

elektor

electrónica: técnica y ocio

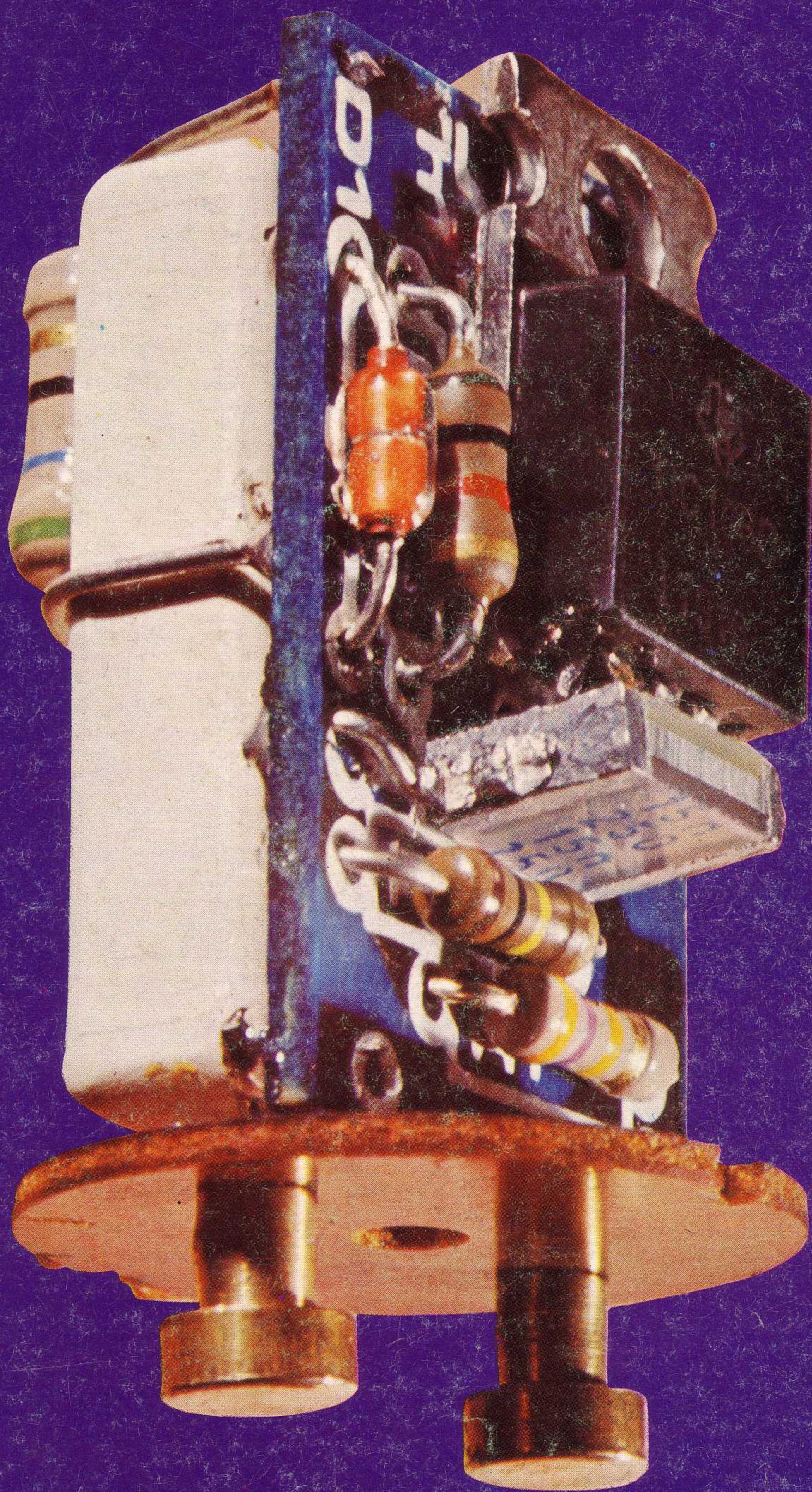
N.º 31
diciembre 1982

200 Ptas.

**control
de tubos
fluorescentes**

**sistema
de telefonía
interior**

**receptor BLU
curso de BASIC (4)**



Selektor	12-01
Ultimas novedades en los sectores de audio y video.	
La BLU: de la B a la U	12-03
Un pequeño adelanto: BLU son las siglas de «Banda Lateral Unica»... pero, por supuesto, no hay que quedarse ahí; a lo largo del artículo le desvelamos lo que hay detrás de estas enigmáticas siglas.	
Regulador universal	12-06
Los reguladores de intensidad luminosa para lámparas incandescentes se han convertido en verdaderos bienes de consumo, circunstancia que contrasta con la casi total ausencia de reguladores para tubos fluorescentes. Como quiera que nuestra debilidad es, precisamente, encontrar soluciones... ¡aquí está nuestro regulador universal!	
Detector de gas	12-11
¡Una sensible y eficaz «nariz» electrónica capaz de protegerle contra los peligrosos efectos de los escapes de gas!	
Sistema de telefonía interior	12-15
En los paseos por el rastro o por los mercadillos de oportunidades, es frecuente encontrarse con receptores telefónicos usados. A partir de ahora es posible que estos saldos le merezcan mayor atención. En esta ocasión, le contamos como construir una instalación interior con hasta 9 puestos telefónicos.	
Curso de BASIC (4.ª parte)	12-21
¡Hemos llegado a la cuarta lección!... en la que seguimos ampliando nuestro vocabulario lingüístico y, por supuesto, poniendo en práctica nuestros conocimientos de programación.	
Teclado sensorial	12-29
Una alternativa económica a los tradicionales teclados de contactos mecánicos.	
Intermitente electrónico	12-33
¡No todo van a traerlo los reyes!... para usted queda la tarea de darle un mayor «encanto» y atractivo a los automóviles de juguete.	
Receptor BLU de onda corta	12-35
De la teoría a la práctica a través de la construcción de un receptor compacto y de calidad.	
Cebador electrónico para fluorescentes	12-41
... la solución electrónica para un encendido sin centelleos.	
Introducción a la familia DMOS	12-44
Cada día aparecen nuevos transistores FET: los VFET, HEXFET, TMOS, SIPMOS... ¡Una verdadera sopa de letras! A lo largo de este artículo vamos a tratar de digerir la consabida sopa, examinando a los transistores FET en general y prestando especial atención a la nueva rama constituida por la familia DMOS.	
Mercado	E-11
Anuncios breves	E-13
Indice de anunciantes	E-13

sumario

SUMMAR
SUMMA
SUMM
SU



El personal de elektor
les desea ...

¡Feliz Navidad!



La modulación

En principio, un emisor no es otra cosa que un oscilador capaz de generar una señal de frecuencia relativamente elevada. Esta señal se envía «al aire» por medio de una antena. Como indica la figura 1a, en la mayor parte de los casos los emisores son algo más complejos y contienen varios componentes además del oscilador. Examinemos el diagrama de bloques de la figura 1a. Una señal de oscilador con una frecuencia de 4 MHz, por ejemplo, entra en un amplificador en donde se intensifica desde un par de mW a 100

La denominación de onda continua no es totalmente exacta ya que, en este caso, la onda no es continua sino «troceada» en pequeños fragmentos por el codificador (o manipulador). Esta forma de modulación recibe el nombre de *modulación por impulsos*. Hay otras clases de modulación. La figura 1b ilustra una de las más conocidas. En este caso, el codificador se ha sustituido por un montaje más gradual, y por consiguiente más flexible, que hace variar la tensión de salida de un amplificador de potencia en proporción a una señal microfónica. En el diagrama de bloques de la figura 1b, una señal de 1 kHz se ha seleccionado como frecuencia de modulación y se constata claramente que la amplitud de la señal de salida tiene tendencia a tomar la forma de esta señal de modulación de 1 kHz. Por ello, a esta forma de proceder, se le ha dado el nombre de *modulación de amplitud* (AM).

Al ser simétrica la modulación, tendremos en la salida una señal simétrica con un valor máximo que es el doble de la magnitud correspondiente de la onda portadora no modulada.

Otra forma de modulación muy conocida es la *modulación de frecuencia* (FM). No vamos a entrar en detalles pero ilustramos su principio básico en la figura 1c. Para esta clase de modulación, no es la amplitud de la portadora la que varía, sino su frecuencia. La señal producida por el micrófono se transforma en tensión de control, que sirve para desplazar la frecuencia del oscilador, en una pequeña magnitud, hacia arriba y hacia abajo. Puede constatarse que la amplitud de la señal de salida se mantiene bastante constante.

Naturalmente, hay otros tipos de sistemas de modulación aparte de los mostrados en la figura 1. La modulación de fase (PM) y la modulación de frecuencia de banda estrecha (NBFM = Narrow Band FM) forman parte de la familia de la FM. Entre los miembros de la familia de AM figuran, entre otros, la modulación de amplitud con portadora suprimida (DSB-SC: Double Side Band-Suppressed Carrier) y la BLU (Banda Lateral Unica). A esta última vamos a dedicarle este artículo.

Las bandas laterales

Los sistemas de modulación DSB-SC y BLU se conocen desde hace tiempo. Los principios fundamentales son los siguientes: Si un emisor AM, como el de la figura 1b, modula, a una frecuencia de audio de 1 kHz, una onda portadora de 4 MHz (= 4.000 kHz) se producen dos bandas laterales (armónicos), una a 3.999 kHz y otra a 4.001 kHz. En la figura 2a se muestra el aspecto de una señal de esta naturaleza cuando se observa bajo la «lupa» de un analizador de espectro. Las dos bandas laterales son la imagen perfecta la una de la otra y contienen la misma información. La portadora, en sí misma, no lleva información alguna, si bien, utiliza la parte más grande de la energía de emisión, tal como se muestra en la figura 2a. En los primeros días de la era radiofónica, alguien tuvo la feliz idea de suprimir la onda portadora en su totalidad y canalizar la energía de transmisión en las bandas laterales que contienen la información. Este método se conoce como DSB-SC que significa

la BLU: de la B a la U

una rápida y demoledora incursión en la tecnología de los transceptores!

En este mismo número de nuestra revista presentamos un artículo sobre la construcción de un receptor BLU, pero es evidente que al oír el término BLU no todos nuestros lectores saben «de qué va». Podríamos curarnos en salud diciendo que BLU es la abreviatura de Banda Lateral Unica que corresponde a SSB (Single Side Band) en la nomenclatura anglosajona, pero no queremos quedarnos ahí y pretendemos descubrir lo que hay tras los bastidores de esta enigmática palabra.

W. A continuación, pasa a través de un filtro que la «limpia» suprimiendo cualquier componente indeseable (interferencia, etc). El filtro asegura también que la impedancia del amplificador y la resonancia de la antena estén bien adaptadas.

La señal que se transmite efectivamente se conoce como una onda portadora. Aun cuando un receptor adecuado fuera capaz de captarla, la onda portadora no proporcionaría nada inteligible. Si se quiere transmitir una información desde el emisor al receptor, es preciso, de una forma u otra, añadir la información a la portadora: la portadora sufre una «modulación». Como su nombre indica, una onda portadora sirve para transportar información.

La forma más elemental de efectuar una modulación es utilizar un interruptor como se ilustra en la figura 1a. Con su empleo, se puede interrumpir periódicamente la portadora emitida y, si se emplea un código bien determinado (el Morse, por ejemplo), se hace posible transmitir una información. Aunque, salvando las distancias, el resultado se parece bastante a la transmisión de mensajes que efectuaban los indios por medio de señales de humo.

El conmutador de la figura 1 a puede considerarse como el codificador de un transmisor de onda continua (CW = continuous wave).

Double Side Band (banda lateral doble) — Suppressed Carrier (portadora suprimida). El resultado ilustrado en la figura 2b es que la potencia de salida efectiva (portadora de información) es doble que la producida en AM.

Algunas investigaciones suplementarias han hecho dar el paso siguiente y han permitido llegar a la BLU (banda lateral única). Como ya hemos indicado, las dos bandas son idénticas y, por consiguiente, se puede suprimir una de ellas sin que por ello se pierda información, ni se produzca inconveniente alguno. Como se muestra, muy elementalmente, en la figura 2c, la potencia eficaz de la BLU es también el doble con respecto a la banda lateral doble. Si se efectúa una comparación entre las figuras 2c y 2a, salta a la vista que la potencia producida por el emisor se utiliza de manera más «eficaz» y útil que en el modo AM.

BLU: los pros y los contras

No es sorprendente que la BLU sea el sistema de modulación más frecuentemente utilizado en el dominio de las ondas cortas. Los radioaficionados que operan en estas bandas utilizan casi exclusivamente esta clase de modulación.

La BLU emplea la energía de forma más eficaz que muchos otros modos de modulación y, en consecuencia, el «radio de acción» del emisor es claramente más importante y el ancho de banda necesario sólo es la mitad del que se necesita para la AM.

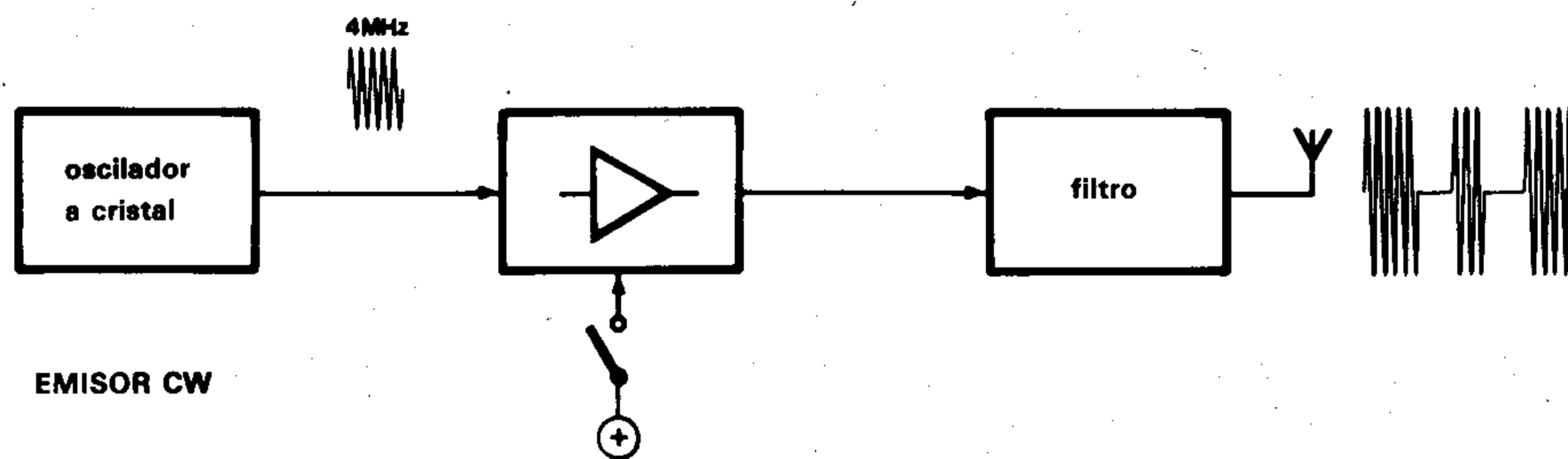
Si la audiofrecuencia es de 3.000 Hz como máximo (lo que permite transmitir la palabra), las bandas laterales se extenderán 3.000 Hz por encima y por debajo de la frecuencia de la portadora, lo que da un ancho de banda de 6 kHz. La banda lateral única de una señal de BLU sólo ocupa 3 kHz de la gama de transmisión. Ello significa que pueden comprimirse el doble de emisores en una banda determinada. En la práctica, este factor 2 es ampliamente superado pues la inexistencia de portadora hace imposible la interferencia recíproca de las portadoras de dos emisores próximos.

Lamentablemente, también hay un par de desventajas asociadas a la BLU. Una de ellas es que un emisor de BLU es más complejo y, por consiguiente, cuesta más caro que uno de AM. El principal inconveniente radica, sin embargo, en el lado del receptor. Habida cuenta de que este último debe sintonizarse en esta banda lateral única, la estabilidad en frecuencia es un criterio claramente más importante para un receptor de BLU que para su homólogo de AM. Y ello hace también más difícil la construcción y la puesta a punto de un receptor de «cosecha propia».

El receptor

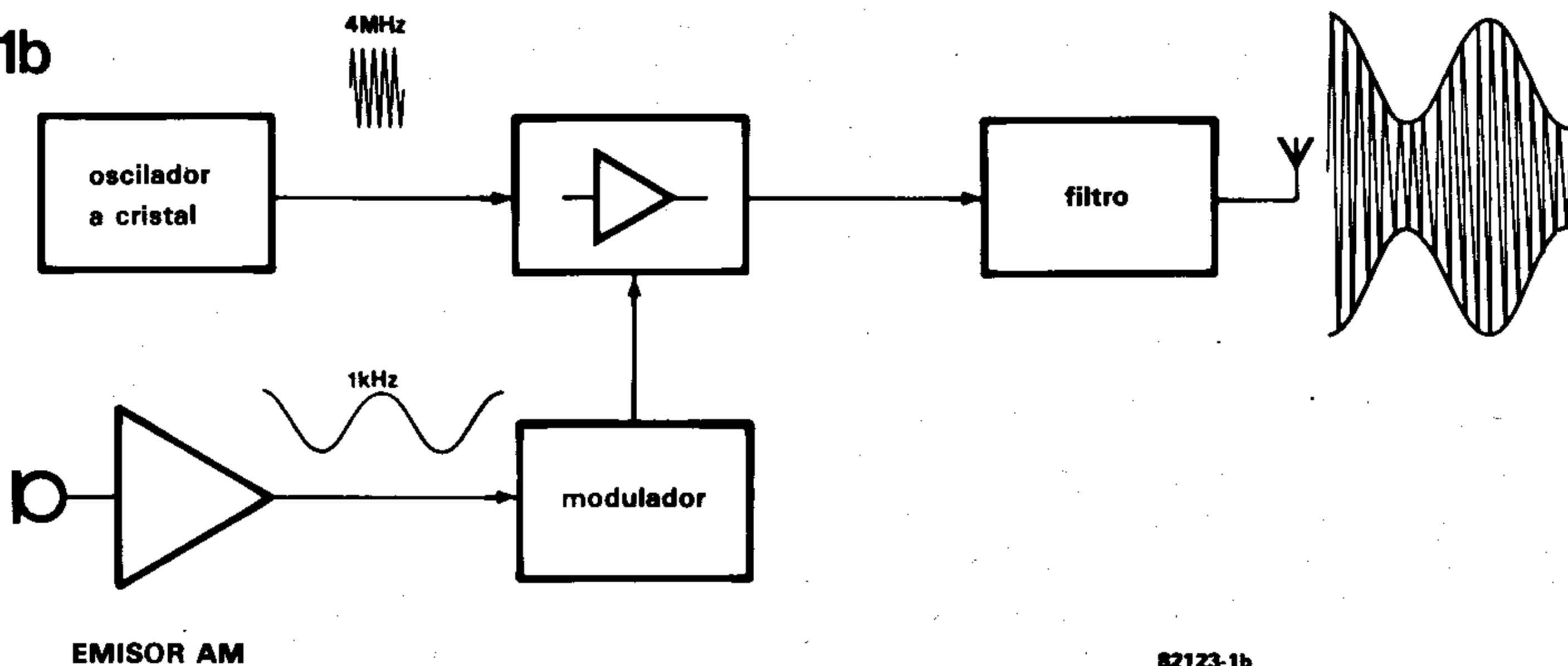
El receptor convierte la información difundida por un emisor en una forma que puede ser comprensible por los oyentes. Para poder realizar esta función han de cumplirse dos requisitos. Ante todo, debe ser capaz de seleccionar la emisora deseada de entre una enorme cantidad de otras señales «en el aire». A continuación, debe extraer la infor-

1a



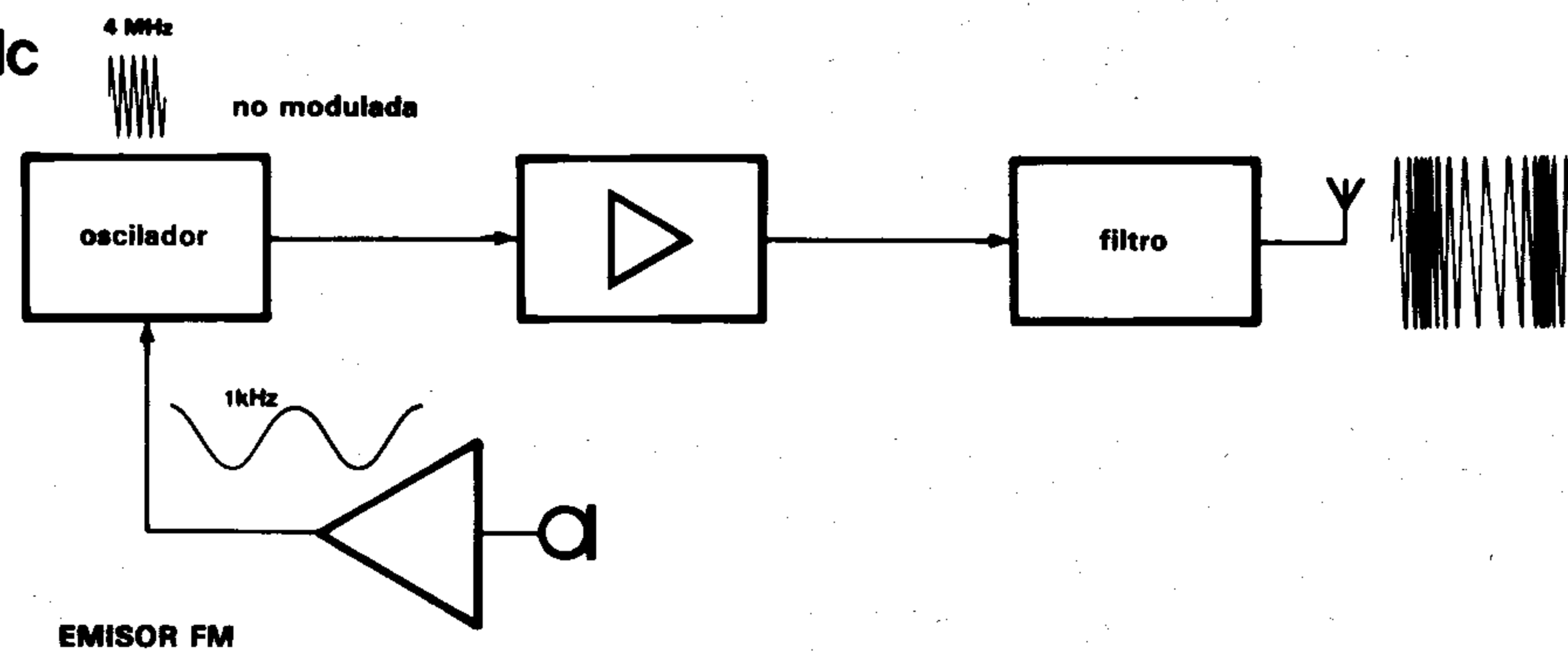
82123-1a

1b



82123-1b

1c



82123-1c

Figura 1. Existen diversas clases de modulación en el campo de los emisores. Ilustramos las más importantes: modulación por impulsos o de onda continua CW (1a), la modulación en amplitud AM (1b) y la modulación de frecuencia (1c).

mación contenida en la señal y traducirla en una señal acústica.

Cuando se trata de captar señales en AM, basta con un receptor a cristal. Este consta, como mínimo, de un circuito L-C sintonizable que permite la selección de la señal deseada, un diodo para recuperar la información de audiofrecuencia a partir de la señal de RF y, finalmente, auriculares para hacer audible la modulación.

Sin embargo, si se requiere una selectividad algo más importante y una sensibilidad superior a la normal, el receptor tendrá que incluir varios circuitos de selección y la señal habrá de amplificarse en A.F. Por esta razón, un receptor de AM simple suele tener la disposición de la figura 3, que corresponde al llamado receptor superheterodino. La señal de entrada se mezcla con la de un os-

cilador. El oscilador se ajusta a una frecuencia algo más alta que la señal de entrada y se sintoniza junto con el circuito de entrada. En consecuencia, la diferencia entre la entrada y las señales del oscilador se mantiene constante (455 kHz) en todo el margen de sintonía del receptor y esta señal de diferencia (que se llama frecuencia intermedia o F.I.) estará disponible a la salida del mezclador.

Ahora, la señal puede filtrarse completamente para proporcionar la selectividad requerida porque, al contrario que para el circuito de entrada, la señal de F.I. está a una frecuencia constante, de modo que los circuitos de filtro ya no han de sintonizarse para cada emisora particular. Después de las necesarias operaciones de filtrado y de amplificación, se detecta la señal de F.I. y se am-

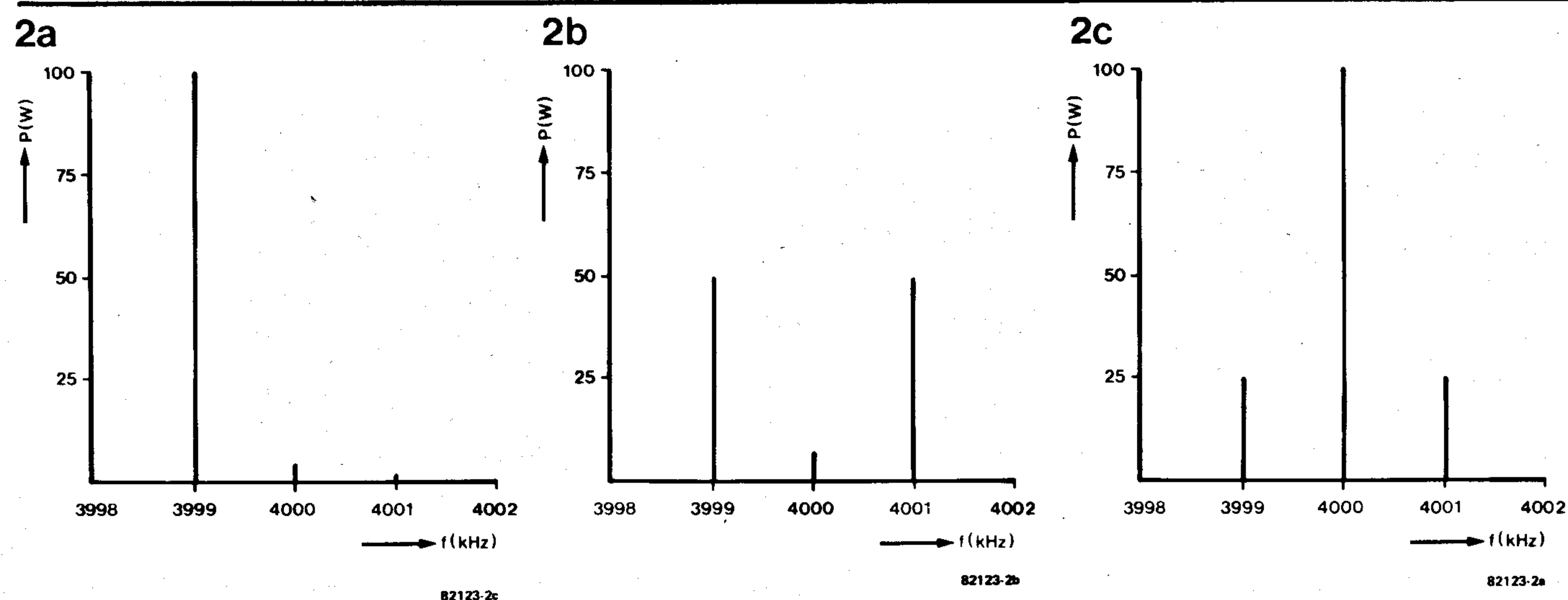


Figura 2. Espectro de frecuencia de una onda portadora de 4.000 kHz (4 MHz) modulada con una señal de audio de 1 kHz en modo AM (2a), en DSB-SC (2b) y BLU (2c), respectivamente.

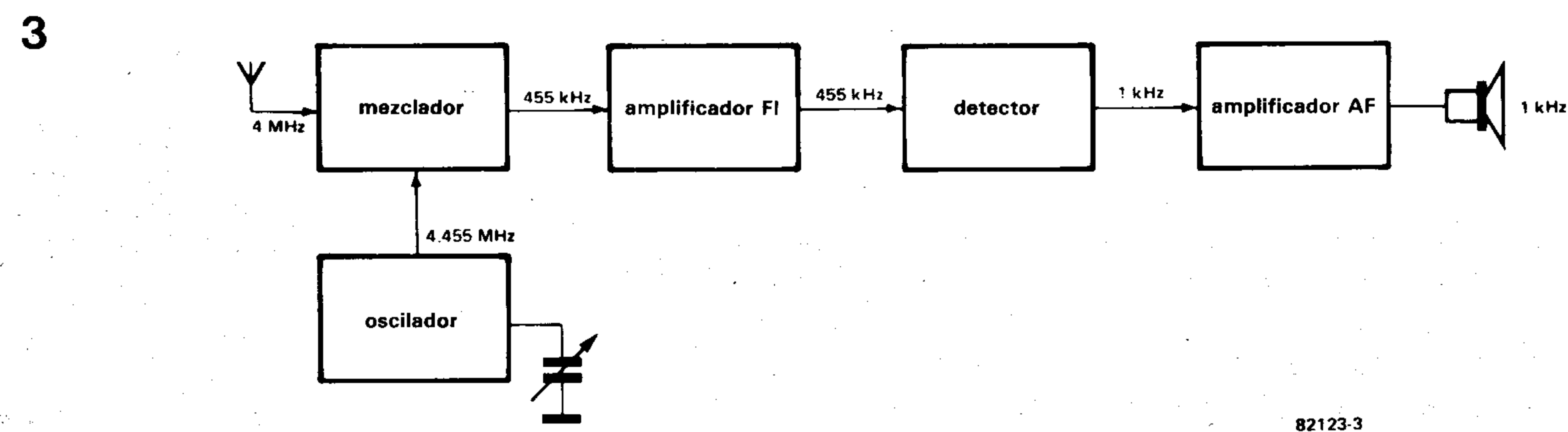


Figura 3. Diagrama de bloques de un receptor AM de onda corta ordinario.

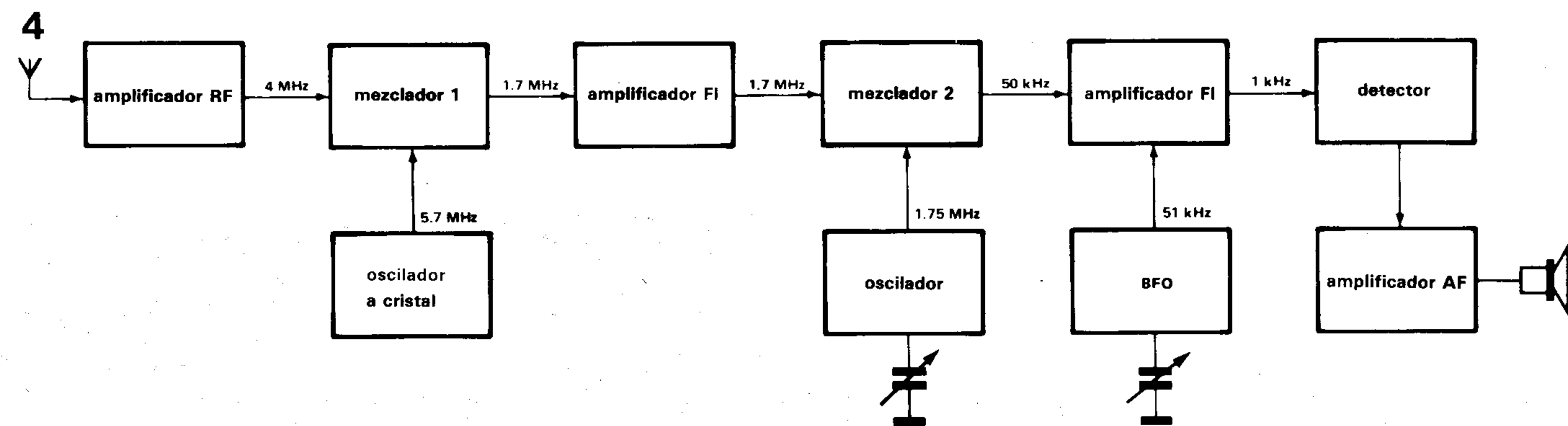


Figura 4. Diagrama de bloques de un receptor de BLU basado en el principio de «doble superheterodino».

plifica en B.F. Entonces, la modulación es audible en el altavoz.

Sabemos bastante por lo que respecta al receptor de A.M.

En principio, un receptor de BLU es muy similar a su contrapartida de AM, pero como las señales con las que trabaja el sistema de BLU tienen un ancho de banda muy pequeño, son mayores las exigencias de selectividad y el esquema ultrasencillo de la figura 3 ya no permitirá (salvo casos muy excepcionales) la obtención de resultados satisfactorios. El receptor del tipo BLU tendrá, con frecuencia, una configuración interna (reflejada en su diagrama de bloques) muy semejante a la descrita en la fig. 4. Se trata de un «doble superheterodino» que compren-

de dos mezcladores y dos frecuencias intermedias distintas.

Así llegamos al punto en que podemos vislumbrar una diferencia importante entre los sistemas de BLU y de AM. Si se quiere poder detectar la señal de F.I. se tiene necesidad de una portadora, que no existe en una señal de BLU. Será preciso encontrar un método para crearla en el receptor y añadirla a la señal. Por regla general, esta adición se hace inmediatamente antes de la detección de la señal con la ayuda de un oscilador sintonizable, el oscilador de batido (BFO: Beat Frequency Oscillator). Si este oscilador de batido se ajusta muy exactamente a la frecuencia que tendría la portadora (si existiera), el detector permite encontrar la frecuen-

cia de modulación original (1 kHz). Esta forma de proceder exige una gran estabilidad de frecuencia a través de todo el receptor y sobre todo en el BFO, pues la más ligera fluctuación da lugar a un desplazamiento de frecuencia en la señal de B.F.

El ajuste de la sintonía del BFO es muy importante y requiere un considerable cuidado y precisión pues, de no ser así, el timbre de cualquier voz puede alterarse para sonar tan aguda como la del «pato Donald» en un extremo de la escala o tan grave como la voz del más potente «bajo» en el otro extremo. Este es el motivo de que la utilización de un receptor de BLU exija paciencia y experiencia, lo que, a su vez, proporciona una cierta «sensibilidad táctil».

**un control de intensidad luminosa
para lámparas de incandescencia...
¡y tubos fluorescentes!**

regulador universal

Los reguladores para lámparas de incandescencia han llegado a ser bienes de consumo muy comunes, si bien, no son capaces de controlar tubos fluorescentes. En realidad, son los tubos los que no se adaptan a los reguladores para lámparas de incandescencia. En este artículo explicaremos cómo modificar los tubos para posibilitar dicha adaptación. Además, describiremos un circuito que puede ser controlado por medio de un dispositivo lógico, tal como el cronoprocador universal. El regulador en cuestión es idóneo para acuarios, terrarios y pajareras, ya que imita satisfactoriamente la luz solar diurna y hace que los animales «olviden» que están en interiores.

Antes de tratar del regulador propiamente dicho, quisiéramos aclarar algunos malentendidos. Suele decirse que los reguladores permiten economizar una cantidad considerable de energía. Ciertamente cuando la potencia de la lámpara sólo se utiliza a la tercera o cuarta parte de su magnitud, la corriente consumida es más pequeña. Pero el rendimiento, es decir, la relación de la luminosidad producida a la corriente absorbida, es inferior.

Dicho de otro modo, es mucho más económico sustituir una lámpara potente por otra de menor potencia, pero de rendimiento mayor que la primera utilizada con un regulador. A título de ejemplo ilustrativo: una lámpara de 100 W regulada hasta 40 W proporciona menos luz que una de 40 W plenamente encendida. En pocas palabras, puede decirse que una lámpara y un regulador

proporcionan una luz más cara que una lámpara de menor potencia y de aquí que el uso del regulador se llegue a considerar un lujo. No obstante, a veces, merece la pena gastar un poco más, pues los reguladores tienen grandes ventajas. Es ideal poder ajustar las luces para cada ocasión: para leer, mirar la televisión o pasar una velada agradable con amigos, etc., sin tener que cambiar las lámparas en cada momento. Piense en la «colección» de lámparas que deberíamos tener, y también en los animales criados en un medio artificial, que soportan mal el paso brusco de la «noche al día» y viceversa.

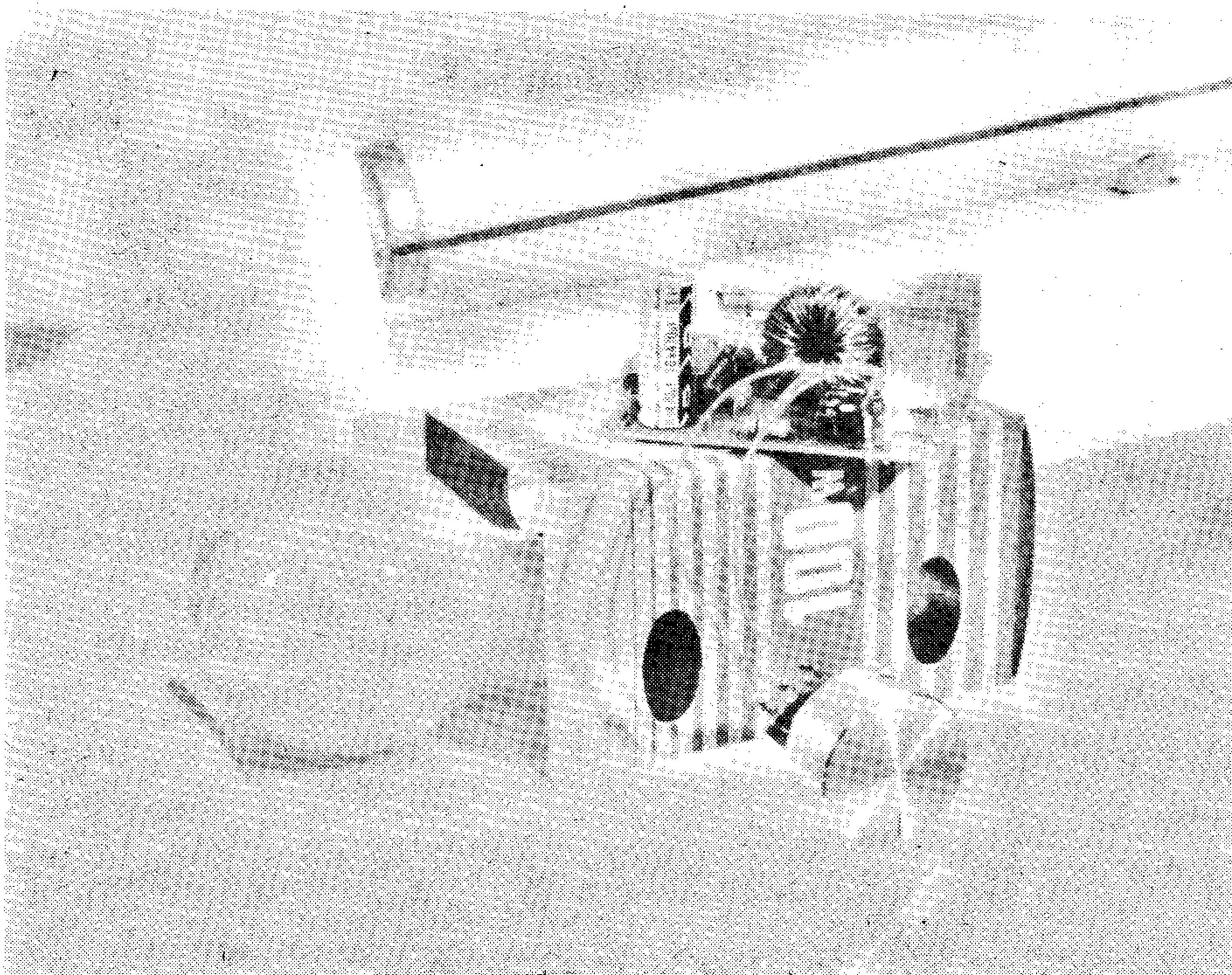
La placa de circuito impreso para el regulador se ha concebido con la posibilidad de varias versiones distintas para adaptarse a diversas aplicaciones:

- a) como regulador para lámparas de incandescencia, cuya luminosidad podrá ajustarse con el empleo de un potenciómetro en una gama determinada por el usuario.
- b) lo mismo que en el apartado anterior pero con la peculiaridad de que puede funcionar manualmente o por medio de un temporizador. El usuario tiene libertad para seleccionar el tiempo de control y el margen de luminosidad variable.
- c) lo mismo que en los apartados anteriores pero con tubos fluorescentes.

El circuito

En la figura 1 se muestra el circuito del regulador objeto de este artículo. Por lo general, para controlar un triac (Tri 1) se utiliza una red R-C combinada con un diac. Sin embargo, en este caso el ángulo de corte de fase se controla con el empleo de un circuito integrado especial: el SL440 de Plessey. Este último tiene la ventaja de permitir un control de la fase durante casi la totalidad del semiciclo de la tensión de red, lo que significa que, en la práctica, la potencia podrá variar continuamente desde el máximo hasta cero.

En la figura 2 se ilustra el principio del corte de fase. Cada vez que el triac recibe del



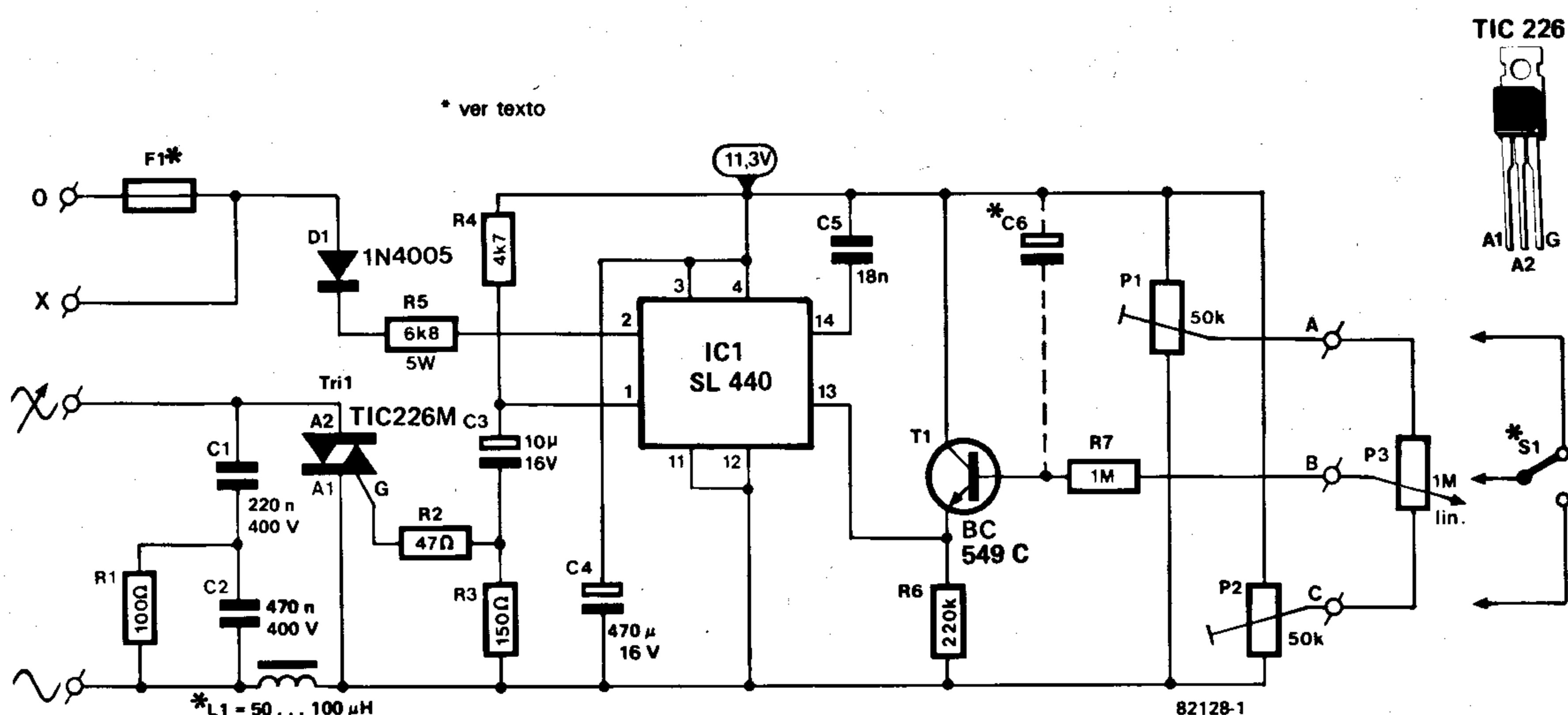


Figura 1. El regulador está concebido sobre la base de un circuito integrado especializado en el control del ángulo de fase. Se trata del SL 440 de Plessey. Los límites de la gama de ajuste pueden determinarse por el usuario con la ayuda de los potenciómetros P1 y P2.

circuito integrado un impulso de disparo (curva *a* en cada punto de cruce por cero de la tensión de la red, se aplicará la totalidad del potencial alterno a la carga (curva *b*). Pero si, por ejemplo, los impulsos de disparo se aplican dos milisegundos después de los puntos de cruce por cero (curva *c*), la carga estará al 95% aproximadamente de la potencia total (curva *d*). Si se acentúa el desplazamiento de fase de los impulsos de disparo respecto a los pasos por cero (curva *e*), se debilitará tanto más el potencial efectivamente aplicado a la carga (curva *f*); en este caso, aproximadamente la mitad de la potencia. Es así como mediante una regulación del ángulo de corte de fase se puede controlar la potencia en la carga, desde la totalidad hasta cero.

Como ya hemos indicado, los impulsos de disparo, o de puerta, son suministrados por el circuito integrado SL440. Entre otras cosas, el circuito integrado incluye un estabilizador de corriente continua, un detector de

paso por cero, un generador de impulsos de retardo variable y un amplificador. El circuito de estabilización asegura la alimentación del circuito integrado a partir de la tensión de la red rectificada en monoalternancia a través de D1 y de R5. El filtrado está asegurado por el condensador electrolítico C4.

El detector de paso por cero capta el punto de cruce por cero de la tensión de la red y dispara al generador de impulsos. Este último es realmente un monostable cuyo intervalo entre dos impulsos es variable. En la patilla 1 del circuito integrado, estos impulsos se suministran amplificados de modo que, a través de C3 y de R2/R3, pueden invertirse; su duración es, entonces, de unos 50 μ s para una intensidad de corriente de 100 mA. El ángulo de corte de fase se controla con el potenciómetro P3 (a través del seguidor de emisor T1) para proporcionar una tensión de control comprendida entre 1,8 V y 8,5 V en la patilla 13 de IC1 (el conmutador S1 y

el condensador C6 se considerarán más adelante). El margen de control está limitado en ambos extremos por medio de la posición del cursor de P1 y de P2.

En el momento en que el triac recibe un impulso de cebado se pondrá a conducir, aunque provocará también la aparición de parásitos de alta frecuencia. Por esta razón, una red de filtro debe incluirse alrededor del triac y está constituida por los componentes L1, C1, C2 y R1. Esta red L-C atenúa la elevación rápida de la corriente, con lo que se limitan los riesgos de parásitos.

Como las lámparas de incandescencia tienen tendencia a producir cortocircuitos en otros dispositivos cuando estos sufren anomalías, el fusible F1 se ha conectado en serie con la alimentación de la red. También sirve para proteger al triac contra las corrientes excesivas.

Regulación de lámparas de incandescencia

En la figura 3 se ilustra la forma (sencilla y directa) de conectar las lámparas de incandescencia con el regulador. Los valores máximos de corriente transitoria, en el momento del encendido de una lámpara fría, pueden tener una magnitud de 10 a 20 veces el valor nominal, por lo que el fusible (F1 en la figura 1) debe ser capaz de soportar estas sobrecargas. En la práctica, se le da de dos a tres veces el valor nominal de la corriente consumida por la lámpara (o sea el número de vatios dividido por el número de voltios de la tensión de la red, multiplicado por dos o tres). Para una lámpara de 100 W, por ejemplo, será preciso un fusible (lento) de un amperio.

Para el ajuste de los límites del margen de control, es necesario disponer de un voltímetro provisto de una escala de tensión alterna de 220 V, como mínimo, que se conectará a los bornes de la lámpara. Poner el cursor de P3 en la conexión marcada A. A continuación, es preciso ajustar P1 hasta que la tensión obtenida por el instrumento ya no

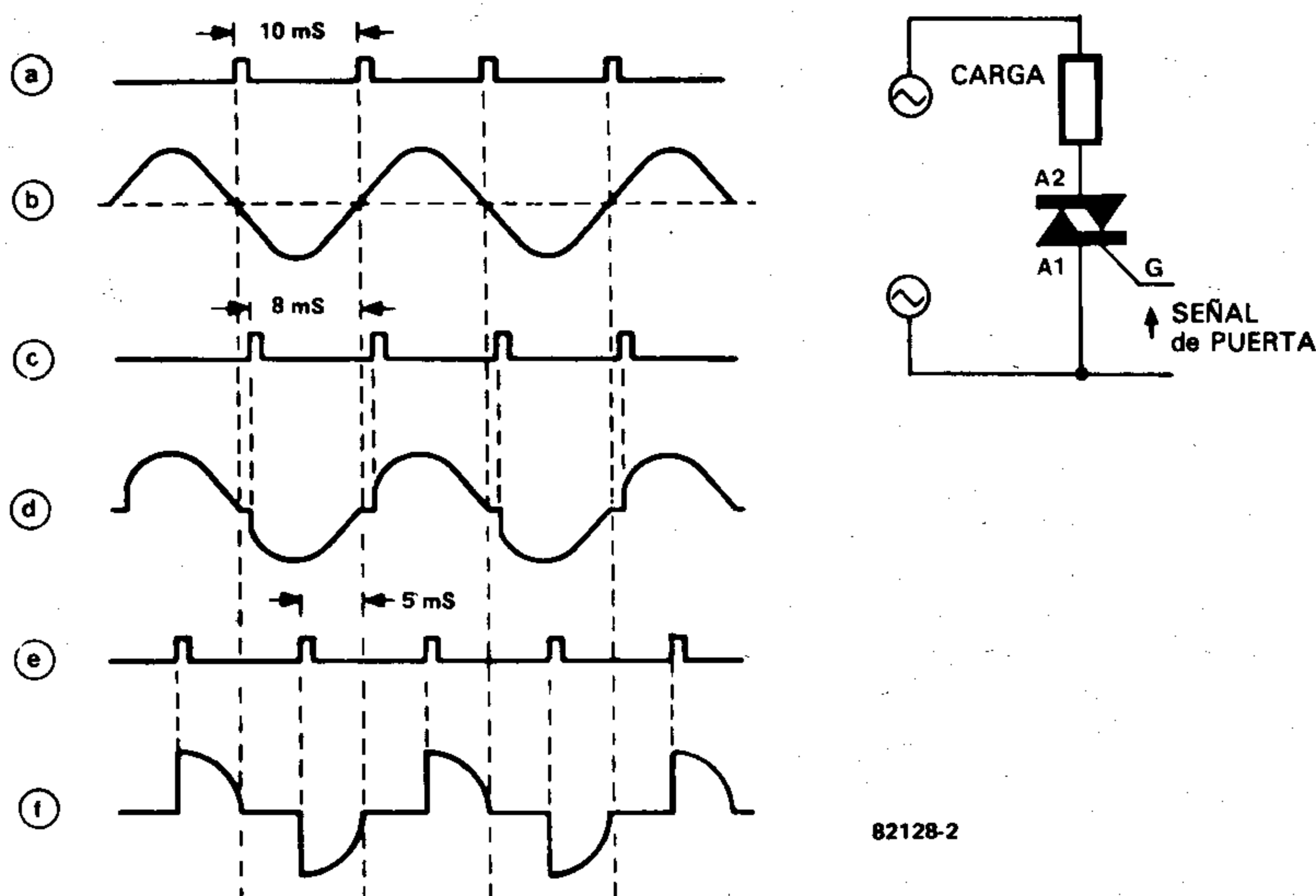
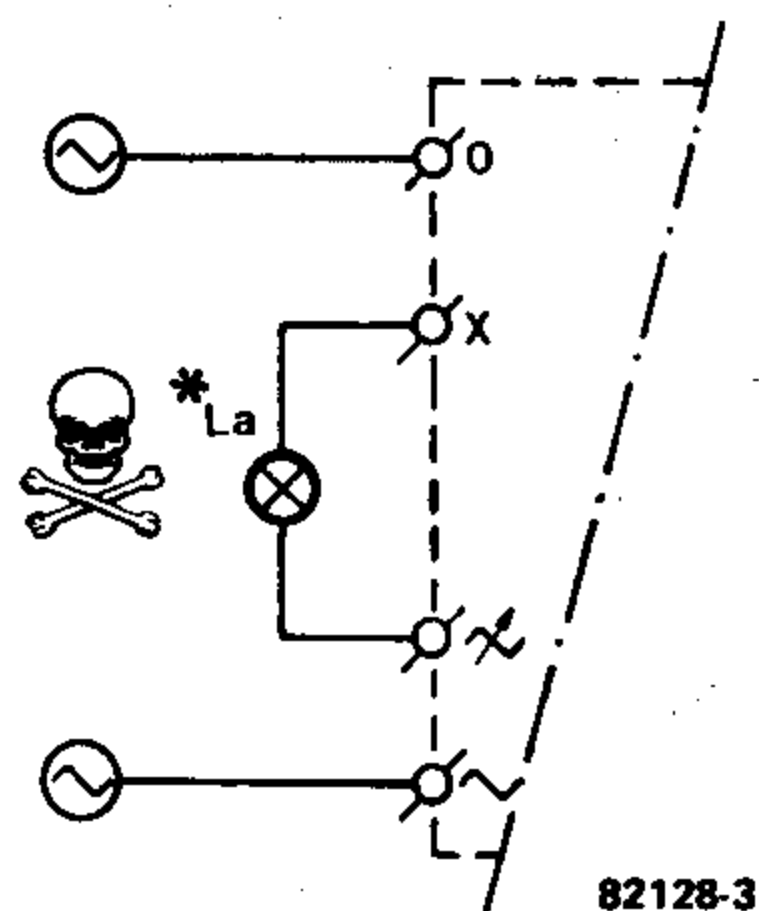


Figura 2. Según el momento en que el triac reciba su impulso de desbloqueo (a, c ó e), proporcionará toda o parte del potencial de la red a la carga (curvas b, d ó f).

3



82128-3

Figura 3. Conexión de una lámpara de filamentos a las salidas del regulador.

disminuya (aproximadamente a 0 voltios). Llevar el cursor de P3 al otro extremo; el instrumento debería indicar una tensión mucho más alta. Ajustar P2 hasta que la tensión alcance su nivel máximo sin que aumente más (o sea, aproximadamente la totalidad de la tensión proporcionada por la red, menos 6,5 V).

No tendría nada de sorprendente que se constatará un «espacio muerto» en el encendido de la lámpara, puesto que es preciso una cierta tensión de umbral antes de que se encienda. Dicho «espacio muerto» podrá suprimirse fácilmente con el empleo de P1, que se ajustará de modo que, cuando P3 esté en la posición de luminosidad mínima (cursor hacia el punto A), la lámpara salga sensiblemente de las «tinieblas».

Hay un inconveniente y es que la lámpara siempre consumirá corriente de la red; dicho de otro modo, siempre hay una tensión a través de la misma.

¡Hay que tenerlo en cuenta a la hora de cambiar las lámparas!

Por supuesto, los constructores tienen libertad para elegir otros márgenes de control de la luz (tales como el 30%... 80% del brillo máximo de la lámpara) según sus necesidades.

Una potencia mínima de 40 W es necesaria para asegurar el buen funcionamiento del regulador que podrá suministrar hasta 200 W con un triac sin disparador. Con tal de que el triac esté suficientemente refrigerado, pueden controlarse lámparas de hasta una potencia total de 1.500 W. Recuerde que la bobina L1 también debe estar adaptada a la potencia de la carga. Para 1.000 W, por ejemplo, a una tensión de 220 V, le será preciso suministrar unos 5 A (1.000 W/220 V).

Regulación de los tubos fluorescentes

Es imposible controlar un tubo fluorescente directamente por el regulador, puesto que éste precisa una tensión de precalentamiento elevada. Una vez que se haya cebado el tubo, la temperatura se mantiene por la descarga, a condición de que no caiga por debajo de cierto umbral. Dicho de otro modo, un tubo fluorescente ordinario no se presta a la regulación a la que nos gustaría, no obstante, someterle.

Pero, afortunadamente, existen tubos de fabricación especial que no requieren una tensión de cebado elevada. Los tubos TLM, como se les denomina, están provistos de una

cinta conductora en la pared exterior, conectada a un electrodo, a través de una resistencia de valor óhmico elevado (ver figura 4). Cuando se pone en servicio este tipo de tubo, la tensión de la red total se aplica entre el extremo desconectado de la banda conductora y el filamento. La mezcla gaseosa entre ellos se ionizará muy rápidamente debido al potente campo eléctrico (=tensión $\times$ distancia). La banda conductora asegura que la nube de iones se distribuya rápidamente a todo lo largo del tubo. En consecuencia, se produce una descarga gaseosa y se ilumina el tubo. Cuanto mayor sea el calentamiento por corriente de los filamentos del tubo, tanto más fácil será el encendido. Con tal de que los filamentos estén su-

ficientemente calentados, el tubo puede encenderse incluso a tensiones muy bajas, en cuyo caso el tubo puede regularse en su luminosidad. En principio, un tubo fluorescente ordinario también puede regularse, con tal de que esté suficientemente precalentado. Sin embargo, no funcionará tan bien como uno del tipo de cebado automático. Los tubos delgados más modernos serán todavía más difíciles de regular. Es preferible optar por tubos fluorescentes de cebado automático, aunque sean algo más caros que sus contrapartidas.

Habrà de encontrarse un método para precalentar los tubos fluorescentes para que puedan regularse. En la figura 5 se ilustra cómo puede conseguirse con un transforma-

4

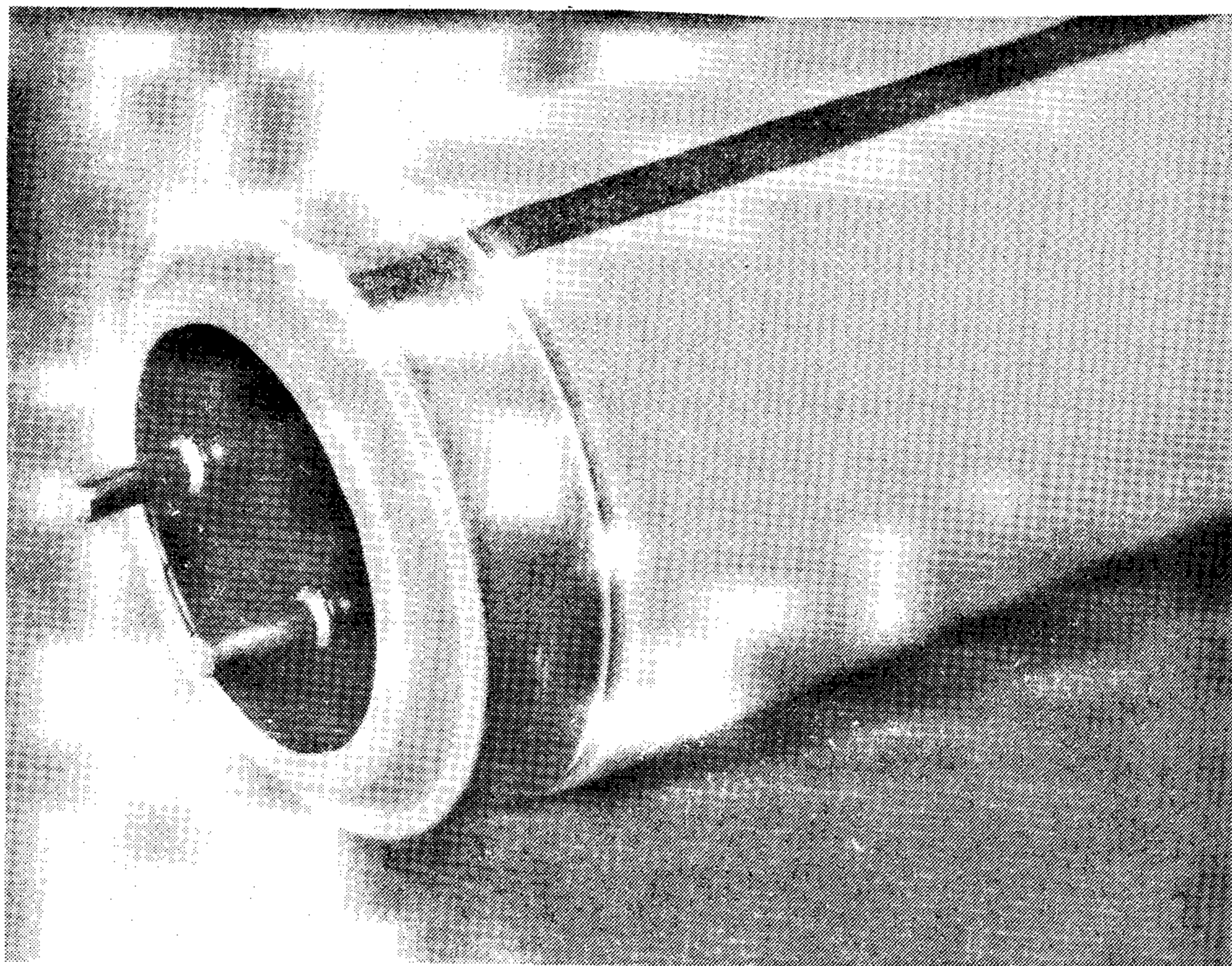
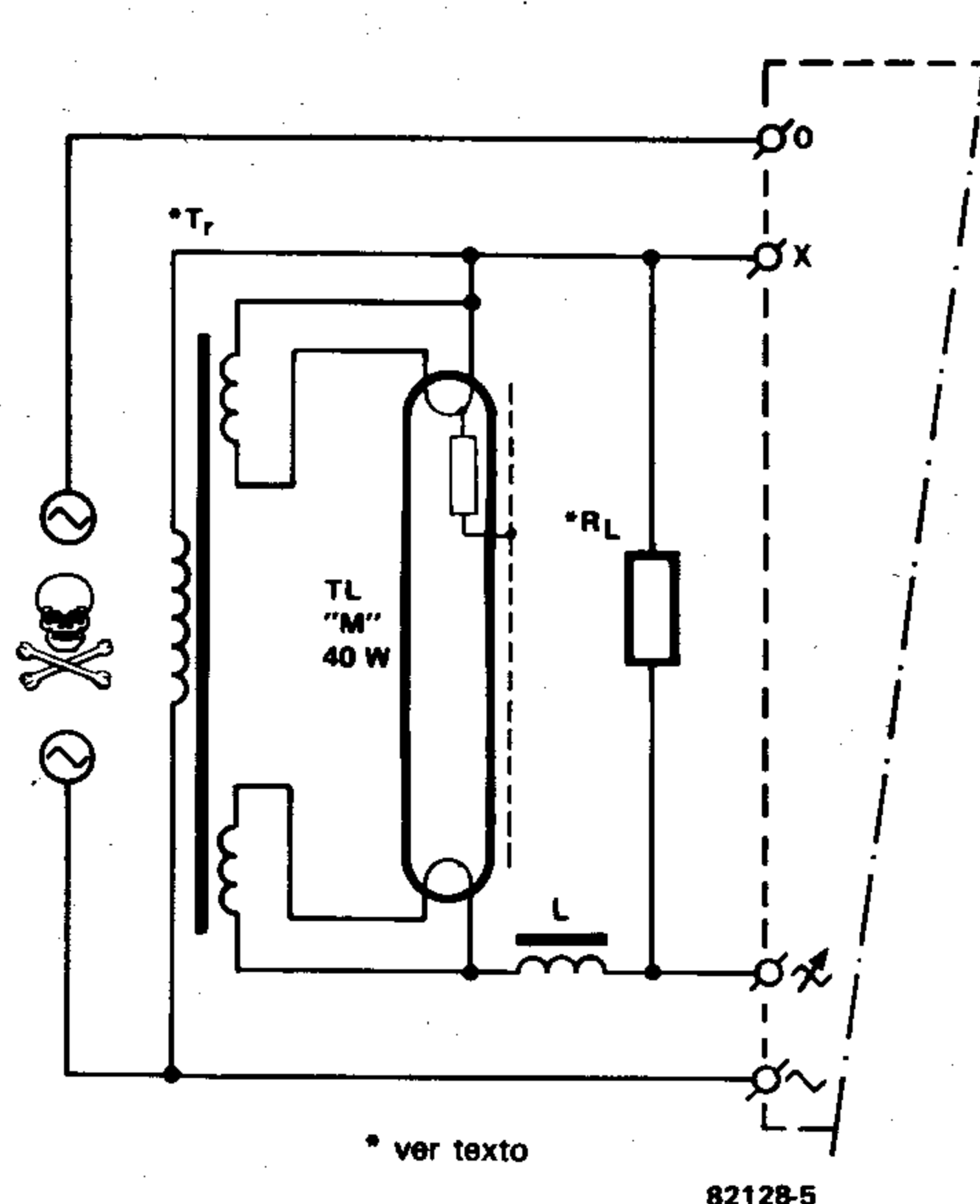


Figura 4. Un tubo TLM está provisto de una cinta conductora en su pared exterior, conectada a uno de los electrodos por una resistencia de gran magnitud. Se distingue, sobre todo, por el hecho de que no requiere más que una tensión relativamente pequeña de cebado.

5



* ver texto

82128-5

Figura 5. Forma de conectar un tubo TLM al regulador. El tubo está continuamente calentado con la ayuda de un transformador (Tr), que permite que el tubo se regule con niveles de tensión muy bajos. Para que el regulador trabaje adecuadamente se requiere una resistencia de carga R_L .

dor. Este último necesita disponer de dos devanados secundarios separados de 3,7 V/0,62 A.

Philips fabrica transformadores de filamentos especiales para este propósito (ver tabla 1), que pueden incorporarse en la caja del cebador del soporte del tubo fluorescente. Como alternativa, también puede ser adecuado un transformador ordinario con dos devanados independientes de 4 V (6 V max.)/0,8 A. De ser necesario, pueden utilizarse dos transformadores de timbre de 3... 5 V/1 A.

Echemos una nueva mirada al esquema circuital de la figura 5. L es una bobina de reactancia de fluorescente normal. El campo magnético que se crea periódicamente en la bobina debe ser capaz de disgregarse rápidamente pues, de no ser así, el triac en el regulador estaría conduciendo demasiado tiempo. Esto se logra mediante la resistencia R_L . Cuanto más pequeño es el valor de R_L , tanto más rápidamente se disgrega el campo magnético y tanto más amplio es el margen de control del regulador. Si se sobrepasa el margen, el tubo fluorescente empezará a centellear. Esta circunstancia debe subsanarse muy rápidamente, pues una componente de corriente continua asimétrica y perjudicial comenzará a circular a través de la bobina. P1 y P2 (figura 1) mantienen el margen dentro de unos límites de seguridad. Hay que ajustar P2 de modo que se encienda el tubo a plena potencia sin parpadeo alguno.

Aunque la selección de un valor pequeño para R_L proporcionará un margen de control más amplio, también es cierto que se perderá más energía. Como una solución de compromiso, es preferible seleccionar un valor de 4k7/15 W para un tubo fluorescente de 40 W. En el caso de una potencia más alta del tubo, o cuando hay que controlar tubos múltiples, R_L debe tener un valor más bajo y una mayor potencia de disipación (para un tubo fluorescente de 80 W, se precisará $R_L = 2k \Omega / 30 W$). Se descubrirá que una lámpara ordinaria de filamento se adapta muy bien a esta finalidad. Una lámpara de 40 W será suficiente para dos o tres tubos

fluorescentes de 40 W de cebado automático. En la figura 6 se muestra cómo pueden conectarse dos tubos fluorescentes a un circuito regulador único.

El transformador ha de tener tres devanados. Dos filamentos de tubos pueden conectarse en paralelo en uno de los devanados. Por supuesto, el devanado debe ser capaz de proporcionar la corriente para ambos (ver tabla 1). También es factible emplear dos transformadores con dos devanados separados cada uno o utilizar tres transformadores con un solo devanado cada uno.

Como en el caso de las lámparas de incandescencia, la carga, cuando los tubos fluorescentes estén regulados, debe ser, como mínimo, de 40 W (potencia en el interior del tubo + R). Con el empleo de un triac no enfriado, el regulador puede operar con una carga de hasta 200 W. Si el triac está refrigerado, el regulador puede trabajar con 1.500 W.

Los tubos fluorescentes del tipo de 40 W de cebado automático son los más fáciles de conseguir. Tienen una longitud de 120 cm., pero si son demasiado largos para instalarse en su acuario utilice una lámpara fluorescente ordinaria. Tal como se indicó anteriormente, no son de calidad plenamente satisfactoria, pero se probó un prototipo en el laboratorio y se obtuvieron resultados aceptables.

Regulación de lámparas de incandescencia o fluorescentes en régimen de conmutación

En los reguladores de conmutación gradual, el potenciómetro de ajuste P3 se sustituye por un conmutador articulado de dos direcciones o por un contacto de relé de dos vías en un cronointerruptor (ver figura 1, puntos de conexión A ... C). El resultado es un regulador de dos vías. El brillo puede preajustarse en ambas posiciones del conmutador por medio de P1 y de P2. Con la adición del condensador C6 al circuito, la tensión en la patilla 13 de IC1 cambiará gra-

dualmente de nivel mientras que se carga o descarga el condensador, aumentando o disminuyendo gradualmente el brillo. Ello permite que la luz se desvanezca, de una forma conmutada, con una apariencia más natural.

En combinación con lámparas de incandescencia y, de ser necesario, con un cronointerruptor, este regulador resulta ser un controlador ideal de la luz en pajareras. La lenta aparición de la oscuridad da a los pájaros mucho tiempo para «prepararse para ir a la cama». En esta aplicación particular, las lámparas de incandescencia tienen una ventaja sobre los tubos fluorescentes por cuanto que disipan una cantidad de calor considerable, lo que los pájaros «agradecen» durante los meses fríos del invierno. Un regulador de esta clase también podría instalarse en los dormitorios de los niños, para casos en que éstos odien despertarse de repente por una luz brillante.

Los poseedores de acuarios encontrarán muchas aplicaciones para este circuito. Aunque los peces se caracterizan por su «sangre fría», casi salta fuera de su habitat cuando la oscuridad se hace repentinamente. Probablemente piensen que han sido engullidos por un feroz predador a pleno día. Por consiguiente, el regulador de conexión/desconexión gradual puede contribuir, en gran medida, a la felicidad doméstica, tanto dentro como fuera del agua.

Es preciso calibrar el margen de regulación con P1 y P2 antes de que se monte el condensador C6. Una vez realizado el ajuste, el condensador puede añadirse al circuito. Es preciso cerciorarse de haber desconectado, de antemano, la alimentación de la red. Debido a la corriente de fugas, el condensador debe tener una tensión de trabajo de 40 voltios. Por lo que respecta a la capacidad, cada μF aporta un retardo aproximado de 5 segundos; pero la relación de la capacidad a la duración de la regulación depende también de la posición de ajuste de P1 y de P2. Se podrá empezar con un valor de prueba de 4,7 μF y modificarle si fuera necesario. Después del encendido, hay que esperar hasta que C6 se cargue hasta el nivel de tensión

6

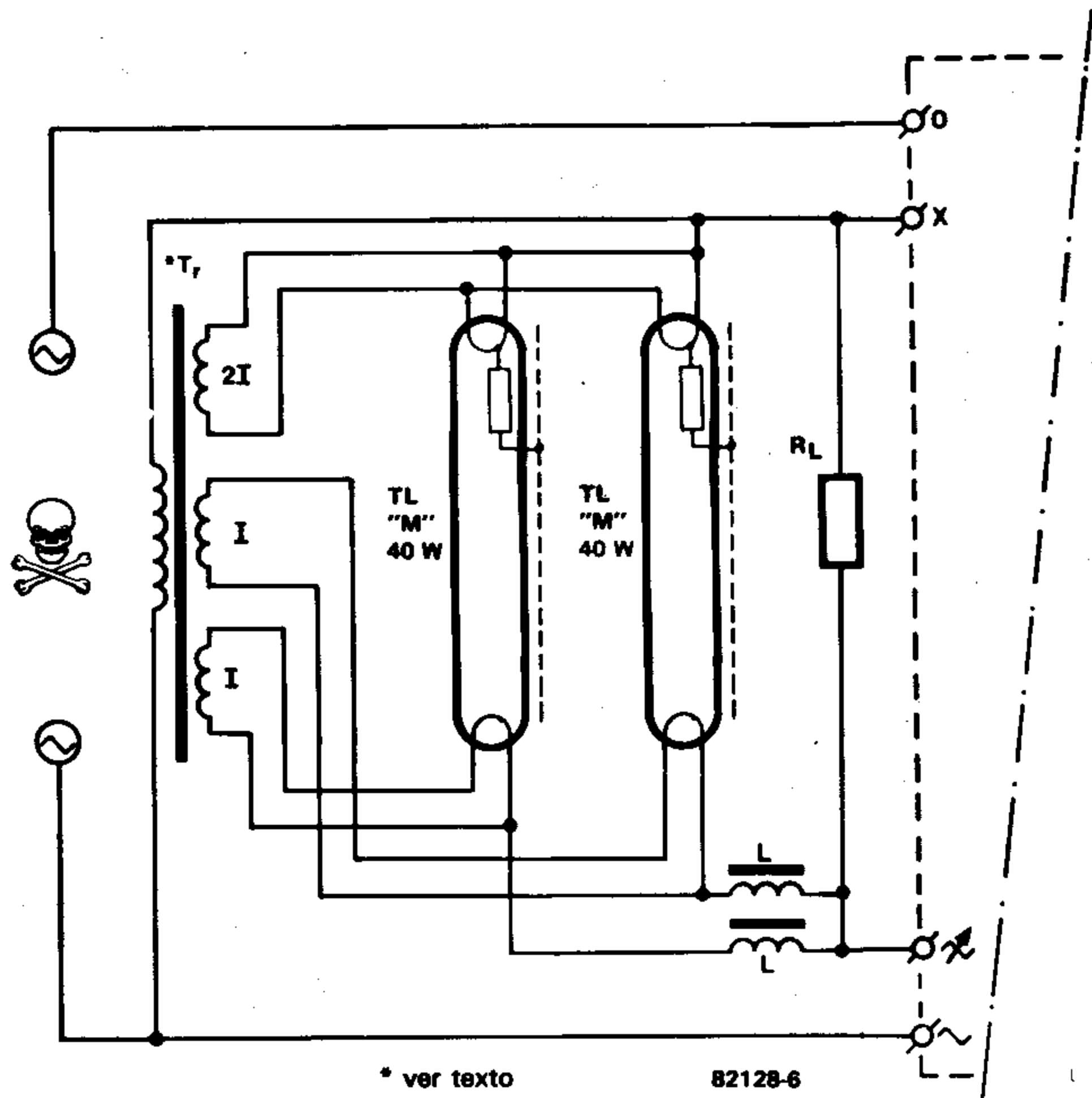


Figura 6. Cómo conectar dos tubos fluorescentes (de cebado automático) a un único regulador de control.

Tabla 1.

Philips dispone de componentes especiales para los tubos TLM y para su adaptación a reguladores de intensidad luminosa:

Transformador:

PMP 42T/05 con tres devanados secundarios independientes: 2 x 3,7 V/0,62 A y 1 x 3,7 V/1,25 A.

Choques amortiguadores:

BTP 40L05L para 1 x 40 W (TL o TLM)

Armaduras completas:

con uno o dos choques y un transformador del tipo PMP 42T/05

TMX 100 - 140 DIM para 1 x 40 W TL(M)

TMX 100 - 240 DIM para 2 x 40 W TL(M)

TMW 060 - 140 DIM para 1 x 40 W TL(M)

«especial acuario»

Tubos TLM:

El tipo más corriente es el de 40 W, disponible en distintas versiones de «calor» (luz más o menos blanca). Los modelos de la serie 80 tienen un rendimiento botánico óptimo y son muy adecuados para aplicaciones domésticas. También existen tubos fluorescentes de 20 y 65 W con reflector incorporado.

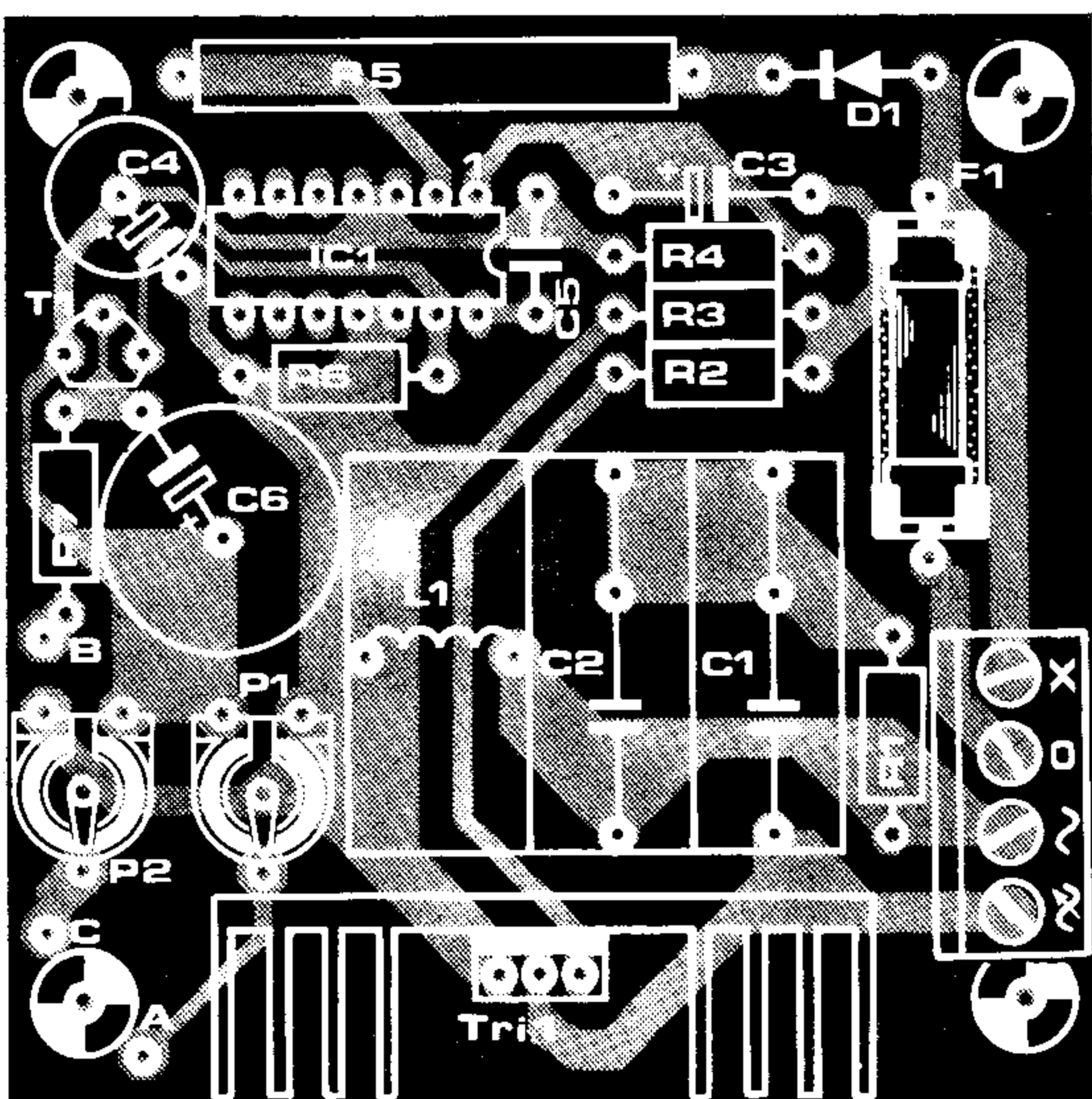
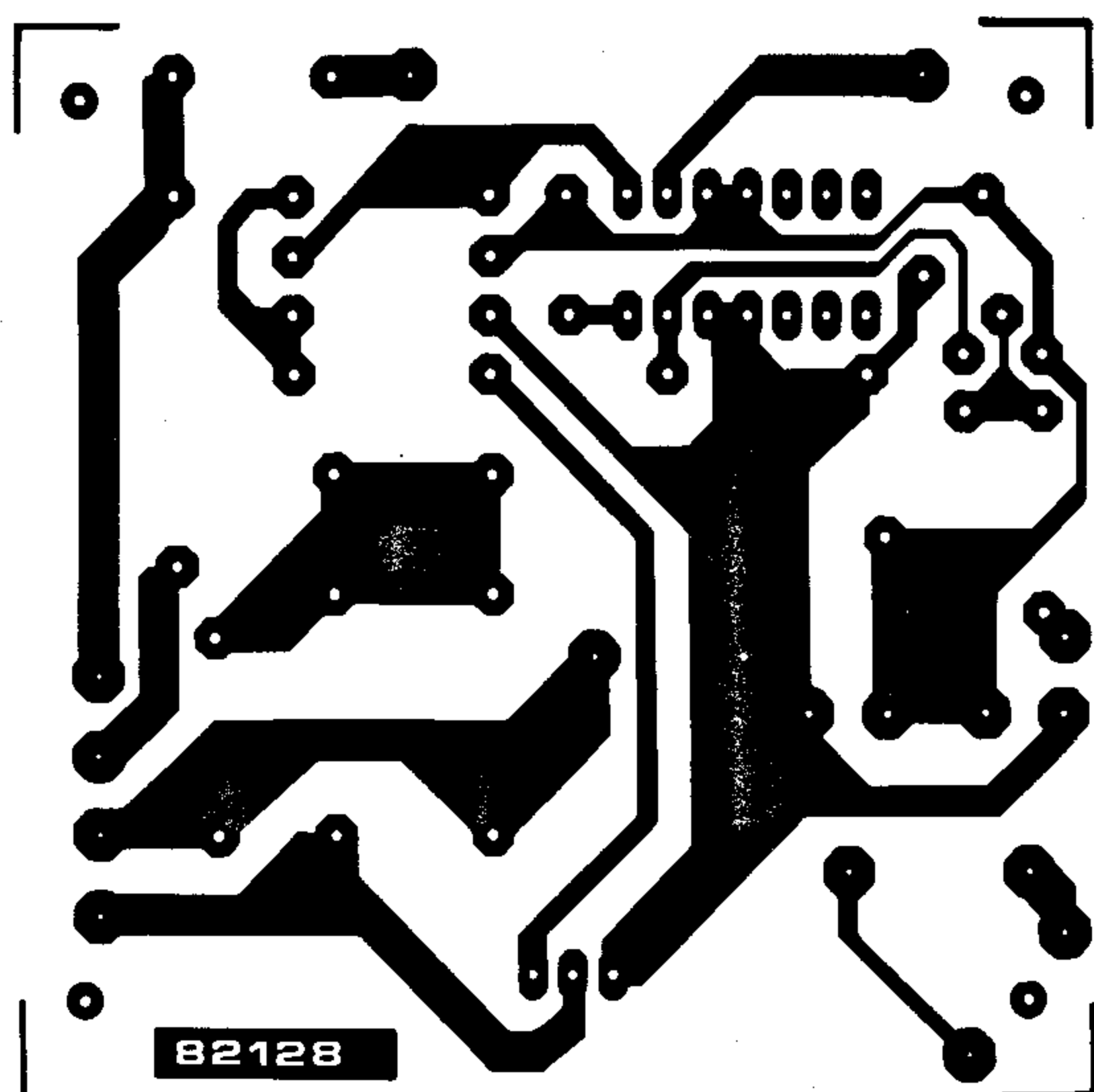


Figura 7. Placa de circuito impreso y disposición de los componentes del regulador universal.

Lista de componentes

Resistencias:

- R1 = 100 Ω
- R2 = 47 Ω
- R3 = 150 Ω
- R4 = 4k7
- R5 = 6k8/5 W
- R6 = 220 k
- R7 = 1 M
- P1,P2 = 50 k ajustable
- P3 = 1 M lineal (ver texto)

Condensadores:

- C1 = 220 n/400 V
- C2 = 470 n/400 V
- C3 = 10 μ/16 V
- C4 = 470 μ/16 V

- C5 = 18 n
- C6 = ver texto

Semiconductores:

- T1 = BC 549C
- IC1 = SL 440 (Plessey)
- D1 = 1N4005, 1N4004
- Tr1 = triac TIC 226M o TIC 226D

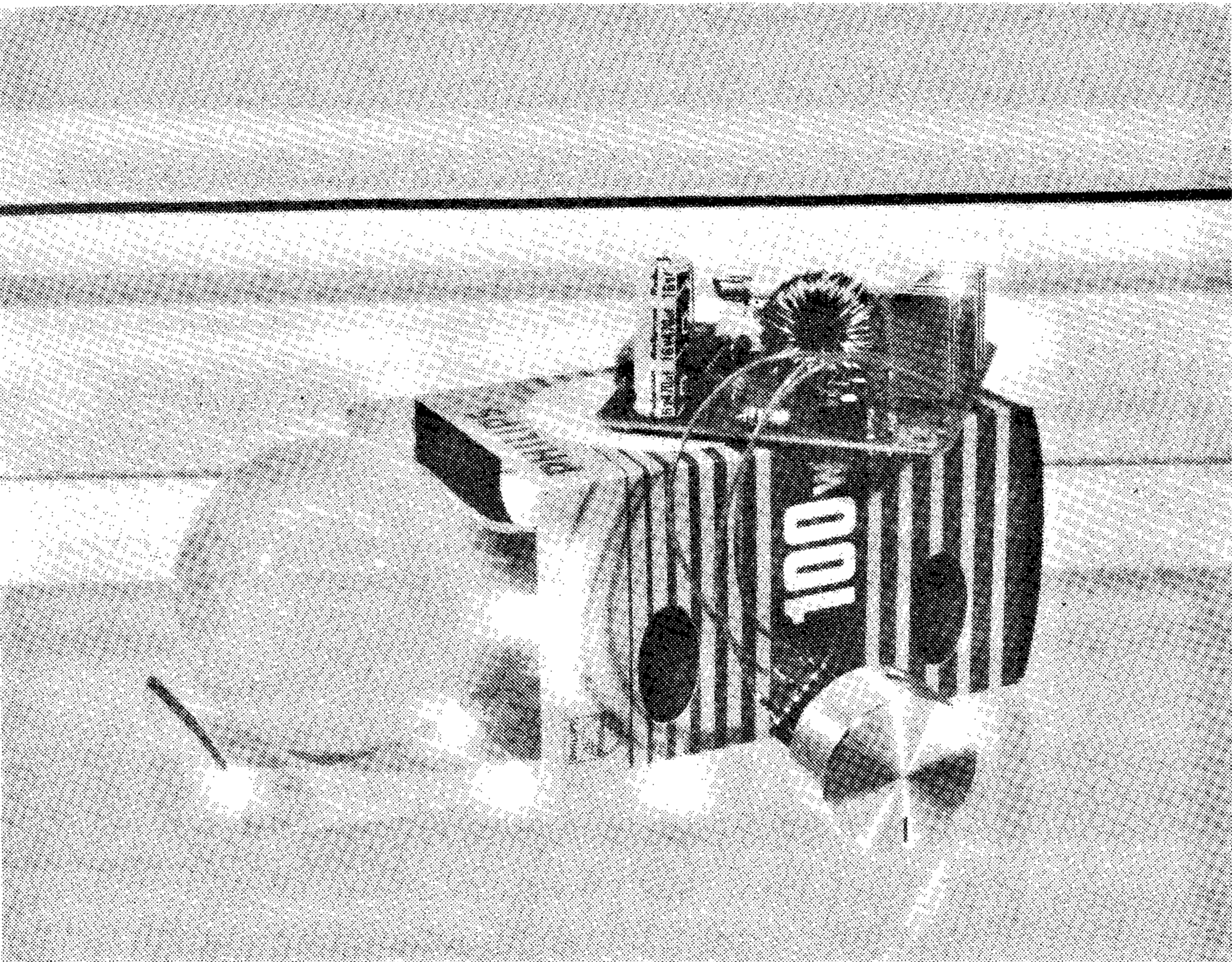
Varios:

- L1 = Choque toroidal 50 ... 100 μH
* ver texto
- S1 = Conmutador unipolar o relé inversor (ver texto).
- F1 = fusible (ver texto)
porta-fusible para circuito impreso

de reposo, antes de comprobar el tiempo de retardo. Un valor capacitivo muy alto (superior a 1.000 μF) puede dar lugar a una corriente de fugas excesivas y plantear problemas.

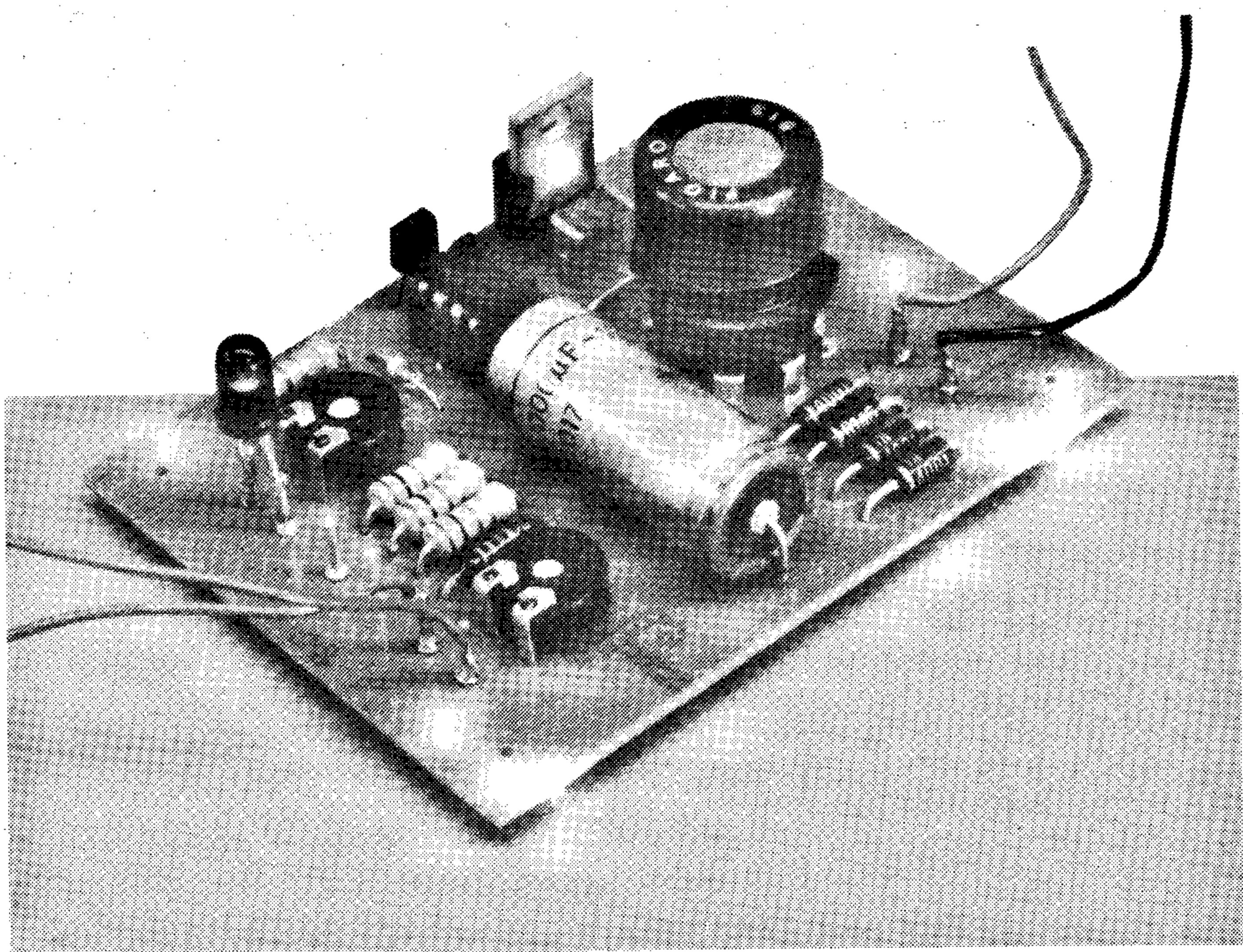
Recomendaciones prácticas

La construcción del montaje sobre la placa de circuito impreso no planteará problema alguno. Sin embargo, si se utiliza un zócalo para IC1 es importante cerciorarse de que C4 se descargue cada vez antes de instalar el circuito integrado. En la práctica, la unidad puede emplazarse virtualmente en cualquier lugar que sea conveniente, incluyendo una caja de conmutación ya existente. En este último caso, hay que observar que el control del regulador no será compatible con el cableado casero normal y tendrán que instalarse cables adicionales. Un total de 4 hilos de conexión más el de tierra debe instalarse entre la caja de conmutación y el soporte de la luz. Además de ello, se requiere también una alimentación tomada de la red. Una alternativa es instalar el sistema electrónico total en el soporte de la luz, si fuera posible dicho montaje. Este método sólo requiere tres hilos de conexión a la unidad de control de conmutación. No es posible indicar con exactitud qué modificaciones se requieren, pues puede ser que no prevalezcan los cableados eléctricos normalizados, sobre todo si vive en una casa antigua. Ni que decir tiene que si no se está familiarizado con las instalaciones eléctricas de la vivienda, puede ser recomendable acudir a un electricista conocido para que preste su asesoramiento al respecto.



**una «nariz electrónica» que protege
contra los escapes de gas!**

detector de gas



Comúnmente, suele subestimarse el peligro de una fuga de gas y el daño que puede causar si permanece sin detectarse durante un largo período de tiempo. ¡Y claro... llega el día en que nos «acordamos de Santa Bárbara porque truena». Este sensor de gases da la alarma con gran rapidez en el caso de producirse un escape, e incluso puede poner en evidencia a los fumadores al hacer patente su contaminación ambiental.

El aire puro es un bien escaso en nuestros días, sobre todo en las zonas urbanas. No es, pues, sorprendente que la «nariz urbana» media se haya hecho insensible a la contaminación. Después de todo, generaciones enteras han estado respirando una mezcla compuesta de escapes de automóvil, derrames cloacales y humos de chimeneas industriales. Como resultado de este caos, se ha producido un cambio en el curso de la evolución y, en muchos de nosotros, se ha destuido nuestra aptitud para olfatear escapes de gases e incendios caseros. Es importante detectar estas posibles catástrofes en su etapa preliminar, cuando aún pueden dominarse de forma económica e inmediata. Un dispositivo detector eficaz y sofisticado es bastante caro. Así pues, en lugar de pagar un alto precio por un equipo que haga inútil a su nariz, parece más oportuno proveerse de un dispositivo útil y barato, aunque no sea tan exacto, para aumentar el olfato humano.

La firma japonesa Figaro Engineering Company (a la que no se le conoce ningún parentesco con el «Barbero de Sevilla») nos ha llevado, con sus precios y su disponibilidad,

a estudiar un circuito no profesional, relativamente sencillo, para su utilización casera. Comencemos por indicar que este captador se comporta como una resistencia NPC (*Negative Pollution Coefficient* o de coeficiente de polución negativo), lo que significa que cuanto más alta es la concentración de gas indeseable en el aire tanto más disminuye la resistencia del captador. La adición de un comparador y de un circuito visualizador da lugar a una «nariz electrónica» razonablemente sensible.

Hay dos tipos básicos de sensores NPC que son adecuados para el circuito descrito: el 812 y el 813. La única diferencia entre ambos es que el 812 es muy sensible al monóxido de carbono, amoníaco, alcohol y bencina, lo que le hace ideal para la detección de humos/incendios y respiraderos, mientras que el 813 es mejor para detectar propano, butano y metano, es decir, para localizar escapes de gases domésticos. La elección se deja a la discreción del constructor.

El circuito

Se pueden decir muchas cosas respecto al circuito de la figura 1, pero nunca que se trata de una nariz complicada. Como puede observarse, se requieren muy pocos componentes.

A la izquierda del esquema hay una fuente de alimentación estabilizada de 5 voltios. Como el consumo de corriente es elevado, la utilización de pilas es crítica. Más adelante describiremos la posibilidad de una alimentación por medio de baterías de plomo/ácido, si bien, cabe recomendar el empleo de alimentación a partir de la red.

La resistencia R1 y el potenciómetro P1 se conectan en serie con el detector de gases GS1. Esta red constituye un divisor de tensión, cuyo funcionamiento está relacionado con la magnitud de la contaminación del aire. Dicho para entendernos, cuanto mayor sea la concentración de gases tanto más alta es la tensión aplicada a la entrada positiva del comparador IC2. La tensión de referencia se aplica a la entrada negativa de IC2 y se determina por medio de P2.

Consideremos que el circuito ha sido «equilibrado» en un ambiente neutro (relativamente limpio). Entonces, la tensión en la entrada inversora es superior a la tensión existente en la entrada no inversora. En caso de alta concentración de gases, se perturba el equilibrio y el transistor T1 se pone a conducir, se enciende el diodo LED D7 y se activa el relé Rel. Evidentemente, el relé puede utilizarse para activar cualquier clase de sistema de alarma elegido, tal como zumbadores, timbres o, por supuesto, ventiladores. Al examinar el circuito puede constatar, a primera vista, que uno de los dos potenciómetros es superfluo. Este no es realmente el caso, ya que es conveniente utilizar P1 para determinar el margen de trabajo del sensor y P2 para ajustar la sensibilidad del dispositivo. El funcionamiento de IC2 es básicamente independiente de los niveles «absolutos» en una u otra de sus entradas, debido a que tiene un margen de modo común. Los componentes D5 y S2 están marcados con un asterisco por motivos que merecen una explicación. Cualquier sistema de alarma que tenga timbres y diodos LED sólo es eficaz mientras que haya alguien que los es-

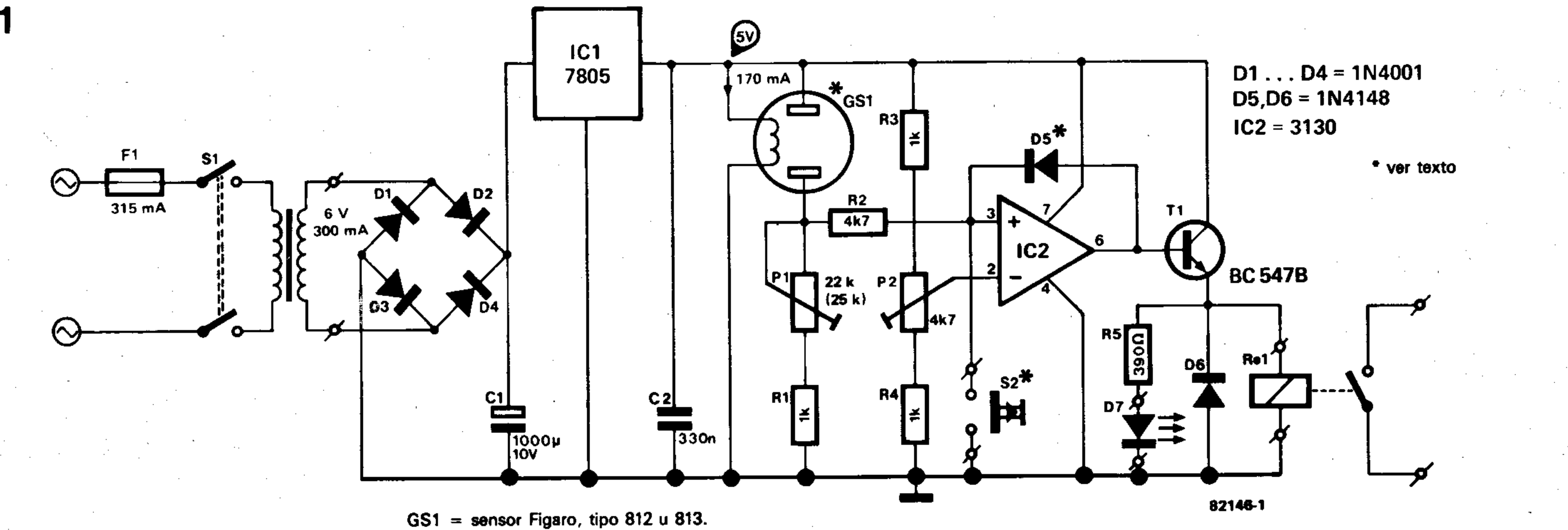


Figura 1. Esquema del circuito detector de gases. Todos los componentes están montados en la placa de circuito impreso con la excepción del transformador y del relé.

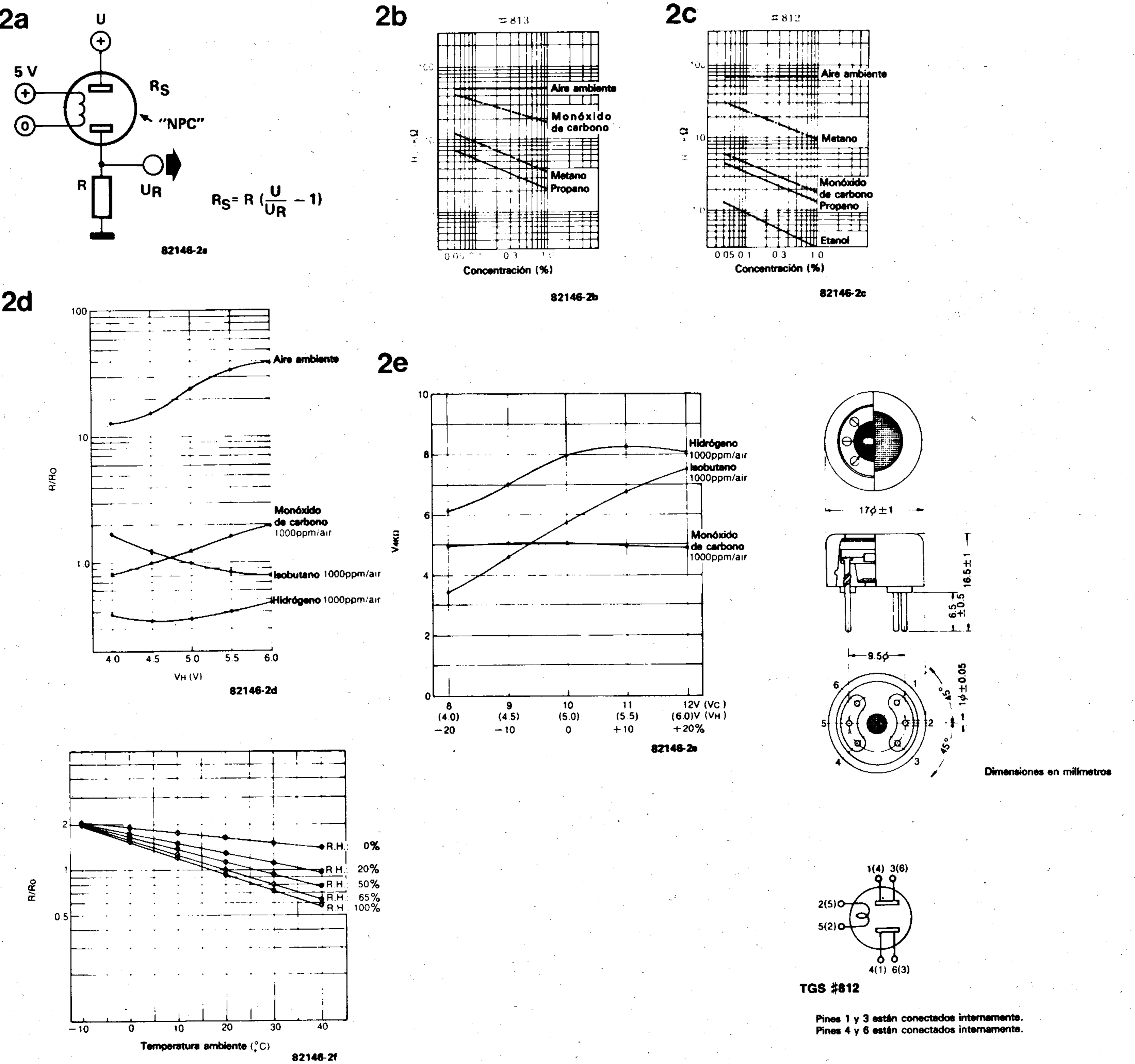


Figura 2. Parámetros y especificadores de los detectores de gas tipos 812 y 813 de FIGARO.

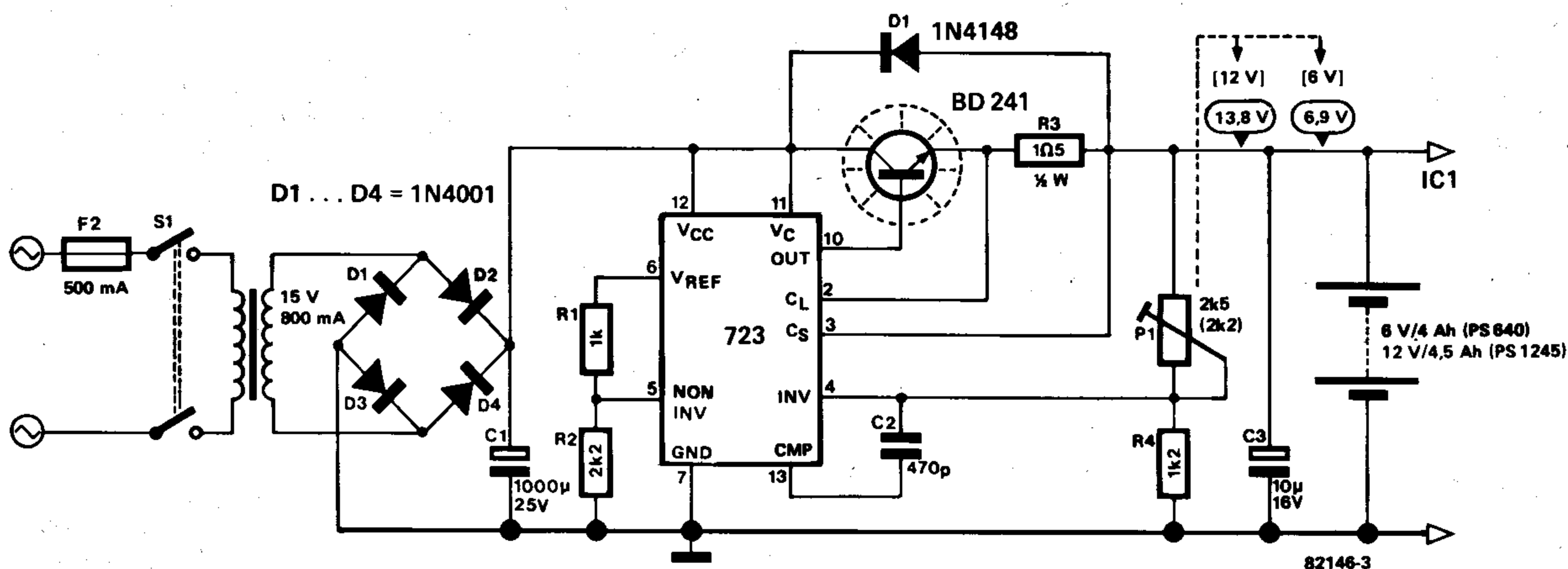


Figura 3. Esquema de la fuente de alimentación a partir de la red y del cargador/alimentador combinados para su empleo con baterías estancas de plomo-ácido. Este circuito no se monta sobre la placa de circuito impreso.

cuche o los vea. Esto no ocurre cuando se está fuera de casa durante el día o se está de vacaciones. En la mayor parte de los casos de escape de gases, y transcurrido un breve intervalo de tiempo, desciende automáticamente la presión real bajo la que escapa el gas. Ello puede dar lugar a que el sensor interrumpa la detección de su presencia. Al regresar a casa, no nos serviría de nada dicha detección inicial. Por ello parece, pues, indispensable una memorización. Y ésta es la justificación de la existencia de IC2 provisto de D5, que se pone a conducir cuando se conmuta el comparador y da lugar a una realimentación regenerativa alrededor de IC2. El comparador se mantendrá en este estado (nivel lógico alto), aunque se produzca cualquier eventual cambio en la zona del sensor. Esta función de «cerrojo» sólo podrá anularse por la pulsación de S2, conmutador de reposición.

Realización del montaje

En la figura 4 se muestra la placa de circuito impreso del montaje. No se indicó nada respecto al montaje del relé ya que su tipo y tamaño suele ser optativo del constructor. Sólo indicaremos que las exigencias funcionales para dicho relé son: una tensión de servicio de unos 4 V y una corriente igual o inferior a 100 mA, con lo que el consumo total del circuito se mantiene inferior a 200 mA.

Aunque la tensión de emisor de T1 conmute de 0 a 4 voltios, no se puede emplear este punto para conexión directa a un TTL, dado que cuando la salida pasa a un nivel lógico bajo, funciona como drenador y por consiguiente, se corre el riesgo de que la caída de tensión en los bornes del relé se haga demasiado importante.

El captador GS1 podrá insertarse en los dos sentidos pues no está polarizado y la disposición de los terminales es simétrica. La acción de los potenciómetros ajustables permite compensar la distinta sensibilidad a los ga-

ses que caracterizan a los captadores del tipo 812 y los del tipo 813. Aunque sean, pues, perfectamente intercambiables, no es inútil insistir en que es preferible elegir el captador más adecuado a las necesidades concretas de detección.

Antes de la calibración definitiva del circuito con P1 y P2, es conveniente hacer funcionar el detector de gases durante un cierto período de tiempo, que permite al captador satisfacer las exigencias funcionales impuestas por el constructor, puesto que GS1 necesita un tiempo de caldeo antes de comportarse correctamente. Recomendamos dejar el circuito bajo tensión permanentemente durante un período de dos a tres días, aunque para la mayoría de las aplicaciones caseras son suficientes unas doce horas.

Es preciso comprobar si GS1 está consumiendo corriente, tocando IC1 y GS1. Si todo es correcto, al «meter el dedo en la nariz electrónica» notará que ambos componentes están calientes. Inicialmente, se ajusta P2 de modo que se mida una tensión en el margen de 1 a 3 V en la unión de P1 y de R2. Transcurrido el período de calentamiento inicial, el circuito puede calibrarse con precisión; por supuesto, en un ambiente lo menos contaminado posible. Si vive en un ambiente de gran ciudad, váyase al campo (¡con la carga de la batería tiene para 20 horas de autonomía!).

Para que el procedimiento de calibración sea satisfactorio es necesario un voltímetro para medir los niveles de tensión en la unión de P1 y de R2 y en la unión de P2 y la patilla 2 de IC2. Los valores absolutos en estos puntos no son importantes pues el objetivo del procedimiento es asegurarse de que el nivel en P2/patilla 2 sea más alto que en P1/R2. La cuestión que se plantea es ¿cuánto? Cuanto más pequeña sea la diferencia tanto más sensible será el circuito y según la ley de Murphy, cuanto más sensible es algo, tanto mayor es la probabilidad de una falsa alarma. Al fin y al cabo, lo que no nos interesa es que la alarma suene cada vez que alguien encienda un cigarrillo o cuando se encienda el calentador de gas.

En teoría, la única forma exacta de calibrar el circuito es utilizar un laboratorio completamente dotado. Lamentablemente, no todos nosotros disponemos de tal equipo y por ello no podemos tener en cuenta todos los factores, tales como la temperatura ambiente y humedad, etc. Por un proceso de eliminación, encontramos que el prototipo trabajaba bien cuando P1 estaba ajustado para proporcionar una lectura comprendida entre 1 y 3 V, pon P2 a unos 0,5 voltios más de tensión.

El circuito funciona siguiendo pautas digitales. Dicho de otro modo, detecta si existe, o no, algo tangible. A modo de experimento, esto puede traducirse en una indicación analógica conectando un polímetro, o un voltímetro, a través de la red constituida por P1 y R1. Se constatará que el nivel de la tensión se eleva en la presencia de una concentración nociva de gas o compuesto.

Puntos de interés e información deseable sobre sensores de gas

Los métodos para la detección de gases por medios electrónicos van desde la cromatografía de gases a circuitos complejos en los que se emplean elementos radioactivos. La firma japonesa Figaro ha probado que también era posible la utilización de semiconductores; en este caso, se trata precisamente de óxido de estaño dopado N (SnO_2).

El principio básico utilizado es que la conductividad de los materiales semiconductores se reduce a medida que se produce la absorción de oxígeno. Como la presión local de oxígeno en el aire ambiente es pasablemente estable, se puede definir una resistencia ideal R_s para el «aire puro». La presencia de gases, tales como los óxidos e hidróxidos de carbono, influye, por supuesto, sobre la relación entre el captador y el oxígeno y, por consiguiente, sobre la resistencia R_s . Para conseguir una respuesta razonablemente rápida, el proceso completo se ace-

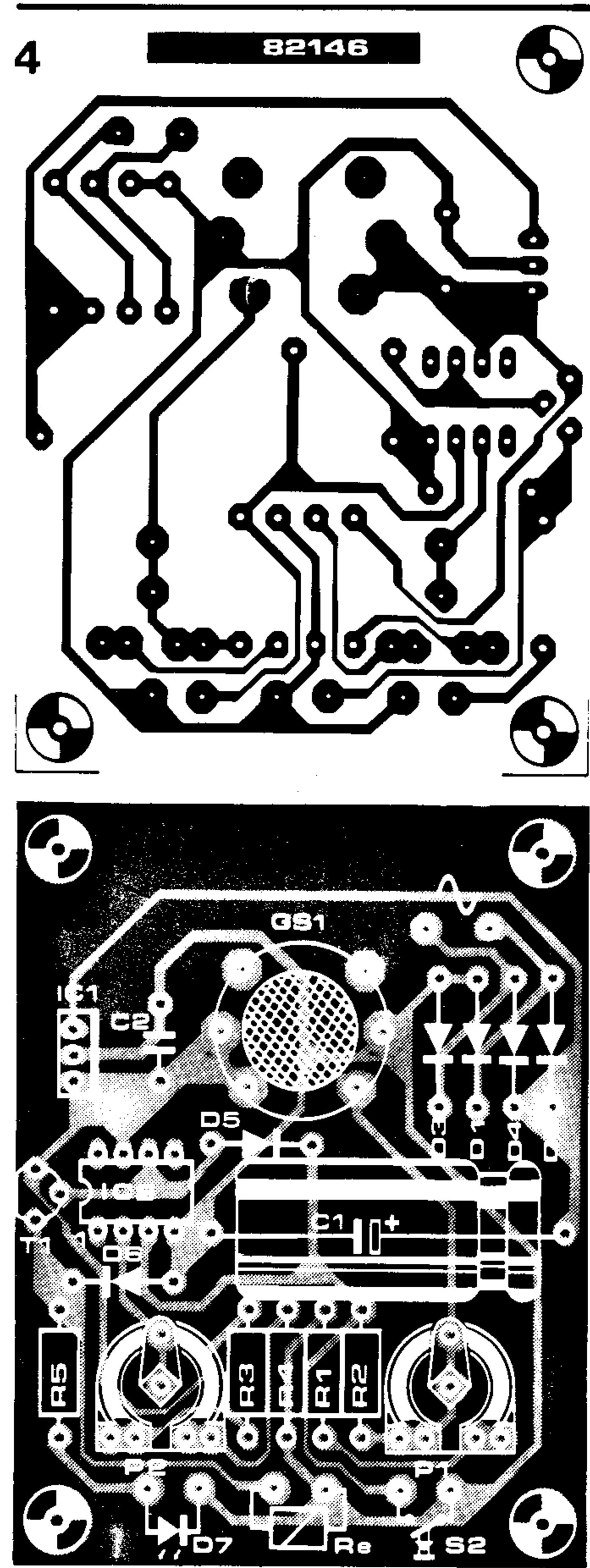


Figura 4. Placa de circuito impreso para el detector de gases.

Lista de componentes

Resistencias:

R1,R3,R4 = 1 k
R2 = 4k7
R5 = 390 Ω
P1 = 22 k (25 k) ajustable
P2 = 4k7 (5 k) ajustable

Condensadores:

C1 = 1000 μ/10 V
C2 = 330 n

Semiconductores:

D1,D2,D3,D4 = 1N4001
D5 (ver texto) ,D6 = 1N4148
D7 = LED
T1 = BC 547B
IC1 = 7805
IC2 = 3130

Varios:

GS1 = Sensor de gas Fígaro, tipo 812 (naranja) u 813 (negro). (ver texto)

Tr1 = Transformador 6 V/300 mA
F1 = Fusible 400 mA
S1 = Interruptor de red
S2 = Pulsador (contacto de trabajo)
Re1 = Relé 5 V/300...650 ohmios (exterior al circuito impreso).

lera calentando la superficie del sensor en varios centenares de grados centígrados. En las figuras b y c se ilustra el comportamiento de la resistencia correspondiente a los tipos 812 y 813 en presencia de algunos gases. Puede constatarse que dicha resistencia disminuye en correspondencia con su exposición a ciertos gases. El valor de referencia R_s tiene un equivalente que puede determinarse con el empleo de la fórmula mostrada en la figura 2a. En este caso, U es la tensión de alimentación constante y U_r es la tensión medida en la unión del sensor y la resistencia R . La sensibilidad y el rendimiento de los detectores se afectan, en gran medida, por la caída de tensión a través del filamento (figura d), por la tensión de alimentación (figura e) y por la temperatura y humedad relativa del ambiente (figura f, tipo 812). Estos diagramas muestran también los límites de utilización casera, dado que no se trata de construir un discriminador preciso y que no suele disponerse de un laboratorio muy bien dotado.

El proceso de absorción en el que se basa el funcionamiento de los captadores de gases, contruidos con semiconductores, tiene una cierta duración; dicho de otro modo, R_s tarda un cierto tiempo en tomar un nuevo valor como reacción a los cambios en las condiciones ambientales y a las modificaciones de las concentraciones de gases. Si el detector no ha estado en uso durante un cierto período de tiempo (que ha de ser corto), debe repetirse el procedimiento completo de calentamiento y de calibración.

¿Qué ocurre si falla la alimentación?

El dispositivo que presentamos consume una cantidad de corriente considerable, lo que ex-

cluye el empleo de pilas. Por el contrario, la alimentación por baterías resulta idónea, a condición de que se trate de tipos modernos, estancas a los gases. Una batería de 6 V/4 amperios-hora de 66 mm. × 3 mm. × 127 mm. pesa unos 700 gramos; el tipo de 12 V/4,5 amperios-hora es más grande y más pesado, pero sigue siendo utilizable.

La carga de estas baterías se realiza con la ayuda de una tensión constante de 6,9 V para una tensión nominal de 6 V (para una carga rápida, la tensión es todavía más elevada). El circuito propuesto (fig. 3), basado en el regulador de tensión integrado 723, puede servir para cargar independientemente la batería o como un cargador/alimentador combinados para un sistema de alarma que active la alimentación de la batería en el caso de fallo de la red. El circuito 723 es también útil como un estabilizador si se desea alimentar el detector de gases por batería solamente. En este último caso, deben omitirse todos los componentes mostrados en la figura 1 a la izquierda de IC1.

Cuando hay que elegir entre una batería de 6 V o una de 12 V, es preciso tener en cuenta la observación siguiente: para una batería de 6 V hay que emplear un circuito integrado LM 2930 de National en lugar de un C.I. 7805 para IC1. Este regulador especial se distingue por el hecho de que continúa funcionando con una tensión diferencial muy pequeña (0,6 V) entre la entrada y la salida, por lo que tiene una baja disipación. El patillaje de los dos circuitos integrados reguladores es compatible. Con una batería de 12 V, el 7805 es perfectamente adecuado con tal de que sea objeto de refrigeración.

Antes de que se conecte cualquier batería al circuito, la tensión de salida del C.I. 723 debe ajustarse a 6,9 V o a 13,8 V. La batería debe tener una vida media de 20 horas (sin recarga); transcurrido dicho período, es preciso volver a aplicar al circuito la alimentación correspondiente.

Característica	812	813
Resistencia del filamento:		
Resistencia del sensor:		
R_s (1.000 ppm)	(38 ± 3) Ω	(30 ± 3) Ω
Disipación de potencia:	1...10 K (isobutano)	5...15 K (metano)
Tensión de uso:	máx. 15 mW	máx. 15 mW
Disipación del filamento:	máx. 24 V	máx. 24 V
Tensión del filamento:	650 mW	830 mW
Período de precaldeo:	5 V 0,2 V	5 V 0,2 V
Color:	2 minutos naranja	2 minutos negro

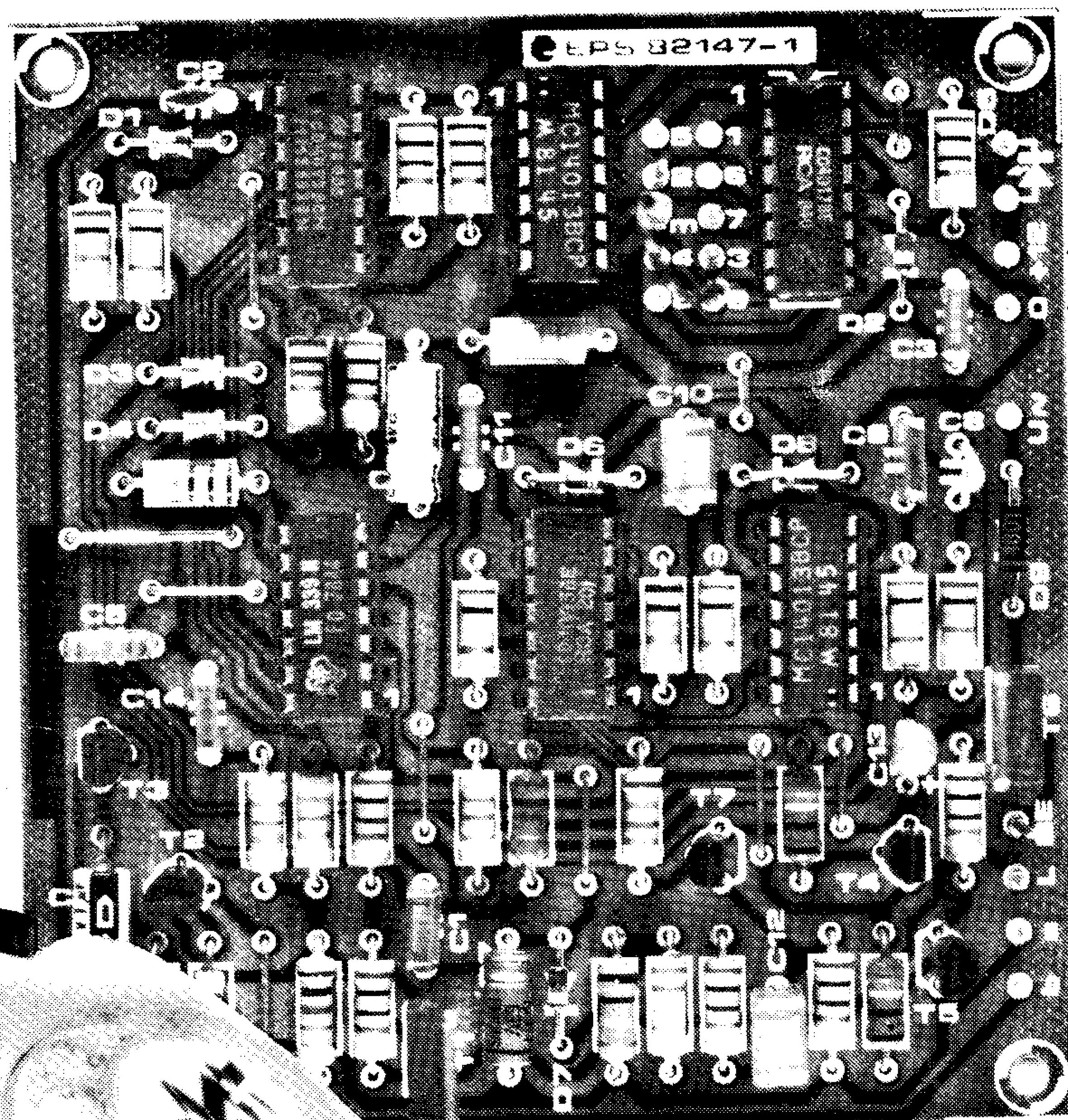
Características más importantes

sistema de telefonía interior

con un máximo de 9 teléfonos.... ¡Y sin centralita!

En cualquier mercadillo de oportunidades y por supuesto en el rastro de las ciudades, es muy común encontrar teléfonos a precio de saldo. Nuestro repetido encuentro con estos aparatos constituye el punto de arranque de una aventura apasionante, que culmina con la realización de la instalación de telefonía interior que hoy les proponemos.

La instalación es de lo más simple y atractiva; no precisa de central alguna y las diversas unidades están construidas con un mismo tipo de circuito impreso.



Dentro de su banalidad cotidiana, el teléfono posee una magia que nadie puede desmentir. ¿Ha paseado ya la vista sobre el esquema de la figura 2? No es tan simple... ¿verdad? Pues bien, es necesario un circuito semejante por cada extensión (puesto de teléfono). En cualquier caso, esta complicación circuital no parece tan terrible si tenemos en cuenta que en nuestra instalación no existe centralita; en consecuencia, los puestos de teléfono deben estar dotados de una cierta inteligencia.

En principio, es necesario que desde cualquier aparato pueda llamarse a cualquier otro; esta posibilidad se consigue por medio de los impulsos producidos por el disco marcador al girar. Una vez que se ha compuesto el número, es preciso activar una señal acústica. Cada extensión dispondrá, en con-

La comunicación entre los distintos puestos se establece a través de un cable de cuatro conductores.

1

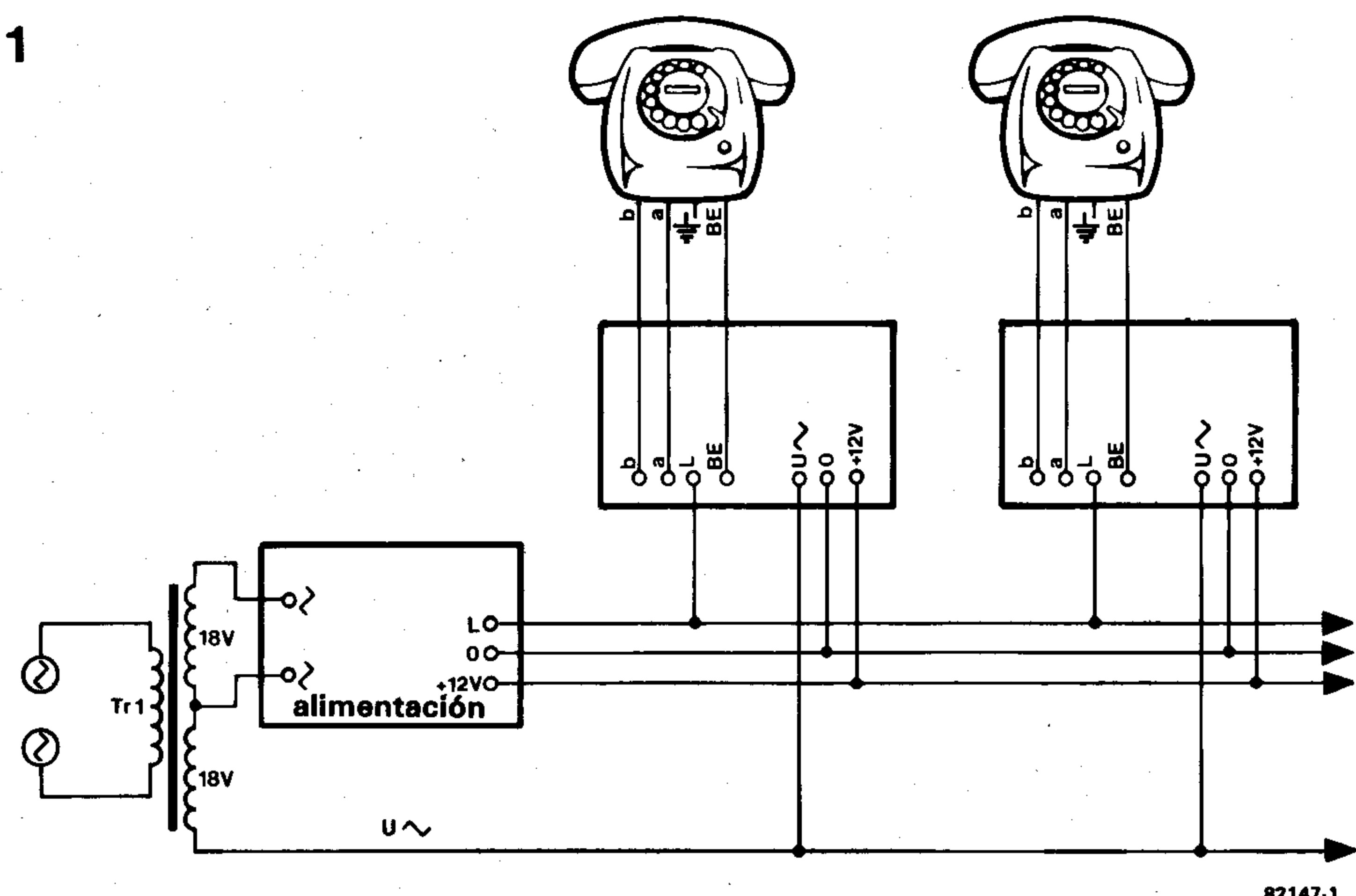


Figura 1. Diagrama del cableado de la instalación de telefonía interior. Los diversos puestos o extensiones (hasta un máximo de 9) están conectados a través de un cable de cuatro conductores.

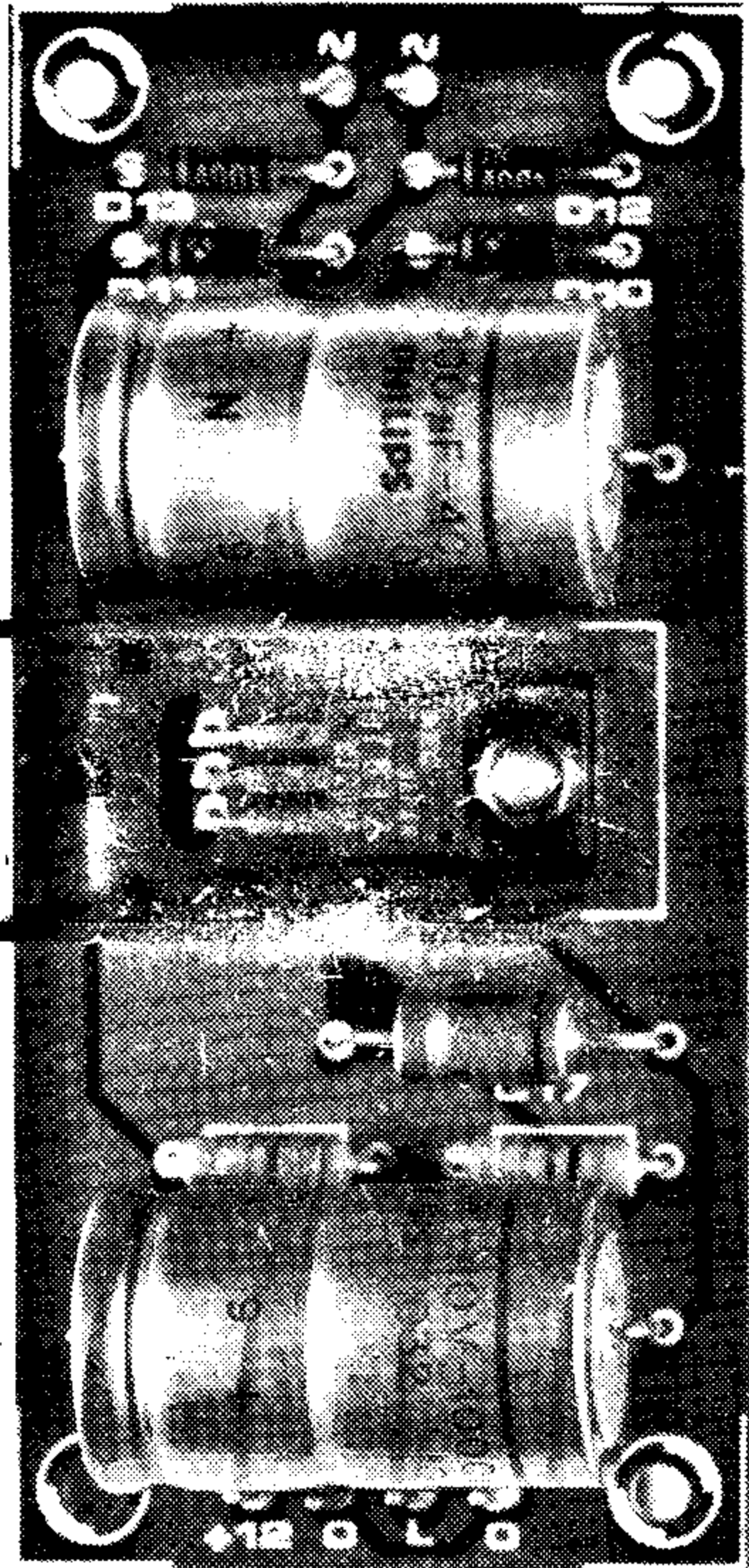
secuencia, de un dispositivo de análisis de la señal de llamada que le servirá para reconocerse. Cuando el número es el correcto para una extensión determinada, se activará en ésta su avisador acústico: en estas condiciones, al descolgar el teléfono, el micro y el altavoz del aparato «llamado» entran en con-

tacto directo con el altavoz y el micro de la extensión origen de la llamada. Naturalmente, es preciso evitar que el contenido de las conversaciones quede al alcance de oídos indiscretos, situación que nuestro diseño resuelve a la perfección, evitando que un tercer aparato pueda «intervenir la línea». Estas son algunas de las características de la instalación de telefonía interior que hemos diseñado, no sin prestar total atención a la lista de componentes con el fin de mantener la factura dentro de un margen razonable.

¿Dónde está la central?

Tal como observarán a lo largo del presente artículo, nuestro sistema de telefonía interior no tiene nada que envidiar a los sistemas que existen en el mercado. Puede recibir hasta 9 puestos o extensiones distintas lo que, sin lugar a dudas, es ampliamente suficiente. En la mayor parte de las instalaciones, sea en su propio hogar o en la oficina, no será necesario disponer de más de 3 ó 4 extensiones. La gran ventaja de nuestro sistema reside en la ausencia de central, elemento que suele constituir la parte más compleja y onerosa de este tipo de sistemas. Así pues, nuestro sistema de telefonía interior constituye una red «axial» y no en «estrella». Los diversos puestos están enlazados por un cable de cuatro conductores (especial para telefonía, o de tipo audio blindado) conforme a las indicaciones de la figura 1.

El consumo de corriente de cada puesto o extensión es muy discreto, lo que permite utilizar una sola alimentación central, común a todas las extensiones. De ella parten cuatro cables que acceden a cada uno de los puestos: una línea de «avisador acústico» (U), una línea de alimentación (+ 5V), una línea de masa (0) y línea de palabra (L). El puesto telefónico propiamente dicho está conectado a su correspondiente circuito impreso a través de los puntos a, b, y BE. Más adelante hablaremos detalladamente de estas conexiones. En todo caso, nos parece oportuno adelantarle que el conexionado no presenta problema alguno.

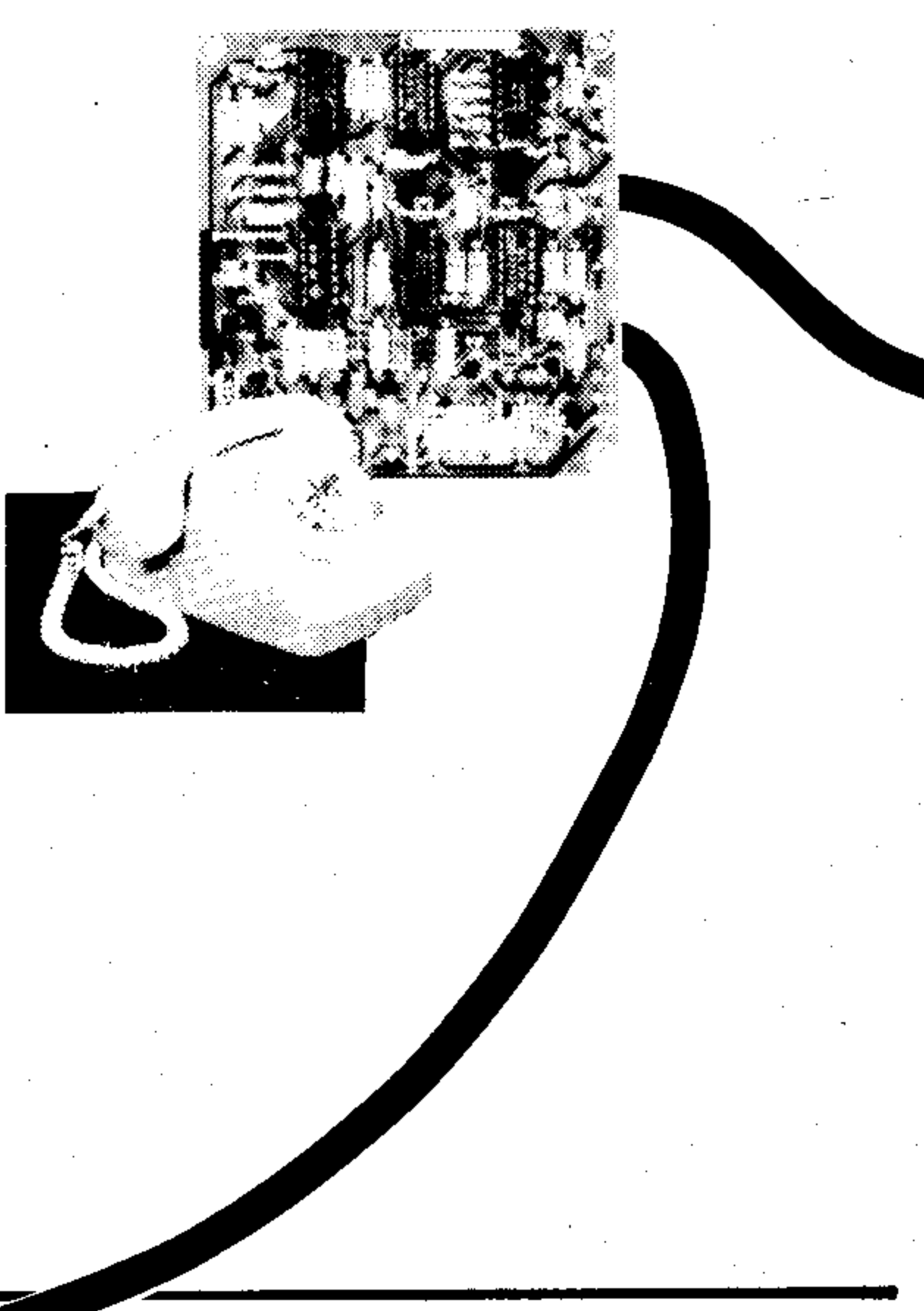


Protección de las conversaciones

A cada uno de los puestos se le atribuye un número comprendido entre 1 y 9. Esta asignación se realiza colocando un puente en su tarjeta de circuito impreso, uniendo una de las salidas de IC1 a la entrada de FFI. Al marcar en el disco el número atribuido a uno de los puestos, se activará el avisador acústico del puesto en cuestión, mientras que en el puesto que se ha efectuado la llamada se percibirá una señal sonora intermitente. Cada puesto está dotado de un LED de «ocupado». *Un tercer puesto del sistema no puede intervenir la conversación en curso.* En muchos casos puede resultar interesante disponer de una opción muy útil: nos referimos a las conversaciones de tipo «conferencia» en las que pueden participar varios puestos a la vez, por supuesto, de común acuerdo. Veamos un ejemplo; usted (¡promotor de la conferencia!) quiere comunicar simultáneamente con las extensiones números 3 y 7, solución: basta con que usted marque en el disco de su teléfono el 3 y luego el 4 ¿Por qué el 4 si lo que pretendo es hablar con el 7?... pues, simplemente, porque la diferencia entre 3 y 7 es 4. Así, por ejemplo, para entrar en comunicación con los puestos 3 y 8, hay que marcar los números 3 y 5; si el segundo comunicante debe ser el puesto 9 en lugar del 8, marcaremos el 3 y el 6 ¡Fácil... no es cierto!

Funcionamiento

Cada puesto o extensión está dotado de un circuito análogo al que aparece en la figura 2: una electrónica algo compleja pero muy fácil de construir sobre la tarjeta de circuito impreso. Cuando el microteléfono descansa sobre el receptor telefónico (situación de reposo), la línea L no consume corriente alguna, de tal forma que la tensión es de



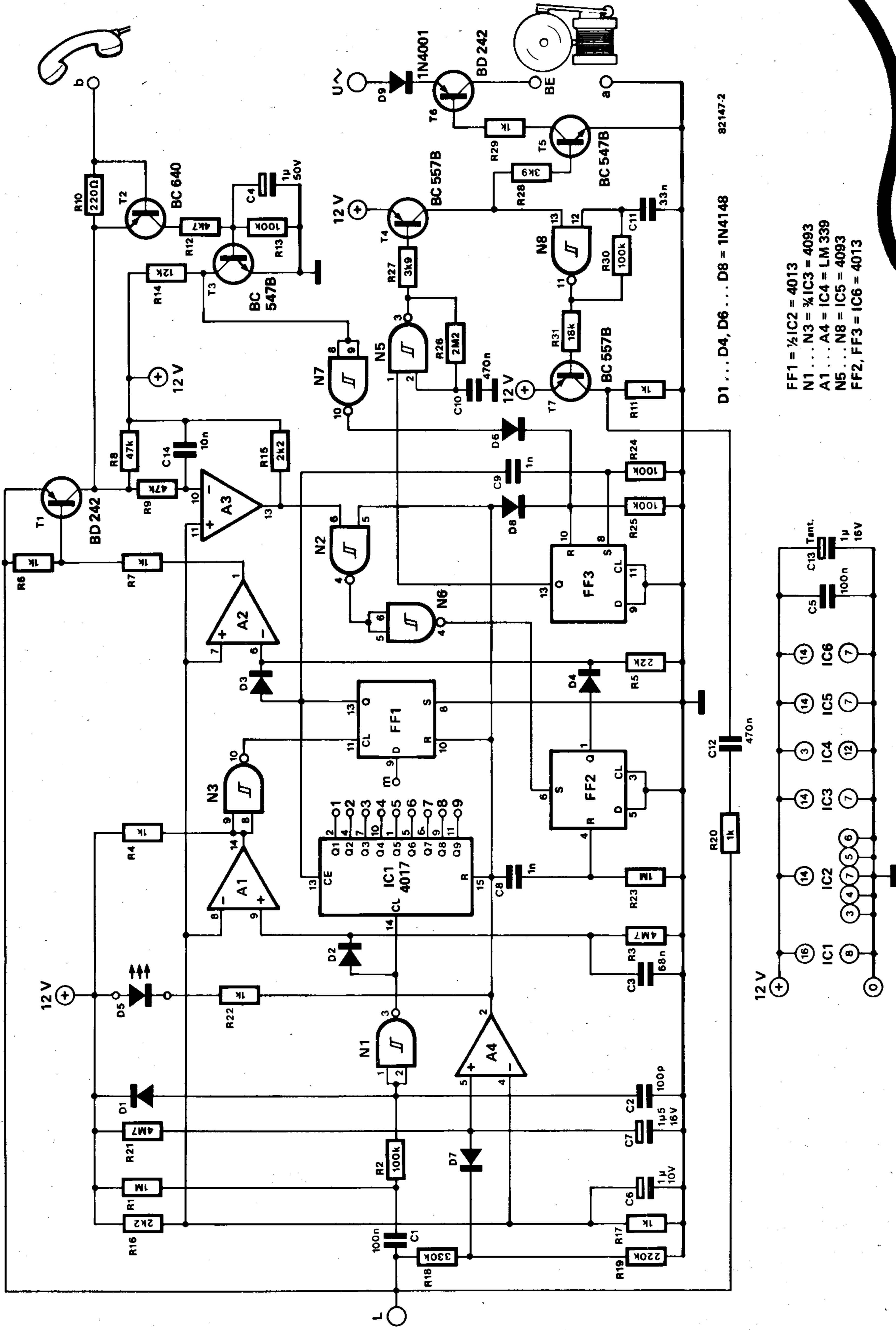


Figura 2. Esquema del circuito necesario para cada uno de los puestos. Su relativa complejidad se debe al hecho de que el sistema prescind

25...30 V (ver alimentación en la figura 3). Al descolgar el microteléfono y marcar un número en el disco, esta línea se ve cortocircuitada: la patilla 10 de A3 pasa a nivel lógico bajo, de tal forma que el biestable FF2 es posicionado por N2 y N6. En este instante, el aparato está en línea. Como quiera que la línea de palabra está cortocircuitada, la salida del comparador A4 va a pasar a nivel lógico bajo, con lo que se activan, a su vez, el contador IC1 y el biestable FF1. Cuando se marca un número, el disco marcador entrega siempre dos impulsos más de la cifra correspondiente: cada impuso dura 38 ms y cada intervalo entre dos impulsos 62 ms. Estos impulsos llegan, a través de la línea L, a los contadores IC1 de todos los puestos o extensiones. Los componentes A1, D2, R3 y C3 constituyen un multivibrador

monoestable redispensible. Unos 0,2 segundos después de la llegada del último impulso, A1 conmuta y la salida Q de FF1 pasa a nivel lógico alto. En consecuencia, el contador del puesto «llamado» se bloquea, aunque, desde luego, su contenido no se pierde. En estas condiciones, el aparato en cuestión entra en línea a través de A2 y T1. Si ahora se descuelga el microteléfono, FF3 es puesto a cero y la señal acústica cesa su actividad. Este biestable fue posicionado en el instante adecuado por el impulso CE del contador IC1. En todos los restantes puestos de la instalación telefónica, FF1 entrega un nivel lógico bajo, de tal forma que las conversaciones no pueden ser oídas desde cualquier otra extensión distinta de la que ha recibido la llamada; el LED D5 permanecerá encendido indicando que la línea está ocupada. Dado que los contadores no han sido puestos a cero, es posible conectar con cualquier otro puesto a partir de la extensión que ha realizado la llamada. Al colgar los microteléfonos, la salida de A4 regresa a nivel lógico alto, con lo que todos los contadores y el biestable FF2 pasan a estado lógico bajo. Adicionalmente, el LED D5 también se apagará. Si el microteléfono del puesto que recibe la llamada no es descolgado, el biestable FF3 pasará a cero en el instante en el que se cuelgue el microteléfono de la extensión que ha realizado la llamada. Esta circunstancia coincidirá con el cese de la señal acústica de llamada.

Alimentación

La figura 3 muestra el esquema del circuito de alimentación diseñado para nuestro sistema de telefonía interior. Como observarán, no puede ser más elemental. Además de la tensión de alimentación de +5 V (proporcionada por un 7805), nos encontramos con las conexiones para las líneas L y U. La primera es la línea de palabra que recibe, a través de R32, R33 y C16, una tensión rectificada aunque no estabilizada. Según el transformador empleado, esta tensión —necesaria para los micros y los auriculares de carbón— será de 25 a 30 V. De los dos devanados secundarios del trans-

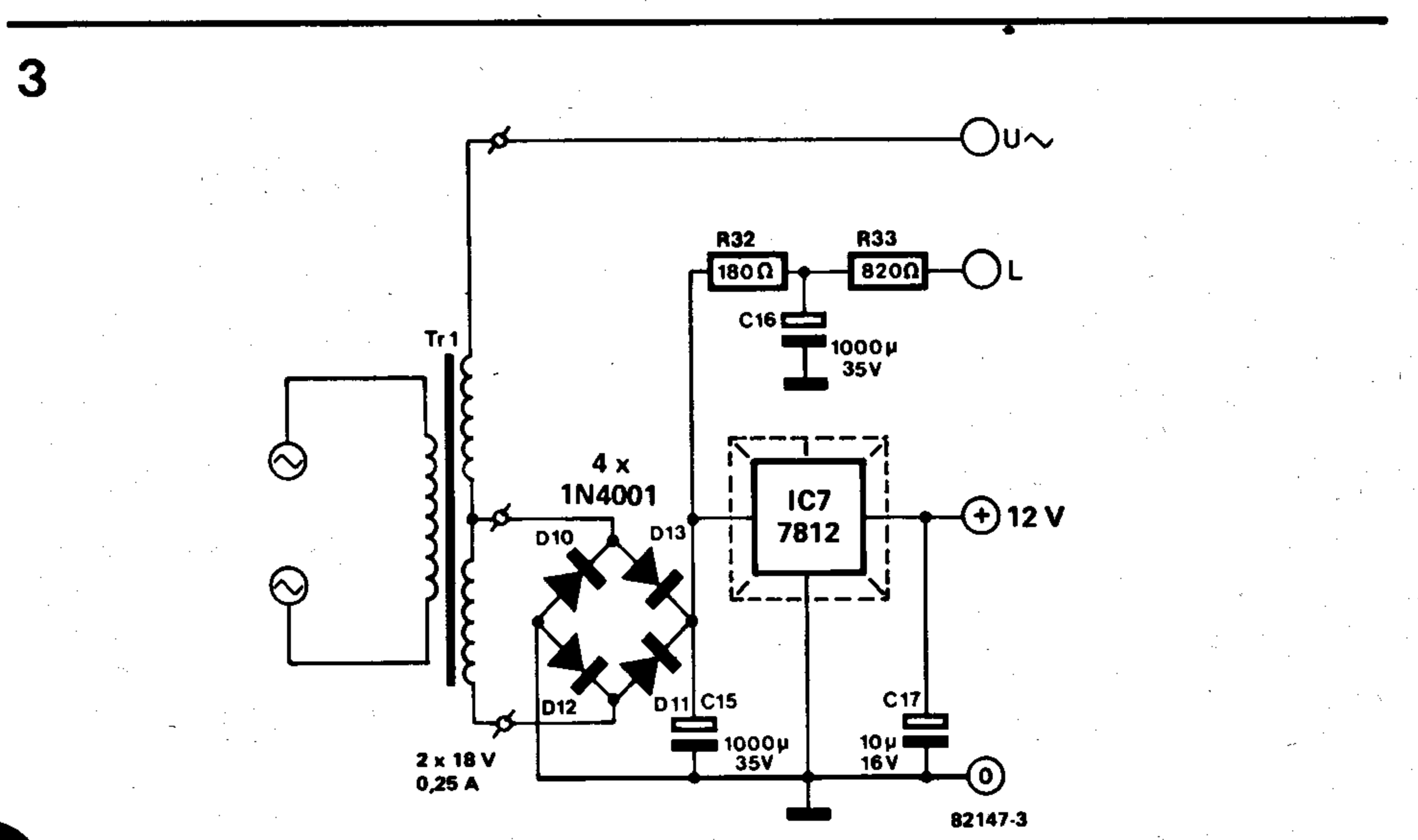


Figura 3. La alimentación es común a todos los puestos del sistema. Además de la línea de +5 V, la fuente de alimentación entrega la tensión L para los micros y auriculares de carbón y la tensión alterna U para los zumbadores de llamada.

formador sólo uno de ellos se utiliza para obtener la alimentación de +5 V. El otro se emplea para los zumbadores de llamada que precisan de una tensión relativamente elevada; éste es el motivo por el que ambos devanados se asocian en serie, con objeto de disponer de una tensión alterna para los zumbadores de 36 V en el punto U.

Montaje

Las figuras 4 y 5 muestran los circuitos impresos serigrafiados que se han diseñado para el montaje de la circuitería de cada puesto y de la alimentación del sistema. Por lo que respecta a los circuitos impresos de «puesto», serán necesarios tantos como extensiones tenga la instalación telefónica (recordemos que el máximo es 9). Las dimensiones del circuito de la figura 4 deben permitir montarlo directamente en el interior del receptor telefónico, a menos que se prefiera alojarlo en una caja-consola independiente del receptor. En cualquier de los casos, el LED D5 debe situarse en un lugar tal que se perciba a simple vista. En el circuito impreso de la figura 4 nos encontramos, entre IC1 e IC2, con los puntos 1...9 y m; las salidas del comparador corresponden a los números de los nueve puestos que admite el sistema, mientras que el punto m corresponde a la entrada de FF1. Este último punto debe conectarse a uno de los identificados con las cifras 1...9. El circuito impreso correspondiente a cada puesto debe incorporar un puente que una el punto m con la salida de IC1 cuyo número coincida con el asignado al puesto en cuestión. Cuando en la instalación intervengan varios puestos, será necesario dotar al regulador IC7 del adecuado radiador.

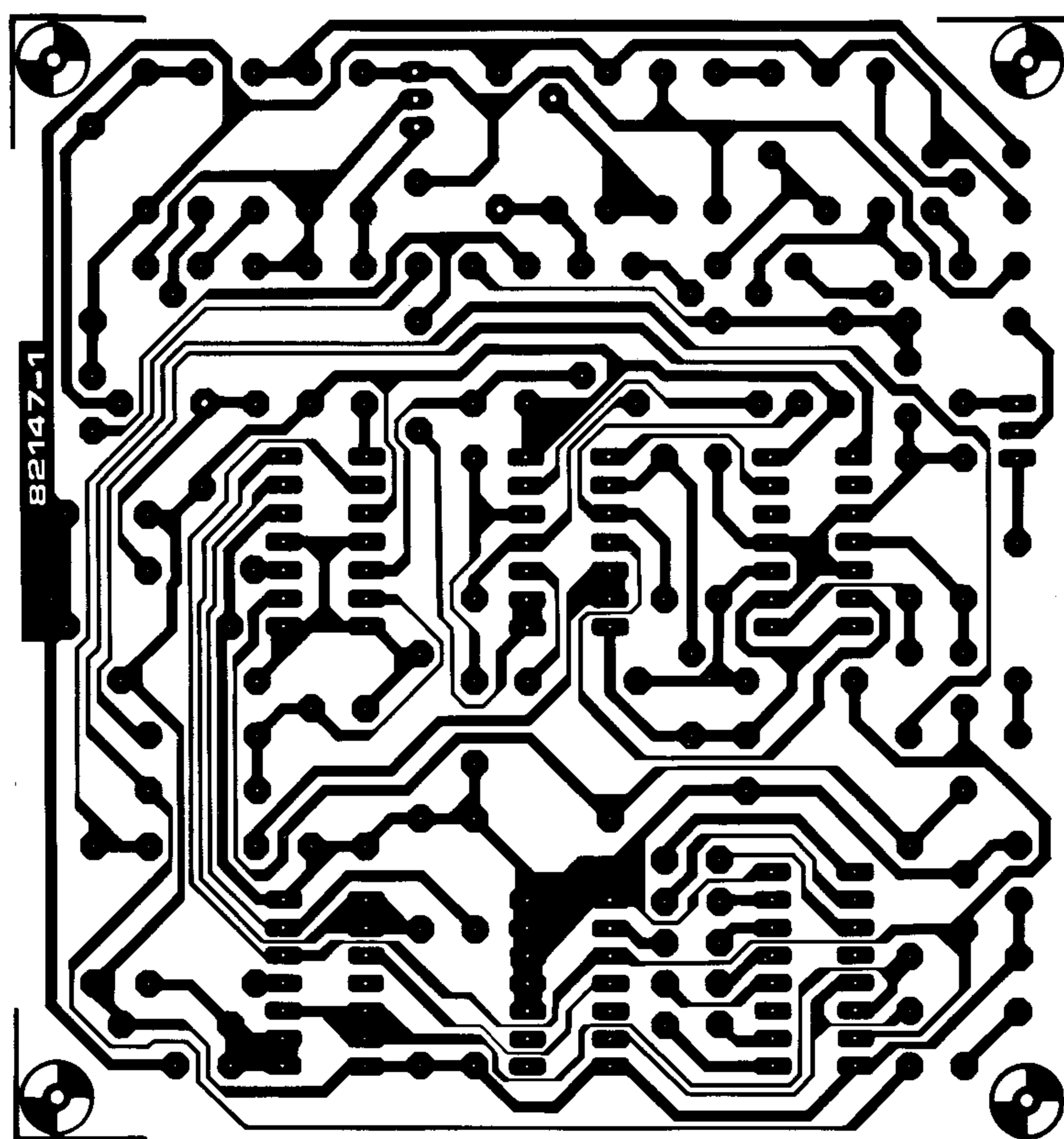
Cableado

En la figura 1 se detalla la distribución del cableado a realizar. Su realización es extre-

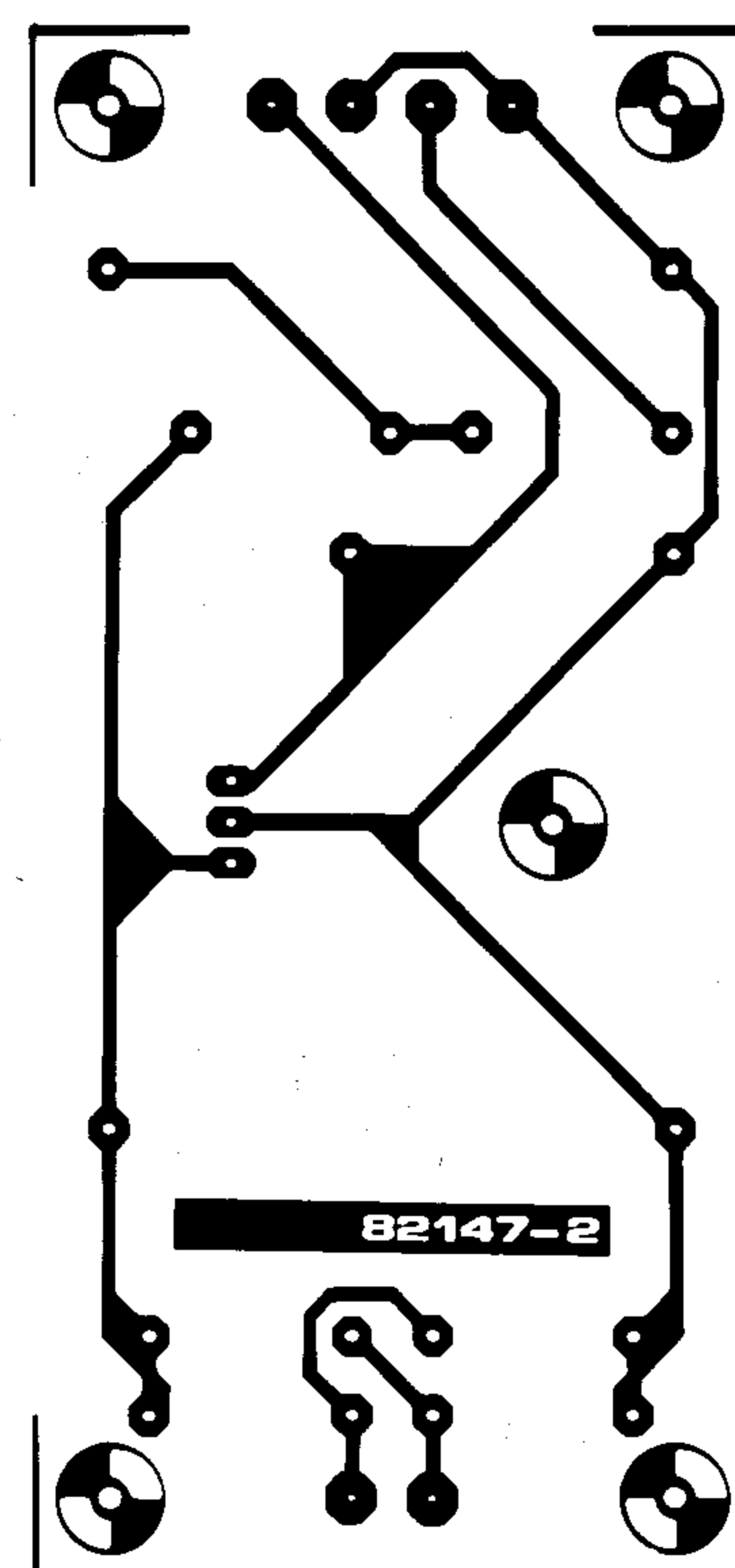
Lista de componentes para las figuras 2 y 4

- Resistencias:
- R1,R23 = 1 M
 - R2,R13,R24,R25,R30 = 100 k
 - R3,R21 = 4M7
 - R4,R6,R7,R11,R17,R20,R22,R29 = 1 k
 - R5 = 22 k
 - R8,R9 = 47 k
 - R10 = 220 Ω
 - R12 = 4k7
 - R14 = 12 k
 - R15,R16 = 2k2
 - R18 = 330 k
 - R19 = 220 k
 - R26 = 2M2
 - R27,R28 = 3k9
 - R31 = 18 k
- Condensadores:
- C1,C5 = 100 n
 - C2 = 100 p
 - C3 = 68 n
 - C4 = 1 µ/50 V
 - C6 = 1 µ/16 V
 - C7 = 1µ5/16 V
 - C8,C9 = 1 n
 - C10,C12 = 470 n
 - C11 = 33 n
 - C13 = 1 µ/16 V tántalo
 - C14 = 10 n
- Semiconductores:
- T1,T6 = BD 242
 - T2 = BC 640
 - T3,T5 = BC 547B
 - T4,T7 = BC 557B
 - D1...D4,D6,D7,D8 = 1N4148
 - D5 = LED
 - D9 = 1N4001
 - IC1 = 4017
 - IC2,IC6 = 4013
 - IC3,IC5 = 4093
 - IC4 = LM 339

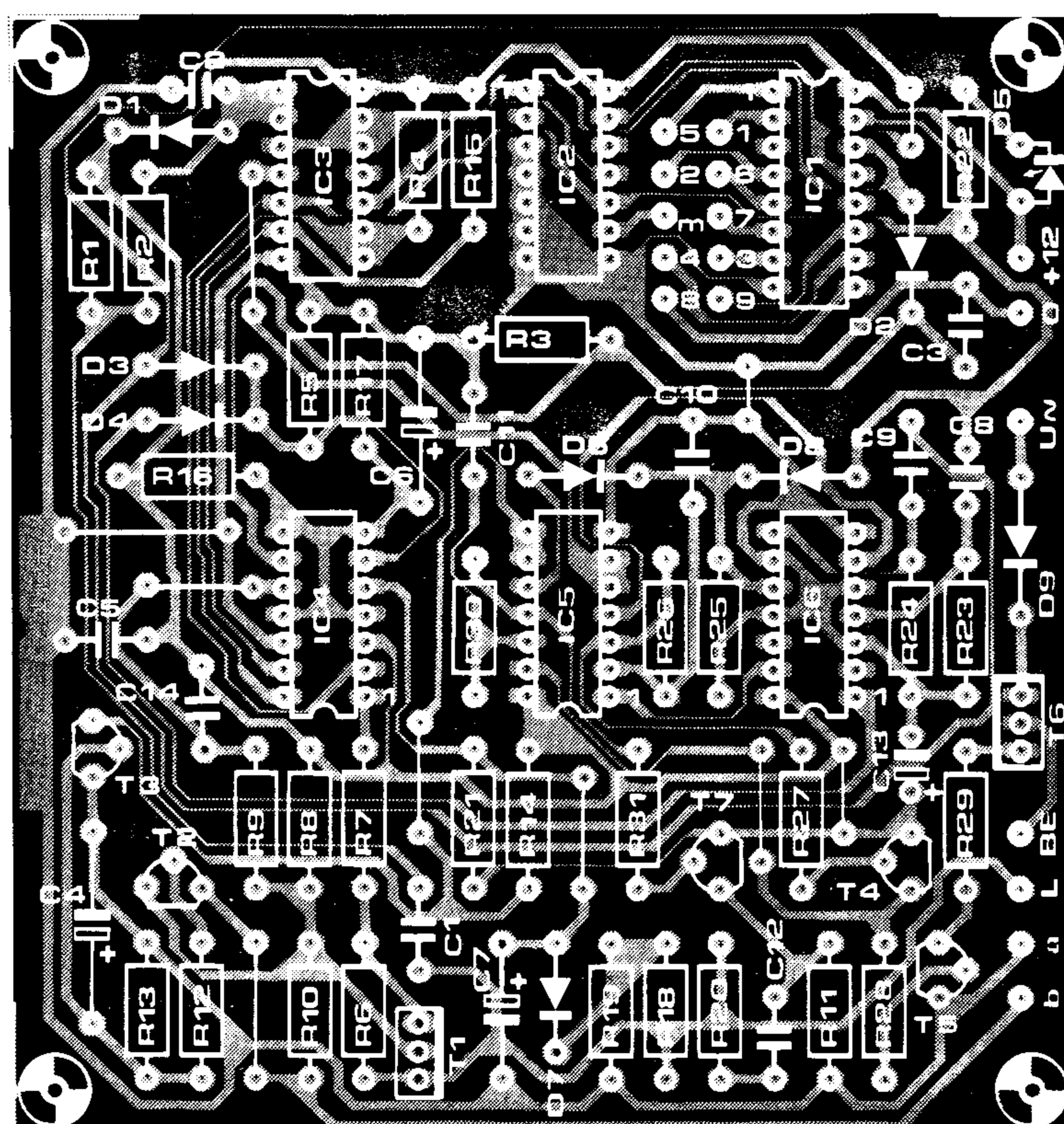
4a



5a



4b



5b

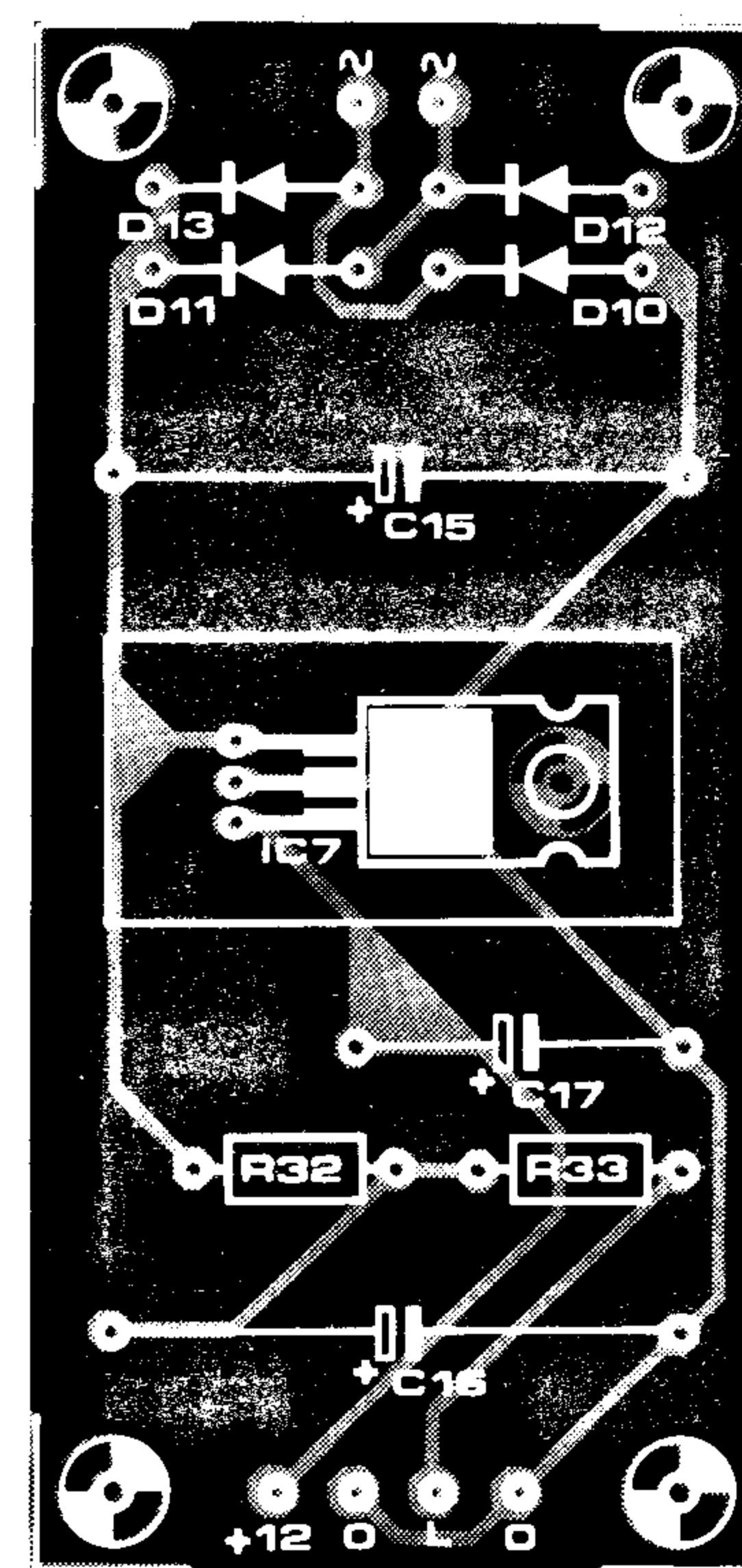


Figura 4. Trazado de las pistas de cobre y serigrafía del circuito impreso de puesto telefónico. ¡Muy importante! Debido a un pequeño error de serigrafía, algunos de los circuitos impresos poseen la indicación +12 en lugar de +5 en el punto de entrada de la tensión de alimentación. Por supuesto, en cualquiera de los casos, la tensión correcta es la de +5 V.

Figura 5. Circuito impreso de la fuente de alimentación común para todo el sistema.

Lista de componentes para las figuras 3 y 5

Resistencias:

R32 = 180 Ω
R33 = 820 Ω

Condensadores:

C15, C16 = 1000 μ /35 V
C17 = 10 μ /16 V

Semiconductores:

D10... D13 = 1N4001
IC7 = 7812

Varios:

Tr1 = transformador 2 x 18 V/250 mA
Radiador para IC7

madamente simple. La única salvedad a reseñar es que el punto U $\sim$ no figura en el circuito impreso de alimentación: su conexión se tomará directamente del secundario del transformador.

Por lo que respecta a los receptores telefónicos que debe adquirir «en los saldos», el tema es un tanto más delicado. Por supuesto, es necesario abrirlos para localizar los puntos a, b y BE. Los puntos a y b deben estar situados en la zona de unión con el microteléfono, ya que a través de ellos se establece la conexión del micrófono y del auricular. El zumbador de llamada suele encontrarse entre los dos puntos referidos, a través de un condensador.

En la figura 6 se reproduce el esquema simplificado de las conexiones interiores de un teléfono de tipo común. Para que sean adaptables a nuestra instalación, es preciso reformar su conexionado interno según se indica en la figura 6b; esto es: el zumbador o campanilla debe disponer de una conexión independiente (BE). En todo caso, la modificación es muy fácil de realizar.

Si por cualquier circunstancia no es capaz de identificar a simple vista las conexiones interiores al receptor telefónico, puede utilizar una tensión alterna de 18 V (extraída del secundario del transformador) y aplicarla entre dos puntos no identificados. Si ello se traduce en la activación de la señal acústica de llamada, estará en condiciones de afirmar que se trata de los puntos a y BE. Si en lugar de la campanilla de llamada percibe un ronquido en el auricular, habrá localizado los puntos a y b. Para evitar que con estos ensayos empíricos pueda dañar al receptor, debe colocar un condensador de 100 μ F/35 V en serie con la tensión alterna de prueba.

Consejos finales

La tensión de referencia de los comparadores A1...A4 es de 2,5 V. A4 actúa correctamente si su salida pasa a nivel lógico bajo en el instante preciso en el que se descuelga el microteléfono; una buena referencia de esta situación la obtendrá comprobando que el LED D5 se ilumina en el momento en el que se descuelga algún microteléfono. Si éste no es el caso, será necesario modificar el valor de la resistencia R17. Si el LED sigue sin evidenciar un correcto funcionamiento, puede ser preciso modificar también la magnitud de R18.

El zumbador de llamada es desactivado por el sensor de corriente construido en torno a T2. En la práctica no tienen por qué existir problemas, dado que la mayor parte de receptores telefónicos necesitan entre 20 y 25 mA. Al descolgar el microteléfono, la tensión en la patilla 10 de A3 cae a cero (o a un valor muy próximo) y el puesto «llamado» entra en línea.

Si está en condiciones de permitirse algún gasto suplementario, puede plantearse la posibilidad de alimentar cada puesto por separado, lo que le permitirá economizar dos conductores del cable que enlaza las diver-

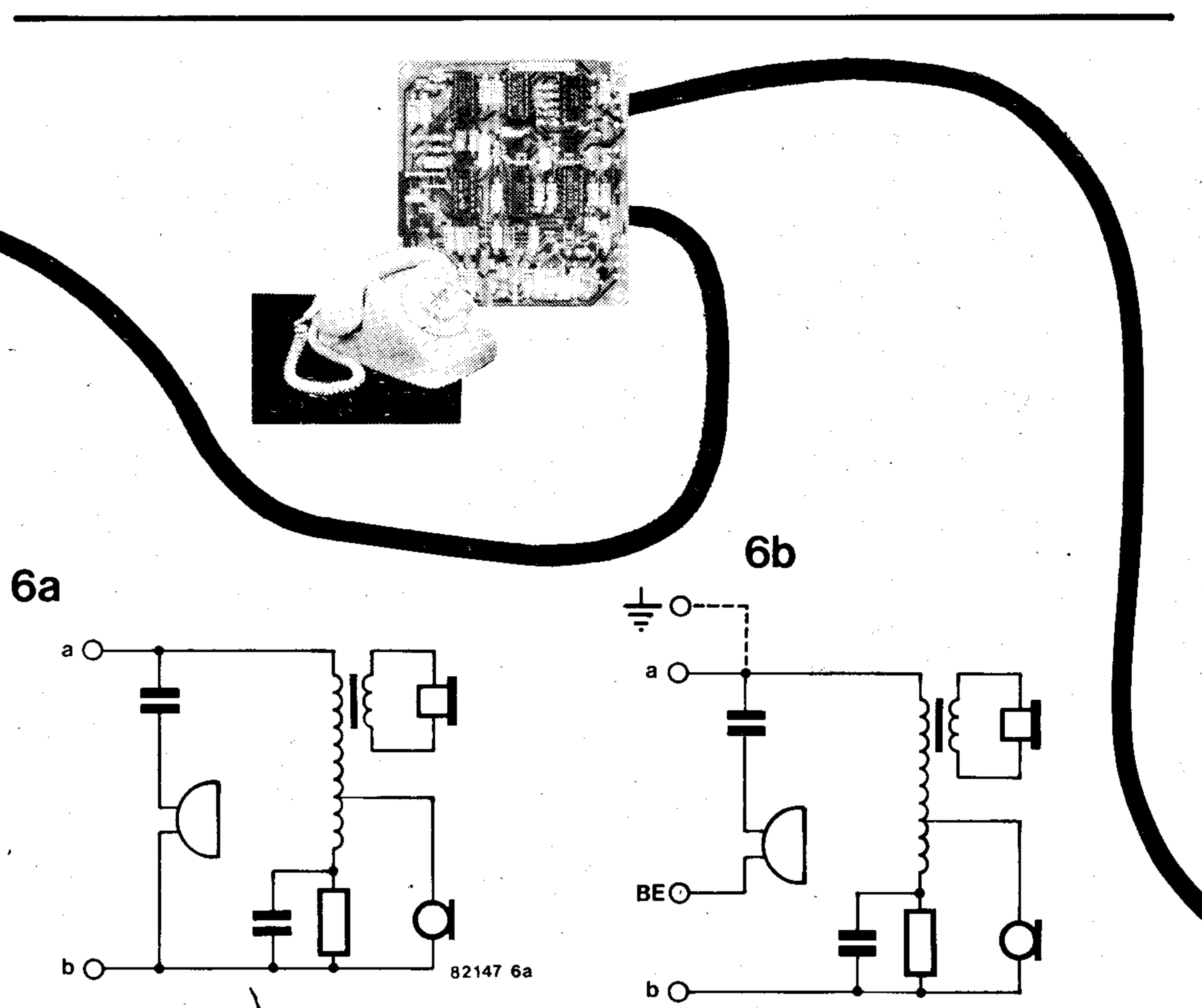
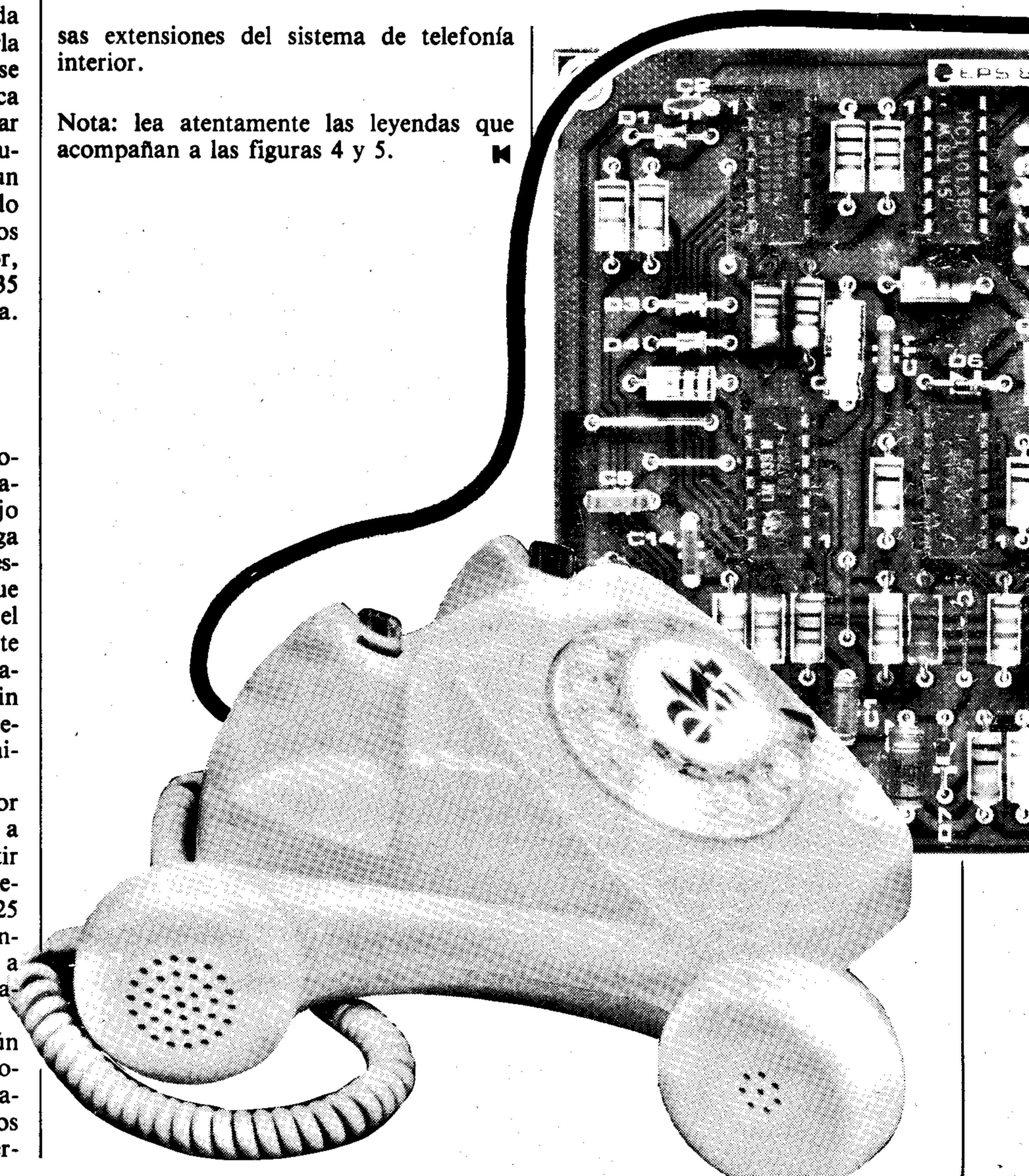


Figura 6. Esquema interior de un receptor telefónico común. Si el zumbador no dispone de una conexión independiente de los puntos a y b, será necesario modificar el conexionado interno conforme a las indicaciones del esquema de la figura 6b.

sas extensiones del sistema de telefonía interior.

Nota: lea atentamente las leyendas que acompañan a las figuras 4 y 5.



BASIC

4ª PARTE

¡Hemos llegado a la cuarta lección! Aunque de forma rudimentaria, a estas alturas somos ya capaces de establecer nuestros primeros «diálogos» con el ordenador en su lenguaje de alto nivel. Para seguir adelante y alcanzar el adecuado nivel de virtuosismo, nada más oportuno que incrementar nuestro vocabulario «lingüístico» con un nuevo grupo de comandos BASIC.

Al igual que en el capítulo anterior, vamos a realizar el estudio de los comandos atendiendo a un orden didáctico. Cada uno de los comandos abordados nos permitirán ampliar nuestro repertorio de instrucciones para crear programas de mayor complejidad o controlar con mayor eficacia las tareas de edición. En esta ocasión, vamos a empezar precisamente con un comando de esta última categoría que facilita la posibilidad de listar el programa en curso —sólo determinadas líneas o en su totalidad— en la pantalla del ordenador.

El comando LIST

LIST	Tipo: comando de control Función: listar
------	---

Definición: El comando de control LIST constituye una instrucción de tipo directo (sin número de línea) que se utiliza para pedir al ordenador que liste en la pantalla del sistema el programa en curso, en su totalidad o sólo determinadas líneas.

Estructura:

LIST	n.º línea primera — n.º línea final
------	-------------------------------------

Ejemplos: LIST
LIST 20
LIST 30—70
LIST 45—
LIST —50

En los ejemplos se observa que el comando LIST permite construir cinco tipos de instrucciones directas, a las que el intérprete BASIC responderá con ejecuciones distintas.

a. LIST

El intérprete BASIC listará por pantalla la totalidad del programa en curso (con el que estamos trabajando en este instante).

b. LIST 20

Se listará exclusivamente la instrucción cuyo número de línea es 20.

c. LIST 30—70

El intérprete listará la zona de programa constituida por las instrucciones cuyo número de línea esté comprendido entre 30 y 70.

d. LIST 45—

En el listado aparecerá desde la instrucción 45 hasta el final del programa.

e. LIST —50

La ejecución se traduce en el listado de la porción de programa comprendida entre la primera instrucción y la 50.

Por si alberga alguna duda al respecto, veamos como responde el ordenador ante las diversas formulaciones de LIST.

```
10 REM ==DEMOSTRACION==
20 REM INSTRUCCION LIST
30 READ A1,Z$
40 PRINT A1;47/A1
50 PRINT Z$
60 PRINT "====="
70 DATA 25.67,"!HALE HOP!"
80 END
```

```
LIST
10 REM ==DEMOSTRACION==
20 REM INSTRUCCION LIST
30 READ A1,Z$
40 PRINT A1;47/A1
50 PRINT Z$
60 PRINT "====="
70 DATA 25.67,"!HALE HOP!"
80 END
```

```
LIST 20
20 REM INSTRUCCION LIST

LIST 30-70
30 READ A1,Z$
40 PRINT A1;47/A1
50 PRINT Z$
60 PRINT "====="
70 DATA 25.67,"!HALE HOP!"
80 END
```

```
LIST 45-
50 PRINT Z$
60 PRINT "====="
70 DATA 25.67,"!HALE HOP!"
80 END
```

```
LIST -50
10 REM ==DEMOSTRACION==
20 REM INSTRUCCION LIST
30 READ A1,Z$
40 PRINT A1;47/A1
50 PRINT Z$
```

BASIC

4ª PARTE

Al ejecutar la instrucción LIST 45— se observa una particularidad característica de este tipo de instrucciones: cuando LIST se utiliza en la variante d), el listado se realiza a partir del número de línea que figura en el argumento, aunque éste no exista. Así, por ejemplo, al no existir la línea 45, el ordenador empezará a listar a partir de la línea inmediatamente siguiente.

Otros formatos de LIST

Para algunos intérpretes BASIC, la formulación de las instrucciones LIST no coincide exactamente con la indicada como norma general en los casos c, d y e definidos anteriormente. A continuación, indicamos el formato alternativo más común, aplicándolo a los tres casos a los que hemos hecho referencia.

- c. LIST 30, 70: La diferencia respecto al caso general radica en que el guión separador se sustituye por una coma.
- d. LIST 45, 70: En esta ocasión, para listar desde una instrucción determinada hasta el final del programa, es preciso actuar como en el caso (c), indicando el número de línea de la última instrucción. Si el 70 se sustituyera por cualquier otro número mayor, la actuación sería idéntica, puesto que el listado concluiría en la última línea de programa (70 en el ejemplo).
- e. LIST 0, 50: De nuevo, para listar desde el principio del programa hasta una instrucción específica hay que delimitar totalmente el bloque, indicando los números de las líneas primera y última.

El intérprete BASIC del Junior Computer pertenece al grupo más común, constituido por los que aceptan la formulación de instrucciones LIST definida como caso general.

El comando NEW

NEW	Tipo: comando de control Función: borrado
-----	--

Definición: Borra por completo el programa en curso y todas las variables definidas.

Estructura: NEW

Ejemplos: NEW

Al igual que el comando LIST, NEW se utiliza exclusivamente constituyendo instrucciones directas (sin número de línea). Por lo demás, no incorpora argumento alguno.

NEW en otros dialectos BASIC

En algunos intérpretes BASIC, el comando que realiza la función de borrado total del programa en curso adopta otra denominación, por ejemplo:

SCRATCH: en su expresión literal (SCRATCH), abreviada (SCR), o en la forma SCRATCH ALL que el ordenador ejecuta borrando la totalidad de la sección de memoria programable por el usuario.

PURGE o DELETE.

Como podrán imaginar, no es necesario incluir ningún tipo de ejemplo para ilustrar el funcionamiento del comando de control NEW... ¡salvo, acaso, mostrarles una pantalla en blanco! En cualquier caso, hay que advertirles acerca de la conveniencia de no utilizar indiscriminadamente el comando NEW, ya que puede ser totalmente nefasto. Para comprobar su temple nervioso, después de invertir un par de horas en la puesta a punto de un complejo programa BASIC... ¡pueden intentar ejecutarlo con el comando NEW en lugar del RUN!

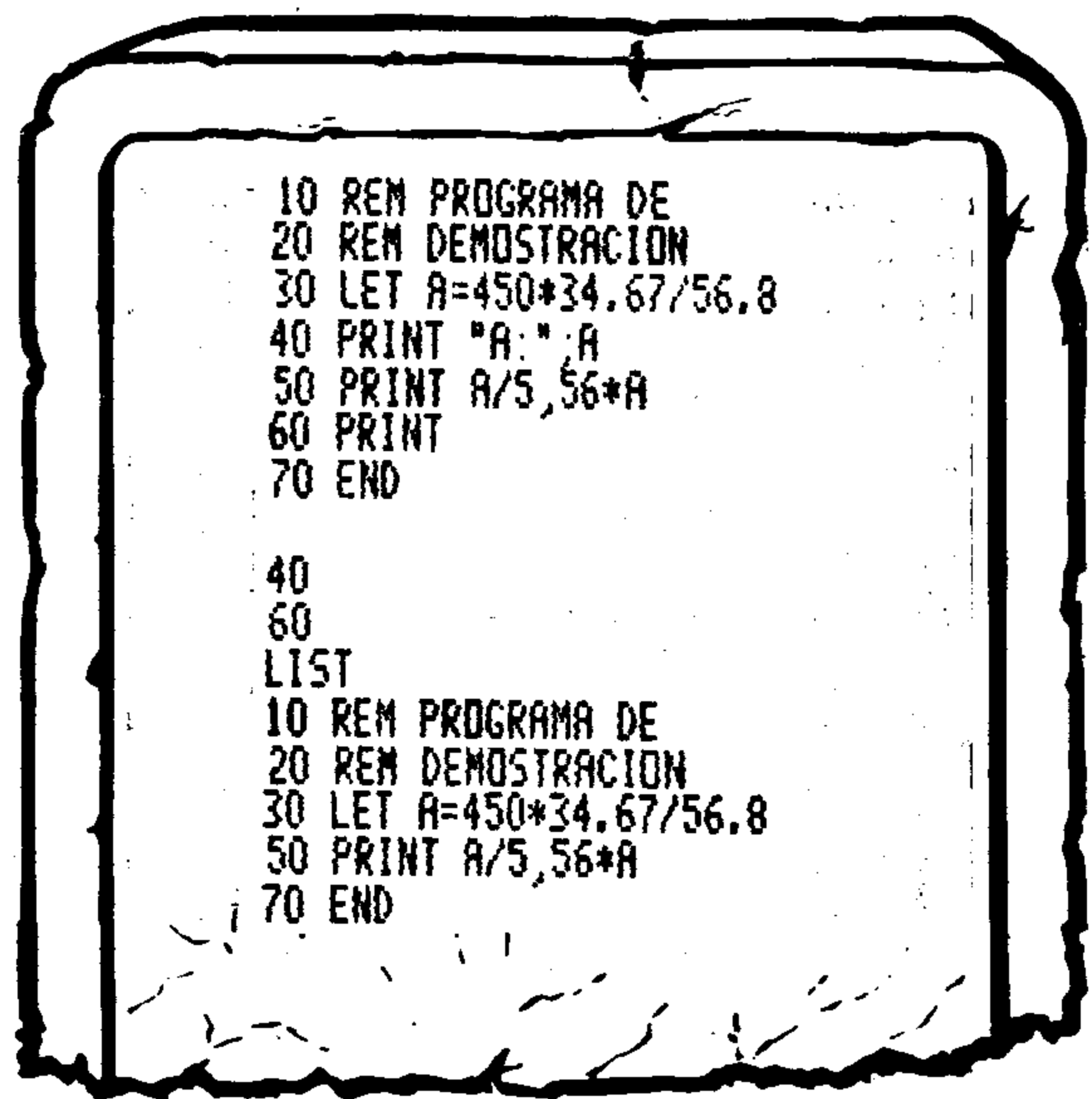
Borrado y modificación de líneas de programa

Los dos comandos precedentes nos permiten listar a voluntad una o varias líneas de programa y borrar el programa completo. No obstante, nos queda por saber aún como borrar o modificar líneas, una a una.

Borrado de una línea de programa

- En la mayor parte de los dialectos BASIC, para borrar una línea del programa basta con escribir el número de la línea a borrar y pulsar la tecla «return» (fin de línea y retroceso de carro).

Veamos:



En el ejemplo observamos que el borrado de las instrucciones 40 y 60 se realiza, simplemente, introduciendo el correspondiente número de línea seguido de una acción sobre la tecla «return» (carriage return).

Aunque el procedimiento habitual es el comentado, algunas versiones BASIC relativamente evolucionadas ad-

BASIC

4ª PARTE

miten determinados comandos específicos para el borrado de líneas, por ejemplo:

DELETE 40: borrado de la línea 40

DELETE 30, 60: borrado de las líneas comprendidas entre la 30 y la 60, incluidas ambas.

Modificación de líneas

- Una línea de programa se modifica escribiéndola de nuevo en la forma oportuna con el mismo número de línea que la instrucción a la que sustituye.

Veamos un ejemplo, partiendo del programa de demostración utilizado en el párrafo anterior.

```
10 REM PROGRAMA DE
20 REM DEMOSTRACION
30 LET A=450*34.67/56.8
40 PRINT "A: " ; A
50 PRINT A/5, 56*A
60 PRINT
70 END
```

A continuación, vamos a modificar las instrucciones cuyo número de línea es 30 y 60:

```
30 LET A=50+68.4*7896.31/6.7
60 PRINT "MODIFICACION DE LINEAS"
LIST
10 REM PROGRAMA DE
20 REM DEMOSTRACION
30 LET A=50+68.4*7896.31/6.7
40 PRINT "A: " ; A
50 PRINT A/5, 56*A
60 PRINT "MODIFICACION DE LINEAS"
70 END
```

El comando INPUT

INPUT	Tipo: comando de entrada/salida Función: entrada de datos
-------	--

Definición: Se utiliza para escribir instrucciones que permiten introducir datos durante el desarrollo del programa. Cuando el ordenador se encuentra con una instrucción de este tipo, pedirá al usuario a través de la pantalla (con el signo «?») que introduzca los datos oportunos.

Estructura:	INPUT	variable/s
	INPUT	comentario; variable/s

Ejemplos: 10 INPUT A2,H\$,R
20 INPUT «POTENCIA»;P

El comando INPUT admite las dos variantes de formulación indicadas en el apartado de estructura.

- a. INPUT A2,H\$,R
- En su forma más habitual, una instrucción INPUT consta del comando y de una zona de argumento constituida por una o más variables (numéricas o de cadenas de caracteres) separadas por el signo (,). Al ejecutar la instrucción, el ordenador presentará en pantalla el símbolo (?) pidiendo al usuario que introduzca los valores a asignar a las variables que aparecen en el argumento. Por ejemplo:

```
10 REM INSTRUCCION INPUT
20 INPUT A2,H$,R
30 PRINT
40 PRINT A2,H$,R
50 END
RUN
? 35,PRUEBA,80.65

35      PRUEBA      80.65
```

- b. INPUT «POTENCIA»; P
- En esta segunda variante de formulación de la instrucción INPUT, el ordenador mostrará el comentario que aparece entre comillas antes de imprimir el carácter (?) solicitando la entrada de datos.

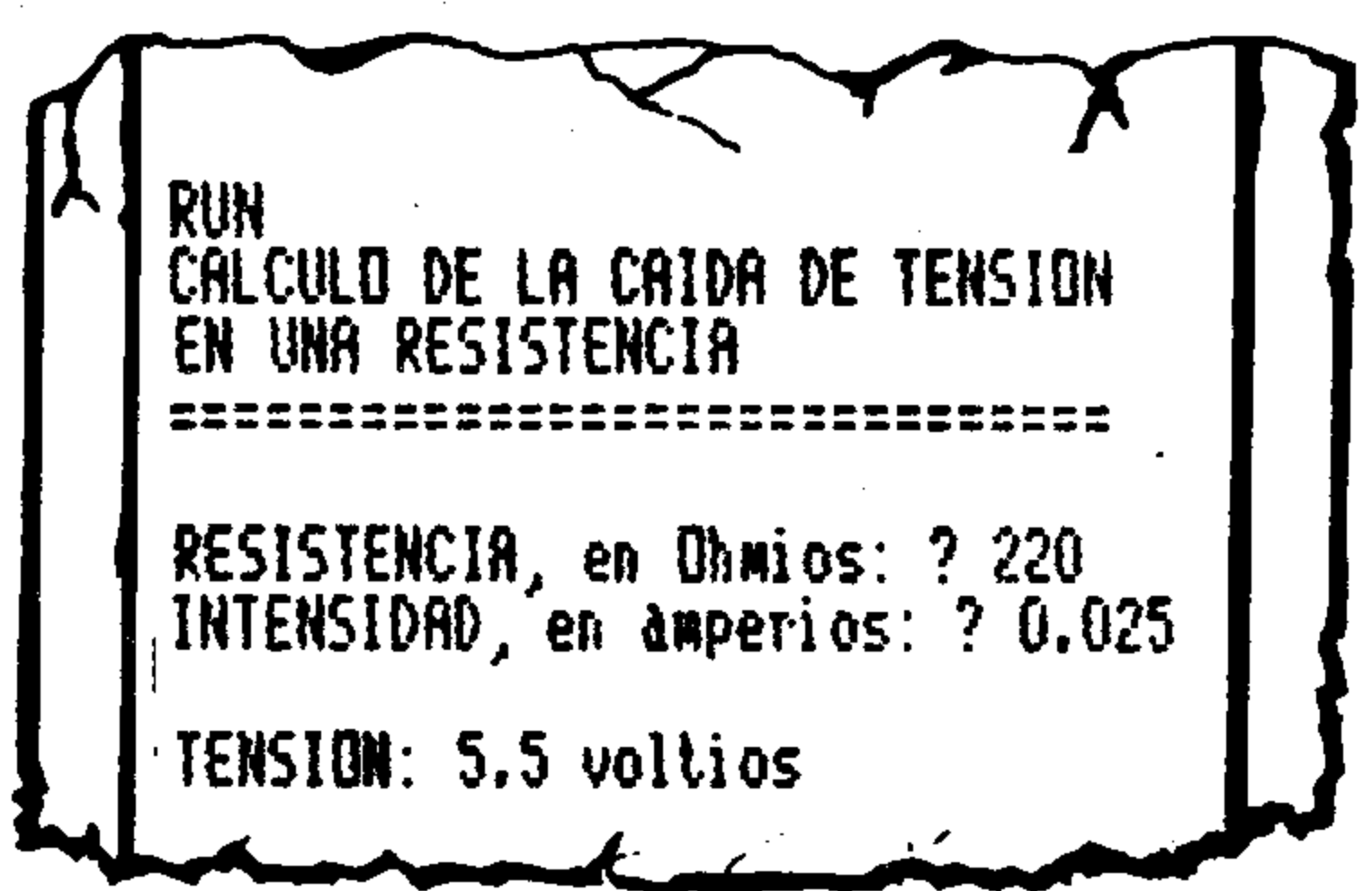
Veamos un ejemplo basado en el programa para el cálculo de la caída de tensión en una resistencia según la «Ley de Ohm»:

```
10 REM ** LEY DE OHM **
20 PRINT "CALCULO DE LA CAIDA DE TENSION"
30 PRINT "EN UNA RESISTENCIA"
40 PRINT "===== "
50 PRINT
60 INPUT "RESISTENCIA, en Ohmios: "; R
70 INPUT "INTENSIDAD, en amperios: "; I
80 PRINT
90 PRINT "TENSION: "; I*R; "voltios"
100 END
```

La ejecución del programa transcurrirá como se indica a continuación:

BASIC

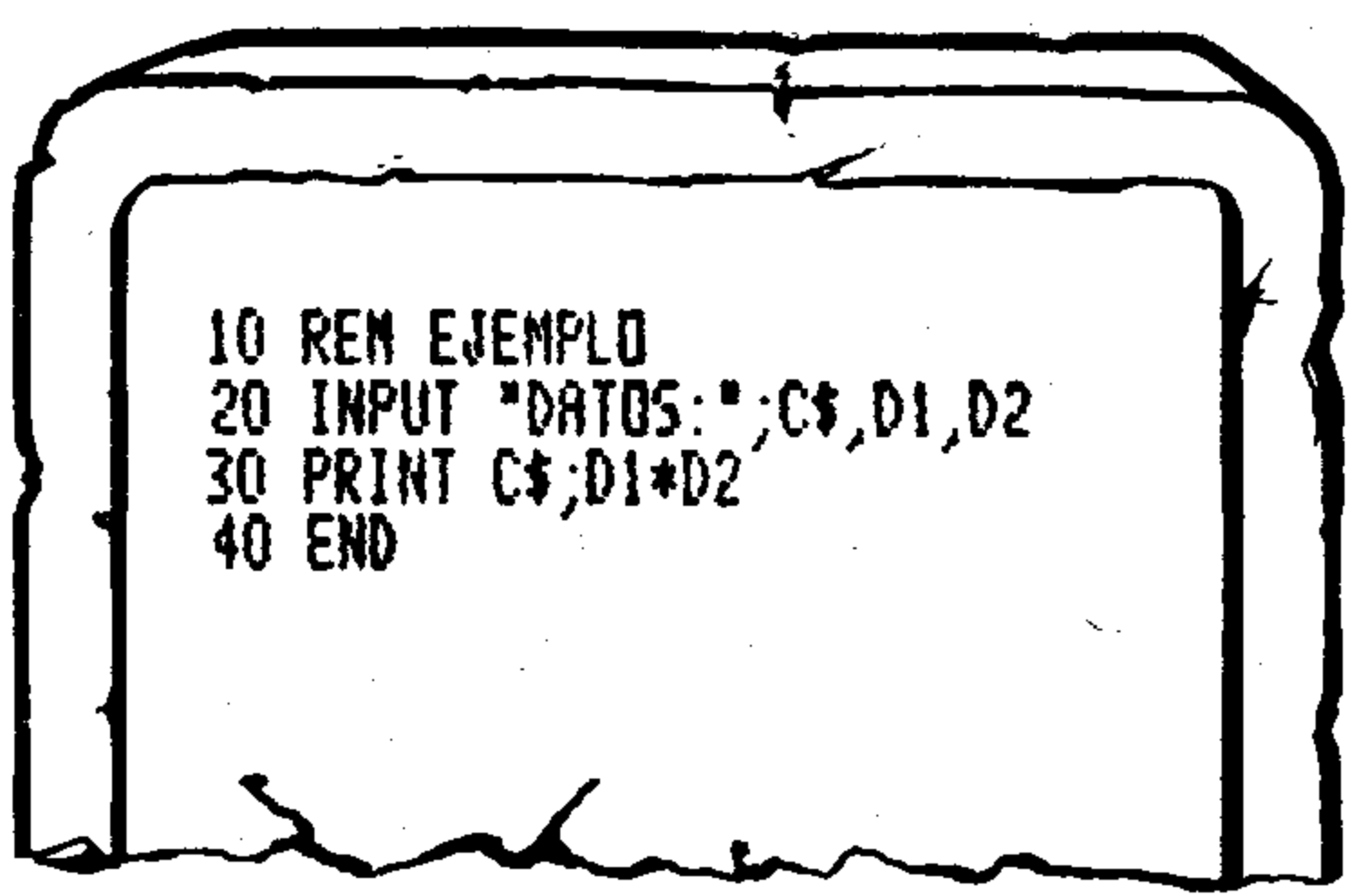
4ª PARTE



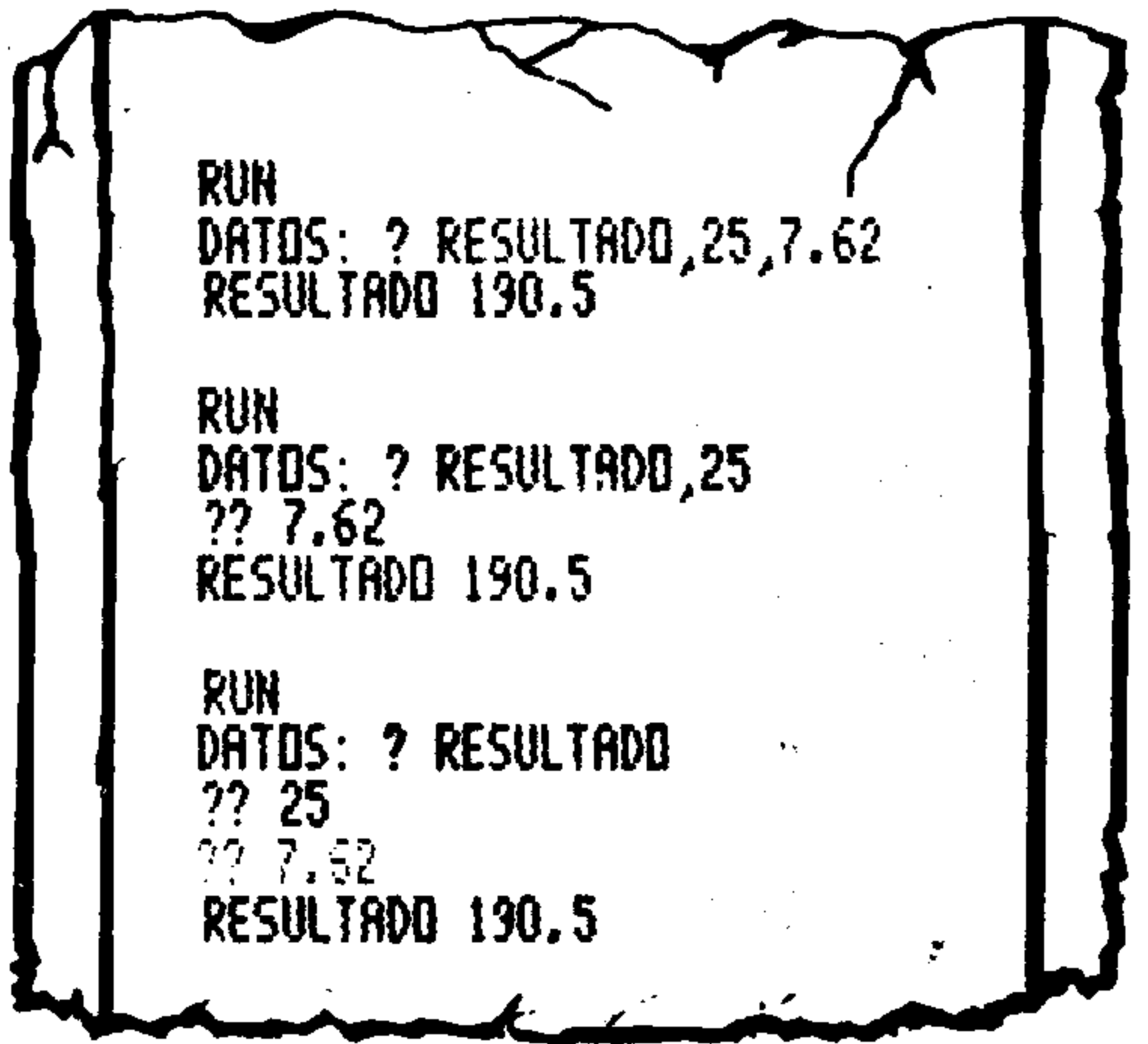
Observaciones

Cabe añadir algunas observaciones adicionales respecto a la actuación de las instrucciones INPUT. A la hora de introducir los datos asociados a la ejecución de INPUT, puede darse el caso de que se introduzcan menos o más datos de los exigidos; veamos que ocurre en ambos casos:

- Si se introducen menos datos de los exigidos por la instrucción INPUT, el ordenador exigirá los datos adicionales mostrando en pantalla el indicativo (?). ¡ahí va un ejemplo!



Para evitar cualquier posible duda vamos a ejecutar el programa repetidas veces, dándole todos los datos exigidos o «dosificándolos» en distinta medida:



Según observamos, el ordenador sigue pidiendo datos hasta que se hayan ingresado en su totalidad... Pero ¿qué

sucede cuando se introducen más datos de los requeridos?

- Si se introducen más datos de los exigidos por la instrucción INPUT, el ordenador utilizará los que precise (según el número de variables que figuren en el argumento), rechazando los datos sobrantes y mostrando en la pantalla el mensaje «EXTRA IGNORED» (ignorado los datos sobrantes).
- Por último, queda por precisar que la introducción de cadenas de caracteres se rige por las mismas condiciones de formato definidas para las instrucciones READ/DATA.

... ¡a practicar!

Al fin y al cabo, la nuestra es una revista de electrónica... por lo tanto, nada más lógico que confeccionar un programa ejemplo destinado a una aplicación electrónica.

Presentación:

Pretendemos (¡a ver si lo logramos!) escribir un programa que nos solvete el diseño de una etapa amplificadora semejante a la que aparece en la figura 1.

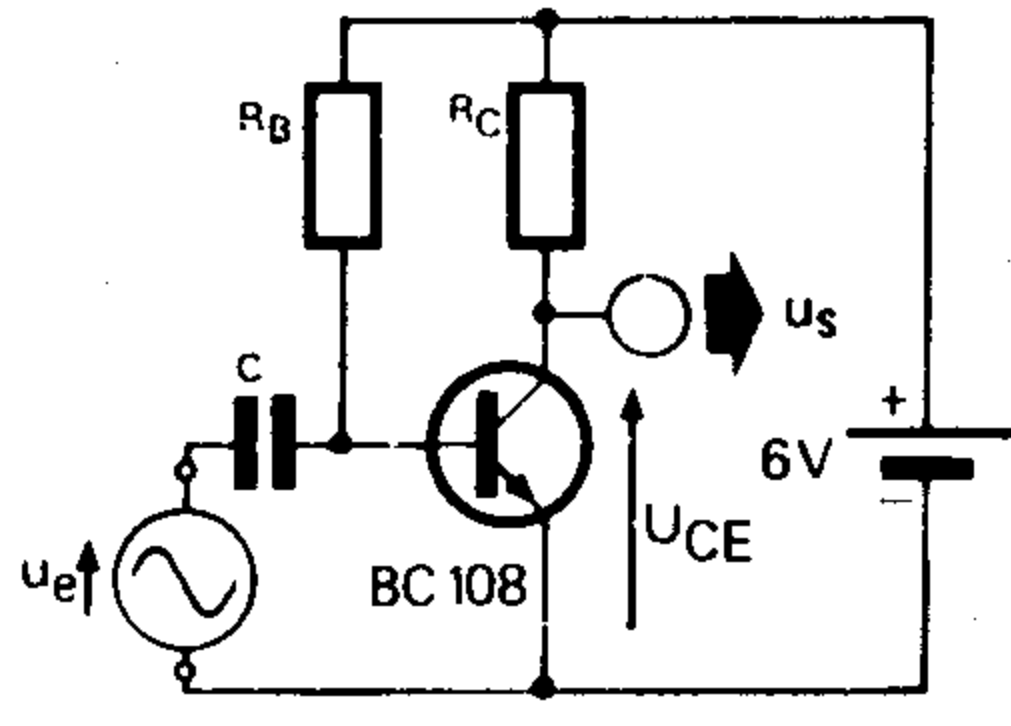


Figura 1. Esquema de la etapa amplificadora cuyo cálculo vamos a programar.

La etapa en cuestión está basada en un transistor NPN del tipo BC108. Además de las indicaciones aportadas en el esquema, partimos de los siguientes datos:

Datos e incógnitas:

- B = 125 (ganancia en corriente del transistor)
- U_A = +6 V (tensión de alimentación)
- I_{CR} = 1 mA (corriente de reposo del colector)

A partir de estos datos, queremos que el ordenador nos calcule los siguientes valores:

1. La magnitud de las resistencias de polarización R_B y R_C
2. Magnitud del condensador: C
3. Resistencia de entrada al circuito: R_E
4. Ganancia en tensión: A_V
5. Tensión de entrada máxima que admite el transistor sin que se produzca sobremodulación: U_{EM}

Como características propias del circuito, cabe citar que suponemos que la señal de salida de la etapa amplificadora (tensión alterna) puede llegar a alcanzar su valor máximo U_{SM} (U_{SM} = U_A/2).

Por lo demás, cabe indicar que esta aplicación está basada en uno de los ejercicios (problema 2a) que apare-

BASIC

1ª PARTE

cen dentro del «Curso técnico de introducción a la electrónica» publicado por Elektor.

Ordinograma:

El diagrama de flujo que ilustra la secuencia de operaciones es el que aparece en la figura 2. Su simplicidad es manifiesta, ya que la descomposición del proceso se ha realizado de una forma muy global, sin detallar una a una las operaciones parciales.

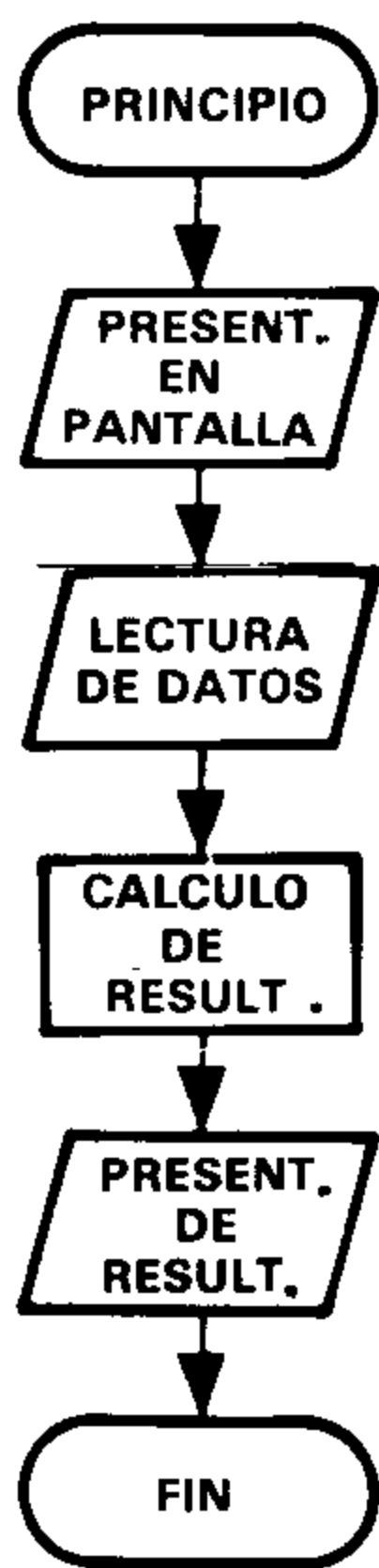


Figura 2. Ordinograma correspondiente al programa para el cálculo de la etapa amplificadora.

Programa:

El programa no presenta ninguna innovación significativa, sino que está basado únicamente en el grupo de instrucciones estudiadas hasta el momento. No cabe duda de que podríamos haberlo realizado con menos de la mitad de las instrucciones que intervienen, no obstante, de confeccionarlo de forma simplificada perdería su interés didáctico ¡Tiempo habrá de obligarles a descifrar programas «super-concentrados»!

Por ejemplo, hemos llenado el programa de comentarios que, sin lugar a dudas, ayudarán al lector a asimilar claramente las posibilidades que el lenguaje BASIC le ofrece. Por otra parte, los cálculos —sometidos a las aproximaciones habituales en electrónica— se operan de forma muy detallada y en sucesivas líneas, para que el aficionado adquiera una idea cabal de las fórmulas que entran en juego.

```
120 PRINT
130 INPUT "Ganancia en corriente (B):";B
140 INPUT "Tension de alimentacion (UA):";UA
150 INPUT "IC de reposo (ICR) en mA:";ICR
160 PRINT
170 REM CALCULO DE LA ETAPA AMPLIFICADORA
180 LET RC = .5*UA/(ICR*.001)
190 LET URB = UA-.7
200 LET IB = ICR*.001/8
210 LET RB = URB/IB
220 LET RE = B/(40*ICR*.001)
230 LET C = 1/(100*RE)
240 LET AV = 40*(UA/2)
250 LET UEM = (UA/2)/AV
260 PRINT "RESULTADO DE LOS CALCULOS"
270 PRINT "-----"
280 PRINT
290 PRINT "Valores de los componentes"
300 PRINT "RB:";RB*.001;"Kilo-ohmios"
310 PRINT "RC:";RC*.001;"Kilo-ohmios"
320 PRINT "C:";C*1000000.0;"micro-faradios"
330 PRINT
340 PRINT "Resistencia de entrada"
350 PRINT "RE:";RE*.001;"Kilo-ohmios"
360 PRINT
370 PRINT "Ganancia en tension"
380 PRINT "AV:";AV
390 PRINT
400 PRINT "Maxima señal de entrada"
410 PRINT "UEM:";UEM*1000.0;"mili-voltios"
420 PRINT
430 PRINT "Para repetir el calculo con cualquier"
440 PRINT "otro grupo de datos, no tiene más que"
450 PRINT "dar de nuevo la orden de ejecucion."
460 END
```

Y aquí está el resultado de la ejecución:

```
RUN

=====
PROGRAMA PARA EL CALCULO
DE UNA ETAPA AMPLIFICADORA
=====

INTRODUZCA LOS DATOS A MEDIDA
QUE YO -EL ORDENADOR- SE LOS
SOLICITE ... ¡ADELANTE!

Ganancia en corriente (B):? 125
Tension de alimentacion (UA):? 6
IC de reposo (ICR) en mA:? 1

RESULTADO DE LOS CALCULOS
-----

Valores de los componentes
RB: 662.5 Kilo-ohmios
RC: 3 Kilo-ohmios
C: 3.2 micro-faradios

Resistencia de entrada
RE: 3.125 Kilo-ohmios

Ganancia en tension
AV: 120
```

```
10 REM ***** CURSO DE BASIC *****
20 REM ***** ELEKTOR *****
30 PRINT
40 PRINT "=====
50 PRINT "PROGRAMA PARA EL CALCULO"
60 PRINT "DE UNA ETAPA AMPLIFICADORA"
70 PRINT "=====
80 PRINT
90 PRINT "INTRODUZCA LOS DATOS A MEDIDA"
100 PRINT "QUE YO -EL ORDENADOR- SE LOS"
110 PRINT "SOLICITE ... ¡ADELANTE!"
```

BASIC

4ª PARTE

Maxima señal de entrada
UEN: 25 mili-voltios

Para repetir el calculo con cualquier otro grupo de datos, no tiene más que dar de nuevo la orden de ejecucion.

Los comandos FOR...TO...STEP

FOR-TO-STEP	Tipo: comando de programación Función: desde - hasta - paso
--------------------	--

Definición: Los tres comandos se utilizan en conjunto para construir instrucciones destinadas a crear bucles o lazos dentro del programa. El lazo del programa está delimitado por la instrucción que establece las condiciones de ejecución (FOR-TO-STEP) y por la instrucción final o de cierre de lazo (NEXT).

Estructura:

FOR	var. = expres.	TO	expres.	STEP	expres.
------------	----------------	-----------	---------	-------------	---------

La estructura general de este tipo de instrucciones es la que se indica, no obstante, la zona STEP es opcional y puede no intervenir en la instrucción. Los dos formatos en los que pueden formularse este tipo de instrucciones son los que se comentan a continuación.

a.

FOR variable = expresión	TO expresión
---------------------------------	---------------------

En la primera zona de la instrucción (FOR variable = expresión), se define el valor inicial de la variable que se utiliza para controlar la ejecución repetitiva del bucle. La expresión cuyo valor se asigna a la variable puede ser un número o una expresión matemática en la que aparezcan valores numéricos y/o variables relacionadas por operadores. Por su parte, la expresión que sigue al comando TO (hasta), define el valor final o valor límite de la variable cuyo valor inicial se ha establecido con el comando FOR (desde).

Así, por ejemplo, la instrucción FOR C = 1 TO 90 equivale a la lectura literal: «desde C igual a 1 hasta C igual a 90». Como quiera que no aparece la zona STEP (paso), que especifica el valor del salto de incrementación o decrementación, la variable (C en el ejemplo) se incrementará en una unidad cada vez que se ejecute la instrucción, hasta llegar al valor final (90 en el ejemplo).

Ejemplos: FOR G = 2 TO 420
FOR B2 = 36.5/N*B TO 4678.7
FOR H = A/2 + 5 TO N*80.5

b.

FOR var. = expres.	TO expres.	STEP expres.
---------------------------	-------------------	---------------------

La estructura de este segundo tipo de instrucción es análoga a la anterior, con la salvedad de que incorpora el

campo STEP. La expresión que acompaña al comando STEP define el valor del paso o salto que, sucesivamente, incrementará o decrementará el valor de la variable, según sea positiva o negativa la expresión que acompaña a STEP.

Así, por ejemplo, FOR I = 4 TO 10 STEP 2 (desde I = 4 hasta I = 10 en pasos de 2) constituye la instrucción inicial de un bucle que se recorrerá 4 veces. En la primera, el valor de la variable I será 4 (el inicial), en el segundo recorrido su valor será 4 + 2 = 6, en el tercero 6 + 2 = 8 y en el cuarto alcanzará el valor final (TO 10): 8 + 2 = 10. El ordenador regresará a la instrucción inicial del bucle cada vez que se encuentre con la instrucción final o de cierre NEXT I.

Ejemplos: FOR I = 2.6 TO 52 STEP 2
FOR S = B/2 TO 0 STEP -3*Y
FOR V3 = N TO 567.8/K STEP C*A

Observamos que en los argumentos de los tres comandos que intervienen en la instrucción, pueden figurar valores numéricos o expresiones constituidas por constantes y variables relacionadas por medio de operadores aritméticos.

El comando NEXT

NEXT	Tipo: comando de programación Función: próximo (vuelta a FOR-TO)
-------------	---

Definición: El comando NEXT, acompañado de la correspondiente variable, constituye la instrucción final de un bucle iniciado con una instrucción del tipo FOR-TO.

Estructura:

NEXT	variable/s
-------------	------------

Ejemplos: 40 NEXT
65 NEXT N
90 NEXT X, Y

El formato más común de este tipo de instrucciones es el de NEXT N (con una sola variable en el argumento) que constituye la instrucción de cierre de un bucle FOR/NEXT creado en torno a la variable N. La mayor parte de los dialectos BASIC admiten otros dos formatos. El primero, constituido únicamente por el comando NEXT sin argumento. Al no aparecer variable alguna, esta instrucción constituye el cierre del bucle cuya instrucción de apertura aparece más próxima. La segunda variante normalmente admitida es aquella en la que aparece más de una variable en el argumento de NEXT (por ejemplo: NEXT H, L). En este caso, la instrucción se utiliza para cerrar varios bucles, justamente aquéllos cuyas variables de referencia figuran en el argumento de NEXT.

Los bucles FOR/NEXT

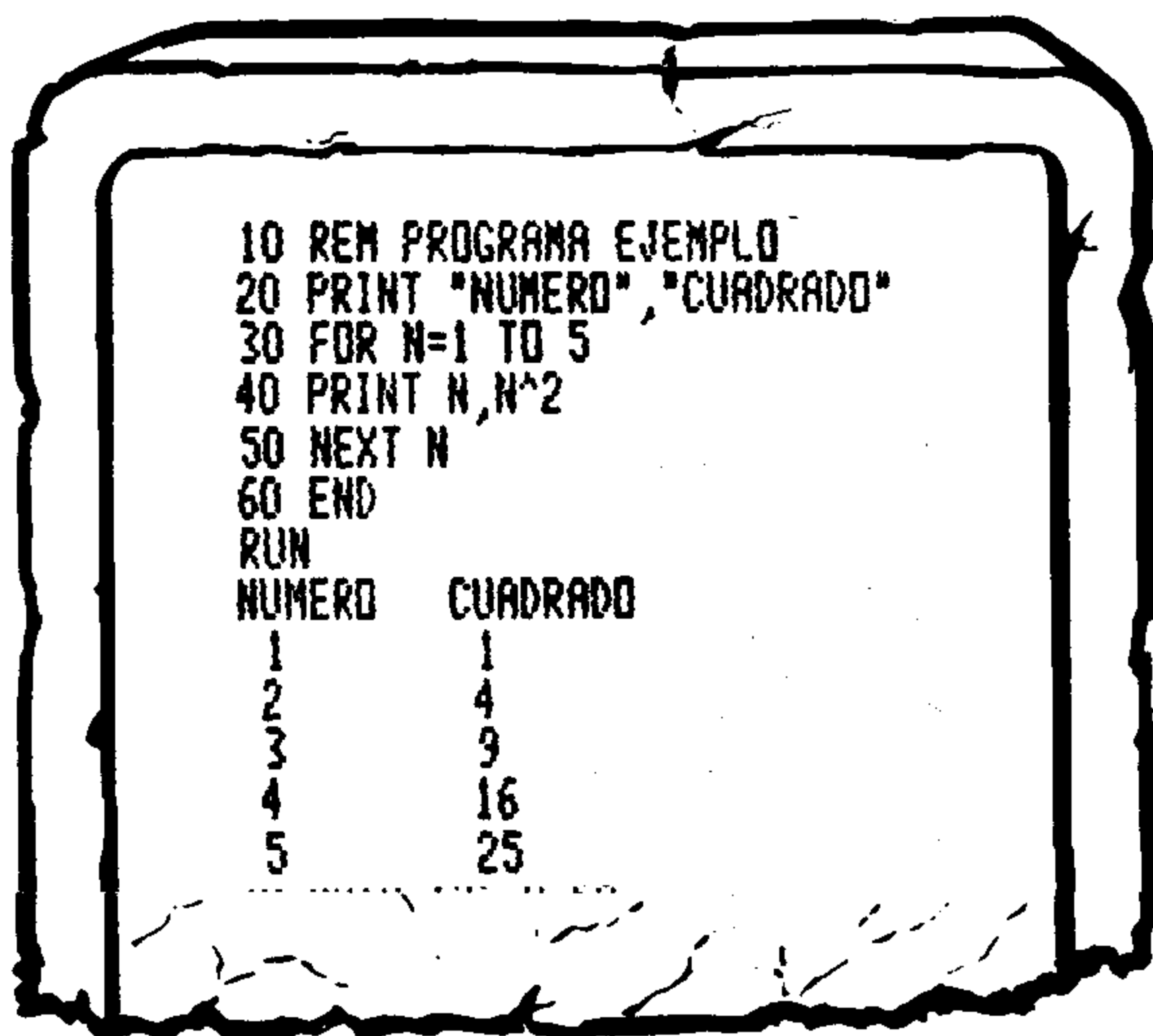
Aunque hemos analizado por separado la estructura de las instrucciones que conforman esta categoría de bucles o lazos de programa, es muy posible que su actuación conjunta no sea aún evidente para el lector. ¡En

BASIC

4ª PARTE

cualquier caso, para eso está nuestro «Curso de BASIC»... para aclarar lo que sea necesario.

Empecemos con un ejemplo muy simple:



Tal como observamos, el programa en cuestión es capaz de calcular el cuadrado de los números 1 a 5. La zona esencial la constituye el bucle comprendido entre las instrucciones 30 (principio del bucle) y 50 (fin del bucle).

30 FOR N=1 TO 5: instrucción inicial del bucle.

Con la misma iniciamos las condiciones de recorrido del bucle, dando a la variable N su valor inicial (desde N=1) y su valor final (hasta N=5).

Dentro del lazo figura una única instrucción:

40 PRINT N, N^2: presentación en pantalla del número (base a elevar, coincidente con el valor de N) y de su cuadrado.

Y, a continuación, nos encontramos con la instrucción de cierre del bucle:

50 NEXT N

Al introducir el comando RUN, el ordenador operará como sigue:

1. Ignorará la instrucción REM, dado que es una simple observación que no debe ejecutarse, y pasará de inmediato a ejecutar la instrucción 20 presentando en la pantalla el texto siguiente:

NUMERO CUADRADO

2. A continuación encontrará la instrucción 30 que ejecutará dando a N el valor 1, «tomando nota» del valor final que esta variable puede adoptar (N=5) y observando que la evolución de N va a ser en pasos de una unidad (ausencia de comando STEP).

3. Como quiera que N posee un valor situado dentro del margen establecido, el ordenador seguirá dentro del bucle, ejecutando la instrucción 40, con lo que aparecerán en pantalla los valores:

1 1

4. Al ejecutar la instrucción de cierre de bucle (50 NEXT N), se incrementará el valor de N y se evaluará si N excede del margen establecido. Al no ser así (N=2),

se pasará a ejecutar de nuevo la instrucción 40 apareciendo ahora en pantalla:

2 4

5. Una vez más se llega a 50 NEXT N. Al ejecutarla, el valor de N se hará igual a 3 y se volverá a ejecutar la instrucción 40, en este caso con N=3.

3 9

6. En el próximo paso por 50 NEXT N, el valor de N se convertirá en 4, de ahí que... ¡vuelta a la instrucción 40!

4 16

7. ¡Quinto paso por la instrucción NEXT!... N será ahora igual a 5. Desde luego el valor 5 es el límite (FOR N=1 TO 5) pero sigue dentro del margen, luego hay que ejecutar de nuevo la instrucción 40:

5 25

8. La ejecución de la instrucción 50 eleva el valor de N a 6, con lo que al compararlo con el valor límite (5) el ordenador interpreta (¡sabidamente!) que el bucle ha concluido. El programa seguirá ejecutándose ahora a partir de la próxima instrucción: 60 END.

¡Y esto es todo!

No nos cabe la menor duda de que ahora les resultará totalmente familiar la actuación de los bucles de programa FOR/NEXT; sin embargo, cabe aún por mencionar algunas particularidades y detalles a tener en cuenta.

Observaciones acerca de los bucles FOR/NEXT

- Cuando no se especifica la opción STEP el salto que afecta a la variable que establece el bucle es de +1.
- Las instrucciones internas del bucle se ejecutarán siempre que se de una de las siguientes condiciones:

- La expresión de STEP es positiva y el valor de la variable es menor o igual que el valor final definido.
- La expresión de STEP es negativa y el valor de la variable es mayor o igual que el valor final. De no cumplirse una de ambas condiciones, la ejecución seguirá a partir de la instrucción siguiente a la de fin de bucle (NEXT).

- Las expresiones que acompañan a los comandos FOR, TO y STEP se operan sólo una vez, concretamente, antes de ejecutar por primera vez las instrucciones interiores del bucle.

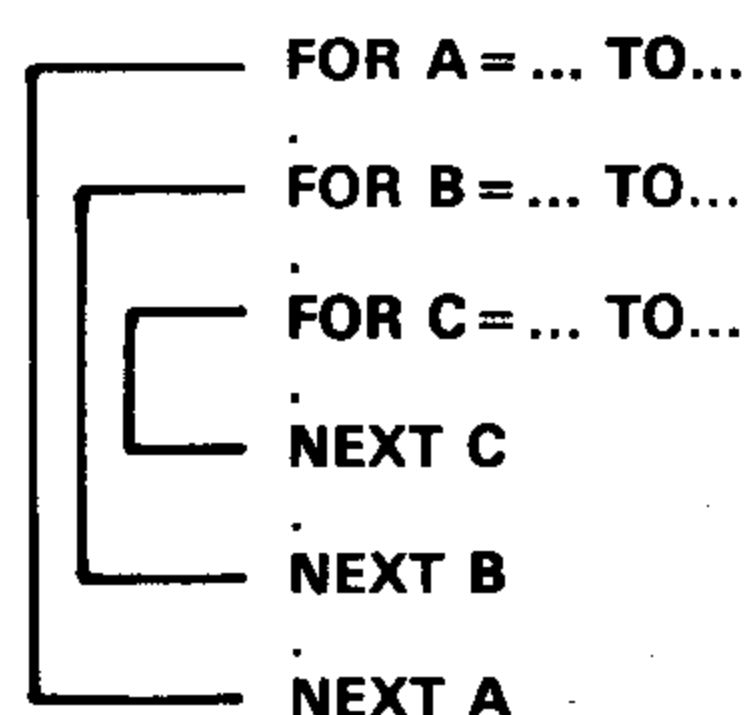
Bucles FOR/NEXT anidados

Dentro de un mismo programa pueden intervenir diversos bucles FOR/NEXT, siempre y cuando su disposición se ajuste a la estructura que se indica en la figura 3.

BASIC

4ª PARTE

BUCLAS CORRECTOS



BUCLAS INCORRECTOS

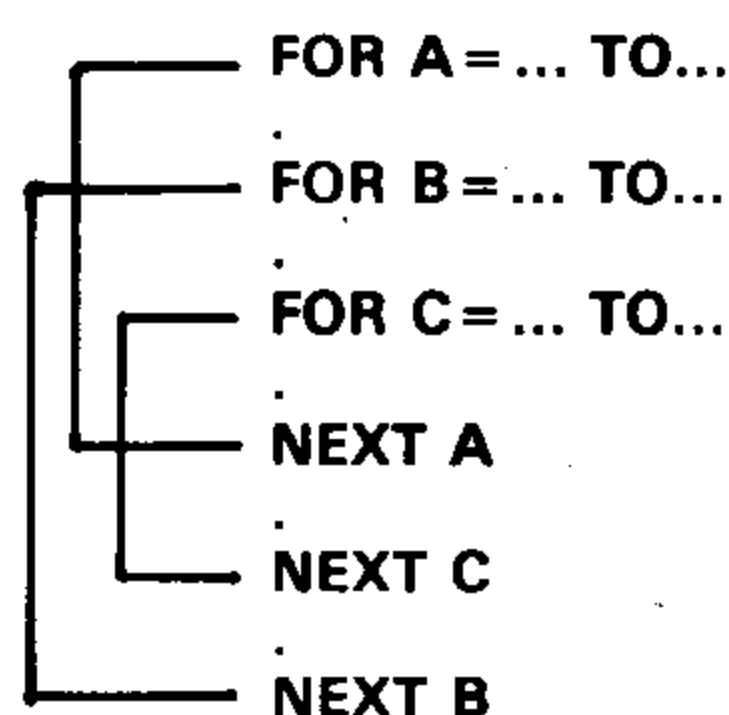


Figura 3. Anidado correcto e incorrecto de bucles FOR/NEXT.

Los bucles pueden anidarse internamente con la condición de que cada uno de los bucles se cierren ordenadamente. Desde luego, si los bucles son independientes no existe inconveniente alguno para introducir el número de éstos que sea preciso.

En cualquier caso, el número de bucles que intervienen en un programa está limitado exclusivamente por la magnitud del espacio de memoria disponible.

Respuestas al CUESTIONARIO de la 3.ª parte

1. Las instrucciones REM se utilizan para insertar comentarios y observaciones dentro de un programa BASIC. Este tipo de instrucciones no son ejecutadas por el ordenador.

2. Veamos lo que ocurrirá en el caso de que la instrucción DATA asociada a 20 READ K1, K2, K3, K4 sea una de las siguientes:

a) 50 DATA 24, 56, BASIC, 578

El ordenador dará indicación de error, ya que se pretende asignar una cadena de caracteres (BASIC) a la variable numérica K3.

b) 60 DATA 78.5, 46, 35, 67, 90

La lectura de datos se realizará correctamente; el último dato (90) será ignorado ya que READ no incorpora una quinta variable.

c) 70 DATA 35, 80, 36

No existen suficientes valores para asignar a las variables que aparecen en la instrucción READ; en consecuencia, el ordenador dará un mensaje de error.

d) 80 DATA 60, 45.3, 67, 100, RESULTADO

Se leerán correctamente los cuatro primeros datos numéricos requeridos por la instrucción READ; el quinto dato (cadena de caracteres) será ignorado ya que sólo hay cuatro variables en la instrucción READ.

3. El resultado de la ejecución del programa BASIC propuesto será el siguiente:

CALCULO MATEMATICO: 576.283398 72.0354247
2881.41669 A=24.5 240.456786

4. Utilizando el comando LET, las expresiones propuestas se convierten en las siguientes instrucciones:

a) $LET X = (A + B) * C / (20 * 4 * D)$ hay que precisar que al proponerles la expresión a convertir, nuestro duende particular sustrajo los paréntesis que encerraban a $(A + B)$.

b) $LET X = A / C + C * D / 6 / E * A * 46$

¡aunque parezca curioso, se resuelve sin paréntesis!

c) $LET X = A^B * C / (D / 5 - A)$

d) $LET X = 50 * A - 4.52 * 25 / C^7$

e) $LET X = B * (C / D) * (X - 1)$

f) $LET X = A^{(2/C)} * B^{(1/A)} * C^{(-20)}$

5. Es evidente que el programa admite muy diversas variantes que, en definitiva, dependen de la capacidad e incluso de la fantasía del programador... ¡en este caso, de Ud. mismo!

A continuación les presentamos una de las posibles soluciones programadas:

```

10 REM ***** EJERCICIO 5 *****
20 REM * CURSO DE BASIC - 2ª *
30 PRINT
40 READ C, R, T
50 PRINT "CON", C, "PESETAS POCO ES"
60 PRINT "EL BENEFICIO QUE VA A"
70 PRINT "OBTENER... ¡OBSERVE!"
80 PRINT
90 PRINT "ESTOS SON LOS DATOS"
100 PRINT "CAPITAL:", C, "PESETAS"
110 PRINT "REDITO:", R, "%"
120 PRINT "TIEMPO:", T, "AÑOS"
130 PRINT
140 PRINT "INTERES:", C * R * T / 100, "PESETAS"
150 PRINT
170 PRINT "¡ENHORABUENA!"
180 DATA 250000, 14.28, 1.5
190 END

RUN

CON 250000 PESETAS POCO ES
EL BENEFICIO QUE VA A
OBTENER... ¡OBSERVE!

ESTOS SON LOS DATOS
CAPITAL: 250000 PESETAS
REDITO: 14.28 %
TIEMPO: 1.5 AÑOS

INTERES: 53550 PESETAS

¡ENHORABUENA!

```

CUESTIONES (4.ª Parte)

- ¿En una instrucción FOR...TO...STEP, el valor final de la variable puede ser inferior al valor inicial?
- ¿Qué ocurre cuando en respuesta a una instrucción INPUT se introducen más o menos datos de los exigidos?
- ¿Pueden cerrarse varios bucles de tipo FOR/NEXT con una misma instrucción NEXT?
- ¿Cuál es el número máximo de bucles FOR/NEXT que pueden intervenir dentro de un programa?
- Escribir un programa BASIC que se reúna las siguientes condiciones:

- Pedirá al usuario que introduzca un número (N)
- Al ejecutarlo, debe construir una tabla en la que figuren los siguientes valores:
 - Primera columna: números enteros comprendidos entre N y N + 19.
 - Segunda columna: valor del cuadrado (N²) de cada uno de los números.
 - Tercera columna: valor del cubo (N³) de cada uno de los números.
 - Cuarta columna: diferencia entre los valores del cubo y del cuadrado (N³ - N²).

teclado sensorial

¡un teclado sin teclas!

Las posibilidades de elección y la gama de precios de los teclados de ordenador disponibles en el mercado, están llegando a un nivel de proporciones ilimitadas. En consecuencia, muchos de nuestros lectores tal vez prefieran construirse el suyo propio. Un sistema que utilice teclas mecánicas, aunque sea sencillo, es relativamente caro. Los teclados activados por contactos capacitivos constituyen una alternativa económica. Consiguen un alto nivel de fiabilidad, sin la necesidad de emplear equipos mecánicos tradicionales de alto coste.



El caso más frecuente de utilización de este tipo de teclado con 10 a 20 teclas es el de un pequeño sistema basado en microprocesador, por ejemplo, en un microordenador mono-tarjeta. La construcción de un teclado realizado con teclas mecánicas puede parecer muy simple, aunque es sorprendentemente caro. El uso normal somete a las teclas a un desgaste mecánico considerable, haciendo necesaria una sustitución frecuente. El empleo de goma conductora y elementos de efecto Hall es una forma de resolver el problema. Si bien, una solución más efectiva es un teclado capacitivo en el que se utilicen teclas accionables por el tacto, en sustitución de los dispositivos mecánicos. Las teclas son, en realidad, pequeñas superficies cuya capacidad cambia al tocarlas. Contrariamente a la primera impresión que se pueda tener, este teclado puede ser minúsculo, tal como se demostrará en este artículo. Al montaje basta con añadirle un pequeño programa que se encargue de efectuar la decodificación e impida la consideración de los efectos de eventuales rebotes.

Principio de funcionamiento

En la figura 1 se ilustra perfectamente el modo de funcionamiento de un teclado capacitivo. Dos condensadores están intercalados en serie entre dos conductores. La unión entre ellos forma el contacto sensitivo. Uno de los condensadores (C_b) está conectado a la entrada de un multivibrador monostable (MVM); el otro condensador (C_a) recibe un impulso. En ausencia de contacto, el impulso es objeto de diferenciación por el conjunto C_a , C_b y R . El multivibrador monoestable reacciona al primer flanco del impulso diferenciado y luego genera, a su vez, un impulso que tiene una cierta duración. El nivel de disparo del multivibrador monostable se elige de modo que el impulso transmitido por intermedio de C_a y de C_b sea suficientemente importante para asegurar, en el límite, el disparo del multivibrador monoestable. Si se pone, ahora, un dedo sobre el contacto sensitivo, se introduce una capacidad y/o una resistencia de fuga hacia tierra. Esto disminuirá la duración del impulso que llega a la entrada del multivibrador a una magnitud inferior a su umbral de disparo de entrada. Por consiguiente, el multivibrador monoestable ya no generará un impulso de salida.

1

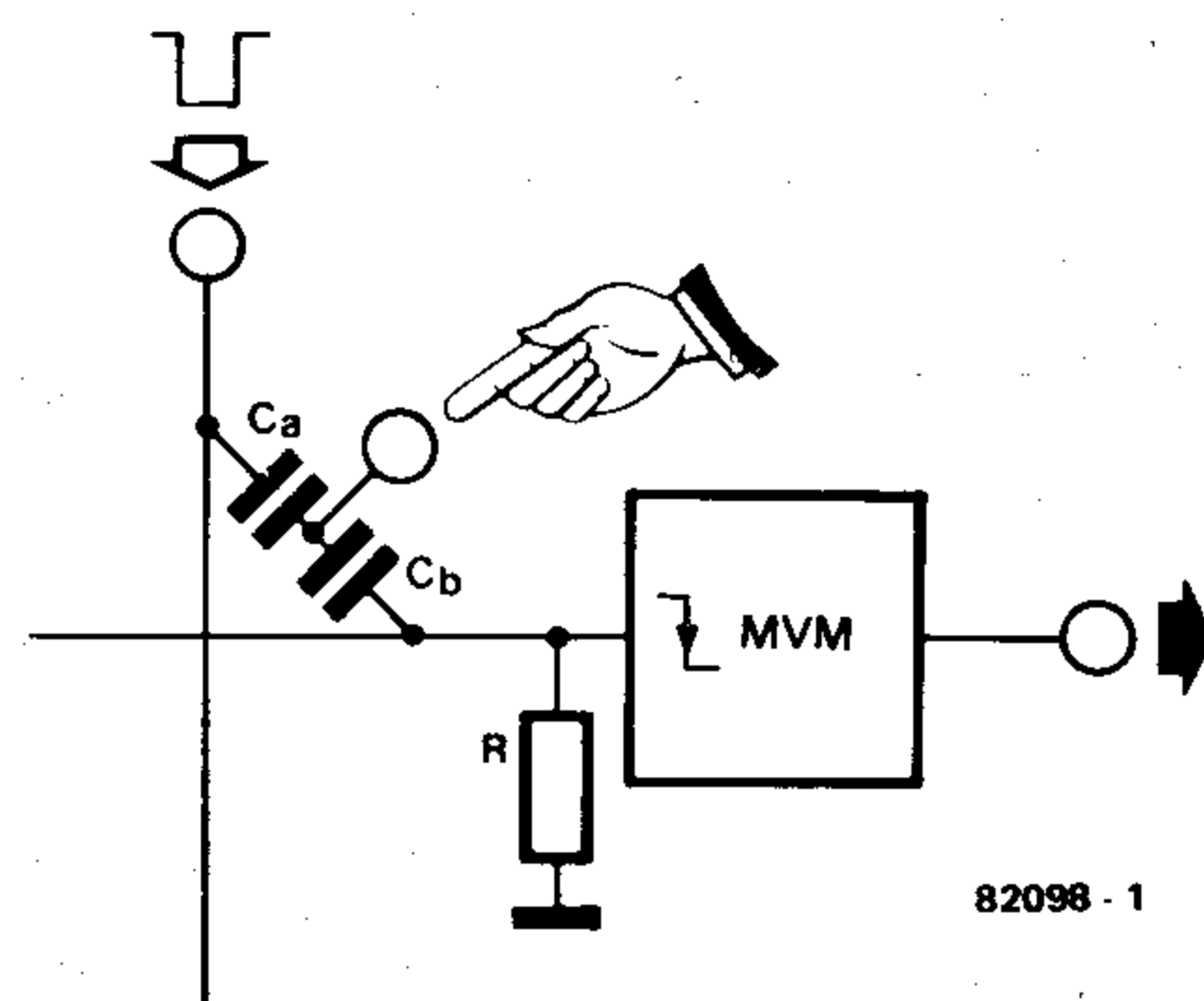


Figura 1. Principio fundamental que permite la realización de un teclado de contactos capacitivos.

Así se puede construir un teclado completo; basta con añadirle un pequeño programa que gestione la generación de impulsos y sea capaz de detectar su presencia o ausencia (v.g. cuando se toca uno de los contactos). El resultado es un teclado económico y duradero.

El circuito

En la figura 2 se muestra el diagrama circuital del teclado capacitivo. Hay 20 contactos sensitivos dispuestos en una matriz de cuatro filas y cinco columnas. Una «tecla» capacitiva está situada en cada nodo de unión entre las filas y columnas.

Las columnas reciben los impulsos a través de los inversores N9...N13, del tipo colector abierto. Ello explica la existencia de resistencias de activación (R1...R5) necesarias para obtener en las salidas una señal de 10 voltios. Los impulsos requeridos de control de las columnas puede generarlos el microprocesador. Evidentemente, es preciso activar las columnas en forma sucesiva. Al final de cada una de las cuatro filas hay un multivibrador monoestable constituido por dos disparadores Schmitt, un condensador y una resistencia. Los circuitos CMOS utilizados tienen una impedancia de entrada notablemente alta, lo que permite hacer abstracción de la misma si se tienen en cuenta el valor óhmico de las resistencias existentes a la entrada. La pequeña tensión de alimentación (3,3 V) produce una histéresis muy reducida, prácticamente despreciable (inferior a 400 mV). Este es un requisito esencial puesto que las «teclas» capacitivas transmiten señales muy bajas. C2 y R14 determinan la duración del impulso de N1 y de N2, C4 y R18 determinan la de N3 y N4 y así sucesivamente. La red R-C conectada a la entrada de cada uno de los multivibradores (C1-R12, C3/R16, etc.) constituye un filtro de paso-alto que está concebido para disminuir la sensibilidad del sistema al zumbido y a otros «parásitos eléctricos». Un transistor (de colector abierto) está conectado a la salida de cada multivibrador actuando como buffer/inversor.

Los potenciómetros P1...P4, en las entradas de las puertas, permiten el establecimiento de un nivel de tensión continua. Se ajusta esta tensión a un valor tal que el impulso transmitido (por una columna) quede apenas por debajo del umbral de disparo inferior (U_{T-}), de manera que el multivibrador monoestable proporcione un impulso cuando no esté activado el contacto de la columna y de la fila correspondiente.

Puesto que N9...N13 invierten los impulsos de entrada positivos, los multivibradores monoestables se dispararán por el primer flanco de bajada de un impulso (en CL1...CL5). Para aclarar estas materias, en la figura 3 se muestran varias formas de onda de las señales. En nuestro ejemplo, no hay acción sobre el contacto en el curso de los dos primeros impulsos, teniendo lugar el contacto en el transcurso de los tres últimos. En la figura 3 también se indica claramente el nivel de tensión de corriente continua que debe utilizarse.

Ajuste

Ante todo, los cursores de P1...P4 se llevan al extremo positivo del potenciómetro (+3,3

2

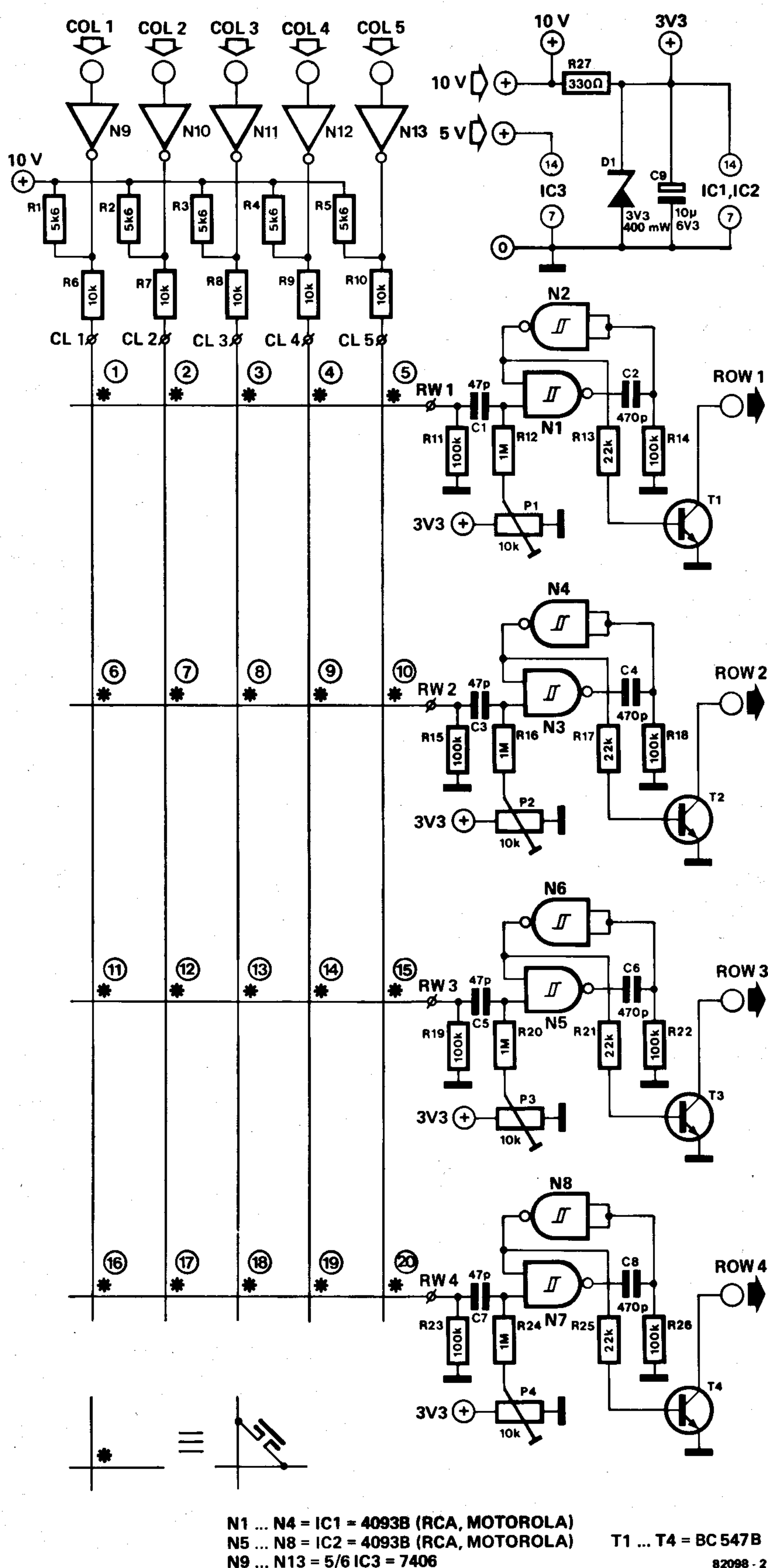


Figura 2. Esquema de principio de un teclado con 20 contactos sensitivos. Se requieren muy pocos componentes, pues el microprocesador se encarga del control y de la detección del teclado.

V). A continuación, una tensión de onda cuadrada se aplica a CL1. P1 se ajusta lentamente hasta que se encuentre la posición en la cual el monoestable correspondiente se dispara a la llegada de un flanco de subida de la señal rectangular (es decir, al flanco de bajada a la entrada de N1). Si, ahora, se toca una tecla perteneciente a esa fila, no debe dispararse el monoestable. Se prosigue ligeramente la rotación del potenciómetro, en el mismo sentido, de forma que el montaje sólo reaccione en caso de acción firme sobre una «tecla». El mismo procedimiento se lleva a cabo, sucesivamente, para cada columna.

Control y detección

En la figura 4 se muestra cómo se generan los impulsos para las columnas, que suelen estar proporcionados por el sistema basado en microprocesador, al que está conectado el teclado. Después del paso del flanco de subida de cada impulso, el microprocesador puede escrutar las salidas de las filas para «ver» si una de ellas ha adquirido el nivel lógico alto («1») durante el período correspondiente a la constante de tiempo del monoestable. Si tal circunstancia se produce, ello significa que se ha pulsado la tecla situada en la intersección de la columna en donde se produjo el impulso y la fila en donde se detectó la falta de impulso.

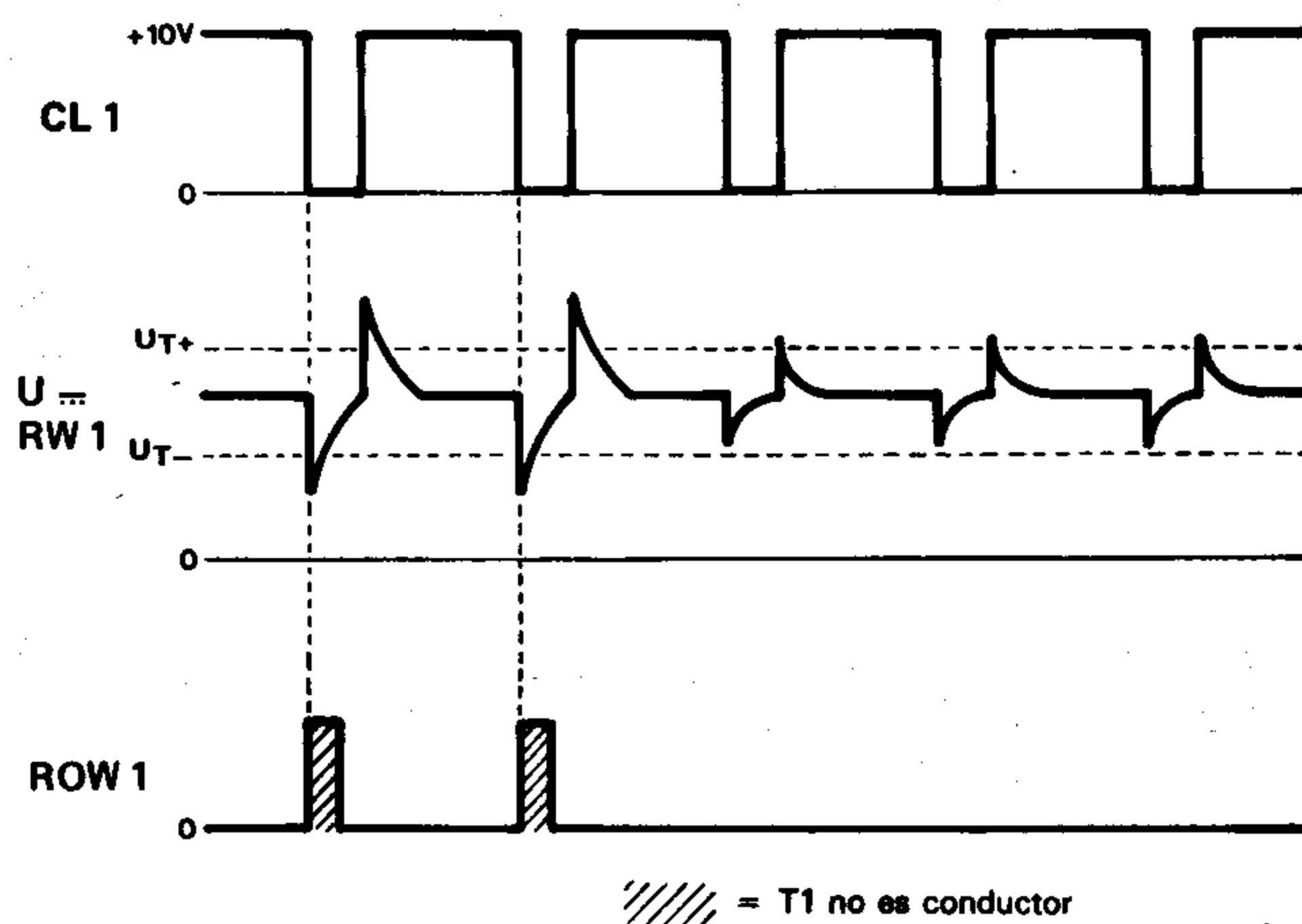
Los valores de los componentes utilizados en el esquema proporcionan una duración de impulso de unos 50 microsegundos. Este valor puede modificarse fácilmente cambiando la capacidad del condensador de 470 pF por la de otro condensador diferente.

Si se quiere obtener un buen funcionamiento del sistema será preciso añadir una función «anti-rebote». De nuevo, es el microprocesador quien puede encargarse de resolver este problema. En la figura 5 se ilustra el diagrama de flujo de un programa hipotético capaz de cumplir esta función. A falta de la acción sobre una tecla, el programa se mantendrá en el bucle B, en donde el ordenador espera 10 ms.

Durante este período, puede realizar otras tareas, tal como la excitación del display. Cuando se genera un impulso en la columna 1, se exploran las líneas y luego, se emite un impulso en la columna 2 y así sucesivamente, hasta que se hayan explorado todas las filas. Una vez que se han explorado las cinco columnas sin encontrar un nivel lógico «1» (o lo que es lo mismo, ausencia de impulso), se efectúa un retorno al comienzo del bucle B, pues evidentemente no se tocaron las teclas.

Si, por el contrario, se detecta un nivel lógico alto, se abandona el bucle B y el microprocesador espera 10 ms antes de verificar si la misma tecla sigue activada; si tal fuera el caso, el programa sigue la flecha y pasa a otro programa que se encarga de utilizar la información de acción sobre la tecla. Cuando esto ha tenido lugar, el programa vuelve a la etiqueta KEY y si, en este momento, la tecla sigue en estado de activación se recorrerá el bucle A hasta que se libere la tecla. A continuación, el bucle B se reinicializa y el procesador espera a que se accione una nueva tecla. Este procedimiento permite eliminar los rebotes durante al menos 10 ms. La práctica nos ha demostrado que fun-

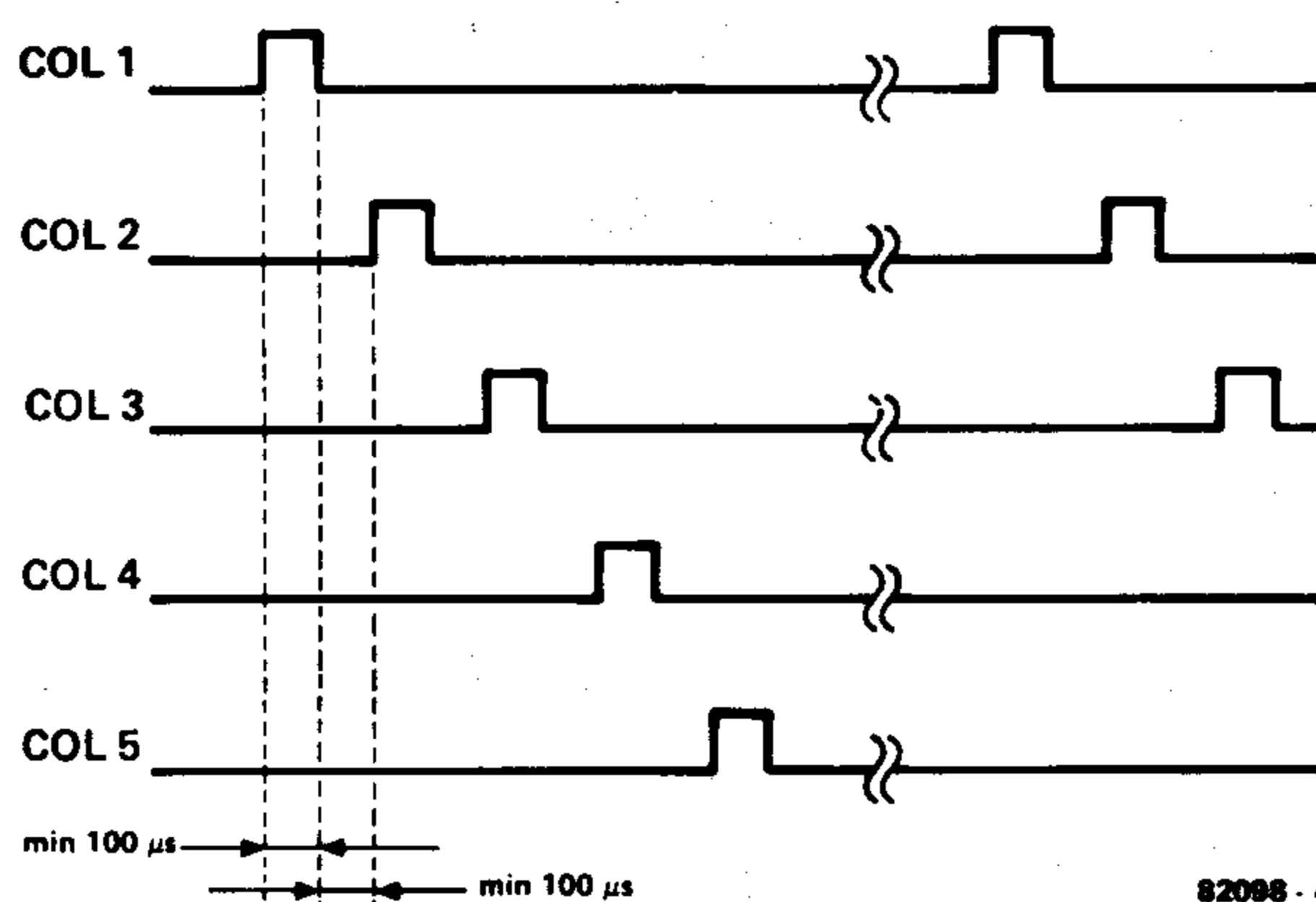
3



82098 - 3

Figura 3. Ejemplo de las formas de señal que se generan cuando se envían impulsos a la columna 1. La tecla 1 no se toca hasta después del segundo impulso.

4



82098 - 4

Figura 4. Procedimiento de control de las diversas columnas.

ciona perfectamente. Examinando más detenidamente el esquema, se constatará que el sistema tiene necesidad, en total, de 9 líneas de entrada/salida (E/S) del microprocesador, de las cuales hay cinco salidas para las columnas y cuatro entradas para las filas.

Las teclas

Un trozo de circuito impreso de doble cara es todo lo que se precisa para obtener la superficie activa del teclado. En una de las caras, que llamaremos superior, se delimitan 4 x 5 pequeñas superficies de cobre cuadradas (de 1,5 x 1,5 cm.) con un espaciado de 5 mm. entre sí (como se muestra muy claramente en la figura 6). En la cara inferior, se delimita un número de superficies idénticas, situadas muy precisamente en la vertical de las existentes en la cara superior, estando estas superficies inferiores divididas en dos por supresión del cobre a lo ancho de una pista. Con la observación del gráfico se facilitará, en gran medida, la comprensión del texto. Las mitades superiores están todas interconectadas por medio de una pista de cobre estrecha, mientras que la mitad inferior de cada cuadrado está provista de un punto de soldadura.

Debe tener gran cuidado cuando se grave y prepare la placa de circuito impreso, ya que el rendimiento satisfactorio del teclado depende completamente del hecho de que todas las teclas tengan la misma capacidad. Para poder conseguirlo, todas deben tener la misma área superficial y han de estar situadas en posiciones relativas (idénticas en ambos lados de la placa).

En la figura 6 se muestra el teclado. Una vez que se haya grabado la placa, los puntos de soldadura de las columnas quedan interconectados en la forma indicada en la figura 7. De nuevo, se exige un cuidado muy notable. Para este montaje puede utilizarse hilo de cobre delgado, procurando obtener conductores lo más rectilíneos posible.

Pasemos, ahora, a la zona superior del circuito impreso, sobre la que se encuentran las superficies sensibles. Para protegerlas y conservar su integridad el tiempo más largo posible, es recomendable recubrirlas con una película plástica autoadhesiva. Antes de proceder a esta operación, se dará un nombre a cada tecla mediante la impresión correspondiente sobre la superficie cobreada (con la ayuda de letras autoadhesivas, por ejemplo). Una vez que se haya aplicado la película plástica, las teclas formarán un teclado capacitivo.

Si la calibración final se hubiera efectuado antes de la colocación de la película plásti-

5

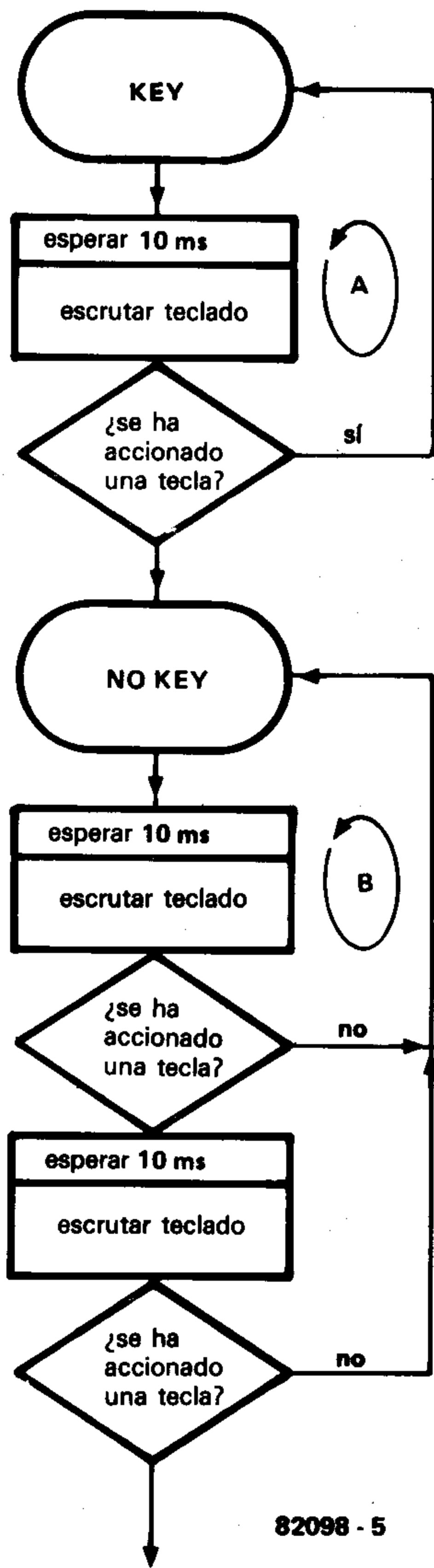


Figura 5. Diagrama de flujo correspondiente a un programa concebido para efectuar la exploración del teclado y para eliminar los rebotes del mismo.

ca, será preciso realizar una nueva puesta a punto. Si se desea proteger al teclado, en la medida que sea posible, contra una eventual influencia exterior, se puede incorporar una placa metálica, de las mismas dimensiones del circuito impreso, a unos 2 cm. por debajo del teclado. La placa debe ser paralela a la placa de circuito impreso pues, de no ser así, algunas teclas pueden resultar más sensibles que otras. La placa metálica debe conectarse a la masa de la caja. Sólo dos pequeños consejos antes de concluir. Recorte al máximo los hilos de conexión de las columnas y de las filas. Asimismo, procure hacer tan pequeño y tan simétrico como sea posible, el circuito impreso que sirve de soporte a los multivibradores. Por supuesto, los lectores tienen libertad para seleccionar cualquier número de teclas que deseen. El funcionamiento del circuito depende, en gran medida, de la precisión con la que se construya el teclado. Se puede experimentar con ajustes de potenciómetros y superficies sensitivas.

6

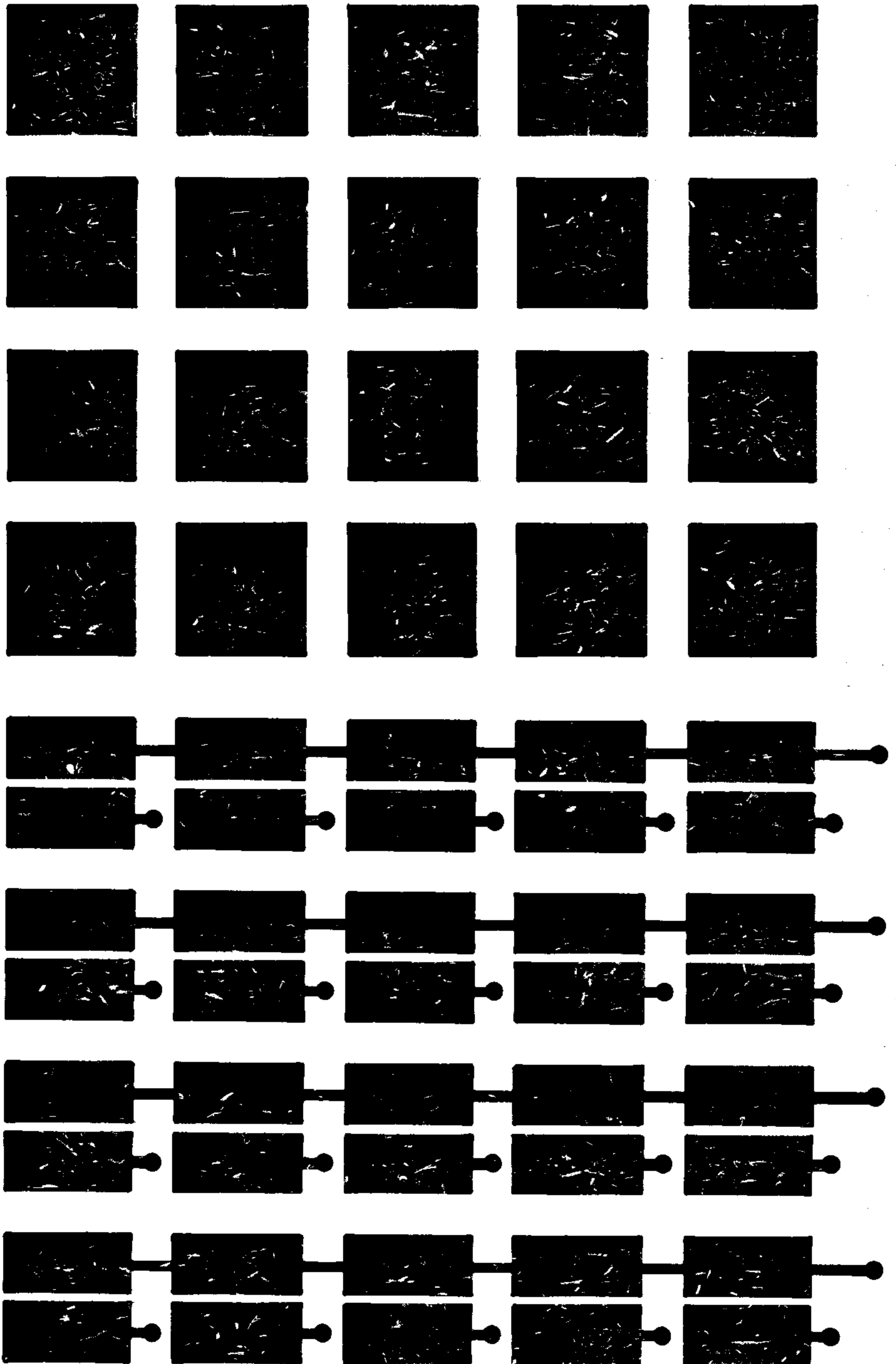


Figura 6. Realización del teclado. Se necesita una placa de circuito impreso de doble cara y mucho cuidado y precisión, puesto que es condición esencial que las capacidades de las diversas teclas sean idénticas.

7

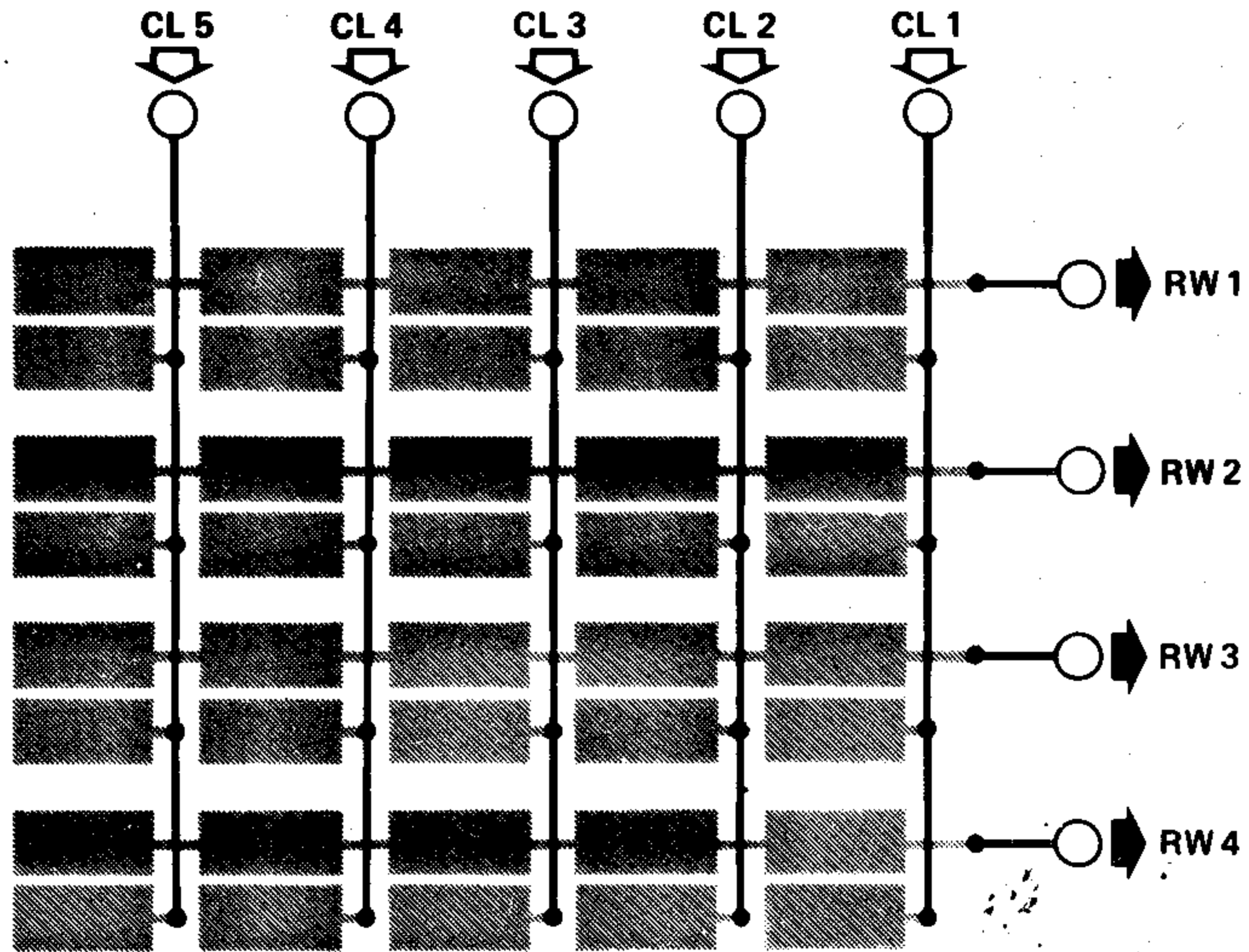


Figura 7. Todas las semisuperficies inferiores de una misma columna están conectadas entre sí con la ayuda de un hilo de cobre de pequeña sección. Las semisuperficies superiores de una misma fila están conectadas por medio de una delgada pista de cobre.

Cualquier automóvil de juguete, con tal que no sea exclusivamente minúsculo, puede alojar sin problemas una pequeña placa de circuito impreso y un par de pilas. El circuito que presentamos añade un «encanto» especial al coche de juguete: un efecto de luces intermitentes muy semejante al de los coches de policía, ambulancias... El efecto está tan perfectamente simulado que no hay necesidad de incluir ninguna pieza móvil.

¡Derechos al grano!

Tal como se ilustra en el diagrama de la figura 1, es posible diseñar cualquier clase de circuito divertido y efectista con un mínimo de componentes. El dispositivo comple-

mento diferenciador. Puesto que R3 está conectada a la toma positiva de la alimentación, la red sólo es sensible a los flancos de pendiente negativa de la onda cuadrada. Estos breves «sobrepulsos» se convierten, entonces, en impulsos utilizables por la puerta N2 para excitar al transistor Darlington T1. A su vez, este transistor enciende la lámpara conectada a su colector durante un breve período de tiempo. Se ha incluido una resistencia (R5) entre el emisor y el colector del transistor para conseguir que la lámpara se mantenga a la temperatura correcta. Esto tiene la ventaja de que la corriente inicial a través de la lámpara es mucho menor que la normal y, por consiguiente, la lámpara tendrá una duración mucho mayor. Para lograr una luz brillante, puede utilizar una lámpara de 6 V y alimentar el circuito con una tensión de 9 V.

La única diferencia entre las secciones primera y segunda del circuito es el hecho de que la segunda puede «programarse» para

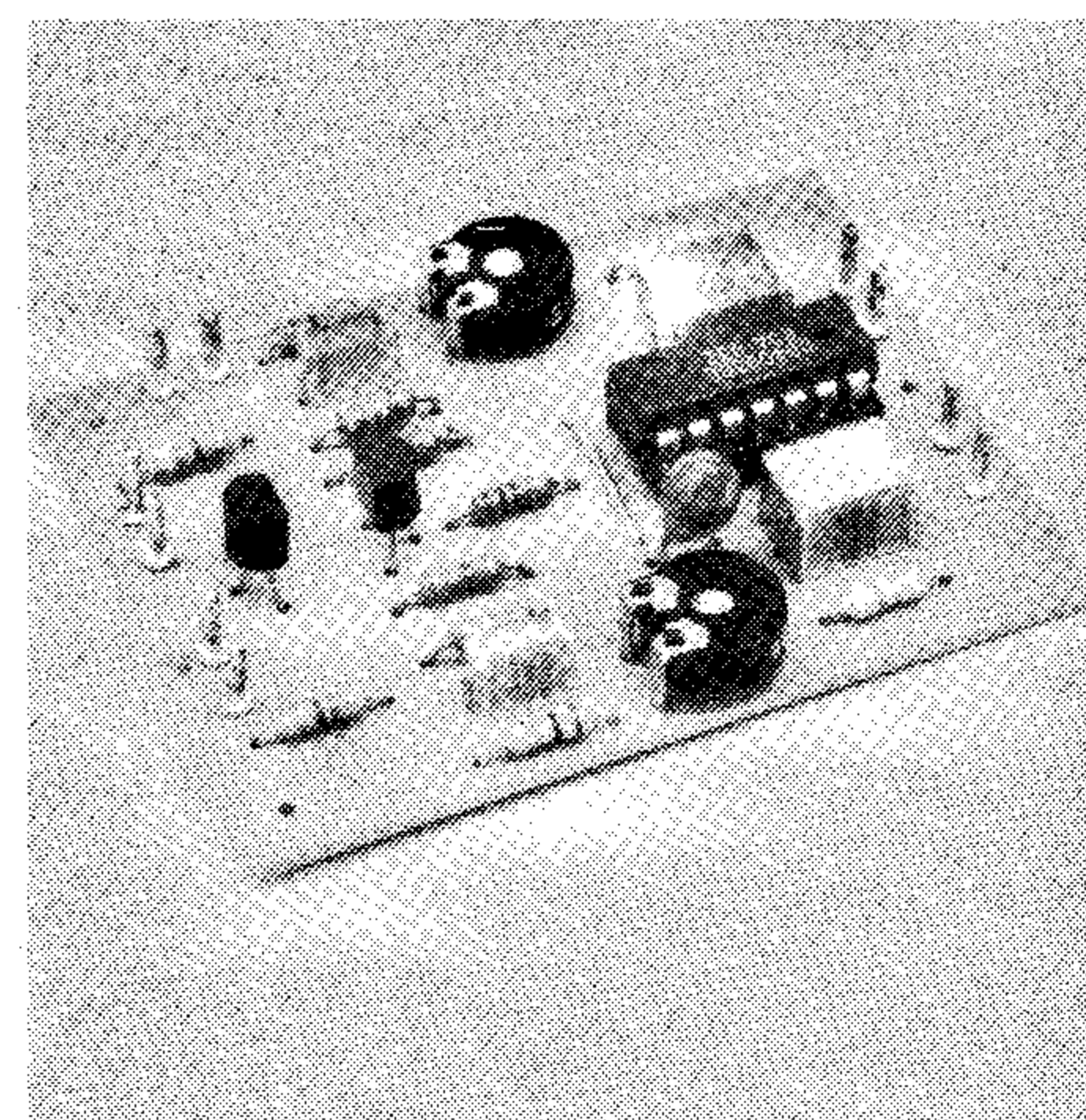
intermitente electrónico

Con este intermitente luminoso no sólo alegrará sus fiestas navideñas sino que también puede incorporarlas a un coche de juguete y obtener un efecto muy similar a las luces de aviso que se ven en las ambulancias, coches de bomberos y vehículos de policía. Con un coste mínimo, puede añadir un complemento fascinante a los vehículos de juguete.

to está constituido por dos circuitos osciladores idénticos de baja frecuencia, cada uno de los cuales controla una pequeña lámpara. El principio de funcionamiento puede describirse con bastante brevedad. Como quiera que ambos circuitos son idénticos, sólo desvelaremos uno de ellos. El oscilador (multivibrador astable) está construido sobre la base del disparador Schmitt N1. El condensador C1 está conectado entre las entradas de la puerta y masa. La salida de N1 se realimenta a la entrada a través de la resistencia R1 y del potenciómetro P1. El condensador se carga, o descarga, a través de estas resistencias, dependiendo del nivel lógico presente en la salida de N1. Siempre que la tensión a través del condensador alcance uno de los niveles del disparador, «conmutará» la salida de la puerta.

En consecuencia, el multivibrador genera una señal de salida de onda cuadrada, cuya frecuencia viene determinada por la relación entre el valor del condensador y la resistencia total de R1 y de P1. La frecuencia puede modificarse ajustando P1.

La red R-C, constituida por C2/R3, conectada a la salida de N1 actúa como un ele-



Lista de componentes:

Resistencias:

R1, R6 = 47 k
R2, R7 = 10 k
R3, R8 = 470 k
R4, R9 = 22 k
R5, R10 = 470 Ω (ver texto)
R11 = 100 Ω
P1, P2 = 1 M ajustable
resistencias de 1/8 W

Condensadores:

C1, C3 = 820 n
C2, C4 = 100 n
C5 = 10 μ /16 V

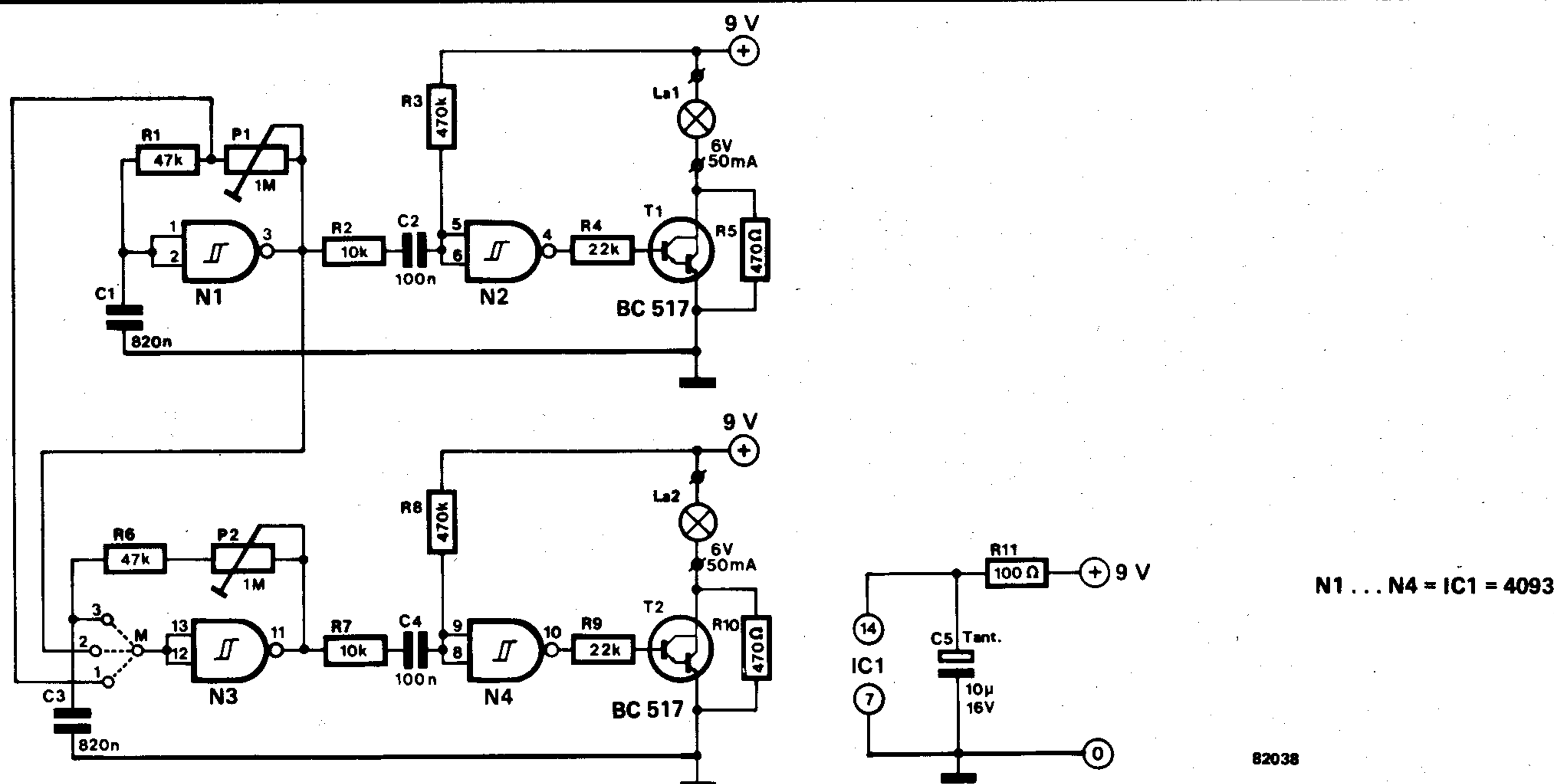
Semiconductores:

T1, T2 = BC 517
IC1 = 4093

Varios:

La1, La2 = lámpara de incandescencia miniatura de 6 V/50 mA (ver texto)

1



N1 ... N4 = IC1 = 4093

82038

Figura 1. El circuito está constituido por dos multivibradores idénticos. Dependiendo del efecto requerido hay tres métodos de conectar los dos circuitos.

2

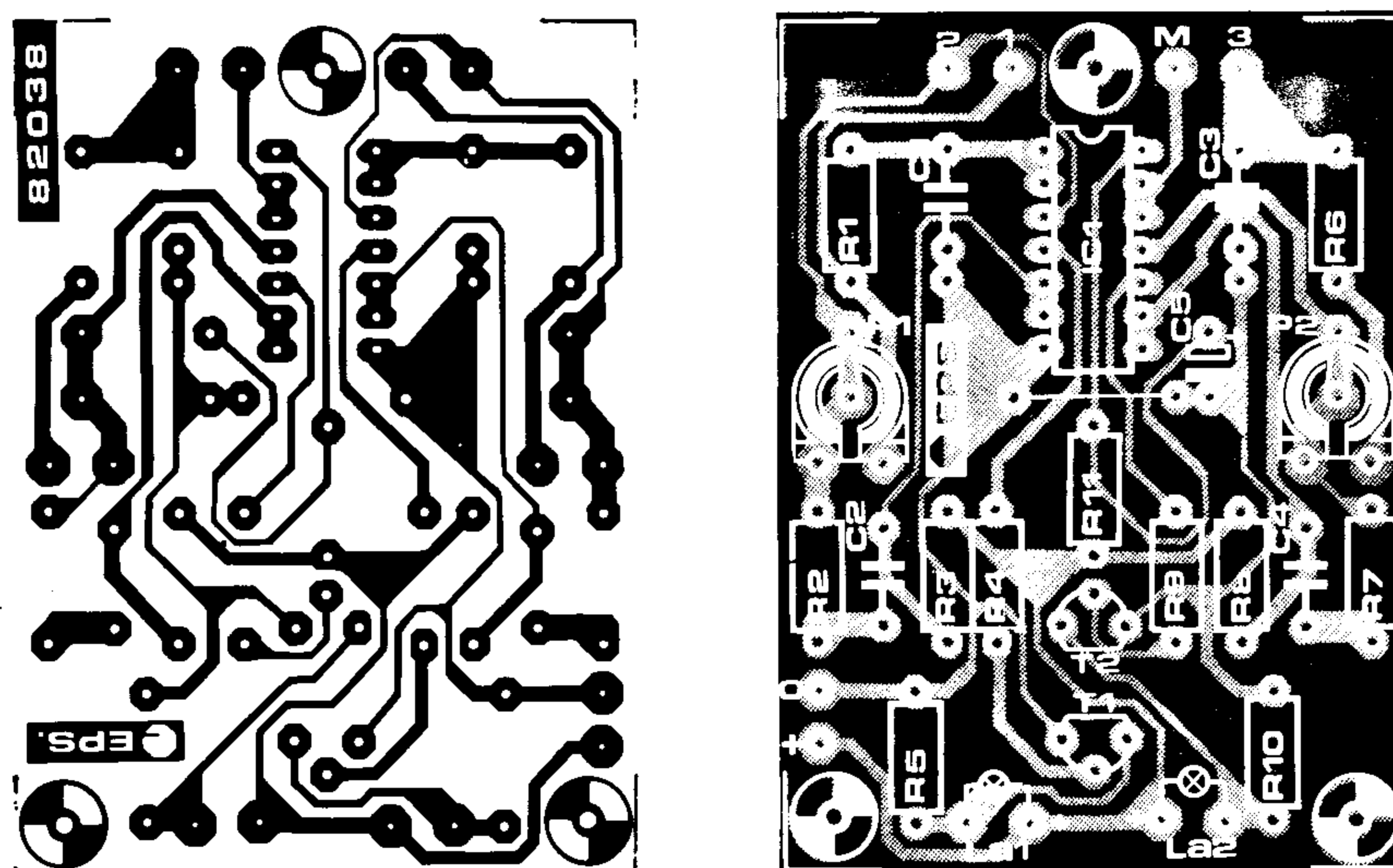


Figura 2. Placa de circuito impreso y disposición de los componentes para el circuito de luces intermitentes. El montaje es tan pequeño que puede incorporarse, con suma facilidad, en un coche de juguete.

funcionar de tres formas distintas, lo que se consigue con la ayuda de un puente móvil en la placa. Puenteando los puntos 3 y M se obtienen dos luces intermitentes completamente independientes. Puenteando los puntos 2 y M, las lámparas se iluminan alternativamente. A continuación, la frecuencia puede ajustarse por medio de P1. Finalmente, si se puentean 1 y M, las dos lámparas se encenderán simultáneamente. De nuevo, la frecuencia viene determinada por medio del ajuste de P1.

El circuito impreso

Los dos circuitos osciladores pueden montarse en la placa de circuito impreso mostrada

en la figura 2. Los controles de frecuencia, P1 y P2, pueden ser potenciómetros normales o simples ajustables. No hay que olvidar establecer el puente entre el punto M y uno de los puntos 1 ... 3. La tensión de alimentación para el circuito ha de estar comprendida entre 3 y 15 V, aunque, como se indicó anteriormente, resulta ideal una pila de 9 V (tipo PP3).

Para obtener un rendimiento óptimo, la tensión nominal de las lámparas debe ser aproximadamente 2/3 de la tensión de alimentación mientras que la intensidad no debe superar el límite de 400 mA.

Los valores óhmicos de las resistencias R5 y R10 deben elegirse en la práctica de tal forma que las lámparas estén a punto de encenderse.

No nos proponemos dar detalles sobre la instalación del montaje acabado en el coche miniatura, ya que depende en gran medida del juguete elegido. Todo lo que se necesita es un par de agujeros para las lámparas y la agilidad suficiente para elegir el lugar idóneo en el que introducir el montaje.

receptor BLU de onda corta

un compacto con grandes prestaciones

El término de banda lateral única (BLU) suele estar asociado con alto coste y complejidad. Pero no siempre tiene que ser así. Con el circuito de conversión directa, la señal de A.F. se convierte directamente en B.F. sin producir una frecuencia intermedia (F.I.). El empleo de esta técnica conduce a un receptor de BLU compacto, de bajo coste y con excelentes prestaciones. Aunque no se trate de un circuito sencillo, tenemos la seguridad de que este montaje introducirá en su hogar el mundo de la banda de los 20 metros, sin los gastos asociados a los equipos disponibles en el comercio.

Los receptores de comunicaciones son invariablemente muy complejos y caros... y a pesar de ello, en algunos casos, esto no garantiza, sorprendentemente, su buen rendimiento. En el artículo sobre los principios fundamentales de un receptor de BLU, que se presenta en este mismo número, se puede constatar que la construcción de un receptor de BLU exige bastante destreza y conocimientos. El diagrama de bloques simplificado, que se indica en el artículo teórico, ilustra un receptor medio que es, de hecho, difícil de construir. Sin embargo, los esquemas correspondientes a los aparatos existentes en el mercado van aún más lejos: son capaces de «poner los pelos de punta» a cualquier aficionado por muy experto que sea. El empleo de circuitos complejos de tipo superheterodino no es la única forma de diseñar un receptor de BLU. Un método más directo consiste en emplear la «conversión directa». Este principio permite la construcción de receptores mucho más simples que, no obstante, presentan unas características muy notables.

La diferencia principal entre un receptor de conversión directa y un superheterodino es el hecho de que el primer tipo no produce una frecuencia intermedia (F.I.). Aquí también, como sucede en el interior del superheterodino, la señal de entrada se mezcla con una señal de batido en un mezclador y como, en este caso, la frecuencia de batido es exactamente la misma que la de la señal de entrada; los productos de suma y de diferen-

cia suministrados por el mezclador quedan limitados a informaciones de B.F. El filtrado de la señal, necesario para obtener una buena selectividad, se efectúa en el subconjunto de B.F. del receptor (bloque LPF = Low-pass filter = filtro paso bajo). El oscilador de recepción de BLU desempeña también la función de oscilador de batido (BFO), pues trabaja a la misma frecuencia que la señal de entrada. Sin embargo, ha de hacerse notar que, en lo que respecta a la construcción y a la puesta en punto, este oscilador será uno de los pilares básicos del montaje, pues las exigencias de estabilidad que se imponen a un oscilador de batido son relativamente elevadas.

Las ventajas principales de un receptor de conversión directa sobre uno de diseño superheterodino son:

- Construcción sencilla y compacta.
- Facilidad de calibración y control.
- Al ser idénticas las frecuencias del oscilador y de la señal de entrada, se eliminan los problemas planteados por las frecuencias imágenes. Solamente los armónicos y los subarmónicos de la frecuencia del oscilador podrían crear algunas anomalías, pero este inconveniente también se presenta en los receptores superheterodinicos.
- Su precio de coste es moderado lo que se debe, sin duda alguna, a la escasa complejidad del montaje y también al hecho de que el filtrado necesario para obtener una selectividad aceptable tiene lugar en la parte de baja frecuencia. Un filtro de A.F. que proporcionara un ancho de banda idéntico a la del filtro de audio utilizado (a saber, -6 dB a 3 kHz y -60 dB a 5,5 kHz) costaría, sin lugar a dudas, en torno a las 6.000 ptas. En cuanto a las desventajas del receptor de conversión directa podemos citar las siguientes:
- Es sensible a lo que se podría denominar interferencia de frecuencia imagen de audio pues, en efecto, es sensible a las dos bandas laterales en lugar de a una, como nos interesaría.
- El margen dinámico es algo menor que el de un superheterodino, ya que la etapa del mezclador podría trabajar como un detector de A.M., si se supera la intensidad de la señal de entrada especificada.

Versatilidad

El receptor descrito es adecuado para la banda de 20 metros que va desde 14,00 a 14,35 MHz. Y no hemos seleccionado esta banda al azar, sino porque es la de uso más frecuente y, por lo tanto, la más interesante para escuchar. Desde hace algunos años, nos encontramos bajo la influencia de una actividad de manchas solares que posibilita que la banda de los 20 metros esté en uso durante las 24 horas del día, al existir condiciones favorables de propagación para las comunicaciones por radio en esta banda. Se trata, pues, de una buena inversión de su tiempo y de su dinero. Y nada nos obliga a limitar la utilización de su receptor a esta sola banda. La adición de un conversor adecuado permitirá la recepción de cualquier otra banda que hayamos elegido.

La gama de frecuencias de 0,5 MHz de ancho restringe un poco este empleo universal; en cualquier caso sólo la banda de ondas cor-



tas más «alta» (de 28 a 29,7 MHz) no es accesible en su integridad. Todas las demás bandas pueden transferirse a la banda de los 20 metros con la ayuda de conversores simples, lo que permite construir un sistema de recepción de comunicaciones completo de muy buena calidad.

Cuando la frecuencia utilizada es inferior a 14 MHz, no es necesario que los conversores desempeñen también la función de amplificadores e incluso, si llega el caso, puede actuar como atenuador. Si se quiere trabajar a una frecuencia más elevada, también es aconsejable emplear un conversor/amplificador que tenga una ganancia de unos pocos dB.

El circuito

En la figura 2 se muestra el diagrama de bloques de nuestro receptor en su forma definitiva. Hay que reconocer que es algo más complicado que el ilustrado en la figura 1, pero ello no debe asustarle sino permitirle comprender más a fondo el funcionamiento del conjunto, pues en dicho diagrama se incluyen todas las «etapas» de nuestro montaje. Si hubiéramos representado el esquema de un receptor de BLU standard con la misma precisión, nos habría ocupado de 5 a 6 páginas. Esta observación está destinada a aquéllos para los que la visión del esquema de la figura 2 haya sido «espeleznante» ¡Si no cunde el pánico, conseguiremos nuestro objetivo!

Comencemos por «sobrevolar a gran altura» el diagrama de bloques. El primer obstáculo que encuentra la señal de entrada es un filtro paso-banda (BPF1) destinado a determinar el ancho de banda de entrada y, por consiguiente, la amplitud de la gama de sintonía. A continuación se encuentra una etapa de amplificación de A.F. y luego, un segundo filtro paso-banda antes de que la señal llegue al mezclador, que es de un tipo muy peculiar.

La señal procedente del oscilador de sintonía se aplica también al mezclador a través de una etapa amplificadora-separadora («buffer»). La señal de salida de B.F., procedente del mezclador, se filtra fuertemente por medio de dos filtros paso-bajo (LPF1 y LPF2). Entre estos dos filtros hay una pequeña etapa de amplificación de B.F., a la que se añade un limitador de ruido.

Se percibe la existencia de una nueva etapa de amplificación de B.F. después del segundo filtro paso-bajo. Se trata de un CAG utilizado para limitar el nivel de entrada al mezclador y proteger, pues, a las etapas de entrada del receptor respecto a las tensiones de entrada excesivas. Una etapa de amplificación simple se añade para alimentar al altavoz, con lo que queda completado el diagrama.

Los componentes

En la figura 3 se presenta el circuito completo. Habremos de estudiar conjuntamente las representaciones de las figuras 2 y 3 para conseguir un conocimiento profundo del circuito.

El filtro paso-banda BPF1 del diagrama de bloques se encuentra ahora en forma circuital L1/C1/C52. El margen de sintonía conseguido con este filtro es de unos 500 kHz (desde 14 a 14,5 MHz), que es suficiente para

cubrir la banda de 20 metros y lo bastante estrecha para impedir la recepción de señales parásitas procedentes de emisores que operan en la banda de 19 metros.

La preamplificación de A.F. se obtiene gracias a un MOSFET de doble puerta (T1). Este transistor desempeña funciones suplementarias. Ante todo, sirve de «buffer» entre el oscilador y la antena, lo que impide cualquier radiación de una señal del oscilador por la antena, también constituye un preamplificador para la entrada y forma una parte activa del CAG, que es una función auxiliar muy importante.

El filtro paso-banda BPF2 está constituido por las bobinas L4 ... L7 y los condensadores C6 ... C13. Ello nos da un filtro relativamente grande que tiene un ancho de banda de unos 3 MHz y una curva de respuesta muy plana en el interior de la banda de 20 metros. Lo cual permite, en su conjunto, evitar un efecto de «detección de fase», que constituye un fenómeno «desagradable» que se caracteriza por una sensibilidad exagerada del receptor a las vibraciones mecánicas. El siguiente componente es el mezclador (T2). El principio de este mezclador pasivo de conductibilidad unilateral se muestra en la figura 4, y con él se asegura que nada en absoluto se alimente a la salida cuando no haya ninguna señal de A.F. a la entrada. Cuando exista una señal de entrada de

A.F., sólo habrá a la salida la señal de entrada y, por consiguiente, no la señal producida por el oscilador. El núcleo fundamental del mezclador es un MOSFET de doble puerta muy versátil. Los elementos que constituyen el mezclador se han elegido de manera que aseguren la mayor linealidad posible. La puesta en funcionamiento y el bloqueo, de forma continua, del mezclador exigen una alta tensión en el oscilador. Para conseguir un alto nivel de estabilidad de frecuencia para el oscilador se volvió a utilizar un MOSFET de doble puerta y de buena calidad: el BF 900 (T3).

Sobre la base de este FET se ha construido una variante de un oscilador de Clapp, con su proverbial estabilidad. La sintonía se obtiene por intermedio de un pequeño varicap (diodo de capacidad variable) D4.

Una tensión de sintonía procedente del regulador de tensión, IC1, se envía a los diodos varicap; dicha tensión de sintonía es regulable gracias a P1. Este potenciómetro es del tipo de 10 vueltas (multivuelta), lo que permite obtener la sintonía del receptor de manera relativamente fácil. A continuación del oscilador hay una etapa amplificadora-separadora (T4); la señal del oscilador se transmite inmediatamente al mezclador a través de C14. Y así llegamos a la zona de B.F. del receptor.

Inmediatamente después del mezclador, en-

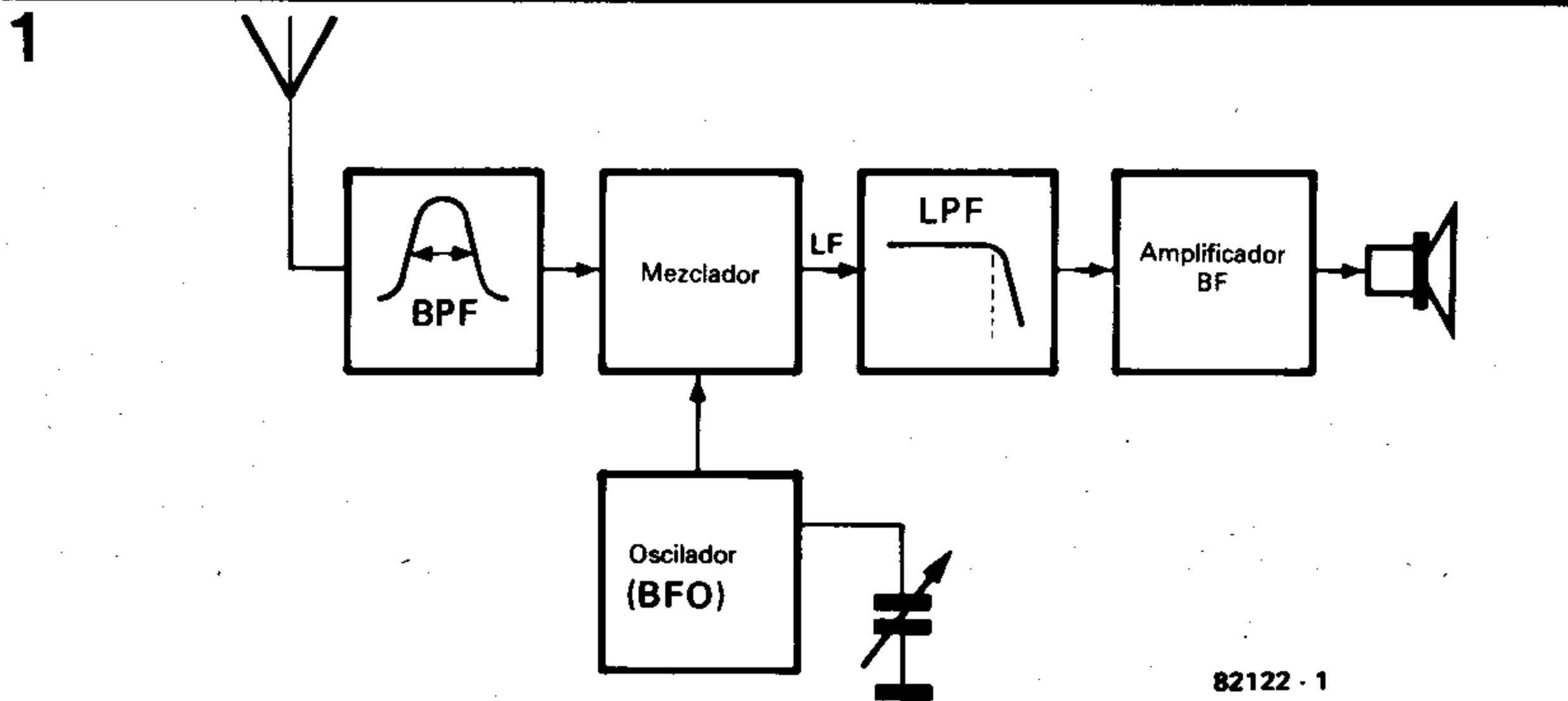


Figura 1. Un receptor que funciona según el principio de la conversión directa y de concepción claramente más sencilla que la de un aparato del tipo superheterodino. La frecuencia de entrada y la del oscilador son idénticas y, por consiguiente, no se produce ninguna frecuencia intermedia. En BLU, el oscilador servirá también de oscilador de batido (BFO).

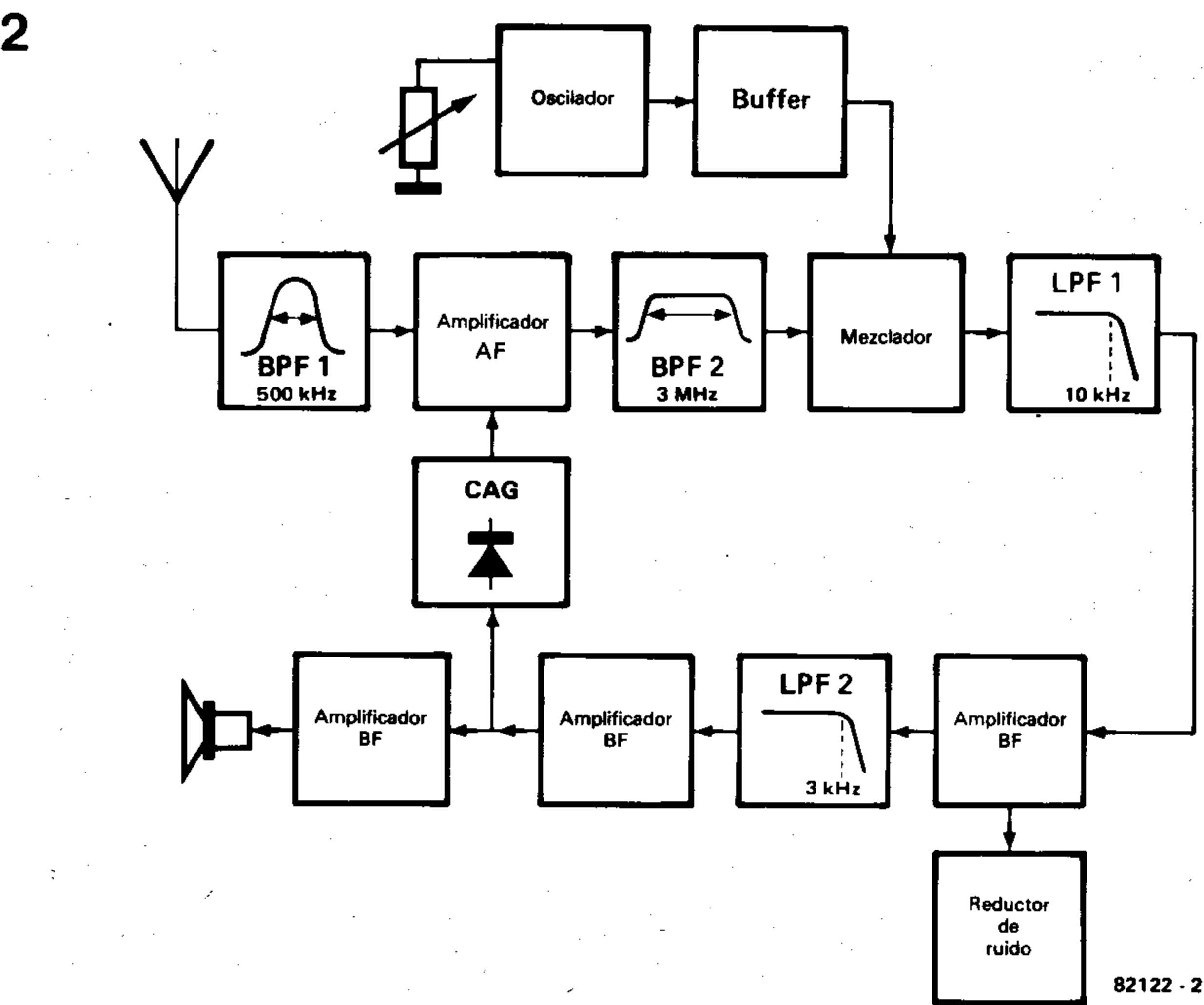
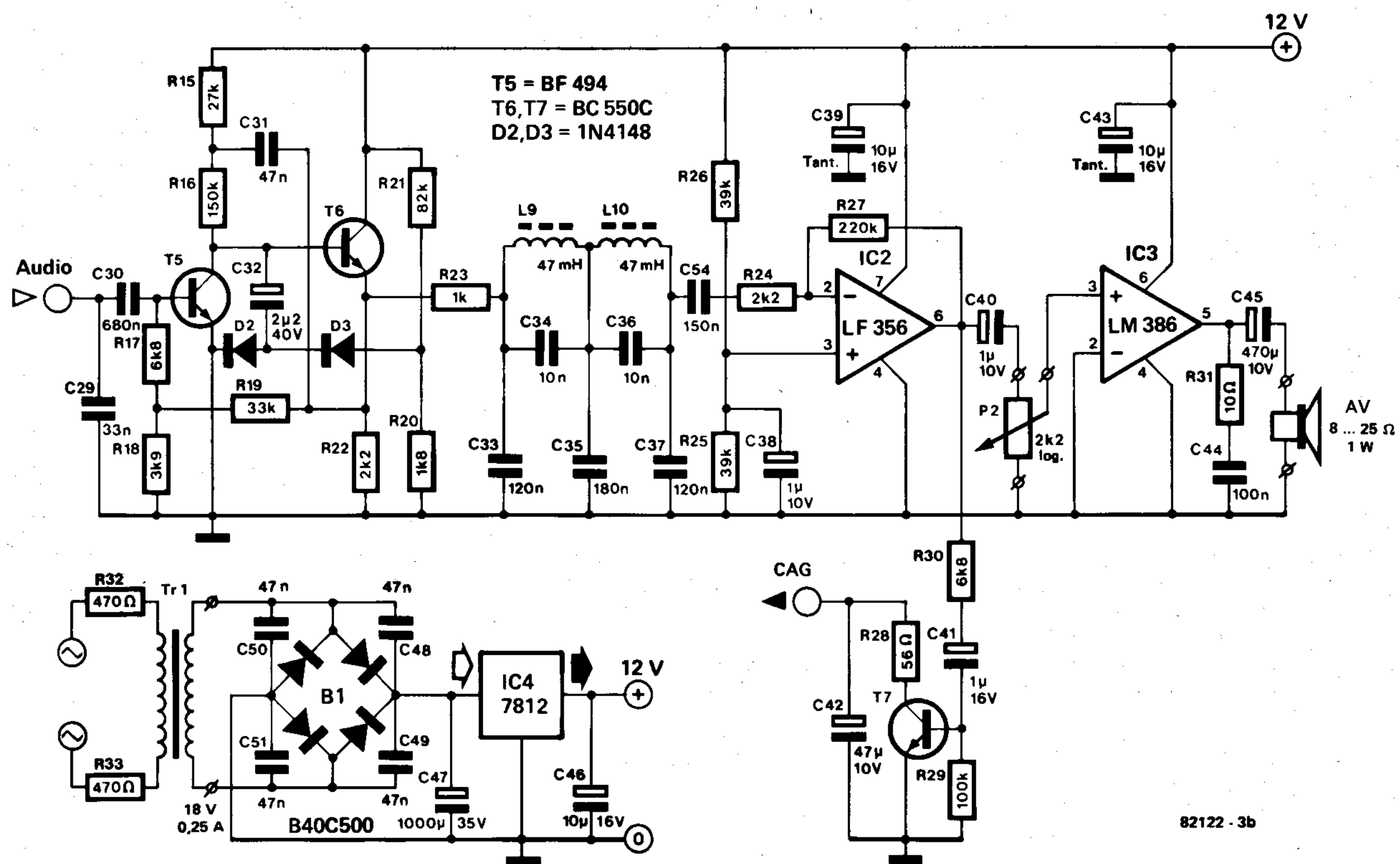


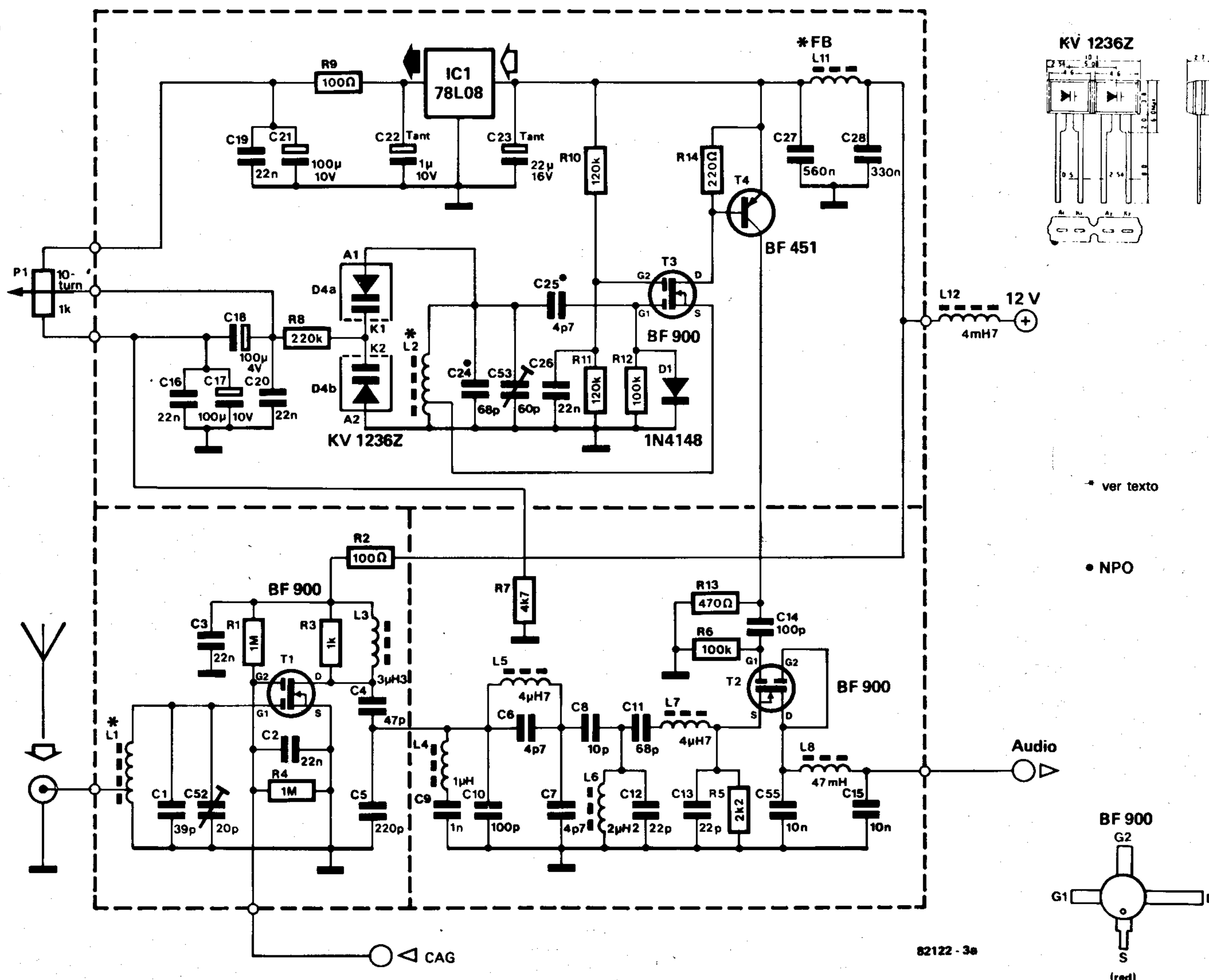
Figura 2. Diagrama de bloques completo de nuestro receptor.

3a



82122 - 3b

b



82122 - 3a

(red)

Figura 3. Esquema de principio. En la figura 3a se ilustra la zona de BF y la alimentación, mientras que la figura 3b muestra la parte de A.F.

4

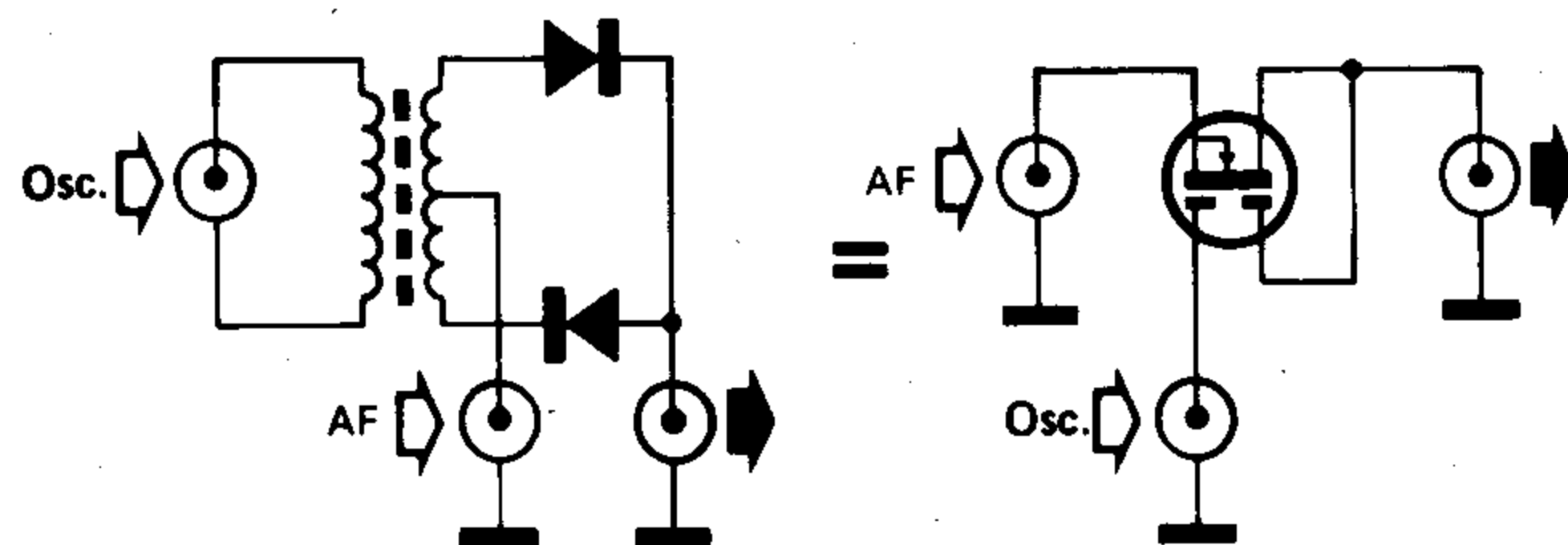
contramos un filtro paso-bajo relativamente sencillo (LPF1). Los elementos que los constituyen son la bobina L8 y los condensadores C15, C29 y C30. Su frecuencia de transición es relativamente elevada, situándose en las proximidades de 10 kHz pues, de no ser así, el limitador de ruido no sería efectivo. Este limitador es esencialmente un recortador a diodos (D2 y D3) ordinario, que también forma parte del amplificador de A.F. (T5 y T6). La señal se amplifica y se filtra, una vez más, por un segundo filtro de paso bajo (LPF2) constituido por L9, L10 y C33... C37. Este último filtro elimina cualquier componente de la señal que esté por encima de 3 kHz (unos 66 dB por octava). El control automático de la ganancia (CAG) está constituido por T7 y sus componentes asociados. Un detector de un sólo transistor (T7) rectifica parte de la señal de salida de B.f. de IC2 convirtiéndola en corriente continua. El nivel de c.c. es proporcional a la intensidad de la señal de B.F. A continuación, esta tensión de c.c. se realimenta a la segunda puerta (G2) de la etapa de A.F. (T1).

Si se supera el umbral base/emisor de T7 (con fuertes señales de entrada), la tensión puerta 2/fuente de T1 disminuye y ello hace más pequeña la ganancia de este FET. La subida se ha elegido rápida y la caída relativamente lenta, de manera que se evite el fenómeno de «bombeo», que caracteriza a algunos de los CAG corrientes.

Acabamos de pasar revista a los diversos subconjuntos que constituyen nuestro receptor de BLU. Sólo falta por mencionar a IC3, que es un amplificador de audio integrado, capaz de excitar directamente un altavoz y que no precisa más que un pequeño número de componentes suplementarios. El potenciómetro P2 permite la regulación del volumen.

Montaje

Para simplificar el trabajo, hemos optado por diseñar dos circuitos impresos: uno para la parte de alta frecuencia y otro para la de baja. En la figura 5 se muestra la sección



82122 - 4

Figura 4. Principio del mezclador pasivo utilizado en este circuito.

de alta frecuencia (A.F.), que también se ilustra como un esquema circuital en la figura 3a. Asimismo, la parte de baja frecuencia de la placa ilustrada en la figura 6 corresponde al esquema de la figura 3b. Con la excepción del transformador, todos los componentes de la fuente de alimentación pueden montarse en la placa de baja frecuencia (B.F.).

Se deja al constructor la elección de si la placa de circuito impreso se deja en una sola pieza o se corta en dos. Para conseguir un producto final lo más compacto posible, en el prototipo de Elektor se separaron las placas y se hizo el montaje en «sanwich». Si se separan, las placas indican claramente las interconexiones correspondientes, tales como señal de B.F., tensión CAG, alimentación y masa. También son fácilmente identificables los puntos de conexión que van a los componentes fuera de la placa (altavoz, tensión de alimentación, potenciómetros P1 y P2). No hay que olvidar la conexión en serie de una bobina de choque L12 al enlazar la alimentación entre las secciones de A.F. y de B.F. Destaquemos que no está señalada la ubicación del montaje de dicha bobina en ninguna de las placas.

Hagamos algunas observaciones sobre la instalación de los componentes en los circuitos impresos. Las dos placas son de doble cara y el lado en que están dispuestos los componentes (que denominaremos cara superior) queda cobreado y sirve de masa. Será preciso, pues, soldar en esta cara superior todos

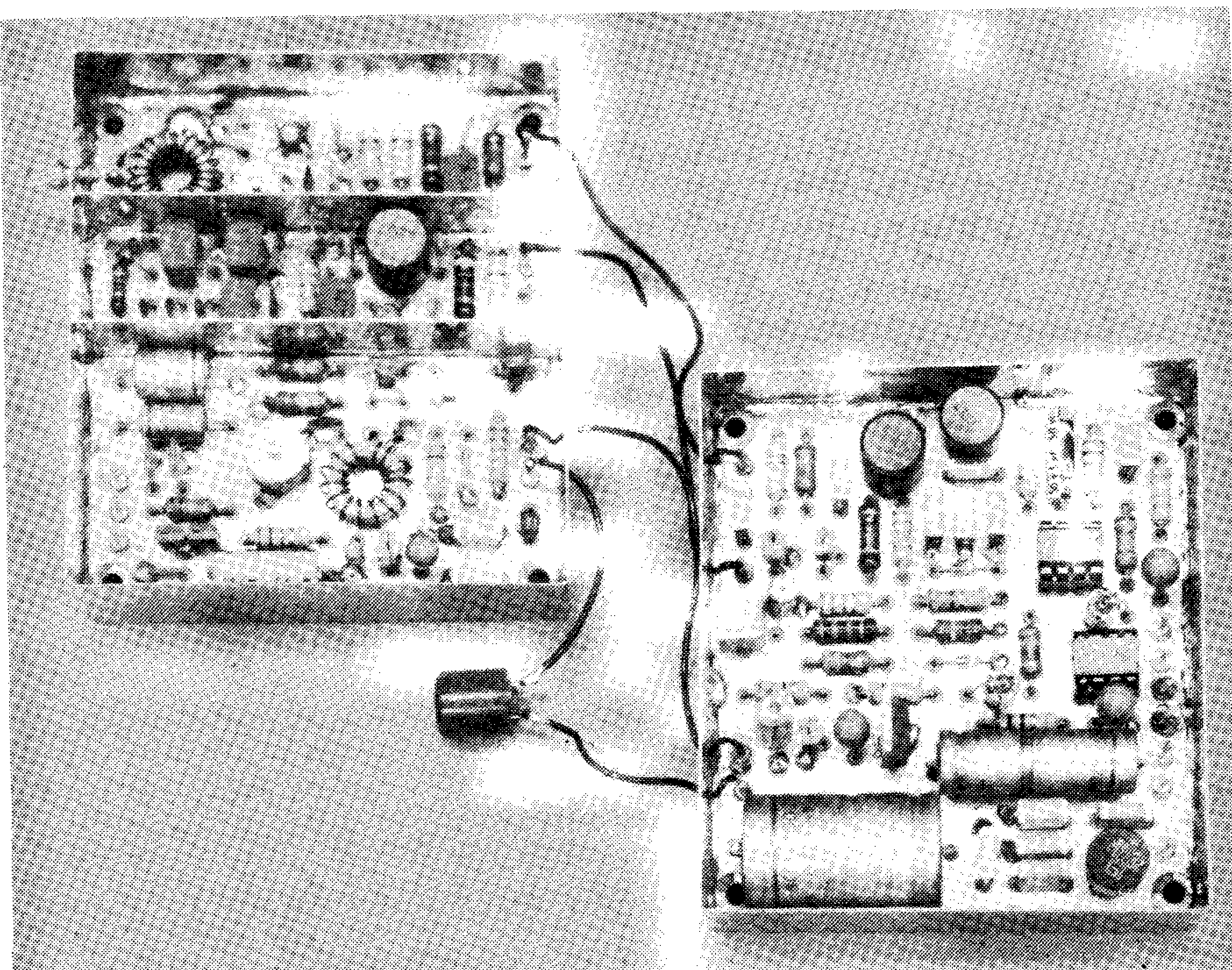
los componentes que hayan de conectarse a masa. Todos los demás puntos de no soldadura están caracterizados por anillos de aislamiento fresados en el cobre alrededor de los agujeros.

En la figura 3 se ilustra claramente el patillaje del FET BF900 y del doble varicap KV1236Z. Hay que tener el máximo cuidado cuando se proceda al montaje de estos componentes, sobre todo de los varicaps. Asimismo, hay que prestar especial atención al montaje de los condensadores ajustables C52 y C53 que están provistos de tres patillas, de las cuales sólo se emplean dos. Si no se conectan en la forma correcta se podría producir un cortocircuito que afectaría a todo el circuito.

Veamos, ahora, las bobinas. Hay quienes son «alérgicos» a la construcción de bobinas y, sin duda, pensarán que hay «demasiados» bobinas en este receptor. Pero, para su alivio, diremos que, en su mayoría, son de tipo standard y basta con comprarlas. Sin embargo, no sólo es importante adquirir las bobinas correctas sino también montarlas en el lugar correcto. Un descuido sería fatal para el funcionamiento satisfactorio del montaje.

L1 y L2 no existen como taleas en el comercio y será preciso construirlas. El bobinado se efectúa sobre núcleos toroidales del tipo T50-6. L1 se obtiene enrollando 21 espiras de hilo de cobre esmaltado de 0,4 mm. en el toroide y disponiendo una toma de derivación exactamente en la penúltima espira antes de la masa. L2 necesita 12 espiras de hilo de cobre de 0,6 mm. de espesor y una toma intermedia a dos espiras de la masa. El espaciado de las espiras se hará de manera que se obtenga una bobina, cuyas espiras estén distribuidas bien simétricamente a lo largo de todo el núcleo. Una vez que el receptor esté terminado y puesto a punto, es recomendable pegar con adhesivo ambas bobinas en la placa de circuito impreso. También pueden fijarse con la ayuda de tornillos y de tuercas, en cuyo caso será preciso perforar también el agujero de fijación en el circuito impreso.

Las etapas del amplificador, mezclador y oscilador de la sección de A.F. deben separarse entre sí mediante una placa de cobre o de estaño. En la placa de circuito impreso y en el esquema se indica claramente en donde han de montarse (ver fotografía 2). Cuando esté concluido el ajuste definitivo, se recomienda encarecidamente que se cierren los compartimentos por arriba, para evitar que una radiación parásita alcance a los componentes. De esta forma impediremos una eventual interacción. Todas estas recomendaciones las damos con el fin de evitar, a toda costa, una reacción del oscilador sobre



5+6

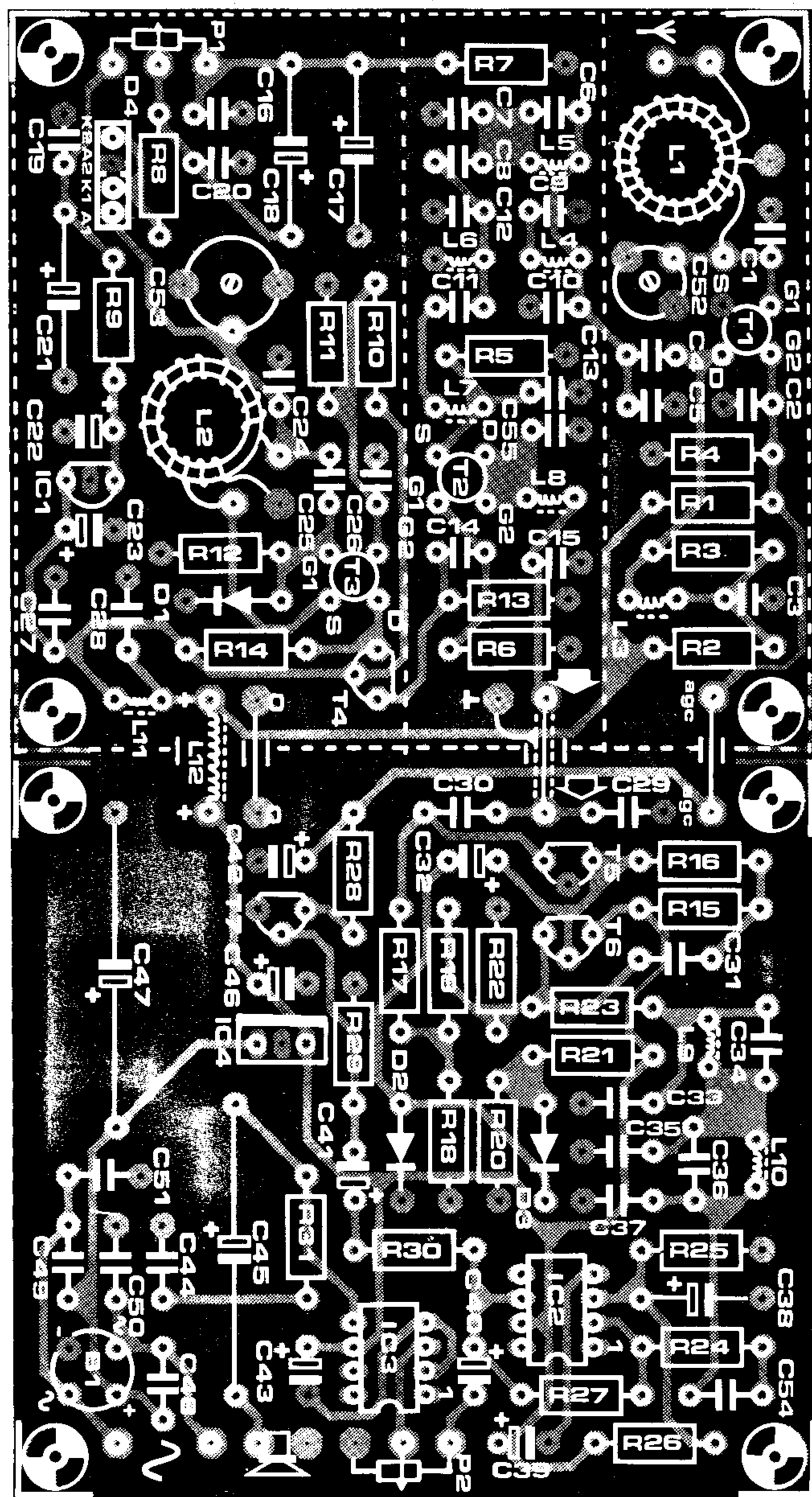
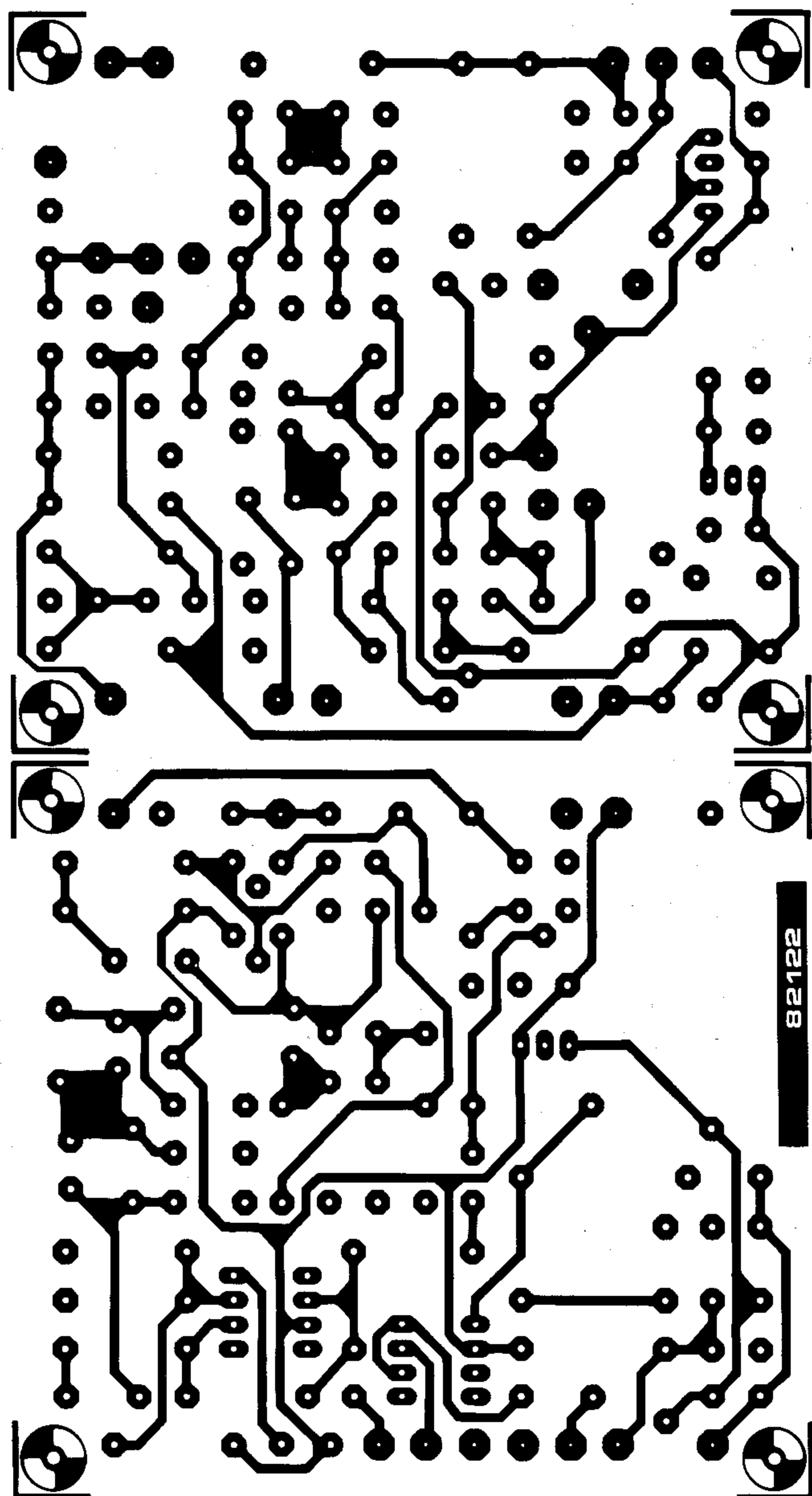


Figura 5. Lado de cobre y lado de componentes de la sección de A.F. La placa de circuito completa tiene doble cara.

Figura 6. Sección de B.F. El lado de los componentes es una sola superficie de cobre amplia.

la antena, lo que se traduce inmediatamente en la producción de zumbidos, de vibraciones microfónicas o de ambas cosas a la vez.

Es preferible una caja metálica para alojar el montaje. Si se han separado, de forma completamente hermética, los diversos compartimentos, nada impide el empleo de una caja de plástico. Los mejores resultados se obtienen colocando trozos de goma espuma entre las placas y las partes laterales y el fondo de la caja. Ha de dejarse para el final la interconexión entre las secciones de A.F. y de B.F. y los demás componentes, con el empleo de hilo de conexión flexible.

Además de los circuitos impresos, se pondrán en la caja los potenciómetros, el conector coaxial de antena, el interruptor de red y un pequeño transformador (18 V/0,25 A). En cambio, es preferible poner el altavoz en una caja separada para evitar toda clase de realimentación indeseable. El altavoz debe ser de buena calidad y vale la pena «rascar-

se un poco el bolsillo» al respecto, pues un altavoz malo no sólo reduce la inteligibilidad de las señales de salida, sino que también, en algunos casos, puede eliminar todas las posibilidades de comprensión de las señales recibidas.

Puesta a punto

La puesta a punto del receptor no requiere grandes conocimientos y el procedimiento a seguir es bastante sencillo.

Comencemos por ajustar el condensador C52 a su posición media (unos 10 pF) y C53 al máximo (posición de capacidad mínima) y luego, llevar P1 a la posición que permita medir una tensión de unos +8 V en su cursor (con un polímetro situado entre el cursor de P1 y masa). Afortunadamente, la tensión suministrada por IC1 es de +8 V por lo que todo lo que se precisa es ajustar P1 al máximo.

Llegados a este punto, nos encontramos con la alternativa de si se dispone, o no, de un frecuencímetro. En el caso de que se disponga de un frecuencímetro el procedimiento es el siguiente:

- conectar el frecuencímetro a la puerta 1 de T2 a través de una sonda de alta impedancia.
- actuar sobre el condensador ajustable C53 con la ayuda de un destornillador (de plástico), de modo que se obtenga una lectura de una frecuencia del oscilador de 14.360 kHz.
- conectar, a continuación, la antena y girar P1 cinco vueltas de manera que se llegue a las proximidades del punto medio de la banda de los 20 metros.
- ajustar C52 de modo que se tenga la señal máxima. Si se tiene dudas en cuanto a la mejor posición de ajuste, queda la posibilidad de medir la tensión del CAG y de actuar sobre C52 de forma que ésta se haga mínima.

Lista de componentes:

Resistencias:

- R1,R4 = 1 M
R2,R9 = 100 Ω
R3,R23 = 1 k
R5,R22,R24 = 2k2
R6,R12,R29 = 100 k
R7 = 4k7
R8,R27 = 220 k
R10,R11 = 120 k
R13,R32,R33 = 470 Ω
R14 = 220 Ω
R15 = 27 k
R16 = 150 k
R17,R30 = 6k8
R18 = 3k9
R19 = 33 k
R20 = 1k8
R21 = 82 k
R25,R26 = 39 k
R28 = 56 Ω
R31 = 10 Ω
P1 = 1 k (10 vueltas)
P2 = 2k2 log.

Condensadores:

- C1 = 39 p
C2,C3,C16,C19,C20,C26 = 22 n ceramico
C4 = 47 p
C5 = 220 p
C6,C7 = 4p7
C8 = 10 p
C9 = 1 n MKM
C10,C14 = 100 p
C11 = 68 p
C12,C13 = 22 p
C15,C55 = 10 n ceramic
C17,C21 = 100 μ/10 V
C18 = 100 μ/4 V
C22 = 1 μ/10 V tantalo
C23 = 22 μ/16 V tantalo
C24 = 68 p coeficiente de temperatura 0
C25 = 4p7 coeficiente de temperatura 0
C27 = 560 n MKM
C28 = 330 n MKM
C29 = 33 n MKM
C30 = 680 n MKM
C31,C48 . . . C51 = 47 n MKM

- C32 = 2μ2/40 V
C33,C37 = 120 n MKM
C34,C36 = 10 n MKM
C35 = 180 n MKM
C38,C40 = 1 μ/10 V
C39,C43,C46 = 10 μ/16 V tantalo
C41 = 1 μ/16 V
C42 = 47 μ/10 V
C44 = 100 n MKM
C45 = 470 μ/10 V
C47 = 1000 μ/35 V
C52 = 20 p ajustable
C53 = 60 p ajustable
C54 = 150 n MKM
atención: C32, C40, C41 y C42 deben montarse verticalmente.

Bobinas:

- L1 = 21 espiras de hilo de cobre esmaltado de 0,4 mm. de diámetro, con toma intermedia a 1 espira del extremo de masa.
L2 = 14 espiras de hilo de cobre esmaltado de 0,6 mm. con toma intermedia a 2 espiras del extremo de masa.
(L1 y L2 bobinadas sobre núcleo toroidal T50-6; Amidón-Amarillo.
L3 = 3.3 μH
L4 = 1 μH
L5,L7 = 4.7 μH TOKO
L6 = 2.2 μH
L8,L9,L10 = 47 mH
L11 = núcleo de ferrita con 4 espiras de hilo de cobre esmaltado de 0,3 mm. de diámetro.
L12 = 4.7 mH (Toko)

Semiconductores:

- D1,D2,D3 = 1N4148
D4 = KV 1236Z (Toko)
T1,T2,T3 = BF 900
T4 = BF 451
T5 = BF 494
T6,T7 = BC 550C
IC1 = 78L08
IC2 = LF 356
IC3 = LM 386
IC4 = 7812
B1 = B40C500 redondo

Varios:

- Tr1 = Transformador de red 18 V/250 mA
LS = altavoz 8 . . . 25 Ω/1 W

Observación: No hay que olvidar que, en este caso, el CAG reacciona con una cierta inercia.

Veamos, ahora, cómo proceder si no se dispone de un frecuencímetro:

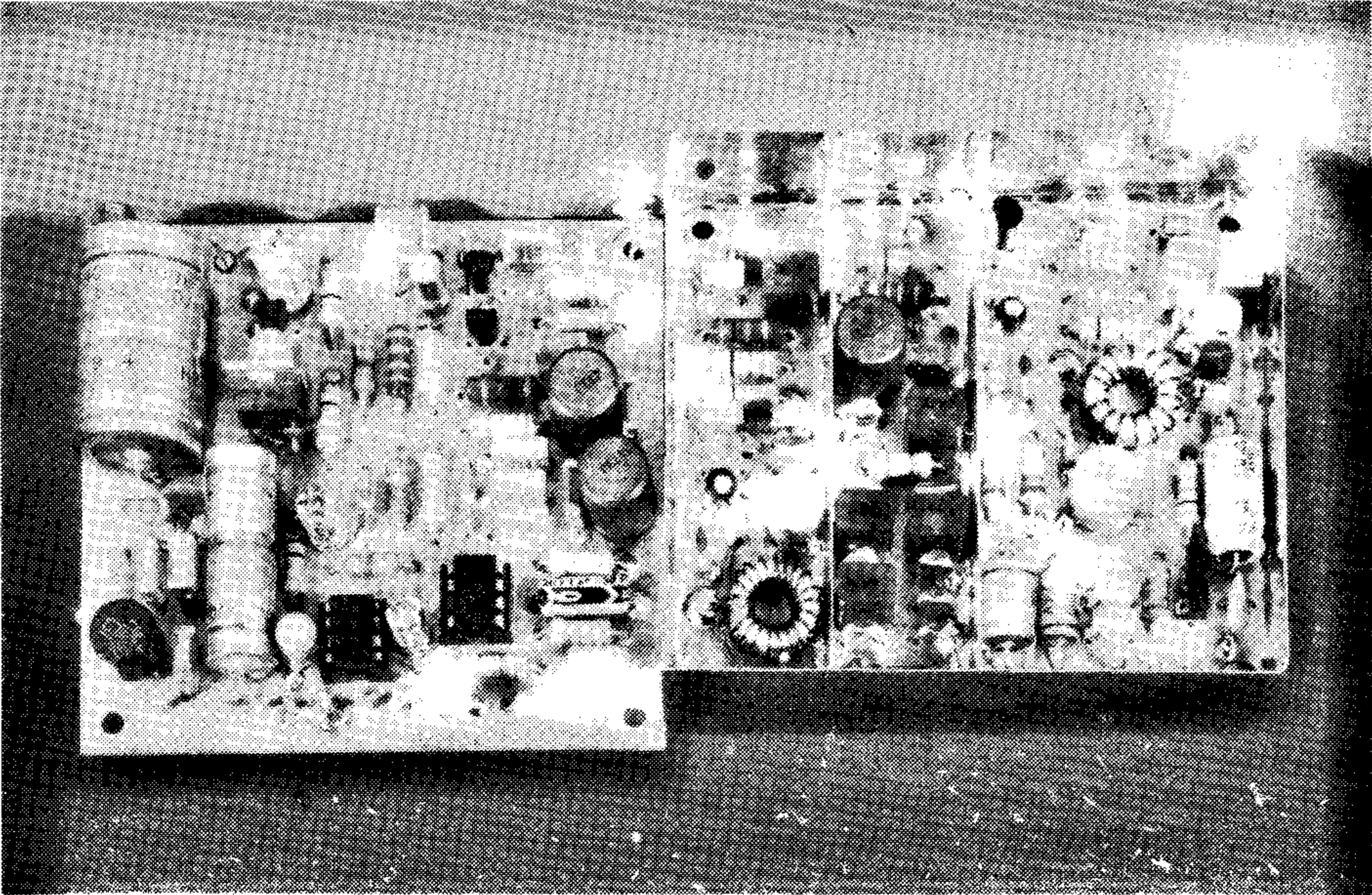
- conectar la antena
- girar C53 hasta que se escuche la deformación de la voz característica de la BLU («voz del pato Donald»). Proseguir, muy lentamente, la rotación (aumentando la capacidad) hasta conseguir que aparezcan señales Morse en el altavoz.
- dar cinco vueltas a P1 y actuar sobre C52 siguiendo el procedimiento anteriormente descrito (P1 en su posición media y C52 para una tensión del CAG máxima).

¡A la escucha!

Unos pocos metros de hilo, estratégicamente colocados, resultan suficientes para obtener una antena muy aceptable. Una antena característica para la banda de los 20 metros es una varilla vertical de unos 5 metros de longitud.

Los principiantes en BLU deberán habituarse a la manipulación de sus equipos para conseguir la sintonía correcta. Pero no hay que preocuparse pues se consigue con rapidez. Las señales de prueba están disponibles en el éter, en abundancia, lo que permite hacerse rápidamente una idea sobre el funcionamiento correcto del aparato. Como ya se ha indicado al principio del artículo, las condiciones de propagación permanecen buenas a lo largo de todo el día y de la noche en esta banda; por ello no se deben tener dificultades para encontrar numerosas emisiones útiles cuando se accione el mando correspondiente. La actividad es algo menor por la mañana, pero siempre quedan algunas emisiones. La banda de los 20 metros suele servir también para la telegrafía, por lo que recomendamos al constructor que refresque sus conocimientos en Morse. Las prestaciones de nuestro receptor son especialmente impresionantes para un aparato que funciona según el principio de la conversión directa. La sensibilidad medida en uno de nuestros prototipos nos da un valor de, como mínimo, 0,15 μV para una relación señal/ruido de 10 dB. En la práctica, esto significa que este pequeño montaje es capaz de estar a la altura de los aparatos existentes en el comercio. No se trata, pues, de un simple montaje de bricolaje.

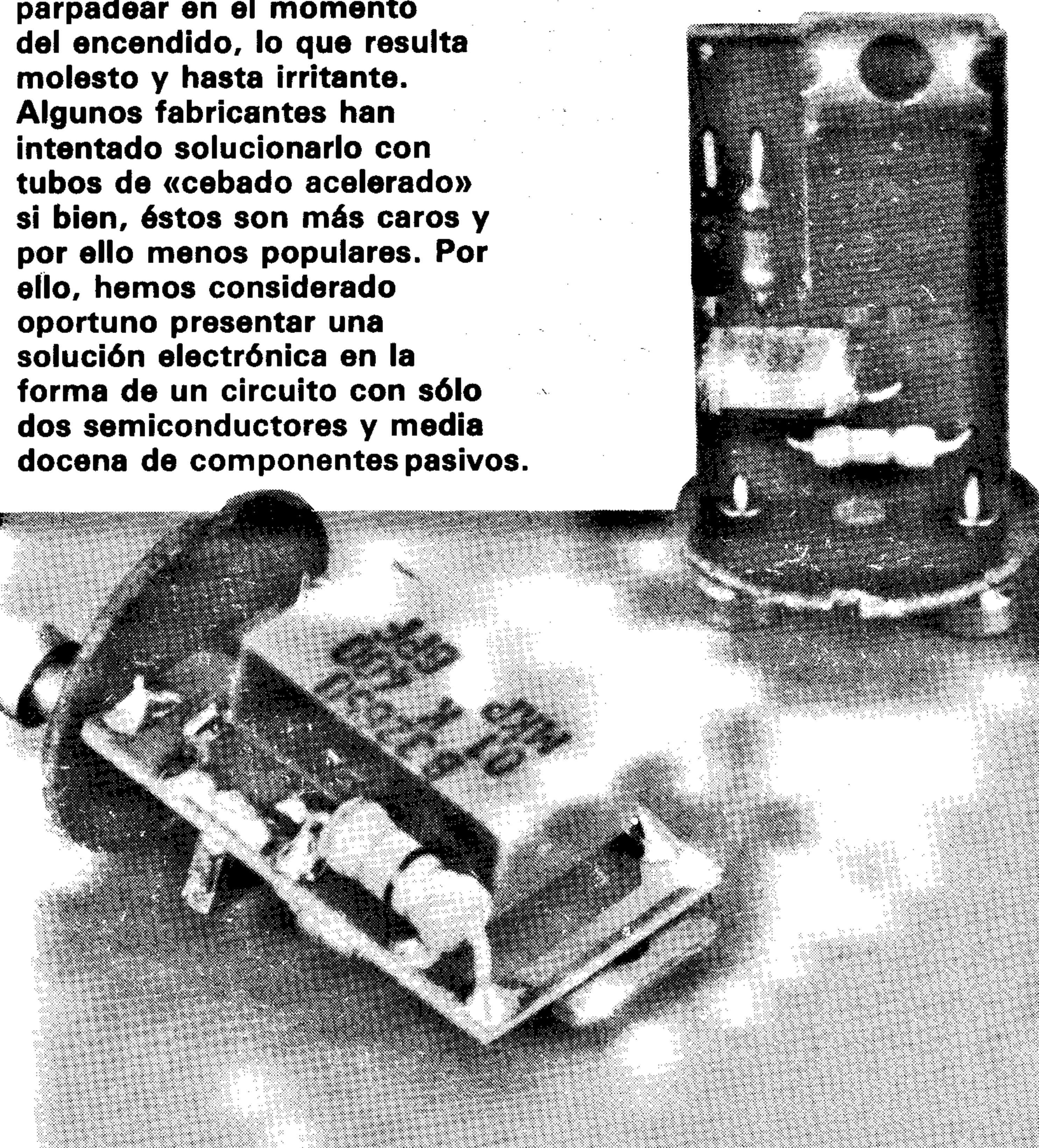
Una observación final. El consumo del montaje no es un modelo de economía, pero es «soportable». Si se pone el potenciómetro de volumen en un valor medio, el consumo no sobrepasará los 40 mA. Esto debería permitir una utilización en modo de funcionamiento portátil con alimentación a pilas. La forma más fácil de conseguir esto último es conectar en serie dos pilas compactas de 9 V y conectar el conjunto en paralelo con c47 (en la alimentación). Las pilas alcalinas de manganeso tienen una capacidad media del orden de 500 mA-h y ello supone una potencia suficiente para un empleo de 10 horas.



**la solución electrónica para un
encendido sin centelleos!**

cebador electrónico para fluorescentes

Los tubos fluorescentes suelen parpadear en el momento del encendido, lo que resulta molesto y hasta irritante. Algunos fabricantes han intentado solucionarlo con tubos de «cebado acelerado» si bien, éstos son más caros y por ello menos populares. Por ello, hemos considerado oportuno presentar una solución electrónica en la forma de un circuito con sólo dos semiconductores y media docena de componentes pasivos.



No se trata de un circuito espectacular, aunque su efecto sí lo es. Resulta sorprendente, desde luego, que los tubos fluorescentes ordinarios se enciendan sin centelleo. Por añadidura, las pequeñas dimensiones del circuito deben permitir montarlo directamente en la caja del antiguo cebador. No es precisa modificación alguna ni en el tubo ni en su soporte.

Empecemos por examinar el sistema tradicional. Esencialmente, el fluorescente es un tubo de cristal que contiene vapor de mercurio a una presión muy pequeña (aproximadamente 0,00001 atmósferas o sea, considerablemente menor que la presión atmos-

férica normal). Cuando este gas se somete a un campo eléctrico adecuado y con potencial suficiente, resulta ionizado y se produce una descarga eléctrica; el gas conduce una cierta corriente eléctrica, emitiendo una luz prácticamente invisible puesto que se sitúa en el espectro ultravioleta. Las paredes internas del tubo están recubiertas de un polvo fluorescente muy fino y es así como la luz ultravioleta se convierte en luz visible. Este polvo actúa, en cierta forma, como convertidor, de forma que las ondas cortas de la luz ultravioleta se alargan de modo que se hacen visibles.

Para facilitar el cebado del tubo, se mezcla un poco de argón (un gas raro) con el vapor

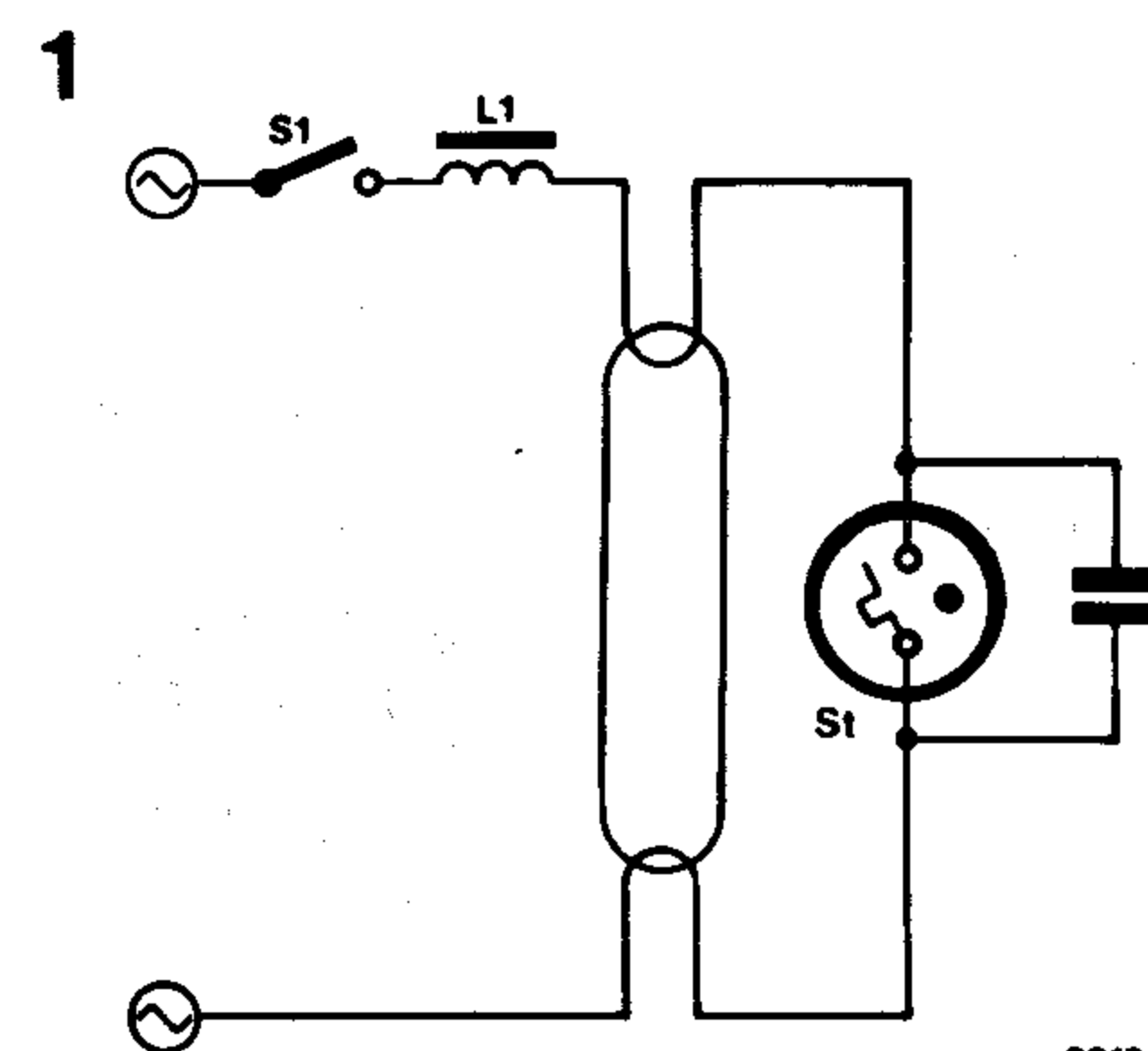


Figura 1. Una fuente de luz fluorescente está constituida por una bobina de amortiguamiento, un cebador mecánico y un tubo fluorescente.

de mercurio, actuando a modo de catalizador. La tensión de cebado está estrechamente ligada con la temperatura del tubo en forma inversamente proporcional. Es decir, cuanto más baja sea la tensión tanto más alta habrá de ser la temperatura del gas. Por este motivo, los filamentos están montados en uno u otro extremo del tubo con el fin de precalentar el gas (ionización previa) para estimular el encendido. Una vez que haya comenzado la descarga (tubo iluminado), la tensión puede bajar hasta una tensión (más débil) de mantenimiento de la descarga. El tubo se comporta como una resistencia negativa a la tensión de descarga; es decir, que disminuye a medida que aumenta la corriente. Es necesaria, pues, una limitación de corriente para evitar la destrucción del tubo. Se utiliza para ello una bobina de amortiguamiento («choque») que, debido a sus propiedades de resistencia inductiva, disipa muy poca potencia en forma de calor. Por lo tanto, se utiliza como bobina de inducción, combinada con un cebador para proporcionar la alta tensión de cebado del tubo. Esta bobina también tiende a suprimir las interferencias de R.F. causadas por la descarga de gas en el interior del tubo. La función de un cebador no se limita a proporcionar una tensión de inducción, sino que también sirve para la aplicación de una corriente a los electrodos de cebado. El tipo más común de cebador consiste en una ampolla llena de helio que contiene un interruptor bimetalico (interruptor térmico), lo que se ilustra en la figura 2. En reposo, los contactos de este interruptor están abiertos. Cuando se cierra el interruptor de red (S1), la tensión de la red se aplica al cebador, lo que es suficiente para encender la lámpara de helio. La estructura de diseño del cebador sólo permite que circule una corriente débil (de unos 0,1 amperios). El calor producido por la descarga en el gas provoca el cierre del interruptor bimetalico. En consecuencia, una corriente intensa circula a través de los filamentos y el tubo entra en su fase de precalentamiento. El cierre del interruptor bimetalico equivale a un cortocircuito interno de la lámpara de helio, que, en consecuencia, se apagará. Transcurrido un breve período de tiempo, la temperatura en el interior de la lámpara descenderá a un nivel tan bajo que se abrirá el interruptor bimetalico (posición de reposo). La interrupción de la corriente es muy brusca, con lo que se induce una alta tensión en la bobina que se aplica a ambos extremos del tubo y, por consiguiente, se encenderá el tubo. Una vez que se haya encendido, la tensión a través de los bornes del cebador es igual a la tensión de descarga del tubo y dicha tensión no es suficiente para producir un nuevo encendido de la lámpara de helio. Puesto que el interruptor bimetalico sigue abierto, el cebador permanecerá inactivo mientras que el tubo se mantenga encendido. Se puede considerar, pues, que el cebador está fuera del circuito una vez que se haya cebado el tubo. En paralelo con el cebador hay un condensador cuya función esencial es suprimir cualquier interferencia de parásitos de R.F. en el tubo. Lamentablemente, es preciso considerar que este primer cebado, que acabamos de describir, no suele ser el definitivo. En efecto, la temperatura del gas no es suficiente para que la descarga pueda mantener-

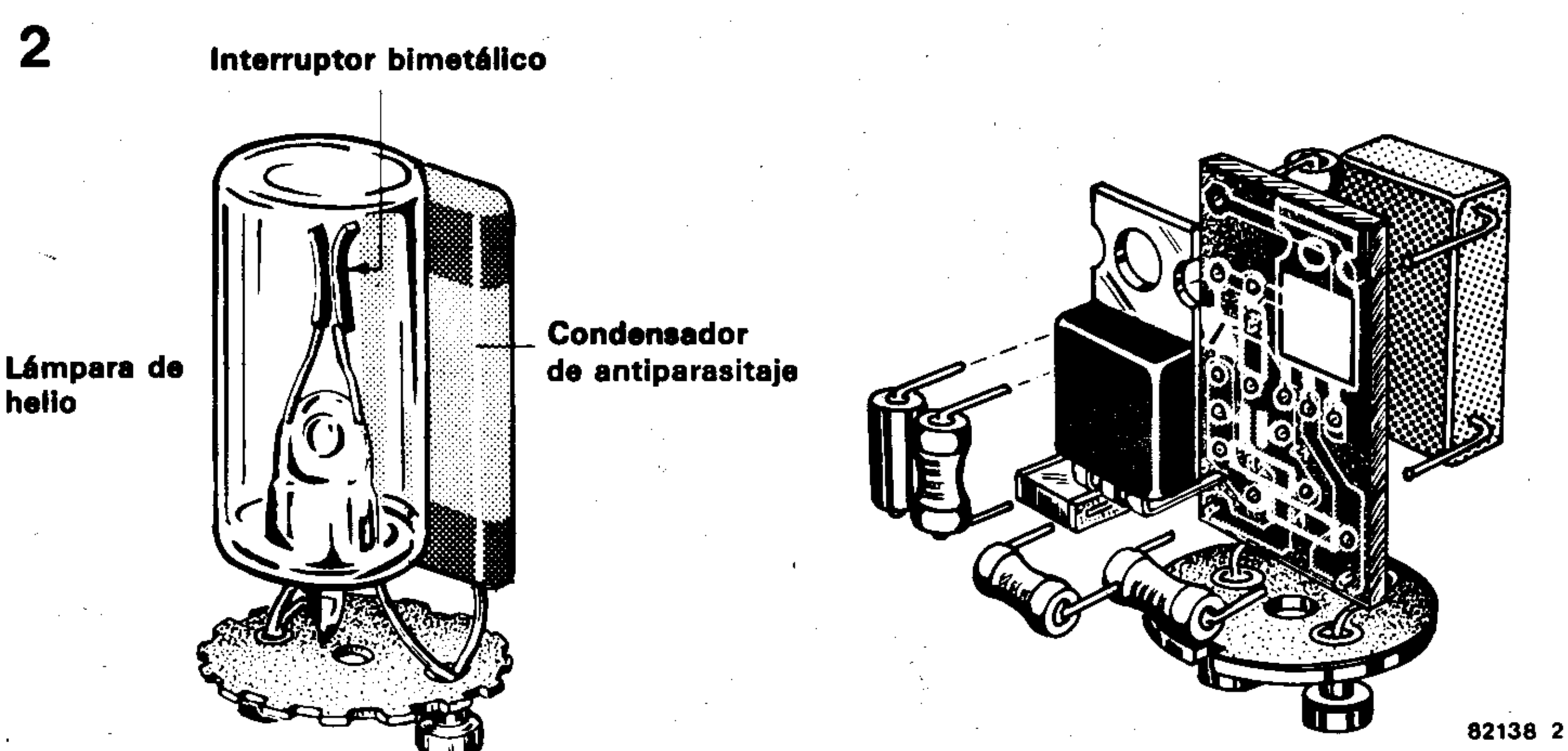


Figura 2. La mayor parte de los cebadores consisten en una ampolla de helio que contiene un interruptor bimetalico. En paralelo con el cebador hay un condensador para suprimir cualquier interferencia de RF producida por la descarga del gas.

se sin interrupción. Además, la corriente podría ser prácticamente nula cuando se abra el interruptor bimetalico, en cuyo caso, no se dispondrá de la tensión de inducción suficiente. El tubo tendrá que cebarse varias veces para que quede definitivamente encendido. Como el cebado se realiza con medios parcialmente mecánicos, y por ello bastante lentos, no hay que extrañarse de que se produzca un retardo visible entre cada proceso de encendido. Ello explica por qué el tubo centellea después de haberse encendido. Si se quieren suprimir los centelleos, tendremos que asegurarnos de que el tubo está suficientemente precalentado y de que los cebados sucesivos no estén separados por intervalos de magnitud perceptible.

Lista de componentes

- Resistencias:**
R1 = 470 k
R2 = 100 k
R3 = 1 k
R4 = 56 Ω
- Condensadores:**
C1 = 15 n (ver texto)
C2 = 100 n/630 V
- Semiconductores:**
D1 = diac ER 900
Th1 = tiristor TIC 106D

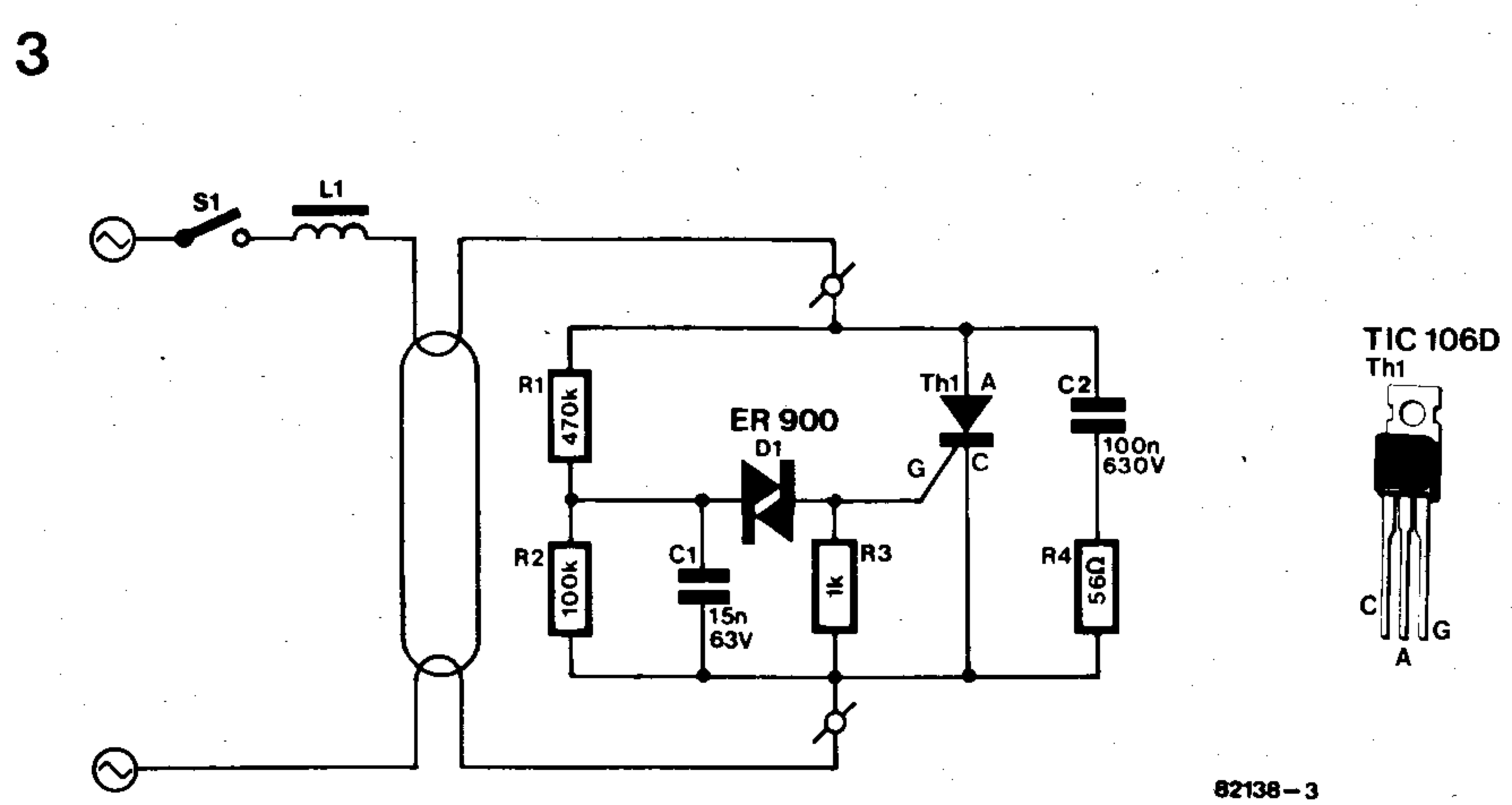


Figura 3. El circuito de cebado electrónico sólo tiene ocho componentes. Garantiza un encendido instantáneo de los tubos fluorescentes, suprimiendo los desagradables centelleos (inevitables hasta ahora).

4

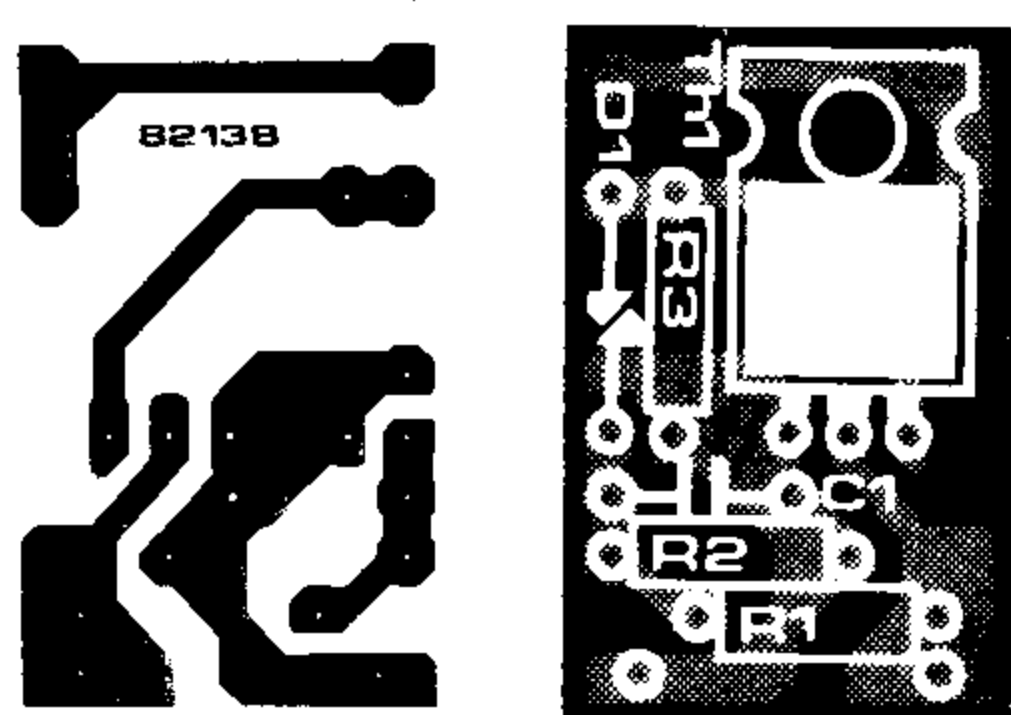


Figura 4. Placa de circuito impreso y disposición de componentes para el circuito del cebador electrónico. La placa montada es suficientemente compacta como para alojarse en el interior de la caja del antiguo cebador mecánico. Por razones de seguridad, no debe utilizarse una caja metálica.

El circuito

En la figura 3 se muestra el esquema del cebador electrónico para tubos fluorescentes. Para comenzar, es preciso considerar que el interruptor S1 está cerrado y que la tensión de ánodo del tiristor es más positiva que su tensión de cátodo. En tanto que el tubo no esté encendido, la tensión en los bornes del cebador electrónico es igual al valor instantáneo de la tensión de la red. Una vez que la carga del condensador C1 (a través del divisor R1/R2) sea suficiente para el cebado del diac (aproximadamente 30 V), su descarga ceba al tiristor cuya conducción da lugar a la aparición de una corriente consiguiente en los electrodos y en la bobina; esta corriente induce, a su vez, un campo magnético en la bobina. Cuando la onda de la red se hace negativa (inversión de polaridad), queda una corriente positiva circulando a través de la bobina, pero sólo durante un breve período de tiempo, hasta que desaparezca el campo magnético. En este momento, el tiristor se bloquea y el valor instantáneo de la tensión de la red aparece en los bornes del tubo, con lo que inmediata y rápidamente se carga C2. Asociado a L1, este condensador forma un circuito resonante que lleva la tensión a los bornes del tubo a un potencial considerablemente más elevado que la tensión de la red (aproximadamente el doble). Este nivel de alta tensión cebará inmediatamente el tubo. Cuando llega el siguiente semiciclo positivo de la onda de la red, el tiristor volverá a conducir y se reiniciará el ciclo, a un régimen de 50 veces por segundo. Después de varios ciclos, el tubo estará suficientemente caliente para permanecer encendido y por ello, la tensión en los contactos del cebador descenderá a la de funcionamiento del tubo (nivel de descarga), que no basta para cebar al diac y, por consiguiente, el tiristor quedará inactivo. En consecuencia, el cebador electrónico quedará en un estado de reposo.

Construcción del circuito

En la figura 4 se muestra la placa de circuito impreso correspondiente. Es suficiente-

mente compacta para poder instalarse en el interior de la caja de plástico (no de aluminio) de un cebador tradicional, sin tener que modificar para nada los soportes de los tubos fluorescentes. Es preciso cerciorarse de que los cables de conexión del tiristor no están en contacto con el disipador de calor metálico. De ser necesario, puede pegarse el tiristor a la placa con un poco de resina epoxídica. Sólo hay 8 componentes, de los cuales R4 y C2 deben montarse en el lado de cobre de la placa de circuito impreso. En la fotografía se muestra el producto acabado desde dos ángulos distintos. Como se mencionó anteriormente, no recomendamos el empleo de una caja metálica por razones de seguridad.

La caja del cebador ha de abrirse con cuidado. Retire la lámpara de helio y el condensador de las clavijas de contacto. No hay que cortar los hilos de conexión del condensador demasiado cortos, pues pueden utilizarse para soldar la placa de circuito impreso a las clavijas de contacto. Monte cuidadosamente el cebador e insértelo en un soporte de tubo fluorescente. El circuito está concebido para fluorescentes en la gama de 20 a 65 W.

En el caso de que un tubo de 20 W no se encienda directamente, hay que bajar el valor de C1 a 10 nF. El valor capacitivo de este condensador depende realmente del tipo de tubo fluorescente utilizado. Incidentalmente, lo mismo se aplica a C2. Si se seleccionan potencias de tubos fluorescentes inferiores a 20 W, puede ser necesario probar valores distintos.

próximamente en elektor

Ordenador para laboratorio fotográfico

Termómetro a LCD

Sintetizador polifónico

El Junior habla BASIC

Interface para unidad de disco flexible

... para el Junior y cualquier sistema basado en microprocesador 6502.

Cargador de baterías

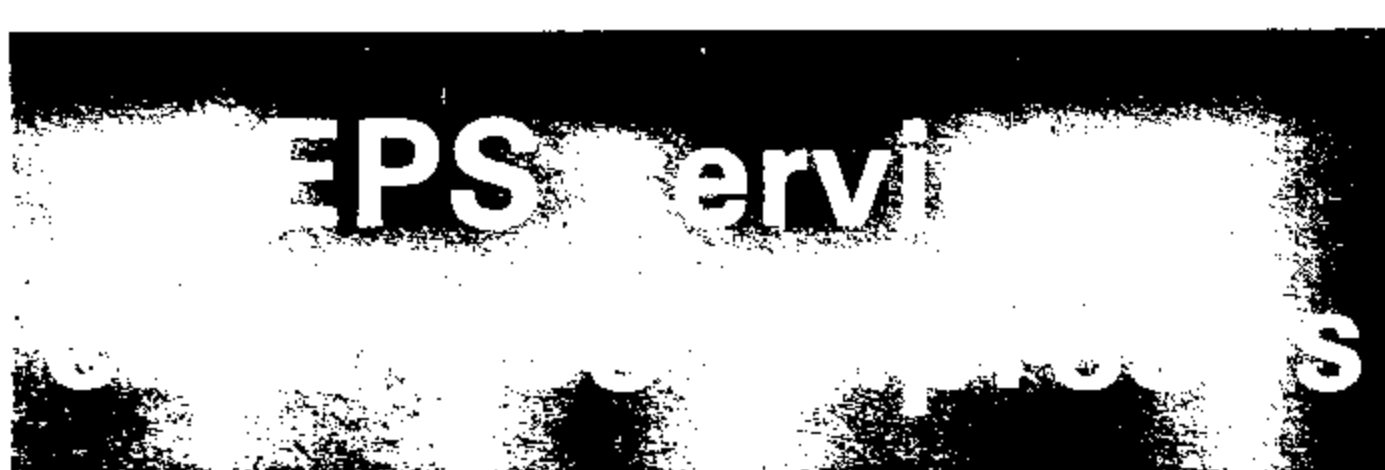
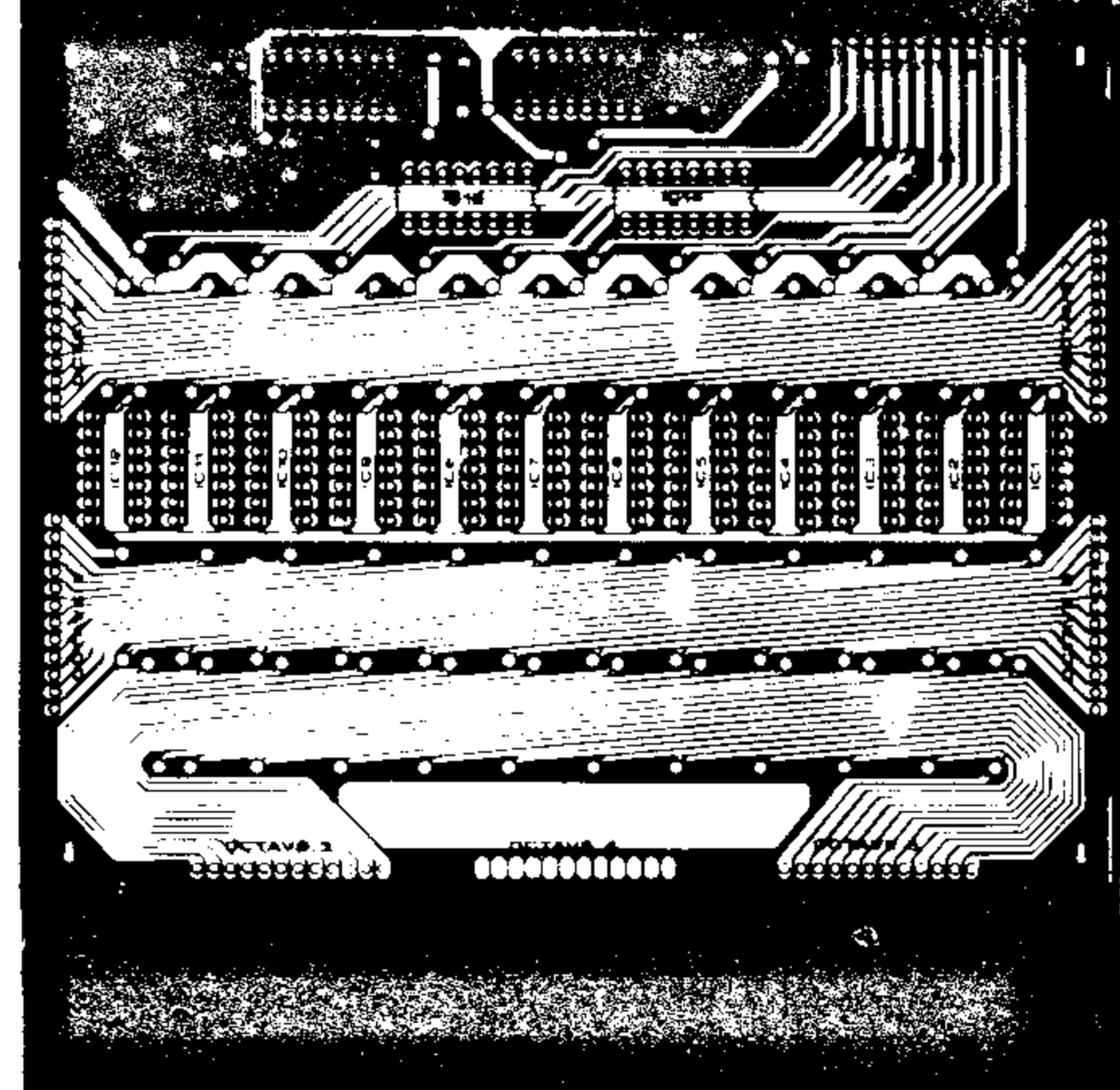
Téster trifásico

Antenas activas

¡Más Curso de BASIC!

próximamente en elektor

SERVICIOS DE ELEKTOR



Muchos de los circuitos de Elektor se acompañan de la correspondiente placa de circuito impreso.

Para algunos de estos diseños se suministra la placa de circuito impreso totalmente acabada, con serigrafía y taladrado.

La sección «servicios de elektor» de cada revista, incluye una relación completa de los circuitos impresos disponibles.

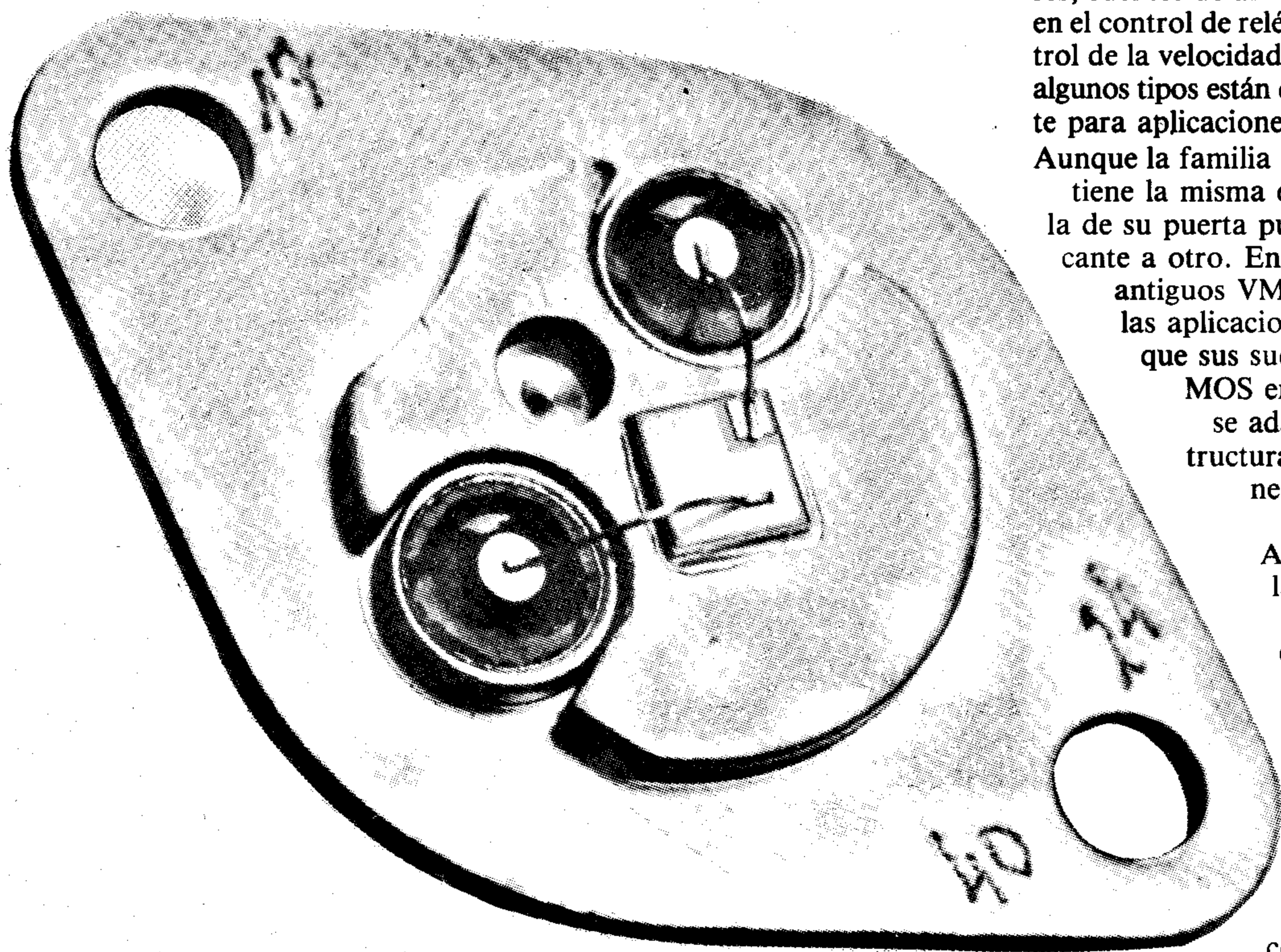
Hay que tener en cuenta que sólo son suministrables las placas de circuito impreso que figuran en la relación EPS. El hecho de publicar en determinado artículo el diseño de su circuito impreso, no implica necesariamente que sea suministrable por Elektor. El pedido de las placas de circuito impreso se realizará por medio de la correspondiente tarjeta en la que, además de la relación de placas, se indicará la forma de pago y los datos del lector.

Dentro de la sección «Quién y Dónde» están referenciados los establecimientos de electrónica que, además de distribuir la revista ELEKTOR, comercializan las placas del servicio EPS.



... los nuevos FETs de potencia

introducción a la familia DMOS



Cada día aparecen nuevos transistores FET de potencia que aumentan la sopa de letras ya existente (VFET, HEXFET, DMOS, TMOS, SIPMOS, etc.) A pesar de esta gran diversidad, todos tienen mucho en común en lo que respecta a sus características, estructura y aplicaciones. En este artículo se examinan los FETs en general y se presta especial atención a la rama constituida por la familia DMOS caracterizada por su rápida conmutación.

El término VFET les sonará familiar a la mayoría de los lectores, aunque probablemente pocos hayan visto uno «en persona». Y es fácil constatar que merecen una mayor popularidad. Si retrocedemos al año 1976, nos encontramos con que se los consideraba como los transistores de salida (casi) ideales para amplificadores de audio. Sin embargo, debido a su alto precio y poca disponibilidad nunca llegaron a ser populares. Se produjo el clásico círculo vicioso de que los componentes no bajan su precio hasta que sean fáciles de obtener y se hagan populares... y viceversa.

Hace aproximadamente un año tuvo una buena aceptación en la familia VFET un nuevo producto: la serie DMOS, también conocida bajo la denominación D-MOS FET o DFET. Esencialmente, son de funcionamiento muy similar a los VFET, pero su estructura es algo diferente y sus tiempos de conmutación son mucho más cortos. Los transistores FET de la familia DMOS se han dado a conocer, sobre todo, como conmutadores rápidos. Se prevé que ocupen una gran parcela del mercado de transistores de potencia y pueden utilizarse en convertidores, fuentes de alimentación conmutadas y en el control de relés y sistemas para el control de la velocidad de motores. Asimismo, algunos tipos están concebidos concretamente para aplicaciones de alta frecuencia.

Aunque la familia DMOS, en su conjunto, tiene la misma estructura fundamental, la de su puerta puede variar de un fabricante a otro. En términos generales, los antiguos VMOS se adaptan mejor a las aplicaciones en altas frecuencias que sus sucesores de la familia DMOS en cambio, estos últimos se adaptan mejor, por su estructura plana, a las aplicaciones con niveles de tensión más altos.

Antes de seguir más adelante, echaremos un vistazo a las características principales de la familia VFET en su conjunto y dejaremos a un lado, de momento, sus propiedades individuales. Ante todo, necesitamos averiguar en qué difieren los FETs respecto a sus contrapartidas bipolares.

En pocas palabras, podemos decir que los FETs cuestan menos que los tipos bipolares, tienen una conmutación más rápida (en unos pocos nanosegundos), proporcionan impedancias de entrada más altas con parámetros de excitación de valor bajo y tienen una amplísima gama de posibilidades circuitales.

Hasta el momento presente, los transistores D-MOS son todavía unos componentes raros (y todo lo que es raro es caro...), no obstante, esperamos, muy fundamentalmente, que esta situación cambie en un futuro no muy lejano.

Los elementos MOSFET

La tecnología MOS, cada vez más familiar gracias a los transistores (MOS FET) y a los circuitos integrados (p-MOS, n-MOS, CMOS), entran en juego dentro de todos los

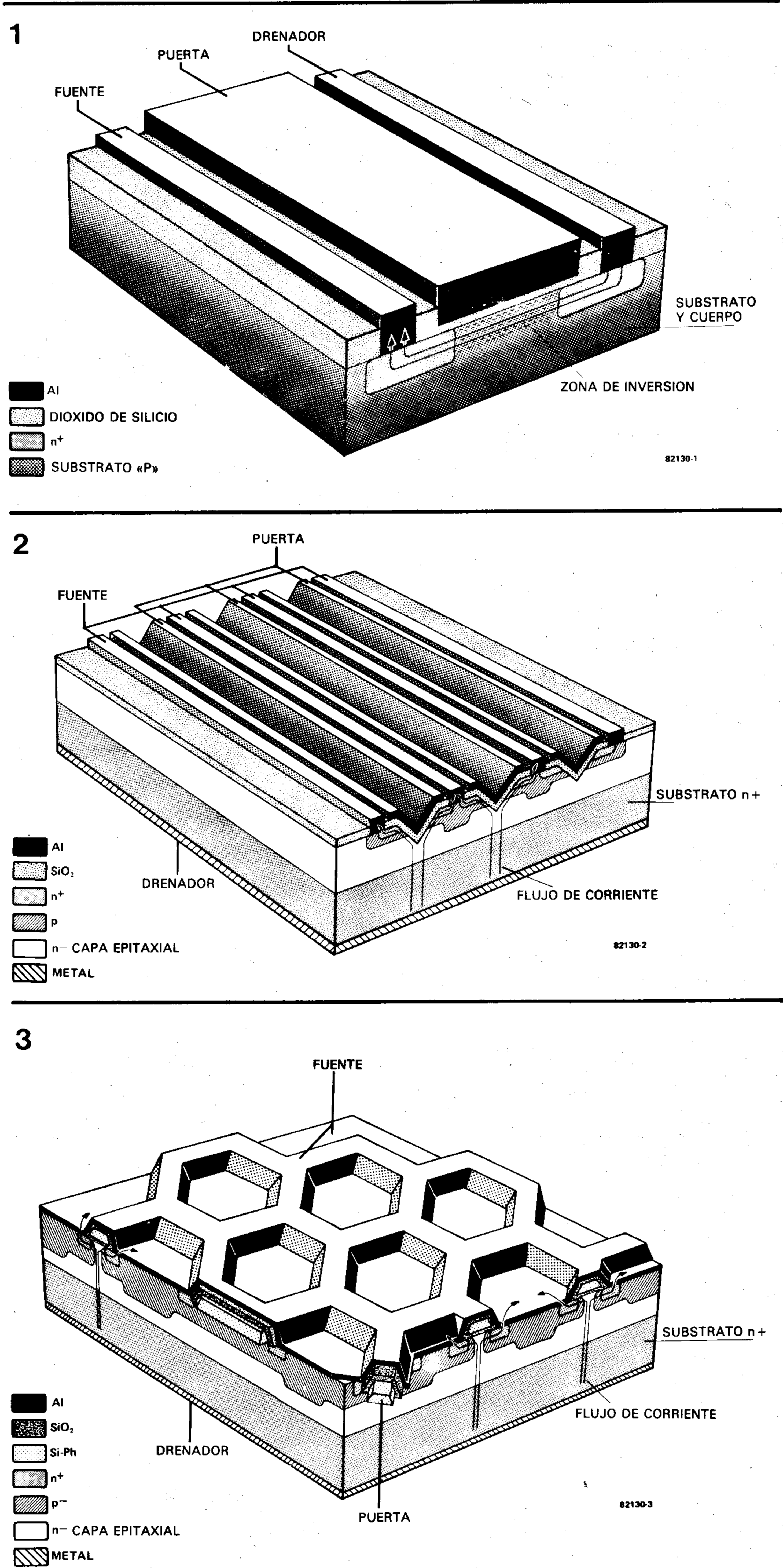
transistores FET de potencia actualmente disponibles. El término MOS es la abreviatura de «Metal Oxyde Semiconductor» (semiconductor de óxido metálico). Se comprende fácilmente de que se trata cuando se observa la estructura de un MOSFET, reproducida en la figura 1: un sustrato de silicio cristalino está recubierto de una capa de metal se obtiene el contacto con la superficie del sustrato. En la tecnología MOSFET, la capa de óxido es interrumpida en dos lugares, en los que se encuentran las conexiones del drenador y de la fuente. Debido a la capa metálica, en este caso de aluminio, la fuente y el drenador están conectados a una zona de conducción N, sobre un sustrato semiconductor tipo P. Ha lugar, pues, a hablar de una estructura NPN, como en el caso de los transistores bipolares. Puede representarse como dos diodos conectados en oposición (anti-serie), de manera que siempre uno de los dos esté bloqueado y que ninguna corriente circule entre el drenador y la fuente.

En un transistor MOSFET, la puerta está constituida por una superficie de aluminio, aislada del sustrato por una capa de óxido de silicio; asociada al sustrato como contraelectrodo, esta puerta se comporta como un condensador. Cuando la puerta está a un potencial más positivo que el sustrato, los electrones de la zona de conducción P del sustrato se desplazan hacia la zona de influencia del electrodo de puerta. La carga de los electrones es, como sabemos, negativa y ésta es la razón de que sean atraídos de esta forma por la puerta positiva. Debido a este enriquecimiento, la zona de conducción P tiene electrones «de sobra», al menos en las proximidades inmediatas de la capa de óxido, en una zona llamada de inversión, pasando de conducción P a conducción N. De este modo, se establece, entre el drenador y la fuente, un canal de conducción N, a través del cual circula una corriente, si se aplica una tensión entre el drenador y la fuente (se trata, pues, de la corriente de drenador del FET). En definitiva, se verifica que cuanto mayor es la tensión a través de la puerta, tanto más ancho será el canal y tanto más pequeña será la resistencia entre fuente y drenador.

Por este enriquecimiento en la zona de la puerta, a veces se le ha aplicado a este tipo de transistores el término de «modo de enriquecimiento» (enhancement mode); este tipo de FET se bloquea automáticamente en ausencia de la tensión de puerta. Pero los hay también que funcionan de manera inversa («modo de deplexión»), los cuales se bloquean cuando se les aplica una tensión de puerta. Este tipo de FETs no tiene interés para nuestro estudio y haremos caso omiso de los mismos.

V-MOS FET

Teóricamente, la diferencia entre el MOSFET y el V-MOS FET no es grande. Si se examina la figura 2 se constata, en efecto, que la estructura es fundamentalmente idéntica a la ilustrada en la figura 1; la fuente y el drenador están conectados a zonas de conducción N, separadas por una zona P. El hecho de que la conexión del drenador se caracterice por dos zonas de conducción N fuerte $n^+ = \text{fuerte}$, $n^- = \text{débil}$) carece de im-

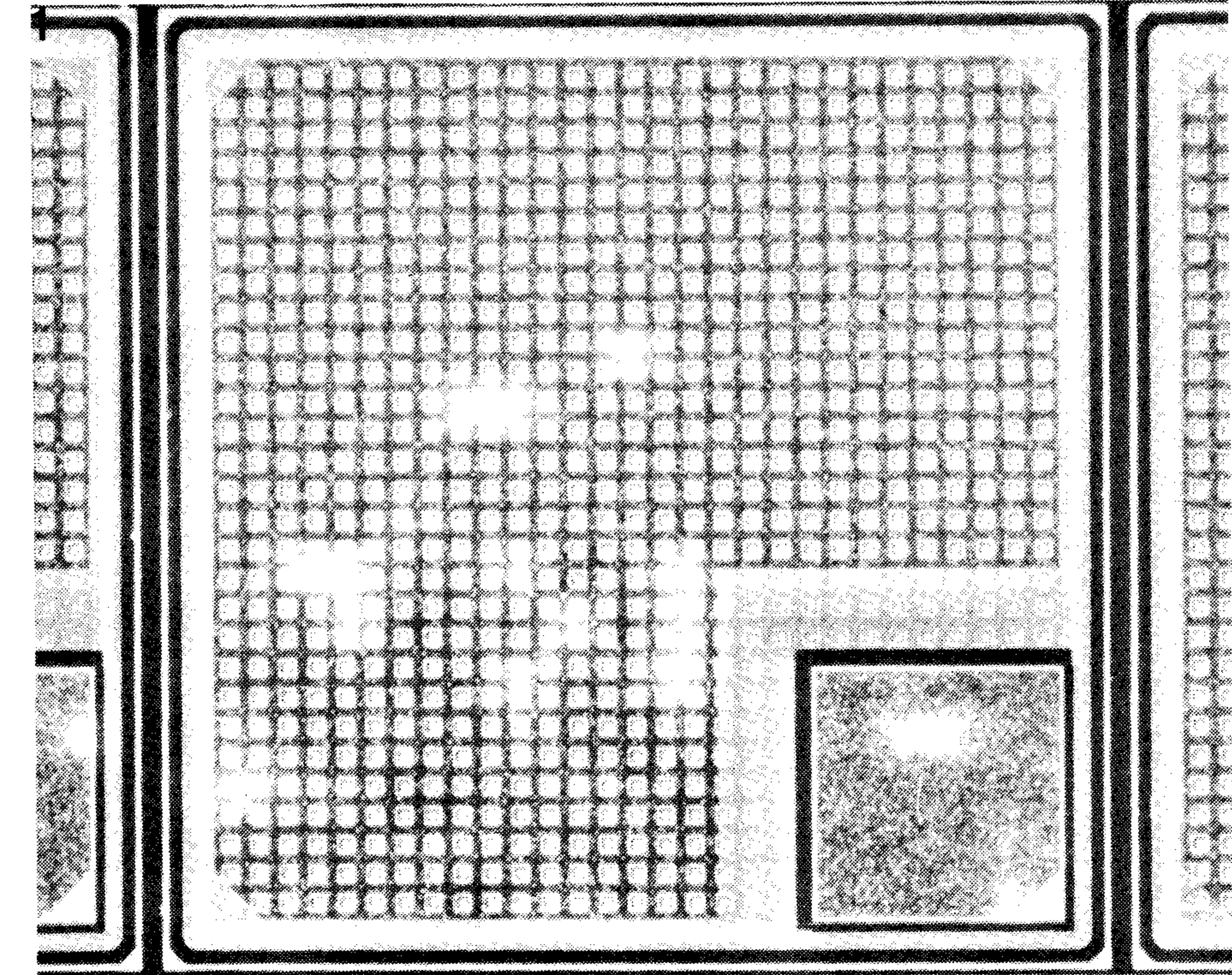


Figuras 1, 2 y 3. Estructuras esquematizadas de un MOSFET ordinario, de un V-MOSFET y de un D-MOSFET. En tanto que el transistor no sea conductor, se le puede considerar como un diodo bloqueado. Cuando sea conductor, desaparece esta función de «diodo» y sólo subsiste una pequeña resistencia óhmica.

portancia en cuanto al principio. Ocurre lo mismo para la conexión de la fuente que está, a su vez, unida a una zona N y a una zona P. Pero, en este caso, uno de los diodos «de explicación» de la figura 1 se ha hecho inútil; éste es el que conduce cuando la tensión de drenador es positiva. Aparte de estos detalles, todo sucede como con el MOSFET de la figura 1. Cuando la puerta se hace positiva, se forma un canal conductor en la zona P, por el cual puede fluir una corriente desde el drenador a la fuente. De aquí el nombre de V-MOS FET (VFET, en contracción). ¿Pero, por qué la «V»? Es la primera letra de «vertical», puesto que el sentido de la corriente a través del sustrato es vertical; no tiene nada que ver con la forma de V de la ranura en el sustrato. Una cuestión más profunda es saber por qué los V-MOS FET son más potentes que los MOS FET ordinarios. Técnicamente, el origen de esta peculiaridad no reside más que indirectamente en la estructura vertical; en efecto, el precio de coste de un semiconductor está ligado a las dimensiones, y más precisamente a la superficie de la pastilla utilizada. Ahora bien, si se obtuviera un FET de potencia con la tecnología MOS tradicional, esta superficie tendría dimensiones prohibitivas. Mientras que con la estructura FET-V se economiza toda la superficie requerida por la conexión del drenador, colocándola pura y simplemente en la parte inferior de la pastilla. Asimismo, los canales están formados por difusión, lo que permite que el VFET trabaje a niveles de tolerancias mucho menos estrictos, lo que, a su vez, facilita la miniaturización.

El transistor FET de potencia está, de hecho, constituido por algunos millares de minitransistores FET conectados en paralelo en una misma pastilla. La potencia no es, pues, solamente una característica dependiente de la estructura VFET, sino también del número de estructuras colocadas en la misma pastilla.

Después de esta introducción, no debería ser difícil comprender en qué consiste la peculiaridad de un verdadero D-MOS FET. La vista en sección de la figura 3 ilustra muy



Fotografía 1. Un DFET está constituido por un gran número de estructuras FET montadas en paralelo. La superficie que aparece arriba a la izquierda es la conexión de la puerta, mientras que el resto de la superficie metalizada de la pastilla corresponde a la fuente.

claramente la estructura de la pastilla. En este caso, la puerta está completamente revestida por la capa de Si O₂ aislante, mientras que la fuente cubre toda la superficie de la pastilla en forma de una capa de aluminio. Mientras que en el VFET la puerta adoptaba la forma de una depresión, en el caso del DFET (D-MOS FET) es protuberante. En la fotografía 1, se reconoce la puerta bajo el aspecto de un pequeño cuadrado, en cuyas proximidades se perciben otras estructuras geométricas como, por ejemplo, los exágonos de los HEX-FET que son posibles de acuerdo con las preferencias de cada fabricante particular.

Lo expuesto se trata de generalidades, comunes a todos los DFET que, a veces, presentan peculiaridades según las aplicaciones a los que estén destinados. La estructura D-MOS, como acabamos de describir, tiene el inconveniente de tener una resistencia interna de puerta, asociada a una capacidad relativamente grande (de varios nanofaradios). No hay que excluir la posibilidad de que cuando se apliquen a la puerta señales del orden del MHz, se caliente hasta un punto en que el FET se haga inutilizable. A este respecto, el VFET, con su puerta de aluminio de baja impedancia, resulta ventajoso. Es, sin duda, la razón por la que los

4

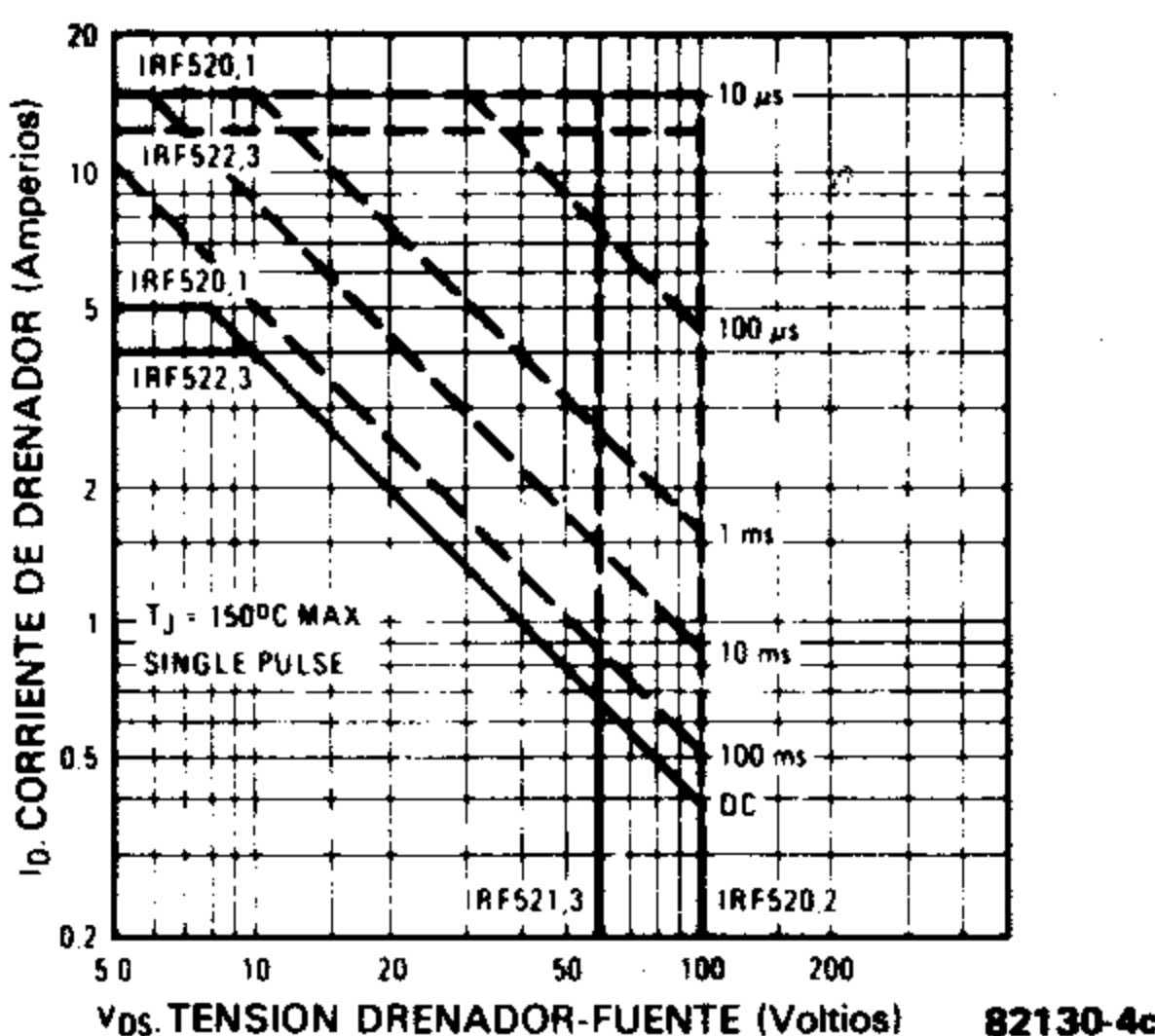
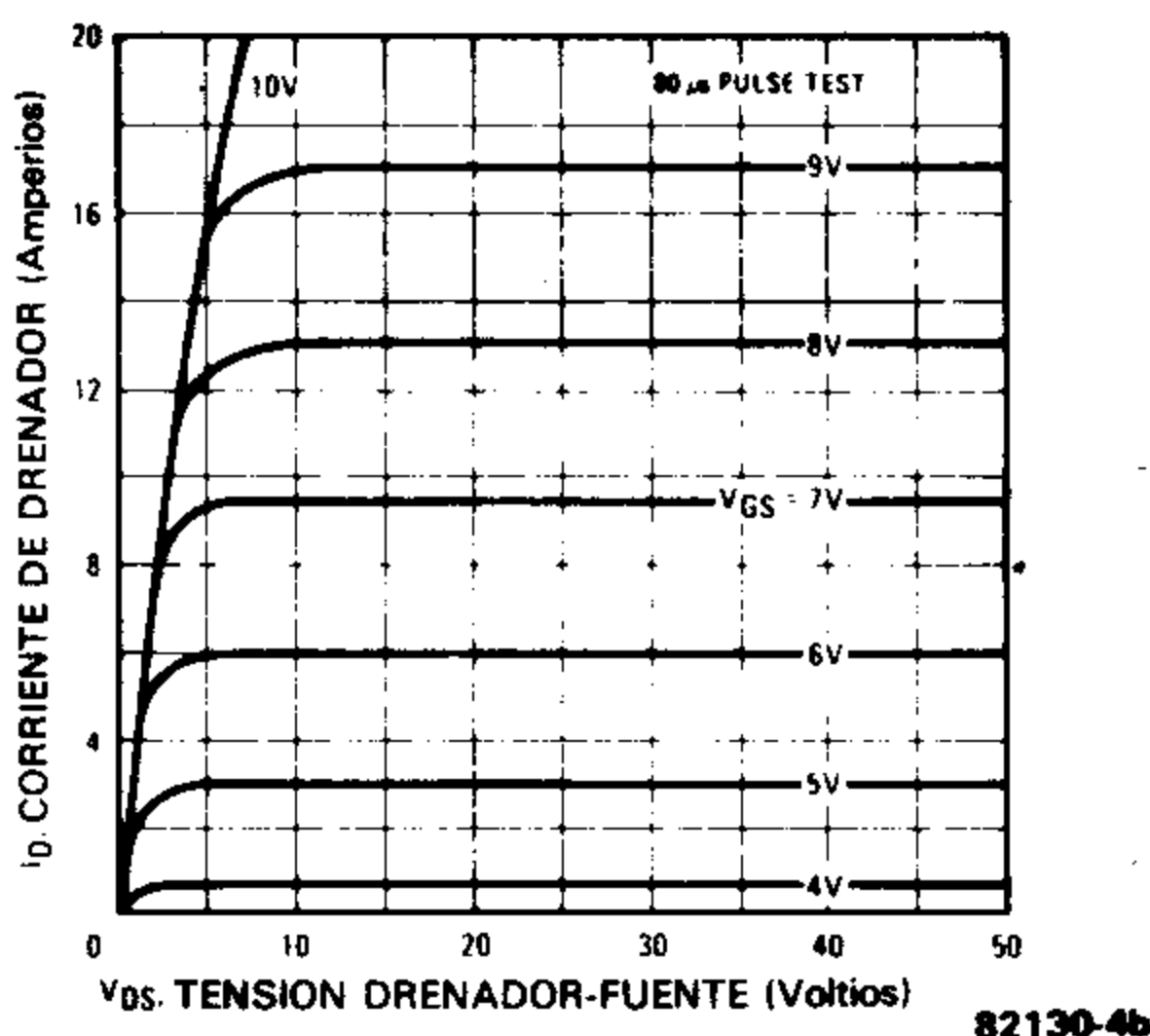
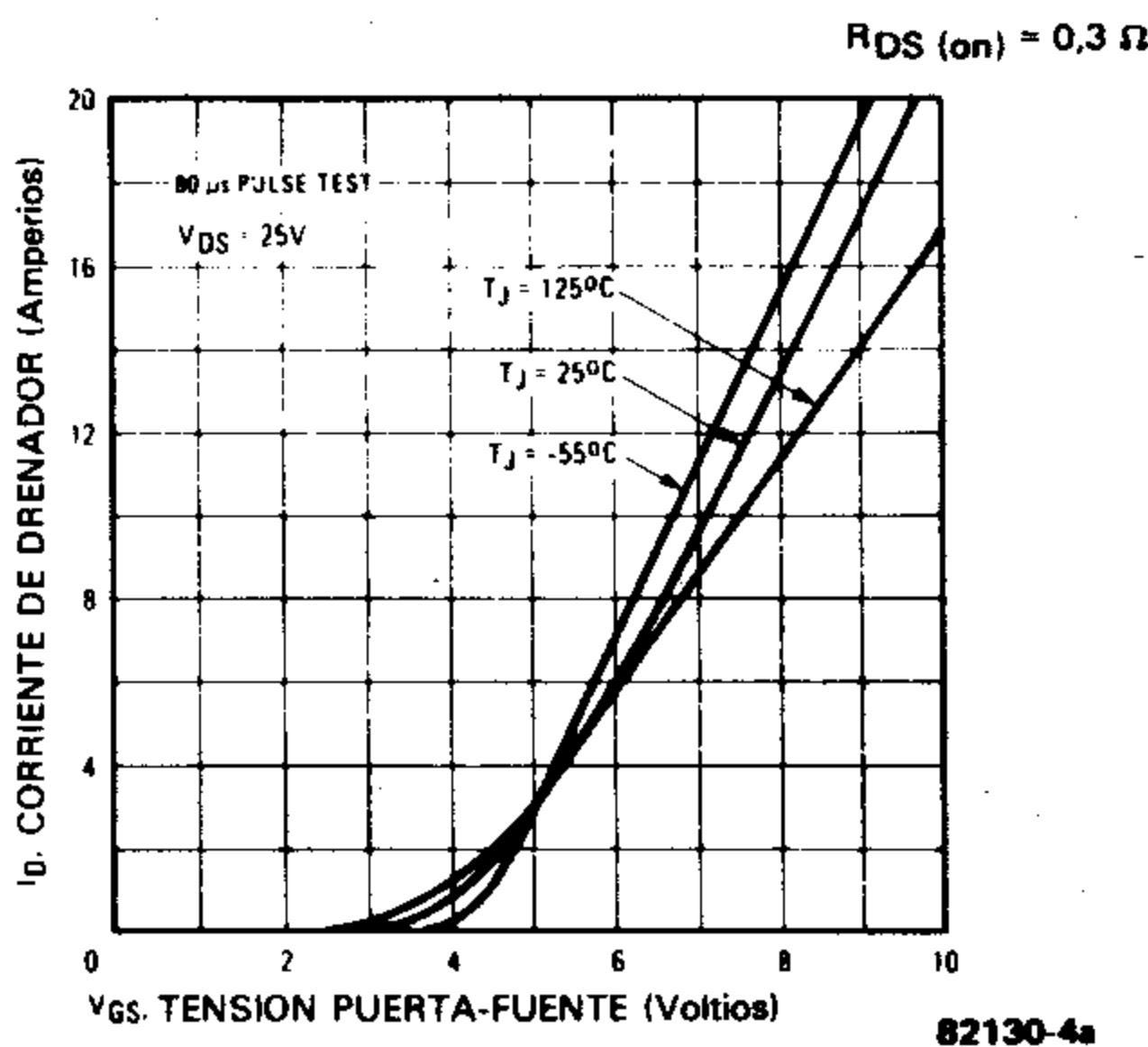


Figura 4. Características de un MOSFET típico (los demás miembros de la familia presentan curvas muy semejantes).

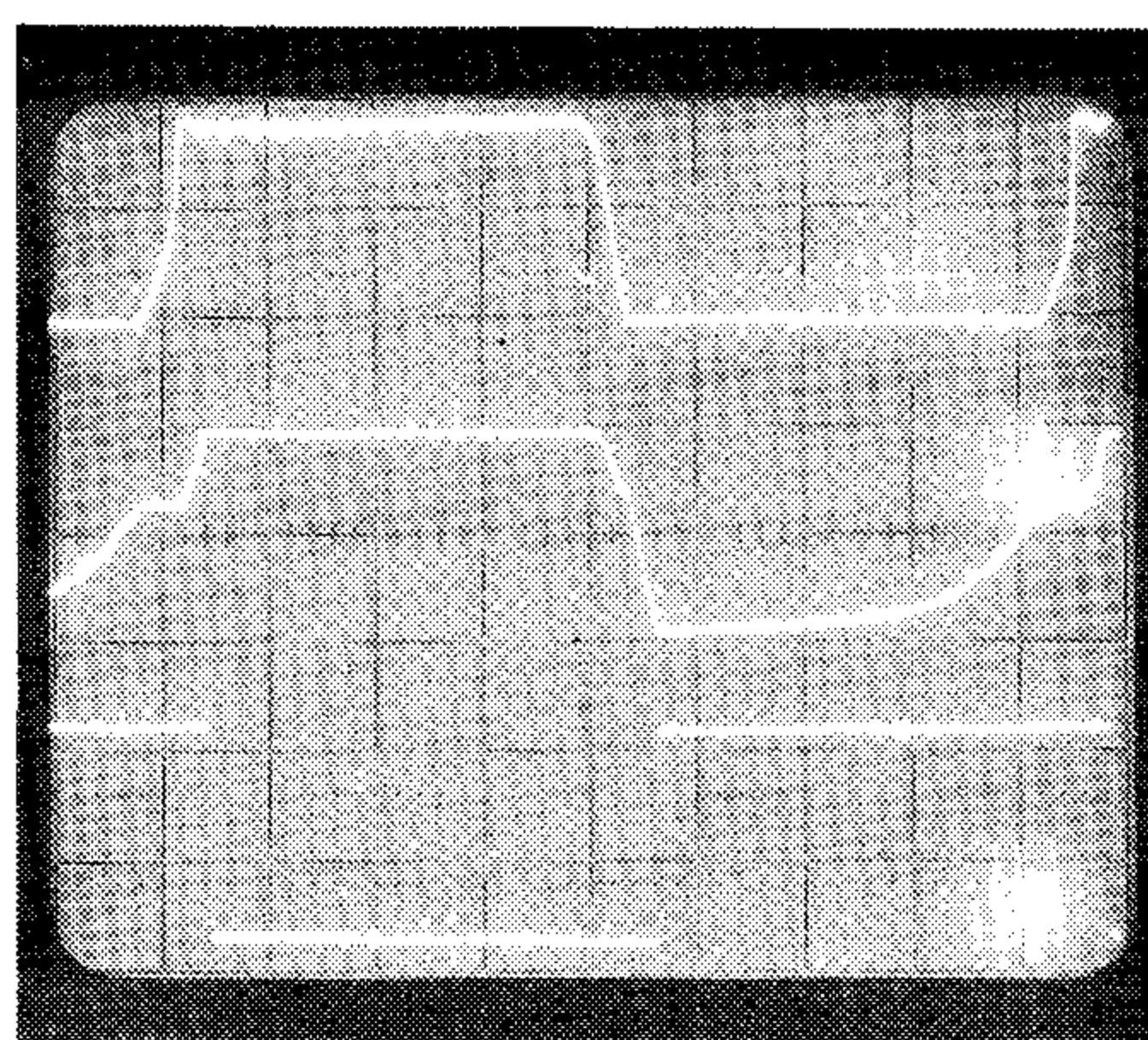
DFET son esencialmente utilizados para la conmutación de potencia, antes que para la amplificación en el campo de las altas frecuencias. Sin embargo, el DFET tiene también ventajas sobre el VFET, cuya estructura en forma de V da lugar a campos eléctricos importantes en el fondo de la garganta, en donde los fenómenos de difusión y de ionización son casi inevitables. Esta vez, la ventaja radica en la «pseudo» estructura DMOS.

Propiedades

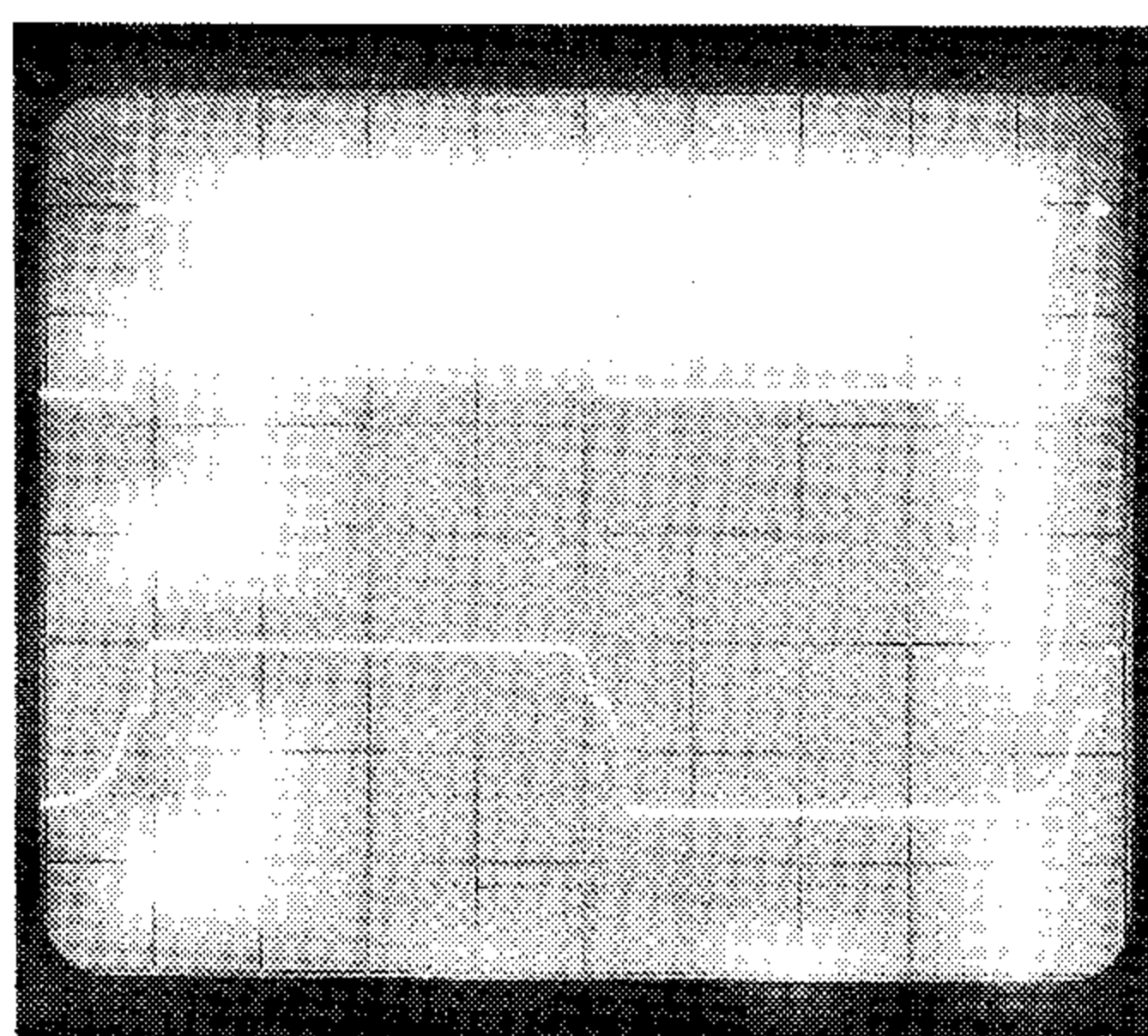
¿Qué se puede conseguir con los nuevos FETs de potencia? En primer lugar, una disipación de potencia máxima comparable con la correspondiente a los transistores de potencia bipolares montados en cápsulas idénticas. Existen FETs que se adaptan a tensiones que pueden alcanzar 1.000 V y otros que pueden conmutar hasta 25 A. Al igual que los transistores ordinarios, soportan corrientes máximas todavía más elevadas, con tal de que sean breves.

Más importante que la corriente máxima indicada por el constructor, merece una especial atención al valor $R_{DS(on)}$, que es la resistencia entre el drenador y la fuente cuando el FET es conductor. Esta es la resistencia que, en la aplicación, determina, en definitiva, la magnitud de la corriente. Los límites varían en función de la disipación de potencia máxima (que no ha de sobrepasarse) y la caída de tensión a través del FET conductor, que se hace prohibitiva cuando la tensión de alimentación y/o la resistencia de carga sean demasiado débiles.

La amplificación del FET no puede expresarse, en este caso, en función de la amplificación de corriente, puesto que se trata de un control en tensión del drenador. La transconductancia constituye, por el contrario, una medida de la amplificación y es el orden de algunos amperios por voltio con una tensión umbral de unos voltios. En la figura 4 se da un ejemplo ilustrativo. Debido al control en tensión, no hay, teóricamente, ninguna corriente de entrada y la amplificación de potencia sería ideal (al ser infinita). En la práctica, no todo sucede de esta for-



Fotografía 2. Cuando se excita un FET de potencia con una puerta CMOS, el tiempo de conmutación se alarga considerablemente, dado que el excitador no puede producir corriente suficiente para una transferencia capacitiva rápida de la puerta del FET. De arriba abajo: señal de control a la entrada de la puerta CMOS, señal en la puerta del FET y curva de transferencia del FET.



Fotografía 3. Mejores tiempos de conmutación se obtienen excitando los FETs de potencia a partir de un buffer con colector abierto TTL. La tensión de control de la puerta puede seleccionarse mejor al doble del nivel de la tensión de puerta requerida para obtener la corriente de drenador deseada. Más allá de este valor, la tensión de puerta ya no aportará ninguna aceleración de la conmutación, sino que, por el contrario, impondrá una potencia de excitación desproporcionada (con el cuadrado de la tensión).

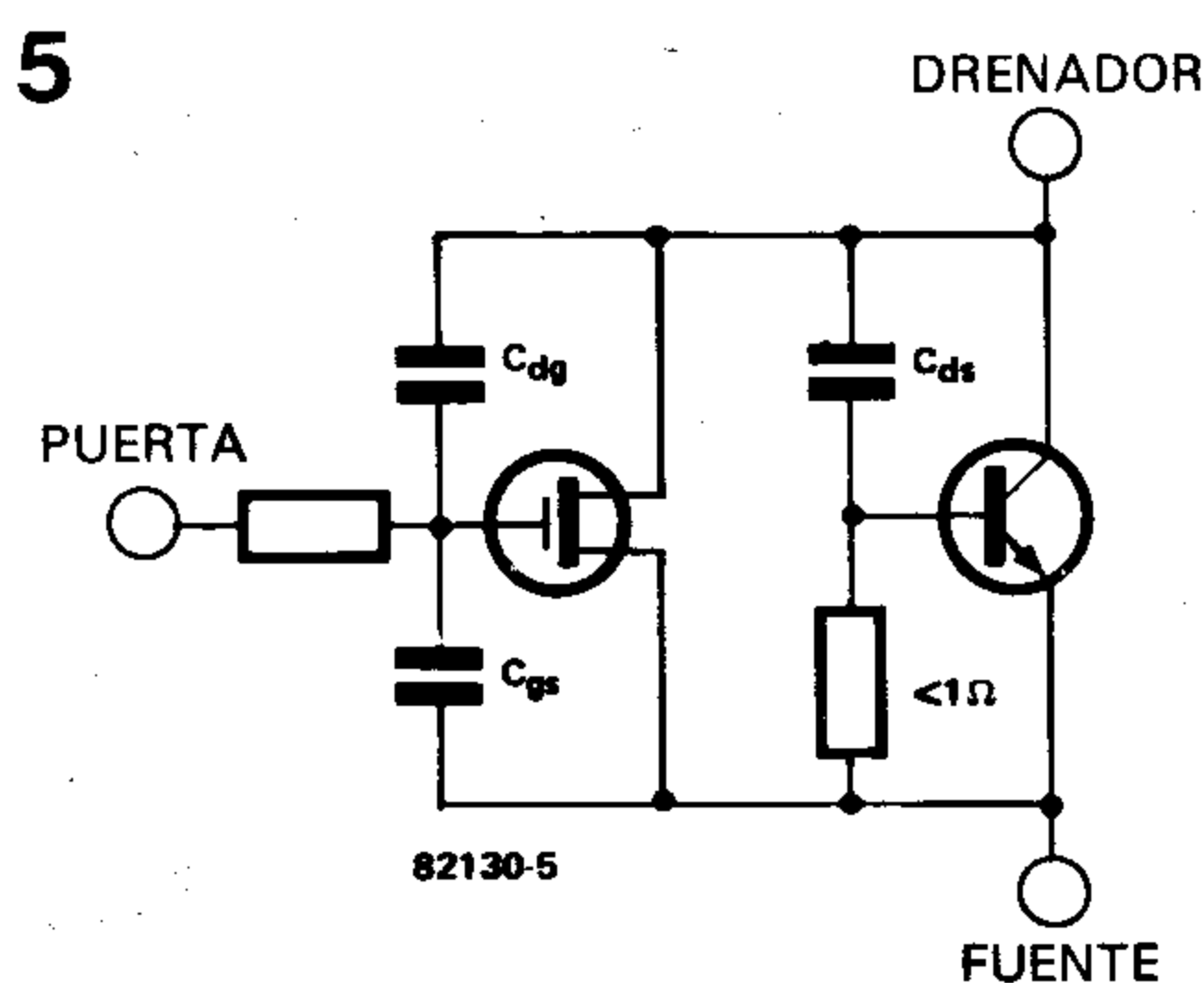
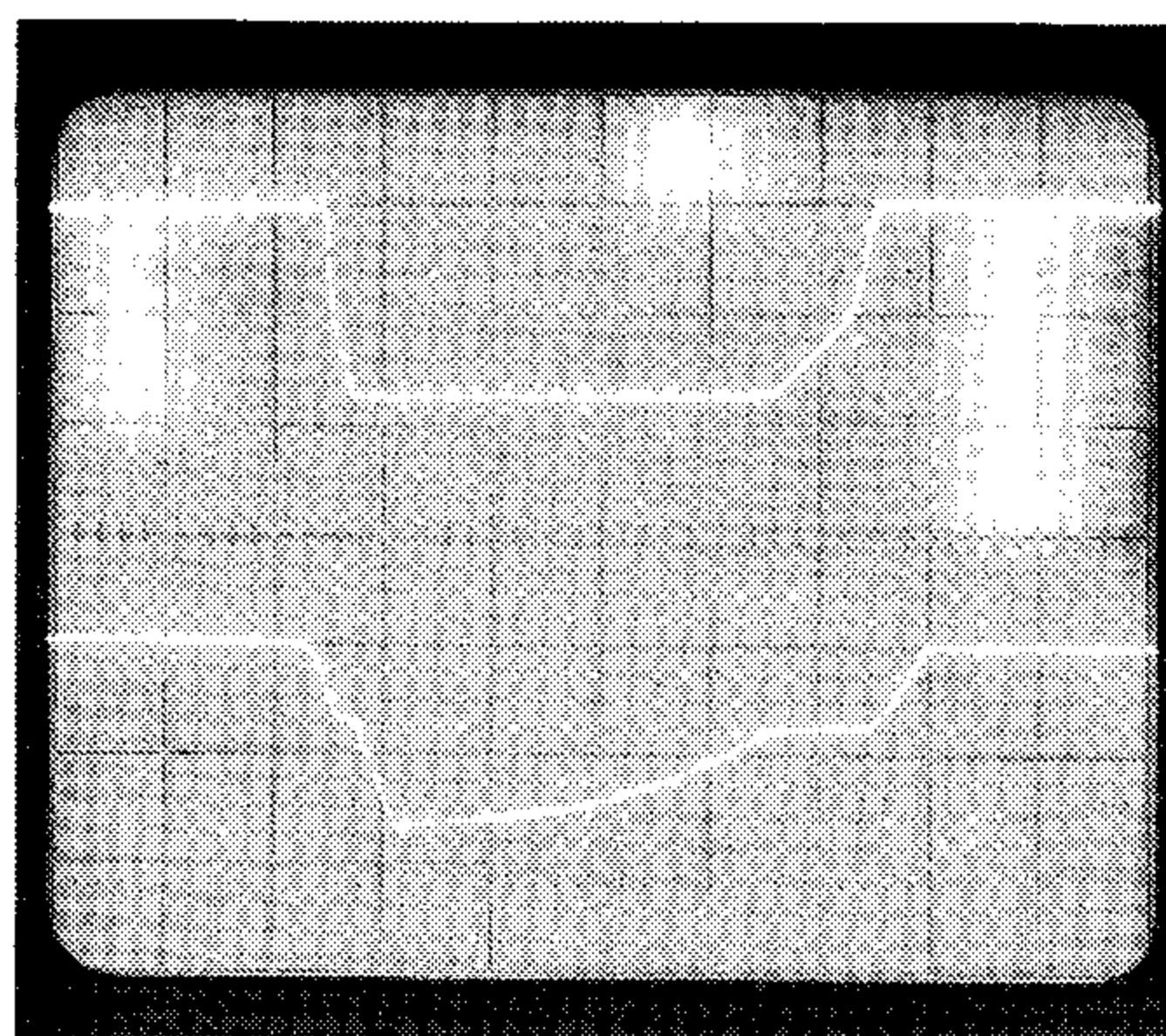


Figura 5. Circuito equivalente simplificado de un D-MOS (D-MOSFET o D-FET) polarización óptima, cuando se excita el transistor, es preciso tener en cuenta la capacidad entre drenador y fuente. Un transistor está conectado en paralelo con el FET y actúa como un diodo (a través de la resistencia de 1 ohmio y la unión base-colector) en el caso de tensiones de drenador negativas. El diodo puede dejar pasar la misma cantidad de corriente que el DFET, aunque, por supuesto, con mucha más «lentitud».



Fotografía 4. Curva de la tensión de puerta durante la conmutación. Es muy difícil prever exactamente el retardo de conmutación, habida cuenta de que la capacidad de puerta efectiva y momentánea depende del valor de la tensión de drenador. Lo mismo se aplica a las exigencias de potencia de excitación y ésta es la razón por la que algunos fabricantes proporcionan gráficos de la tensión de puerta para diversas tensiones de drenador.

ma. Lamentablemente, en el momento de la conmutación, se precisa en el FET una determinada potencia de control para la transferencia de la carga capacitiva (en nanofaradios) de la puerta. Si el control se efectúa a alta impedancia, la transferencia durará demasiado. Si se realiza en buenas condiciones, el FET es aceptablemente rápido: un FET de potencia normal conmuta una corriente de algunos amperios en menos de 100 microsegundos, a condición de que la señal de control sea perfectamente cuadrada y de que proceda de una fuente de baja impedancia de salida. En la práctica, esta condición casi nunca se cumple y la fotografía 2 da una idea «algo exagerada» de lo que ocurre realmente. La señal en onda cuadrada de la parte superior es la adecuada y pasa a través de un inversor CMOS del tipo 4049 para excitar directamente la puerta de un transistor D-MOS del tipo BUZ 10. Los flancos de la señal dejan bastante que desear y tienden a formar escalones «a modo de espinazo», a medio camino del mismo. La curva inferior es la correspondiente a la corriente de drenador del FET.

El inversor CMOS tarda cierto tiempo en modificar la tensión en la puerta del FET, pues la capacidad de la puerta sólo puede transferirse con un par de miliamperios. Como el 4049 es un buffer de adaptación a los niveles TTL, puede proporcionar más corriente hacia masa que hacia el potencial positivo de la alimentación. No es sorprendente, en efecto, que el flanco de bajada sea mucho más «escarpado» que el flanco de subida. Queda aún por diagnosticar el origen de los escalones de nuestra señal.

Una vez más, la causa radica en la capacidad puerta-drenador. En la figura 5 se da el esquema de un circuito equivalente simplificado de un DFET. Los especialistas en válvulas de vacío (¡si es que existen todavía!) pensarán en el efecto Miller. La elevación del potencial de la puerta tiene, como consecuencia, una caída de tensión en el drenador que, en virtud de la capacidad puerta-drenador, afecta retroactivamente a la puerta, cuyo potencial es frenado en su curso ascendente. Esta situación continúa hasta que la tensión de drenador ya no pueda disminuir más. El efecto es claramente visible en la fotografía 2, en donde la tensión de puerta es relativamente constante mientras se altera la tensión de drenador. Además, hay casi siempre una cierta cantidad de inductancia en la conexión de la fuente y ello refuerza el efecto, haciendo algo negativa la fuente. A una tensión de alimentación más alta, se acentúa el efecto y la capacidad puerta-drenador tarda más en transferirse. Y reiteramos aquí la afirmación hecha al principio de nuestros comentarios: el circuito de control desempeña un papel determinante en la velocidad de conmutación del FET de potencia. Los factores influyentes son: la tensión drenador-fuente (cuanto más elevada sea, tanto más lenta será la conmutación), la capacidad de puerta (que depende del FET) y las características del circuito de control o excitador (que depende del usuario). En la fotografía 3 se muestra un FET activado por un TTL, que es mucho más rápido. Sin embargo, la conmutación a alta velocidad trae consigo algunos inconvenientes: cuando el FET interrumpe el flujo de una corriente de algunos amperios en unos nanosegundos, basta una mínima autoinduc-