

The background image is a collage of electronic components. At the top, a microchip is visible with the text 'HCF4016BP', 'S991A9514', and 'MALLAYSTA' printed on it. Below this, there's a circuit board with various components. On the right side, there's a vertical strip showing a server rack with multiple units and blue indicator lights. The bottom right corner shows a close-up of a server unit with yellow cables.

Electricité et électronique de base

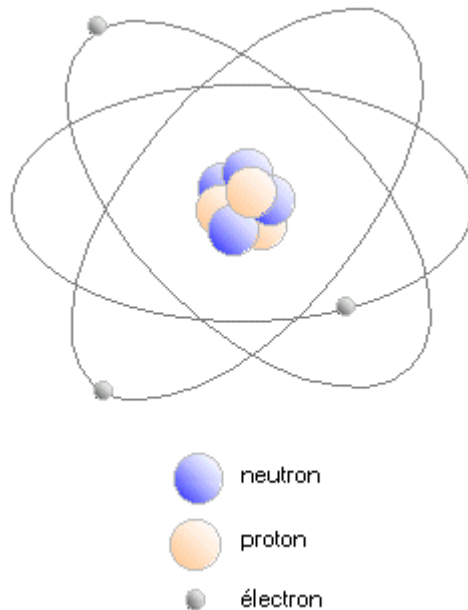
De Benedetti

- Structure de la matière:

La matière qui nous entoure est constituée de molécules, elles même constituées d'atomes.

Un atome comporte deux parties distinctes : le noyau et le nuage d'électrons.

Le noyau est la partie centrale de l'atome. Il est composé de protons qui ont une charge positive, et de neutrons qui ont une charge nulle.



Globalement la charge du noyau est positive.

Autour du noyau gravite un nuage d'électrons chargés négativement.

Les charges négatives compensent exactement les charges positives du noyau :

La matière est électriquement neutre .

Electrisation d'un solide

Dans les solides l'organisation des noyaux est très rigide. Les noyaux ne peuvent pas être déplacés (ce n'est pas le cas à l'état liquide ou gazeux).

Par contre, il est possible de déplacer des électrons.

La matière étant électriquement neutre :

- Si un atome reçoit un électron il se charge négativement
- Si un atome perd un électron, il se charge positivement.

Cette opération s'appelle l'électrisation.

Un atome auquel on aura enlevé 1 électron sera dit présenter une charge de 1 e.

Un atome auquel on aura enlevé 2 électrons sera dit présenter une charge de 2 e.

Un atome qui aura reçu 1 électron sera dit présenter une charge de - 1 e.

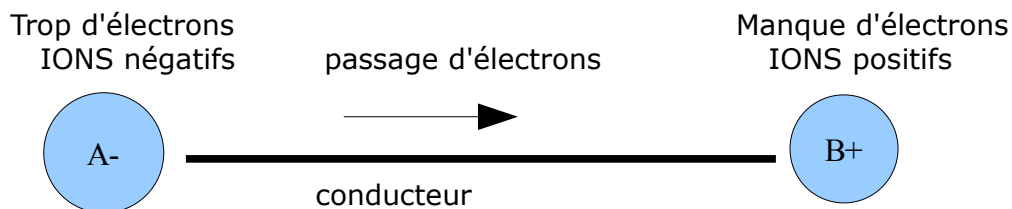
Un atome qui aura reçu 2 électrons sera dit présenter une charge de - 2 e.

Un solide chargé positivement est attiré par un solide chargé négativement. Deux solides de même charge (positives ou négatives) se repoussent.

– **Le courant électrique :**

L'électron, à cause de sa mobilité, est donc l'agent du courant électrique, il suffit de :

- l'arracher à son ensemble, l'atome
- lui offrir un passage, le conducteur



Si on réunit A et B par un canal constitué par un corps ayant des électrons périphériques facilement détachables, les électrons circuleront de A vers B.

Le canal est dit **conducteur du courant électrique**.

Au contraire, si on réunit A et B par un canal constitué par un corps ayant des électrons périphériques stables, il n'y aura pas de courant électrique.

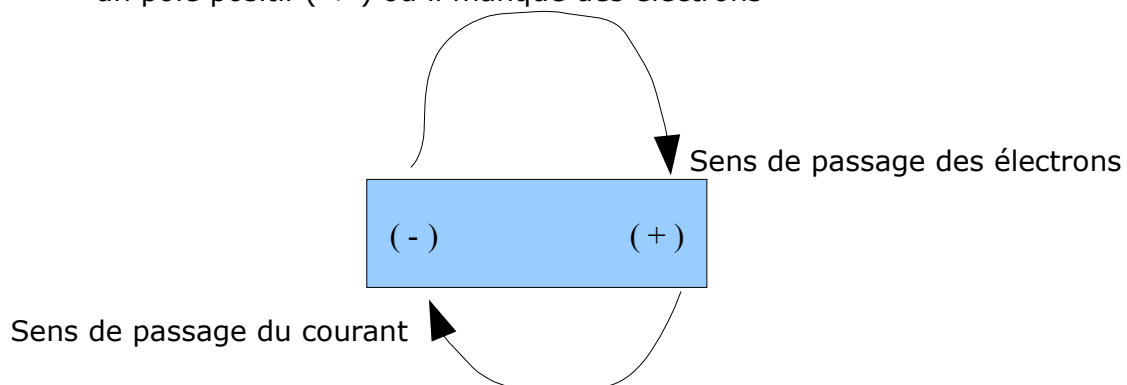
Le canal est dit **très résistant ou isolant**.

– **Générateur ou source de courant :**

On appelle générateur tout système capable de mettre des électrons en mouvement.

Un générateur comporte donc deux BORNES ou POLES:

- un pôle négatif (-) ou les électrons sont en surnombre
- un pôle positif (+) ou il manque des électrons



On adopte comme sens conventionnel du courant électrique, le sens borne positive vers borne négative. (inverse du passage d'électrons).

– **Circuit électrique :**

En réunissant les deux bornes d'un générateur avec un conducteur, on réalise un **circuit fermé**, c'est à dire dans lequel un courant peut circuler (en général des métaux).

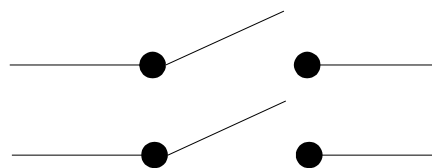
Si on interpose dans ce circuit un corps isolant (impropre au passage du courant), on réalise un **circuit ouvert**.

L'air est un isolant, pour ouvrir ou fermer un circuit, on utilise des interrupteurs.



Interrupteur unipolaire

Interrupteur bipolaire (agissant sur les deux fils)



Dans un circuit électrique nous trouverons aussi le récepteur, celui qui reçoit et transforme l'électricité (tout les appareils électriques).

– **Les effets du courant électrique :**

Dégagement de chaleur.

Un conducteur dans un circuit fermé s'échauffe, il peut aller jusqu'à rougir ou fondre suivant sa section.

Les électrons qui circulent à travers les atomes du conducteur subissent des chocs et des frottements. L'échauffement constaté correspond à une perte d'énergie par frottement.

Ex : ampoule à incandescence, radiateur...

Action magnétique.

Si nous plaçons sous le conducteur une aiguille aimantée parallèle au fil (une boussole) : Quand le courant passe dans le circuit, l'aiguille dévie, si on inverse le courant, l'aiguille dévie dans l'autre sens. Une aiguille aimantée s'oriente, (placée dans une zone ne contenant ni aimant, ni circuit électrique) dans le champ magnétique terrestre : l'extrémité qui pointe vers le nord est appelé pôle nord.

Il y a action magnétique.

– Les grandeurs électriques

Unité de quantité d'électricité : Le Coulomb (C).

L'électron porte une charge électrique, que l'on appelle charge élémentaire , e.
Le passage d'un nombre n d'électrons, correspond au passage d'une quantité d'électricité, ne.

Comme unité de **quantité d'électricité**, on ne prendra pas la charge élémentaire (trop faible), mais une unité beaucoup plus grande appelée le **Coulomb** (symbole **C**).
Dans ces conditions, la charge d'un électron est de $1,6 \times 10^{-19}$ coulombs.

Intensité d'un courant : (I)

L'intensité d'un courant est mesurée par le nombre de Coulombs débités par seconde.

Symbole : I

Unité d'intensité de courant : l'Ampère (A)

L'ampère est l'intensité d'un courant qui débite un coulomb en une seconde.

$1A \rightarrow 1C/s$

ex: Si un courant débite 100C pendant 20 sec, son intensité est de $I=100/20 =5A$

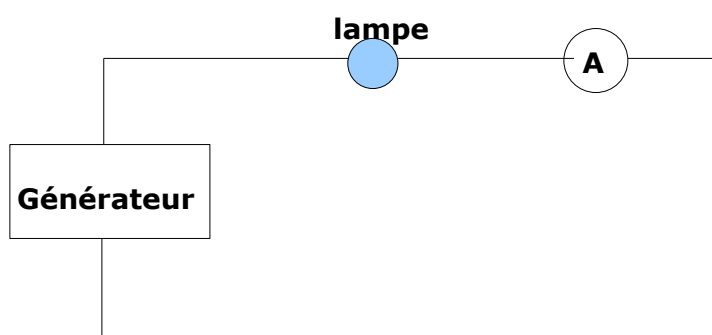
ex : Si un courant a une intensité de 2A pendant 60 sec, il débite $2 \times 60 = 120C$

Si l'on représente la quantité d'électricité en Coulomb par Q, l'intensité en Ampère par I, le temps en s par t, on a :

$$Q = I \times t$$

Mesure d'une intensité :

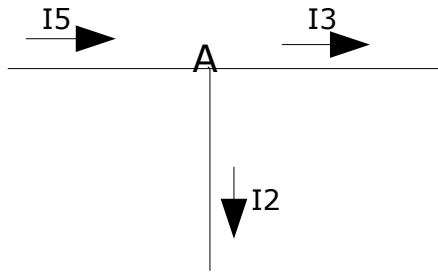
On utilise un Ampèremètre, qui se branche en série dans le circuit:



- **Plusieurs récepteurs montés en série, sont parcourus par le même courant.**
- **Le courant d'un circuit principal est la somme des courants des circuits dérivés qu'il alimente.**
Conséquence : la loi des nœuds.

On appelle nœud un point commun à plusieurs portions de circuit.

Ex : A



Comme il ne peut y avoir accumulation d'électrons en un point du circuit, **La somme des courants arrivant a un nœud est égale à la somme des courants qui en repartent.**

Donc, dans l'exemple : $I_5 = I_3 + I_2$

En général on résume cette loi par la formule : $\sum (I) = 0$

Ce qui veut dire : **la somme algébrique des intensités relatives au nœud est nulle .**

- **Unité industrielle de quantité d'électricité : L'ampère-heure (Ah)**

Le temps d'utilisation du courant électrique se chiffre plutôt en heures qu'en secondes, on utilise en industrie une unité plus grande que le coulomb, l'Ampère-heure (Ah).

C'est la quantité d'électricité transportée par un courant durant 1 heure.

Donc : **1Ah = 3600 Coulombs**

C'est le genre d'unité que l'on retrouve sur les batteries d'accumulateurs.

Exemple : si une batterie de voiture a une capacité de 75Ah, cela veut dire qu'elle peut fournir 75 Ampères par heure, ou 7,5 Ampères durant 10h, ou 5A durant 15h....

– **La différence de potentiel (d.d.p.) : Symbole U**

Analogie hydraulique :

Pour comprendre certaines propriétés du courant électrique, il est intéressant de le comparer à de l'eau s'écoulant dans un circuit de tuyaux. Le générateur peut alors être vu comme une pompe chargée de mettre en pression ce liquide dans les tuyaux.

La différence de potentiel, ou tension, ressemble alors à la différence de pression entre deux points du circuit d'eau. Elle est notée U, et exprimée en volts (V).

L'intensité du courant électrique peut être rapprochée du débit d'eau dans le tuyau. Elle mesure le nombre de charges qui passent chaque seconde à un point du circuit ; elle est souvent notée I, et mesurée en ampères (A).

La résistance d'un circuit électrique serait alors analogue au diamètre des tuyaux. Plus les tuyaux sont petits, plus il faut de pression pour avoir le même débit ; de façon analogue, plus la résistance d'un circuit est élevée, plus il faut une différence de potentiel élevée pour avoir la même intensité. La résistance électrique mesure donc la faculté de freiner plus ou moins le passage du courant. Elle est notée R et, elle est exprimée en ohms (Ω).

Il est possible de pousser cette analogie beaucoup plus loin mais il est important de garder à l'esprit qu'elle a ses limites et que certaines propriétés du courant électrique s'écartent sensiblement de ce modèle à base de fluide, de tuyaux, et de pompes.

Définition :

La d.d.p. Entre deux points d'un circuit, a pour mesure le rapport de l'énergie perdue entre ces deux points, a la quantité d'électricité qui y passe.

$$U = W / Q$$

U : d.d.p.

W : énergie perdue en joules

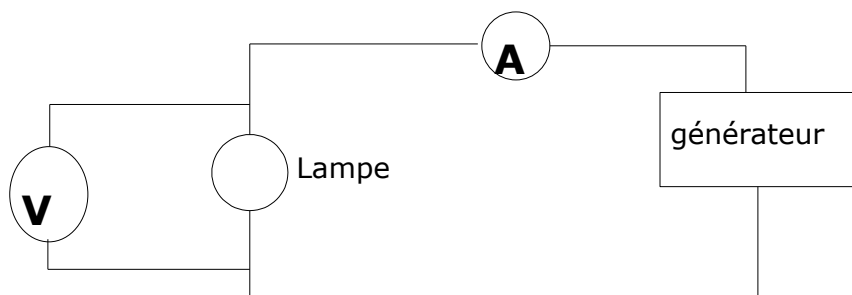
Q : quantité d'électricité en Coulombs

Unité : le Volt (V).

Il existe une différence de potentiel ou tension de 1 Volt entre deux points d'un circuit, quand chaque coulomb qui y passe abandonne une énergie de 1 joule.

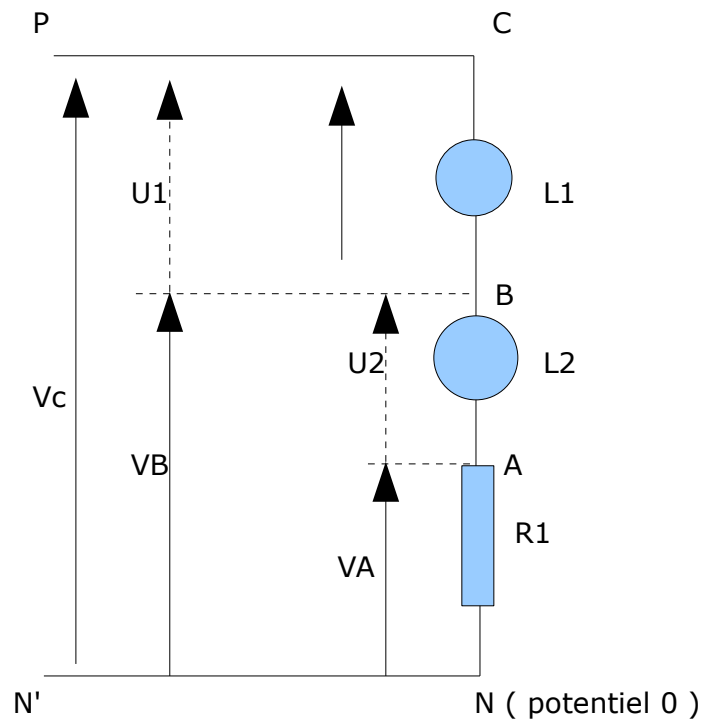
Mesure d'une différence de potentiel : appareil, le Voltmètre

Un voltmètre se branche en parallèle (ou dérivation) sur le circuit.



Règles :

- Plusieurs récepteurs groupés en parallèle sont soumis à la même différence de potentiel.
- Si des récepteurs sont groupés en série, la différence de potentiel aux bornes de l'ensemble est la somme des différences de potentiel successives .



On considère comme potentiel de référence le point N (potentiel 0, masse)

donc : $V_N = 0$

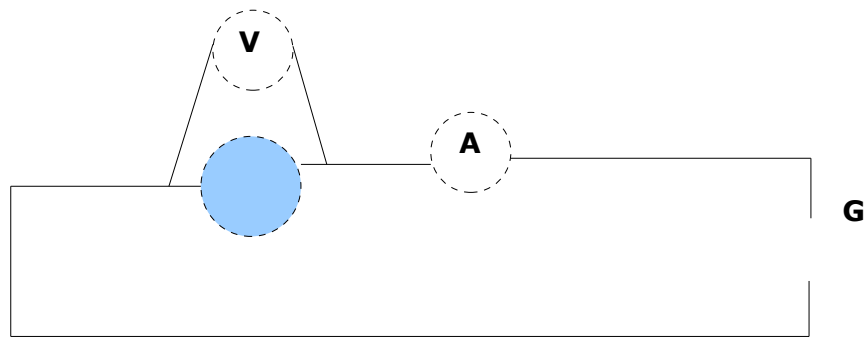
$$U_1 = V_C - V_B$$

$$U_2 = V_B - V_A$$

U_1 et U_2 ne dépendent pas du potentiel de référence.

$$U_1 + U_2 = (V_C - V_B) + (V_B - V_A) = V_C - V_A$$

La puissance électrique : symbole P



Dans ce montage, un générateur de courant alimente une ampoule électrique. Un ampèremètre est branché en série dans le circuit, un voltmètre est branché en parallèle sur l'ampoule.

Le voltmètre indique 220V et l'ampèremètre 0,4 A

On dit que l'énergie dépensée dans l'ampoule est de $220 \times 0,4 = 88$ joules
Ces 88 joules représentent la puissance de l'ampoule, soit 88 Watts

La puissance électrique dépensée par un récepteur est proportionnelle à la d.d.p. À ses bornes et à l'intensité du courant qui le traverse.

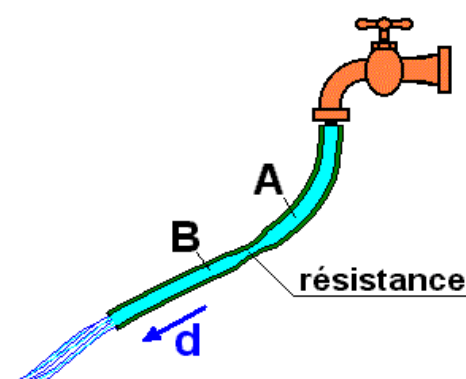
La puissance se mesure en WATT (W)

$$P = U \times I$$

La loi d'Ohm

La notion de résistance

Lorsque qu'un tuyau transporte de l'eau il suffit d'écraser un peu le tuyau pour que le débit d diminue et que la pression augmente en amont **A** du point d'étranglement et diminue en aval **B** de celui-ci. Plus la résistance au point d'étranglement sera importante et plus sera grande la différence de pression entre les points A et B. En même temps que la résistance augmente, le débit diminue. Si le tuyau est complètement écrasé, l'eau ne passe plus, le débit est nul et la résistance est infinie. On imagine également que s'il y a deux points d'étranglement l'un à la suite de l'autre (en série) sur le tuyau, la résistance globale sera plus grande et le débit encore plus faible.



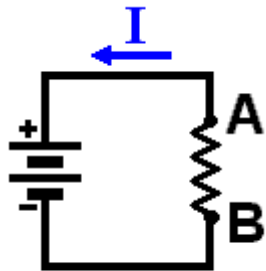
La résistance électrique

En pratique, dans un circuit électrique, la résistance au courant électrique peut être un composant appelé résistance (ou resistor dans certains manuels) ou plus simplement la résistance qu'oppose les conducteurs eux-mêmes au courant électrique. Un petit câble laissera moins passer le courant qu'un gros câble. Certains métaux sont moins bons conducteurs pour l'électricité que d'autres. Par exemple la résistivité du fer est plus grande que celle du cuivre ou de l'argent.

L'unité de résistance électrique est l'ohm (symbole Ω , oméga). Les valeurs communes de résistance vont de quelques milli-ohms à des dizaines de mégohms. Les multiples communs de l'ohm sont :

- kilohm : $1k = 1000$ ohms
- mégohm : $1 M = 1000 k = 1000\ 000$ ohms

Le symbole de la résistance en tant que grandeur électrique est **R**.



La loi d'Ohm

C'est une loi fondamentale de l'électricité. Elle exprime la relation qui existe entre l'intensité **I** dans une portion de circuit de résistance **R** et la différence de potentiel **U** aux bornes de cette portion de circuit et s'énonce :

"La différence de potentiel en volts aux bornes d'une résistance est égale au produit de la valeur en ohm de cette résistance par l'intensité en ampères qui la traverse". Ce qui se traduit par la formule :

$$U = R \cdot I$$

Note : R doit être une résistance pure, c'est à dire ne transformant l'énergie électrique qu'en énergie calorifique. Cette remarque est particulièrement importante en courant variable.

Deux autres formules très utiles en découlent :

$$R = \frac{U}{I} \quad \text{et} \quad I = \frac{U}{R}$$

Pour conserver l'analogie avec notre tuyau d'arrosage du début nous dirons que la différence de pression entre l'aval et l'amont d'un étranglement dans une conduite d'eau est proportionnel à la résistance de l'étranglement et au débit d'eau dans le tuyau. La comparaison s'arrête là, elle n'a pas d'autre but que d'aider à la compréhension du phénomène.

Association de résistances en parallèle

Résistances en parallèle

Lorsqu'un générateur débite un courant d'intensité **I** dans un groupement de résistances en parallèle, il se décompose en i_1 , i_2 et i_3 tels que :

$$I = i_1 + i_2 + i_3$$

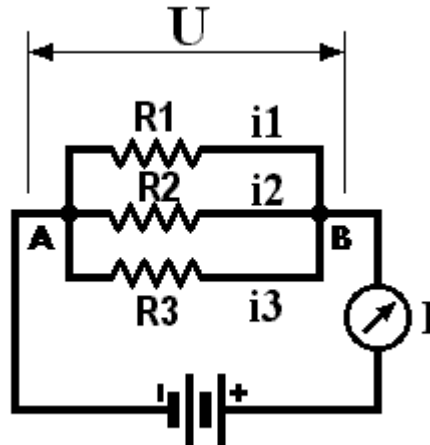
Selon la première loi de Kirchhoff, la somme des courants qui entrent dans le noeud A est égale à la somme des courants qui en sortent.

Autre remarque : la tension aux bornes des trois résistances est la même :

$$U = R_1.i_1 = R_2.i_2 = R_3.i_3$$

La valeur de la résistance équivalente à R_1 , R_2 et R_3 en parallèle peut être calculée par la formule :

$$\frac{1}{R} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3}$$



Cas de n résistances identiques en parallèle

Pour réaliser une charge de 51 ohms et de puissance 10 watts il suffit de mettre 10 résistances de 510 ohms en parallèle.

La résistance R_e d'un groupement parallèle de n résistances identiques R est donnée par la formule :

$$R_e = \frac{R}{n} \quad \text{d'où} \quad n = \frac{R}{R_e}$$

Il est facile de déterminer le nombre de résistances de 2200 ohms nécessaires pour obtenir une résistance équivalente de 30 ohms :

$$n = R/R_e = 2200/30 = 73 \text{ résistances}$$

Cas de 2 résistances différentes en parallèle

C'est le cas le plus fréquent : une résistance R_1 est en parallèle avec une résistance R_2 , quelle est la résistance équivalente R ?

Une formule facile à démontrer permet ce calcul simple :

$$R = \frac{R_1 \cdot R_2}{R_1 + R_2}$$

exemple : $R_1=10$ ohms et $R_2=30$ ohms, $R=(10 \cdot 30)/(10+30) = 300/40 = 7,5$ ohms.

Dans tous les cas la résistance équivalente à un groupement de résistances en parallèle est inférieure à la plus faible de ces résistances.

Pour préciser que R_1 et en parallèle avec R_2 on écrit souvent $R_1//R_2$.

Association de résistances en série

Résistances en série

Doubler la longueur d'un conducteur équivaut à doubler sa résistance électrique (voir résistivité et résistance d'un conducteur).

Exemple : un câble de longueur 1m a une résistance de 1 ohms. Mettre en série deux tels câbles donne un câble de longueur 2m et de résistance 2 ohms. Autrement dit, lorsqu'on met en série deux résistances de valeurs respectives R_1 et R_2 on obtient une résistance équivalente de valeur R telle que :

$$R = R_1 + R_2$$

Dans le cas général où plusieurs résistances sont en série la résistance équivalente R est égale à la somme des valeurs de chacune des résistances du circuit :

$$R = R_1 + R_2 + \dots + R_N$$

Exemple pratique

Question : Un circuit comporte en série trois résistances de valeur 1,2k , 2,2k et 3,3k. Quelle est la résistance équivalente du circuit ?

Solution : $R = 6700$ ohms.

Le diviseur potentiométrique

Ce type de circuit est très fréquent. Il permet d'obtenir une tension quelconque U_2 à partir de la tension U d'un générateur. La tension U_2 étant comprise entre 0 et U volts. C'est ce principe qui est utilisé dans les potentiomètres.

Dans le circuit ci-contre la tension U_2 est égale à U multipliée par le rapport $R_2/(R_1+R_2)$.

Exemple :

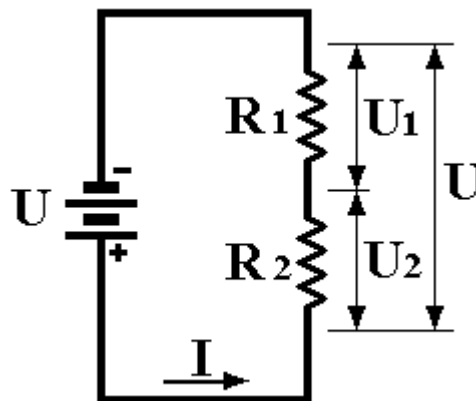
$U = 12$ volts

$R_1 = 100$ ohms

$R_2 = 20$ ohms

donc $U_2 = 12 \times 20/(100+20)$

$U_2 = 2$ volts.



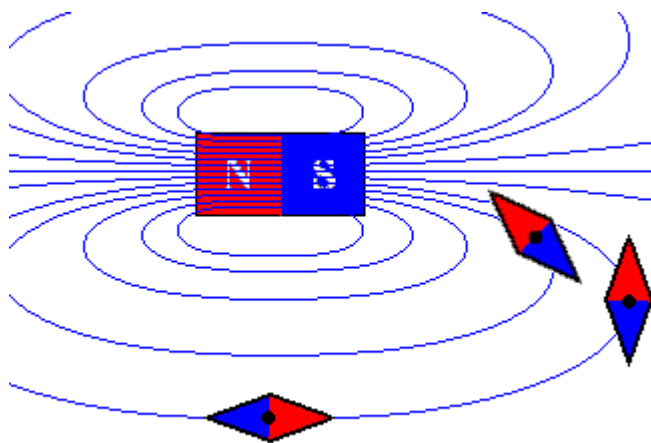
L'électromagnétisme et ses applications

Les aimants

Le terme de "magnétisme" vient du nom de l'île grecque de Magnésie où l'on trouve la magnétite, oxyde de fer (Fe_3O_4) qui est une forme d'aimant naturel. Ce minéral de fer a la propriété d'attirer les objets en fer ou en acier.

On peut réaliser facilement un aimant artificiel en frottant un barreau en acier contre un autre aimant.

La forme habituelle des aimants est celle d'un barreau droit, d'une aiguille de boussole ou d'un fer à cheval mais beaucoup d'autres formes se rencontrent comme les aimants cylindriques de haut-parleurs, par exemple.



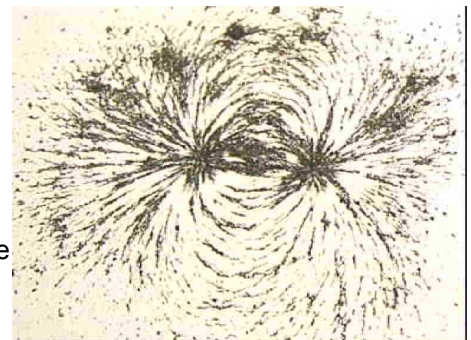
Les deux pôles de l'aimant

Un aimant possède un pôle nord et un pôle sud. La Terre peut être considérée comme un gigantesque aimant dont les pôles magnétiques sont relativement proches des pôles géographiques.

Quand on ramasse de la limaille de fer avec un aimant, celle-ci s'agglutine principalement au niveau des pôles. On ne peut séparer le pôle nord du pôle sud d'un aimant. Couper un aimant en deux revient à fabriquer deux aimants plus petits. Placée dans un champ magnétique une aiguille de boussole s'oriente selon les lignes de force. (lignes d'induction)

Le champ magnétique

Les lignes de force magnétiques traversent l'aimant de part en part suivant l'axe des pôles et se referment à l'extérieur de celui-ci. L'ensemble des lignes de force constitue le champ magnétique que l'on peut visualiser facilement à l'aide de grains de limaille de fer. La photo ci-contre représente le *spectre magnétique* d'un petit aimant placé sous une feuille de papier. Chaque grain de limaille se comporte comme un petit aimant avec ses deux pôles nord et sud qui attirent les pôles sud et nord des grains voisins.



De ces deux expériences, nous pouvons tirer les conclusions suivantes :

l'induction en un point est caractérisée par :

une direction (celle de l'aiguille aimantée)

un sens : celui qui va du nord au sud

une intensité : avec la même aiguille, pour deux aimants différents, il faudra appliquer une force différente pour lui faire quitter sa position d'équilibre. Nous dirons que le champ le plus intense est celui qui nécessite la plus grande force.

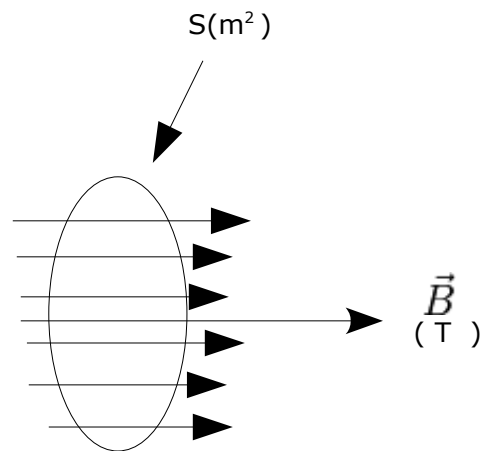
L'induction magnétique est donc une grandeur vectorielle :

$\vec{B}$

L'unité d'induction est le Tesla (T)

Flux d'induction magnétique : symbole Φ (phy)

Considérons une surface S placée perpendiculairement aux lignes d'induction d'un champ uniforme :



Les lignes d'induction sont parallèle (champ uniforme)
La surface est perpendiculaire au vecteur d'induction B .

On appelle flux d'induction à travers une surface le produit de l'intensité du champ magnétique par la surface considérée.

$$\Phi = B \times S$$

Φ en Wb
 B en T
 S en m^2

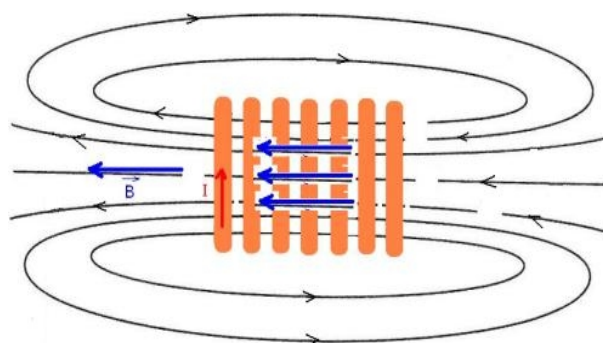
L'unité de flux est le Weber (Wb)

Action magnétique du courant :

Une bobine de fil de cuivre (de longueur supérieure au diamètre = bobine longue), est alimentée en courant continu, un rhéostat permet de faire varier l'intensité dans la bobine.

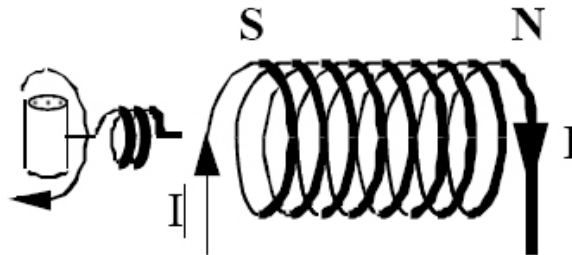
Cette bobine s'appelle un **solénoïde**

Quand un solénoïde est parcouru par un courant, il produit un champ magnétique dans son voisinage, et plus particulièrement à l'intérieur de l'hélice où ce champ est quasiment uniforme.



Caractéristiques du vecteur d'induction dans le solénoïde :

- il a pour direction l'axe de la bobine
- pour sens , celui du tire bouchon de maxwell : un tire bouchon placé dans l'axe de la bobine et tournant dans le sens du courant, se déplace dans le sens du vecteur d'induction . (B)



- pour valeur : $B = (4\pi / 10^7) \times (N \times I)/l$
avec B en T, N = nombre de spires de la bobine, I intensité en ampère, l = longueur de la bobine en mètre.

Aimantation du fer et de l'acier :

Si nous alimentons le solénoïde du montage précédent, il présente une aimantation insuffisante pour soulever des clous.

Plaçons un un noyau en fer dans cette bobine : les clous sont alors attirés

Coupons le courant, les clous retombent, c'est une aimantation temporaire.

Explication : Le fer canalise les lignes d'induction, on dit qu'il est plus perméable que l'air

Le champ magnétique de la bobine crée, dans tout milieu, un champ d'induction magnétique dont l'intensité est représenté dans l'air par B_0 , et dans le fer par B.

Nous constatons que B est supérieur à B_0 .

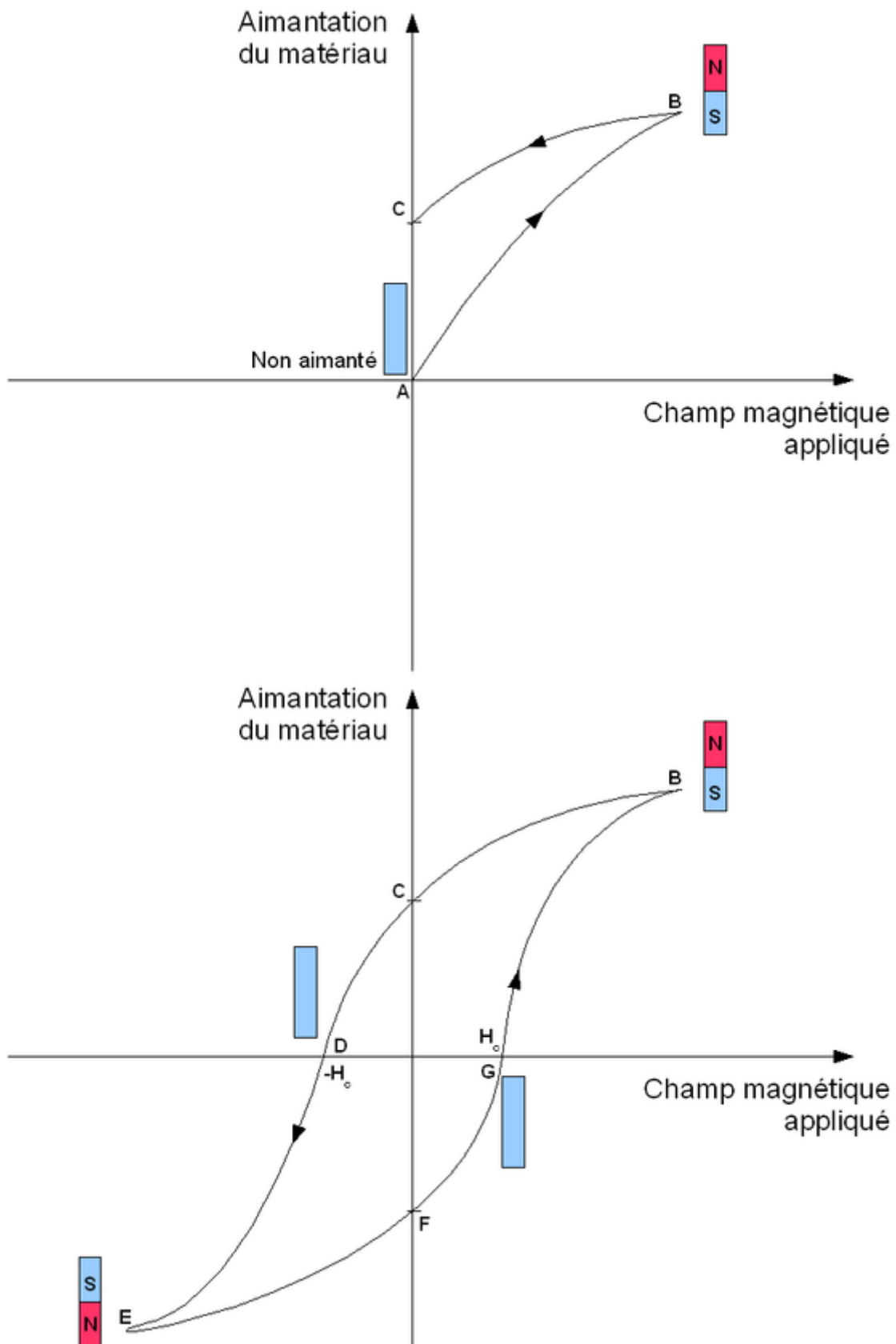
$$B/B_0 = \mu_r \text{ (lettre grecque mu)}$$

μ_r mesure la perméabilité relative (c a d par rapport à l'air) du noyau de fer.

Le cycle d'hystérésis

Lorsqu'un matériau ferromagnétique est soumis à un champ magnétique, il conserve une certaine aimantation, il ne pourra pas perdre cette aimantation, sauf en lui appliquant un champ magnétique dans l'autre sens, d'une intensité appropriée, appelé champ coercitif.

La courbe d'aimantation d'un ferromagnétique est appelée courbe d'hystérésis . C'est grâce au comportement décrit par cette courbe, et en particulier du fait qu'ils restent aimantés même en l'absence de champ extérieur, que les matériaux ferromagnétiques sont intéressants pour l'enregistrement.



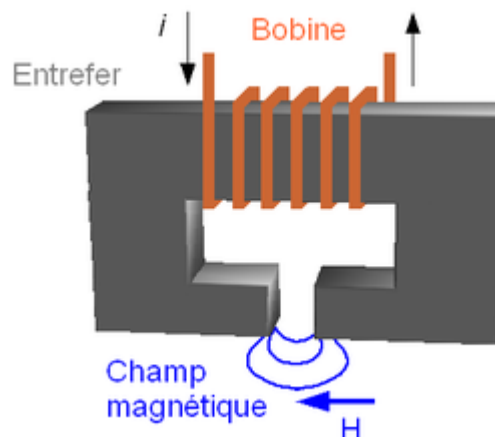
En haut : au départ (point A) le matériau n'est pas aimanté. L'application d'un champ magnétique H va lui faire prendre une aimantation (point B). Lorsque le champ magnétique appliqué est retiré, le matériau conserve son aimantation (C) :

c'est devenu un aimant.

En bas : le cycle d'hystérésis complet. Il faut appliquer un champ magnétique $-H_c$, appelé champ coercitif, pour faire perdre son aimantation au matériau (D). L'application d'un champ $-H$ plus intense va aimanter le matériau dans l'autre sens (E). Là encore, lorsqu'on retire le champ magnétique appliqué le matériau conserve son aimantation (F), et un champ coercitif $+H_c$ est alors nécessaire pour lui faire perdre son aimantation (G).

L'application d'un champ $+H$ plus intense inversera de nouveau son aimantation (B). Le cycle peut être parcouru à volonté, aimantant ainsi le matériau dans un sens ou dans l'autre.

L'écriture (sur bande magnétique ou disque dur) est généralement assurée par un électroaimant . Le principe est basé sur l'électromagnétisme le plus simple : tout courant (soit un déplacement de charge électrique) provoque l'apparition d'un champ magnétique. La bobine est donc soumise à un courant électrique, qui induit un champ magnétique dans l'entrefer. Dans les cassettes audio, c'est le son enregistré qui module le courant et donc le champ magnétique ; la bande sera ainsi plus ou moins aimantée au cours de l'enregistrement. On parle d'enregistrement analogique car le matériau peut être plus ou moins aimanté de manière continue. Par opposition, dans les disques durs c'est un enregistrement numérique, c'est-à-dire utilisant le langage binaire -des 0 et des 1. Selon que le courant passe dans un sens ou l'autre dans la bobine, le champ magnétique induit sera orienté dans un sens ou dans l'autre, permettant ainsi de choisir le sens d'aimantation du matériau.



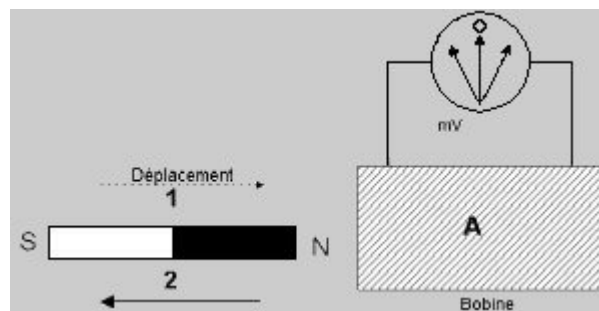
Pour les dispositifs utilisant des domaines magnétiques assez étendus (bandes magnétiques, cassettes audio et vidéo, les premiers disques durs...), la lecture est assurée par le même dispositif, mais en utilisant le principe d'électromagnétisme inverse : toute variation de champ magnétique provoque un champ électrique. Le champ magnétique du matériau va ainsi induire un courant électrique dans la bobine, qui pourra être amplifié et exploité par un circuit électronique.

Étude de l'induction électromagnétique:

Petite histoire:

Les phénomènes d'induction ont été découverts en 1831 par le physicien Anglais FARADAY. Dans la même années le Français AMPERE, l'américain HENRY, le Russe LENTZ, et l'allemand NEUMANN précisèrent la théorie et entrevirent quelques applications du phénomène.

Expérience:



D'après le schéma, les extrémités d'une bobine "A" sont reliées aux deux pôles d'un millivoltmètre, à courant continu ayant son "0" au milieu du cadrant.

lorsqu'on introduit le pôle Nord d'un aimant dans cette bobine (direction 1), l'aiguille du millivoltmètre dévie à gauche, puis revient à 0 quand le mouvement cesse.
quand on retire (direction 2) le Nord de l'aimant, l'aiguille dévie à droite et revient à 0 quand le mouvement cesse.

Avec le pôle Sud les déviations sont inversées.

plus le mouvement est rapide, plus la déviation de l'aiguille est importante.

Si on remplace l'aimant par une bobine alimentée elle aussi en courant continu, on constate les mêmes phénomènes. La déviation de l'aiguille augmente si l'intensité dans la bobine augmente.

Interprétation des phénomènes:

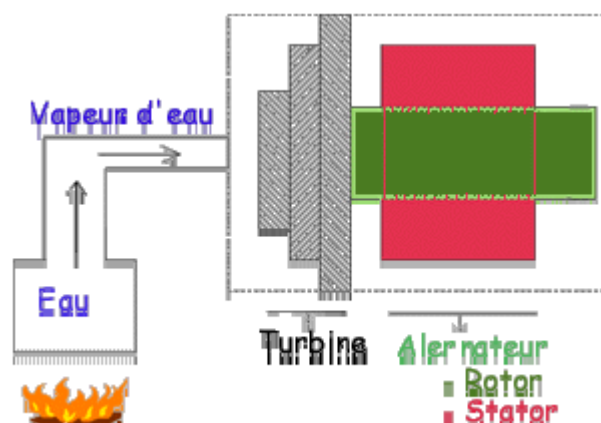
La déviation du millivoltmètre est produite par un courant qui prend naissance dans la bobine "A". Ce courant est appelé Courant induit.

La bobine "A" s'appelle bobine induite ou induit.

La bobine "B" (ou l'aimant) constitue la bobine inductrice ou inducteur.

Le phénomène est l'induction électromagnétique

Dans les centrales électriques, les aimants sont remplacés par des électroaimants (bobines parcourues par un courant). Leur rotation est obtenue grâce aux turbines mises en rotation par un fluide. Exemple: centrale thermique



Un alternateur est une machine rotative qui convertit l'énergie mécanique fournie par un moteur (turbine, diesel, éolienne) en énergie électrique à courant alternatif.

Plus de 95 % de l'énergie électrique est produite par des **alternateurs synchrones** : machines électromécaniques fournissant des tensions de fréquences proportionnelles à leur vitesse de rotation.

Cette machine est constituée d'un **rotor** (partie tournante) et d'un **stator** (partie fixe).



Le **rotor** est l'inducteur

- Il peut être un aimant permanent, la régulation de la tension de sortie se fera en régulant la vitesse de rotation de l'alternateur, (la fréquence du courant variera également). C'est le principe de la *dynamo* de vélo, qui est en fait un petit alternateur.
- Plus couramment un électroaimant assure l'induction. Ce bobinage est alimenté en courant continu, soit à l'aide d'un collecteur à bague rotatif (une double bague avec balais) amenant une source extérieure, soit par un excitateur à diodes tournantes et sans balais. Un système de régulation permet l'ajustement de la tension et de la phase du courant produit.

Le **stator** est l'induit . Il est constitué d'enroulements qui vont être le siège de courant électrique alternatif induit par le champ magnétique créé par l'inducteur en mouvement.

Ce type de machine produit du **courant alternatif**.

Le courant alternatif est un courant électrique qui change de sens.

Ce courant alternatif est dit périodique s'il change régulièrement et périodiquement de sens.

Un courant alternatif périodique est caractérisé par sa fréquence , mesurée en hertz (Hz). C'est le nombre d' « aller-retours » qu'effectue le courant électrique en une seconde

Un courant alternatif périodique de 50 Hz effectue 50 « aller-retours » par seconde, c'est-à-dire qu'il change 100 fois (50 allers et 50 retours) de sens par seconde.

La forme la plus utilisée de courant alternatif est le courant sinusoïdal, essentiellement pour la distribution commerciale de l'énergie électrique

La fréquence du courant électrique distribué par les réseaux aux particuliers est généralement de 50 Hz en Europe et 60 Hz en Amérique du Nord .

On doit distinguer :

Les courants purement alternatifs dont la valeur moyenne est nulle, qui peuvent alimenter un transformateur sans danger.

Les courants alternatifs à composante continue non nulle qui ne peuvent en aucun cas alimenter un transformateur

Avantages

Contrairement au courant continu, le **courant purement alternatif** peut voir ses caractéristiques (tension/courant) modifiées par un transformateur à enroulements. *Dès qu'il existe une composante continue non négligeable, un transformateur est inutilisable.*

Grâce au transformateur :

- Le courant transporté par des lignes à haute tension subit des pertes par effet Joule beaucoup plus faibles. En divisant simplement par 10 l'intensité du courant transporté, on divise par 100 les pertes dues à la résistance des câbles électriques, la puissance dissipée dans une résistance étant proportionnelle au carré de l'intensité du courant. ($P = RI^2$)
- À puissance constante, on peut réduire fortement l'intensité d'un courant alternatif en augmentant sa tension.
- On abaisse ensuite la tension afin de fournir une alimentation en basse tension près du lieu de distribution, afin de pouvoir l'utiliser à des fins domestiques (attention le danger est bien réel, même en basse tension).

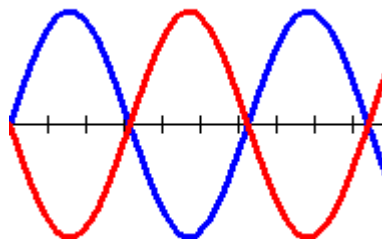


Figure 2b

Voici un exemple de signaux sinusoïdaux.

On dit de ces deux signaux qu'ils sont identiques mais déphasés de 180°. Entre leurs deux équations, il y a donc seul le déphasage (ou phase à l'origine) qui diffère.

Fonctionnement des récepteurs électriques.

Les récepteurs électriques habituellement utilisés peuvent se ranger en deux catégories:

a) Les appareils qui fonctionnent directement sous la tension du secteur: les appareils de chauffage (four, réchaud, fer à repasser), d'éclairage (lampe à incandescence, tubes luminescents) ou les récepteurs contenant des moteurs (machine à laver, réfrigérateur, tondeuse à gazon, mixeur, perceuse...)

b) Les récepteurs électroniques: ordinateurs, chaîne Hi fi, radio-cassettes, qui ne fonctionnent pas en alternatif mais en courant continu.

D'ailleurs, certains de ces appareils peuvent fonctionner sur piles et possèdent un adaptateur externe permettant de remplacer l'énergie très coûteuse des piles par celle du secteur. Dans un ordinateur de bureau ou une chaîne Hi fi, l'adaptateur est interne, il constitue la partie alimentation de l'appareil.

L'adaptateur permet d'obtenir une basse tension continue à partir du 230V alternatif de la prise du secteur.

La transformation d'une tension alternative.

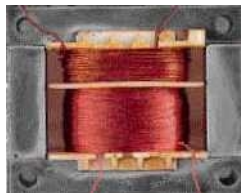
Rôle d'un transformateur

Un transformateur sert à modifier la valeur efficace d'une tension alternative. Il peut l'abaisser ou l'élever.

Description d'un transformateur

Un transformateur est constitué de 2 bobines de fil de cuivre isolé montées sur une armature en fer doux.

Remarque: Le fer doux est du fer pur, alors que l'acier est un alliage de fer et de carbone. Le fer doux et l'acier s'aimantent lorsqu'ils sont placés dans le champ magnétique d'une bobine, mais lorsqu'on interrompt le courant dans la bobine, le fer doux cesse d'être aimanté alors que l'acier conserve son aimantation.



La bobine d'entrée est appelée primaire, celle de sortie, secondaire. Les 2 bobines sont indépendantes. Il n'existe aucune liaison électrique entre elles.

L'armature en fer doux passe à l'intérieur des bobines et se referme à l'extérieur. Elle est constituée de plaques superposées pour diminuer les pertes.

Le fil de cuivre est isolé par un vernis transparent qui pourrait laisser croire que le fil est nu.

Fonctionnement d'un transformateur

En déplaçant un aimant près d'une bobine, on crée une tension variable dans la bobine (voir alternateurs). La tension induite dans la bobine est due à la variation du champ magnétique de l'aimant que l'on déplace.

Ici, c'est la variation du champ magnétique créé par le courant variable circulant dans la bobine primaire qui induit une tension variable dans la bobine secondaire.

Remarque: Un transformateur ne fonctionne pas en courant continu (pas de variation du champ magnétique), de même qu'un alternateur ne fournit aucune tension si on ne le fait pas tourner.

Si le primaire est soumis à une tension alternative, le secondaire sera soumis à une tension alternative de même fréquence.

La tension efficace obtenue au secondaire dépend du nombre de spires des bobines.

Rapport de transformation

Le rapport de transformation k est le quotient de la tension au secondaire U_s et de la tension au primaire U_p :

$$k = U_s / U_p$$

Exemple: Un transformateur qui fournit une tension de 24 V à la sortie lorsque l'entrée est soumise à 240V a un rapport de transformation:

$$k = 24V / 240V = 1/10 = 0,1$$

Si on compte le nombre de spires des bobines de ce transformateur on observera que le rapport est environ 1/10

Le rapport d'un transformateur (supposé sans pertes) est égal au quotient du nombre U_s de spires au secondaire et du nombre U_p de spires au primaire.

$$k = N_s / N_p$$

Il suffit que le nombre de spires au secondaire soit dix fois plus petit qu'au primaire pour obtenir un rapport $k = 1/10$

Cette règle suppose le transformateur idéal, sans perte. Elle est assez bien vérifiée à vide, c'est-à-dire lorsque rien n'est branché au secondaire. Mais lorsque le courant passe dans cette bobine, on constate une chute de tension.

Remarque: Pour minimiser les pertes, le nombre de spires doit être suffisant. Il n'est pas possible de réaliser un transformateur de rapport 1/10 avec 20 spires au primaire et 2 spires au secondaire. On compte souvent une dizaine de spires par volt soit environ 2000 spires pour un primaire relié au secteur 230V.

Redressement d'une tension alternative.

Redressement d'une alternance

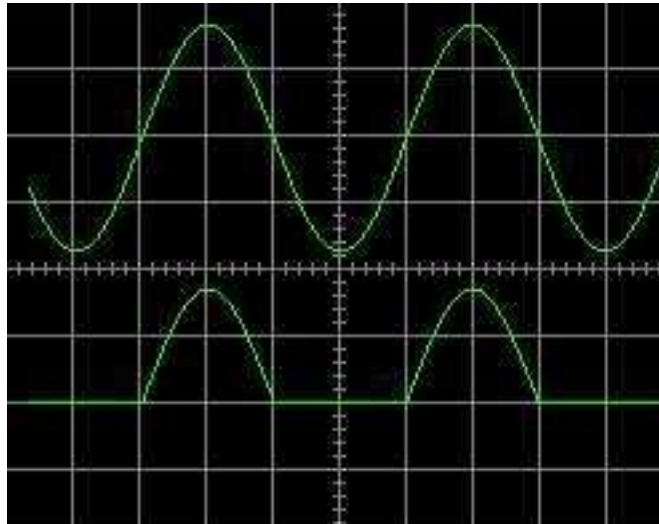
On utilise habituellement des diodes à jonction pour le redressement des tensions alternatives.

Les diodes à jonction sont constituées par deux petits morceaux de semi-conducteurs (en général du silicium). L'un des morceaux est de type N (négatif) car il a été dopé par adjonction d'une impureté qui lui donne une majorité de porteurs de charges négatifs (électrons) tandis que l'autre morceau de type P a été dopé pour avoir des porteurs majoritaires positifs (trous). Ces deux morceaux sont soudés pour former une jonction.

Une diode ne laisse passer le courant que dans un sens.

Si une diode est placée en série dans un circuit soumis à une tension alternative, le courant ne passera que pendant l'une des deux alternances: il sera redressé.

Remarque: la première diode (1905 inventée par John Fleming) était un tube électronique (diode à vide)

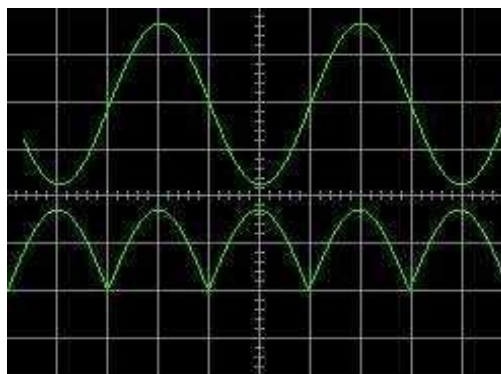


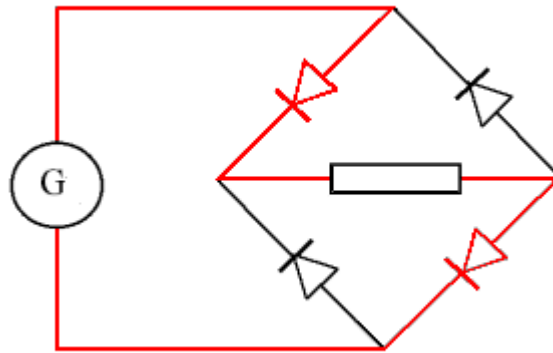
Redressement des 2 alternances

La tension redressée mono alternance est moins efficace que la tension alternative, puisque le courant ne circule que la moitié du temps.

En utilisant 4 diodes habilement connectées, on peut redresser les 2 alternances et augmenter ainsi l'efficacité.

Remarque: Lorsqu'une diode est traversée par le courant on observe une chute de tension de l'ordre de 0,7V à ses bornes. Dans le pont de Graetz, la chute de tension sera donc de $2 \times 0,7V = 1,4V$





Pont de 4 diodes (pont de Graetz)

Lissage d'une tension redressée.

Une tension redressée a toujours le même signe mais elle n'est pas continue puisqu'elle varie de 0 à U_m .

Pour obtenir une tension continue, il reste une étape: le lissage. Il consiste à empêcher les variations brutales de tension.

Fonctionnement d'un condensateur

Un condensateur est un réservoir à charges électriques.

Il est constitué de 2 armatures (surfaces conductrices) séparées par un isolant (diélectrique)

Un condensateur peut être réalisé par deux plaques métalliques séparées par de l'air. Certains condensateurs sont réalisés par des feuilles métalliques séparées par une couche d'isolant ou par une film isolant sur les faces duquel on a déposé deux couches métalliques. L'ensemble est enroulé en cylindre pour limiter l'encombrement.



Charge d'un condensateur

Si on relie un condensateur à un générateur continu, on observe le passage d'un courant dont l'intensité diminue rapidement et s'annule après un temps en général assez bref. Le condensateur se charge.

La caractéristique essentielle d'un condensateur (comme celle d'une réserve) est sa capacité C . Elle s'exprime en farads (F)

Un condensateur qui possède une charge q (coulombs) lorsque la tension à ses bornes est U (volts) a une capacité C (farads):

$$C = q/U$$

Exemple: Un condensateur de $5 \mu\text{F}$, soumis à une tension de 10 V a une charge

$$q = C \times U = (5 \times 10^{-6} \text{ F}) \times 10 \text{ V} = 5 \times 10^{-5} \text{ coulombs}$$

Décharge d'un condensateur

Si on branche une DEL aux bornes d'un condensateur chargé de forte capacité, on observe le fonctionnement de la DEL pendant quelques secondes: le condensateur se décharge.

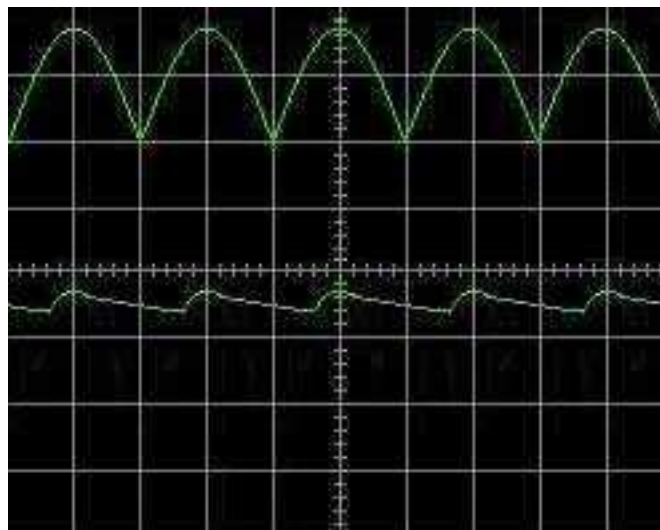
Filtrage de la tension redressée

Un condensateur est placé en dérivation à la sortie du pont de redressement.

Lorsque la tension augmente, le condensateur se charge.

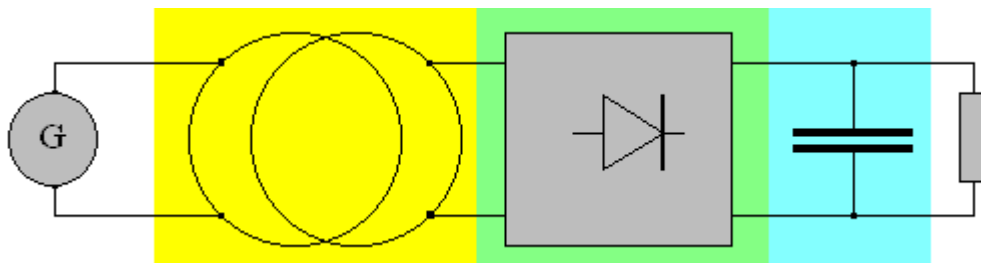
Lorsque la tension à la sortie tend à diminuer, le condensateur se décharge ce qui réduit fortement la chute de la tension.

Si le condensateur a une capacité suffisante, les variations de la tension peuvent être négligeables, la tension est quasiment continue.



En résumé.

Un adaptateur secteur est constitué d'un transformateur qui abaisse la tension alternative, suivi d'un pont de diodes qui redresse cette tension et d'un condensateur qui lisse la tension redressée.



Electronique des semi-conducteurs

Définissons tout d'abord le mot « électronique » : l'électronique est la science et la technologie du passage de particules chargées dans un gaz, dans le vide ou dans un semi-conducteur. On note donc qu'un mouvement de particules dans un métal n'est pas considéré comme de l'électronique mais comme de l'électricité.

L'histoire de l'électronique se divise essentiellement en deux parties, le passé et le présent. Le passé est l'ère des tubes à vide ou à gaz, on considère qu'elle a commencé en 1895 lorsque H.A. Lorentz découvrit théoriquement l'électron, deux ans plus tard J.J. Thompson le trouve expérimentalement et la même année Braun construit le premier tube électronique. Le présent commence en 1948 par l'invention du transistor, c'est cette partie qui nous intéresse ici.

Définition du semi-conducteur :

Se dit d'un corps non métallique qui conduit imparfaitement l'électricité, et dont la résistivité décroît lorsque la température augmente. (petit Larousse illustré).

En simplifiant, on dira qu'ils sont isolants à basse température et conducteurs à température élevée.

Les semi-conducteurs usuels, le germanium (Ge) et le silicium (Si) appartiennent à la colonne IVB de la classification périodique des éléments, ils ont quatre électrons périphériques.

Dopage d'un semi-conducteur :

Le dopage est l'introduction dans un semi-conducteur d'un corps étranger appelé dopeur. Les dopeurs utilisés sont des éléments voisins de la colonne IIIB (bore, aluminium....) ou VB (azote, phosphore....).

IIIB	IVB	VB
B	Si	N
Al	Ge	P
Ga		As
In		Sb

La quantité de dopeur introduite est très faible, de l'ordre d'un atome de dopeur pour

un million d'atome de semi-conducteur.

Après dopage, la conductibilité est essentiellement due à la présence du dopeur : la conductibilité est **extrinsèque**.

Il existe deux types de dopage, le type N et le type P.

Le type N

Le dopeur appartient à la colonne VB, il possède 5 électrons périphériques, et comme tous les atomes, il est électriquement neutre.

L'atome du dopeur s'intègre dans le cristal du semi-conducteur, mais pour assurer les liaisons entre atomes voisins, seuls quatre électrons sont nécessaires, le cinquième est donc en excès : **il n'y a pas de place pour lui.**

Notons que cet électron est en excès au point de vue place, mais pas en tant que charge, le cristal reste neutre.

L'électron excédentaire est disponible et à température ambiante il peut participer à la conduction.

En quittant son atome, l'électron excédentaire y laisse une charge positive liée au noyau :

Un ion positif mais pas de place libre pas de trou

Le type P

Le dopeur appartient à la colonne IIIIB, son atome possède 3 électrons périphériques et est électriquement neutre.

L'atome du dopeur s'intègre dans le cristal du semi-conducteur, pour assurer les liaisons entre atomes voisins, quatre électrons sont nécessaires alors que le dopeur n'en apporte que trois : une place d'électron est donc inoccupée et il y a un **trou disponible** au voisinage de l'atome dopeur.

Notons que comme dans le type N, le trou est inoccupé, mais sans charge, le cristal reste neutre.

Le trou apporté par le dopeur est disponible et susceptible de recevoir un électron.

Lorsque le trou du dopeur est occupé par un électron, nous aurons un électron excédentaire par rapport au noyau, donc un **ion négatif**.

La jonction P-N

Si nous rapprochons deux semi-conducteurs dopés, l'un de type N et l'autre de type P, de chaque côté de la surface de contact, les deux semi-conducteurs gardent leurs caractères propres.

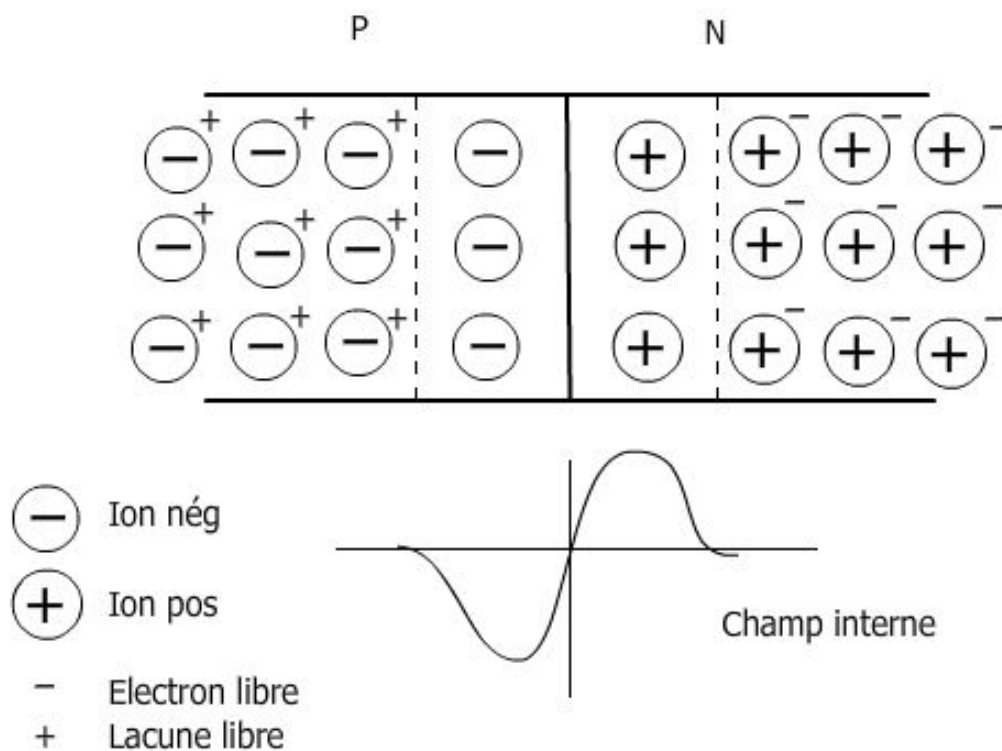
Au contraire, si les deux composants SCP et SCN sont réalisés dans un cristal unique (donc sans discontinuité de matière), il se produit des modifications interne dans une région de faible épaisseur, nous avons créé **une jonction**.

Fonctionnement d'une jonction P-N

- nous avons un grand nombre d'électrons libres dans la zone N et un grand nombre de lacunes vides (trous) dans la zone P.
- les électrons libre de N vont diffuser vers P, les lacunes libres de P vont diffuser vers N, ils vont tous deux se recombiner.
- Du côté N, lorsque les électrons ont quitté cette zone, ils ont laissé des ions +, il en résulte la création d'une charge spatiale positive près de la jonction.
- Inversement du côté P les lacunes libres ont laissé des ions -, cause d'une charge spatiale négative près de la jonction.
- Ces charges ne peuvent pas se recombiner puisque les ions sont **fixes**.

Cette double charge d'espace fait apparaître un champ électrique à cause interne.

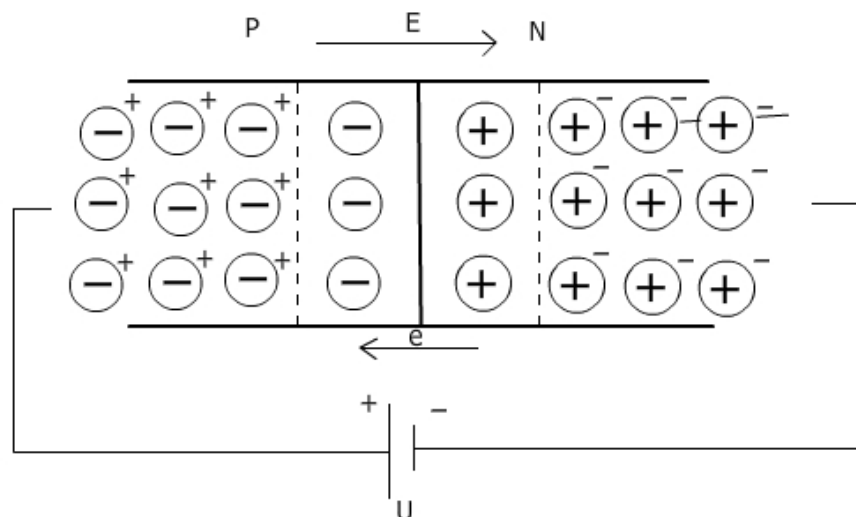
Ce champ électrique s'oppose à la diffusion des lacunes vers N et des électrons vers P → nous avons une barrière de potentiel qui refoule les lacunes vers P et les électrons vers N.



Jonction P-N reliée à un circuit extérieur

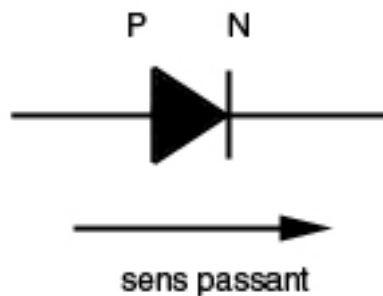
1/ sens passant :

La tension U crée un champ électrique E du + vers -, il s'oppose au champ interne e . Comme $E \gg e$, la barrière de potentiel est surmontée et les électrons peuvent diffuser vers les lacunes $\rightarrow$ la jonction est polarisée dans le sens passant, le courant traverse la jonction.



Si on inverse l'alimentation, E vient renforcer le champ interne et aucun courant ne passe la jonction est polarisée dans le sens bloqué.

Ce genre de jonction est un composant électronique que l'on appelle une **diode** .



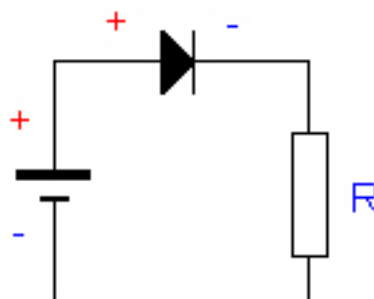
Caractéristique d'une diode.

Les caractéristiques courant-tension en direct et en inverse se représentent sur le même schéma, mais pas à la même échelle.

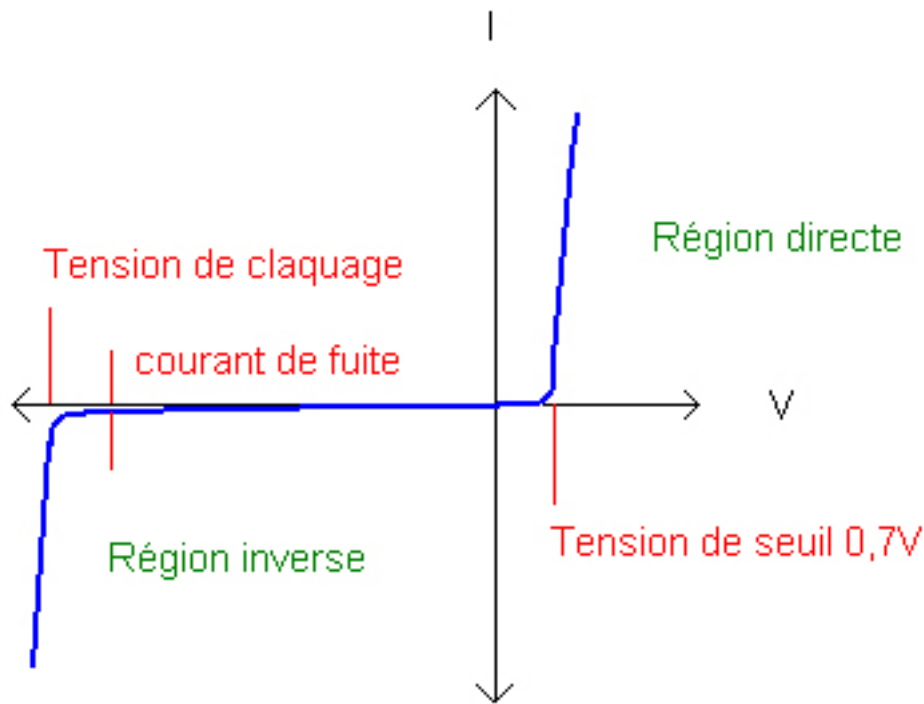
Voici le montage.

Remarquez que nous avons connecté le + de l'alimentation à l'anode de la diode et le - par l'intermédiaire de la résistance à la cathode. Ce branchement provoquera la circulation du courant, on dira que la diode est polarisée pour le sens passant.

Si nous avions adopté l'autre sens (le + sur la cathode) nous aurions polarisé notre diode en inverse et aucun courant n'aurait circulé, notre diode aurait été bloquée et polarisée pour le sens non passant.



Nous allons faire varier la tension du générateur de 0 à + V_{cc} (nouveau terme indiquant la tension maximum d'alimentation continue) en relevant à chaque fois le courant qui circule dans le circuit et la tension aux bornes de la diode. Une fois ceci effectué, nous inverserons les pôles du générateur et pratiquerons de même. Ces relevés nous permettront d'établir graphiquement la caractéristique tension-courant de la diode.



Quand la tension aux bornes de la diode est inférieure à 0,7 V, aucun courant ne circule dans le circuit, c'est comme si nous avions un interrupteur ouvert. A 0,7 V, brutalement le courant apparaît. Si nous augmentons la valeur de la tension fournie par le générateur, la tension aux bornes de la diode reste sensiblement constante et égale à 0,7 V. On appellera cette tension, la tension de seuil. Cette tension de seuil est de 0,7 V pour le silicium et 0,2-0,3 V pour le germanium.

Passons dans la région inverse. Nous constatons que la diode, polarisée en inverse ne conduit pas et donc qu'aucun courant ne circule dans le circuit hormis un léger courant de fuite de quelques μA que l'on pourra négliger. (du aux porteurs minoritaires)

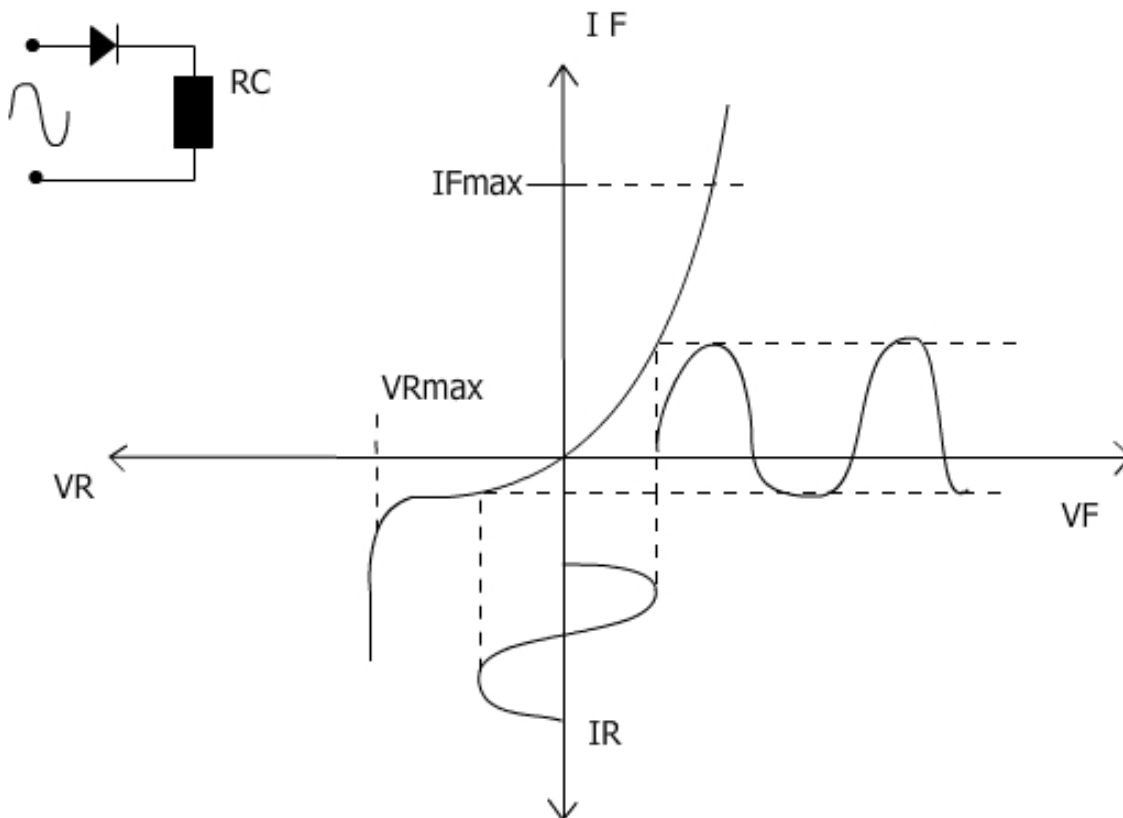
Brutalement, la diode, toujours polarisée en inverse se met à conduire et le courant circule. La tension à partir de laquelle une diode polarisée en inverse conduit s'appelle la tension de claquage. Sur une diode non prévue pour cela, c'est destructif, sachez toutefois que cet effet est exploité dans les diodes Zener.

Pour calculer le courant circulant dans le circuit, on applique la lois d'ohm.

Si $V_{CC}=10\text{V}$, $R=100\Omega$

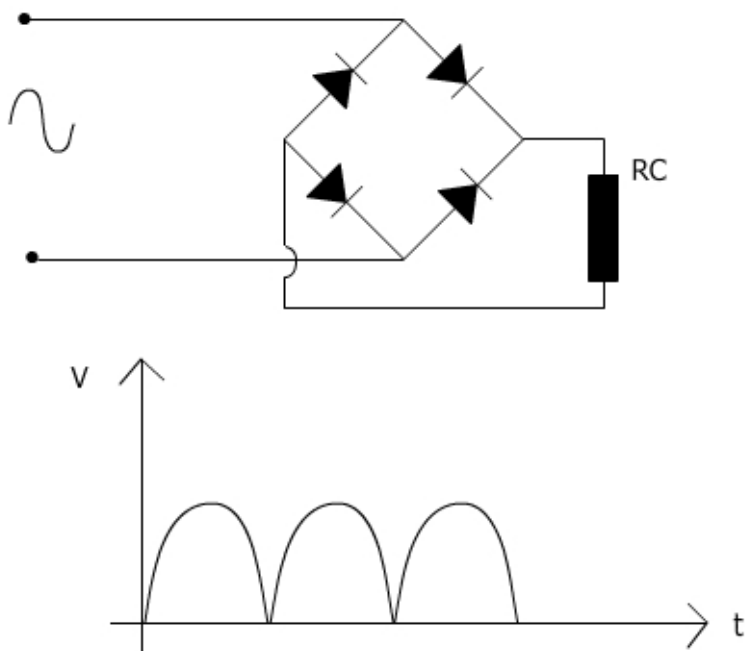
$$I=U/R \quad I=(10-0,7)/100 = 0,093 \text{ A}$$

La diode en redressement mono alternance.

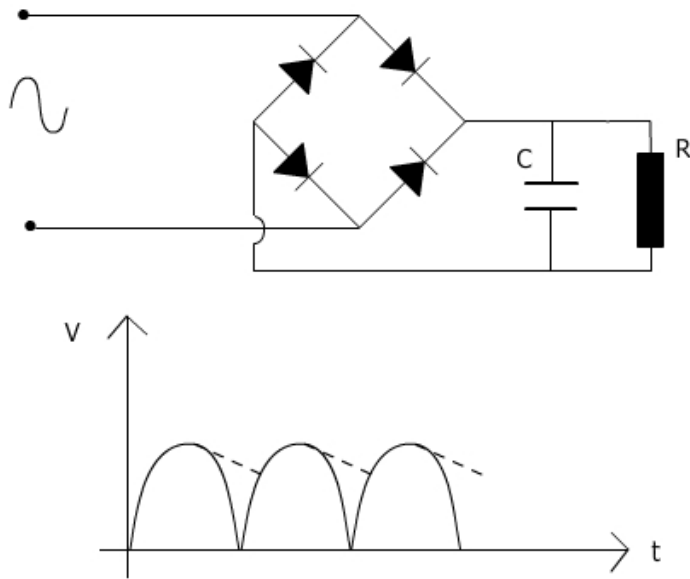


l'amp
litude de la tension alternative ne doit pas dépasser V_{Rm} .
L'alternance négative est supprimée.

Redressement double alternance.



Filtrage en double alternance.



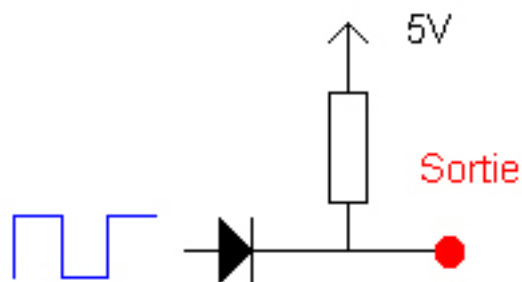
Le composant fondamental du filtrage est le condensateur, il joue le rôle de réservoir de tension, il se charge lorsque la tension qui lui est appliquée est supérieure à sa tension, il se décharge dans le circuit dans le cas contraire.

Comme on le voit dans le schéma ci-dessus (en pointillé), le condensateur C se charge durant la phase montante et se décharge durant la descente.

En plaçant plusieurs condensateurs différents en parallèles, on peut obtenir une tension continue.

Utilisation en limiteur

Vous avez des créneaux de tension de 10 V que vous voulez réduire à une amplitude proche de 5V. Voici une méthode simple :

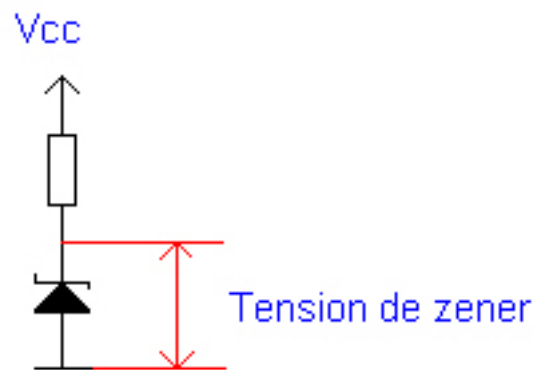


Vous aurez en sortie des créneaux allant de 5.7 à 10 V d'amplitude soit une excursion de 4,3V

La diode comme stabilisateur de tension

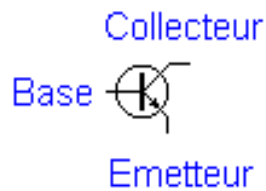
Nous ferons appel à une diode spéciale appelée diode Zener.

Remarquez que cette diode se polarise en inverse pour exploiter la tension de claquage.



Le transistor.

Symbole :



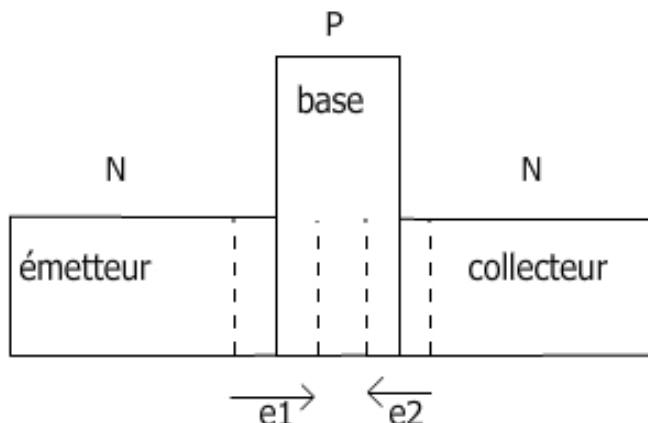
Un transistor est formé de 2 jonctions de polarités opposées placées en série.
Pour que l'effet transistor se manifeste, il faut que l'ensemble soit constitué par le même monocristal et que la partie centrale (la base) soit très mince.

Un transistor peut être formé de deux zones N séparées par une zone P : c'est un transistor du type NPN (le plus utilisé)

Soit de deux zones P séparées par une zone N : c'est un transistor PNP.

Fonctionnement.

Comme dans le cas de la diode, il s'établit à chaque jonction une barrière de potentiel,
Les deux champs internes e_1 et e_2 sont de sens opposés.



Si l'on branche une source extérieure entre le collecteur et l'émetteur, il ne se passera rien, en effet il y aura toujours une jonction polarisée dans le sens bloquant.

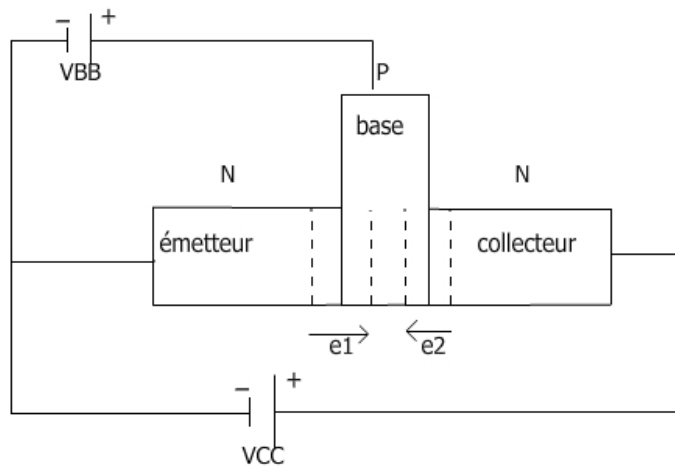
Pour obtenir le fonctionnement du transistor, il y a lieu de brancher 2 alimentations (polariser le transistor).

VCC : pôle positif au collecteur (tension beaucoup plus élevée que VBB)

Pôle négatif à l'émetteur

VBB : pôle positif à la base (tension beaucoup plus basse que VCC)

Pôle négatif à l'émetteur



La jonction B-E doit toujours être polarisée dans le sens passant.

L'émetteur est une zone fortement dopée, son rôle est d'injecter des électrons dans la base. La base est faiblement dopée et très mince, elle transmet au collecteur la plupart des électrons venus de l'émetteur.

En fonctionnement, l'émetteur injecte donc des électrons dans la base, celle-ci étant mince et faiblement dopée, les électrons atteignent le collecteur, ils sont aidés par le champ électrique créé par VCC, c'est l'effet transistor, une jonction polarisée en inverse est franchie par des électrons à cause de la proximité d'une jonction polarisée dans le sens passant (émetteur / base).

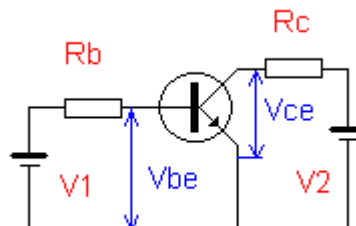
On voit donc que le courant de collecteur est commandé par le circuit émetteur / base.

La plupart des charges (98%) émises par l'émetteur sont entraînées dans le circuit collecteur (I_C) le reste forme le courant de base (I_B).

$$I_E = I_C + I_B$$

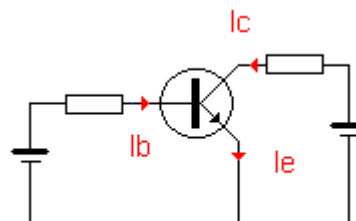
Le branchement en émetteur commun

L'émetteur et les deux sources de tension sont réunies à la masse.



Gain en courant

C'est le rapport entre le courant de collecteur et le courant de base :

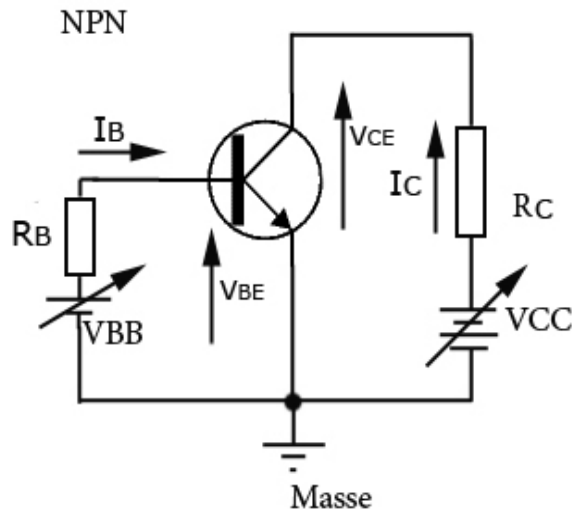


$$\beta = I_C / I_B \quad I_C = \beta \cdot I_B$$

Il est aussi appelé gain en courant et souvent désigné par h_{FE} .

Caractéristiques des transistors.

Montage :

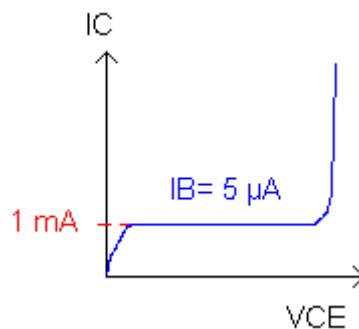


Caractéristique de sortie :

Les sources de tension sont variables, nous les ferons varier et nous mesurons à chaque fois U et I .

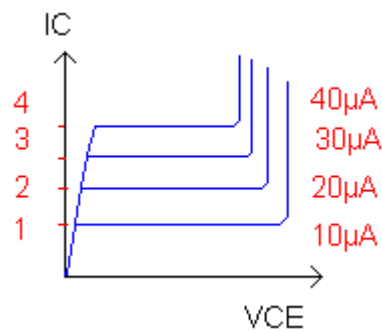
Premier test : Fixons le courant de base à $5\mu A$, nous obtenons un courant de collecteur de $1mA$.
Faisons maintenant varier V_{CC} , à chaque fois on mesure I_C et V_{ce} .

Résultat : Quand on fait varier la tension de collecteur, le courant I_C reste constant sur une grande plage.



Deuxième expérience : On refait la même chose, mais en augmentant à chaque fois le courant de base.

Résultat : On constate que le gain en courant β est sensiblement constant (I_C/I_B) et que plus V_{ce} croit, plus la partie rectiligne diminue.

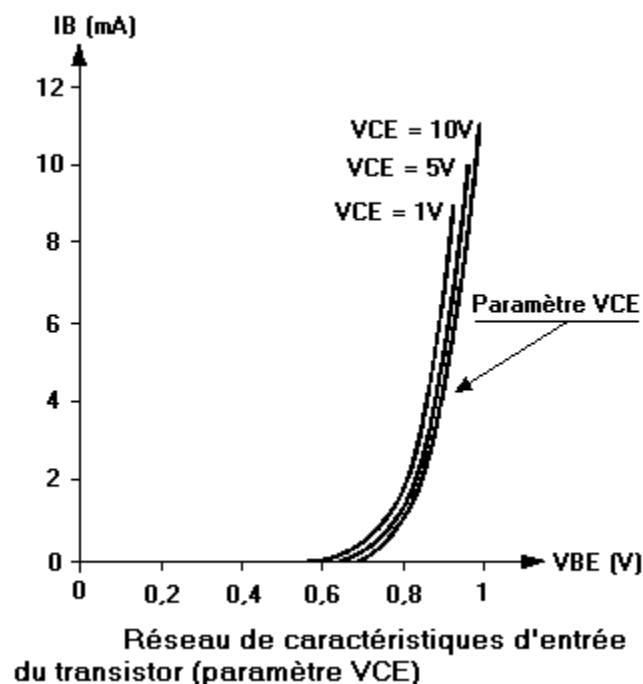


Ce dernier schéma représente la caractéristique du collecteur. (ou de sortie)

Nous allons détailler l'examen de ce réseau.

- Toutes les caractéristiques partent du **point origine 0**. Cela signifie que lorsque **VCE** est nulle, le courant **IC** est également nul et cela quel que soit le courant **IB**.
- Ensuite, les caractéristiques présentent deux parties. La première est commune et pratiquement verticale ; la seconde est pratiquement horizontale **et est fonction du courant IB**.
- La première partie signifie que lorsque la tension **VCE** varie légèrement, le courant **IC** augmente dans des proportions importantes et d'autant plus que le courant **IB** est **élevé**.
- Quand la tension **VCE** atteint un certain seuil, relativement bas pour chaque caractéristique, il y a **un coude** et à ce moment-là **la courbe devient pratiquement horizontale**. En fait, elle est légèrement relevée, c'est-à-dire que pour chaque valeur de **IB** donnée, le courant **IC** augmente légèrement quand la tension **VCE** augmente.
- La position d'une caractéristique est ici fonction du courant **IB**, autrement dit le courant **IC** est en relation étroite avec le courant **IB**.

Caractéristique de base (ou d'entrée) :
même montage



Dans un premier temps, on fixe la tension **VCE** .

Dans un deuxième temps, on relève la caractéristique pour différentes valeurs de **IB**. Pour cela, on fait varier **VBB** et pour chaque valeur de **IB**, on relève la tension **VBE**. Cela détermine une caractéristique. Pour une seconde caractéristique, on change la valeur de **VCE** et on recommence la même série de mesures. On trace ainsi plusieurs caractéristiques .

Ces caractéristiques sont analogues à celle d'une diode. En effet, la jonction base-émetteur est polarisée en direct.

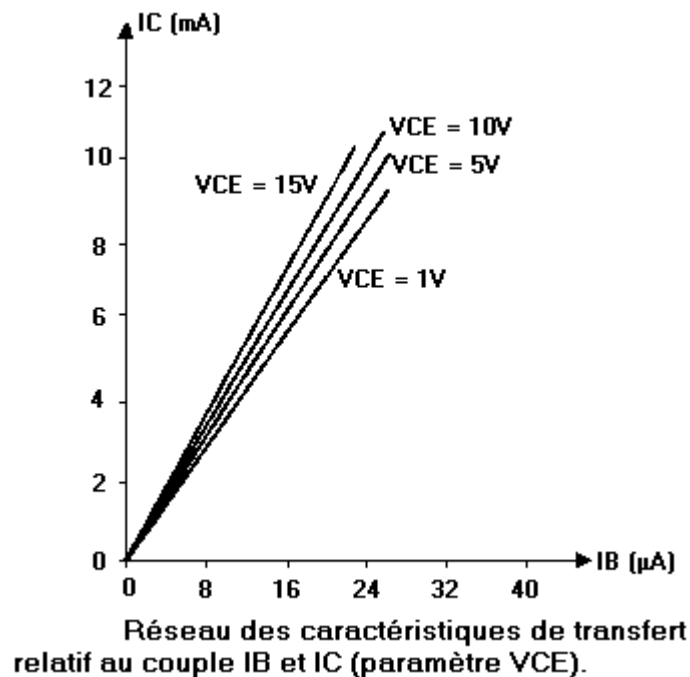
Toutefois, quand **IB** est nul (base en l'air), la tension **VBE** n'est pas nulle, mais vaut environ **0,6 volt**.

Vous constatez également que la tension **VCE** a très peu d'influence sur la tension **VBE**. Si pour un même courant de base **VCE** passe de **1 volt à 10 volts**, la tension **VBE ne varie que de quelques dizaines de mV**.

En conséquence, les fabricants se contentent de fournir une seule caractéristique d'entrée, correspondant à une valeur moyenne de **VCE**.

Caractéristique relative à IC et IB avec VCE comme paramètre
même montage

Dans un premier temps, on fixe la tension **VCE** . Puis dans un second temps, on règle **VBB** de façon à faire varier le courant **IB**. Il suffit de relever la valeur du courant **IC** pour chaque valeur particulière du courant **IB**.



Ces caractéristiques sont appelées **caractéristiques de transfert** car elles mettent en relation le **courant de sortie IC** avec le **courant d'entrée IB**.

Vous remarquez que **IC** est proportionnel à **IB** (à **VCE** constante).

Plus VCE est élevée, plus le coefficient de proportionnalité entre IC et IB est élevé.

En réalité, les caractéristiques de transfert ne sont pas exactement des droites, mais pratiquement on peut les assimiler à des droites.

DIAGRAMME GÉNÉRAL DES CARACTÉRISTIQUES D'UN TRANSISTOR

Il est possible de réunir les trois réseaux de caractéristiques vus précédemment.

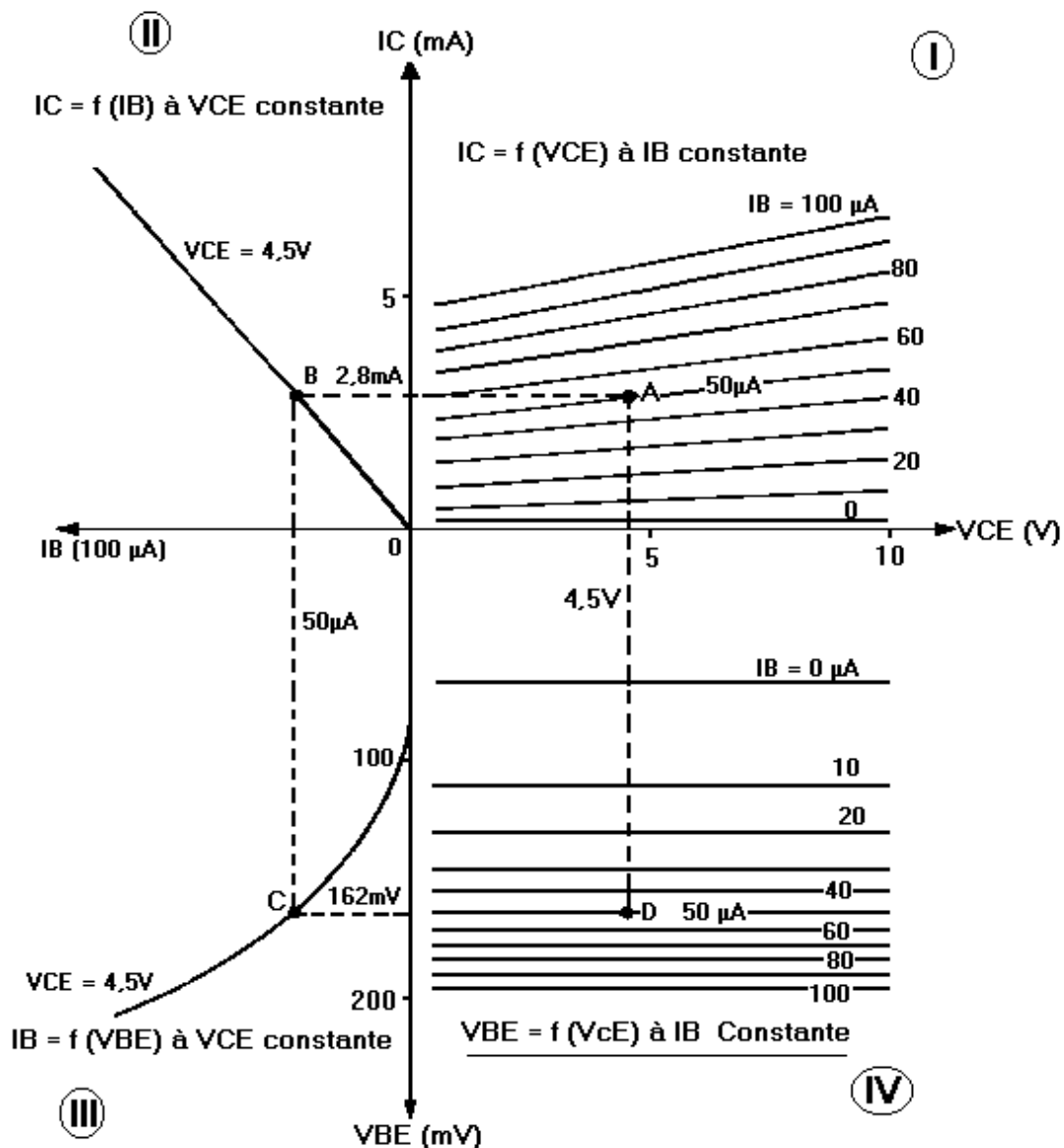


Diagramme des caractéristiques d'un transistor NPN monté en émetteur commun.

Polarisation des transistors.

Notre transistor, pour fonctionner, a besoin d'être "polarisé". Cela signifie qu'on doit appliquer sur ses connexions les tensions correctes et en amplitude et en polarité pour qu'il effectue la fonction qu'on lui demande.

Quand nous parlons de polarisation, nous parlons uniquement de tensions continues, et ce sont ces tensions continues qui vont permettre le fonctionnement correct en alternatif.

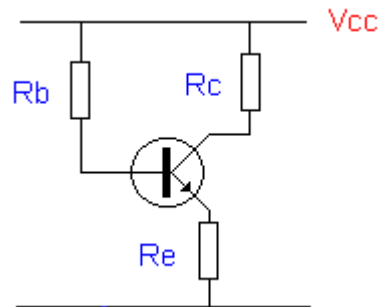
Quand nous utiliserons la fonction amplification par exemple, nous appliquerons un signal alternatif à l'entrée et nous le récupérerons agrandi à la sortie, ceci ne sera possible que si les tensions continues sont présentes.

Tout d'abord une chose importante :

Le gain en courant β du transistor est fortement affecté par la température. Quand la température du transistor croît, le gain β croît. Ce phénomène peut conduire à l'emballement thermique (β croît donc I_C croît, la température du transistor croît, ce qui provoque une augmentation de β etc.)

Le plus gros problème contre lequel nous devons lutter pour polariser correctement un transistor utilisé en amplificateur linéaire est la variation du gain en courant β avec les variations de température.

Pour remédier à ce problème, nous utiliserons la résistance d'émetteur.



(Dans ce type de schéma, le rail supérieur représente la tension positive d'alimentation notée V_{cc} , le rail inférieur représente la référence du 0V, c'est à dire la masse.)

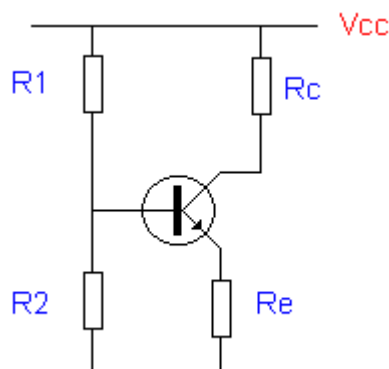
Supposons que pour une raison quelconque, le courant I_c croisse.

Le courant I_e ($I_e = I_c + I_b$) va croître également.

la chute de tension aux bornes de la résistance R_E va augmenter puisque $V_{re} = R_e \cdot I_e$
Or $v_{re} = V_e$ (tension émetteur – masse), V_b (tension base – masse) est constante et ne dépend que de R_b , résumons : V_e augmente, V_b est fixe donc V_{be} diminue. Puisque V_{be} diminue, I_b diminue aussi, ce qui fait décroître I_c !!

En montage émetteur commun, la diode émetteur est polarisée en direct et la diode collecteur est polarisée en inverse.

Polarisation par diviseur de tension:



Le courant circulant dans la base est faible par rapport au courant circulant dans R_1 et R_2 , on utilise donc le théorème du pont diviseur pour obtenir la tension aux bornes de R_2 :

$$V_2 = V_{cc} \cdot \frac{R_2}{(R_1 + R_2)}$$

nous avons sur la base une tension dictée par le pont diviseur R_1/R_2 .

Sur l'émetteur, nous retrouvons cette tension diminuée de V_{BE} soit inférieure de 0,7V (c'est une diode).

Le courant $I_E = I_C + I_B$, négligeons I_B qui est de toute façon très petit, on peut approximativement dire que $I_C = I_E$.

Appliquons la loi d'Ohm pour la résistance d'émetteur :

Le courant qui traverse R_E sera égal à la tension à ses bornes divisée par la valeur de sa résistance, donc

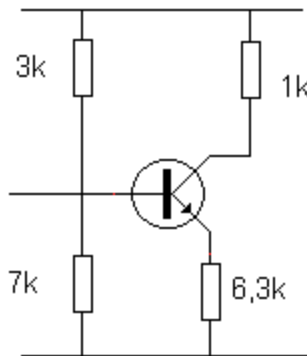
$$I_E = U_{RE} / R_E \text{ soit}$$

$$I_E = V_E / R_E$$

C'est donc R_E qui règle le courant de collecteur.

En pratique :

La tension d'alimentation est de 10V



1 - calculons la tension délivrée par le pont résistif R_1, R_2

Nous savons que

$$V_b = \frac{R_2}{R_1 + R_2} \cdot V_{cc}$$

ce qui donne

$$V_b = \frac{7}{10} \times 10 = 7V$$

2 - calculons la tension présente sur l'émetteur du transistor

$$V_e = V_b - V_{be}$$

$$V_e = 7 - 0,7 = 6,3 V$$

3 - calculons le courant d'émetteur qui sera égal au courant collecteur

$$I_e = U_e / R_e$$

$$I_e = V_e / R_e$$

$$I_e = 6,3 / 6300 = 1 \text{ mA}$$

Le courant collecteur étant à I_b près égal à I_e

4 - calculons la chute de tension aux bornes de la résistance de collecteur

$$U = R_c \cdot I_c$$

On prendra $I_c = I_e$

$$U_{rc} = 1000 \times 0.001 = 1V$$

5 - calculons V_{ce}

$$V_{ce} = V_{cc} - R_c \cdot I_c - R_e \cdot I_e$$

$$V_{ce} = 10 - 1 - 6.3 = 2,7 V$$

6 - notre point de repos est positionné comme suit

$$I_c = 1 \text{ mA}$$
$$V_{ce} = 2.7 \text{ V}$$

7 - dessinons notre droite de charge en déterminant les deux points caractéristiques

Equation globale de collecteur :

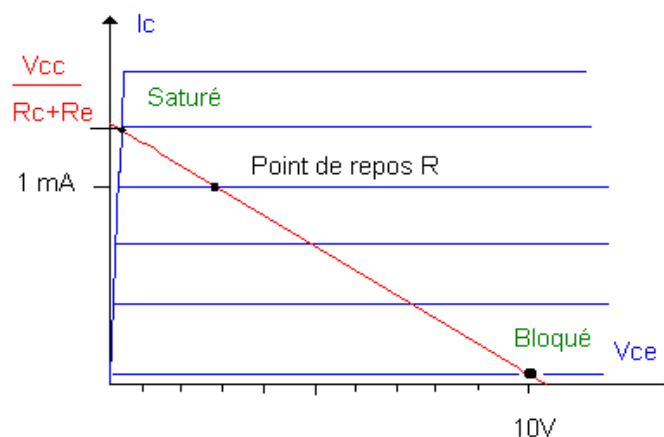
$$V_{cc} - (R_c I_c + V_{ce} + R_e I_e) = 0$$

Courant de saturation

$$\text{Quand } V_{ce} = 0, I_c = V_{cc} / (R_c + R_e)$$

$$I_c \text{ pour } V_{ce}=0 = 10 / 7300 = 1.37 \text{ mA}$$

$$\text{Quand } I_c = 0 \text{ } V_{ce} = V_{cc} \text{ soit } 10 \text{ V}$$



On constate que notre transistor n'est pas idéalement polarisé et qu'il serait souhaitable de faire descendre le point de repos sur la droite de charge de manière à ce que celui-ci soit positionné au milieu de celle-ci.

Nous pouvons réaliser cela de différentes manières, la plus simple consiste à augmenter la valeur de R_e de manière à diminuer le courant I_C . On pourrait également recalculer le pont diviseur pour obtenir le même résultat.

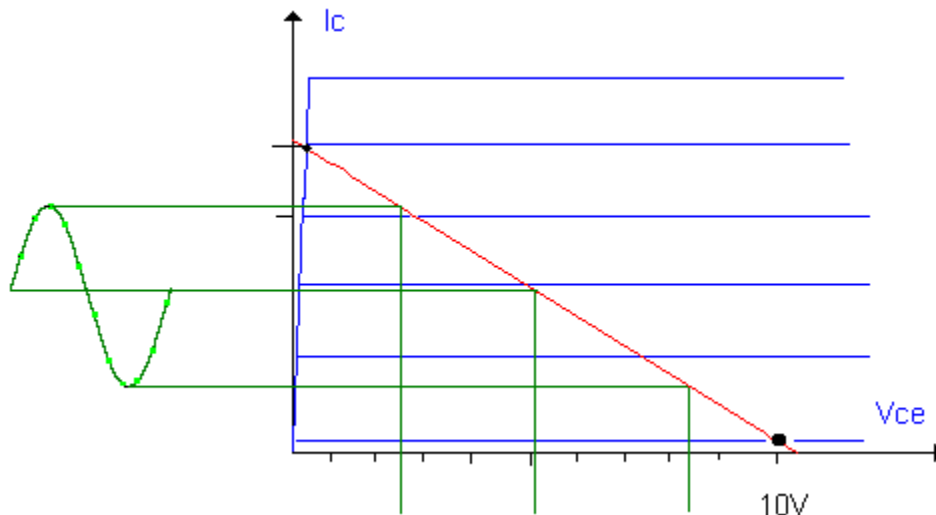
Le montage émetteur commun en amplification :

Si nous faisons bouger notre point de repos en suivant la droite de charge, le courant de collecteur variera continuellement, comme il y a courant, il y a chute de tension dans R_C . Cette chute de tension suivra le courant. C'est cette tension image du courant que l'on va récupérer.

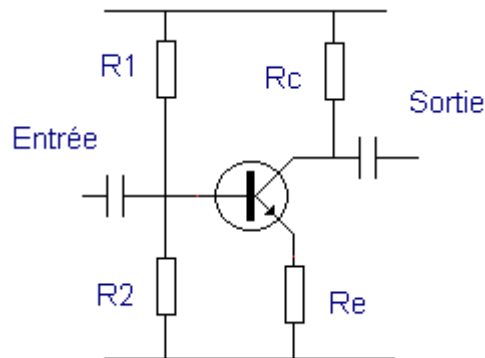
Pour déplacer le point de repos sur la droite de charge, on fait varier la tension de base.

On applique une tension alternative à la base, celle-ci va se superposer à la tension de polarisation, on va donc augmenter la tension de polarisation de la jonction base émetteur, ce qui fera croître I_B et donc I_C .

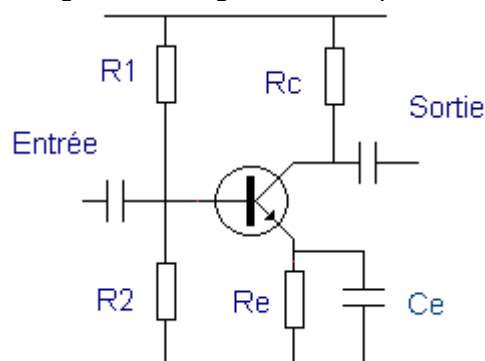
Les alternances négatives se retrancheront de I_B et donc diminueront I_C .



Pour envoyer un signal alternatif (sortie d'un micro par exemple) à l'entrée d'un transistor, on va utiliser un condensateur qui a la particularité de laisser passer l'alternatif et de bloquer le continu, et on récupérera, à la sortie, les variation d'entrée filtrées par un autre condensateur. Ces condensateurs seront appelés condensateurs de couplage.



Un deuxième condensateur est utilisé, il sera chargé de dévier les signaux alternatifs vers la masse, ce qui aura pour effet d'augmenter le gain de l'ampli.



En observant le montage, et en imaginant que Ce est déconnecté, vous remarquerez que I_e varie comme I_c . Cette variation de I_e provoque bien entendu une variation de la tension aux bornes de Re ($U_{re} = I_e \times Re$).

Cette variation tend à diminuer la polarisation de la jonction V_{be} au rythme des variation de I_e .

Branchons C_e , le condensateur élimine complètement la composante alternative, la tension U_{re} est stable, le gain croît.

la tension à amplifier est superposée à la polarisation continue.

V_{be} augmente ce qui fait croître I_c .

Quand I_c croît la chute de tension $R_c \times I_c$ croît également.

Parallèlement si $R_c I_c$ croît, la tension V_{ce} diminue.

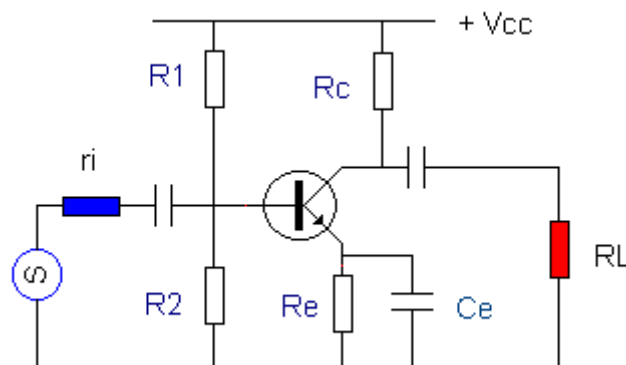
Au demi-cycle suivant c'est l'inverse qui se produit, $R_c I_c$ diminue, V_{ce} augmente.

On constate donc qu'à une augmentation de la tension d'entrée, correspond une diminution de la tension de sortie.

Attention, notez que la tension de sortie est beaucoup plus élevée que la tension d'entrée, car nous avons réalisé un amplificateur. Ici nous parlons de la phase du signal pas de son amplitude.

Le montage Emetteur Commun pour les raisons que nous venons d'expliquer déphase le signal de 180°

Montage complet :



Les montages impulsifonnels (ou logiques)

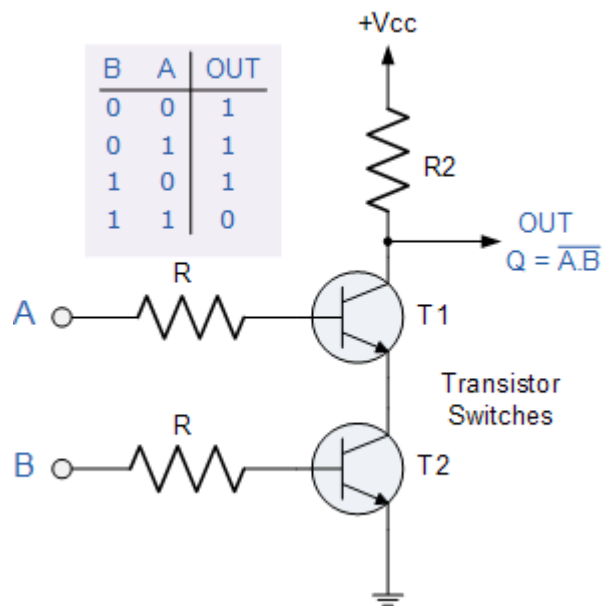
Les transistors peuvent être utilisés en mode numérique (tout ou rien, saturés, bloqués)

Exemple, la porte **nand**



C'est probablement la porte la plus utilisée en logique.

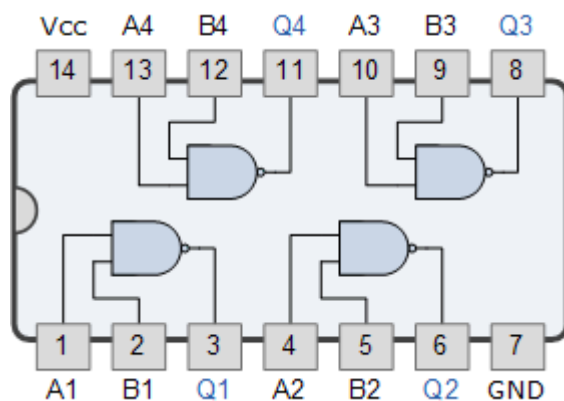
Elle est composée de deux transistors :



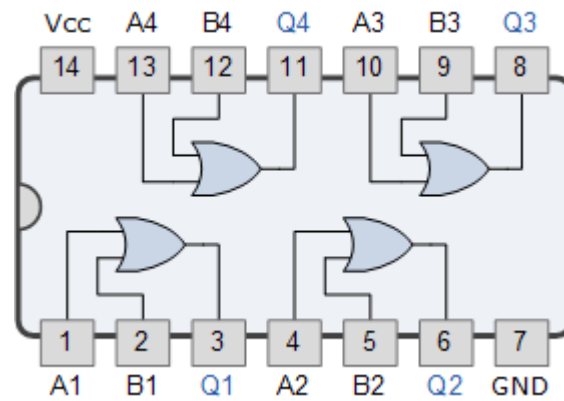
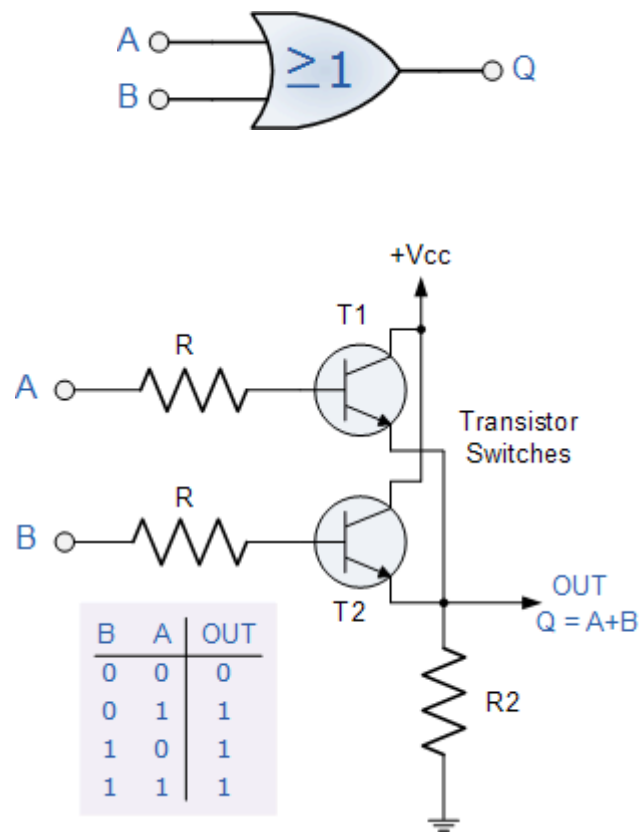
Il est assez facile de comprendre que tant que les deux bases ne sont pas alimentées, aucun courant ne traverse R2, donc la tension en Q = Vcc (niveau 1).

Si les deux bases sont alimentées, un courant traverse R2 et les deux transistors et Q = 0 (niveau 0).

Circuit intégré de type 7400

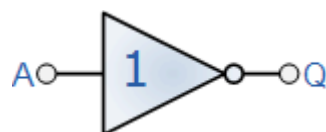


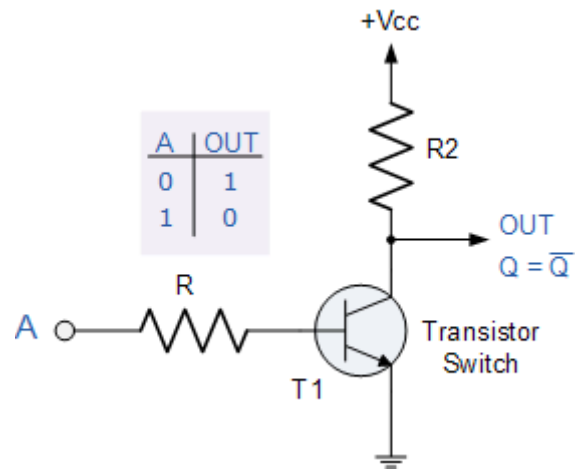
La porte **OR** (ou)



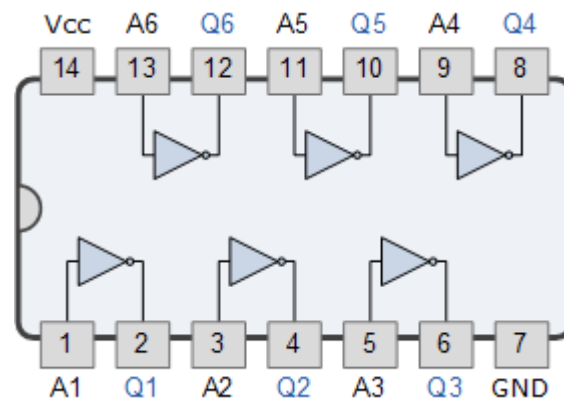
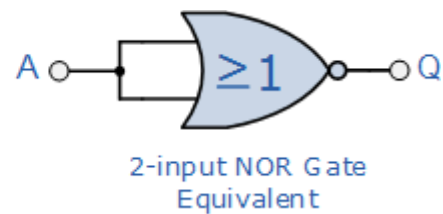
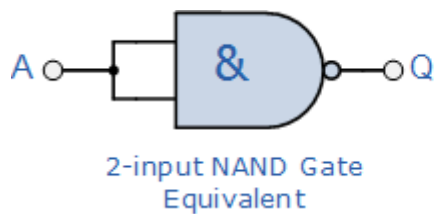
Circuit 7432

La porte **NOT** (inverseur logique)



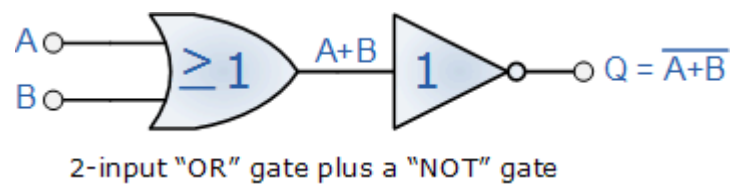
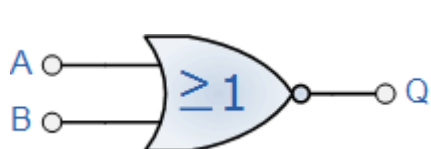


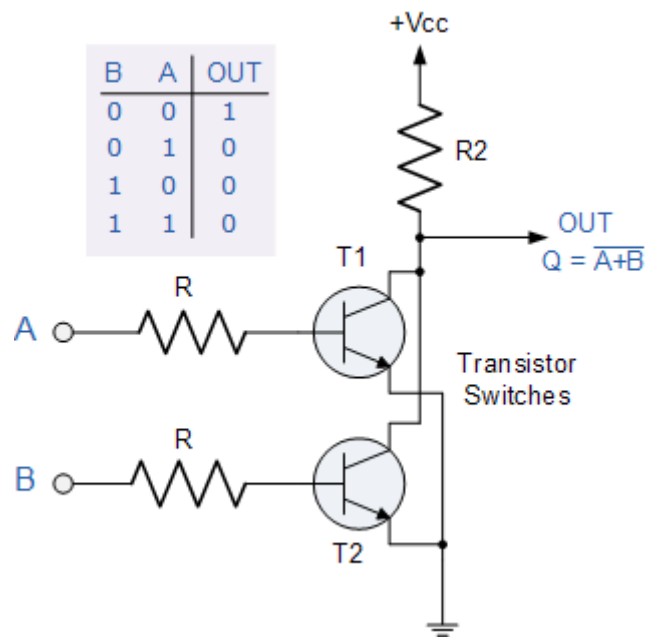
Cette fonction (comme la plupart des autres) peut être aussi réalisée à partir de portes NAND ou OR.



Circuit 7404

La porte **NOR** :





Il en existe de nombreuses autres....

Logique séquentielle .

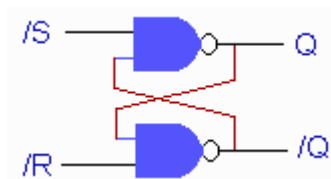
Dans la partie précédente, nous avons vu de la logique combinatoire, c'est à dire que le niveau de sortie ne dépendait que du niveau des entrées.

En logique séquentielle, le niveau de sortie dépend aussi du niveau passé des entrées.

1/ les bascules bistables :

On les appelle bistables car ces bascules ont deux états stables ce qui signifie que s'il n'y a pas d'intervention sur la bascule, celle-ci est verrouillée dans son dernier état.

Exemple, la bascule RS:



Voici une bascule RS constituée à partir de portes NAND.

La sortie Q est reliée à l'entrée /R, la sortie /Q est reliée à l'entrée /S

S (set) consiste à mettre à 1

R (reset) consiste à mettre à 0

la bascule RS est donc munie de deux entrées (R et S) et d'une sortie Q. (la sortie complémentée /Q est également disponible). Une impulsion sur S provoque le passage de Q à 1 et une impulsion sur R provoque la passage de Q à 0

Q _n	S	R	Q _{n+1}
----------------	---	---	------------------

0	0	0	0
0	0	1	0
0	1	0	1
0	1	1	?
1	0	0	1
1	0	1	0
1	1	0	1
1	1	1	?

Q_n représente l'état antérieur de la sortie, S et R sont les entrées et Q_{n+1} la sortie après action sur les entrées.

Les valeurs "?" sont des états d'incertitude.

Constat :

Quand S et R sont à 1, il y a indétermination

Cette bascule réalise une mémorisation de l'état de sortie jusqu'à une nouvelle action sur la mise à 1 ou à 0.

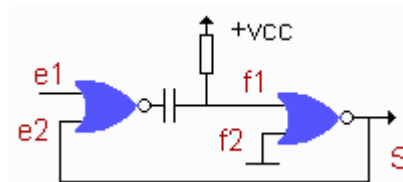
2/Les bascules monostable

Ici, il n'y a qu'un seul état stable.

Au bout d'un temps "t" ces bascules repassent sur l'état stable.

On passe de l'état stable vers l'autre état appelé "quasi stable" par une action sur une entrée.

voici une bascule monostable composée à partir de portes NOR. La constante de temps est fixée par RC. Augmenter la valeur de C et/ou de R revient à augmenter la constante de temps, donc le retour à l'état stable.



Au repos mettons l'entrée e_1 à 0, et supposons e_2 à 0, la sortie de cette porte NOR sera égale 1. L'entrée f_1 est à 1 puisque reliée à V_{cc} par l'intermédiaire de R, f_2 est à 0 puisque directement reliée à la masse, la sortie est donc à 0. Si l'on ne change rien au niveau de l'entrée e_1 , rien ne bougera, l'état est stable.

Ci-dessous l'état des entrées et sorties à l'état repos.

e1	e2	S1	f1	f2	S2
0	0	1	1	0	0

Appliquons un niveau logique 1 sur l'entrée e_1 .

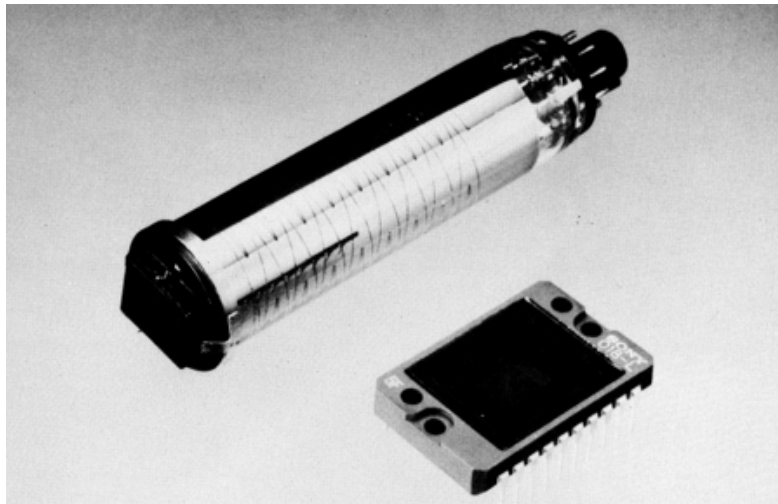
L'entrée e_2 est toujours à 1, la sortie passe S1 passe à 0 ce qui simultanément applique un 0 sur f_1 . Le condensateur se charge à travers R avec une constante de temps = RC.

La sortie S2 passe à 1.

Le condensateur s'est chargé au bout d'un temps "t", l'entrée f1 retrouve un niveau 1 et provoque le basculage de S2 vers 0, c'ad l'état stable.

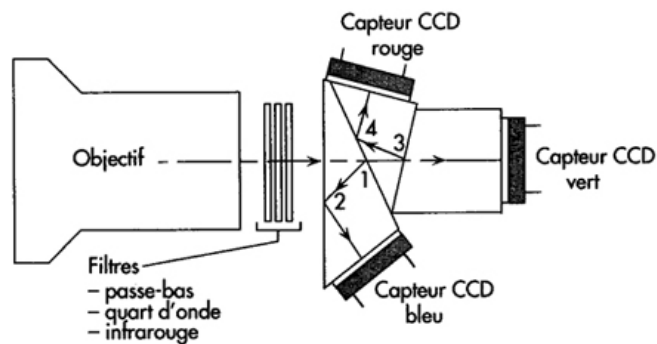
Il existe comme dans les portes, de nombreuses autres bascules comme les trigger de Schmitt, les JK et j'en passe...pour ceux que cela intéresse, le web regorge de données.

Les capteurs CCD

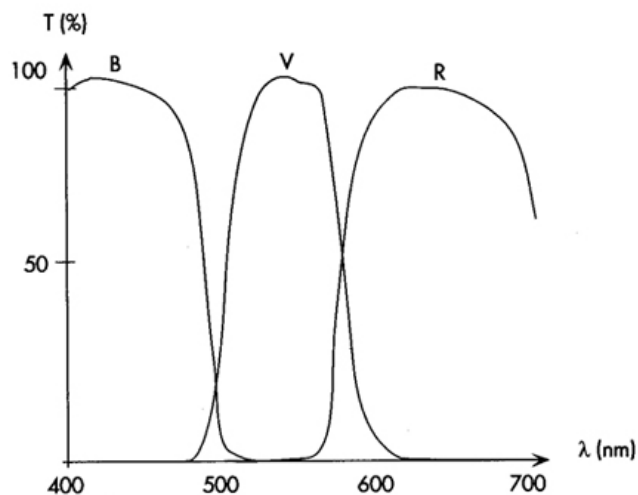


Au dessus un tube de prise de vue de type Plumbicon
En dessous un capteur de type CCD

Le séparateur optique décompose l'image en trois composantes rouge, verte et bleue.



Caractéristiques spectrales d'un séparateur optique.



Un capteur CCD est constitué de la juxtaposition de cellules élémentaires (pixels) communiquant entre elles par des portes.

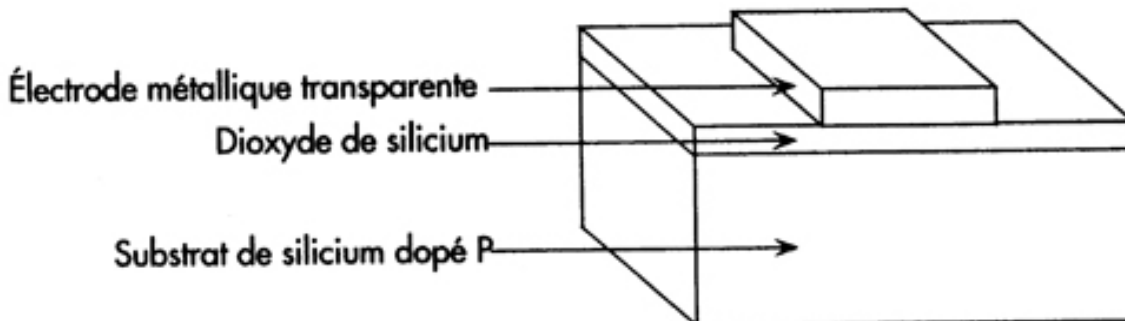
Il se présente sous la forme d'un circuit intégré présentant à sa surface une zone image, La surface de cette zone varie suivant le type de caméra :

- 12,8x9,6mm (1 inch)
- 8,8x6,6mm (2/3 inch)
- 4,3x3,2mm (1/3 inch)
- 3,2x2,4mm (1/4 inch).

Fonctionnement d'un pixel

Les premiers CCD que nous allons analyser seront de type MOS (métal oxyde semi-conducteurs).

L'élément photosensible est constitué d'un substrat en silicium dopé P sur lequel est déposé une fine couche d'isolant (oxyde) elle-même surplombée d'une électrode métallique transparente.



L'électrode est utilisée pour polariser la cellule de manière à créer dans le substrat un champ électrique interne repoussant les charges positives (trous) dues au dopage P du substrat, dans la profondeur de la cellule .

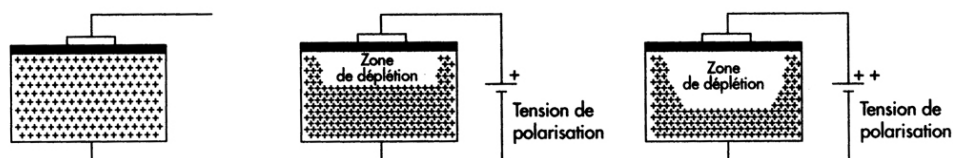
Le puit de potentiel (ou zone de déplétion) est d'autant plus important que la tension de polarisation est élevée.

C'est dans cette zone que seront attirés les électrons libérés par l'effet photoélectrique.

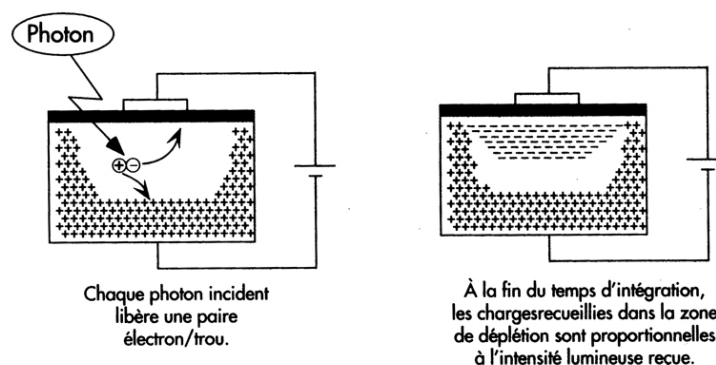
En effet, lorsqu'un rayon lumineux (composé de photons) pénètre dans le silicium, il libère une paire électron/trou.

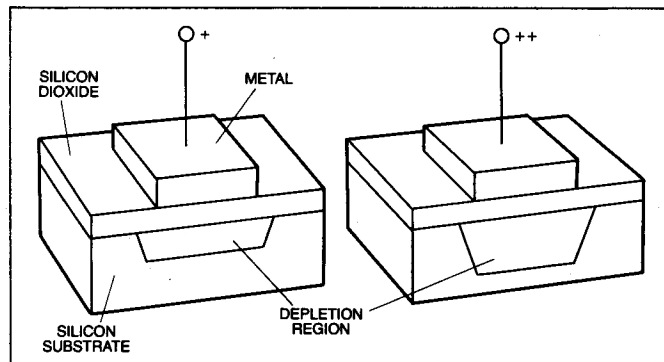
L'électron et le trou se séparent, du fait de la polarisation le trou est repoussé en profondeur et l'électron attiré par l'électrode (qu'il ne peut rejoindre du fait de la couche isolante) et reste dans la zone de déplétion.

A la fin du temps d'intégration (durée d'une trame) le nombre d'électrons accumulés est directement proportionnel à la lumière reçue par ce pixel.



La zone de déplétion est d'autant plus grande que la tension appliquée est élevée.





Le transfert des charges

Il reste maintenant à récupérer ces charges et rendre la cellule de nouveau apte à capter les charges de la trame suivante.

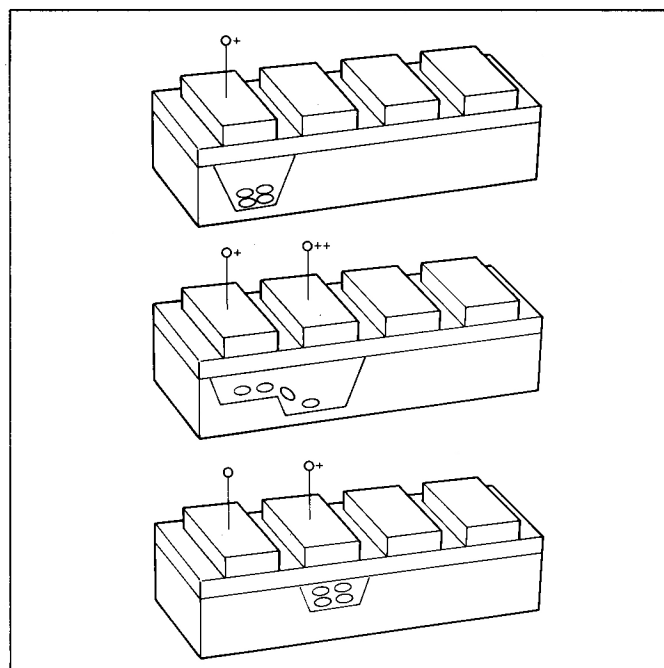
Si l'on applique une tension plus élevée à une cellule voisine, sa zone de déplétion plus grande attirera les électrons de la première cellule.

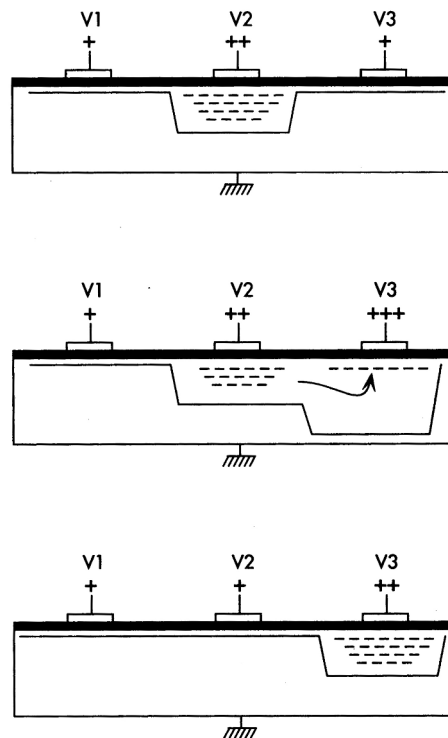
Il suffit donc d'alterner les phases d'accumulation et de transfert avec des tensions appliquées de façon séquentielle appropriée au temps et durée d'un signal vidéo.

Il existe plusieurs manières de transférer ces charges, donc plusieurs structures de CCD.

En MOS, nous avons les IT (Interline Transfert), les FT (Frame Transfert) et les FIT (Frame Interline Transfert).

Viennent ensuite les CCD de type HAD (Hole Accumulated Diode).





Principe du transfert de charge.

La structure IT

Dans cette structure, chaque cellule photosensible est accolée à une cellule non sensible à la lumière servant au stockage et au transfert (registre vertical) .

Nous avons dans ce cas des colonnes de photocapteurs alternant avec des colonnes de cellules de stockage.

Les cellules photosensibles sont séparées par des stoppeurs de canal appelés CSG (Channel Stopper Gate) empêchant la diffusion d'électrons entre cellules sensibles.

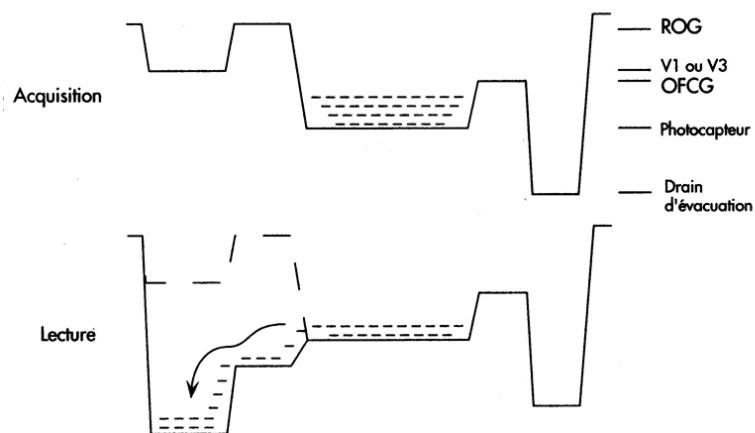
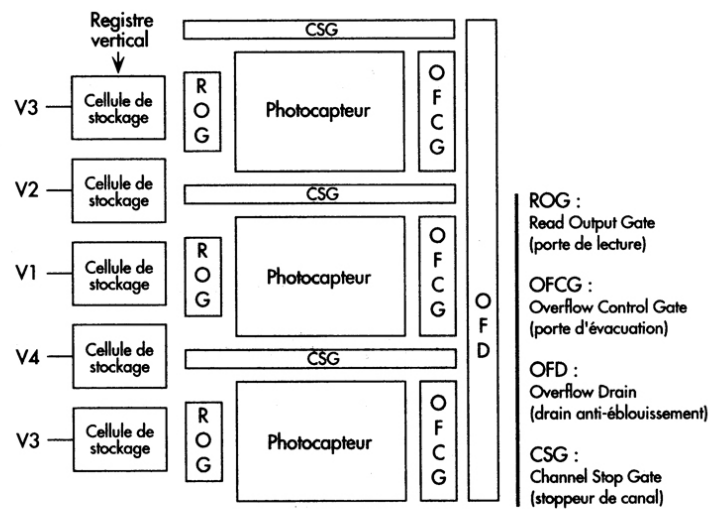
Elles sont aussi séparées d'un drain d'évacuation des charges excédentaires (OFD OverFlow Drain) par des portes OFCG (OverFlow Control Gate) .

Et enfin par des portes ROG (Read Out Gate) qui permettent de transférer les charges de la zone d'acquisition (zone sensible à la lumière) à la zone de stockage .

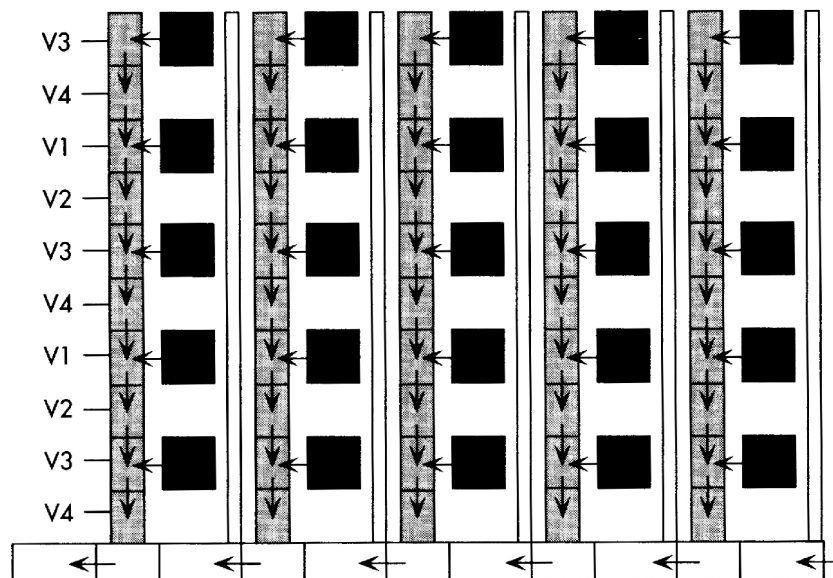
Le drain est utile en cas de forte illumination (10 fois le niveau nominal), il élimine les charges excédentaires.

Durant la période utile de la trame, l'énergie lumineuse fournie par l'optique est convertie en énergie électrique, puis au cours de l'intervalle de suppression de trame (blanking) une impulsion est appliquée aux ROGs, ce qui provoque un déplacement latéral simultané de toutes les charges accumulées (de la zone d'acquisition à la zone de stockage), les zones d'acquisitions sont donc aptes à se recharger durant la trame suivante.

Durant la durée active de la trame, à chaque synchro ligne, les charges des registres verticaux (colonne de zone de stockage) se décalent d'un cran vers le bas jusqu'au registre horizontal de sortie placé sous les registres verticaux .



CCD IT



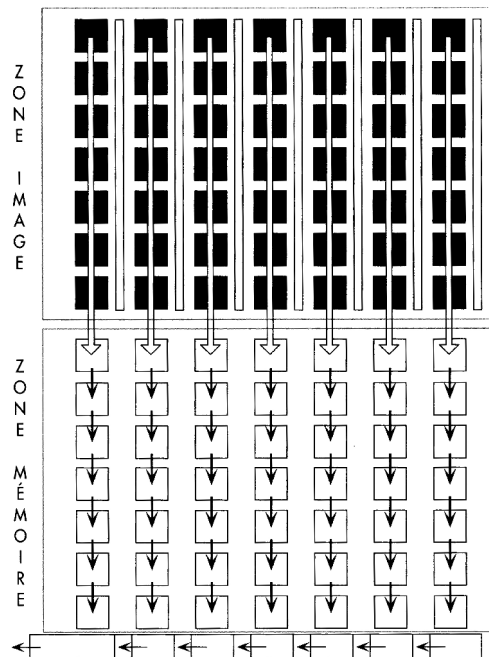
On remarque que dans un CCD de type IT, les portes registres verticaux et drain occupent une place importante à la surface du CCD, la surface sensible n'est que de 22% de la surface totale

Le SMEAR

Le smear est caractéristique des capteurs CCD (surtout IT), il se traduit par une ligne verticale blanche ou rouge traversant une zone trop lumineuse (projo , soleil...). Il est provoqué par la pollution du registre vertical par des électrons parasites générés par un excès de lumière venant s'ajouter aux données utiles au cours de leur transfert.

Deux raisons à cela : le drain est saturé et les électrons générés par des rayons lumineux de longueur d'ondes élevées (proche de l'infrarouge) peuvent pénétrer en profondeur dans le substrat et s'introduire par le bas dans les registres verticaux.

Le CCD FT



La zone image d'un capteur FT n'est constituée que de photocapteurs (pas de registre vertical).

En dessous de cette surface photosensible se trouve la zone de stockage de capacité équivalente à la zone image, à l'extrémité de laquelle se trouve le registre à décalage horizontal.

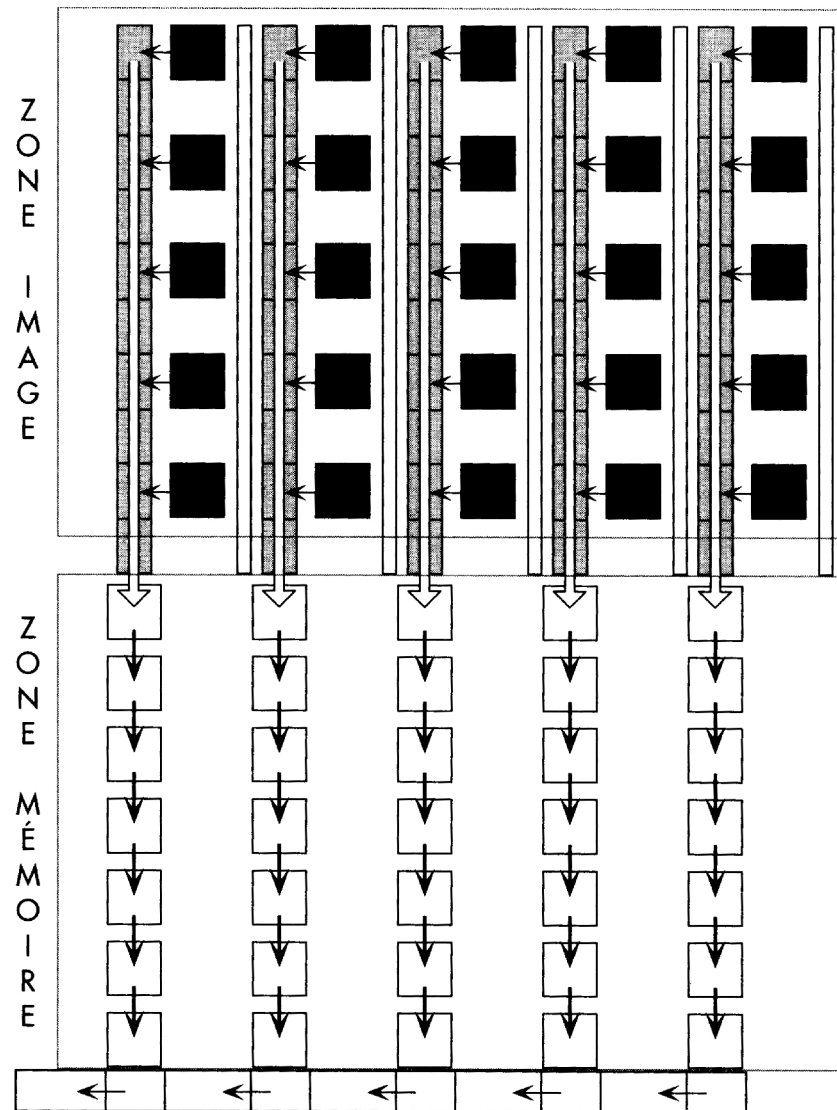
Durant la période utile de la trame, il y a accumulation d'électrons dans la zone d'acquisition, ensuite, durant le blanking, toutes les charges descendent simultanément dans la zone de stockage, ensuite, durant l'acquisition suivante, les charges sont transférées dans le registre horizontal à la fréquence ligne.

Cette structure implique de masquer les zones d'acquisition durant le transfert vertical à l'aide d'un obturateur (svt mécanique comme en cinéma).

Le grand avantage des FT est que presque toute la zone se situant derrière l'objectif est sensible.

Autre avantage, la suppression des registres verticaux induit l'absence de smear.

Le CCD FIT



La structure d'un FIT est une combinaison des IT et des FT , elle associe les registres verticaux et la zone mémoire tampon.

Les charges accumulées durant la durée utile d'une trame sont transférées durant le blanking dans les registres verticaux et immédiatement transférées dans la zone tampon, ensuite durant le temps d'acquisition, elles sont transférées dans le registre horizontal à la fréquence ligne.

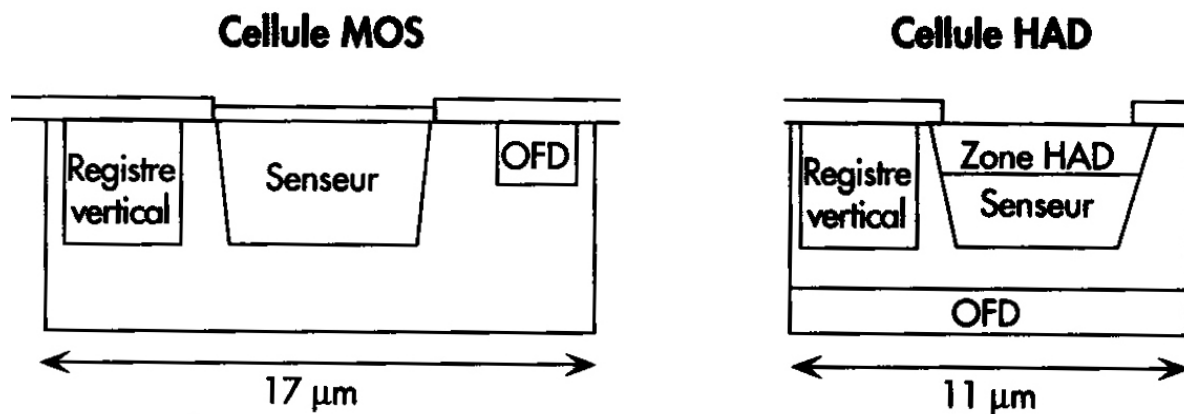
Ce type de structure rend l'obturateur inutile et réduit considérablement le smear qui était surtout causé par la lenteur du transfert vertical.

Les CCD HAD

Les CCD de type MOS que nous avons vu jusqu'ici avaient quelques inconvénients, surtout en prise de vue studio.

Le CCD type HAD (Hole Accumulated Diode) développé par Sony a permis de développer des caméras CCD très haut de gamme.

Chaque cellule d'un capteur HAD renferme, une couche photosensible en dioxyde de silicium, sur laquelle est déposée une couche intermédiaire dopée P (couche HAD), le tout est déposé sur une couche dopée N (le drain).



Comme on peut le remarquer un des grands avantages de cette structure est le renvoi dans la profondeur du capteur du drain, d'où gain de place à la surface pour la partie sensible (on passe de 22 à 32%), la largeur du pixel étant réduite, on peut en loger plus.

Dans une cellule Had, l'électrode ne se trouve plus en surface mais dans la profondeur, l'absence de cette électrode améliore la sensibilité du capteur qui est relativement faible dans les bleus.

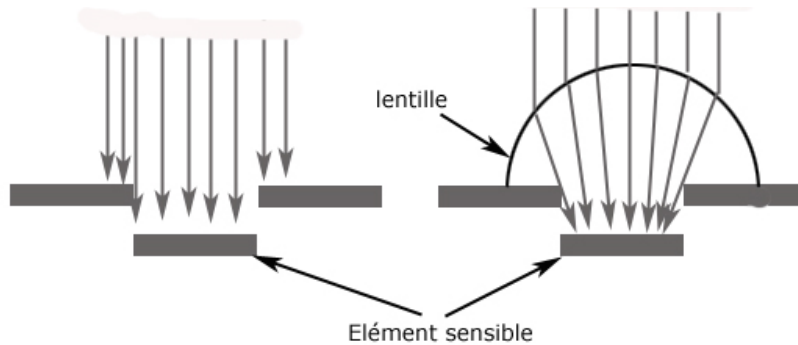
Le smear est en outre fortement réduit, en effet, les électrons provoqués par des rayons proche de l'infrarouge (ceux qui pénètrent en profondeur dans le substrat) sont attirés vers le drain (qui est au fond de la cellule).

La couche HAD sert à compenser la grande sensibilité des photocapteurs aux variations de température, elle absorbe les électrons libres générés par des impuretés à la surface du capteur par la chaleur, et qui se traduisent dans une cellule MOS par un « courant de noir » .

Les CCD Hyper HAD

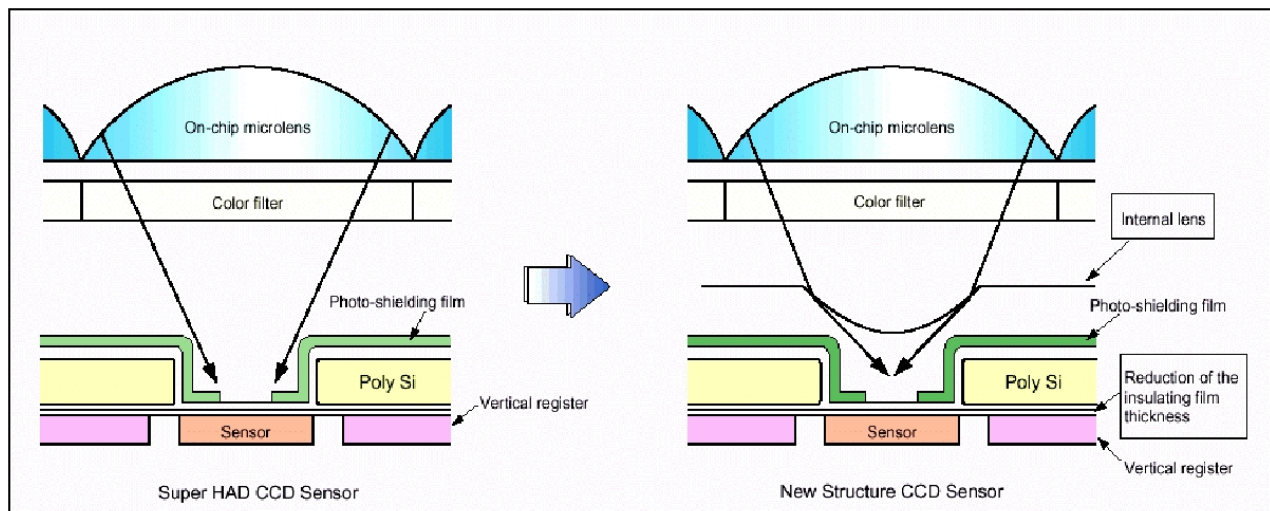
Pour augmenter leur sensibilité, on place sur les CCD un réseau de microlentilles en résine, afin de concentrer les faisceaux lumineux.

Par ce procédé, on double la sensibilité du capteur (on gagne un diaph).



Les CCD Super HAD

L'ajout d'une lentille interne permet de gagner encore en sensibilité en concentrant de façon optimale la lumière arrivant sur le capteur vers le photosite.



Le shutter

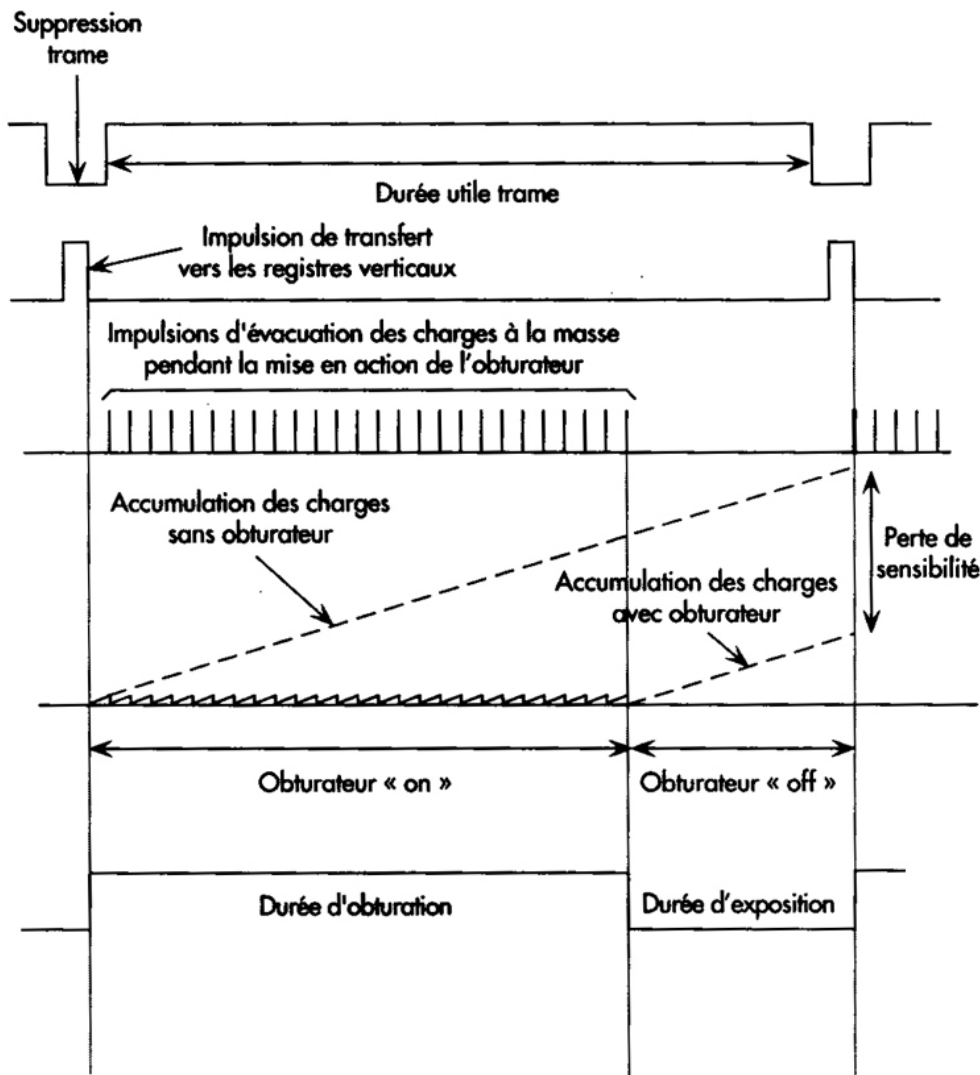
La technologie de type CCD permet de changer la durée d'impression de chaque trame. En fonctionnement normal, le temps de chargement de la zone d'acquisition vaut $1/50$ de sec soit la durée d'une trame.

Avec un shutter, on peut par exemple choisir de ne capter la lumière que durant $1/100$ de sec durant chaque trame, la première moitié de chaque trame n'est donc pas gardée. Le shutter peut être assimilé à un obturateur d'appareil photo mais pour chaque trame d'un signal vidéo.

Le fonctionnement est très simple, durant la partie de la trame non enregistrée, une série d'impulsions ouvre l'OFCD ce qui vide les charges dans le drain.

A partir de la durée voulue, l'OFCD se ferme et les charges sont gardées et ensuite transférées.

Le schéma suivant montre clairement la perte de sensibilité lors d'un temps d'intégration inférieur à $1/50$ de sec.



Le défaut d'aliasing

On a vu que sur un pixel, seul 32% de la surface est sensible à la lumière, c'est un peu comme filmer à travers une grille faite de barreaux très épais.

Le défaut d'aliasing apparaît lorsque l'on capte une scène plus définie que la définition des pixels.

Le résultat sera un signal n'ayant rien à voir avec la scène filmée.

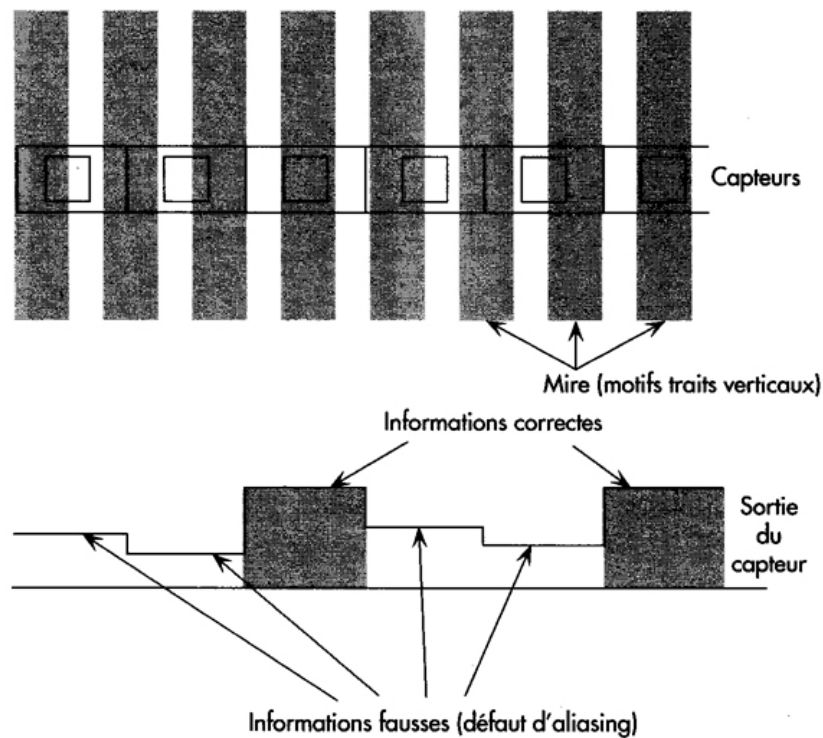
Il faut donc supprimer toutes les fréquences spatiales supérieures à la moitié de la fréquence d'échantillonnage (lois de Shannon).

$$\text{Féch} = \frac{\text{Nbre de points par ligne}}{\text{Durée d'une ligne}}$$

Donc pour un capteur de 786 pts/ligne
Durée active d'une ligne = 52 μ s

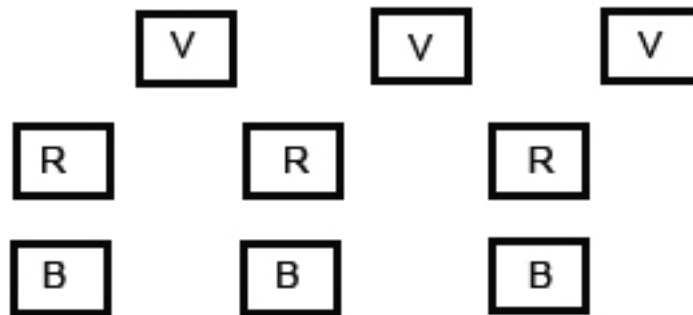
→ Fréq éch = 15 MHz

→ il faut donc supprimer toute fréquence supérieure à 7.5 MHz



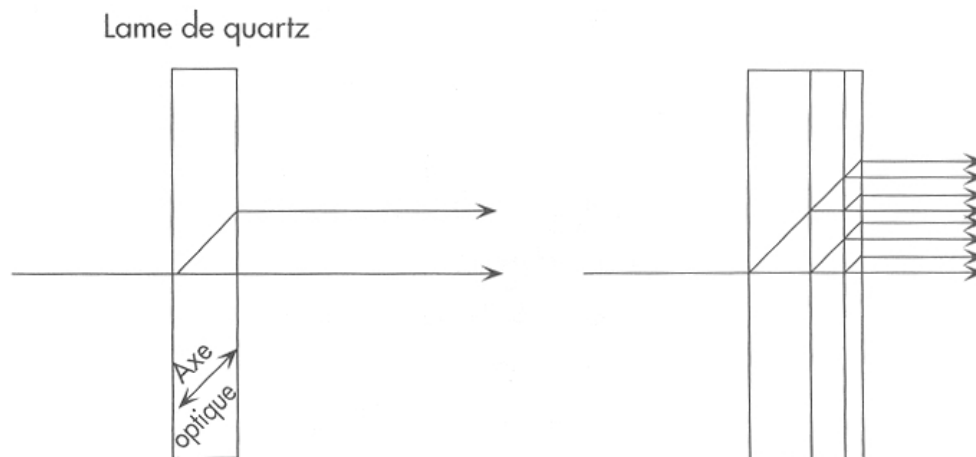
Pour augmenter cette fréquence, on va utiliser le décalage spatial.

→ Le CCD vert est décalé d'un demi pixel horizontalement.
Le canal de luminance étant surtout obtenu par le vert, on double (virtuellement) la fréquence d'échantillonnage à 30 MHz, on doit donc supprimer toute fréquence supérieure à 15 MHz.



Pour éviter cela, on place entre l'objectif et le CCD, un filtre passe-bas fait de lames de quartz. La lame de quartz a la particularité d'avoir un axe optique à 45°, elle laisse passer la lumière en ligne droite et lui ajoute une composante à 45°. (voir schéma page suivante) un point devient donc flou.

Ce filtre empêche les signaux de fréquences trop élevée d'atteindre le CCD.
Après les CCD, un générateur de contour rendra sa netteté à la scène.



Le capteur CCD à couleurs primaires

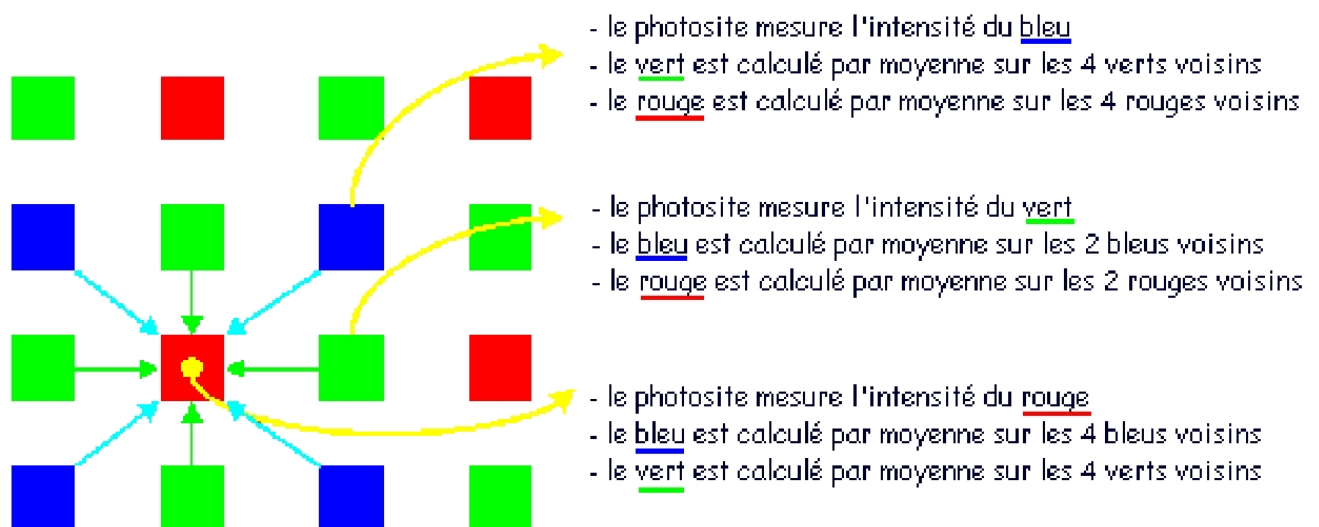
Si, pour des raisons de coût, on ne souhaite utiliser qu'un seul capteur, il faudra surmonter chaque photosite d'un filtre pour l'affecter à une couleur donnée.

L'introduction d'un filtre RVB de Bayer appelé aussi filtre en mosaïque conduira inévitablement à une perte de définition au niveau de chaque couleur :

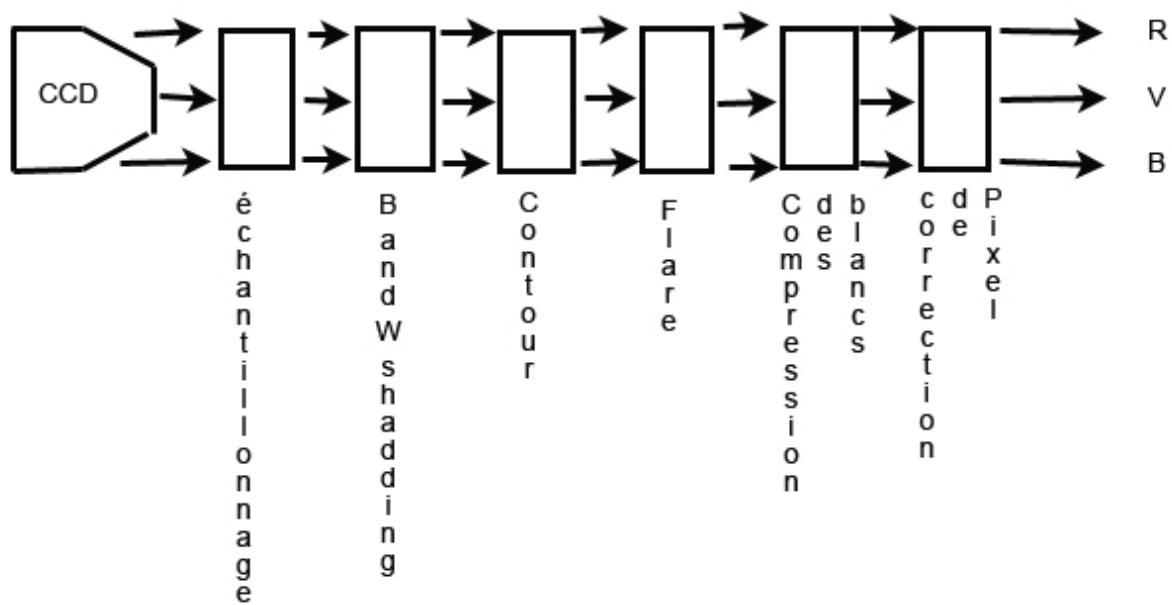
le nombre de photosites sensibles au vert est deux fois plus élevé que ceux sensibles au bleu ou au rouge, ce qui correspond à la sensibilité de l'œil

le dématricage permet d'abord de séparer les informations correspondant aux 3 couches de couleurs

une étape suivante d'interpolation utilisant des algorithmes mathématiques plus ou moins élaborés permet alors d'affecter une valeur RVB à chaque pixel



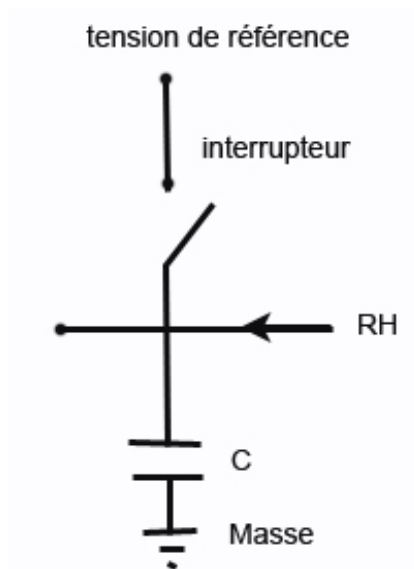
Traitement du signal dans une caméra :



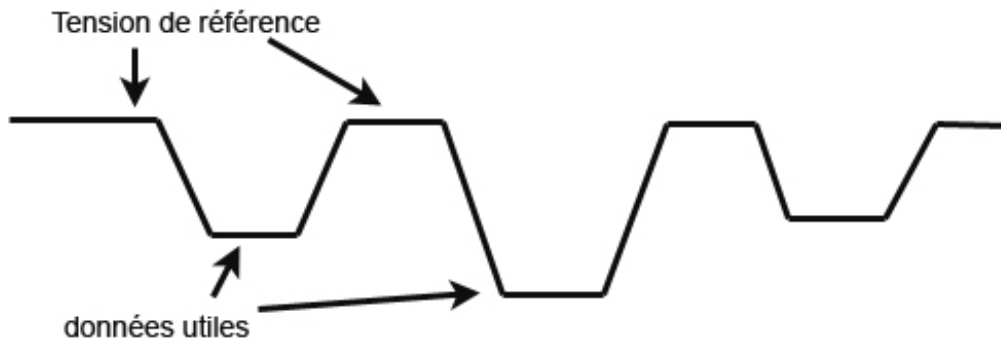
1/ Echantillonnage :

les charges sortant du registre horizontal sont inexploitables telles quelles, il faut les remettre à niveau.

- on charge C sous une tension de référence
- on ouvre l'interrupteur
- le signal issu du registre horizontal décharge C d'un niveau dépendant du niveau de charge
- on referme l'interrupteur



Le tout à la fréquence point.



-On échantillonne ensuite le tension de référence et on échantillonne la donnée utile.
 → la différence entre les deux donne le signal vidéo.

2/ black shading (tache au noir)

le black shading est provoqué par des variations de courant d'obscurité du à la température, cela donne des coloration en horizontal, vertical ou en parabole.

On y remédie en rajoutant une composante de forme inverse sur la couleur désirée.
 Cette correction s'effectue diaph fermé.

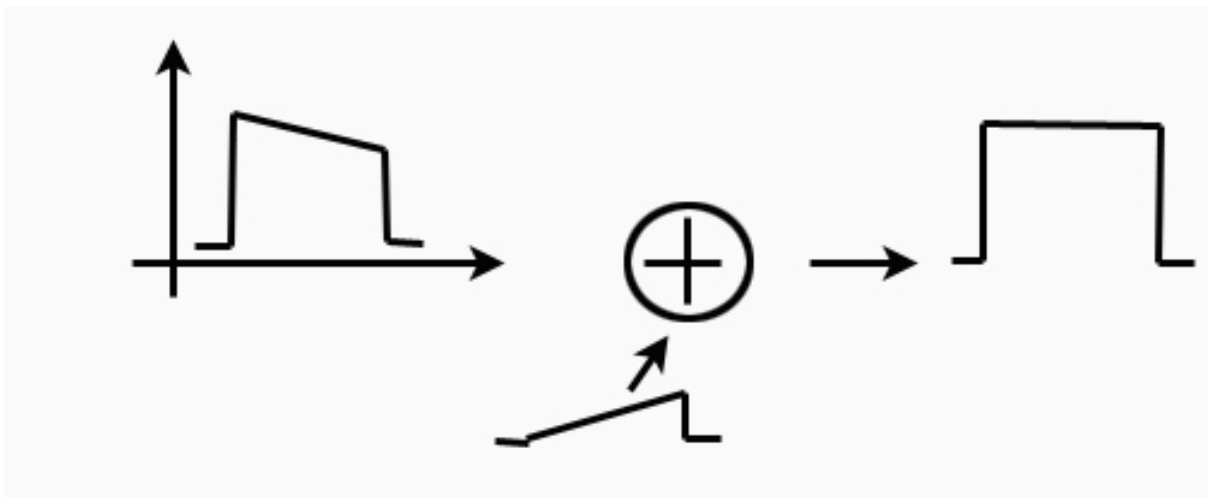
3/ white sadding (tache au blanc)

problème du à l'éclateur optique et à l'objectif.

Plus l'angle d'incidence est grand plus les courbes se décalent vers les faibles longueurs d'ondes.

Si on analyse une surface blanche parfaitement éclairée, on peut avoir une coloration en V en H ou en parabole.

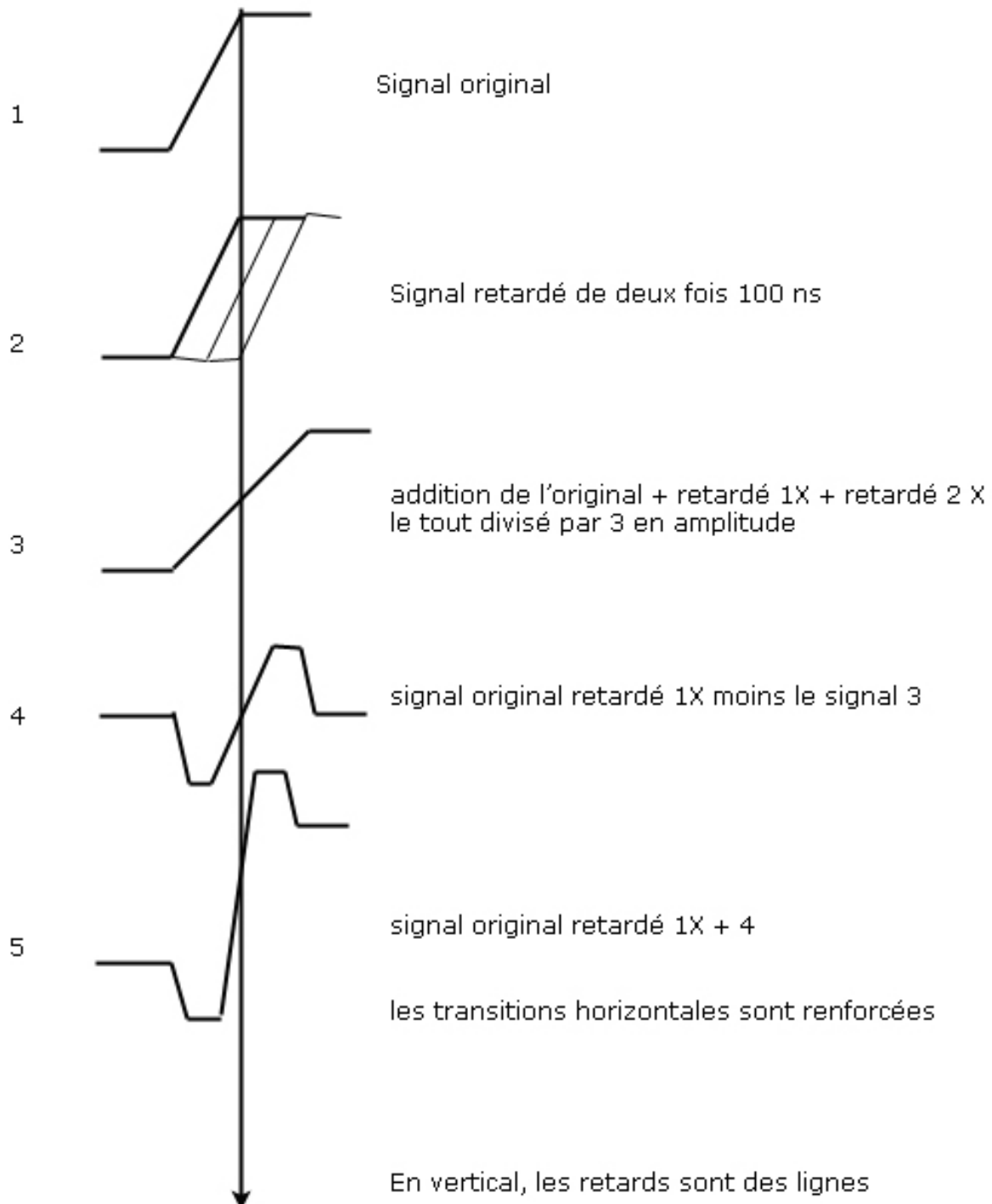
Comme pour le black shading, le correction s'effectue en rajoutant une composante de forme inverse.



4/ correction de contour

Le but de la correction de contour est de rajouter du piqué à une image rendue molle par le filtre passe-bas.

Principe : extraire les hf, les amplifier, les réinjecter.



5/ correction de flare.

Causé par le diffusion parasite de lumière à l'intérieur de l'objectif, plus précisément quand l'ouverture ou la focale varie .

En pratique on utilise une mire blanche avec un carré de velours noir au centre.

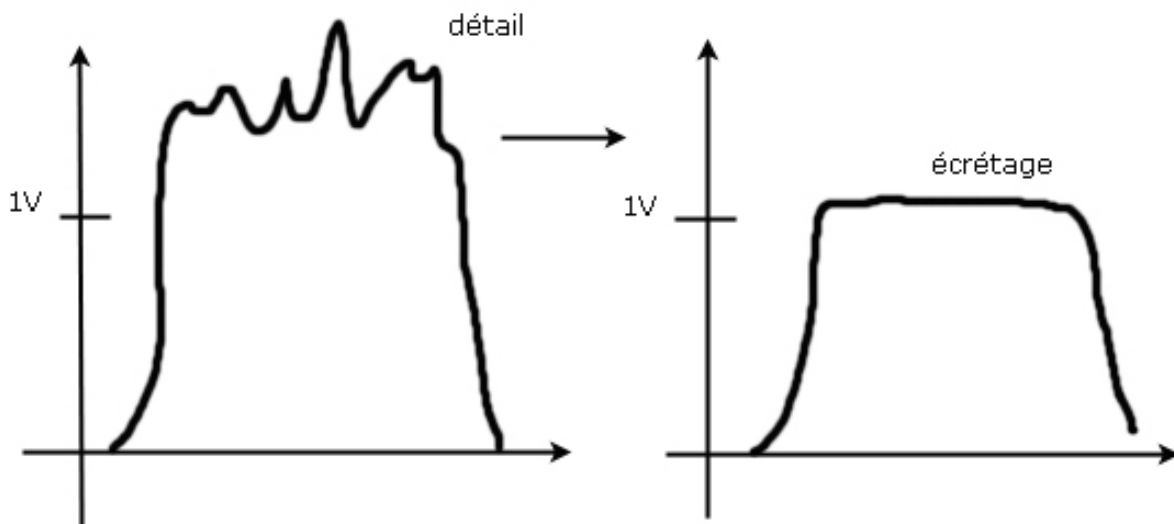
- on zoome pour que le carré noir occupe 90% de l'image
- on règle le niveau de noir
- on dézoome.

Si le niveau de noir remonte, on rajoute une composante continue proportionnelle aux variations, on obtient un niveau de noir constant.

6/ la compression des blancs

Les capteurs CCD sont capables de restituer des niveaux 6 fois supérieurs au niveau nominal d'un signal vidéo (6V), ce qui est formidable du point de vue dynamique, ce qu'il est moins, c'est que le reste de la chaîne (mix, vtr) rabote le signal à un volt, d'où, perte d'informations.

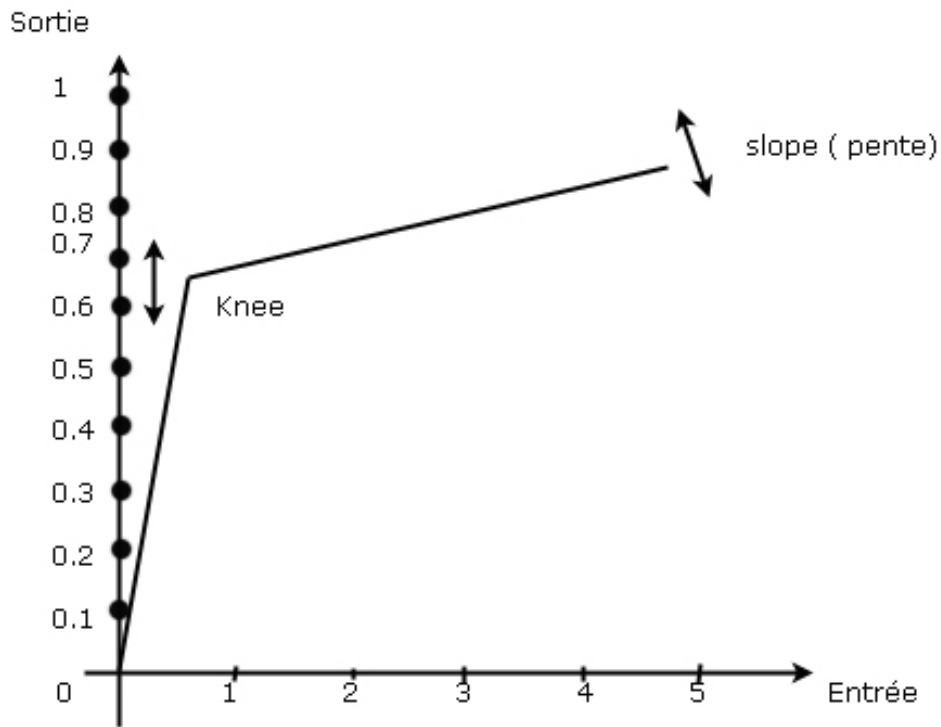
Exemple :



On va donc atténuer les amplitudes des zones sur illuminées, pour cela, nous allons employer deux paramètres :

- le knee ou seuil : niveau à partir duquel la compression entre en action.
- le slope : qui est la valeur de la pente

pour bien faire, cette compression doit agir sur les trois canaux (question de colorimétrie)



7/ correction des pixels défectueux :

lorsqu'un CCD prend de l'âge, certains pixels peuvent devenir plus sombres ou plus actifs, en tout cas, visibles !

dans ce cas deux types de correction sont possibles :

a/ - Acquisition des pixels défectueux

- mémorisation de leur position dans la tête de caméra
- les remplacer par le pixel de la ligne précédente

(cette opération se fait en maintenance)

b/ Correction dynamique

analyse permanente de l'état des pixels, comparaison de leur valeur avec celles de ses 8 voisins.

Si des valeurs inhabituelles sont détectées, on compare avec les deux autres CCD aux mêmes coordonnées.

Si un seul CCD varie, on applique la correction.

Si la variation apparaît dans les trois CCD, il s'agit d'un détail coloré.

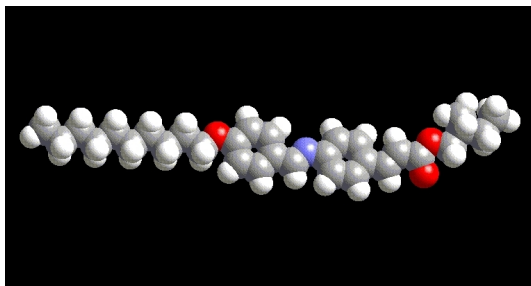
LES ECRANS

Les écrans LCD

Les cristaux liquides

Les cristaux liquides ont été découverts en 1888 par le botaniste autrichien H. Reinitzer. Comme leur nom ne l'indique pas, les cristaux liquides ne sont pas des cristaux :

la plupart des cristaux liquides sont des composés organiques avec des molécules allongées en forme de bâtonnets
ils ne présentent ni les arêtes vives, ni les facettes lisses et brillantes qui caractérisent ce qu'on appelle communément les cristaux.
en revanche, ils possèdent des propriétés d'organisation des molécules qui se traduisent par des caractéristiques optiques particulières et qui les rapprochent de l'état cristallin



Molécule de cristal liquide et exemple d'alignement

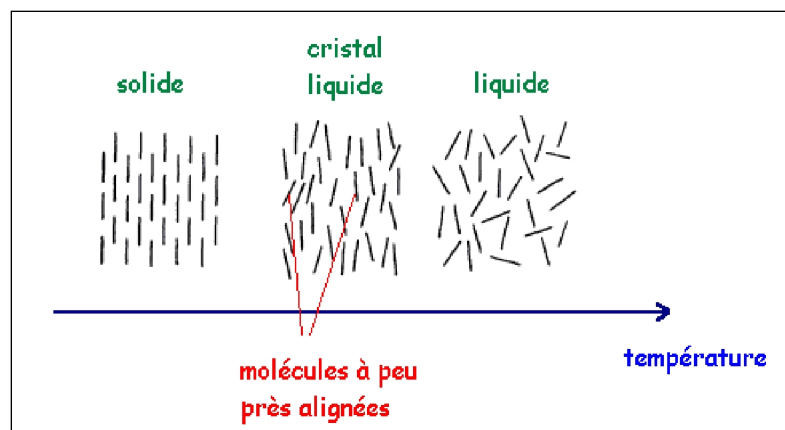
Dans l'état naturel, les molécules se disposent spontanément d'une manière ordonnée avec leurs axes grossièrement parallèles.

Prenons l'exemple du Cholesteryl Myristate, composé de carbone et d'hydrogène, et utilisé par Reinitzer lors de ses expériences en 1888:

à 20 degrés on est en phase solide, les molécules du cristal sont rangées avec un ordre de position et d'orientation

à 71 degrés, le solide fond, mais le "liquide" résultant est trouble, l'ordre positionnel a disparu mais toutes les molécules ont plus ou moins gardé leur orientation d'origine.

à 85 degrés, le liquide trouble devient clair, on a obtenu une phase liquide classique où ne subsiste aucun ordre



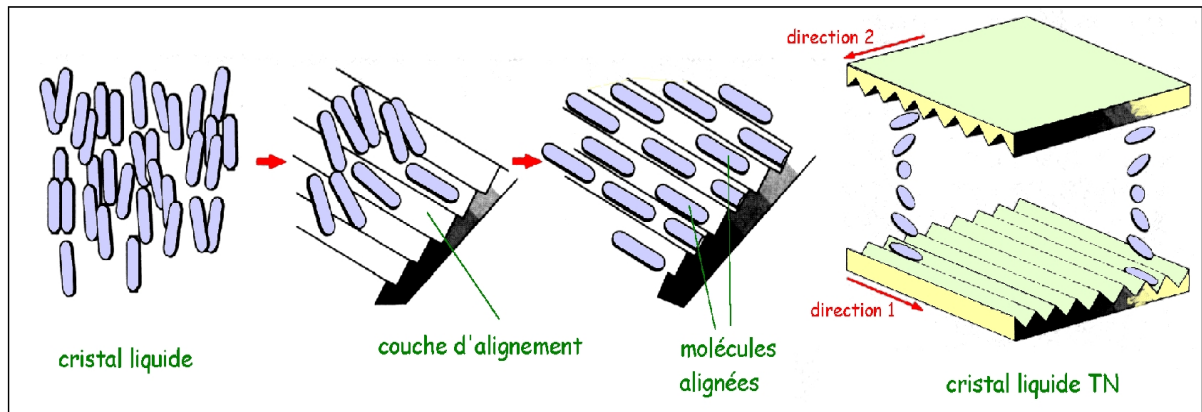
Pour les afficheurs à cristaux liquides, on utilise des substances qui sont à l'état de cristaux liquides à la température ordinaire.

L'orientation des cristaux liquides

Il est possible de "piloter" la direction des molécules de cristaux liquides par divers moyens :

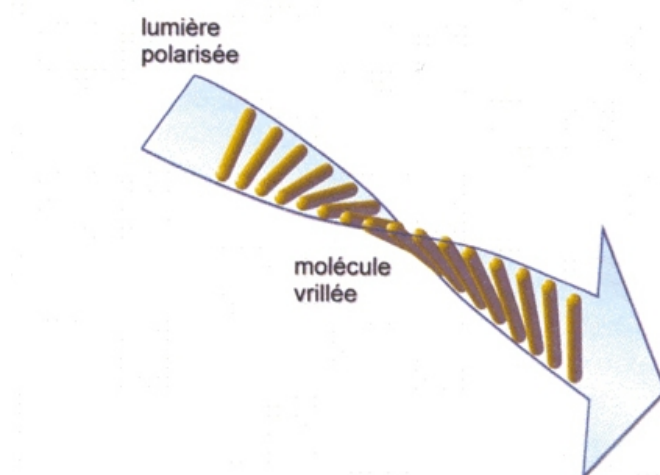
par un **dispositif mécanique**

Lorsqu'on dispose une couche de cristal liquide sur une plaque gravée de fins sillons parallèles (couche d'alignement), les molécules s'orientent parallèlement à ces sillons



Lorsqu'on enferme une couche de cristaux liquides entre deux plaques gravées de sillons orientés dans deux directions différentes, l'orientation des molécules (à l'état de repos) passe progressivement de la direction (1) à la direction (2). Elle fait donc apparaître une torsion.

C'est ce qu'on appelle un **cristal liquide "Twisted Nematic"**, qu'on pourrait traduire par nématique tordu. Ces molécules alignées vont avoir des propriétés spéciales dans le domaine optique puisqu'elles vont faire tourner la direction de polarisation de la lumière.

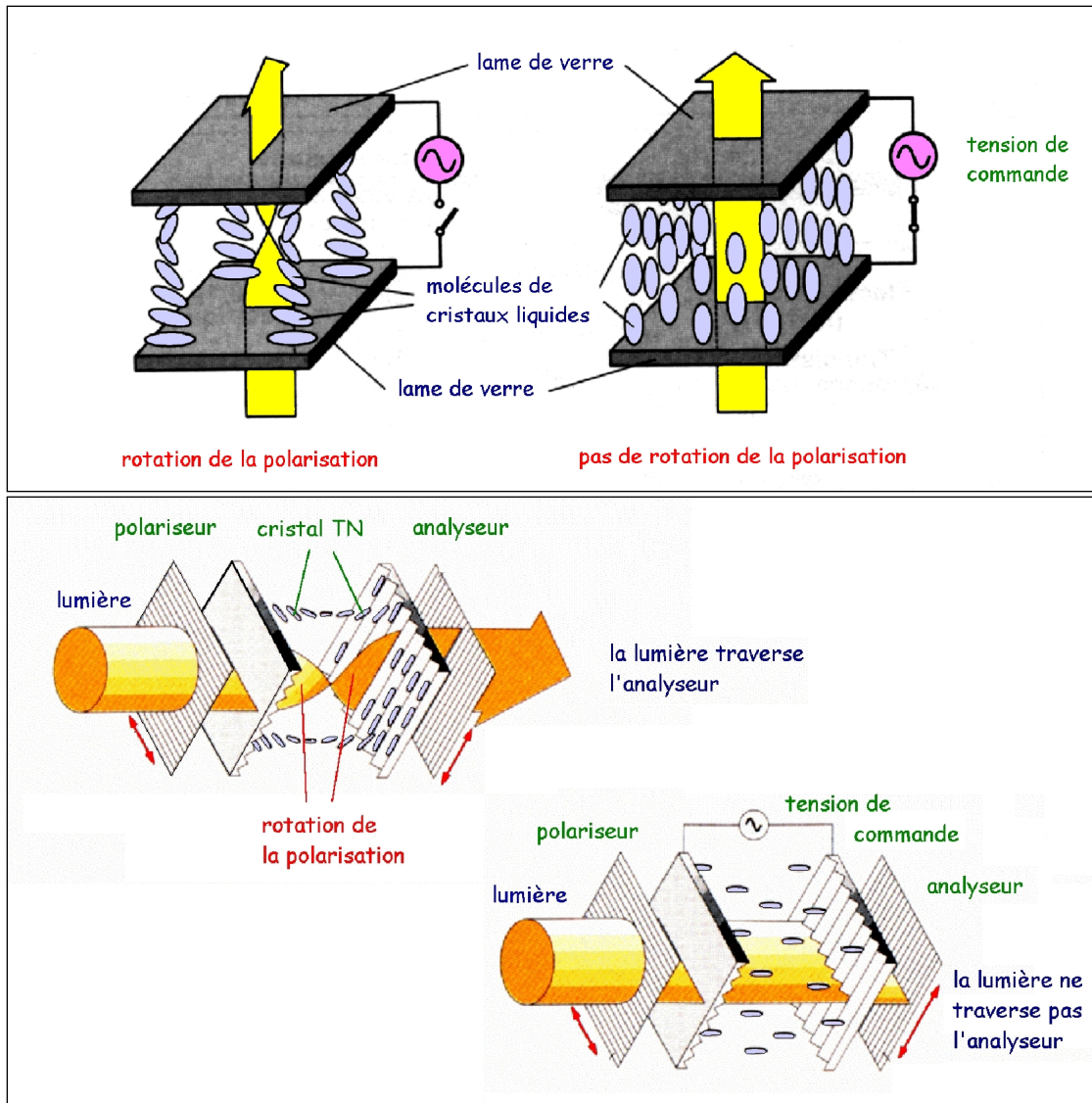


Rotation de la lumière par la molécule nématique polarisée

sous l'action d'un **champ électrique**

Si on soumet le cristal liquide à une tension électrique, l'orientation des molécules se modifie sous l'action du champ électrique.

Les chaînes moléculaires s'orientent parallèlement aux lignes de champ et se retrouvent perpendiculaires aux électrodes. Dans cette situation, elles ne modifient plus la direction de polarisation de la lumière.



Remarque : en jouant sur la valeur de la tension appliquée, il est possible d'obtenir des situations intermédiaires dans lesquelles seule une partie des rayons incidents voit sa direction de polarisation modifiée.

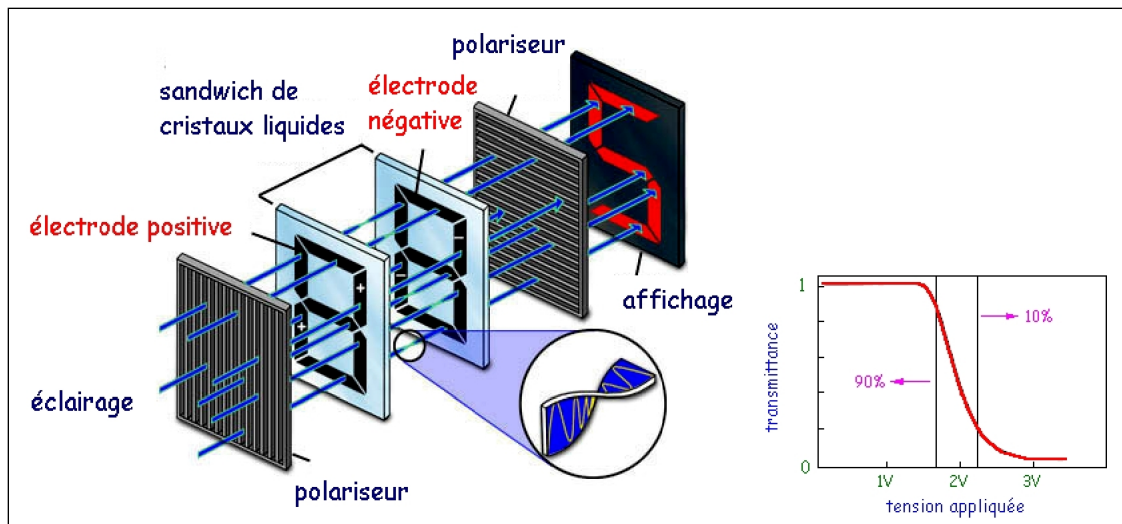
En intercalant une cellule de ce type entre deux polariseurs croisés, on peut fabriquer un dispositif à transmittance optique variable :

au repos, les cristaux liquides font tourner l'axe de polarisation de la lumière qui peut donc traverser le dispositif l'application d'une tension entre les deux plaques de verre aligne les molécules sur le champ et empêche la rotation de la polarisation : la lumière ne traverse plus le dispositif

La variation de la transmittance en fonction de la tension de polarisation appliquée met en évidence une zone où la transmittance varie rapidement.

On pourra donc utiliser l'interrupteur optique ainsi constitué pour afficher sur un écran des zones sombres ou claires :

- qui reproduisent la forme des électrodes : c'est le principe utilisé dans les afficheurs à points ou à segments.
- qui peuvent être adressées de façon matricielle en x,y pour former une image quelconque : c'est la technique utilisée dans les écrans LCD graphiques .



Sur un écran à structure matricielle, les pixels sont repérés par les lignes et les colonnes :
la cellule se comporte comme un condensateur que la commande doit charger ou décharger
cette charge doit être rapide, la commande doit donc être capable de fournir un certain courant
la cellule conserve sa charge et possède donc une mémoire intrinsèque
grâce à cet effet mémoire, l'écran LCD ne présente pas le défaut de papillotement des tubes

l'affichage des niveaux de gris :

L'affichage des différents niveaux de luminance de l'image est basé sur le principe suivant :

en l'absence de champ électrique appliqué (tension de polarisation nulle), les molécules se placent en position « twistée » pour permettre le passage de la lumière.

en présence d'un champ électrique faible (tension de polarisation faible), quelques molécules s'orientent en direction du champ, les autres restant « twistées » : la transmittance a légèrement diminué

en présence d'un champ fort (tension de polarisation élevée), la plupart des molécules sont orientées dans le sens du champ et la transmittance de la cellule est très faible

On pourra donc commander la luminosité du pixel en jouant sur la tension de polarisation qui lui est appliquée : les afficheurs à cristaux liquides nécessitent donc une [commande analogique](#).

Les couleurs sont obtenues en utilisant des filtres de couleurs vert rouge et bleu, les trois couleurs primaires, comme sur un tube cathodique.

Les sources d'éclairage

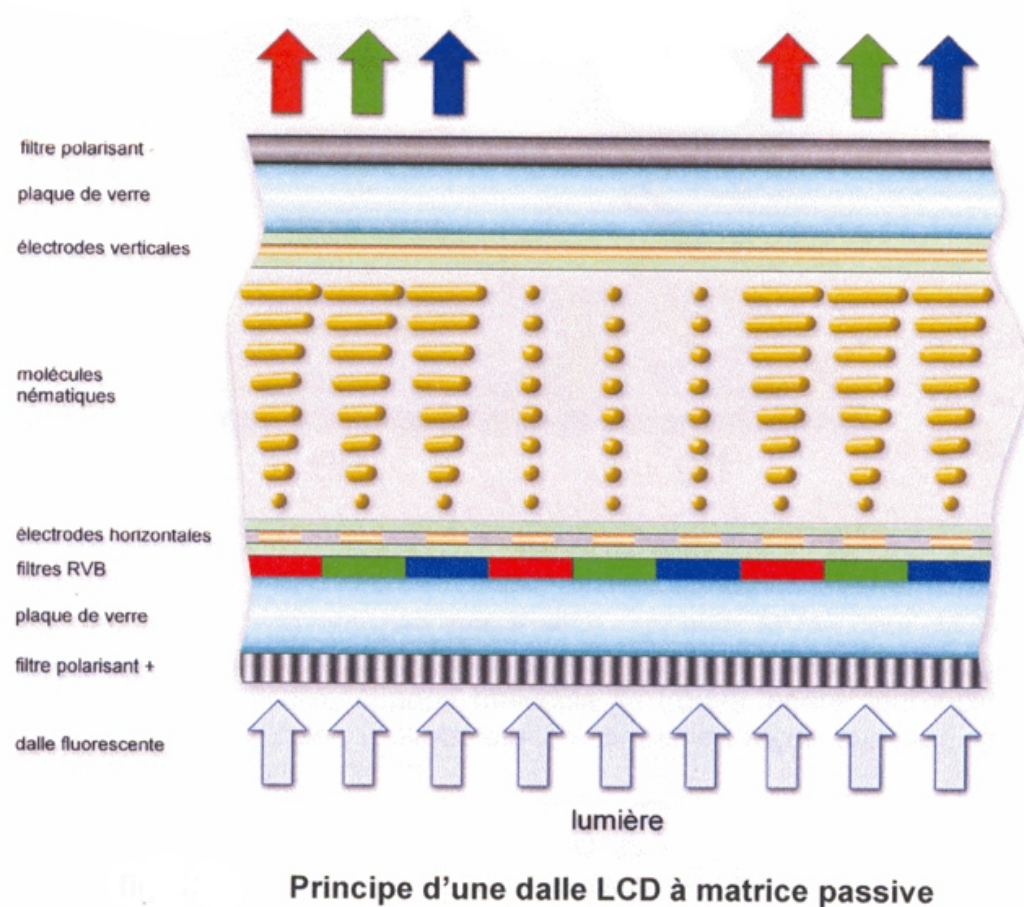
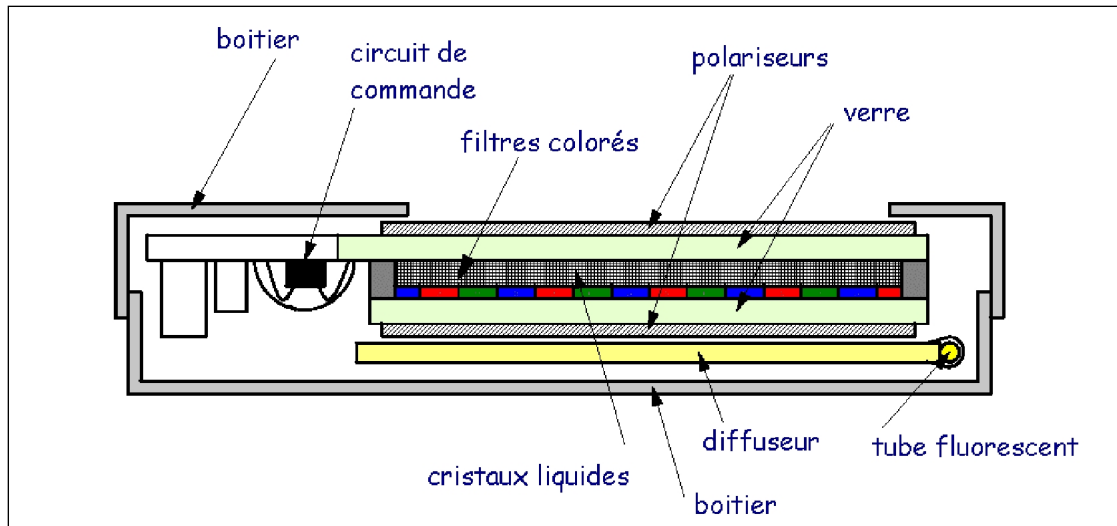
Les deux configurations des écrans LCD sont les écrans réfléchifs et ceux utilisant un rétro-éclairage.

les [écrans réfléchifs](#) utilisent la lumière ambiante.

Au repos, la lumière traverse le cristal avant de subir une réflexion dans le miroir et de repasser à travers le cristal : on observe un point blanc. Quand on applique une tension, on observe un point noir (la lumière est bloquée avant la réflexion).

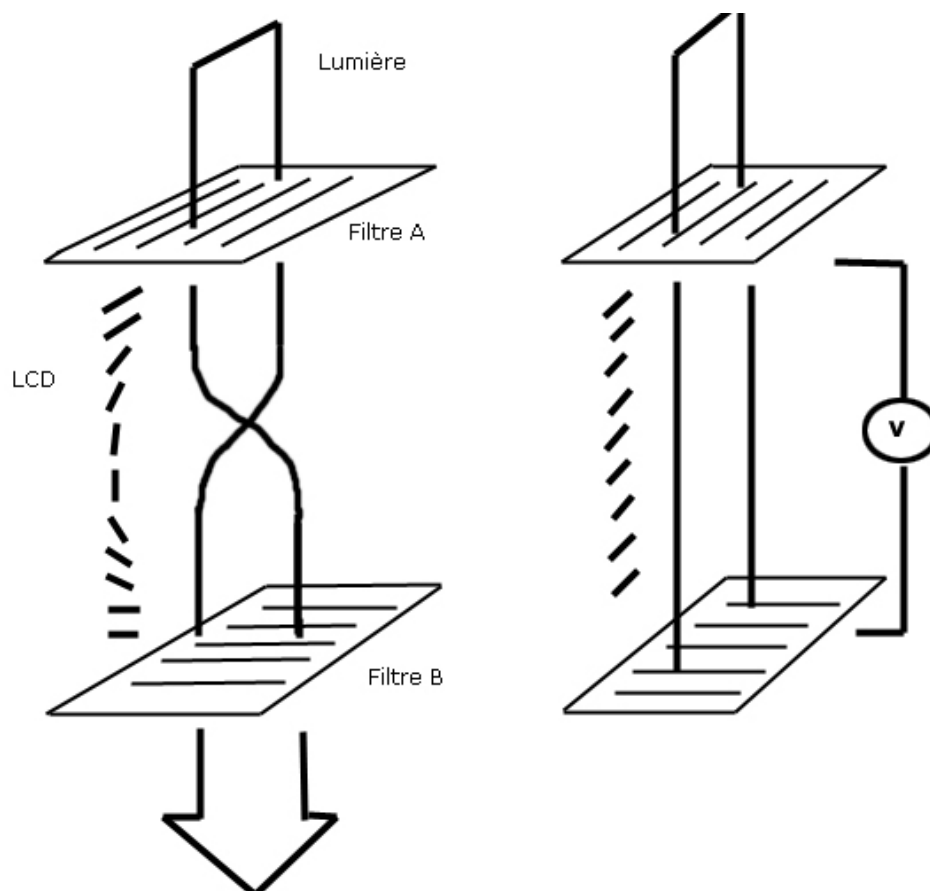
les écrans avec rétro-éclairage (LED, tube fluorescent).

Cette lumière traverse ou non le cristal en fonction de la polarisation appliquée. Ce type d'écran LCD donne une meilleure intensité de l'écran, surtout avec peu de lumière ambiante, mais consomme plus d'énergie qu'un écran réflectif.



Résumé sur les LCD :

- le principe est de mettre en sandwich des cristaux liquides entre deux plaques polarisatrices orientées à 90°
- Les molécules LCD au repos passent progressivement d'une orientation à l'autre
- L'écran est rétro éclairé par des néons polarisés parallèlement à la première plaque.
- la lumière est guidée par les molécules LCD et après une orientation de 90° elle passe le deuxième filtre polarisant
- si on applique une tension de commande, les LCD s'orientent parallèlement à la première plaque et la lumière est bloquée
- chaque pixel de l'image est constitué d'une cellule de ce type avec un filtre R, G ou B



Avantages : faible consommation
Absence de rayonnement
Géométrie parfaite
Exempt de scintillement
Peu encombrant

Inconvénients : l'observation optimale se fait dans l'axe
(à cause de la polarisation)

contraste plus faible (les lcd ne bloquent pas tout, un noir devient sv
gris)

Le tube cathodique

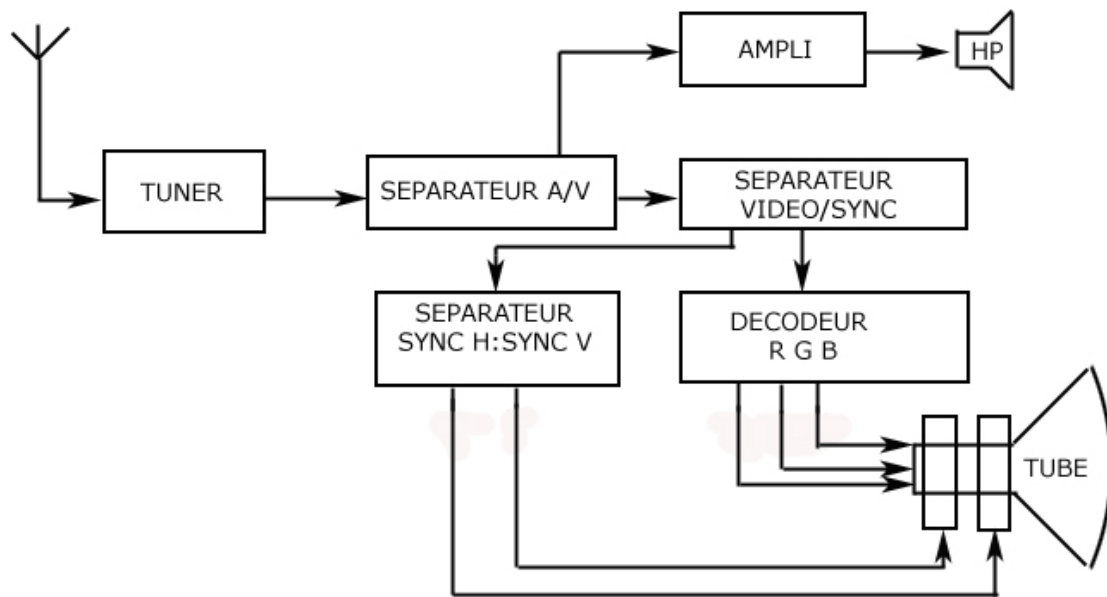
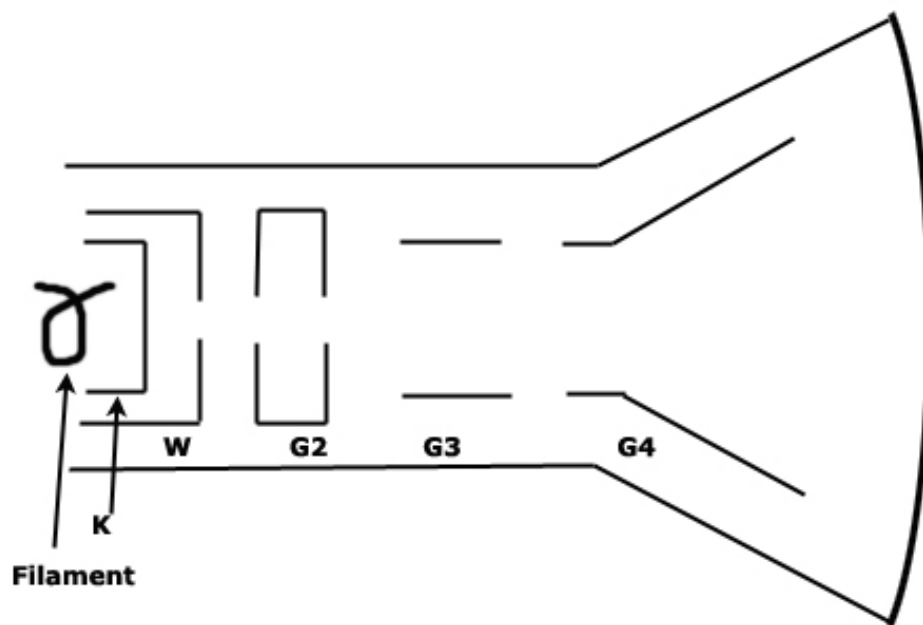


Schéma block d'un téléviseur



Le faisceau d'électrons est obtenu en chauffant, avec le filament, la cathode (K) recouverte d'un oxyde émissif.

Il y a ensuite concentration électrostatique du faisceau par des « grilles » auxquelles on applique des tensions.

Le débit est contrôlé par G1 (W ou Wehnelt)
G2 et G3 assurent la focalisation électronique et G4 l'accélération finale.

Le balayage est assuré par des bobines placées autour du col du tube (déviation électromagnétique).
Une bobine pour la déviation H et une pour la déviation V.

Les tubes trichromes

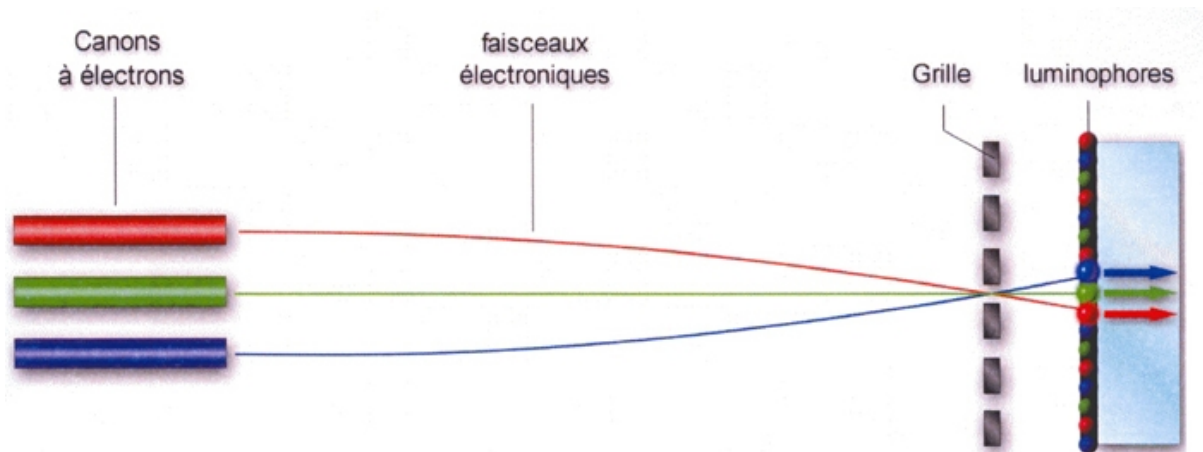
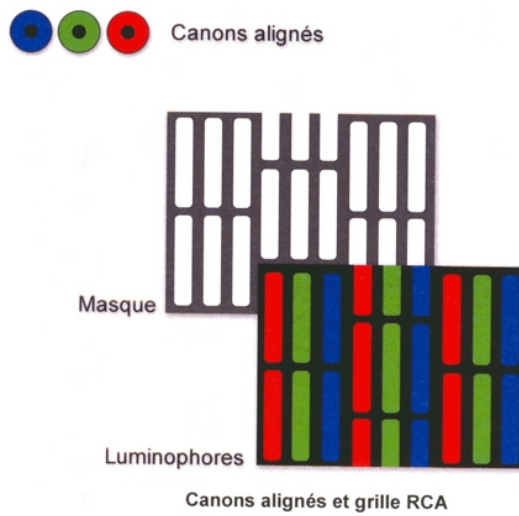
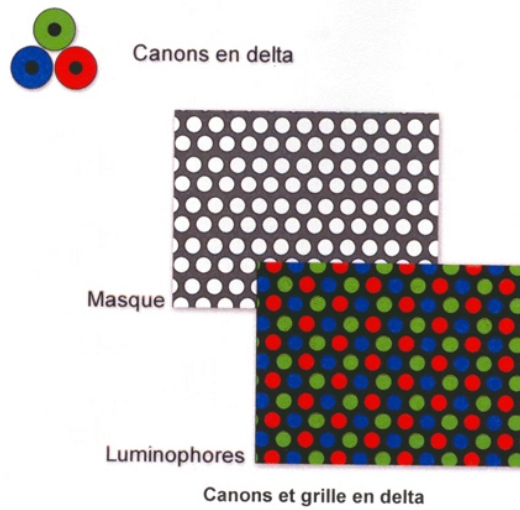


fig. 123. Principe des canons à électrons en mode alignés

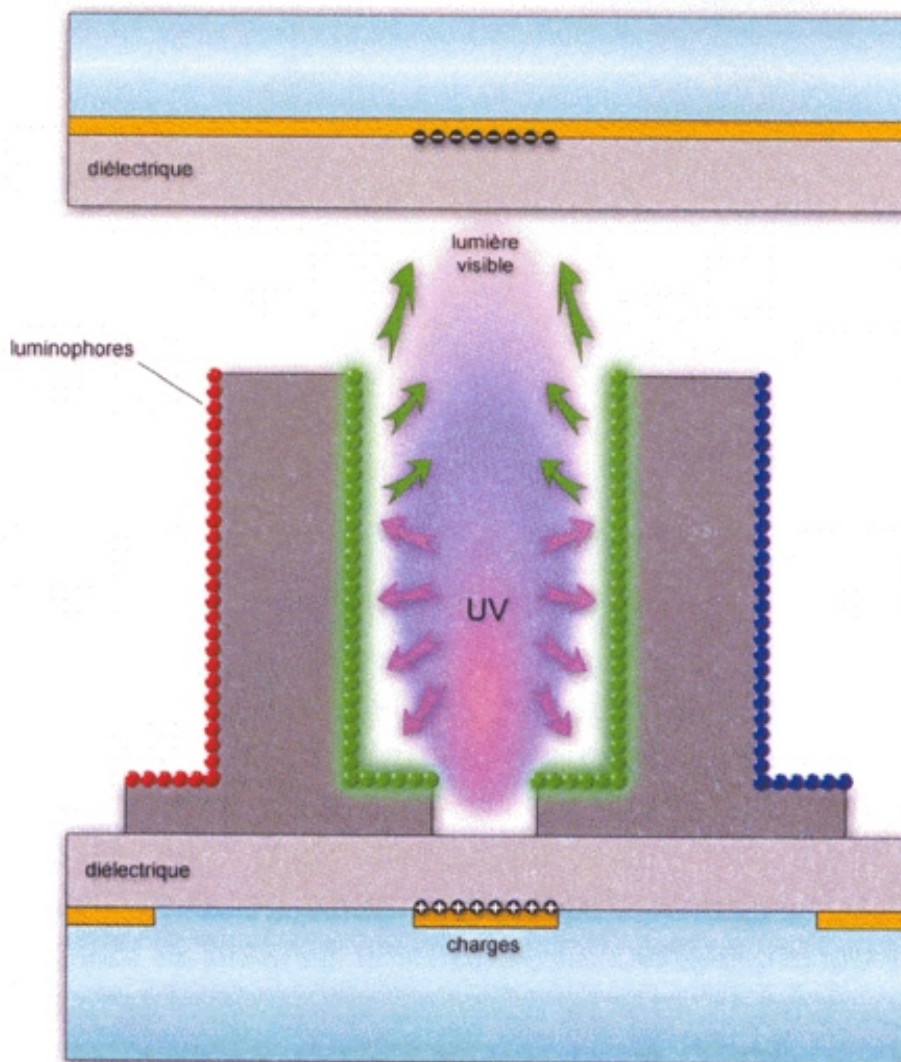


Les écrans à plasma

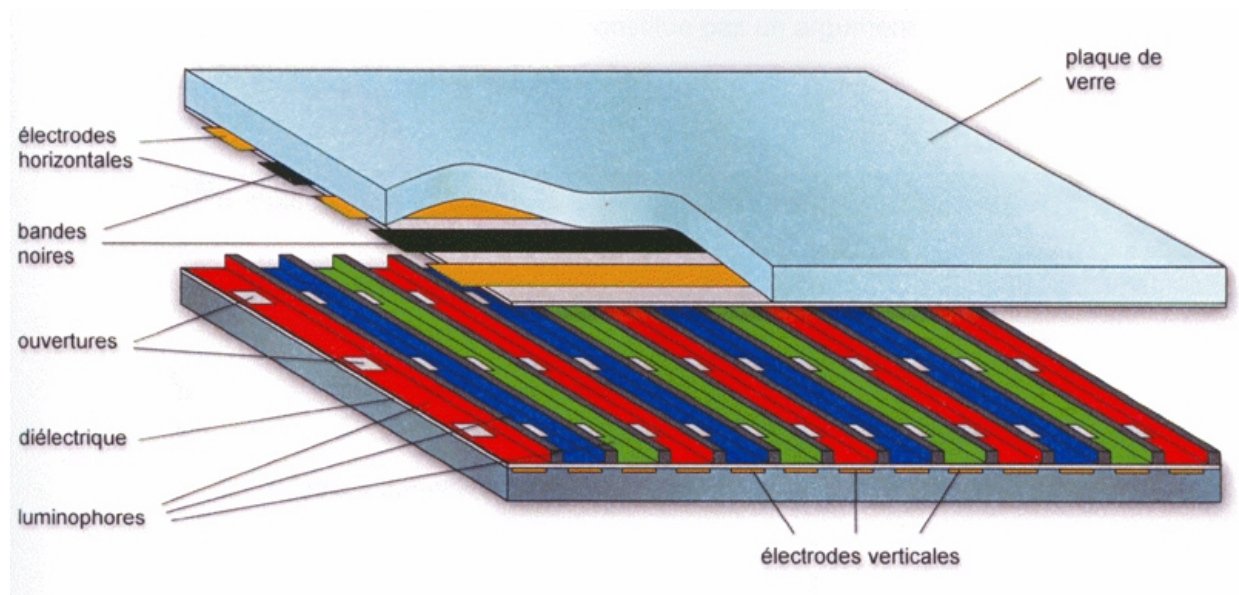
PDP :plasma display panel

- Le « plasma » constitue l'état particulier de la matière caractérisé par un désordre maximal, état dans lequel les atomes sont en grande partie ionisés (noyau et électrons sont séparés)
- Le principe de ces écrans consiste à amorcer et entretenir une décharge électrique dans une cellule contenant un gaz rare (xénon, néon, argon ...) sous basse pression.
- Comme dans un tube au néon, la décharge émet un rayonnement ultraviolet (invisible)
- Ce rayonnement ultraviolet est utilisé pour exciter des luminophores qui vont eux-mêmes émettre une lumière visible.

Donc un écran plasma est constitué d'une multitude de minuscules tubes fluorescents, 1 par pixel, pas besoin de rétro éclairage.



Principe d'émission de la lumière dans un plasma



Structure d'un écran plasma

Avantages : contraste et luminance élevés
Angle de vision indépendant de l'axe

Inconvénient : très gourmand en énergie