimiento del transformador.

En su versión sencilla, un transformador está constituido por:

- a) Un arrollamiento primario.
- b) Uno o varios arrollamientos secundarios.
- c) Un circuito magnético en el cual se encuentran arrollados los devanados primarios y secundarios.

Se puede decir que la potencia eléctrica consumida en el primario se transforma en magnética, volviendo a la forma primitiva de potencia, inducida en el secundario.

Cuando se devana una bobina en el núcleo de hierro y se conecta a una fuente de energía de C.A., se produce un campo magnético variable, cuya intensidad dependerá del número de espiras de la bobina y de la intensidad de la corriente que la atraviesa.

Si en el mismo núcleo se coloca otra bobina, se obtiene una corriente inducida en ella, en virtud de encontrarse en el campo magnético de la primera, o sea, de la bobina excitadora, y cuya f.e.m. inducida dependerá del número y relación de espiras entre una y otra.

Podemos decir que la tensión inducida es proporcional al número de espiras. Por lo tanto, si una bobina excitadora tuviera 660 espiras y la tensión de la red de alimentación fuera de 220 V, y la segunda bobina tuviera 1200 espiras, la tensión en bornes de esta bobina será de:

$$\frac{660}{220} = 3 \quad \text{y de esta forma} \quad \frac{1200}{3} = 400 \text{ voltios}$$

Para comprender mejor esto, tenemos la siguiente relación:

$$\frac{V_1}{V_2} = \frac{N_1}{N_2}$$

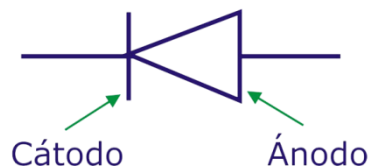
En la que V_1 es la tensión en bornes del primario y V_2 la del secundario, siendo N_1 el número de espiras del primario y N_2 el de las del secundario. Esta relación, llamada relación de transformación, sólo es válida cuando el secundario no está conectado a otro circuito, es decir, está abierto. Si reemplazamos los valores podemos calcular V_2 .

Tenemos entonces:

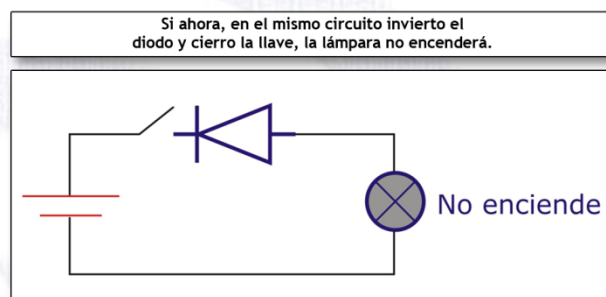
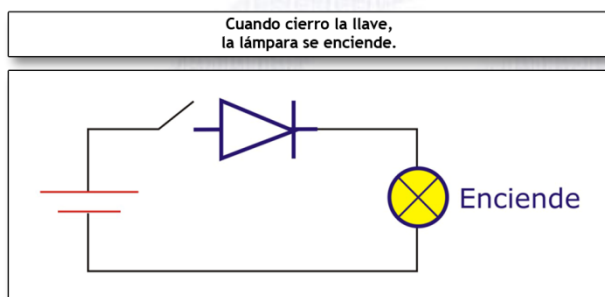
$$V_2 = \frac{1200 \times 220}{660} = 400 \text{ voltios}$$

Diodos

La palabra diodo, deriva de “dos electrodos” y su misión es dejar pasar corriente en un sólo sentido. Su símbolo electrónico es:



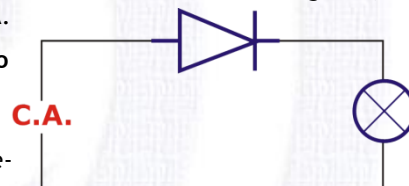
Generalmente se utilizan los diodos para dejar pasar corriente en un sólo sentido. Para entender mejor esto, veamos el circuito de la figura en donde el ánodo está conectado al borne positivo de la batería.



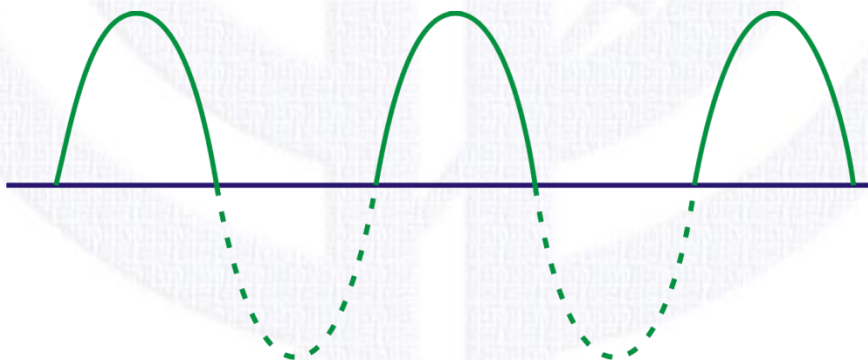
Como vemos el diodo deja circular corriente sólo si el ánodo está conectado a borne positivo, en ese caso se dice que el diodo se ha polarizado en forma **DIRECTA**, en cambio si conectamos el ánodo del diodo al borne negativo, no circulará corriente en nuestro circuito porque está polarizado en forma **INVERSA**.

¿Qué ocurrirá si en vez de alimentar el circuito con tensión continua, lo hago con alterna?

Sabemos que la corriente alterna toma valores positivos y negativos, por lo tanto el diodo conducirá en los instantes en que quede polarizado en forma directa y no lo hará cuando está polarizado en forma inversa.



Para entender mejor lo que ocurre, graficaremos las ondas de tensión



Este circuito se conoce como **RECTIFICADOR DE MEDIA ONDA**, puesto que sólo deja pasar media onda o ciclo.

Lo que hemos conseguido es convertir una onda alterna (que tomaba alternadamente valores positivos y negativos) en una onda que siempre es positiva.

Como no cambia de polaridad, se dice que es una **ONDA CONTINUA PULSANTE**, pues, como vemos, son solamente pulsos positivos los que quedan a la salida.

Si quisiéramos tener una onda continua pura deberíamos tratar de “alisar” los pulsos y ese trabajo lo conseguimos poniendo un capacitor en la salida.

Durante el ciclo positivo el capacitor se carga al valor pico de la tensión alterna, luego durante los picos negativos en los que el diodo no conduce, el capacitor se descarga.

La rapidez con que se descarga depende del tamaño del capacitor, mientras más grande sea el capacitor más tiempo tardará en descargarse. En forma exagerada la onda rectificada presentaría la siguiente forma.

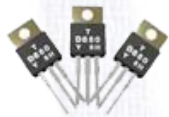


¿Qué valor tengo de corriente continua? Dijimos que el capacitor se carga al valor pico (máximo), por lo que este es el valor que tendremos de C.C. Por ejemplo, si tenemos un transformador de 220 v a 12 v (que sabemos que son eficaces), el valor de tensión de C.C. estará dado por:

$$V_{ef} = \frac{V_{max}}{\sqrt{2}} \rightarrow V_{max} = V_{ef} \times \sqrt{2} = 12 \times \sqrt{2} = 16,97$$

Existen otros tipos de rectificadores que usan dos diodos y se denominan rectificadores de onda completa, ya que cada diodo conduce durante medio ciclo y la onda rectificada luego es filtrada por un capacitor, obteniéndose una onda de C.C. pura. Otro tipo de rectificador es el que usa cuatro diodos y se denomina rectificador de onda completa tipo puente, puesto que los diodos están conectados en forma de puente.

Transistores



El transistor es, actualmente, un componente fundamental en cualquier circuito electrónico que realice funciones de amplificación, control, radio, TV, estabilización de corriente, etc. El transistor es un elemento semiconductor (**normalmente de tres terminales**) que tiene la propiedad de poder gobernar a voluntad la intensidad de corriente que circula entre dos de sus tres terminales, a través de la acción de una pequeña corriente, mucho más baja que la anterior aplicada al tercer terminal.

Los dos primeros se llaman emisor y colector y el tercero recibe el nombre de base.

Básicamente existen dos tipos de transistores, los **NPN** y **PNP** cuyo símbolo se muestran en la figura.

El efecto descrito es una amplificación de corriente ya que gracias a la acción de una débil intensidad que puede tener cualquier forma de variación en el tiempo, tales como señales de **TV**, **radio**, **audio**, etc., se consigue obtener la misma forma sobre una corriente mayor, proporcionada ésta por un circuito de alimentación, lo que permite el poder realizar, en las sucesivas etapas, la transformación de una señal debilísima, en otra lo suficientemente fuerte como para ser capaz de producir sonido en un parlante, imagen en un TV., etc.



La palabra transistor se obtuvo de la composición de otras dos (**TRANS**fer - **reSISTOR**) que describen su aplicación más inmediata, la de transferencia de resistencia. Se descubrió en el año 1948 por **Shockley** como resultado de los trabajos efectuados, previamente, por **Bardeen** y **Brattain** sobre fenómenos eléctricos en la superficie de los semiconductores. Los tres científicos recibieron el premio Nobel de Física en el año 1956.



Válvulas termoiónicas

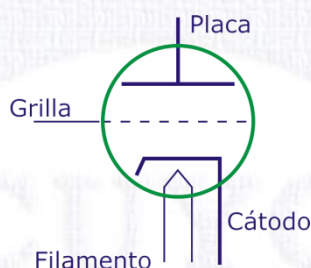


Estos dispositivos electrónicos basan su funcionamiento en la emisión de electrones por calentamiento de un electrodo. Sabemos que si aumentamos la temperatura de un metal, sus electrones adquirirán una mayor velocidad, cuando el metal alcanza una temperatura suficientemente elevada, algunos electrones adquieren tal velocidad que “**escapan**” de la superficie del metal. En la válvula se hace uso de ésta acción para producir la emisión electrónica necesaria, pero el metal se calienta en el vacío.

Esto se consigue colocando los electrodos dentro de una ampolla de vidrio o una cubierta metálica, a la que luego se le extrae el aire. Diremos que **una válvula consta de un cátodo**, que emite electrones, **y de uno o más electrodos adicionales**, los que controlan y reciben esos electrones.

El cátodo, que tiene forma de tubo, se calienta por medio de un **filamento**, que está en el interior del cátodo.

Los electrones emitidos por el cátodo, son atraídos por otro electrodo que está a un potencial positivo y que recibe el nombre de placa. Si se coloca un tercer electrodo, en forma de alambre arrollado (llamado **grilla**), entre medio de la placa y el cátodo tendremos una válvula de **tres electrodos** o **Tríodo**. Su símbolo es como indica la figura:



Si la grilla se conecta a un pequeño potencial negativo, no dejará pasar la totalidad de los electrones desde el cátodo a la placa. Como vemos, podemos regular el flujo de electrones que circula desde el cátodo a la placa. Si ponemos en la grilla un potencial negativo elevado, los electrones que se emiten desde el cátodo no pasarán a la placa y por lo tanto no circulará corriente. En este caso se dice que la válvula está al corte.

Existen otras válvulas además del tríodo, como son los **tetrodos** (cuatro electrodos), **tetrodos de haces dirigidos**, **pentodos** (cinco electrodos), **heptodos** (siete electrodos).

Cuando es necesario utilizar potencias elevadas en radiofrecuencia (500 watts ó más), la válvula aún no ha sido superada por el transistor. Es por ello que el aspirante a radioaficionado podrá ver que los amplificadores lineales usados en las bandas de HF, están contruidos con válvulas termoiónicas.

TEMA: 8

TRANSMISIÓN RADIOELÉCTRICA

Introducción

La limitación inherente a la comunicación por medio de ondas sonoras, a medias y grandes distancias, es un hecho conocidos por todos, y ha constituido un verdadero caballo de batalla de muchos científicos e investigadores, que dedicaron gran parte de sus esfuerzos a la búsqueda de otras formas de enlace entre las personas que, de algún modo, salvaran las dificultades de la comunicación directa.

Así, en la década de **1830**, **Samuel F.B. Morse** puso en práctica la **comunicación telegráfica**, y no fue hasta el año **1876**, en que **Alexander Graham Bell** construyó su primer teléfono, resolviendo el problema de la comunicación hablada entre dos puntos lejanos.

Con estos descubrimientos, el problema de la comunicación a distancia aún no estaba suficientemente resuelto, pues tanto el telégrafo como el teléfono exigen que un cable comunique los aparatos transmisor y receptor. En el año **1888**, el físico alemán **Heinrich Hertz** comprueba la existencia real de las ondas electromagnéticas (que desde entonces llevan el nombre de ondas hertzianas), demostrando que tenían todas las propiedades de la luz: **reflexión y refracción, interferencias, difracción, polarización y velocidad de propagación**.

El francés **Edouard Branly** observó, en **1890**, que la conductividad de las virutas de metal encerradas en un tubo de vidrio, podía ser alterada por las ondas hertzianas y en **Pontecchio**, cerca de Bolonia (**Italia**), un hombre de 22 años de edad experimentaba, en **1897**, un invento producto de su ingenio, que supondría el inicio de lo que más tarde sería la revolución más grande de la humanidad en materia de comunicación. Nos referimos, como el aspirante puede suponer, a la radio y a su inventor **Guillermo Marconi**.

Casi simultáneamente en **Kronstadt**, Rusia, el profesor **A.S. Popov** mejoró el cohesor de **Branly** puesto que las virutas de metal una vez que habían recibido una onda, su resistencia caía y se mantenía baja paralizando el sistema hasta que las virutas se liberaran otra vez. Branly lograba esto, golpeando la mesa, pero Popov usó la misma señal para restablecer la sensibilidad del receptor.

Al igual que Hertz, Popov tuvo una vida muy corta. Murió en 1906 a la edad de 45 años.

A partir de aquí, la comunicación hablada con ondas hertzianas, que por aquel entonces se llamó telegrafía sin hilos, y que luego adoptó el apelativo de radiotelefonía, fue problema de pocos años.

A principio de los años veinte apareció la radio y cientos de emisoras proliferaron por todo el mundo. Las ondas hertzianas, pertenecientes a la familia de las ondas electromagnéticas, presentan dos grandes ventajas frente a las ondas sonoras, en lo que a comunicación se refiere. En primer lugar, **no necesitan de ningún medio “físico” para trasladarse** (al contrario de las ondas sonoras, que requieren el aire, el agua o algún otro soporte de transmisión); la prueba resulta patente: la luz del sol y de las estrellas (que, como luz, son ondas electromagnéticas), llegan hasta nosotros a través del vacío del espacio interestelar, lo que sería poco menos que imposible si, para su propagación, necesitaran un medio concreto.

La segunda ventaja es su **velocidad de propagación**; mientras las ondas sonoras se propagan en el aire a una velocidad de **340 m/seg.**, las ondas electromagnéticas lo hacen a **300.000 Km/seg.** (Cerca de un millón de veces más rápido).

A pesar de esta clara diferenciación entre los dos tipos de ondas, ambas presentan muchos puntos de utilización comunes, razón por la que hemos insistido tanto en los medios empleados en la comunicación sonora. Por ejemplo, las ondas electromagnéticas también se atenúan con la distancia y, aproximadamente, en la misma proporción que lo hacen las ondas sonoras.

No obstante, el aspirante a radioaficionado, se preguntará aún por qué razón, si ambos tipos de ondas tienen similares características de atenuación, lleguen más lejos las radiaciones electromagnéticas. La respuesta a esta cuestión se obtiene haciendo un sencillo cálculo, ya que cualquier receptor de los actuales es 100 veces más sensible que el oído humano mejor dotado (cada uno en su campo), y además hay emisoras que transmiten con una potencia 10.000 veces mayor que el grito más alto que pueda emitir una garganta del hombre. Sólo queda multiplicar y sacar conclusiones.

La modulación

Todos los argumentos en torno a las ventajas de las ondas electromagnéticas no servirían absolutamente para nada si todo se quedara ahí, ya que una onda electromagnética “pura” no es más que eso: un vehículo que no transporta ningún pasajero, es decir, una radiación que no conlleva ningún tipo de información. Es lo mismo que si nos asomamos una noche a una ventana de nuestra vivienda, nuestro oído será capaz de percibir un cierto ruido, un rumor más o menos lejano, que no nos dirá más que “algo” o “alguien” lo está produciendo: se trata, pues, de ondas sonoras sin información útil para nosotros.

La primera idea que se nos ocurre para que una onda electromagnética lleve algo de información (“porte” información), es la de interrumpirla a intervalos más o menos frecuentes. De esta forma se pondrán en evidencia dos situaciones perfectamente diferenciables: o hay una señal o no la hay. Incluso podemos avanzar un poco más, en el sentido de hacer que los intervalos en los que existe señal puedan tener duraciones diferentes. De esta forma seremos capaces de distinguir no sólo la existencia o no de la señal, sino su mayor o menor duración.

Así fue como se realizaron las primeras transmisiones de radio, aprovechando un código que Morse había inventado años antes (el “**código Morse**”), consistente en asignar a cada letra, número y/o signo ortográfico uno o varios intervalos de distinta duración (**conocidos como rayas y puntos**). Espaciando adecuadamente la transmisión de los códigos correspondiente a cada letra, puede enviarse un mensaje inteligible que puede ser descifrado por el receptor.

El código Morse puede asimilarse al lenguaje hablado, siempre y cuando exista un acuerdo preestablecido y conocido por el que lo emite y por el que va a recibirlo, haciendo así posible un entendimiento por medio de sonidos.

Ya sabemos, por tanto, como imprimir información a una onda electromagnética. Sin embargo, este sistema de comunicación requiere del aprendizaje del código Morse, de la misma forma que el lenguaje hablado precisa conocer el significado de cada palabra y la sintaxis correspondiente. Además el código Morse jamás sería capaz de transmitir una información de tipo musical o visual, que son las que más fácilmente traducimos. Por todas estas razones, es evidente que este sistema no resulta satisfactorio para la transmisión de la información que, habitualmente, el hombre utiliza.

El concepto de modulación comienza ya a visualizarse, y para comprender mejor su significado recurriremos a un ejemplo que ilustra mejor la idea.

Se llaman vocales, en lenguaje hablado, a las letras que se pronuncian con tan sólo emitir la voz. En castellano son cinco, y resulta muy fácil comprobar que una suena de distinta forma que la otra tan solo con cerrar más o menos la boca, y con poner la lengua en una u otra posición. En estos casos se dice que modulamos la voz, porque lo único que hacemos es modificar una corriente de aire, salida de nuestra garganta, para que una determinada vocal suene de una forma u otra.

A la corriente de aire podríamos darle el nombre de portadora, pues se trata del vehículo que llevará o “**portará**” la modulación que le imprimamos con la boca y/o lengua. Evidentemente, esta portadora, por si sola, no es sonido, o dicho de otra forma, no transporta lo que venimos llamando información.

Esto tan sencillo y que tan habituados estamos a hacer a diario es lo que acontece en una emisora de radio. Se crea una onda portadora (el equivalente a la corriente de aire) que se lleva a un modulador (la boca y la lengua), para que sobre ella se imprima la información que se desea transmitir, la cual se materializa en forma de una señal moduladora (lo que expresamos con el lenguaje).

El resultado final es la onda modulada que será radiada por el espacio para que otra u otras personas puedan recogerla e interpretarla.

Tipos de modulación

El aspirante a radioaficionado ya conoce dos de los parámetros más importantes de una onda, que son su amplitud y su frecuencia. Recordemos que la amplitud da una idea del valor máximo que adoptará la onda, mientras que la frecuencia nos indicará los ciclos completos que se repiten en cada segundo.

La forma de las ondas hertzianas es idéntica a la del sonido. La única diferencia está en la frecuencia, ya que mientras en las ondas sonoras audibles rara vez se llega a los 20 KHz, las de radio pueden llegar a tener frecuencias superiores a los 10 GHz.

Aunque no se ha dicho explícitamente, ya hemos mencionado una forma de modulación cuando hablamos del sistema Morse. Se trata de dar amplitud máxima a la onda (cuando se transmite un punto o raya) y amplitud cero (intervalo entre puntos y/o rayas). Este método de modulación está basado en el “todo” o “nada”.

Se advertirá que en este caso lo que se modifica es la amplitud de la onda. Todos los sistemas que provocan una variación de la amplitud de la portadora como consecuencia del proceso de modulación, reciben el nombre de moduladores de amplitud, y la señal modulada se denomina señal de **amplitud modulada**, que se designa abreviadamente “**AM**”.

De la misma forma que hemos introducido la información en la portadora modificando uno de sus parámetros, la amplitud, también podríamos hacerlo actuando sobre el otro parámetro antes mencionado, es decir, la frecuencia. Así, si nuestra portadora tuviera una frecuencia de, por ejemplo, 500 KHz. cuando quisiéramos transmitir un punto o una raya, podríamos cambiar la frecuencia 510 KHz. Esa variación de 10 KHz. en la frecuencia la interpretará el receptor como la transmisión de un punto o de una raya, según su duración. La modulación que provoca una variación de la frecuencia de la onda portadora como consecuencia de la influencia de la señal moduladora, recibe el apelativo de modulación de frecuencia, y a la señal modulada se la llama señal de **frecuencia modulada** que se designa abreviadamente como “**FM**”.

Existen otros tipos de modulación, como son la de fase, por impulsos, etc., pero a los efectos de la transmisión radioeléctrica, nos interesan los dos que acabamos de citar. lo que es verdaderamente importante es hacerse a la idea

de que para modular una onda, hay que modificar algunos de sus parámetros que la definen. Una onda portadora no posee en sí misma ningún tipo de información que pueda resultarnos útil.

La modulación de amplitud

Fijándonos en que la señal de audio (**señal de B.F.** procedente de un micrófono o de cualquier reproductor musical debidamente amplificada) es variable, podemos hacer variar la amplitud de la portadora al ritmo de la de audio, diseñando el reproductor de forma tal que sea sensible sólo a dicha variación de amplitud y, en consecuencia, sea capaz de recoger de forma adecuada la información contenida en la onda modulada en AM.

Este es el fundamento de la modulación de amplitud utilizada en los sistemas de comunicación por radio, en los que la portadora es sólo el vehículo que transporta al conductor, que es la señal de audio, y que no nos sirve nada más que para conseguir que el conductor llegue más lejos o recorra más distancia que la recorrería por sus propios medios.

Aunque el proceso de modulación en sí es tan sencillo como acabamos de exponer, al analizar la señal modulada que se va a transmitir en forma de onda electromagnética, aparecen algunas complicaciones. Pongamos un ejemplo para verlo más claro, suponiendo que la portadora tiene una frecuencia de 600 KHz., y que la modulamos en amplitud con una señal de audio de 5 KHz.

Se puede demostrar matemáticamente, que la onda modulada es equivalente a tres señales: una de 600 KHz, otra de 605 KHz y una tercera de 595 KHz, de tal forma que la frecuencia de la primera coincide con la de la portadora, mientras que las otras dos se obtienen sumando y restando, respectivamente, la frecuencia de la portadora y de la moduladora. Si la frecuencia de la onda moduladora (la señal de audio) fuera de 20 KHz, la señal modulada seguiría siendo equivalente a otras tres, ahora con unas frecuencias de 600, 620 y 580 KHz.

Si ahora queremos transmitir una reproducción musical (que está compuesta por señales de frecuencia de hasta 20 KHz), la onda modulada se podrá descomponer en una componente central de 600 KHz y en dos conjuntos de frecuencias laterales o dicho de otra forma en dos bandas laterales, una a cada lado de la frecuencia central, ocupando cada una, una anchura de 20 KHz.

Ahora bien, cabe preguntarse ¿cuánto se puede hacer variar la amplitud? Indudablemente los extremos de tal variación son claros: uno de ellos corresponderá al instante en que la amplitud de modulación alcance la amplitud de la portadora; el otro, por el contrario debemos situarlo cuando no se produce modulación.

En el primer caso se dice que la profundidad o porcentaje de modulación es de 100 %, mientras que el segundo es cero. Entre ambos extremos se obtienen todos los valores posibles de profundidad de modulación. Si se pretende modular la portadora con una profundidad superior al 100%, la forma de la señal será diferente de la original, produciendo una distorsión en la parte receptora.

Una emisora de radio que transmita en AM es indudable que consumirá una cierta potencia en transmitir las ondas de radio. Puesto que la amplitud media de la portadora es siempre la misma, la potencia invertida en enviar la portadora ha de ser necesariamente constante.

Sin embargo, la potencia necesaria para las bandas laterales dependerá de la profundidad de modulación. Cuando dicha profundidad sea nula, no habrá modulación, ni bandas laterales, no siendo necesaria por tanto, potencia alguna para las citadas bandas laterales.

Sin embargo, con una profundidad de modulación del 100 %, se producirán dos bandas laterales iguales, siendo inevitable emplear en cada una de ellas una potencia equivalente a la cuarta parte de la invertida para la portadora.

Para valores intermedios comprendidos entre el 0 y el 100 % de profundidad de modulación, la potencia requerida en cada banda lateral está comprendida entre el 0 y el 25 % de la empleada para la portadora. Por ejemplo, para una profundidad de modulación del 30 %, la potencia necesaria para cada banda representa poco más del 2 % de la correspondiente a la portadora.

Si tenemos en cuenta que la “**información**” que queremos transmitir va contenida en las bandas laterales, será fácil comprender que estamos derrochando energía, puesto que apenas un 5 % de la potencia gastada por un transmisor de AM que module con un 30 % de profundidad va a sernos útil, ya que el otro 95 % se emplea para la portadora que no contiene información. En el mejor de los casos (modulación al 100 %), la potencia empleada para las dos bandas representa tan sólo el 50 % de la consumida por la portadora.

Para conseguir un mejor rendimiento en el proceso de transmisión, puede suprimirse la portadora, emitiéndose tan sólo las dos bandas laterales. Esto da lugar a la denominada modulación en doble banda lateral, o **DSB** (siglas inglesas correspondientes a “**Double Side Band**”).

Incluso puede suprimirse también una de las bandas laterales, ya que toda la información que queremos transmitir va contenida por igual en ambas; de esta forma se obtienen los sistemas de modulación en banda lateral única, conocidos de forma abreviada como **BLU** o **SSB** (del inglés: “**Single Side Band**”). Indudablemente, este último sistema es el de mayor rendimiento, ya que toda la potencia consumida en el transmisor se aprovecha para transmitir la información.

Todavía existen otros tipos de modulación de amplitud, con supresión y sin supresión de portadora. Entre ellos cabe destacar el utilizado por la señal de imagen de TV (video). Aparte de la ventaja que supone un mayor rendimiento

en la transmisión, los sistemas que suprimen una parte o la totalidad de una de las bandas laterales ahorran también ancho de banda ocupado, ya que sólo se emplea el suficiente para poder transmitir la información útil.

Diagramas en bloque

Una forma sencilla de entender el funcionamiento de los circuitos, es estudiarlos agrupando las distintas etapas que forman el circuito en bloques, sin necesidad de dibujar los componentes que componen cada etapa. Si se desea conocer más a fondo sobre una determinada etapa, sí es conveniente estudiar el circuito electrónico. Para nuestro caso, es suficiente con entender los diagramas en bloque de los distintos transmisores y receptores.

Emisión en telegrafía (a1a)

El primer diagrama que estudiaremos es el de un transmisor de telegrafía (Figura 1). La primera etapa que observamos es la del:

OSCILADOR: Su función es la de generar la señal de **Radio Frecuencia (RF)** que, en este caso será del valor de la frecuencia a transmitir, es decir que si deseo operar en la banda de 80 metros el oscilador deberá generar una señal de RF que este comprendida entre 3.500 a 3.750 KHz.

SEPARADOR: Esta etapa tiene la función de aislar el oscilador del amplificador, de forma que aunque la carga sea variable para el separador, el oscilador tenga una carga constante, lo que garantiza que la frecuencia no variará.

AMPLIFICADOR: Esta etapa es, como su nombre lo indica un amplificador de RF, que lleva el nivel de la señal que entrega la etapa separadora a un nivel adecuado para poder emitirla por la antena.

MANIPULADOR: interrumpe la portadora generada por el oscilador y de esta manera, emitimos portadora o no la emitimos.

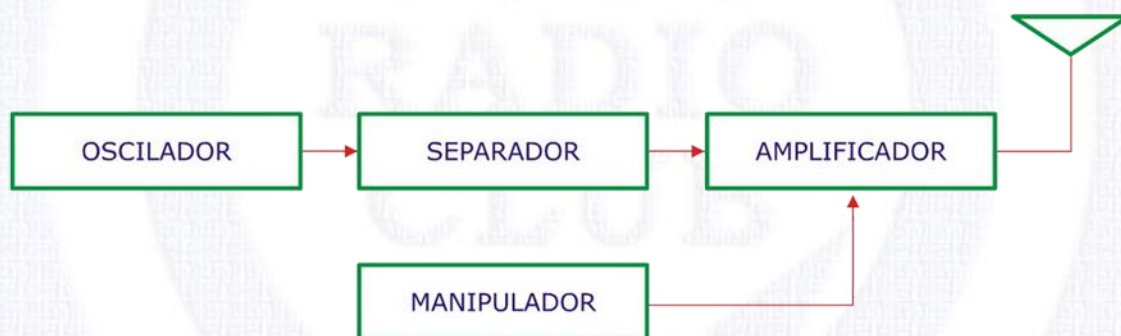


Figura 1: Diagrama de transmisor de telegrafía

Emisión en Amplitud Modulada (A3E)

Se trata de uno de los sistemas más antiguos de emisión (Figura 2).

La concepción más normal de las emisoras de AM es la de variar la tensión de alimentación del amplificador de RF al ritmo de la señal de audio.

Las emisoras con esta concepción constan de las siguientes etapas:

a) Sección de Radiofrecuencia:

- **OSCILADOR:** Es el corazón de la emisora, puesto que genera la señal de RF.
- **SEPARADOR:** También llamado primer amplificador, que recibe la señal del oscilador y la amplifica; al mismo tiempo, establece una cierta independencia entre el oscilador y el resto de la emisora para que las variaciones de carga de la misma no repercutan en la frecuencia del oscilador, que tiene que ser rigurosamente constante.
- **AMPLIFICADOR DE POTENCIA DE RF:** que es el que proporciona la energía que ha de pasar a la antena

b) Sección de Audiofrecuencia

- **MICRÓFONO:** Es el transductor que genera la señal de baja frecuencia (BF)
- **PREAMPLIFICADOR DE AUDIO:** encargado de aumentar el nivel de la señal de audio hasta el valor adecuado para excitar al amplificador de potencia del paso siguiente.
- **AMPLIFICADOR DE POTENCIA DE AUDIO:** cuya misión es la de obtener la potencia necesaria para excitar adecuadamente el amplificador de RF, en el que está incluido el proceso de modulación.

La principal ventaja de la amplitud modulada es su sencillez constructiva, razón por la cual los radioaficionados recién iniciados prefieren este modo para armar su primer transmisor.



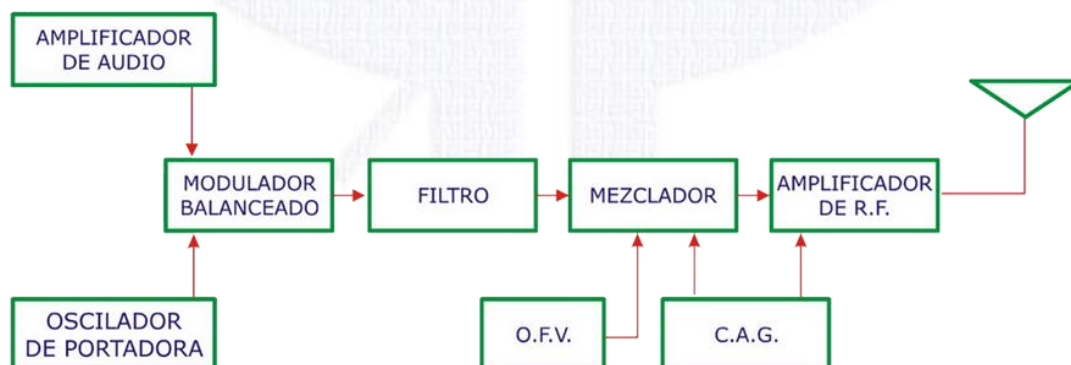
Figura 2: Diagrama de AM (A3E)

Emisión en Banda Lateral (J3E)

La emisión en banda lateral, no es más que una variante de la emisión en amplitud modulada, siendo la señal emitida (en AM), de un ancho de banda de unos 6KHz. Como ya sabemos, la portadora no contiene ninguna información, siendo solamente las bandas laterales las fuentes de la misma. La portadora no es necesaria, pero si el receptor no recibe las bandas laterales conjuntamente con la portadora, la recepción resulta ininteligible. Esto puede subsanarse, y de hecho así se hace, incorporando la portadora en el receptor, por introducción de la misma en el detector, el cual, por recibir la mezcla de la señal emitida y la de la portadora propia, generada por un oscilador de batido, recibe el nombre de detector de producto. Uno de los sistemas más simples para la supresión de la portadora consiste en mezclar la portadora de radiofrecuencia de frecuencia fija, generada por el **OSCILADOR DE PORTADORA**, con las señales de audio provenientes del amplificador de audio. Esta mezcla se realiza en la etapa denominada, **MODULADOR EQUI-LIBRADO O BALANCEADO**.

Además de esto, el modulador balanceado debe mantener dentro de un mínimo los productos indeseables del proceso de modulación. A la salida del modulador balanceado, tenemos una señal de doble banda lateral, para eliminar una de las dos bandas laterales, se utiliza un filtro mecánico de características especiales. A la salida de esta etapa ya tenemos señal de banda lateral única.

Es necesario aclarar que el modulador balanceado no trabaja a la frecuencia de emisión, sino que se emplea una frecuencia fija, por lo tanto para obtener la frecuencia de trabajo mezclamos la señal de banda lateral con una señal proveniente de un **OSCILADOR DE FRECUENCIA VARIABLE (OFV)**, y la resultante se amplifica y se inyecta a la antena para su emisión. Insistimos en que la principal ventaja de la transmisión de BLU es que se aprovecha mejor la potencia de emisión, además de la reducción del ancho de banda, que es de unos 3 KHz aproximadamente (la mitad del AM).



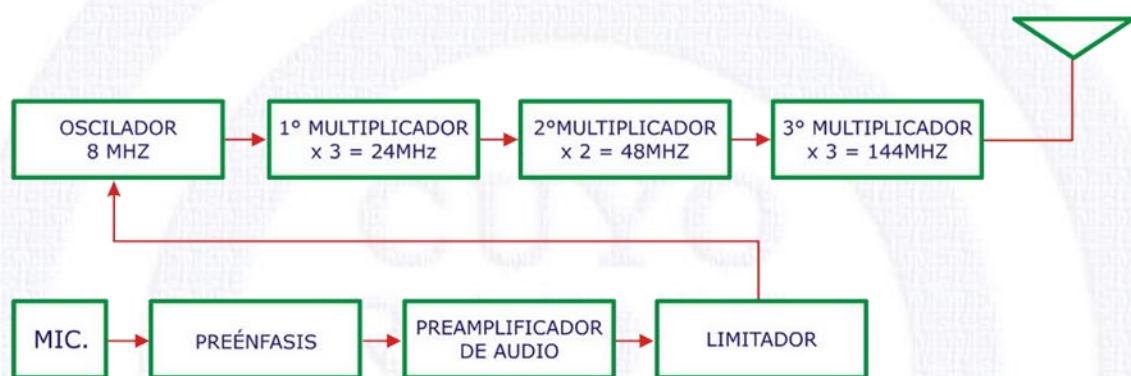
Emisión en Frecuencia Modulada (F3E)

En un **emisor de FM**, la señal de audio del micrófono pasa por un filtro que atenúa los agudos (llamado **circuito de preénfasis**), con lo que resultan amplificados los graves.

La justificación de este proceso, extraño a primera vista es la siguiente: debido a los circuitos multiplicadores de frecuencia que existen en este transmisor, las señales de, por ejemplo, 100 Hz, después de ser multiplicadas varias veces producirían una desviación de 400 Hz, mientras que una señal de 2 KHz produciría una desviación de 8 KHz, y una de 3 KHz produciría una de 12 KHz, que es casi el límite ya que es difícil de encontrar una voz que supere la frecuencia de 3 KHz. Es evidente que los sonidos graves producen menos desviación que los agudos y por lo tanto modularán menos la onda portadora, lo que significa que cuando la señal sea débil no oiremos los graves y si los agudos con lo que se perdería calidad. Luego sigue un preamplificador de audio, al que le sigue un **LIMITADOR**, que amplifica las señales del micrófono, impidiendo la existencia de señales que superen el umbral del limitador. Esta señal se mezcla con la de un oscilador variable de 8 MHz. La frecuencia obtenida a la salida del oscilador se multiplica en varios pasos hasta conseguir la frecuencia de trabajo. Para el caso de la **figura 7**, para pasar de los 8 MHz del oscilador a los 144 MHz a la salida, ha sido necesario multiplicar la frecuencia inicial por 18, para lo cual se han empleado dos triplicadores y un doblador de frecuencia.

El ancho de banda de una transmisión de **FM** en la banda de 2 metros (144 MHz) es de 16 KHz.

Por último la señal se aplica a un **AMPLIFICADOR** para obtener la potencia necesaria en la antena.



Receptores de radio

El receptor de radio es el encargado de recibir y de "interpretar" las ondas emitidas por un transmisor. Es el que debe extraer la información que se transmite por cualquiera de los medios vistos anteriormente. Desde los primeros aparatos, hasta nuestros días, los receptores de radio, han sufrido una profunda evolución, pasando de ser unos equipos de bastante complejidad y excesivamente voluminosos, a los más modernos receptores de pequeño tamaño, capaces de ser transportados en el bolsillo.

Receptor elemental

Un receptor de radio de AM, en su versión más sencilla, se compone de las siguientes partes:

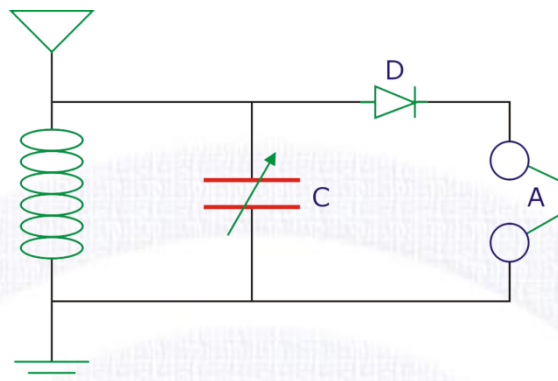
- **Antena**
- **Tierra**
- **Circuito de sintonía**
- **Circuito detector**
- **Transductor de sonido**

Un sencillo aparato, compuesto únicamente de los elementos que hemos enumerado, es el que se conoce con el nombre de receptor de "**galena**", apelativo que se remonta a la primera época en la que se utilizó dicho material como base del circuito detector.

El sistema aún sigue teniendo vigencia, sobre todo entre los que recién comienzan a experimentar en este apasionante mundo de las ondas, habiéndose sustituido el detector de "galena" por un diodo de germanio, que posee características similares a su antecesor.

Su funcionamiento es muy simple: Las señales de radio son captadas entre la antena y la tierra y enviadas al circuito de sintonía, el cual está formado por la bobina L y el condensador variable C.

Mediante este dispositivo se seleccionará la emisora deseada, ya que el circuito presentará una baja impedancia para el resto de las frecuencias, derivándolas hacia tierra. La señal obtenida llega al diodo detector D y se aplica al transductor, que es el elemento más simple capaz de transformar la señal detectada en sonido audible, es decir, unos auriculares (indicados con la letra A).



Lo más destacable a considerar en este receptor es la ausencia de amplificadores, lo que significa que el mismo no requiere ningún tipo de alimentación.

Explicaremos las etapas que forman parte de un receptor de AM, FM y BLU.

Recepción de AM

Una vez que es captada por la antena, las señales de radio se llevan a un **AMPLIFICADOR DE RADIOFRECUENCIA**, de esta forma se aumenta la sensibilidad del sintonizador, ya que permite hacer llegar a la etapa siguiente señales que son muy débiles en la antena. No obstante, si las señales de radio son muy fuertes, pueden resultar perjudicial darle mayor ganancia (hasta el punto de distorsionarla), por lo que se suele incorporar un mando que permite eliminar la amplificación de este paso, para utilizarlo sólo con las señales débiles. Un buen amplificador de RF, nos brindará un receptor con buena sensibilidad, entendiéndose por sensibilidad, la capacidad que tiene un receptor de escuchar las señales muy débiles.

OSCILADOR LOCAL: Genera una señal cuya frecuencia se varía con el mando de sintonía. Dicha frecuencia es siempre mayor que la de la señal que envía la etapa anterior, en una cantidad igual al de la llamada frecuencia intermedia, cuyo valor es elegido previamente por el fabricante. Aunque no existe normalización alguna de dicho valor, los más usuales son los de 455 KHz para AM y 10,7 MHz para FM.

MEZCLADOR: Se encarga de mezclar la señal procedente de la etapa de RF con la que genera el oscilador local. El proceso de la mezcla da como resultado la aparición de dos señales, una cuya frecuencia es igual a la suma de las respectivas frecuencias que se mezcla, mientras que la otra tiene una frecuencia equivalente a la diferencia de las mismas. El fenómeno de mezclar dos señales se conoce como "**heterodinar**", por lo que a éstos receptores se los conoce con el nombre de "superheterodinos".

La única señal que se aprovecha es la de frecuencia diferencia, que es precisamente la frecuencia intermedia antes mencionada, apelativo debido a que dicha frecuencia tiene un valor comprendido entre la de radio frecuencia que se sintoniza y la de audio (**o baja frecuencia**), que es el producto que deseamos obtener para aplicarlo al parlante. El proceso llevado a cabo en esta etapa no varía para nada la información que lleva impresa la onda de radio, sino tan sólo su frecuencia.

Por ejemplo: Si recibimos por la antena una señal de 680 KHz, se mezclará luego con una señal generada en el oscilador local de 1135 KHz, a la salida del mezclador tendré dos señales: una es la suma, es decir 1815 KHz, y otra es la resta o sea 455 KHz. La frecuencia suma se descarta y sólo se usa, como dijimos, la frecuencia diferencia. Si ahora recibimos una señal de 720 KHz, tendré que generar con el oscilador local una frecuencia tal que la diferencia de 455 KHz. Y así para cualquier frecuencia que sintonice. Esto se consigue usando un capacitor variable doble, de manera que cuando se gire el mando para sintonizar una emisora, se varía al mismo tiempo el capacitor que gobierna al oscilador local. De esta manera siempre se generará una señal cuya frecuencia diferirá 455 KHz de la sintonizada. Por lo tanto las etapas siguientes se ajustan a la frecuencia intermedia y no a la frecuencia de la señal recibida en la antena. Esto trae una gran ventaja que: mejora la selectividad del receptor.

Nos queda definir la Selectividad como la capacidad que tiene el receptor de evitar las interferencias de estaciones potentes y cercanas a la frecuencia de trabajo.

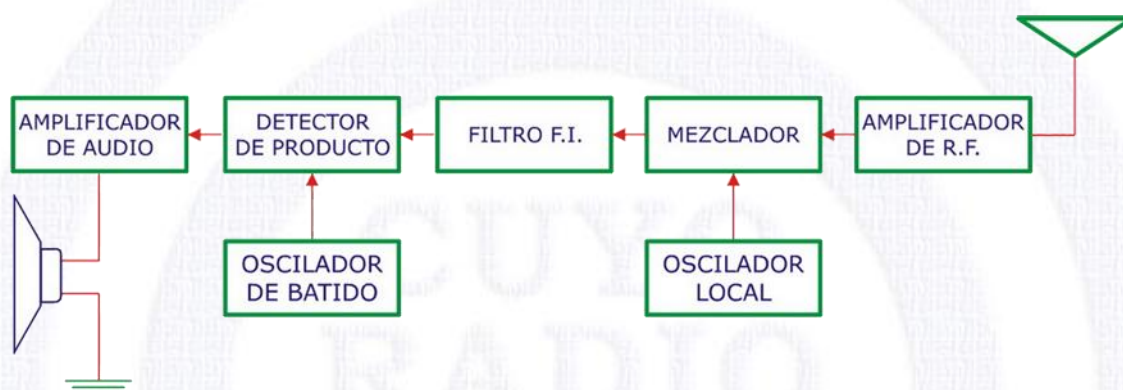
AMPLIFICADOR DE F.I.: En esta etapa se amplifica la señal hasta el nivel requerido por la etapa siguiente.

DETECTOR: También llamado **DEMODULADOR**. Es la etapa encargada de extraer la modulación contenida en la señal. Se llama detector porque permite detectar la información contenida en la señal de radio. Aquí se elimina la portadora y queda la señal de audio.

AMPLIFICADOR DE AUDIOFRECUENCIA: La señal de la etapa anterior es muy débil para excitar un parlante, por lo que debe agregarse esta etapa para que el nivel de audiofrecuencia llegue al valor adecuado para excitar un parlante o auriculares.

Recepción en FM

El diagrama en bloque de la **figura 6**, nos muestra un receptor de FM, en el que vemos casi las mismas etapas que el receptor de AM. Excepto que existe una etapa llamada **LIMITADOR** que no existe en el receptor de AM. Como sabemos, la modulación en una onda de FM va contenida en las variaciones de su frecuencia y no en las de su amplitud. La misión de esta etapa es la de recortar o "limitar" la amplitud de la señal, para eliminar cualquier variación que pudiera existir en dicha amplitud y que podría dar lugar a distorsiones en la etapa siguiente. Aunque la portadora está modulada en frecuencia, puede llegar al receptor acompañada de otras que la modulan en amplitud, y cuyo resultado es la introducción de parásitos indeseables que producen ruidos molestos en la recepción.



El demodulador de FM

La detección de señales moduladas en frecuencia se lleva a cabo por medio de un doble proceso:

- 1 - Las señales de FI, moduladas en frecuencia y de amplitud constante, se aplican a un circuito llamado discriminador, el cual proporciona a su salida variaciones de amplitud proporcionales a las variaciones de frecuencia, es decir, que la señal de salida está modulada tanto en frecuencia como en amplitud.
- 2 - Las señales proporcionadas por el discriminador se someten a un proceso de rectificación y filtrado, similar al de AM, con el que se detectan las variaciones de amplitud que constituyen la señal de baja frecuencia o señal de audio.

Así pues un detector de FM consta de un discriminador y de un detector de AM, aunque se suele dar al conjunto de los dos circuitos el nombre de discriminador.

Receptor de BLU

Un receptor de **BLU** por el tipo de señal que debe manejar utiliza etapas especiales, el diagrama en bloque se muestra en la **figura 5**.

Las primeras etapas son las ya vistas. La particularidad de estos receptores radica en el detector o demodulador que recibe el nombre de **DETECTOR DE PRODUCTO**.

La señal de BLU que sale de la etapa de FI, carece de portadora, por ese motivo para poder demodularla, se debe inyectar la portadora.

La portadora, que la generamos en el receptor con el **OSCILADOR DE PORTADORA**, se inyecta en el detector de producto y se mezcla con la señal de FI. La salida del detector es señal de audiofrecuencia que se aplica a la etapa amplificadora de audio.

Otra etapa que poseen estos receptores es la de **CONTROL AUTOMÁTICO DE GANANCIA (CAG)**.

La misión de esta etapa es la de tomar una muestra de la señal de audio a la salida y reinyectarla en el amplificador de radiofrecuencia. Si la señal a la salida es muy grande, inyecta en el oscilador una tensión para que disminuya la ganancia del oscilador.

Por el contrario, si la señal a la salida es muy baja, inyecta en el oscilador una tensión tal, que hace que aumente la ganancia del oscilador.

