

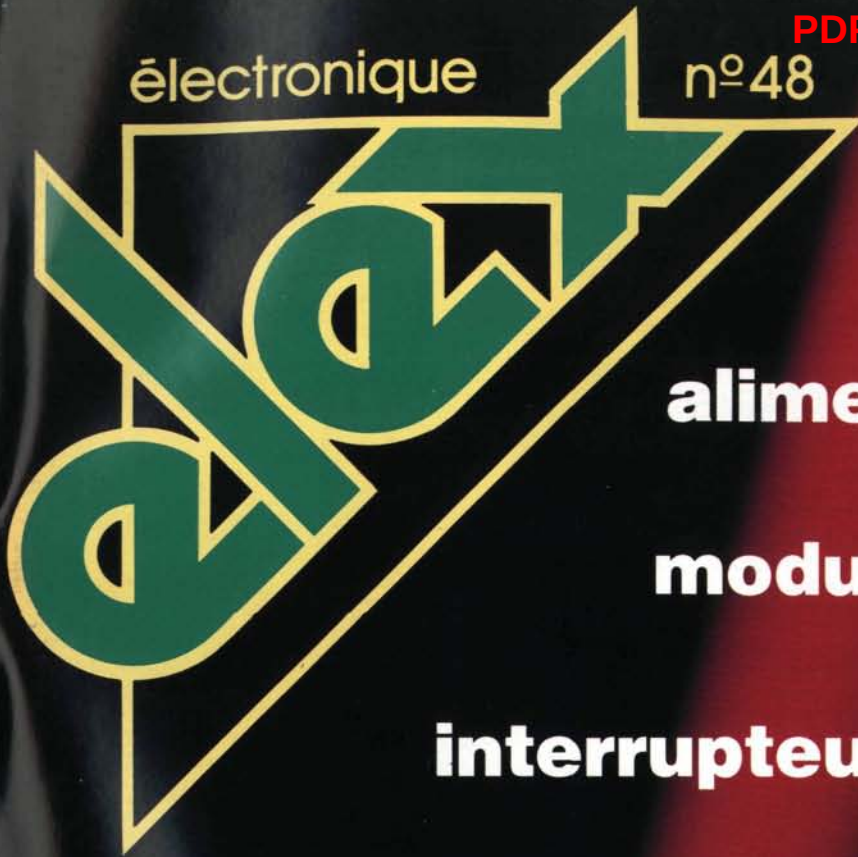
électronique

n°48

octobre 1992

23 F/168 FB/7,80 FS

mensuel



alimentation d'école
avec circuit imprimé

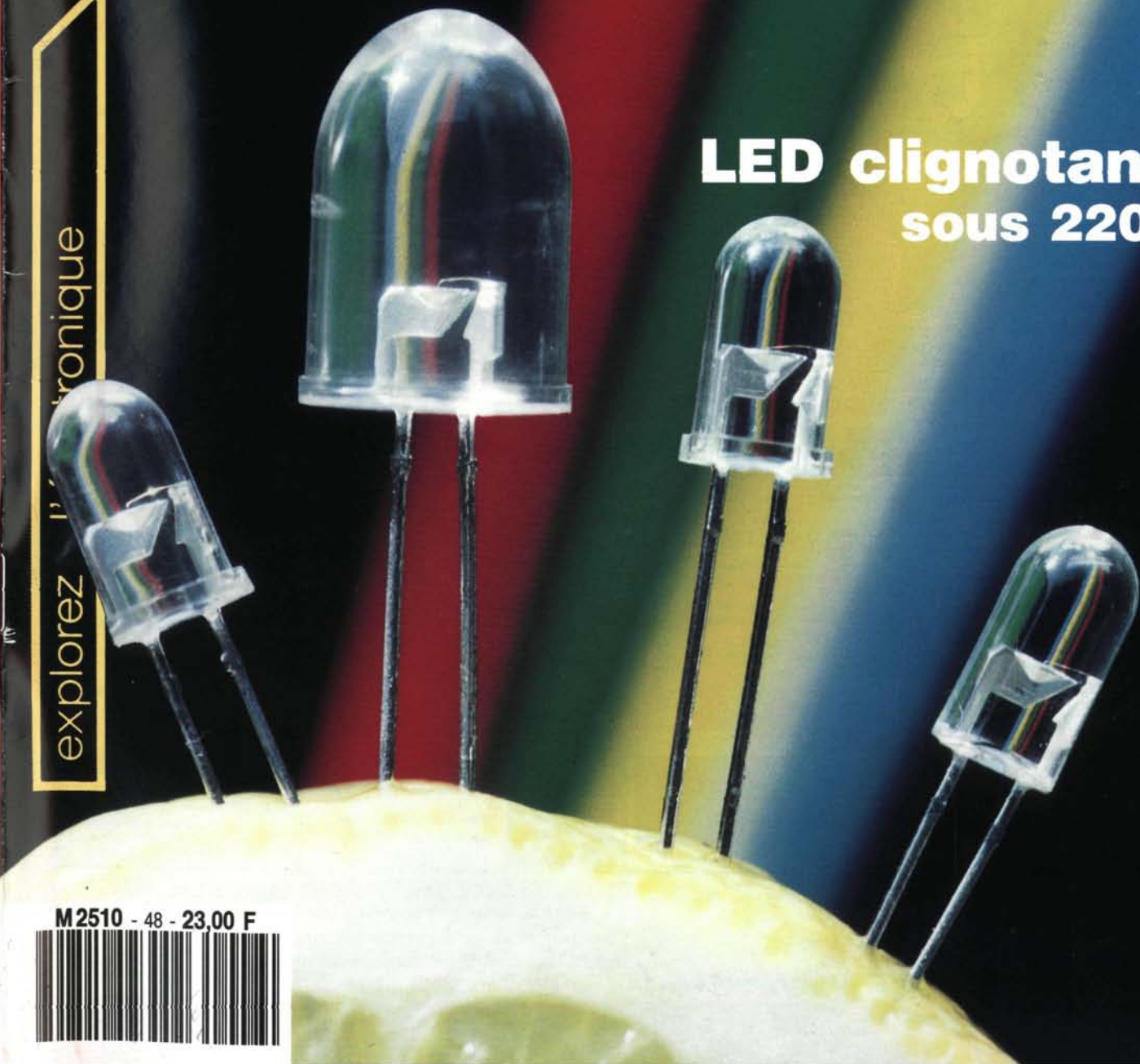
module capacimètre
avec circuit imprimé

interrupteur crépusculaire
pour lampes fluo
avec circuit imprimé

LED clignotante
sous 220 V

explorez l'électronique

M2510 - 48 - 23,00 F



SOMMAIRE ELEX N°48

60 ➡ mots croisés

60 ➡ petites annonces gratuites

I . N . I . T . I . A . T . I . O . N

4 ➡ Rési & Transi : bande dessinée

8 ➡ les lampes fluorescentes

16 ➡ les moteurs électriques

28 ➡ système K : combinaisons

50 ➡ la doc ad hoc :

le régulateur de tension 723

R . É . A . L . I . S . A . T . I . O . N . S

13 ➡ une machine à rêves
(générateur de bruit de vagues)

20 ➡ un module capacimètre

26 ➡ un interrupteur crépusculaire
pour lampes fluorescentes

32 ➡ un détecteur de fin de course

34 ➡ un dé électronique
pour faire plus ample connaissance avec le circuit intégré 4029

37 ➡ un compte-pose pour agrandisseur

42 ➡ un stéthoscope audio

43 ➡ une alimentation d'école autour du 723

52 ➡ un ionomètre

55 ➡ clignotant à LED sous 220 V

avec
circuits
imprimés

Annonceurs : ARQUIÉ COMPOSANTS p. 49 - BÉRIC p. 48 - B.H. ÉLECTRONIQUE p. 49 - CENTRAD p. 63 - CIF p. 61 et 62 -
 COMPOSANTS DIFFUSION p. 49 - COMPOSIUM p. 49 - ELC p. 63 - ELECTRON SHOP p. 49 - EURO COMPOSANTS p. 49 -
 EUROTECHNIQUE p. 7 - EXPOTRONIC p. 23 et p. 59 - HB COMPOSANTS p. 49 - J. REBOUL p. 48 - LAYO FRANCE p. 49 -
 LOISIRS ELECTRONIQUES p. 48 - MAGNÉTIQUE FRANCE p. 31 - MICROPROCESSOR p. 49 - PUBLITRONIC pp. 6, 23, 61 et 62 -
 SÉLECTRONIC pp. 2, 61, 62 et 64 - SPESYS p. 49 - SAINT-QUENTIN RADIO p. 49 - SVE ELECTRONIC p. 49 - TSME p. 49 -
 URS MEYER p. 49

lampes fluorescentes

éclairage économique

Le titre pourrait laisser penser que nous allons traiter d'un nouveau type de lampe. Il n'en est rien. Les lampes dont nous allons parler sont les tubes fluorescents, qui éclairaient peut-être votre lecture (tubes dits TL), et les lampes, coûteuses à l'achat, mais moins gourmandes en énergie et d'une plus grande longévité que celles à incandescence, connues sous les noms de PL et SL. Comme ce numéro décrit un montage à réaliser autour d'une lampe SL, il était bon de regarder de plus près comment elle fonctionne.

Les tubes fluorescents, les tubes au néon⁽¹⁾, les lampes à vapeur de sodium ou à vapeur de mercure à haute pression appartiennent à la même famille. Pour toutes ces lampes, la lumière est due, directement ou indirectement, à une décharge électrique dans un gaz, d'où leur nom de "lampes à décharge".

Qu'est-ce qu'une décharge électrique dans un gaz et pourquoi cela produit-il de la lumière ? Voyons sur la **figure 1** comment est constitué un tube fluorescent ordinaire (tube TL). C'est un tube de verre (7) aux extrémités duquel se trouvent des filaments (1) portés à incandescence, les électrodes. Le tube lui-même contient un mélange de gaz, vapeur de mercure et argon entre autres, et il est intérieurement recouvert d'une couche de poudre fluorescente.

Supposons pour commencer que la lampe éclaire, nous verrons ensuite comment l'allumer. Les électrodes sont chauffées comme le filament d'une lampe à incandescence. Plus on chauffe un métal, moins les électrons à sa surface sont retenus : ils tendent à "s'évaporer". S'ils sont soumis à un champ électrique, ils "s'évaporent" dans la direction (une "direction" suppose deux sens) du champ, don-

- | | | |
|--|-------------------|----------------|
| 1 filament incandescent | 2 électrons | 3 atome de gaz |
| 4 radiations ultraviolettes nées des collisions | 5 lumière visible | |
| 6 poudre fluorescente produisant, sous l'influence du rayonnement ultraviolet, de la lumière visible | 7 tube de verre | |

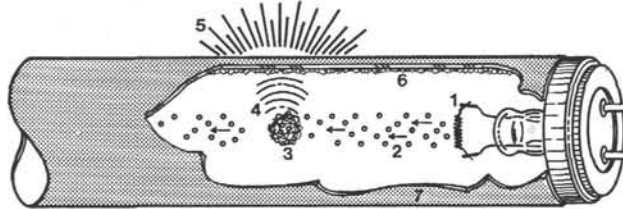


Figure 1 - Intérieur d'un tube fluorescent.

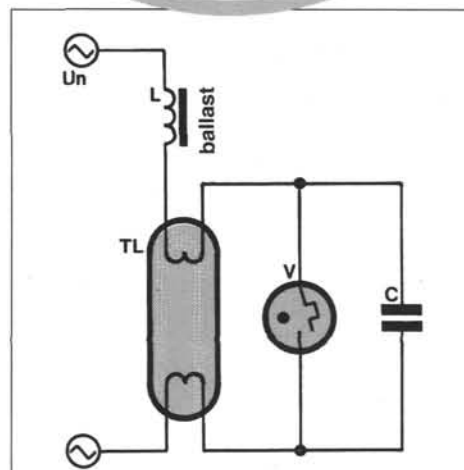
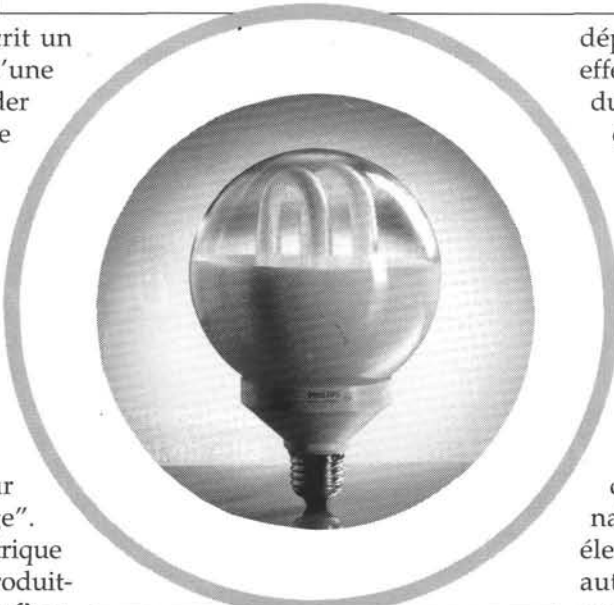


Figure 2 - Une bobine (le ballast), un démarreur et son condensateur d'anti-parasitage, et naturellement un tube fluorescent, sont les constituants de ce que l'on appelle improprement un "tube au néon".ballast

nant naissance à un courant (2). Comme la voie n'est pas libre puisque le tube dans lequel ils circulent contient des gaz, les électrons rencontrent des atomes sur leur chemin ce qui donne lieu à des transferts d'énergie cinétique : c'est du billard. De ces transferts d'énergie résultent trois types de réaction : heurtés par les électrons les atomes se

déplacent plus vite, ce qui n'a pour effet que d'augmenter la température du gaz (la chaleur est un mouvement désordonné). Ce chauffage se fait en pure perte, puisque le but de l'opération est d'éclairer. Le deuxième effet est plus intéressant : sous l'effet de la collision, les atomes perdent des électrons qui, ainsi libérés, participent à la course et augmentent le courant. Tout se passe comme s'il y avait une avalanche d'électrons. Le troisième effet est celui que nous attendons : des collisions naît la lumière. Dans un atome, les électrons ont des orbites stables autour du noyau, caractérisées par une certaine énergie. Pour qu'un électron change d'orbite et s'éloigne du noyau, il faut lui fournir de l'énergie. On dit alors que l'atome est excité. En rencontrant les atomes, les électrons libres qui circulent dans le tube les mettent quelquefois dans cet état instable d'excitation. Cet état ne dure pas et, au bout d'un certain temps, un électron déplacé retombe sur son orbite d'origine en libérant, sous la forme d'un photon (grain de lumière), l'énergie qui lui avait permis de s'éloigner. Les photons ont une longueur d'onde, caractéristique de l'atome dont ils sont issus ; dans le cas des atomes de mercure, dont est en partie constitué le gaz des tubes fluorescents⁽²⁾, ce sont des photons ultraviolets. Nous sommes bien avancés puisque le rayonnement ultraviolet est non seulement invisible mais en

(1) Les tubes au néon sont des tubes luminescents qui donnent une lumière rouge. Ils sont surtout utilisés pour les enseignes publicitaires ou la décoration.

plus dangereux⁽³⁾. C'est alors que le revêtement pulvérulent, qui opacifie de l'intérieur le tube, intervient : il absorbe les photons ultraviolets et libère en échange des photons de lumière visible.

amorçage d'une lampe à décharge

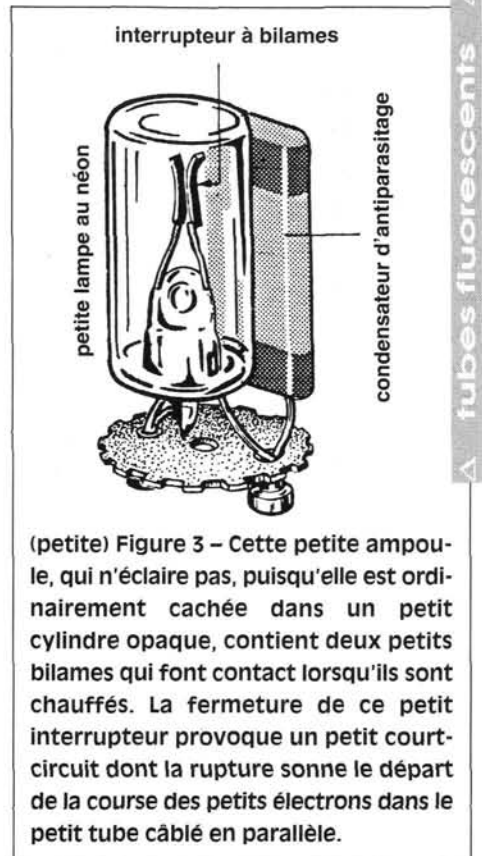
Tout le monde sait ça : un tube fluorescent ne s'allume pas immédiatement après sa mise sous tension. Un certain nombre de conditions doivent être remplies pour mettre en route le courant d'électrons. Il faut pour commencer les arracher à leur support : la tension entre les électrodes situées à chaque extrémité du tube doit donc être suffisante (compte tenu de leur distance) pour que le champ électrique leur permette de se libérer. Cette condition est remplie plus tôt si l'on chauffe les filaments. L'obtention de la tension relativement élevée de démarrage est due à un ballast, en fait une bobine, placé en série avec le tube et "commandé" par un démarreur⁽⁴⁾ en parallèle (figure 2). Voyons ce qu'est le starter.

Les démolisseurs invétérés, nombreux parmi les lecteurs (et les rédacteurs) de la revue savent que l'enveloppe de plastique du starter contient un ampoule de verre pourvue de deux connecteurs, qui ressemble fort à une lampe, et un condensateur d'antiparasitage en parallèle. Une observation plus attentive montre que chaque connecteur est relié à une plaquette métallique. Ce que l'on ne voit pas est que l'ampoule de verre contient du néon, un gaz plus rare dans les tubes que dans le langage.

Les deux plaquettes sont des bilames, constituées donc chacune de deux

métaux dont les coefficients de dilatation sont différents. Ils sont disposés de façon à se rapprocher l'un de l'autre, jusqu'à se toucher, s'ils sont chauffés. Le chauffage des bilames intervient dès la mise sous tension : cette sorte de tube au néon éclaire, puisqu'il y a décharge, et chauffe. Dès que les bilames sont assez chauds, ils font contact et court-circuit, si bien qu'un courant important circule. Comme on le voit sur la figure 2 le courant traverse le ballast et les deux filaments du tube qui chauffent à leur tour. Cependant, dès que les bilames ont établi le contact, la décharge et la production de chaleur cessent dans le petit tube au néon d'amorçage. Les deux bilames se refroidissent et finissent par couper le contact. Le courant est interrompu brutalement dans le ballast et les filaments. Comme le ballast est une bobine, son comportement à la coupure du courant est celui, assez singulier, des bobines. La figure 4 décrit précisément ce qui se passe alors. Récapitulons à l'aide de celle-ci.

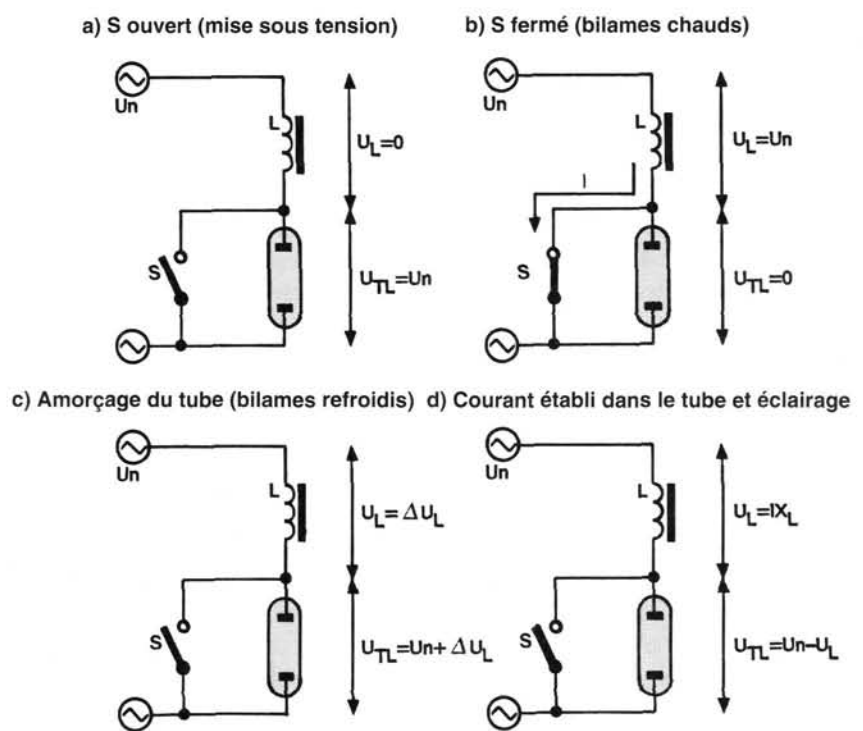
À la mise sous tension, lorsque la lampe n'est pas encore allumée et que les bilames de l'ampoule d'amorçage n'établissent pas encore le contact, le courant qui circule est minuscule. La tension aux bornes du tube est celle du secteur puisque la chute de tension dans le ballast est pratiquement



(petite) Figure 3 – Cette petite ampoule, qui n'éclaire pas, puisqu'elle est ordinairement cachée dans un petit cylindre opaque, contient deux petits bilames qui font contact lorsqu'ils sont chauffés. La fermeture de ce petit interrupteur provoque un petit court-circuit dont la rupture sonne le départ de la course des petits électrons dans le petit tube câblé en parallèle.

nulle (4a). Ensuite, les bilames se touchent et le circuit est fermé : le courant qui circule est maintenant très grand et, comme le tube est court-circuité, le ballast a toute la tension du secteur à ses bornes (4b). Dès que les bilames se refroidissent et ouvrent le contact, le courant est coupé brutalement : une force électromotrice importante, pro-

Figure 4 – Allumage d'une lampe à incandescence en quatre étapes.



(2) Le mercure est dangereux et il est bon de ne pas casser les tubes, même lorsqu'ils sont usagés.

(3) Les radiations ultraviolettes s'étagent entre le violet visible et les rayons X mous, entre 400 nm et 10 nm de longueur d'onde. Le soleil en produit beaucoup qui ne nous parviennent heureusement pas puisqu'elles sont bloquées dans la haute atmosphère par la fameuse couche d'ozone.

(4) On l'appelle starter pour éviter aux gens de prendre leurs tubes fluorescents pour des véhicules. Pour les mêmes raisons, les techniciens parlent de self parce qu'auto-inductance s'abrégierait en "auto"... Que ces stupidités ne vous empêchent pas de comparer le démarrage d'une voiture à l'amorçage d'une lampe à décharge.

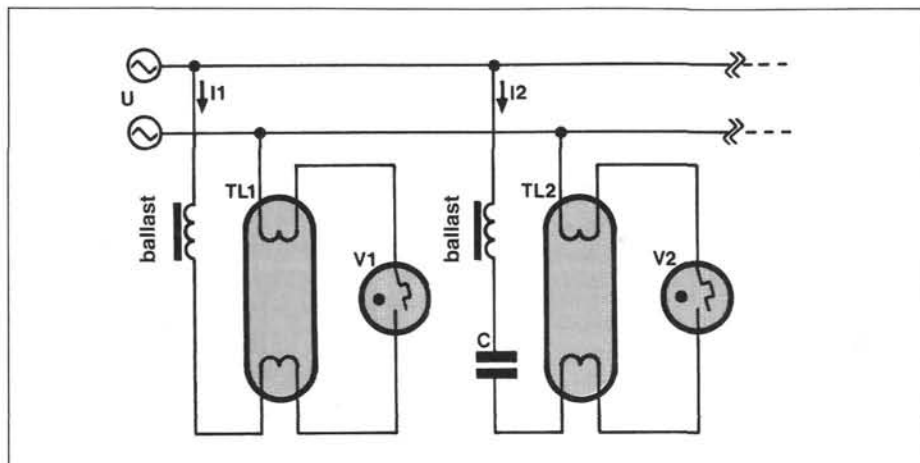


Figure 5 – Les tubes fluorescents sont, le plus souvent, câblés deux par deux en parallèle : en duo.

en attendre en retour, puisque seule la consommation du tube est prise en compte.

Pour les bureaux, les magasins, les usines, ce courant prend des proportions respectables. Dans de telles installations, le fournisseur d'électricité oblige le consommateur à prendre des mesures pour réduire au maximum la consommation de **puissance réactive** à laquelle un tel courant donne lieu : les spécialistes disent que le cosinus phi ($\cos\phi$, pour les intimes) doit être maintenu au-dessus de 0,857. Nous ne nous étendrons pas sur ces termes inconfortables qui nécessiteraient de trop longs développements et qui concernent surtout les gros consommateurs.

Puisque la bobine est **consommatrice** de puissance réactive, à cause du déphasage en retard du courant sur la tension qu'elle occasionne, on peut, pour compenser, inclure dans le circuit un **générateur** de puissance réactive : un condensateur pour lequel le déphasage du courant est en avance sur la tension. Ceci n'est pas fait pour étonner les lecteurs d'ELEX, puisqu'ils savent qu'en continu le courant, dans un circuit contenant un condensateur, est maximum à la mise sous tension, puis diminue quand la tension aux bornes du composant augmente, pour s'annuler lorsqu'elle est maximale. Si l'on câble deux tubes en parallèle en ajoutant un condensateur, comme sur la **figure 5**, on modifie le déphasage pour le second tube de telle façon qu'il compense celui apporté par le premier. Cette façon de procéder présente un autre avantage important : un tube fluorescent clignote, à 100 Hz il est vrai et cela ne nous est pas perceptible⁽⁶⁾. Il est cependant possible qu'une machine tournante, éclairée

proportionnelle à l'inductance de la bobine et à la vitesse de variation du courant, tend à maintenir celui-ci dans le circuit. Il n'y a là rien de miraculeux. Si "inductance" vous fait penser à "inertie" et "bobine" à "volant d'inertie", vous n'êtes pas loin de la vérité. La bobine a en effet emmagasiné de l'énergie magnétique, comme un volant entraîné par un moteur emmagasine de l'énergie cinétique. La bobine a donc à ses bornes une tension élevée lorsque le court-circuit cesse, et cette tension se retrouve, avec celle du secteur, aux bornes du tube (4c). Le champ électrique est maintenant suffisant pour permettre aux électrons des filaments de se libérer et de circuler : le courant est donc établi entre les électrodes du tube. Les collisions de ces électrons libres primaires avec les atomes de gaz en libèrent d'autres comme nous l'avons vu plus haut, et le courant atteint rapidement son intensité de "croisière". Cette intensité est bien sûr limitée par la bobine qui lui évite d'atteindre une valeur destructrice (nous sommes en alternatif et l'impédance de la bobine est toujours en série avec le tube).

Les filaments incandescents manquent sur la **figure 4**. Théoriquement, ils ne sont pas nécessaires, mais le courant d'électrons s'amorce avec une tension beaucoup moins élevée si les électrodes sont chauffées. Pour des

tubes dont la longueur est supérieure à 50 cm, la différence est considérable.

charge complexe

Avec son ballast, un tube fluorescent est une charge dite complexe pour le secteur. Une telle charge peut poser des problèmes aux fournisseurs d'électricité comme EDF. – Pourquoi ? – Tout simplement parce qu'à cause de la bobine le courant qui traverse la lampe est déphasé de 90° en retard sur la tension du secteur. Tout se passe comme s'il y avait circulation de deux courants alternatifs dont les maximums ne se superposaient pas. Un des deux courants, en phase avec la tension du secteur, permettrait à la lampe de brûler, pendant que l'autre engendrerait dans la bobine un champ magnétique puis serait coupé. Ce dernier courant, que nous appellerons "aveugle", en quadrature retard sur la tension, est source de tous les maux pour les centrales électriques⁽⁵⁾. Celles-ci doivent en effet fournir de l'énergie à la bobine au moment où elles n'en ont pas de disponible (lorsque la tension s'annule) et en recevoir en retour lorsqu'elles n'en ont pas besoin. À cela s'ajoute que ce courant circule et chauffe donc les lignes par effet joule. La centrale doit donc compenser la perte d'énergie due à ce courant "aveugle" sans rien

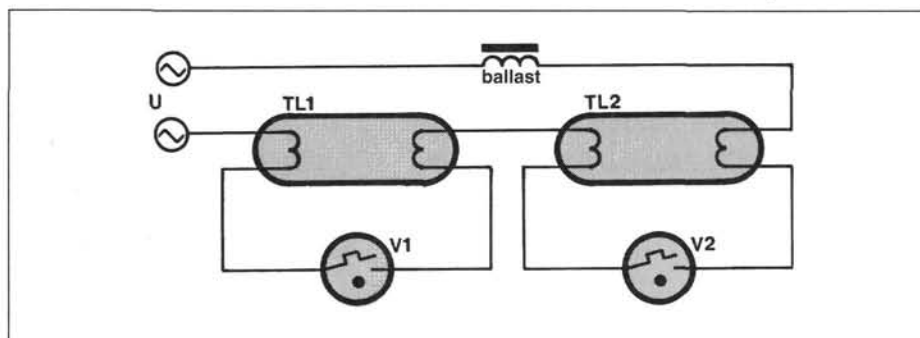
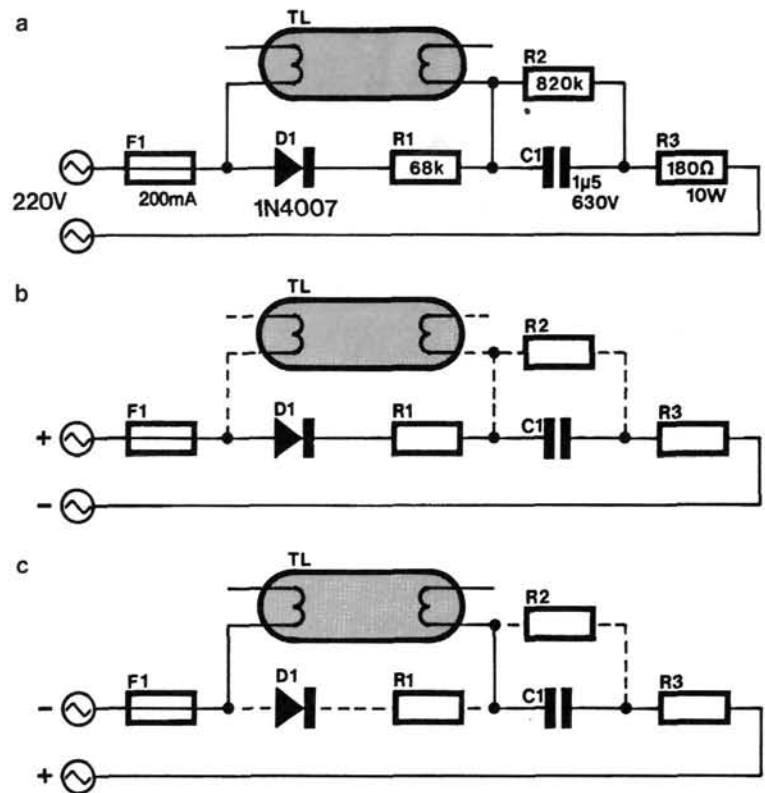


Figure 6 – Le câblage en tandem permet de faire l'économie d'un ballast : les starters fonctionnent cependant sous 110 V, ce qu'il faut préciser lors de leur achat, et la puissance nominale du ballast doit correspondre à la puissance totale des deux tubes.

⁽⁵⁾C'est l'intensité "dewattée" ou "magnétisante". Elle n'arrange les affaires de personne puisqu'elle correspond à une puissance, dite réactive, qui n'est pas utilisable.



Figures 7 et 8 – Pour des tubes de faible puissance (jusqu'à 8 W), comme ceux utilisés pour les baladeuses, il est possible de se passer de ballast.



par un seul tube, paraisse tourner plus lentement, voire même être arrêtée avec les conséquences dangereuses que vous imaginez (effet stroboscopique de *strobos*, "rotation" et *skopein* "examiner", "observer"). Avec le branchement en duo de la figure 5 et le condensateur de compensation, les tubes s'allumeront à peu près à tour de rôle et l'éclairage sera quasiment constant.

autres branchements

Les circuits décrits sur les figures 2 et 5, ne diffèrent guère quant au fonctionnement. Ce ne sont pas les seules manières de brancher des tubes fluorescents (TL ou PL). Nous allons en passer quelques unes en revue, en laissant de côté les solutions électroniques particulières, souvent trop compliquées à réaliser soi-même. Le branchement en tandem est des plus simples. Particulièrement adapté aux aquariums, aux vasques, à la décoration, il présente un certain avantage puisqu'il n'utilise qu'un ballast pour

deux lampes (figure 6). Le fonctionnement est en gros le même que dans un branchement normal mais ici les deux tubes et leur starter respectif sont en série. Il ne faut donc pas que les tensions d'amorçage soient trop élevées, sinon la tension du secteur ajoutée à la force électromotrice d'auto-induction n'y fera pas face. Ceci ne pose pas de problème si les tubes sont de petites dimensions (jusqu'à 20 W environ). Une précision cependant concernant les starters : quand vous les achetez, dites bien qu'ils fonctionnent sous 110 V au lieu de 220 V ou que ce sont des "starters-série". Le ballast doit d'autre part correspondre à la puissance totale des tubes.

Le procédé suivant (figures 7 et 8) est utilisé pour l'amorçage des baladeuses à tube fluorescent et ne nécessite pas de ballast, ce qui représente un avantage de poids et de prix. Le branchement est le même que celui décrit sur la figure 2 à une différence près, l'absence de bobine. Pour ces tubes il n'est évidemment plus question de tension d'amorçage élevée ce qui s'explique par leur petite dimension. La tension crête du secteur est alors suffisante. Reste le problème de la limitation de courant, résolu de façon très élégante : le cordon qui relie la lampe au secteur est résistif. Sa longueur est telle que la résistance

en série avec la lampe soit de 1 kΩ. Le calcul montre que l'intensité du courant qui traverse une lampe de 8 W alimentée par ce cordon est limitée à 175 mA. Le cordon à lui tout seul consomme 30 W en effet joule (dégagement de chaleur) ce qui n'est malgré tout sensible au toucher qu'après quelques heures de fonctionnement de la baladeuse. Au total donc le dispositif consomme 38 W. La lampe à incandescence qui pourrait fournir la même quantité de lumière consommerait, elle, 40 W : on ne peut pas parler ici d'économie d'énergie. La baladeuse présente cependant quelques avantages : elle tient le choc, ce qu'on ne peut pas dire de sa collègue à incandescence, et la chaleur qu'elle diffuse est mieux répartie.

Le dernier circuit dont nous parlerons peut faire l'objet d'un bricolage amusant et montrer qu'un tube fluorescent défectueux n'est pas forcément inutilisable. Comme le montre la figure 8, jusqu'à 8 W, les tubes ne contiennent pas de filament incandescent. On supprime ainsi une cause de panne puisque ce filament est la partie la plus fragile d'un tube fluorescent. Il faut bien entendu que la tension d'amorçage soit supérieure à ce qu'elle serait si le tube était chauffé : ce circuit en tient compte. En effet, avant que la lampe fonctionne, pendant l'alternance positive, le conden-

(6) À moins de disposer d'un PC et du logiciel "CHRONOPC CHUTE LIBRE" proposé par l'Union des physiciens (50 F à l'ordre de L'Union des physiciens, chez Philippe Baffert, 3, rue Hector Berlioz 94370 Sucy en Brie). Ce logiciel, qui permet d'ailleurs des mesures plus sérieuses que celle-ci (ELEX numéro 44, page 10), est indispensable à tout enseignant de physique ou de mécanique un tant soit peu sérieux.

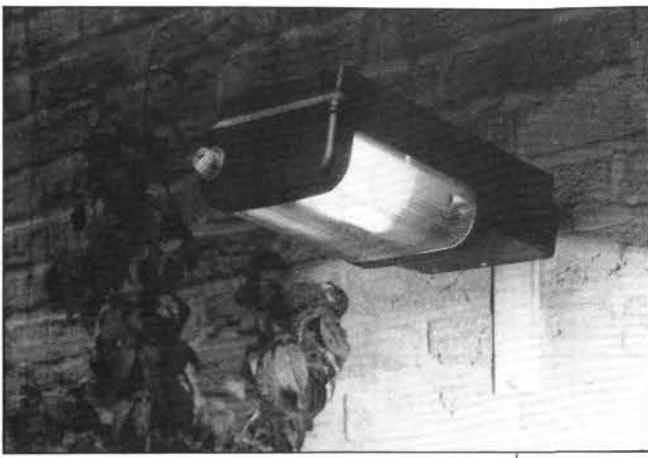
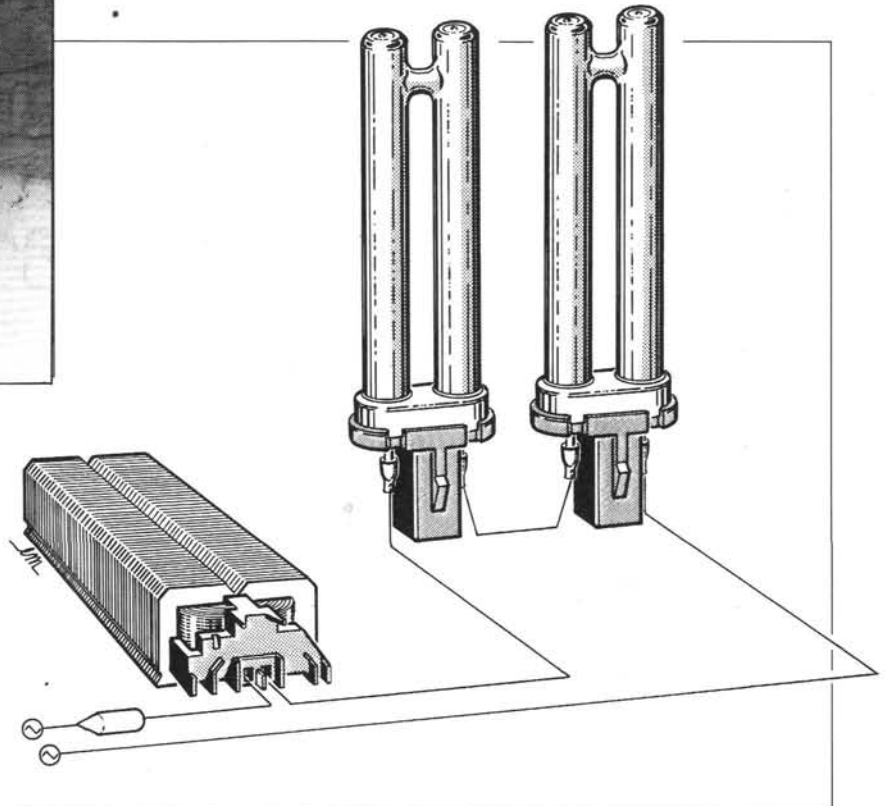


Figure 9 – La forme recourbée des tubes contenus dans les lampes SL et PL implique des conditions particulières d'ionisation aux gaz qu'ils contiennent.

Figure 10 – Il est possible de récupérer le ballast d'une lampe SL usagée pour amorcer une lampe PL. Il faut bien sûr qu'il y ait correspondance entre les puissances.



lampes PL et SL

sateur C1 (figure 8b) se charge à travers D1 et R1 à la tension crête du secteur, soit un peu plus de 300 V. Pendant l'autre alternance, le tube a entre ses électrodes la tension du secteur à laquelle s'ajoute la tension aux bornes du condensateur, soit plus de 600 V crête. C'est plus qu'il ne faut pour amorcer un tube de 8 W sans chauffer ses électrodes. Dès que la lampe est allumée, le courant ne passe pratiquement plus par R1 et D1 (il est inférieur à 1 mA) et la tension aux bornes de ces deux composants est assez proche de 0 V pour qu'on les considère hors circuit. Restent en course R3 et C1 qui limitent l'intensité du courant dans la lampe. À la différence de ce qui se passe dans la baladeuse décrite plus haut, cette limitation ne consomme presque pas de puissance en dehors des 5 W dissipés dans R3. – Et R2 ? – Comme vous l'avez sans doute remarqué, dans les cas 8b et 8c, cette résistance est représentée hors circuit par des pointillés. Elle ne joue de rôle qu'à la mise hors tension. Si le circuit est coupé du secteur à un moment où le condensateur est chargé, R2 lui permet de se décharger assez rapidement pour que les broches du cordon puissent être touchées sans danger.

Nous parlons des tubes fluorescents ordinaires, les TL et ce sont les PL et SL qui sont depuis quelques temps à l'ordre du jour. Il est de fait que toutes ces lampes fonctionnent suivant le même principe. Quelques différences méritent cependant d'être notées, puisque les dernières arrivées sur le marché ont nécessité un certain nombre de recherches avant leur commercialisation.

Le principal problème était leur forme repliée. Il fallait littéralement trouver quelque chose comme "le fusil à tirer les électrons dans les coins", ce qui n'était pas facile. Les électrons se déplacent en effet en ligne droite (pour les courtes distances où nous les considérons) même si cela doit signifier pour eux la traversée du verre et de l'air. Un mélange de gaz spécial, tout simplement ionisé, leur facilite le trajet, puisqu'il est conducteur, et leur évite donc de sortir du tube. Pour le reste, rien de changé, qu'elles soient PL ou SL, ces lampes nécessitent toujours une bobine et un starter. La différence entre les deux est que ces composants sont inclus dans les SL, qui sont donc un peu plus chères, alors que les PL ne contiennent que le starter. Lorsqu'une lampe SL rend l'âme, après quelques années,

c'est parce que le tube ou le starter sont hors service. Au lieu de tout jeter, les petits malins peuvent récupérer le ballast et l'utiliser pour une PL, à condition de tenir compte de leurs puissances respectives. Si vous démontez par exemple une SL de 18 W, son ballast servira à alimenter deux PL de 9 W en série (figure 10) ou une PL de 18 W. Le composant un peu exotique visible sur la figure 10 protégeait la lampe SL contre la surchauffe et peut donc être conservé.

conseils et précautions

Si vous avez l'occasion de bricoler le montage décrit sur la figure 10, n'oubliez pas les consignes de sécurité. Ne branchez jamais une lampe PL ou TL sans intermédiaire sur le secteur : le ballast est indispensable et ses caractéristiques doivent être telles que la tension aux bornes de la lampe ne soit pas trop élevée. Utilisez un ballast dont la puissance corresponde à celle que la lampe consomme. Effectuer un branchement avec une résistance ou un condensateur en série pose vraiment de grandes difficultés (dimensionnement de ces composants) et nous ne saurions trop vous le déconseiller.

886108

ersatz électronique du valium



MACHINE À RÊVES

pseudo-dent de scie

- Le stress est une de ces maladies modernes qui s'attaquent aux « décideurs », aux « hommes » politiques, à tous ceux, en général, qui portent une charge trop lourde pour leurs épaules, sans voir de possibilité de fuite. Imaginez celui qui vous précédait dans la file du péage, au volant comme vous, par trente degrés à l'ombre, se demandant quand le bouchon allait avancer, ne se demandant pas pourquoi il devait aller reprendre son boulot le lendemain : il faut bien payer les traites de la roulotte, de la voiture qui la traîne, l'essence qu'elle engloutit toute l'année, il faut bien revenir bronzé pour pouvoir regarder en face le chef de bureau, il faut bien partir en vacances pour évacuer le stress de l'année... Certains cherchent la fuite dans les drogues, avec ou sans ordonnance médicale ; nous vous proposons de vous évader avec une machine à rêves, sans ordonnance ni contre-indication.
- L'homme a besoin de repos. Un besoin vital de repos. Il ne s'agit pas tant du repos du corps – quelques heures de sommeil chaque nuit suffisent à remettre l'organisme en état de marche – que de celui du cerveau.
- Le repos cérébral est au moins aussi important, peut-être plus ! Il ne suffit pas de s'affaler le soir devant la télé pour se reposer les méninges.
- Il est indispensable de se vider par moments de tous ses soucis, qu'ils

soient professionnels, domestiques ou autres, de porter le regard au loin (à dix pas, comme on dit chez les militaires) et de mettre le cerveau au point mort (comme on fait...). Il semble que beaucoup de gens n'y arrivent qu'à grand peine, si on se reporte aux statistiques de consommation de calmants. En fait il n'est pas nécessaire de vous jeter sur votre flacon de valium chaque fois que vous voulez vous détendre un peu : il est connu que certains bruits procurent une détente bénéfique, et qu'ils ne produisent pas d'accoutumance. C'est le cas, par exemple, du bruissement du vent dans les branches, dans une forêt, ou du bruit des vagues au bord de la mer. Pour vous qui habitez en ville, il est encore possible de trouver un arbre qui ne soit pas empoisonné par les gaz d'échappement, mais quant à une forêt ou même un bosquet... nib ! Si encore vous trouviez plusieurs arbres et qu'il y ait du vent, le bruit serait largement couvert par celui des voitures, des camions, des bus et des pétrolettes de toutes sortes. Comme tout le monde ne peut pas habiter au bord de la mer (vous avez vu ce que c'est en juillet-août, au moment de la transhumance des hordes bataves et teutonnes), il ne vous reste qu'à prendre votre fer à souder et à fabriquer votre générateur de bruit de vagues individuel.

La caractéristique la plus marquante du bruit des vagues est sa montée lente, quand une vague arrive, et sa chute relativement brutale, quand elle s'écroule. Comme il s'agit d'un **bruit**, c'est un mélange aléatoire de toutes les fréquences possibles. Vous n'allez donc pas être étonné d'apprendre que notre générateur de bruit de vagues est constitué de deux parties : un générateur de dents de scie et un générateur de bruit. C'est la dent de scie qui reproduit le mieux la croissance lente et la décroissance rapide du bruit de la marée.

Il est probable que le coup d'oeil curieux que vous avez déjà jeté au schéma de la **figure 1** ne vous a pas permis de reconnaître les deux parties du montage. Nous allons l'examiner en détail. Commençons par le générateur de dents de scie construit autour du circuit intégré IC1. Il ne s'agit en fait que d'un générateur de pseudo-dents de scie : dans une vraie dent de scie, la croissance de la tension est **linéaire**, alors qu'ici nous avons affaire à une charge et une décharge du condensateur C1. Les connaisseurs savent que ce phénomène suit une courbe dite exponentielle*. Il n'en reste pas moins que la forme d'onde obtenue est presque assimilable à une dent de scie et qu'elle est parfaitement utilisable pour notre montage.

* Ce qui l'apparente aux vagues.

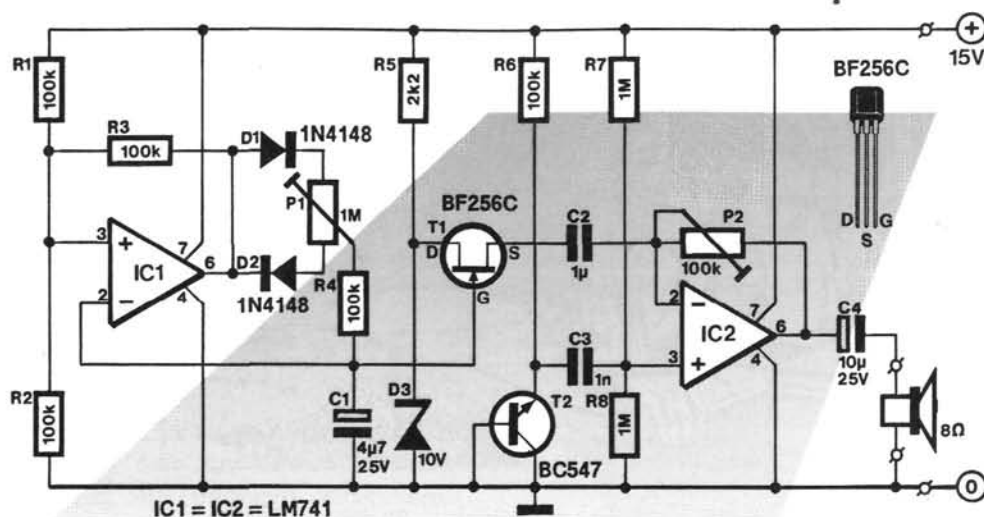


Figure 1 - Le générateur de bruit de vagues est constitué d'un générateur de (pseudo) dents de scie (IC1), d'un générateur de bruit (T2) et d'un amplificateur (IC2). La tension en dent de scie commande le gain de l'amplificateur par l'intermédiaire du transistor à effet de champ T1.

hystérésis

L'amplificateur opérationnel IC1 est monté en comparateur. Sa sortie ne peut donc prendre que deux états : haut, proche de la tension d'alimentation positive, quand la tension de l'entrée non-inverseuse (+) est supérieure à celle de l'entrée inverseuse (-) ; ou bas, proche de la tension d'alimentation négative (ici la masse), quand la différence de potentiel entre les entrées change de signe.

Supposons que le condensateur C1 est déchargé ; cette supposition n'est pas gratuite, C1 a toutes les chances de rester déchargé aussi longtemps que le montage n'est pas sous tension. Ce condensateur déchargé présente à l'entrée inverseuse du comparateur une tension nulle, donc forcément inférieure à celle de l'entrée non-inverseuse, fixée par le diviseur R1/R2 à la moitié de la tension d'alimentation (nous verrons plus loin le rôle de R3). La sortie du comparateur prend sa valeur la plus haute. La tension d'alimentation est appliquée, par la diode D1, la partie « supérieure » de P1 et R4, au condensateur C1 qui se charge. La charge dure jusqu'à ce que la tension dépasse celle de la broche 3 du comparateur. À ce moment la sortie bascule (elle prend l'état bas), ce qui permet au condensateur de se décharger, toujours à travers R4, mais cette fois à travers la moitié « inférieure » de P1 et la diode D2. La décharge dure jusqu'à ce que la tension de l'entrée inverseuse devienne inférieure à celle de l'entrée non-inverseuse, que le comparateur bascule à nouveau, et que le condensateur recommence à se charger.

En pratique, notre oscillateur risque de ne pas donner satisfaction : le changement d'état de la sortie se produit pour une différence de tension entre les entrées de quelques millivolts seulement. Si nous donnons à la tension de référence de l'entrée non-inverseuse une valeur fixe, l'oscillation de la sortie sera très rapide et de faible amplitude. Ce n'est pas ce que nous voulions. La solution tient dans la résistance R3, que nous avons écartée tout-à-l'heure.

Quand la sortie du comparateur est à l'état bas, R3 se trouve connectée en parallèle sur R2 (puisque leurs extrémités sont aux mêmes potentiels). La tension de l'entrée non-inverseuse est donc de 5 V environ, et non de 7,5 V.

[calcul pour les accros : $R2/(R3)/(R1+R2/R3) \times 15V$] La tension du condensateur devra donc diminuer jusqu'aux environs de 5 V avant que le comparateur bascule. Dès que la sortie prend l'état haut, R3 se trouve en parallèle sur R1 ; un calcul similaire montre que la tension du condensateur doit atteindre 10 V pour que le comparateur bascule à nouveau. L'amplitude de l'oscillation sur C1 est de 5 V, ce qui est suffisant. De plus, la durée du cycle de charge-décharge s'allonge, autrement dit la fréquence diminue. L'écart de 5 V entre les deux seuils de basculement du comparateur s'appelle hystérésis, un mot grec qui ne signifie rien d'autre que retard. Le rapport de résistance entre les deux parties de P1 détermine le rapport entre les temps de charge et de décharge du condensateur, donc la pente des courbes correspondantes. C'est le réglage de P1 qui donnera son aspect à la dent de scie.

générateur de bruit

Dans la plupart des appareils électroniques d'usage courant, comme les amplificateurs ou les récepteurs de radio, le bruit est l'ennemi qu'on essaie de supprimer à grand renfort de schémas *pinailés* et de composants triés. Ici aussi, nous avons besoin d'un composant sélectionné pour son niveau de bruit, mais pour un niveau aussi élevé que possible. Ce sera un transistor au silicium banal, utilisé dans des conditions particulières. Le BC547 repéré T2 est monté comme une diode, collecteur et base court-circuités, polarisée en inverse. Comme la tension d'alimentation est relativement élevée, la jonction base-émetteur se comporte comme une diode zener et laisse passer un courant limité par la résistance R6. Les diodes zener en général sont bruyantes, le transistor utilisé en diode zener est particulièrement bruyant.

C'est l'amplificateur opérationnel IC2 qui amplifiera le bruit du transistor jusqu'à un niveau suffisant pour le rendre audible. Il est monté en amplificateur inverseur, d'une sorte particulière. Le signal à amplifier n'est pas appliqué à l'entrée inverseuse, mais superposé à la tension de polarisation continue de l'entrée non-inverseuse (fixée à 7,5 V par le diviseur R7/R8). Abstraction faite du condensateur C2, l'amplificateur est monté en suiveur de tension puisque sa sortie est reliée à l'entrée inverseuse. C'est vrai pour ce qui est des tensions continues, le gain est de un. Le condensateur intervient pour les tensions alternatives, il referme pour elles la boucle de contre-réaction. Le gain en alternatif est déterminé par le rapport entre P2 et la

résistance drain-source du transistor à effet de champ T1. Le trajet drain-source d'un transistor à effet de champ peut être considéré comme un canal dont la largeur – et la résistance – varie en fonction de la tension appliquée à la grille (un FET est commandé par une tension, contrairement aux transistors bipolaires, commandés en courant).

La grille du transistor à effet de champ est reliée au condensateur C1 ; le gain de l'amplificateur IC2 va donc varier en fonction de la tension de la dent de scie. Le bruit du transistor-zyner, audible dans le haut-parleur HP1, reproduit la montée lente et la descente brusque du bruit de la marée.

la construction

Deux amplificateurs opérationnels, deux transistors et quelques bricoles trouveront place facilement sur une platine d'expérimentation de format 1, **figure 2**, ou sur le petit circuit imprimé de la **figure 3**. Ni l'une ni l'autre méthode ne doit poser de problème. Dans le cas de la platine d'expérimentation, il faut simplement veiller à ne pas oublier de pont de câblage : dans les deux cas, respecter la polarité des diodes et condensateurs, et de monter les circuits intégrés sur des supports.

La tension d'alimentation, de 15V environ, sera fournie par un bloc secteur. La valeur exacte n'est pas importante, pas plus que la stabilisation, mais la tension doit être de 15V au moins pour que le bruit de T2 soit suffisant. Il n'est pas rare que les blocs-secteur à quatre sous délivrent une tension à peine filtrée. Si c'est le cas du vôtre, il suffit d'ajouter un condensateur de 470 μ F/35V en parallèle sur les bornes d'alimentation, sous la platine d'expérimentation, ou à la place réservée (C5) sur le circuit imprimé.

Le réglage n'est pas critique, c'est une simple question de goût. Le potentiomètre P1 sert à régler le rapport de durée entre les deux flancs de la dent de scie, autrement dit entre le temps de « montée » des vagues et le temps de « chute » (inutile de vous appuyer mille bornes par l'autoroute et le king-macquickburger* à la pause pour aller étalonner au bord de la mer).

*C'est encore plus mauvais quand ça remonte que quand ça descend, contrairement aux bananes, recommandées par tous les marins.

liste des composants

R1 à R4, R6 = 100 k Ω

R5 = 2,2 k Ω

R7, R8 = 1 M Ω

P1 = 1 M Ω variable

P2 = 100 k Ω (voir texte)

C1 = 4,7 μ F/25 V

C2 = 1 μ F (non polarisé)

C3 = 1 nF

C4 = 10 μ F/25 V

C5 = 470 μ F/35 V (voir texte)

D1, D2 = 1N4148

D3 = zener 10 V/400 mW

D4 = 1N4001 4007 (voir texte)

T1 = BF256C

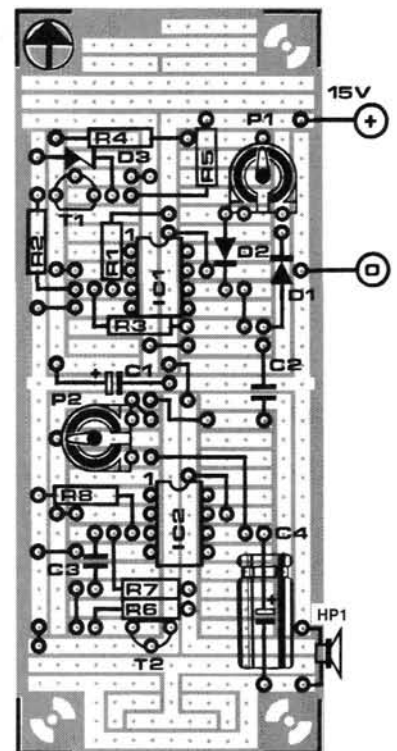
T2 = BC547

IC1, IC2 = 741

HP1 = haut-parleur 8 Ω

platine d'expérimentation de format 1 ou circuit imprimé

Figure 2 - Tous les composants trouvent une place sur la platine d'expérimentation de format 1.



Vous pouvez décider de régler le volume une fois pour toutes, le potentiomètre P2 sera alors un modèle miniature installé sur le circuit imprimé ; ou bien de le rendre variable, dans ce cas P2 sera un potentiomètre ordinaire dont l'axe sortira du coffret.

utilisation

En plus de son utilisation anti-stress (dans la chambre à coucher ou comme un baladeur dans les embouteillages), la machine à rêves peut trouver d'autres applications, comme les effets sonores pendant les projections de films, de diapositives, ou des représentations théâtrales. Vous pouvez aussi le construire juste pour le plaisir et dissimuler le haut-parleur à proximité de votre bocal à poissons rouges, ça leur fera des vacances.

886106

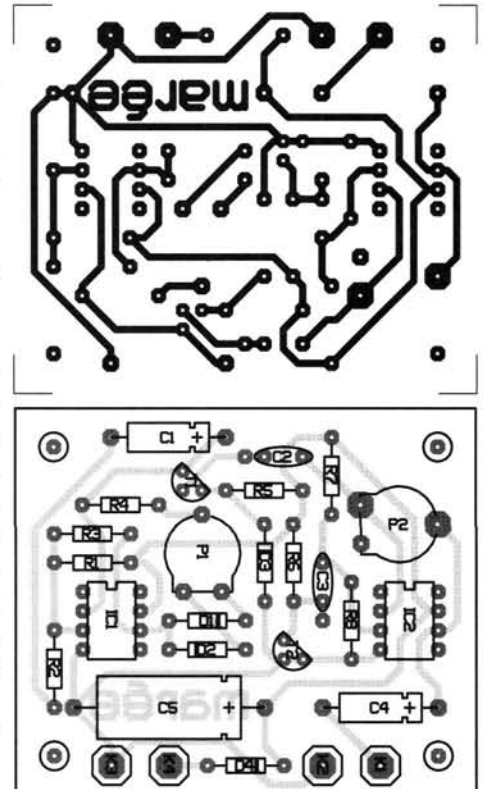


Figure 3 - Le circuit imprimé dessiné pour ce générateur de bruit de vagues. Le condensateur de filtrage supplémentaire y est prévu, de même que la diode D4 pour les étourdis qui raccorderaient à l'envers les deux fils du bloc secteur. Si vous montez un potentiomètre extérieur pour le volume, prévoyez des picots à souder dans les deux grosses pastilles de P2.

MOTEURS ELECTRIQUES

Plus personne ne s'étonne aujourd'hui de voir un appareil se mettre en mouvement quand il enfonce une fiche dans une prise de courant. Les moteurs électriques sont devenus quelque chose de banal. Ce n'est pas étonnant car il y a plus de cent ans que l'électricité fait tourner des machines. Ce qui peut rester un peu mystérieux, c'est la façon dont ce courant d'électrons invisibles peut provoquer un mouvement.

sont de même signe, ce qui a tendance à repousser le fil. La force symbolisée F est la somme de l'attraction et de la répulsion que subit le fil. Une force exactement identique et de sens opposé s'exerce sur l'aimant. Si l'aimant est fixé ou assez lourd, que le fil est mobile ou assez souple, il va s'écarter vers l'extérieur de l'aimant, s'éloigner du champ magnétique. C'est ainsi que le passage du courant électrique est converti en un mouvement. Dans cette expérience, le mouvement est de très faible amplitude et la force produite inutilisable pratiquement.

Le déplacement du fil cesse aussitôt qu'il se trouve en dehors du champ. Il va donc trouver une position d'équilibre, « au bord » du champ de l'aimant. Le phénomène serait plus intéressant si nous pouvions obtenir, au lieu du déplacement linéaire (la translation), un mouvement circulaire (une rotation). Dans ce cas, le fil resterait dans le champ magnétique. Tout cela est vite dit, mais la réalisation pose un certain nombre de problèmes.

Le phénomène qui permet à un moteur électrique de tourner est simple : il s'agit des forces d'attraction et de répulsion qu'exercent des aimants. Un conducteur quelconque, parcouru par un courant électrique, se transforme en aimant. Le passage du courant fait naître un champ magnétique autour du fil (figure 1).

Les lignes de force du champ sont orientées en fonction du sens du courant. C'est la règle dite du tire-bouchon qui permet de connaître cette orientation. Si nous imaginons que le tire-bouchon s'enfonce dans le fil dans le même sens que le courant, le

sens de rotation du tire-bouchon est le même que celui des lignes de force du champ magnétique.

Nous savons que deux aimants exercent une force l'un sur l'autre ; ils s'attirent si les pôles opposés se font face, ou se repoussent si les pôles identiques se font face. Le fil parcouru par un courant, entouré d'un champ, subit lui aussi une force s'il est placé dans un autre champ magnétique, comme celui d'un aimant permanent. Le phénomène est représenté par la figure 2 ; la partie 2b est une vue en coupe de la partie 2a. Elle représente le champ qui se construit autour du fil, et son orientation par rapport au champ de l'aimant permanent. Sur la gauche de la figure, les champs sont de sens opposé, donc le fil est attiré vers l'extérieur ; sur la droite, les champs

Figure 1 – Un champ magnétique se forme autour d'un conducteur parcouru par un courant. La règle du tire-bouchon permet de connaître l'orientation des lignes de force.

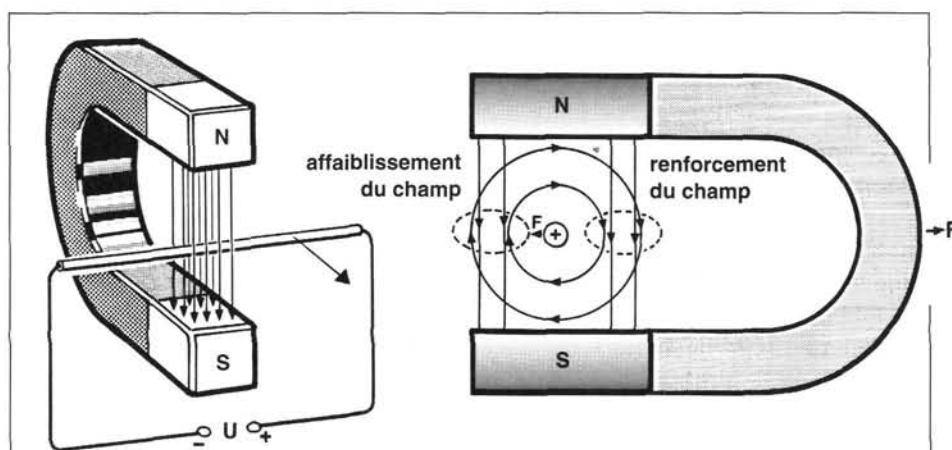
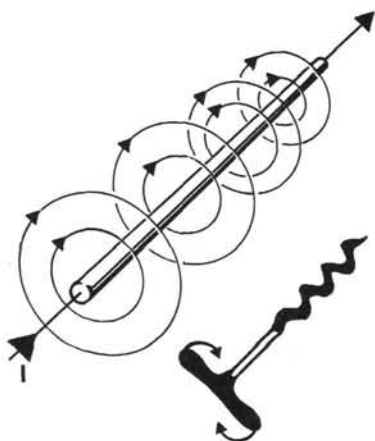


Figure 2 – Comme les deux champs se renforcent mutuellement à droite et s'opposent à gauche, le conducteur subit une force dirigée vers la gauche.

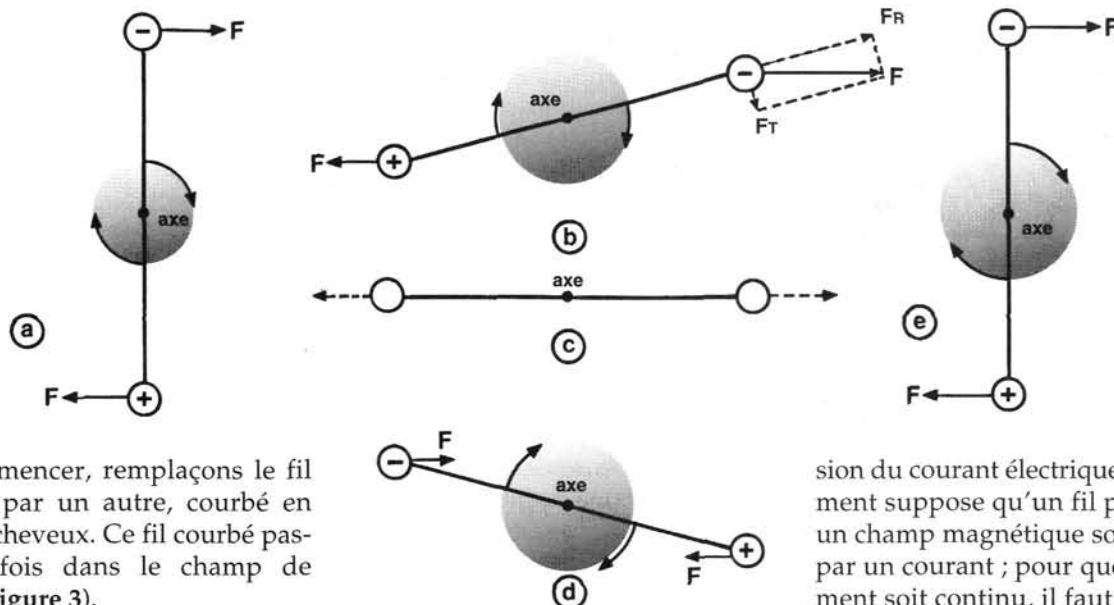


Figure 4 – Pour que l'épingle à cheveux tourne continuellement, il faut que le sens du courant, donc la polarité de la tension d'alimentation, change exactement au moment où le fil est dans un plan horizontal.

Pour commencer, remplaçons le fil rectiligne par un autre, courbé en épingle à cheveux. Ce fil courbé passe deux fois dans le champ de l'aimant (figure 3).

Si un courant parcourt ce petit cadre, il est de sens opposé dans chaque moitié du fil : de gauche à droite dans le fil inférieur, de droite à gauche dans le fil supérieur. Les conséquences mécaniques sont faciles à déduire. La moitié inférieure va être repoussée vers l'extérieur comme dans le cas du fil simple. La moitié supérieure va être attirée vers l'intérieur puisque le courant et le champ magnétique qu'il produit sont de sens opposé. Si l'épingle est montée sur un axe passant par le milieu de l'entrefer de l'aimant, elle va tourner autour de cet axe, sous l'effet du couple de forces antagonistes (figure 3b).

La figure 4a représente l'épingle sans les différentes lignes de champ. Si l'épingle est libre de se mouvoir, elle arrivera en position horizontale. À ce moment le mouvement s'arrêtera puisque les forces s'exerceront dans le plan de l'axe (figure c). Les choses en resteront là si rien ne se passe. On peut imaginer de couper le courant aussitôt que l'épingle arrive à l'horizontale. La force d'inertie lui permet alors de continuer sa course pour

revenir dans la position d'origine. Malheureusement les forces de frottement auront tôt fait d'arrêter le mouvement, sans parler de la charge mécanique éventuelle du moteur. Pour assurer le mouvement, il faut inverser le sens du courant dans l'épingle pour créer des forces orientées dans le bon sens, à condition que la position horizontale soit légèrement dépassée (figure 4d). La rotation reprend pour un demi-tour, en passant par la position de la figure 4e (identique à celle de la figure 4a), pour finir comme en 4c. Il faut à nouveau inverser le sens du courant.

Le mouvement continue, nous avons créé un moteur électrique, constitué comme les autres d'une partie fixe, le stator, et d'une partie mobile, le rotor. Nous tirons deux conclusions importantes de ce qui précède : la conver-

sion du courant électrique en mouvement suppose qu'un fil plongé dans un champ magnétique soit parcouru par un courant ; pour que le mouvement soit continu, il faut que le sens du courant soit inversé à chaque demi-tour. Cette dernière condition peut être remplie de deux manières. La première est d'alimenter le moteur en courant alternatif. Il faut alors que l'inversion de polarité se produise au bon moment, ce qui se passe bien en pratique, car le moteur a tendance à tourner plus vite que ne le permet la fréquence de la tension alternative. Si le fil de notre moteur (figure 3) a dépassé l'horizontale avant l'inversion du sens du courant, les forces produites le freineront pour le « remettre au pas », le synchroniser. Avec une tension à 50 Hz comme celle du secteur, le moteur ne peut tourner qu'à la vitesse de 50 tours par seconde, soit 3000 tours par minute. Si la charge mécanique ou les frottements sont trop importants, le moteur s'arrête.

Si nous voulons alimenter notre moteur en courant continu, il faut trouver un moyen d'inverser la polarité à chaque demi-tour. Ce moyen s'appelle **collecteur**, il est représenté sur la figure 5. Le collecteur est constitué, pour notre rotor à deux pôles, de deux demi-bagues conductrices, isolées l'une de l'autre, montées sur l'axe du moteur. Elles sont reliées chacune à une extrémité de l'enroulement du rotor, et reçoivent le courant par deux frotteurs, les **balais**. Les balais sont soumis à une tension continue, leur polarité ne change pas. Ce qui change à chaque demi-tour, c'est l'extrémité de l'enroulement qui est alimentée par un balai donné.

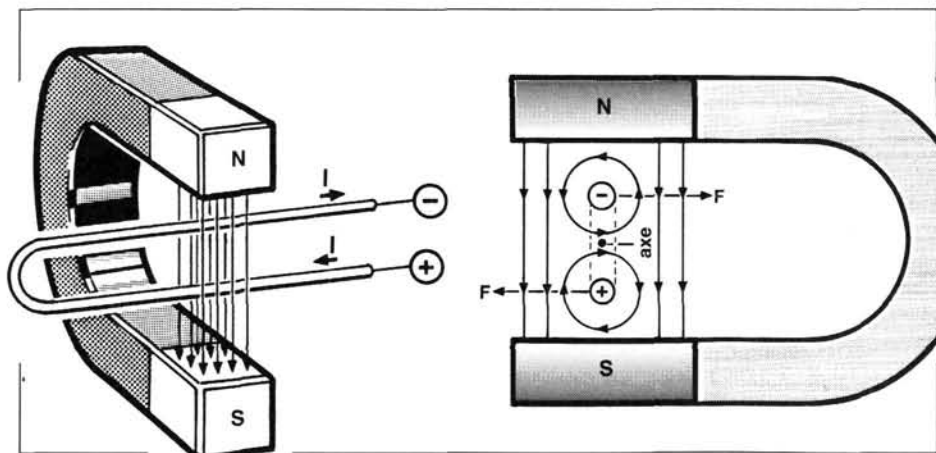


Figure 3 – Le fil courbé en épingle à cheveux tourne sous l'effet du champ magnétique.

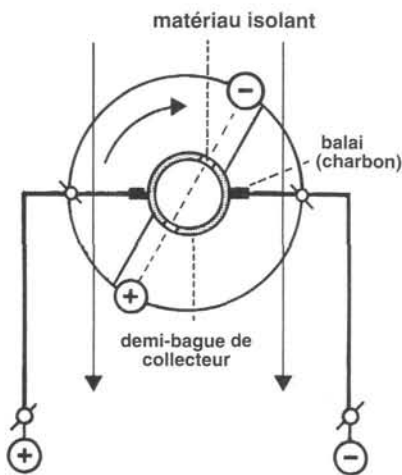
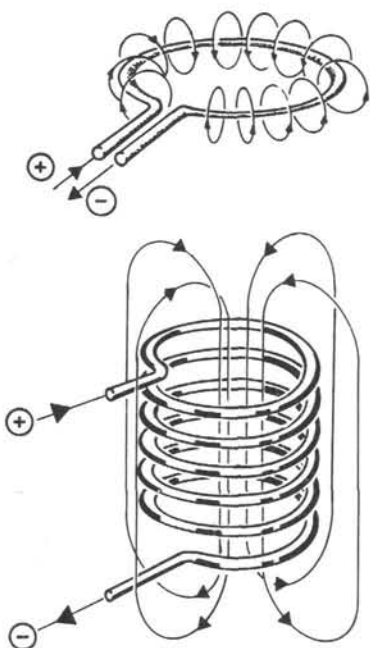


Figure 5 – Deux demi-bagues connectées aux extrémités de la bobine sont montées sur l'axe. Les balais, ou charbons, viennent en contact avec le collecteur que forment ces deux demi-bagues. Cette disposition permet d'inverser la polarité au moment exact où la bobine est horizontale.

L'inversion se produit maintenant quand le rotor est exactement à l'horizontale.

L'avantage de ce moteur à courant continu par rapport à son homologue à courant alternatif est que la vitesse de rotation n'est plus liée à la fréquence de la tension d'alimentation, ni du côté des basses vitesses ni du côté des hautes vitesses. Cet avantage est lié au fait que c'est le rotor lui-

Figure 6 – Le champ du rotor augmente si le bobinage comporte plusieurs spires au lieu d'une seule.



même qui commande l'inversion de polarité. Le régime (ou vitesse de rotation) n'est plus déterminé que par la tension d'alimentation et la charge mécanique du moteur.

Jusqu'à présent, nous avons considéré un moteur dont le rotor est constitué d'une seule boucle, même si nous l'avons appelé enroulement. Tout le monde a déjà vu les entrailles d'un moteur électrique, que ce soit un modèle industriel ou un moteur de train miniature ; tout le monde sait donc que les enroulements ne comportent pas une spire unique. Les noyaux portent toujours une bobine au nombre de spires important. Toutes ces spires parallèles ajoutent les champs qu'elles produisent, pour former un champ total beaucoup plus important (figure 6). On peut se demander pourquoi un champ plus intense est nécessaire, alors que le moteur peut tourner avec une seule spire. En fait, il ne suffit pas que le moteur tourne, il faut aussi qu'il dispose de la force nécessaire pour entraîner autre chose. La force que produit le moteur (le couple) est d'autant plus importante que sont intenses les champs magnétiques qui produisent les forces d'attraction et de répulsion. On pourrait aussi produire un champ plus intense en augmentant l'intensité du courant qui traverse l'enroulement. Cette équivalence est utilisée, entre autres, pour définir les bobines de relais : le champ magnétique nécessaire pour attirer la palette mobile ou le contact à lames souples est mesuré en **ampères-tours**. Une bobine de 1000 spires traversée par un courant de 1 milliampère produit le même effet qu'une « bobine » à une seule spire parcourue par un courant de 1 ampère. La construction d'un moteur résulte d'un compromis entre le nombre de spires qu'on peut loger sur le noyau, la section et la résistance du fil, l'intensité disponible, la tension d'alimentation...

La force produite par les champs magnétiques dépend autant du champ du rotor, que nous venons d'examiner, que de celui du stator, qu'il est beaucoup plus difficile de faire varier. Les aimants puissants et de grandes dimensions sont difficiles à fabriquer, et forcément coûteux. Ils sont remplacés dans les gros moteurs par des électro-aimants, d'autres enroulements. Le champ nécessaire

est obtenu aussi facilement que pour le rotor. Les deux éléments, rotor et stator, peuvent être connectés de trois façons différentes à la source d'alimentation. Ils peuvent être reliés en série, en parallèle, ou alimentés par des sources séparées (figure 7).

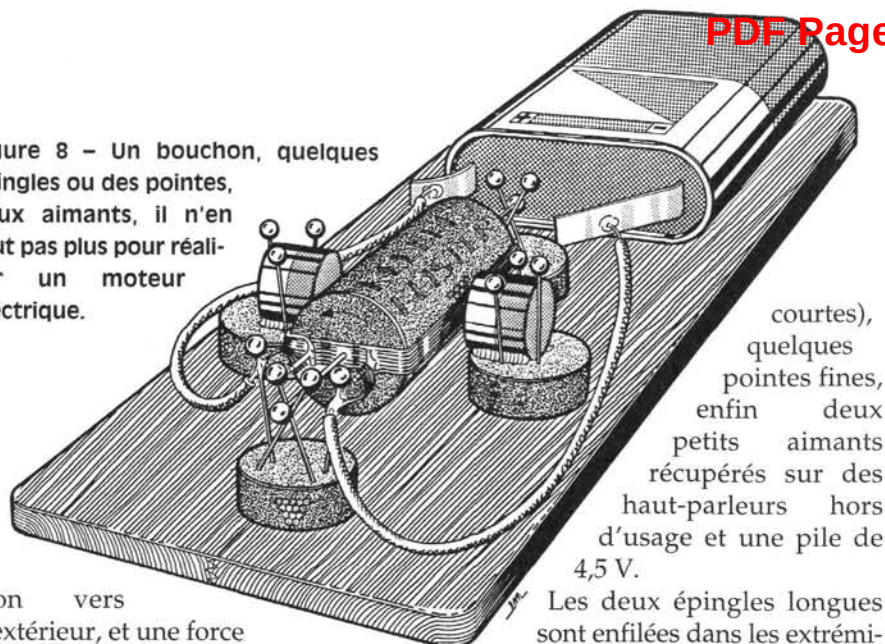
Les avantages et inconvénients de chaque type de moteur font préférer l'un ou l'autre pour chaque type d'application. Le moteur série peut être alimenté aussi bien en continu qu'en alternatif, d'où son autre appellation de **moteur universel** qui remonte à un temps (que les moins de vingt ans ne peuvent pas connaître) où l'électricité était distribuée par des sociétés indépendantes. Il n'y avait pas de normalisation sur la tension (110 V, 127 V, 220 V), sur la nature (continu ou alternatif), sur le nombre de phases... Il (le moteur série) présente l'autre avantage de développer son couple maximal à l'arrêt, ce qui le destine à la traction de véhicules. Pourquoi le moteur universel peut-il fonctionner aussi bien sur courant alternatif que sur courant continu ? Le sens de rotation est déterminé par le sens relatif des champs du stator et du rotor. Si nous inversons le sens du courant, ce sont les deux champs magnétiques qui changeront de sens en même temps ; comme leur sens relatif reste le même, le sens de rotation ne change pas. Pour vous en convaincre, essayez de changer sur la figure 2b le sens de toutes les flèches qui indiquent un champ, l'addition et la soustraction se produisent du même côté, la force a le même sens.

Le moteur parallèle ou *shunt* s'accommode aussi d'une tension alternative, mais supporte moins bien que le moteur série la commande de vitesse par variation de tension. Il fournit en revanche un couple plus important à vitesse élevée.

Le moteur à excitation séparée est comparable au moteur à aimant permanent : le champ du stator, même s'il est variable, est de sens constant. Cette caractéristique permet de commander le changement de sens de rotation en inversant la polarité de l'alimentation.

Revenons pour finir à la figure 4. Chaque fil est soumis à une force horizontale qui peut se décomposer, pour presque toutes les positions du cadre, en deux autres forces : une force radiale F_R dirigée de l'axe de rota-

Figure 8 - Un bouchon, quelques épingles ou des points, deux aimants, il n'en faut pas plus pour réaliser un moteur électrique.



tion vers l'extérieur, et une force perpendiculaire au rayon (tangente au cercle) Ft. Seule la force tangente au cercle sert à mettre le rotor en mouvement, la force radiale tend seulement à écarter les fils l'un de l'autre ou à les rapprocher (4b, 4d). La force utile, celle qui fait tourner le moteur, est maximale dans les positions a et e, minimale dans les positions b et d (nous supposons que le rotor n'est pas alimenté dans la position c). L'ingénieur qui conçoit un moteur cherchera donc à avoir toujours le rotor dans la position que nous appelons verticale. Malheureusement, ce n'est pas possible avec une seule bobine. C'est pourquoi la plupart des moteurs que vous pouvez trouver comportent un minimum de trois pôles, quelquefois huit ou douze. Cette multiplication des pôles permet d'avoir toujours au moins un pôle du rotor dans la position la plus favorable. Naturellement, le collecteur comporte un nombre de plages conséquent, les balais restant le plus souvent au nombre de deux. Plus le nombre de pôles est grand, plus régulière est la rotation, plus faciles les démarrages.

un petit moteur maison

Pour ne pas rester sur cette théorie un peu sèche, nous allons voir comment bricoler un petit moteur qui illustre le principe de fonctionnement. Il nous faut pour cela un bouchon, quelques mètres de fil de cuivre émaillé, quatre épingles à tête (deux longues et deux

courtes), quelques pointes fines, enfin deux petits aimants récupérés sur des haut-parleurs hors d'usage et une pile de 4,5 V.

Les deux épingles longues sont enfilées dans les extrémités du bouchon, parfaitement dans l'axe si possible : elles forment l'axe du rotor. Les deux épingles courtes sont fichées sur une des extrémités, sur un même diamètre de part et d'autre de l'axe ; elles forment le collecteur. Vient le moment de réaliser le bobinage, vingt spires de fil sur la longueur du bouchon, suivant deux génératrices du cylindre. Chaque extrémité du bobinage est soudée à l'une des épingles courtes (voir la figure 8).

Le rotor terminé vient reposer sur deux paliers constitués chacun de deux pointes croisées enfoncées en biais. Il doit pouvoir y tourner librement. Les aimants seront fixés par des pointes ou collés sur la planchette, de part et d'autre du rotor. L'un doit présenter son pôle sud au rotor, l'autre son pôle nord (dans cette position, ils s'attirent). Il reste à installer les balais du collecteur (ou charbons). Pour ce faire, vous dénuderez 1 centimètre des fils qui vont à la pile et vous les fixerez sur la planchette de telle façon qu'ils touchent les épingles quand la bobine est horizontale. C'est terminé, le moteur tourne. Il n'est d'ailleurs pas capable de faire grand chose d'autre.

Si vous voulez le faire fonctionner sur courant alternatif, il faut relier chaque extrémité de l'enroulement à une des épingles qui forment l'axe et non au col-

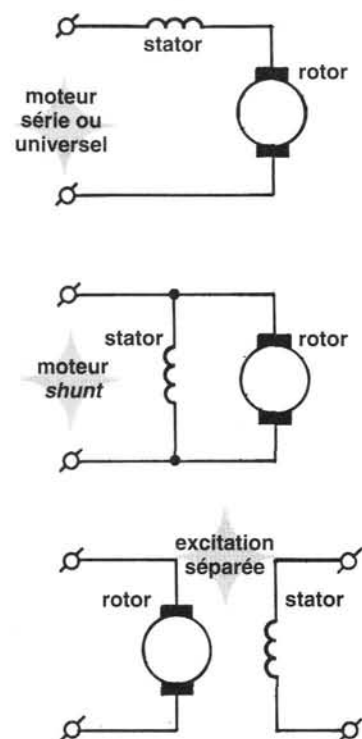
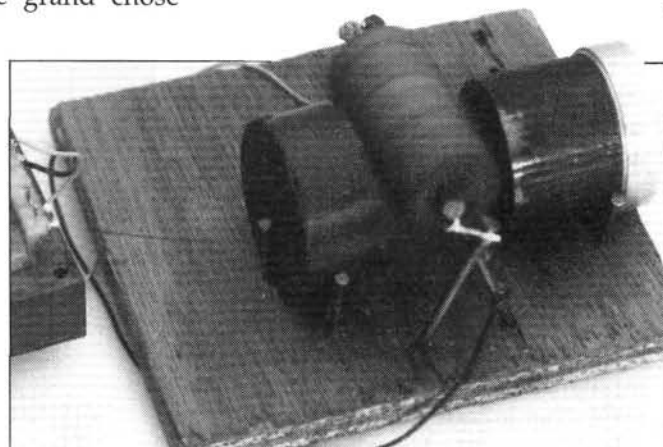
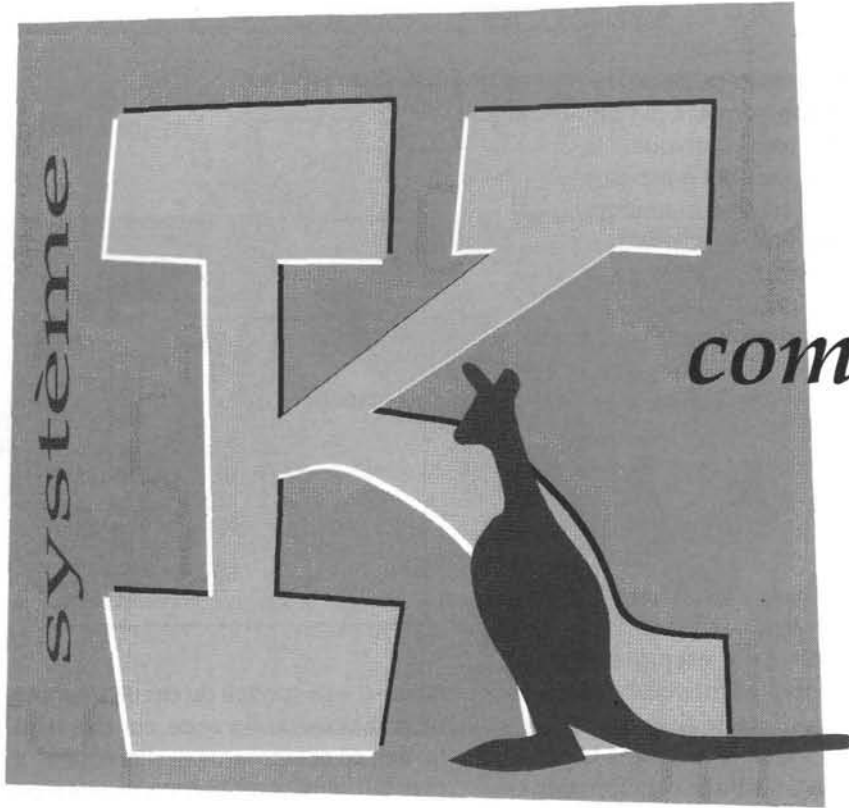


Figure 7 - Il existe trois façons de raccorder les enroulements des électroaimants d'un moteur électrique. Dans les deux premiers types, un changement de polarité n'a pas d'influence sur le sens de rotation ; ces moteurs peuvent donc être alimentés par un courant alternatif. Le troisième type, à excitation séparée, permet une inversion du sens de rotation par l'inversion de polarité sur le rotor seulement ou sur le stator seulement.

lecteur. L'alimentation se fera par les paliers, avec un petit transformateur de 6 V à 12 V. Attention, ce moteur ne démarre pas de lui-même, il faut le lancer à une vitesse suffisante pour qu'il se synchronise avec la fréquence du secteur. Ne laissez pas trop se prolonger la démonstration, car le moteur a tendance à chauffer. 87674

Figure 9 - Le moteur en pleine vitesse. Le transformateur montre que le moteur à bouchon peut fonctionner aussi avec un courant alternatif. C'est aussi un exemple de ce qu'il ne faut jamais faire : laisser à nu les connexions du secteur.





combinaisons

Les platines du système K commencent à se faire assez nombreuses.

Les expérimentateurs ingénieux ne se contenteront pas de les faire fonctionner individuellement : ce sont les pièces d'un jeu de construction, faites pour être utilisées ensemble.

Nous vous présentons ici quelques exemples de synergie multi-modulaire appliquée à la résolution des problèmes par la dynamique de groupe.

étage par étage

L'ingénieur expérimenté cherche, lors de la conception d'un circuit, à maintenir aussi faible que possible le prix de revient des composants. Ce souci conduit, par exemple, à donner deux fonctions à une même étage. Un transistor peut servir à la fois d'amplificateur et d'inverseur, un amplificateur opérationnel peut être à la fois amplificateur et sommateur, ou encore un haut-parleur peut avoir dans le même circuit sa fonction normale et celle de microphone... Il devient difficile pour un œil peu exercé de lire le schéma, de reconnaître le rôle de chaque composant, et de voir comment fonctionne

l'ensemble. Ce genre de schéma est à déconseiller au débutant et à l'amateur, c'est pourquoi nous préférons des montages « à plat », constitués d'une suite d'étages bien distincts. Un nombre limité de fonctions permet de réaliser à peu près n'importe quelle fonction complexe.

Si vous feuillotez les différents numéros d'ELEX, vous constatez que tous les systèmes d'alarme, par exemple, utilisent le même principe : un signal se propage de l'entrée à la sortie en ligne droite, sans dérivation ni rétro-action.

Les étages peuvent être ceux-ci :

- Étage d'entrée (capteur de chaleur, pression, lumière ou son)
- Traitement du signal d'entrée (comparateur, bascule RS, bascule astable)
- Étage de sortie (émission du signal d'alarme sonore ou lumineux : LED, vibreur piézo, haut-parleur, commande de relais).

L'assemblage de ces différents étages est visible sur la **figure 1**. Nous avons, pour les deuxième et troisième étages, des platines à relais, à bascule astable et bistable, amplificateur, etc. Il suffit d'une platine d'entrée pour que les

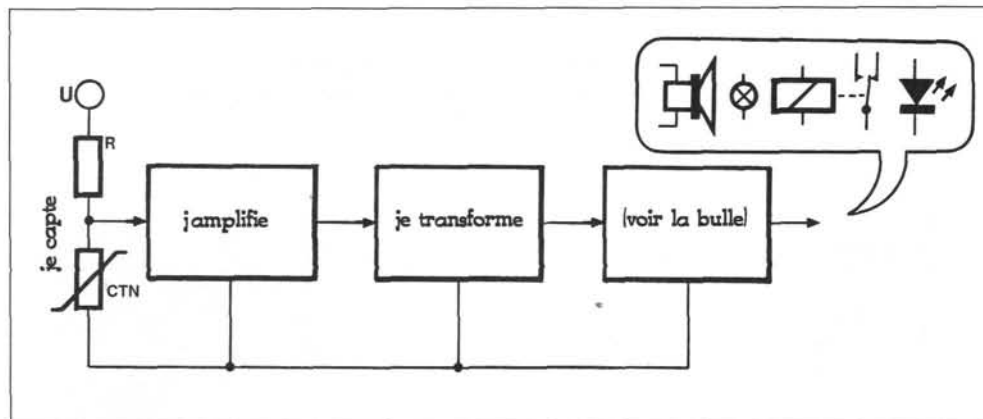


Figure 1 - Un schéma est parlant s'il est divisé en plusieurs étages. Beaucoup de circuits décrits dans ELEX (et ailleurs) se ramènent à quatre fonctions, remplies chacune par un groupe de composants. C'est l'étage d'entrée qui établit le lien entre le circuit et le monde extérieur. Nous n'avons pas encore présenté d'exemple d'étage d'entrée dans le système K. Ce sera fait avec les capteurs de température et de lumière qui permettront d'utiliser les platines existantes pour des applications pratiques.

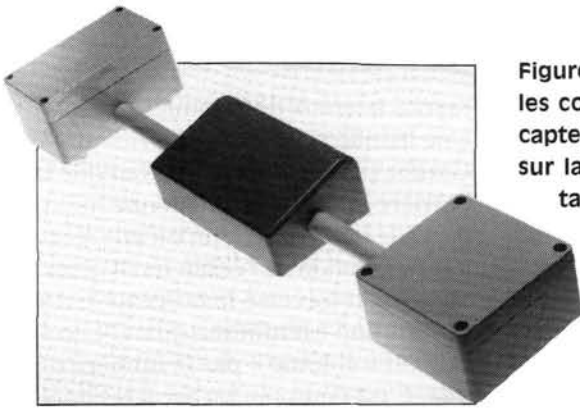


Figure 2 - Pour économiser la place et les composants, nous monterons les capteurs de température et de lumière sur la même platine, celle des résistances. Le montage et le dessin du circuit imprimé sont assez simples pour se passer de plan d'implantation.

platines existantes soient utilisables dans des applications pratiques intéressantes comme des barrières lumineuses ou des thermostats. Cette platine d'entrée peut être réalisée très simplement grâce à la platine à résistances d'ELEX n°40 de janvier 1992. Il suffit de remplacer une résistance par un capteur de température, de son ou de lumière.

chaud et lumière

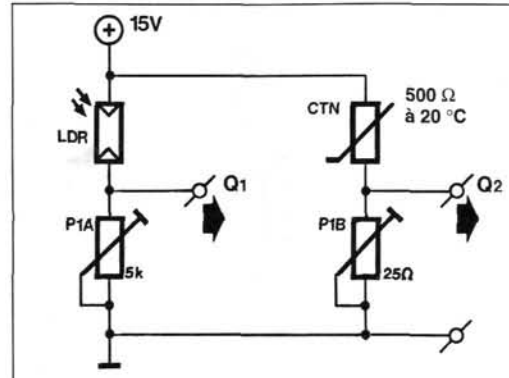
La température et la lumière sont les deux grandeurs physiques qui se convertissent le plus facilement en grandeurs électriques. Nos expériences feront appel à une photo-résistance (LDR pour *Light Dependent Resistor*) et à une thermistance à coefficient de température négatif (CTN). La figure 2 montre les symboles correspondants et les composants auxiliaires nécessaires.

La résistance d'une LDR diminue quand la lumière frappe à travers une fenêtre en plastique transparent ses pistes en zigzag ; elle diminue d'autant plus que l'éclairement est plus intense. Une thermistance à coefficient de température négatif (CTN) se comporte de la même manière quand sa température augmente. Une deuxième résistance, fixe, montée en série avec le capteur choisi, forme un diviseur de tension. La tension relevée au point nodal de ce diviseur (Q1 ou Q2) est une mesure de la température ou de l'éclairement, suivant le capteur. La variation de la résistance montée en série (par un potentiomètre) permet de fixer la plage de variation de la tension de telle façon qu'un transistor devienne conducteur (ou se bloque) dès qu'un seuil critique de lumière ou de température est atteint. Une tension minimale de 0,6 V sur la base est nécessaire pour que le transistor laisse passer un courant de son collecteur vers son émetteur (voir le

transistor en commutation dans ELEX n°44 page 37). L'association de la platine à diviseur de tension et de la platine à relais permet de commander celui-ci par un rayon lumineux ou une variation de température. Équipez une platine à résistances selon le schéma de la figure 2 et les premières expériences peuvent commencer.

barrière lumineuse

La barrière lumineuse utilise la platine « transistor plus lampe » du n°44 ; son entrée est reliée à la sortie de la platine du capteur comme le montre la figure 3. La source de courant est la partie positive de l'alimentation double de 15 V (ELEX n°37 d'octobre 1991). Pour le réglage de la barrière lumineuse, il faut un peu de doigté et une pièce obscure : en plus de la tension minimale de 0,6 V, il faut un certain courant pour alimenter la base du transistor. La plage de variation de la résistance de la LDR s'étend, entre la pleine



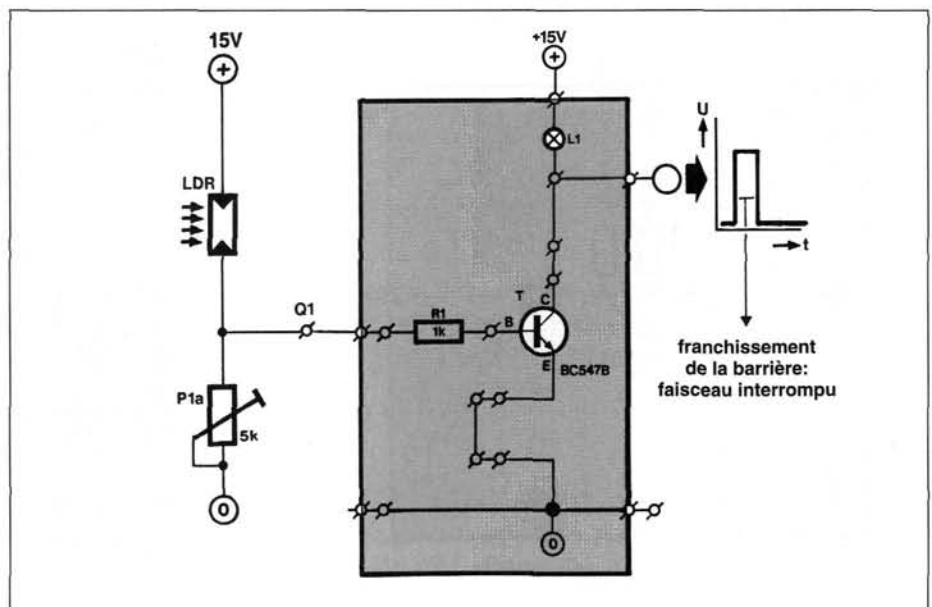
liste des composants

P1a = potentiomètre miniature 5 kΩ
P1a = potentiomètre miniature 25 Ω
1 LDR
1 CTN (500 Ω à 20°C)
platine à résistances système K

lumière et l'obscurité complète, de 1 kilohm à plusieurs mégohms. Le transistor a un gain de 250 environ ; pour qu'un courant de 50 mA traverse la charge, il suffit donc d'un courant de base de 0,2 mA, obtenu pour une résistance de 40 kΩ de la photorésistance. Cette valeur donne une tension supérieure à 0,6 V si le potentiomètre de 5 kΩ est en position médiane (2,5 kΩ).

Autre exemple de calcul pour le choix de la valeur de P1a : la résistance de la LDR est supérieure à 100 kΩ dans l'obscurité. La chute de tension sur P1a

Figure 3 - Le transistor et la lampe se chargent de l'amplification. Comme le transistor se comporte en même temps en interrupteur, la lampe ne peut être qu'allumée ou éteinte. La tension présente entre le collecteur et la masse peut être utilisée comme un signal logique pour commander un ou plusieurs autres étages.



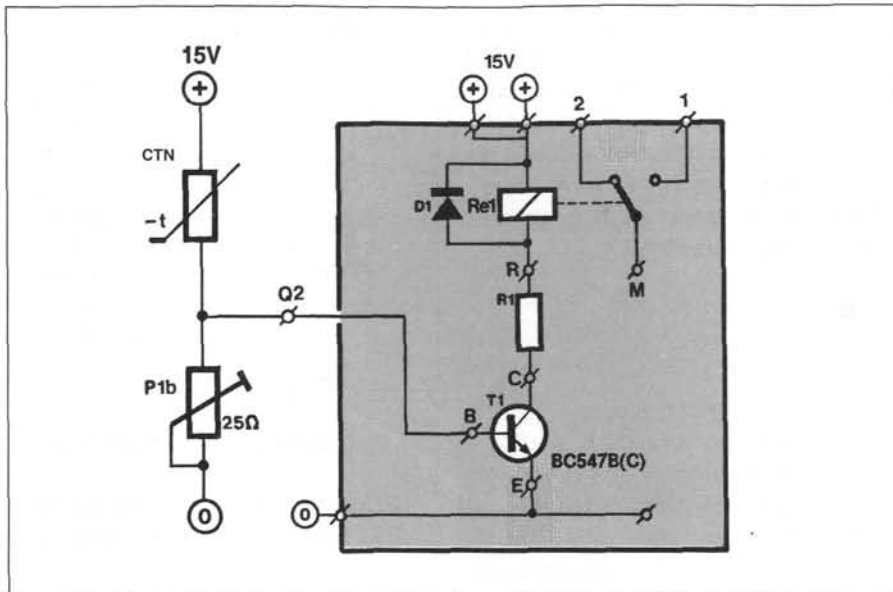


Figure 4 - Un circuit thermostatique complet. Si la température augmente trop, un contact de relais se ferme (ou s'ouvre, suivant le contact choisi). Ce contact peut commander un radiateur électrique, par exemple.

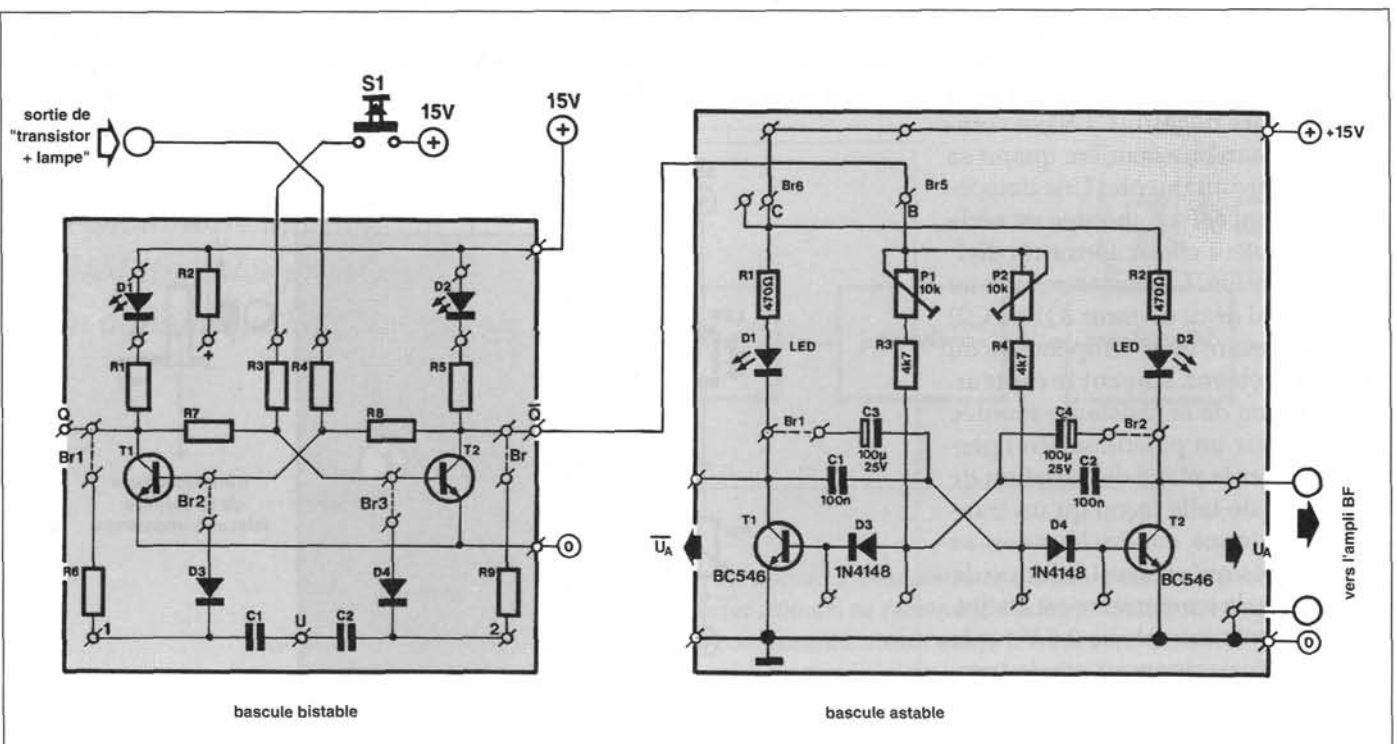
rayons latéraux de la lumière du jour. Une lentille peut rendre service aussi. Gardez présent à l'esprit le fait que la barrière lumineuse fonctionne moins bien à la lumière du jour, car elle réagit à la diminution de l'éclairement quand quelqu'un traverse le faisceau. Cette diminution est minime quand la cellule est « éblouie » par la lumière du jour. C'est pourquoi les barrières lumineuses de fabrication industrielle utilisent un principe différent : la lumière est modulée et un filtre (électronique) incorporé au « récepteur » ne laisse passer que les variations correspondant à la fréquence d'émission. De plus, la lumière utilisée est de l'infrarouge et la lumière du jour est rejetée par un filtre optique. L'étude de ces perfectionnements nous emmènerait trop loin du domaine de l'expérimentation auquel nous voulons nous limiter.

thermostat

Notre deuxième exemple traite de la commande d'un relais par la chaleur, plus précisément par l'élévation de température. Ce montage permet de réaliser un thermostat domestique ou une alarme d'incendie. Le capteur de température, constitué par une thermistance et un potentiomètre, ressemble fort à notre barrière lumineuse : il comporte aussi une résistance variable en fonction d'un phénomène extérieur, en série avec un

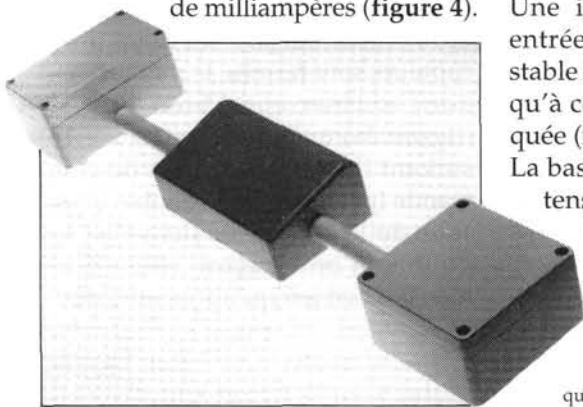
(2,5 kΩ) est alors environ le quarantième de la tension d'alimentation, soit 0,3 V, nettement en dessous du seuil de 0,6 V. La résistance de la LDR tombe en-dessous de 5 kΩ quand elle est éclairée, donc la tension au point Q1 passe au tiers de la tension d'alimentation ($2,5 = (R_{POT} + R_{LDR})/3$), soit 5 V. Comme le seuil de la jonction base-émetteur du transistor est de 0,6 V, c'est à cette valeur que se limitera la tension au point Q, le courant à travers la LDR s'établissant en conséquence. La pièce obscure est nécessaire pour éviter que la lumière du jour perturbe

la mesure. Le réglage se fait avec une lampe de poche placée à quelques centimètres de la LDR. Si le montage ne réagit pas au rayon lumineux, vérifiez la tension sur la base du transistor : elle doit varier au-dessus et en dessous du seuil de 0,6 V suivant que la LDR est éclairée ou non. Vous pouvez essayer d'éloigner la source de lumière, en retouchant au fur et à mesure le réglage du potentiomètre pour obtenir la sensibilité maximale. Pour obtenir la meilleure efficacité du montage, il faut monter la photo-résistance dans un tube en carton qui arrêtera les



un relais, et puis ?

potentiomètre (figure 2). La tension aux bornes du potentiomètre augmente quand la résistance du capteur diminue. Les valeurs sont choisies de telle façon que la tension soit juste en dessous du seuil de 0,6 V à la température ambiante. Pour cela il faut que P1b ait une valeur de 15 Ω environ (vers le milieu d'un potentiomètre de 25 Ω). La valeur de la thermistance diminue de 6% par degré Celsius. Cette réponse rapide permet de construire une alarme d'incendie qui réagit à une élévation de température de 5 à 6°C. Dans notre exemple, nous utilisons une résistance CTN de 500 Ω à 20°C. Si P1b est réglé à 15 Ω , la tension de la base est de 0,44 V. Une élévation de température de 5° provoque une diminution de 30% de la résistance, ce qui donne une tension de 0,62 V sur la base. L'augmentation de la valeur de P1b permet d'augmenter encore cet intervalle, jusqu'à ce que le seuil de commutation se trouve exactement au milieu. Le gain du transistor est indifférent car la faible valeur de la thermistance laisse circuler vers la base un courant de plusieurs dizaines de milliampères (figure 4).



L'appareil commandé par le relais dépend de l'utilisation pratique envisagée. S'il ne s'agit que d'expériences, un dispositif « son ou lumière » est suffisant. Il peut s'agir d'un sirène, d'une sonnette ou d'un générateur d'éclairs. Si vous utilisez des lampes alimentées par le secteur, veillez à ne pas dépasser les limites de tension et de courant du contact du relais. Un radiateur électrique de 2000 W* pompe une dizaine d'ampères, ce qui n'est pas à la portée du premier relais miniature venu.

La commande du relais change aussi suivant la nature du montage. Dans le cas du thermostat, le relais reste excité aussi longtemps que la thermistance est chaude. Le radiateur reste coupé jusqu'à ce que la température ait diminué, ou le gyrophare alimenté jusqu'à l'arrivée des pompiers.

Il en va autrement pour la barrière lumineuse : le passage d'une personne ou d'un objet dans le faisceau qui illumine la LDR est très bref, peut-être trop bref pour que le signal soit perceptible. La solution est à portée de main : la platine à bascule bistable du n°47, montée comme sur la figure 5. Une impulsion brève à l'une des entrées donne un état logique haut stable à la sortie. Cet état persiste jusqu'à ce qu'une impulsion soit appliquée (manuellement) à l'autre entrée. La bascule réagit à la remontée de la tension, quand le passage du faisceau lumineux est libéré après

le passage d'un client. Le relais actionne la sonnette ; quand le vendeur quitte le labo où il était en train de bricoler, il appuie sur le poussoir S1 pour arrêter la sonnette. La sonnette peut être remplacée par le multivibrateur astable du n°46, suivi par un amplificateur BF et un haut-parleur. Si un deuxième client entre, on peut arrêter le signal par un deuxième poussoir connecté en parallèle sur le premier. Il est possible aussi d'utiliser une bascule monostable pour faire suivre chaque impulsion d'entrée, si brève soit-elle, d'une impulsion de sortie de durée constante.

Ce dernier article ne prétendait pas exposer toutes les possibilités des platines du système K, mais vous montrer que vous pouvez concevoir et construire vous-mêmes les montages que vous imaginez. N'importe quelle cause peut produire n'importe quel effet, il suffit de trouver le bon capteur. Le reste est une question de logique pour l'assemblage des différentes fonctions. Une fois le fonctionnement correct obtenu, vous regrouperez en un seul les différents schémas utilisés, vous grouperez éventuellement différentes fonctions pour simplifier l'ensemble et vous pourrez passer au dessin de votre circuit imprimé. 86687

Les lois de la thermodynamique étant ce qu'elles sont, il faut consommer au moins 10000 W d'énergie issue du pétrole, du charbon ou de la fission de l'atome pour fournir 2500 W à la prise. En brûlant autant de pétrole ou de gaz chez soi, on récupère 9000 W de chaleur. Dommage que le gaz soit facturé en kilowatt !

MAGNETIC-FRANCE

Circuits intégrés, Analogiques, Régulateurs intégrés, Interfaces, Micro-Processeurs, Mémoires RAM Dynamiques Statiques, EPROM et EEPROM, Quartz, Bobinage, Semi-Conducteurs Transforiques, Filtres, Ligne à retard, Leds, Supports de CI, Ponts, Opto-Electronique, etc.
Et de nombreux KITS.

Bon à découper pour recevoir le catalogue général

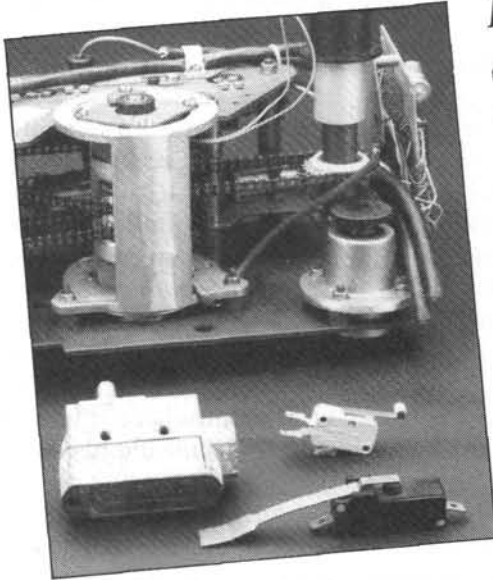
Nom

Adresse

Envoi : Franco 35 F - Vendu également au magasin

11, Place de la Nation, 75011 PARIS **43793988**

Télex 216 328 F - Ouvert de 9 h 30 à 12 h et de 14 h à 19 h
Fermé le Lundi.



Les interrupteurs de fin de course, micro-rupteurs ou "micro-switches", sont indispensables pour éviter qu'un moteur dépasse les limites d'un dispositif mécanique télécommandé comme un bras de robot, une antenne orientable, une porte de garage, etc.

fin de course pour arrêter à temps

Nous vous avons proposé dans *ele* numéro 46, page 43, un système d'indication à distance de la position d'un axe télécommandé. Il s'agissait là d'un système de surveillance ; comme il n'est pas exclu que le manipulateur s'endorme en sursaut ou que le dispositif de commande défaille, un système de sécurité n'est pas superflu. Quel que soit l'objet déplacé par le moteur, il est évident qu'il y a des bornes physiques à son déplacement. Comme le dit le proverbe en usage chez les géomètres, plagés par un radioteur du dimanche matin, « passé les bornes, il n'y a plus de limite ». Prenons l'exemple du portail ou de la porte de garage télécommandés. La porte, une fois ouverte ou fermée, ne peut pas et ne doit pas aller plus loin, à moins de casser le mur ou d'arracher les gonds. Dans la suite, nous appellerons moteur tout l'équipage mobile.

Le but des dispositifs de fin de course est de garantir que le moteur s'arrêtera assez tôt pour éviter de détruire des organes fragiles. Les interrupteurs

de fin de course sont actionnés par un ergot quand le moteur arrive au bout du déplacement autorisé. Leur ouverture interrompt le circuit électrique du moteur et l'empêche de continuer de tourner (voir la figure 1). Il faut cependant que le moteur puisse repartir dans l'autre sens, ce qui donne lieu à des montages relativement compliqués pour les moteurs à courant alternatif. Pour le courant continu au contraire, une paire de diodes suffit le plus souvent.

les diodes

Le principe des interrupteurs de fin de course est repris par la figure 2. Le circuit comporte le moteur M, une source de tension (symbolisée par une simple pile), deux diodes (D1 et D2) et enfin les inter-

rupteurs de fin de course S1 et S2. Il s'agit d'interrupteurs à ouverture, c'est-à-dire fermés au repos. En fonctionnement normal, les deux interrupteurs sont fermés, le moteur peut donc tourner aussi bien à droite (figure 2a) qu'à gauche (figure 2b) suivant l'état des organes de commande (par exemple le double inverseur de la figure 3). Jusqu'ici les diodes n'ont aucun rôle. Elles n'entrent en jeu qu'après un arrêt.

arrêt

Le moteur a atteint l'extrémité de sa course en tournant à gauche : l'interrupteur S1 est actionné. Le courant, qui circulait comme le montre la flèche de la figure 2a, est interrompu. Il ne peut plus circuler à gauche puisque l'interrupteur S1 est ouvert, ni à droite puisque la diode D2 est polarisée en inverse. Par conséquent le moteur est arrêté. Il est possible de faire tourner le moteur dans l'autre sens en changeant la polarité de l'alimentation : nous

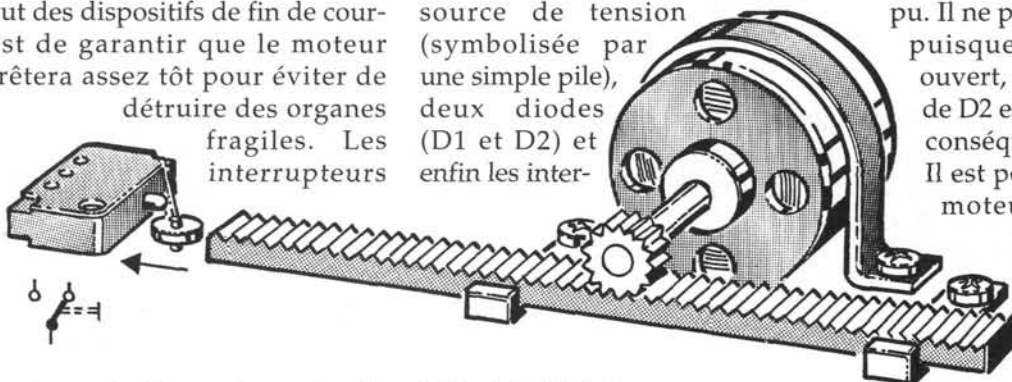
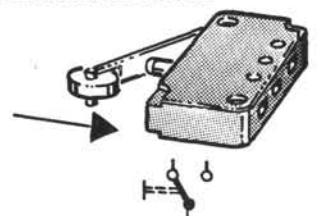


Figure 1 - L'un des interrupteurs est actionné à l'extrémité de la course, pour signaler que le moteur doit s'arrêter.



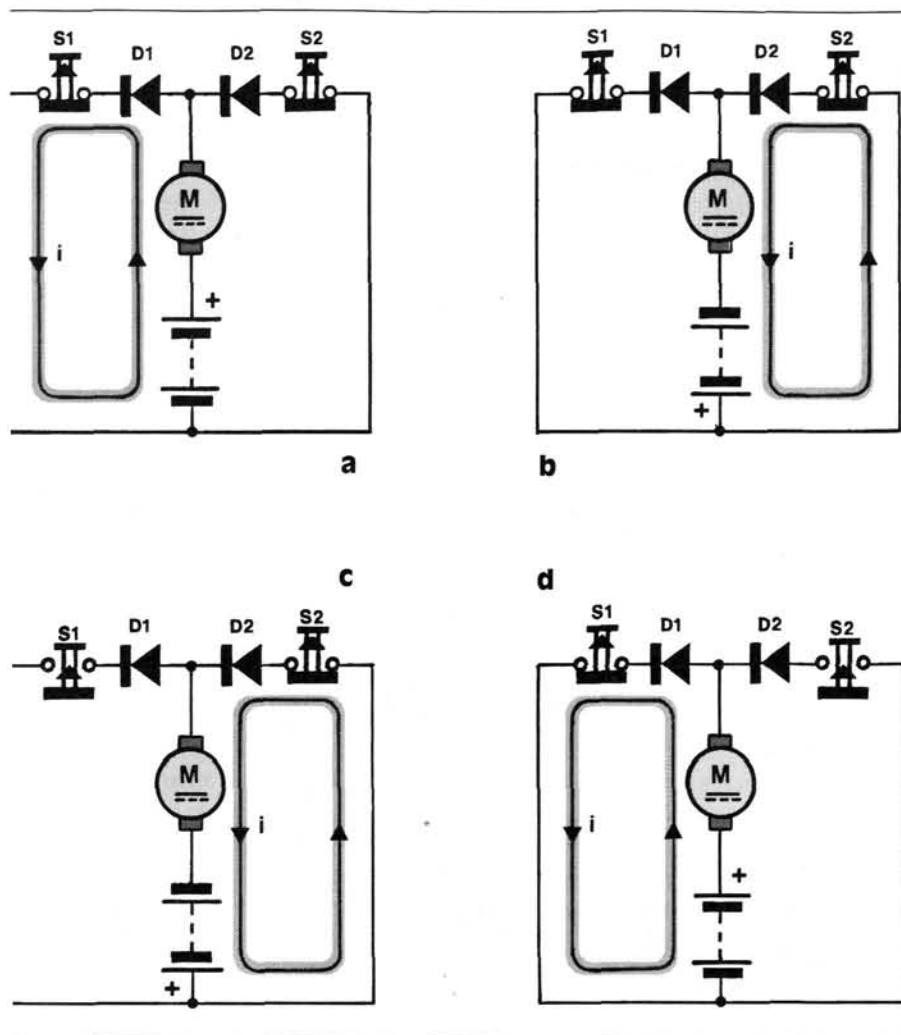


Figure 2 - Les deux interrupteurs et les diodes nécessaires pour assurer la sécurité mécanique. Les flèches indiquent le sens du courant ; elles montrent que la rotation est possible dans les deux sens si aucun interrupteur n'est ouvert. Que l'un des interrupteurs s'ouvre, alors le mouvement n'est plus possible que dans un seul sens.

quelques mots sur les moteurs

Ce qui précède suppose que le sens de rotation des moteurs peut être changé par inversion de la polarité de la tension d'alimentation, sans autre forme de procès. Vous pouvez vous demander avec raison si c'est possible avec tous les moteurs. Voyez à ce sujet l'article *Les moteurs électriques* page 16 de ce numéro. Ce qui est sûr, c'est que la plupart des petits moteurs utilisés en télécommande, sur les perceuses miniatures, dans les jouets, peuvent fonctionner ainsi : ils ont tous un inducteur, ou stator, à aimant permanent. C'est vrai aussi des moteurs fournis avec les jeux de construction de différentes marques.

Si vous voulez utiliser, pour disposer de plus de puissance, des moteurs à stator bobiné, du genre moteur d'essuie-glace, c'est un peu plus compliqué : il faudra alimenter séparément le rotor et le stator pour inverser le sens de rotation. Reportez-vous à l'article cité au début pour les détails sur l'alimentation du rotor ou du stator par un pont redresseur.

87669

nous trouvons dans le cas de la figure 2c. Le courant circule maintenant à travers S2 et D2, polarisée en sens direct.

Quand le moteur arrivera à l'extrémité opposée de sa course, il actionnera S2, ce qui nous placera dans la situation de la figure 2d : le moteur est arrêté et il faut inverser la polarité de l'alimentation pour le faire redémarrer dans l'autre sens. Dans les deux cas, le moteur ne peut pas causer de dégâts en dépassant les limites prévues, mais il peut repartir dans l'autre sens, grâce à deux simples diodes.

pratique

Comment réaliser en pratique ces interrupteurs de fin de course ? Il n'y a ni recettes ni règles fixes, car tout dépend de l'application. Nous pouvons cependant vous donner quelques trucs généraux. La première solution est d'utiliser des *micro-switches*, en français micro-interrupteurs, destinés spécialement à cet

usage et actionnés par des ergots disposés judicieusement.

Il est possible aussi de concevoir une barrière lumineuse avec une LED et un phototransistor : quand le moteur atteint la fin de sa course, une plaque interrompt le faisceau lumineux, le phototransistor commande un relais qui coupe le circuit électrique.

Enfin, vous pouvez alimenter le moteur par des frotteurs et des rails conducteurs. Arrivé en fin de course, le frotteur trouvera une plage isolante sur le rail conducteur. L'exécution pratique peut faire appel simplement à une plaque d'isolant cuivré pour circuit imprimé, débarrassée de son cuivre à certains endroits.

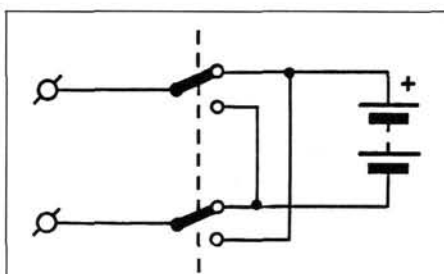
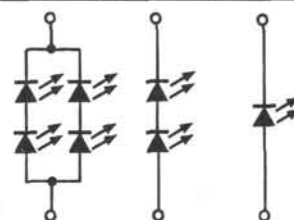


Figure 3 - Ce montage de double inverseur permet de changer le sens de rotation d'un moteur à courant continu. Il n'y a pas de position médiane car l'arrêt ne doit se faire qu'en fin de course, où il est assuré par les interrupteurs ad hoc.



décimal	Q2	Q1	Q0
9	0	0	1
10	0	1	0
11	0	1	1
12	1	0	0
13	1	0	1
14	1	1	0



dé électronique

Le meilleur moyen de faire connaissance avec un circuit intégré est d'en étudier une utilisation. Un numéro déjà ancien d'ELEX (en juin 1989) décrivait un dé électronique basé sur un compteur plutôt simple d'emploi, le 4017. Nous utilisons ici un circuit un peu plus compliqué, le 4029, compteur décompteur programmable.

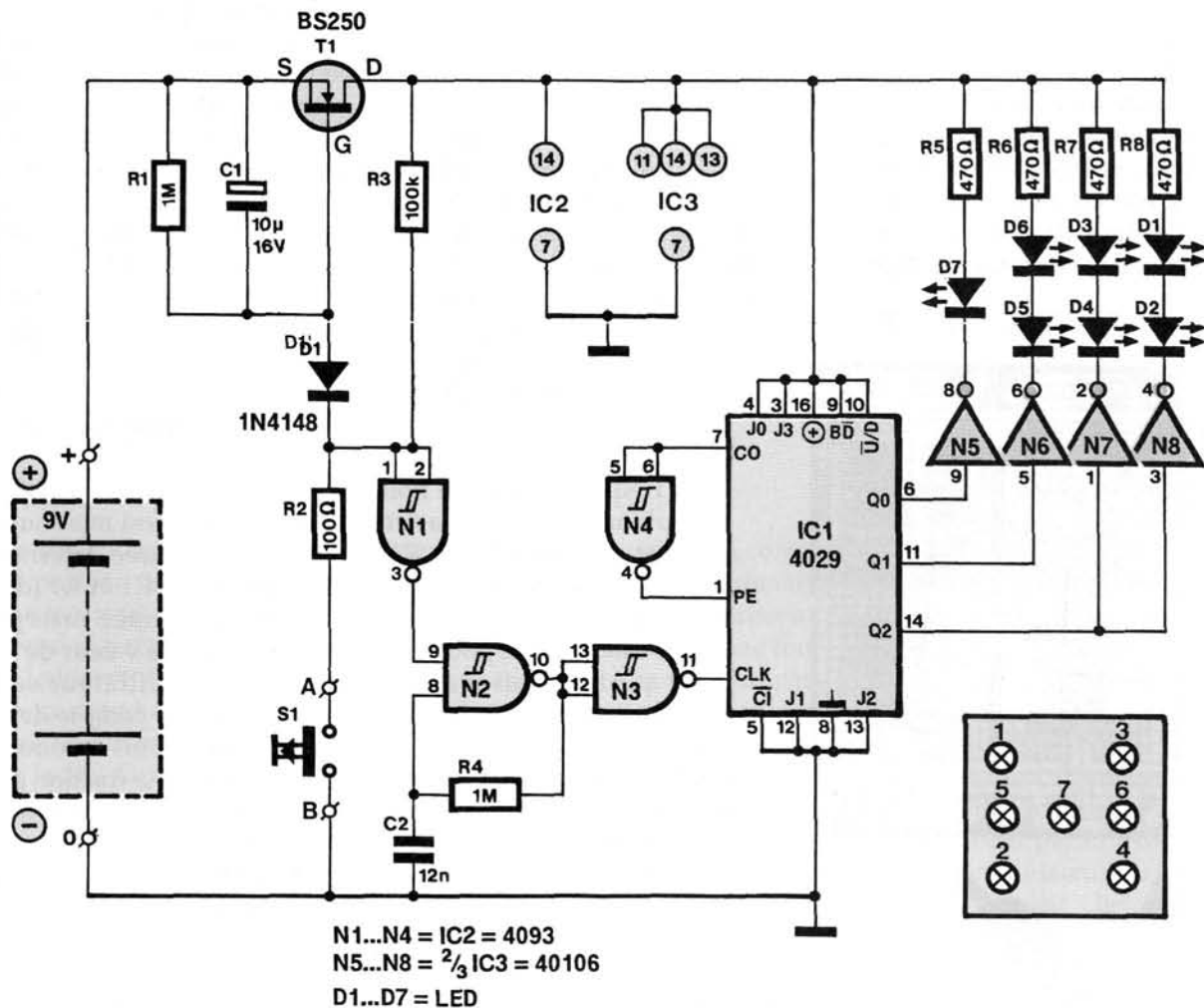
Le cœur du dé est donc un compteur binaire dont trois des quatre sorties sont utilisées (figure 1). Ces sorties commandent, par l'intermédiaire d'inverseurs-tampons, trois séries de LED qui représentent les points du dé. Seules trois sorties sont câblées, puisque pour éclairer de un à six points, disposés comme ils le sont sur un dé ordinaire, il ne faut que trois ensembles de LED comprenant respectivement un, deux et quatre de ces composants comme le montre le tableau 1. Nous voilà mal partis puisque de cette façon, le dé électronique n'a pas six mais huit faces : une face nulle, représentée par l'extinction de toutes les LED, et une face sept, où les trois ensembles brillent simultanément, donnant donc sept points. Quelques particularités du compteur utilisé permettent d'exclure ces situations extrêmes.

Pour commencer, il dispose d'une sortie de retenue Co (carry out), validée par un état bas (0) sur l'entrée Ci (carry in). Cette entrée est donc portée au potentiel de la masse pour nous permettre d'utiliser les effets

de la sortie Co. Quels sont-ils, et d'abord comment compte le 4029, dans le cas le plus simple ? Au premier front montant sur son entrée d'horloge, il met à 1 sa sortie Q0. Au second, cette sortie passe à 0 et Q1 passe à un ($1 + 1 = 10$). Au troisième, la sortie Q1 reste à 1 et la sortie Q0 repasse à 1 ($10 + 1 = 11$). Ces trois premières impulsions d'horloge se traduisent donc en sortie du compteur par 11, c'est-à-dire 1 sur Q0 et 1 sur Q1, mais ce n'est pas fini. Le circuit, puisqu'il a quatre sorties, peut aller jusqu'à 1111. Si nous comptons les fronts montants d'horloge en numération décimale ("normalement"), ce 1111 binaire est obtenu lors du quinzième front montant. Aussitôt la sortie Co (carry out, retenue) manifeste que toutes les sorties sont à 1, par un état bas (0) qui dure jusqu'au seizième front montant. Si cette sortie Co est reliée à l'entrée d'horloge d'un second compteur, lors du seizième front montant, toutes les sorties du premier compteur passent à 0 et la sortie Q0' du second compteur à 1 : le compteur pose "zéro" et retient "un". Nous

donc au niveau logique 1, les entrées J1 et J2 sont à celui de la masse, qui est un 0 logique. La sortie Q3 est donc à 1 et qu'avons-nous maintenant sur les autres sorties ? Au lieu de "111", la première ligne du **tableau 1**, imposée par le programme écrit sur les entrées J3 à J0. Pour les LED, le compte est bon : les données présentes sur l'ensemble des sorties Q2, Q1 et Q0 sont successivement 001, 010, 011, 100, 101, 110, et retour à la ligne départ programmée, 001. Le compteur lui, tient compte - c'est le cas de le dire ! - du 1 de Q3 et compte de 1001, soit 9 en numération décimale, jusqu'à 1110, soit 14, après quoi il revient à 9. Les sept LED ne peuvent donc pas briller toutes ensemble puisqu'au moment où les sorties du compteur seraient toutes à 1, la sortie de retenue Co commande l'entrée PE qui les **REMET** non pas à 0, mais à la valeur Programmée sur les entrées J0 à J3 dites *JAM* (de blocage). Cette entrée, active au niveau logique 1, est appelée *preset enable* ou entrée de validation de la programmation.

Figure 1 - Un oscillateur donne sa cadence à un compteur dont les sorties alimentent trois ensembles de LED, pour simuler un dé à jouer. Il est aussi possible de se servir de ce circuit pour tester ses réflexes, en diminuant la "vitesse de rotation" du dé, donc la fréquence de comptage : il suffit d'augmenter la capacité de C2 ou la résistance de R4 pour « Abolir le hasard ». Voyez-vous l'économiseur de pile ? Sa constante de temps augmente ou diminue si l'on modifie les valeurs de C1 et R1.



liste des composants

R1, R4 = 1 M Ω *
 R2 = 100 Ω
 R3 = 100 k Ω
 R5 à R8 = 470 Ω
 C1 = 10 μ F/16 V*
 C2 = 12 nF*

D1' = 1N4148
 T1 = BS 250
 D1 à D7 = LED rouge

IC1 = 4029 (compteur/décompteur, synchrone, programmable, binaire/par décade)
 IC2 = 4093 (quadruple opérateur ET-NON à 2 entrées et à trigger de Schmitt)
 IC3 = 40106 (sextuple inverseur à trigger de Schmitt)

S1 = bouton poussoir ouvert au repos
 platine d'expérimentation de format 2

*Voir le texte

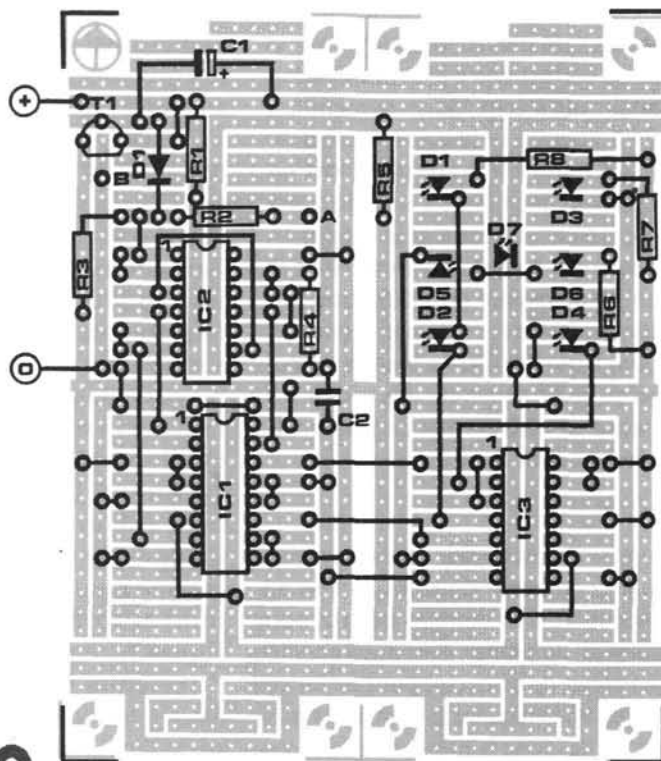


Figure 2 - Disposez les LED comme les points d'un dé ; l'éventuel coffret doit être prévu en conséquence.

D'autres entrées ont aussi leur importance, voyons-les tour à tour. Pour $\overline{Ci}$ (*carry in*, pas de problème, c'est l'entrée, active au niveau logique bas, qui permet à la sortie de retenue Co de fonctionner. L'entrée de comptage, ou de décomptage (*Up/Down*), permet de compter à rebours si elle est au potentiel de la masse. Comme elle est au niveau logique haut, le compteur compte.

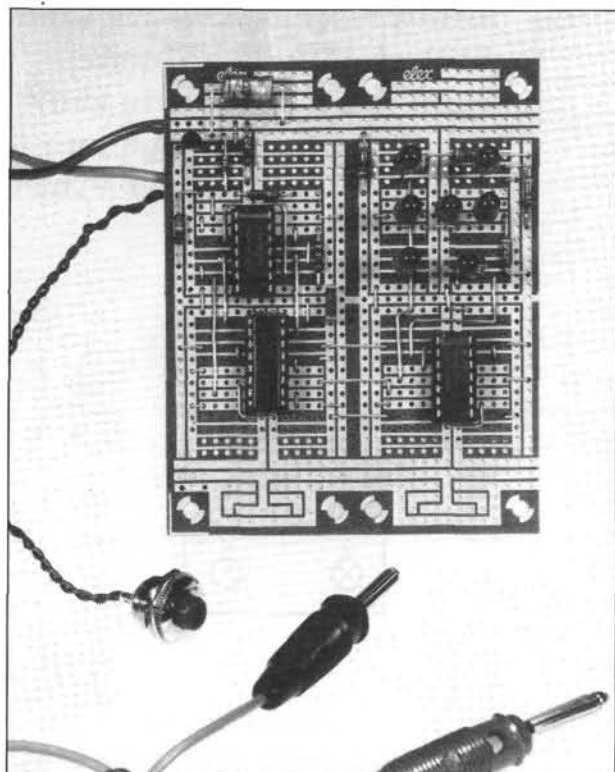
Ensuite, l'entrée accessible sur la broche 9 portée à la masse permet, dans d'autres applications, un comptage par décade : de 0 à 9 (de 0000 à 1001), ou décomptage de 9 à 0. Dans ce cas, la sortie Co passe à 0 au huitième front montant de l'horloge si l'entrée $\overline{Ci}$ ne l'inhibe pas (si elle est donc au niveau logique 0). L'entrée CLK est bien sûr une entrée de *clock*, d'horloge, horloge dont il est d'ailleurs temps de parler. L'oscillateur, qui donne son rythme au compteur, est bâti autour d'un des opérateurs, *trigger* ET-NON, que contient le 4093. Sa fréquence dépend non seulement de la capacité C2 et de la résistance R4, mais aussi de la marque du circuit utilisé. Elle doit être assez élevée pour éviter que l'œil puisse suivre l'allumage et l'extinction des LED : il doit les voir briller toutes en même temps. Pour augmenter le cas échéant cette fréquence, il faut diminuer C2 ou

R4. Voyons maintenant l'alimentation : une pile de 9 V suffit pour un circuit semblable et tient assez longtemps si l'on n'insiste pas lourdement sur le poussoir S1 (ce qui est aussi utile que de lancer le dé par la fenêtre). Elle durera d'autant plus longtemps qu'un économiseur de pile est prévu. Le FET T1, après un laps de temps défini par R1 et C1, ouvre le circuit et la pile ne débite plus. Il ne reste plus maintenant qu'à trouver à la platine un coffret digne du jeu.

fonctionnement

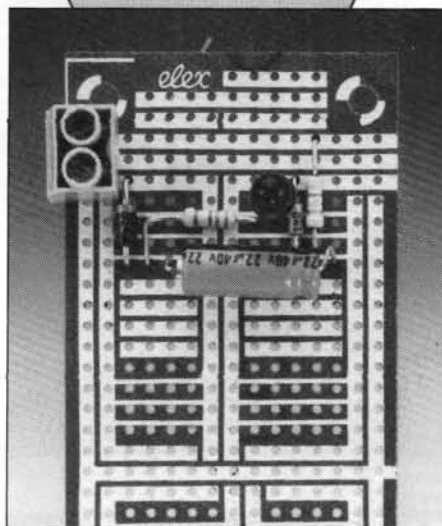
Appuyez sur la touche S1 pour jeter le dé. Le circuit est mis sous tension pendant une durée déterminée par les composants C1 et R1 (durée qui peut être augmentée, mais pas indéfiniment, avec la valeur de ces composants). L'oscillateur se met en branle et le 4029 compte de 1 à 6 (de 9 à 14 en fait), puis recommence. Il ne lui faut qu'une fraction de seconde. Dès que S1 est relâché, le dé a fini de rouler et les LED qui correspondent à la face tournée vers le haut sont allumées.

86693



led clignotante sous 220 V

L'interrupteur qui permet de faire disparaître l'obscurité s'y perd souvent. Il faut parfois tâter la moitié du mur avant de mettre le doigt dessus. Pour peu que l'on soit ivre de sommeil ou dans une pièce que l'on ne connaît pas, le maudit bouton reste introuvable. On éviterait cette mésaventure si une petite lampe en signalait l'emplacement.*



point lumineux dans l'obscurité

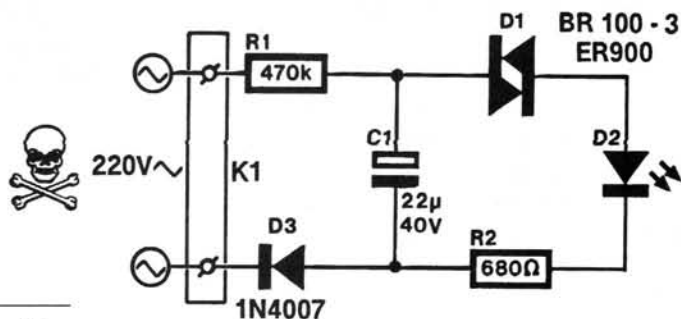
étrange diode

Le problème ainsi posé n'est pas pour autant résolu puisque une telle lampe doit satisfaire deux exigences : comme elle reste allumée nuit et jour, sa consommation sera minime ; elle ne doit pas trop éclairer, pour ne pas perturber le sommeil de ceux qu'elle veille, si elle est par exemple installée dans une chambre à coucher. Une LED répond à toutes ces conditions. Si en plus elle clignote, sa consommation sera encore réduite et son repérage plus facile : un clignotant attire plus facilement l'attention qu'un objet qui éclaire continûment.

Le nombre de composants nécessaires à la fabrication d'un tel circuit n'est pas excessif, six en tout, comme vous le constatez sur la **figure 1**. L'un d'entre eux retient plus particulièrement l'attention : il est un peu étrange et c'est lui qui commande la manœuvre. Nous voulons parler de D1, un DIAC. C'est une sorte de double diode, comme le dit son nom (*Diode Alternating Current* qu'on appelle aussi *bilateral trigger diode*), ou de transistor symétrique, comme le montrerait sa constitution (NPN ou PNP) à deux jonctions identiques. On peut aussi le comparer

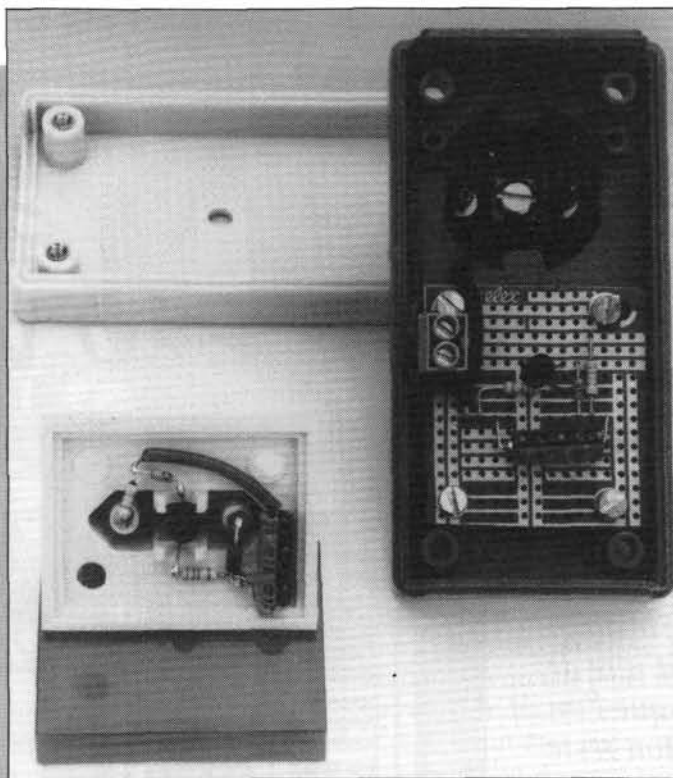
à un interrupteur ne laissant passer le courant que lorsque la tension entre ses contacts dépasse une certaine valeur. Cette valeur, la tension de déclenchement, peut être de 32 V (le plus souvent), 40 V ou 60 V suivant les DIACS, qui restent conducteurs une fois que l'avalanche est déclenchée. Si l'avalanche cesse, c'est-à-dire lorsque le courant descend au-dessous d'un seuil, dit courant de maintien, "l'interrupteur" s'ouvre. C'est un peu comme une porte battante dans le couloir d'un train : le premier voyageur à la franchir doit mobiliser ses forces pour

Figure 1 – Le composant bizarre noté D1 est une sorte de diode double ou de transistor dont les deux jonctions seraient symétriques : un DIAC.



*Dans le meilleur des cas, voir Théodore chercher des allumettes du cher Courteline.

diac et led



elex-abc

diac

Un DIAC est un semiconducteur à avalanche contrôlée qui ne conduit que lorsque la tension à ses bornes dépasse une certaine valeur, dite tension de déclenchement (environ 30 V pour les plus courants). Le DIAC ne conduit que si le courant reste supérieur à un seuil dit courant de maintien.

LED

Une Light Emitting Diode, en français, diode électroluminescente, est une diode enfermée dans une capsule transparente, qui émet de la lumière lorsqu'elle est passante. Cette lumière peut être orange, rouge, jaune ou verte suivant le type (les diodes bleues sont assez peu répandues). L'émission n'a lieu que si l'intensité du courant qui traverse la LED est suffisante. Cette intensité, comprise entre 15 mA et 25 mA, peut descendre à 3 mA pour certaines fabrications spéciales. Une LED éclaire beaucoup moins qu'une lampe à incandescence mais sa durée de vie est beaucoup plus grande. Il faut noter que sa tenue en tension inverse est beaucoup plus faible que celle des diodes ordinaires.

l'ouvrir et si les suivants (ou leurs bagages) sont assez proches les uns des autres, la porte reste ouverte ; si leur flux cesse elle se ferme. Comme ces portes s'ouvrent dans les deux sens, la comparaison est presque parfaite. Un DIAC est le plus souvent utilisé pour commander la gâchette d'un thyristor ou d'un triac, dont nous n'avons rien à faire ici puisque la charge est très petite. Voyons comment il fonctionne.

Nous câblons donc notre DIAC en série avec une petite lampe et nous connectons l'ensemble à une source de tension continue que nous nous réservons le droit de faire varier. Comme prévu, tant que la tension n'est pas assez élevée, rien ne se passe et la lampe reste éteinte. Aux environs de 30 V, elle finit pourtant par s'allumer. À ce moment-là, nous faisons varier la tension dans l'autre sens : le courant continue de circuler, de moins en moins, au fur et à mesure que la tension baisse mais la lampe continue de briller comme si le DIAC était absent. La lampe ne s'éteint que lorsque le courant passe en dessous du courant de maintien puis-qu'alors le DIAC se ferme et reste fermé tant que la tension reste inférieure au seuil de déclenchement.

schéma

La petite expérience qui précède décrivait précisément la façon dont fonctionne notre dispositif. Il est aussi question ici d'une source de tension "continue" variable (une tension redressée donc) que nous fabriquerons à partir du secteur. Notre source de tension, c'est C1 qui se charge lentement à travers R1 et D3 jusqu'au niveau de la tension de déclenchement. Lorsque ce niveau est atteint, le champ électrique aux bornes du DIAC est suffisant pour que les porteurs de charge traversent les jonctions, et grossièrement dit, se multiplient : le courant passe et la LED s'allume. Le courant qui traverse D2 est en grande partie fourni par le condensateur. Celui-ci se décharge progressivement et la tension à ses bornes diminue. Il arrive un moment où l'intensité du courant tombe en dessous du courant de maintien, provoquant la fermeture du DIAC : la LED s'éteint. Le condensateur se recharge à travers R1 et tout recommence.

construction

Comme il s'agit d'un circuit qui fonctionne directement sur le secteur, il est bon de respecter certaines règles pour sa fabrication. La première, puisqu'il est quand même possible de le câbler sur une platine d'expérimentation, est d'arracher quelques pistes comme sur la figure 2. La distance entre deux conducteurs sera de cette façon ce que la norme exige. C'est assez facile avec un fer bien chaud : vous avez sans doute déjà dû constater que les pistes trop chauffées tenaient moins bien sur leur support. Une fois la platine ainsi préparée les composants peuvent y être implantés. Veillez à ne pas vous tromper dans le sens de branchement de C1, D2 et D3. Pour D1, il n'y a pas de danger puisque un DIAC n'est pas polarisé. On place ensuite la platine dans un boîtier en matière plastique pourvu, autant que possible, d'une prise moulée. Cette limitation des usages du circuit aux interrupteurs voisins d'une prise de courant améliore la sécurité : de deux maux, il faut choisir le moindre ! Pour les mêmes raisons, utilisez de la visserie en nylon pour fixer le circuit dans la boîte

dépannage

puisque la distance des trous aux pistes est inférieure à 6 mm. Si vous avez la chance de disposer du boîtier dont nous donnons les références dans la liste des composants, cette précaution n'a plus lieu d'être. Il est en effet possible d'y fixer la platine de telle façon que les vis restent inaccessibles de l'extérieur.

Il existe une autre solution qui consiste à récupérer une de ces veilleuses au néon que l'on place dans les chambres d'enfants. Comme le boîtier est trop petit pour contenir la platine, il faut ici câbler "en l'air", en prenant évidemment soin de bien isoler les fils, pour éviter les courts-circuits. Sur une des photos, il est visible que nous ne l'avons pas fait, pour que les liaisons apparaissent mieux : vous n'avez aucune raison de procéder de cette façon. Pour finir, vous collerez le capot coloré afin d'éviter aux enfants de l'enlever.

liste des composants

R1 = 470 k Ω

R2 = 680 Ω

C1 = 22 μ F/40 V

D1 = diac (ER 900, BR100-3 par exemple)

D2 = LED

D3 = 1N4007

K1 = bornier à 2 contacts
pour circuit imprimé

Boîtier en matière plastique avec prise moulée
(OKW 9010465 par exemple)

Platine d'expérimentation de format 1

À la mise sous tension, il faut quelques 10 s avant que commence le clignotement : c'est le temps que met le condensateur C1 pour se charger aux environs de 30 V. La LED s'allume et s'éteint avec une fréquence comprise entre 2 Hz et 5 Hz. Si cette fréquence est beaucoup plus élevée ou si la luminosité est continue et de très faible intensité, c'est vraisemblablement parce que le courant de maintien du DIAC est trop élevé. Ceci veut dire que le courant est coupé trop tôt pour que la LED en reçoive suffisamment. Il suffit alors de changer le DIAC ou de prendre un condensateur de capacité plus élevée.

Le problème est plus grave lorsque le circuit ne marche pas du tout. Dans ce cas, retirez la prise du secteur, vérifiez votre câblage et testez les composants : sens de branchement de D2, D3 et C1 en particulier ; vérification du bon fonctionnement des diodes à l'ohm-

mètre. Si tous ces points sont en ordre, branchez votre voltmètre aux bornes de C1 avant de remettre sous tension. Si la mesure vous donne une valeur beaucoup plus élevée que les 30 V prévus, la LED étant branchée dans le bon sens, il y a des chances pour que le DIAC soit fichu. Il est aussi possible que la résistance de R2 soit trop élevée : c'est le cas lorsqu'elle fait 680 k Ω au lieu des 680 Ω prévus !

Pardonnez-nous d'insister encore une fois pour conclure sur les dangers que présente la tension du secteur : soyez prudent, nous ne souhaitons pas ouvrir de rubrique nécrologique (et le plus vieux de nos lecteurs n'a encore que quatre-vingt-dix ans).

886031

Figure 2 – Contrairement à notre habitude et malgré l'utilisation de la tension du secteur, le circuit est câblé sur une platine d'expérimentation : les courants qui circulent sont petits, d'une part, et quelques pistes ont été arrachées pour que les distances de sécurité entre deux conducteurs véhiculant le 220 V soient respectées. Si vous câblez "en l'air" pour gagner de la place, ne laissez pas de fils nus.

