

Electronique Vidéo

1) La diode

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.

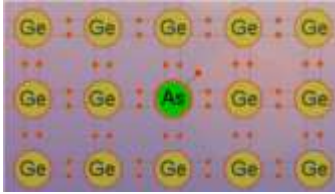
Il existe deux types de dopage, le type N et le type P.

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.

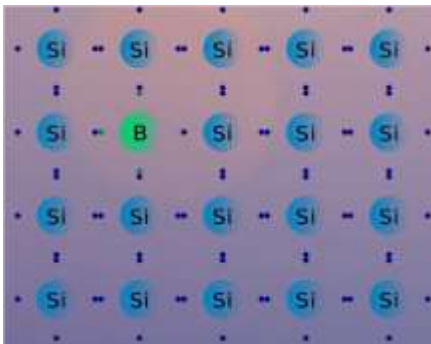
En quittant son atome, l'électron excédentaire y laisse une charge positive liée au noyau :

Un ion positif mais pas de place libre pas de trou

Le type P

Le dopeur appartient à la colonne IIIB, 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.

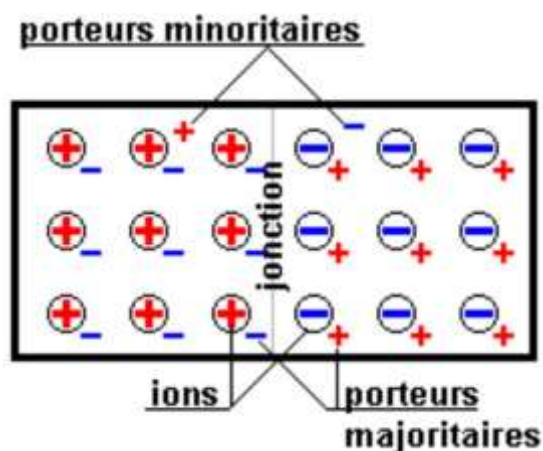


Notons que comme dans le type N, le trou est inoccupé, mais sans charge, le cristal reste neutre.

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.

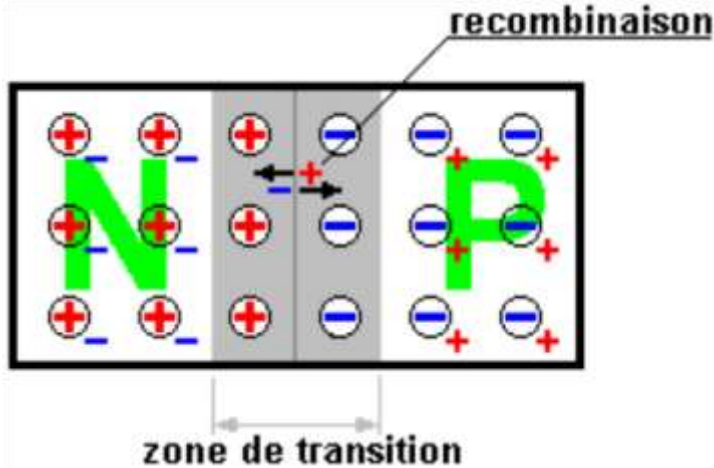
Fonctionnement d'une jonction P-N

- 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.

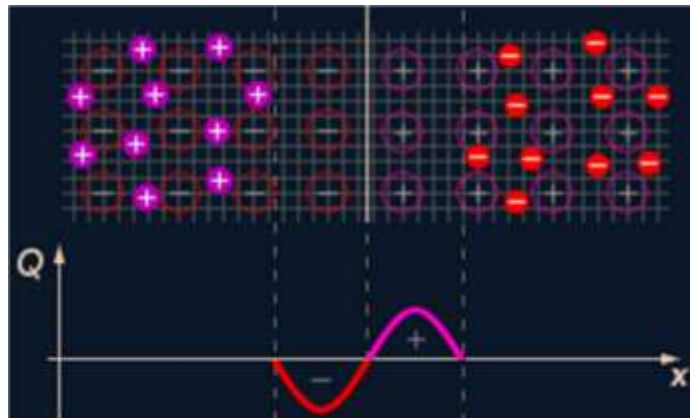


Cette double charge d'espace fait apparaître un champ électrique à cause interne.

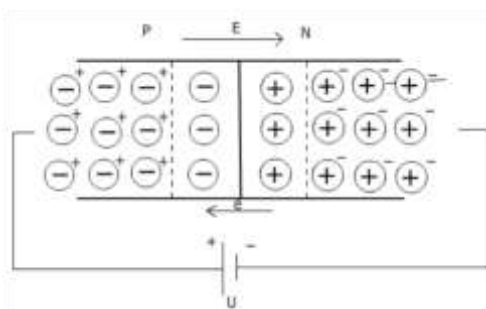
Ce champ électrique s'oppose à la diffusion des lacunes vers N et des électrons vers P.

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



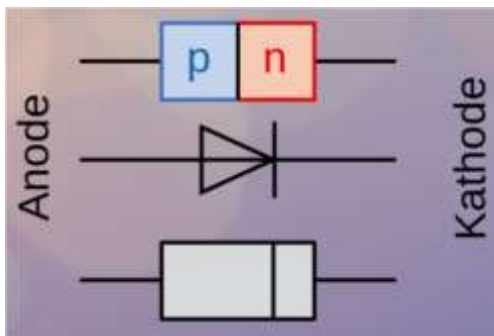
sens passant :



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. la jonction est polarisée dans le sens passant, le courant traverse la jonction.

Sens non passant :

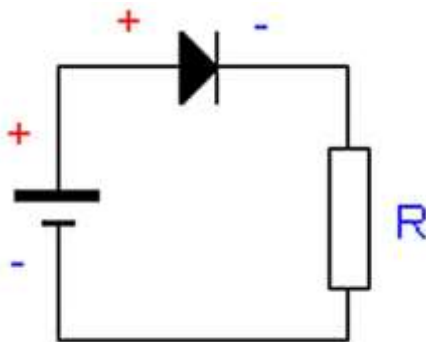


Si on inverse l'alimentation, E vient renforcer le champ interne et aucun courant ne passe la jonction est polarisée dans le sens bloqué. Ce genre de jonction est un composant électronique que l'on appelle une diode.

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.

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 avons 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.

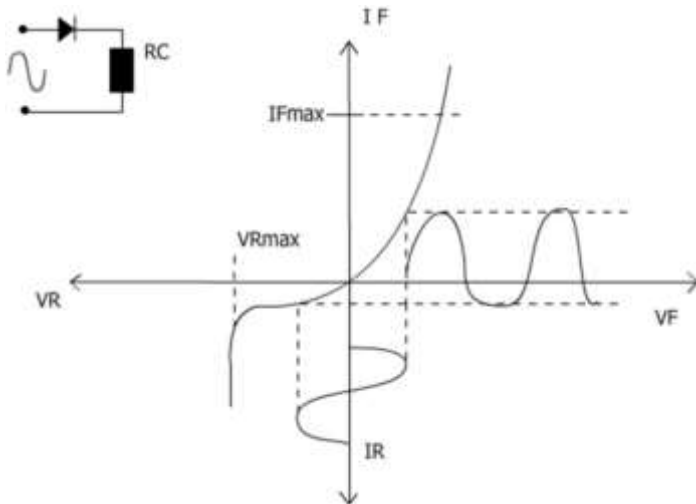
Nous allons faire varier la tension du générateur de 0 à $+V_{cc}$ (nouveau terme indiquant la tension maximum d'alimentation continue) en relevant à chaque fois le courant qui circule dans le circuit et la tension aux bornes de la diode.

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.

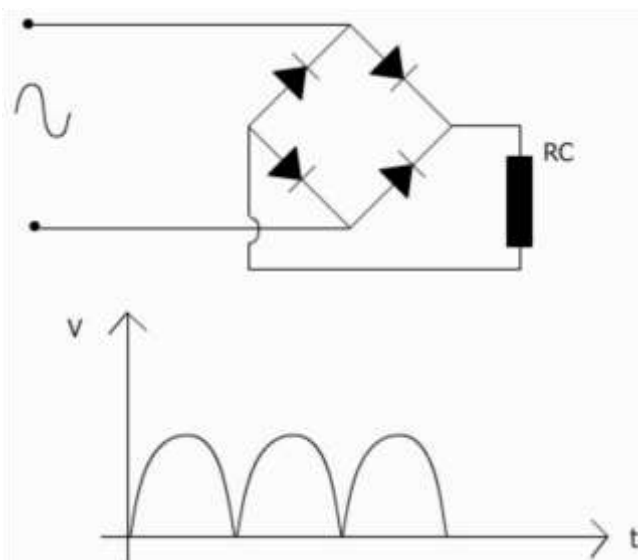
Quand la tension aux bornes de la diode est inférieure à 0,7 V, aucun courant ne circule dans le circuit, c'est comme si nous avions un interrupteur ouvert. A 0,7 V, brutalement le courant apparaît. Si nous augmentons la valeur de la tension fournie par le générateur, la tension aux bornes de la diode 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.

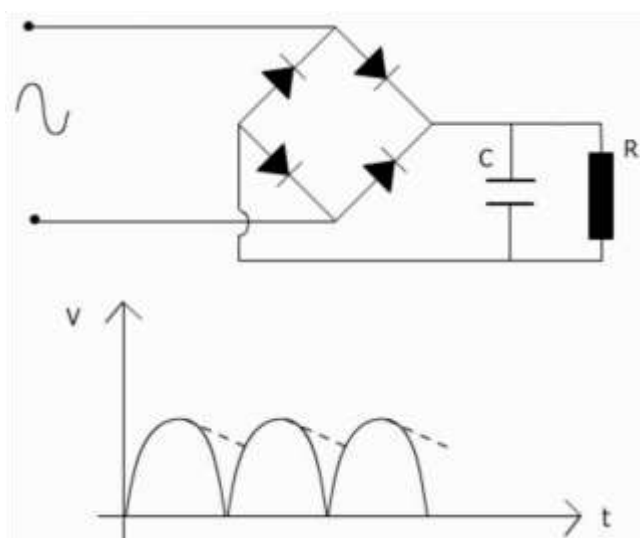
La diode en redressement mono alternance.



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 cidessus (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.

2) Le transistor bipolaire

potentiel,

Symbole :

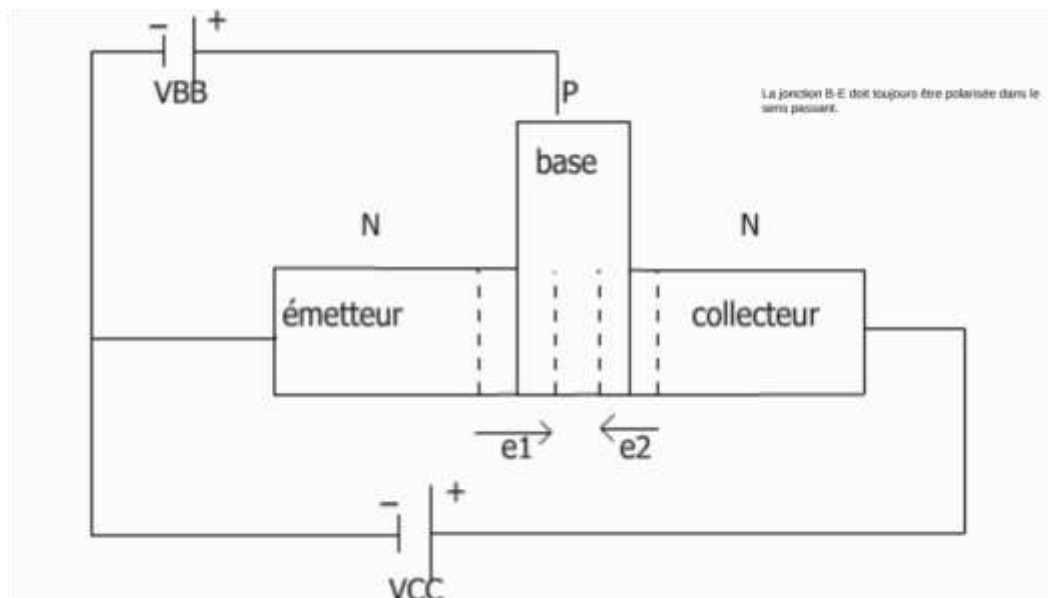
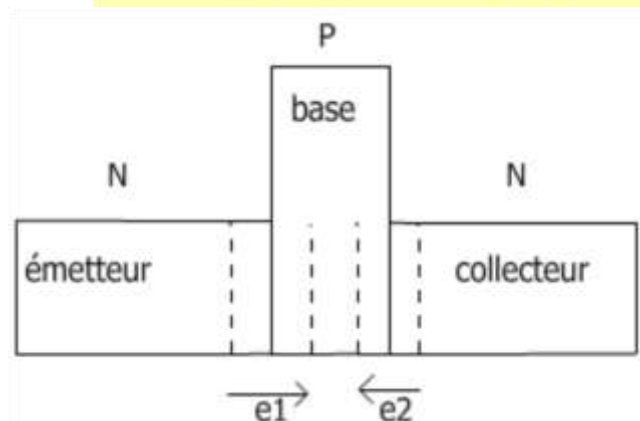
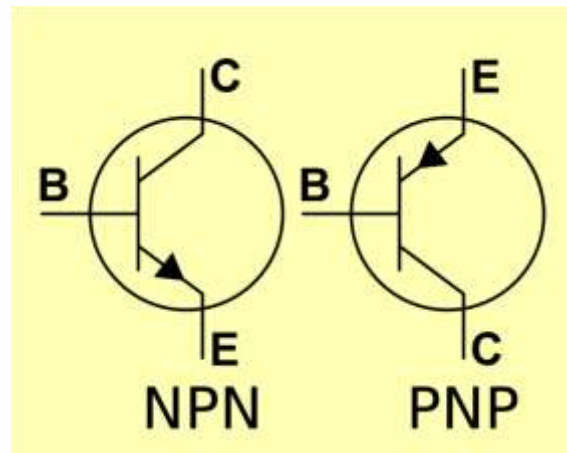
Un transistor est formé de 2 jonctions de polarités opposées placées en série.
Pour que l'effet transistor se manifeste, il faut que l'ensemble soit constitué par le même monocristal et que la partie centrale (la base) soit très mince.

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.

Comme dans le cas de la diode, il s'établit à chaque jonction une barrière de

Les deux champs internes e_1 et e_2 sont de sens opposés.

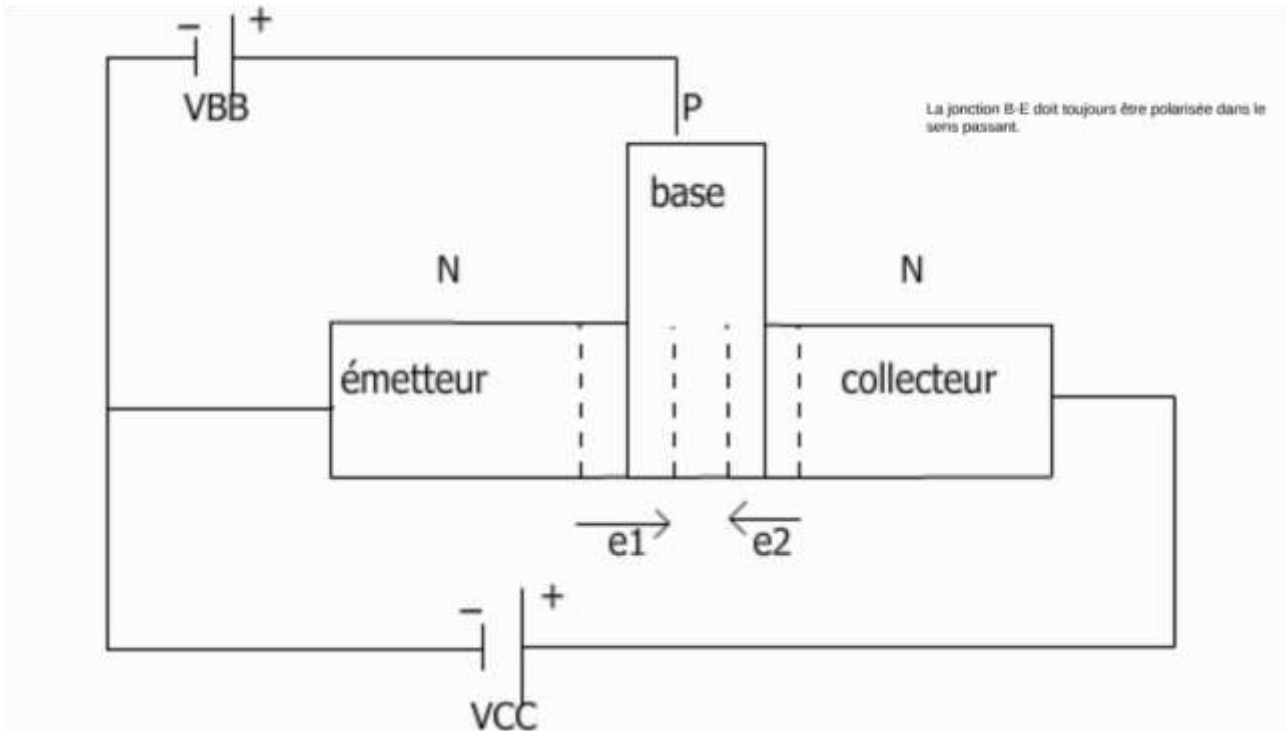


Si l'on branche une source extérieure entre le collecteur et l'émetteur, il ne se passera rien, en effet il y aura toujours une jonction polarisée dans le sens bloquant.

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

V_{CC} : pôle positif au collecteur (tension beaucoup plus élevée que V_{BB}) Pôle négatif à l'émetteur

VBB : pôle positif à la base (tension beaucoup plus basse que VCC) Pôle négatif à l'émetteur



L'émetteur est une zone fortement dopée, son rôle est d'injecter des électrons dans la base.

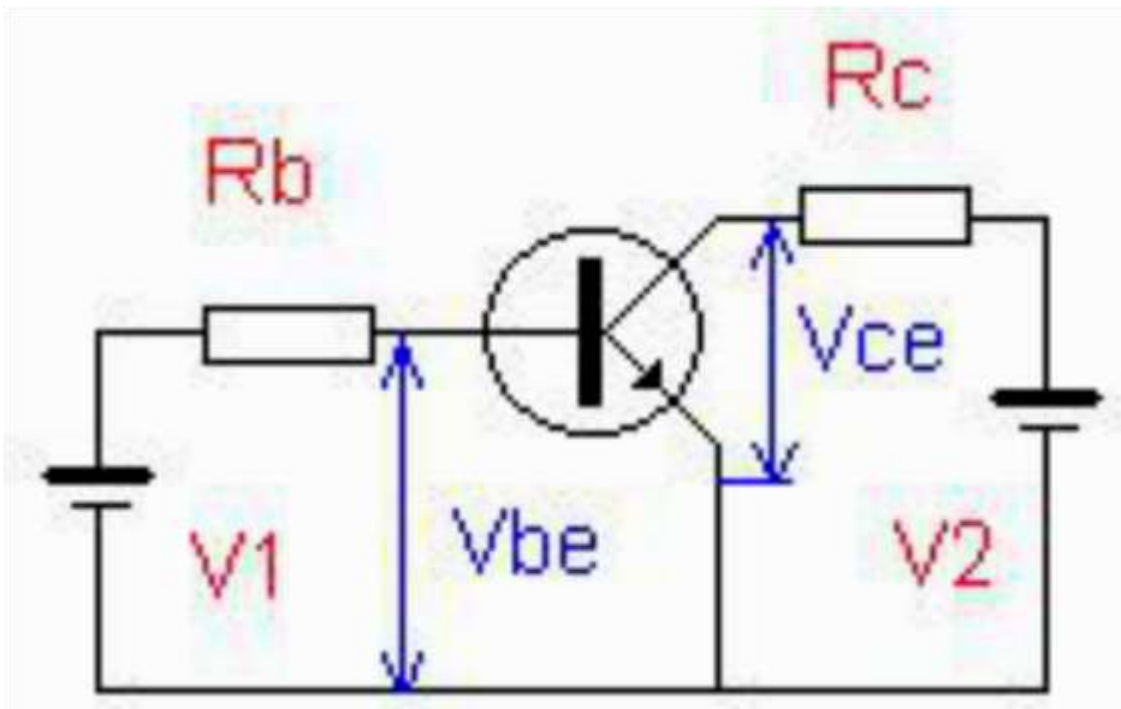
La base est faiblement dopée et très mince, elle transmet au collecteur la plupart des électrons venus de l'émetteur

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).

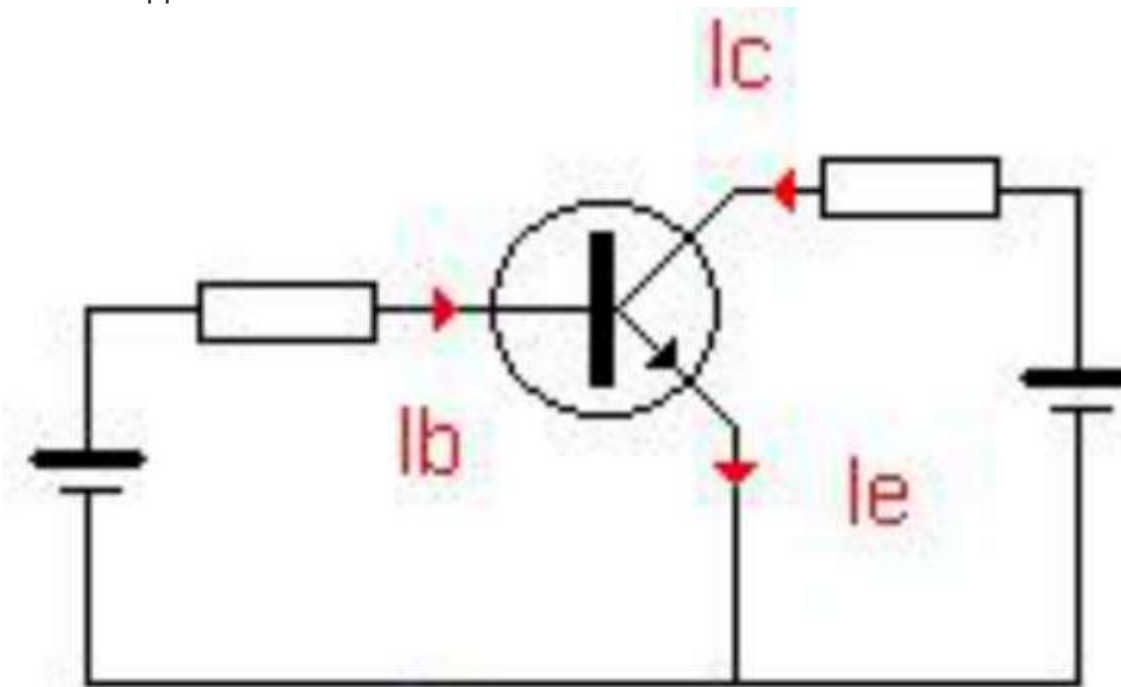
$$I_E = I_C + I_B$$

Le branchement en émetteur commun l'émetteur et les deux sources de tension sont réunies à la masse.



Gain en courant

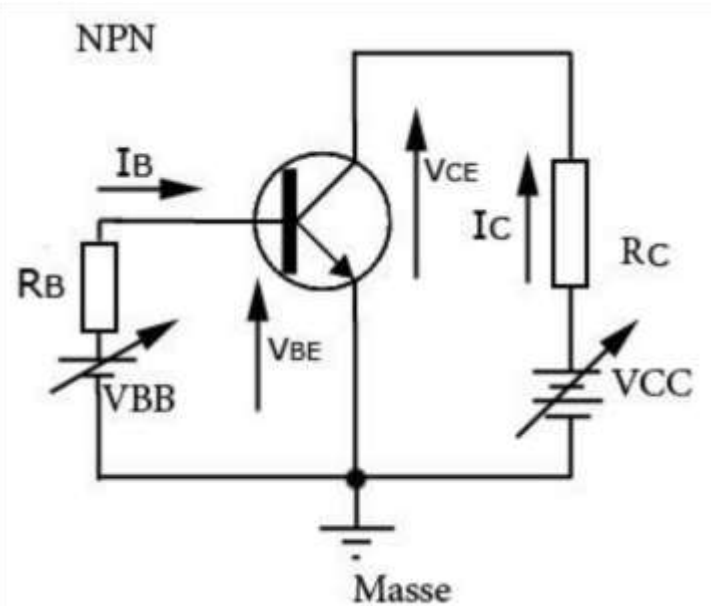
C'est le rapport entre le courant de collecteur et le courant de base :



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

Caractéristiques des transistors.

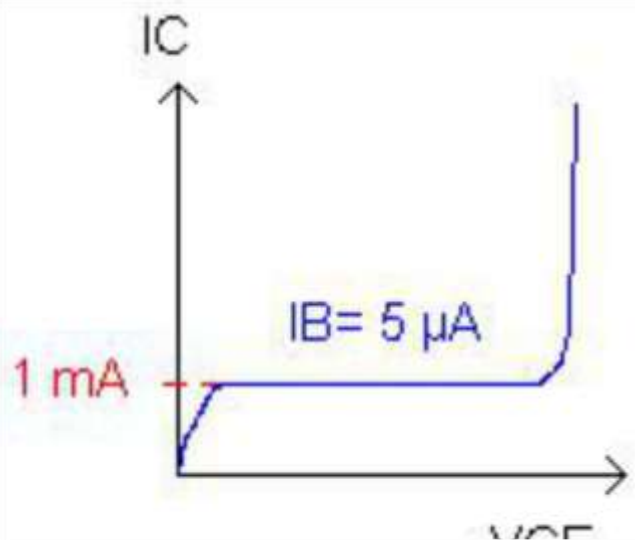
Les sources de tension sont variables, nous les ferons varier et nous mesurerons à chaque fois U et I .



Premier test : Fixons le courant de base à $5\mu\text{A}$, nous obtenons un courant de collecteur de 1mA .

Faisons maintenant varier V_{CC} , à chaque fois on mesure I_C et V_{ce} .

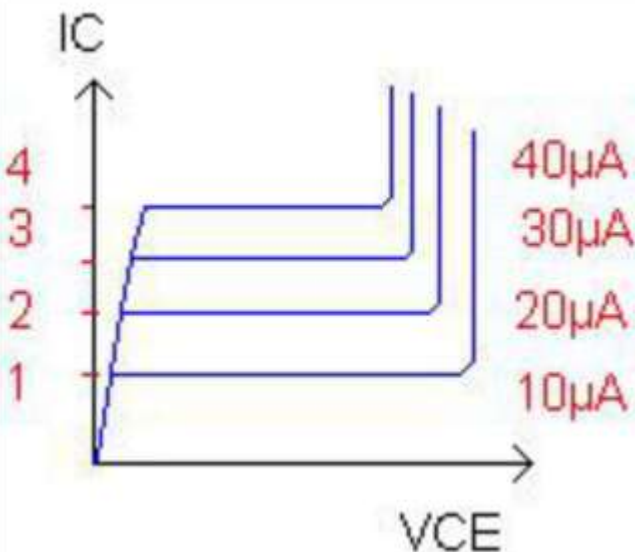
Résultat : Quand on fait varier la tension de collecteur, le courant I_C reste constant sur une grande plage.



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} croît, plus la partie rectiligne diminue.

Toutes les caractéristiques partent du point origine 0. Cela signifie que lorsque V_{CE} est nulle, le courant I_C est également nul et cela quel que soit le courant I_B .



Ensuite, les caractéristiques présentent deux parties. La première est commune et pratiquement verticale ; la seconde est pratiquement horizontale et est fonction du courant I_B .

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.

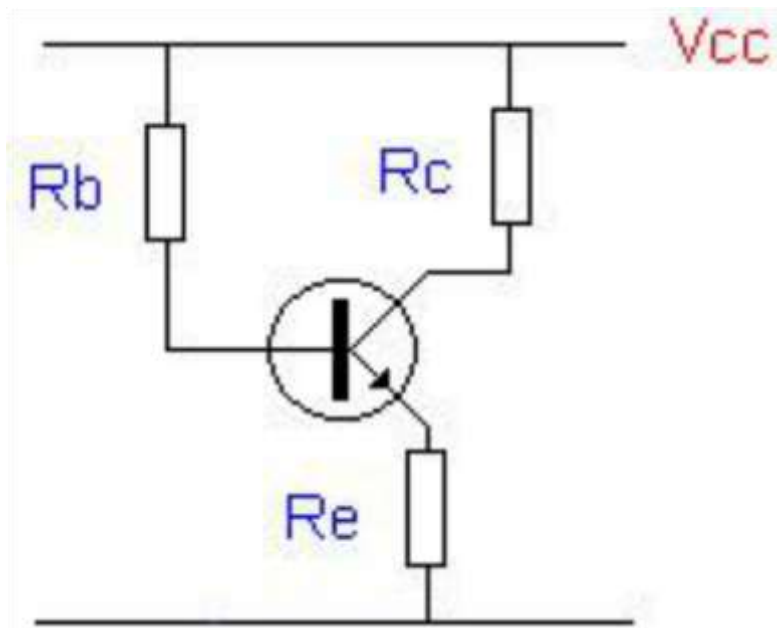
La position d'une caractéristique est ici fonction du courant IB, autrement dit le courant IC est en relation étroite avec le courant IB.

Polarisation des transistors.

Notre transistor, pour fonctionner, a besoin d'être "polarisé". Cela signifie qu'on doit appliquer sur ses connections 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.

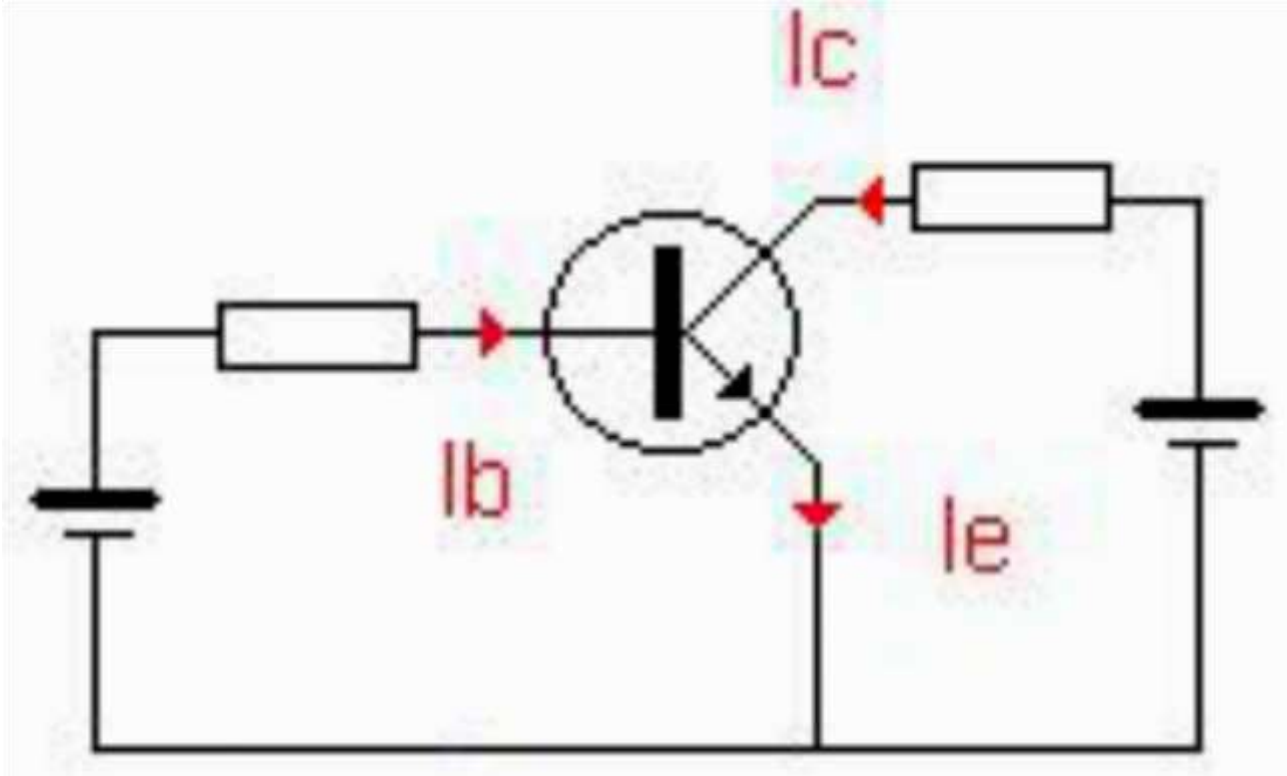
(Dans ce type de schéma, le rail supérieur représente la tension positive d'alimentation notée V_{cc} , le rail inférieur représente la référence du 0V, c'est à dire la masse.)

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 R1 et R2, on utilise donc le théorème du pont diviseur pour obtenir la tension aux bornes de R2:

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

Nous avons sur la base une tension dictée par le pont diviseur R1/R2.

Sur l'émetteur, nous retrouvons cette tension diminuée de VBE 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 :

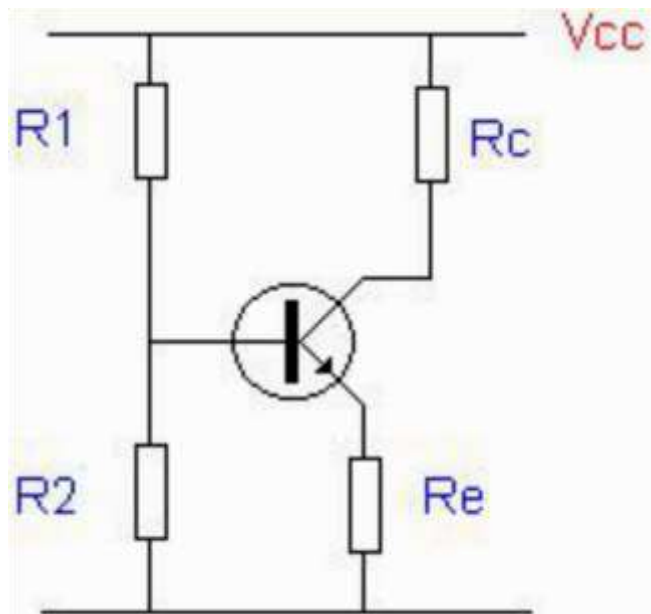
Le courant qui traverse RE sera égal à la tension à ses bornes divisée par la valeur de sa résistance, donc

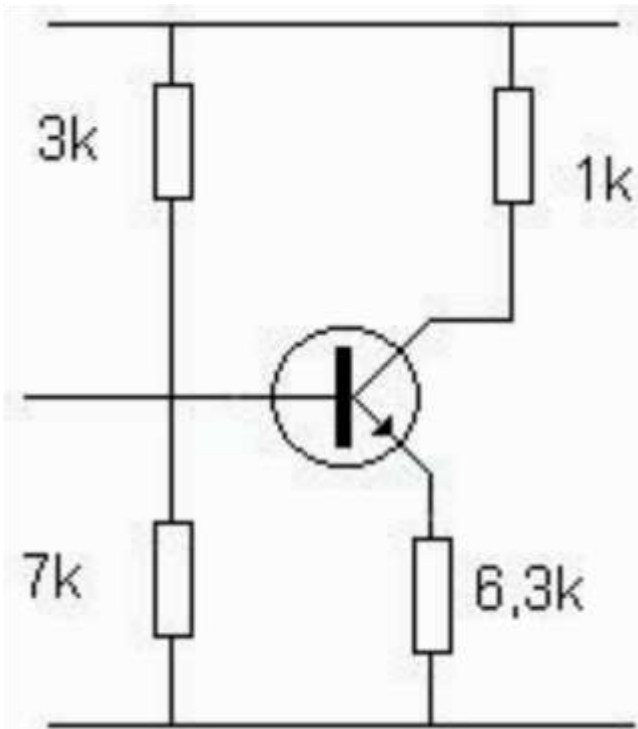
$$I_E = U_{RE} / R_E \text{ soit } I_E = V_E / R_E$$

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

En pratique :

La tension d'alimentation est de 10V





1 - calculons la tension délivrée par le pont résistif R1.R2 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 $V_e = V_b - V_{be}$

$$V_e = 7 - 0,7 = 6,3V \quad \text{0,7 car diode silicium}$$

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

$$I_e = V_e / R_e$$

$$I_e = V_e / R_e$$

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

Le courant collecteur étant à Ib 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 $I_c = I_e$

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

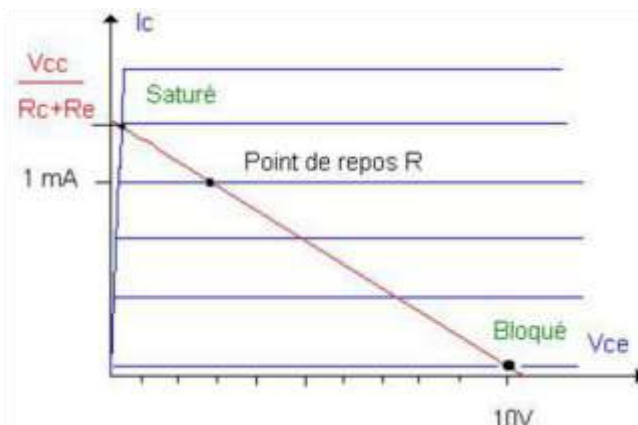
5 - calculons V_{ce}

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

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

6 - notre point de repos est positionné comme suit $I_c = 1 \text{ mA}$

$$V_{ce} = 2.7V$$



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

Equation globale de collecteur :

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

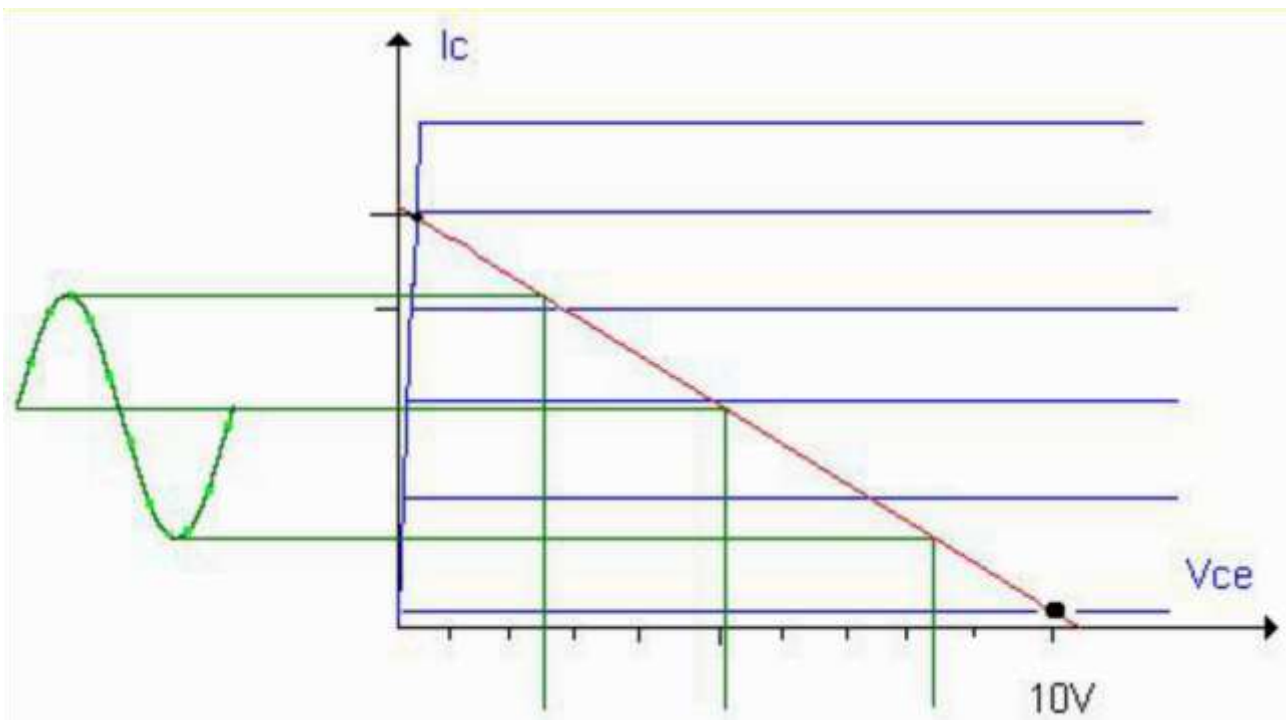
Courant de saturation

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

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

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

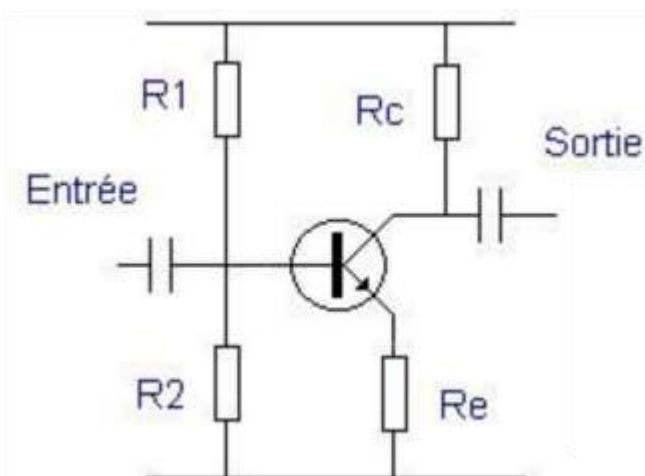
Le montage émetteur commun en amplification



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.

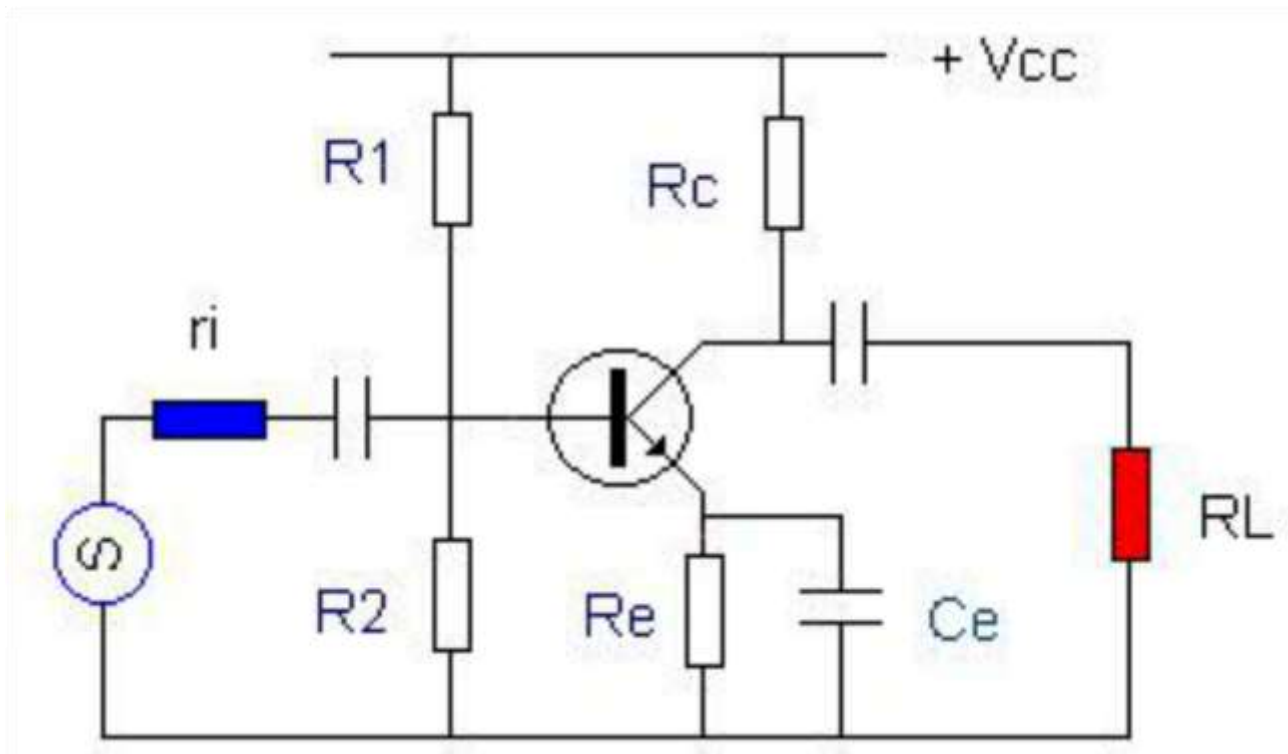
Pour déplacer le point de repos sur la droite de charge, on fait varier la tension de base. On applique une tension alternative à la base, celle-ci va se superposer à la tension de polarisation, on va donc augmenter la tension de polarisation de la jonction base-émetteur, ce qui fera croître I_B et donc 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 variations d'entrée filtrées par un autre condensateur.

Ces condensateurs seront appelés condensateurs de couplage.
Montage complet



le montage complet limite les variations de hautes fréquences mais il nous en fait cadeau.
Le S est la source contenant le courant alternatif.

Les condensateurs ne laissent passer que l'alternatif. Les condensateurs empêchent que le continu viennent se décharger dans la source par exemple.

La barre supérieure et la barre inférieure sont liées pour former un générateur.

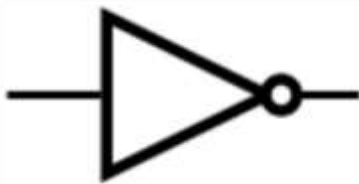
3) Les montages impulsionsnels (ou logique)

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

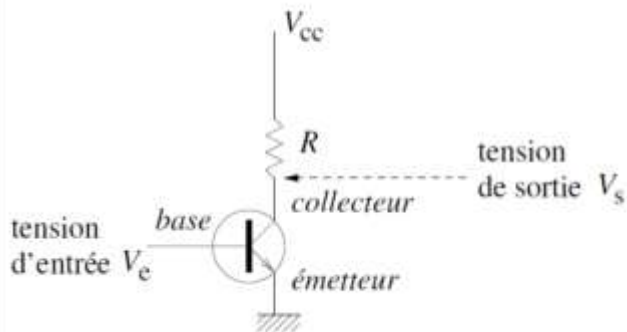
L'état logique peut être noté « 0 » ou « 1 » ou encore « L » (Low) ou « H » (High).

On parle aussi de niveau logique haut ou niveau logique bas, état haut ou état bas.

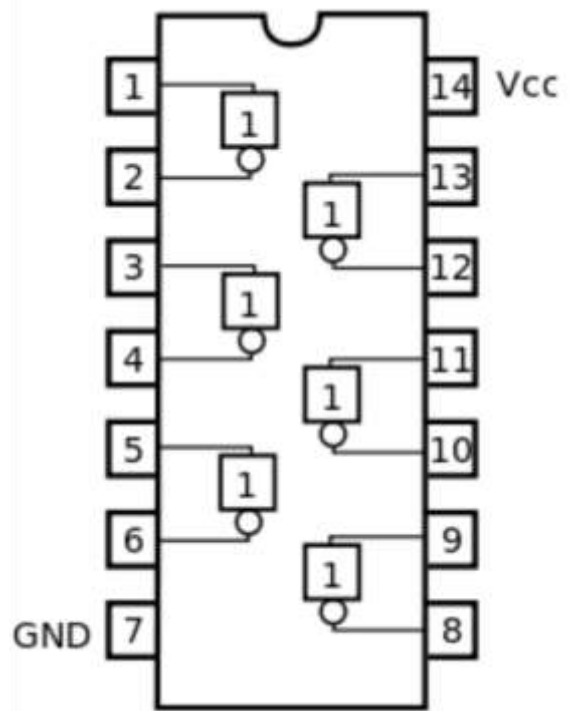
Exemple : la porte " NOT "
inverseur logique



IN	OUT
1	0
0	1

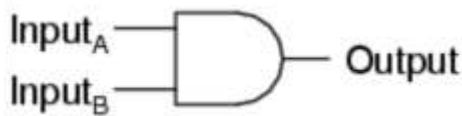


GND est la masse
VCC est l'alimentation

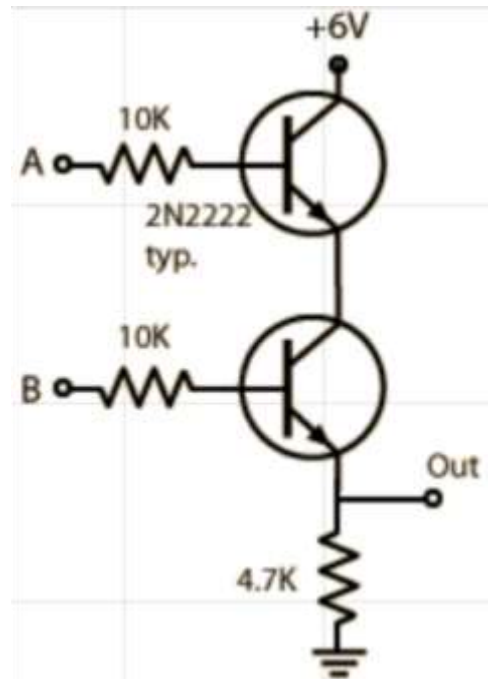
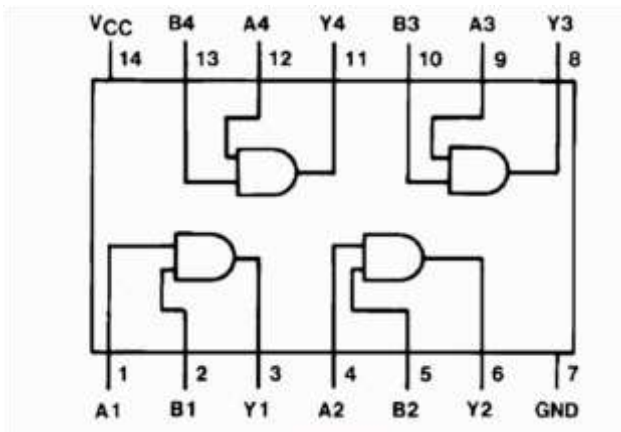


la porte " AND "

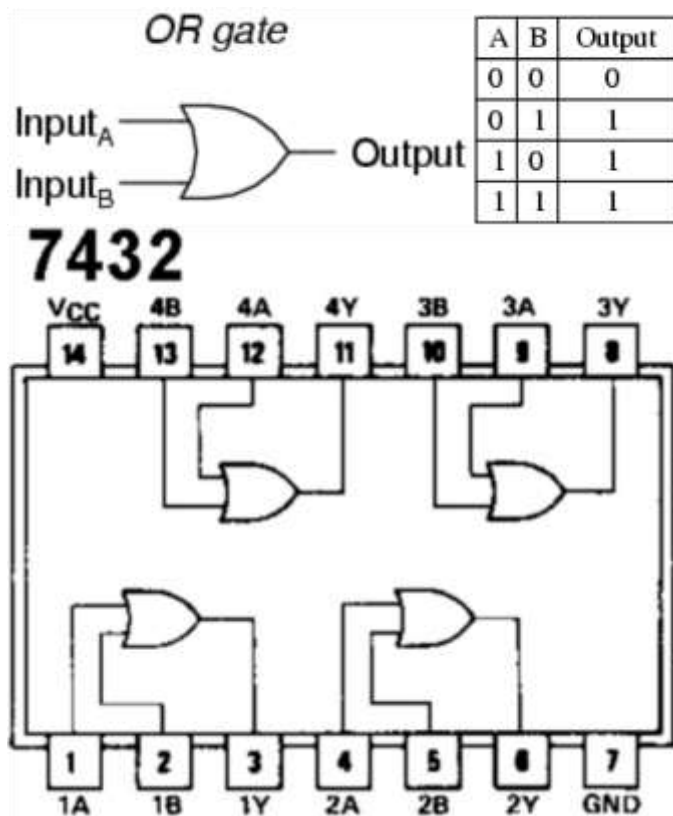
2-input AND gate



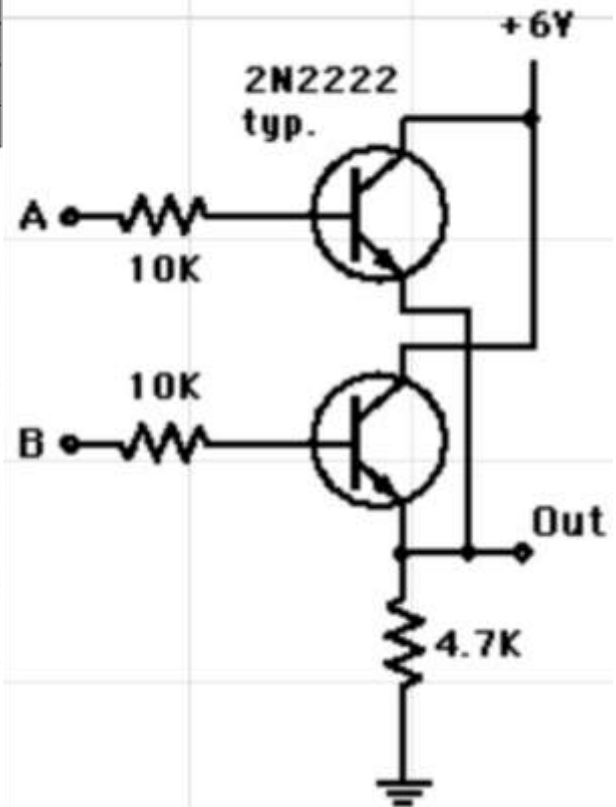
A	B	Output
0	0	0
0	1	0
1	0	0
1	1	1



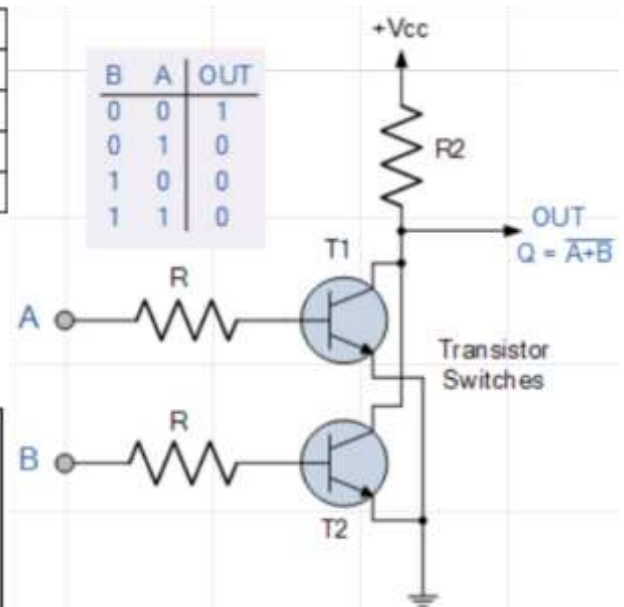
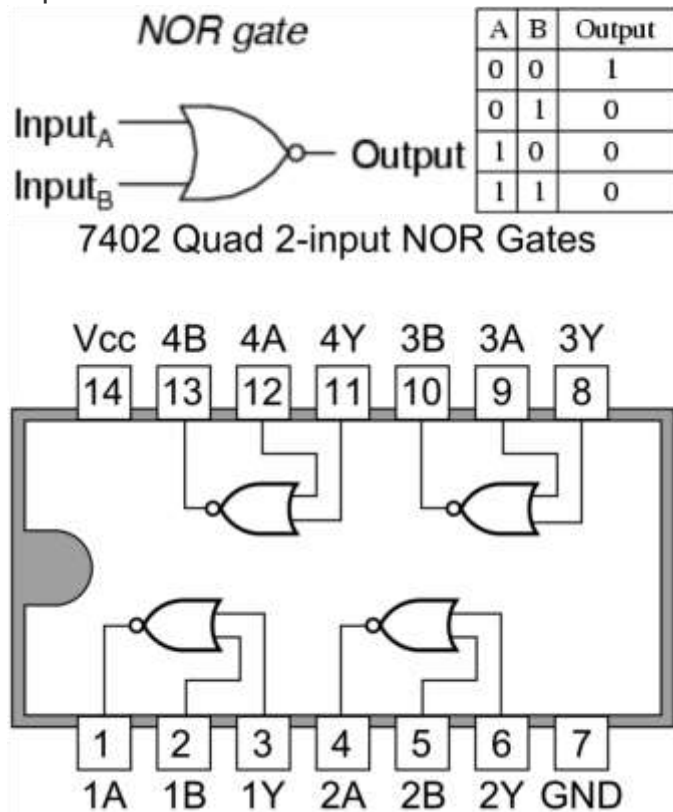
la porte
" OR "



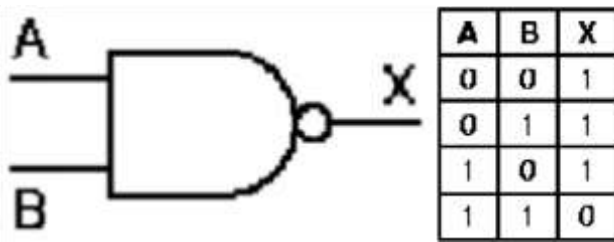
Transistor OR Gate



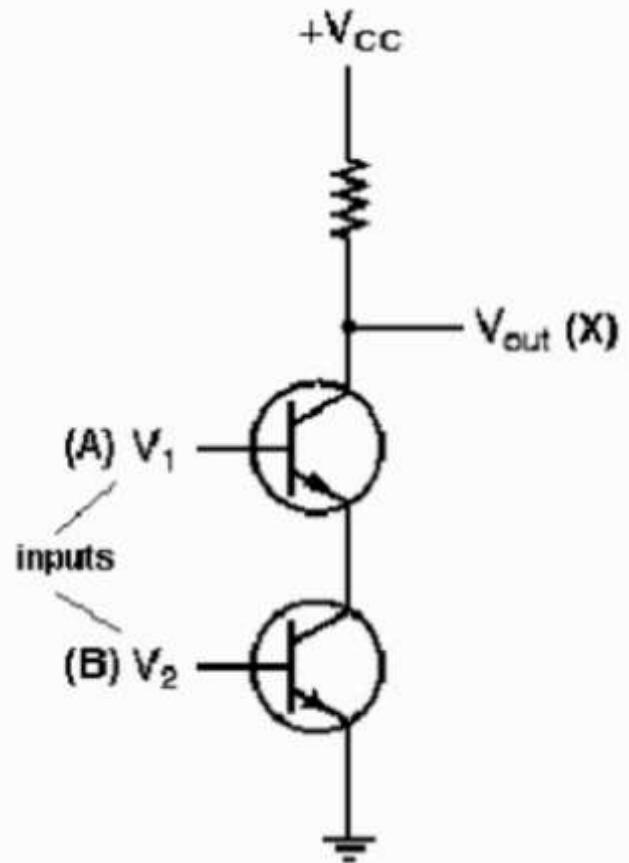
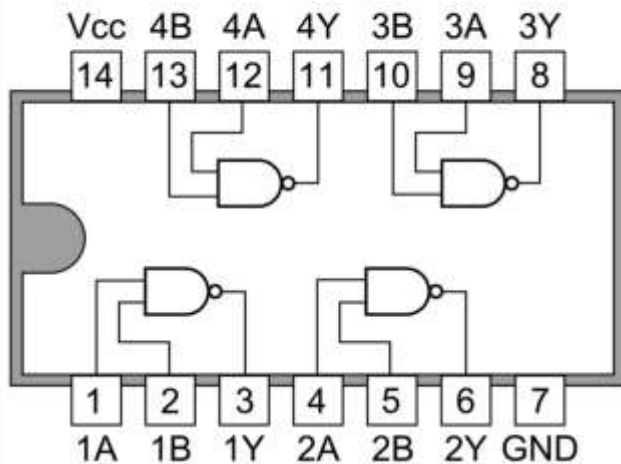
la porte " NOR "



la porte " NAND "



7400 Quad 2-input NAND Gates



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.

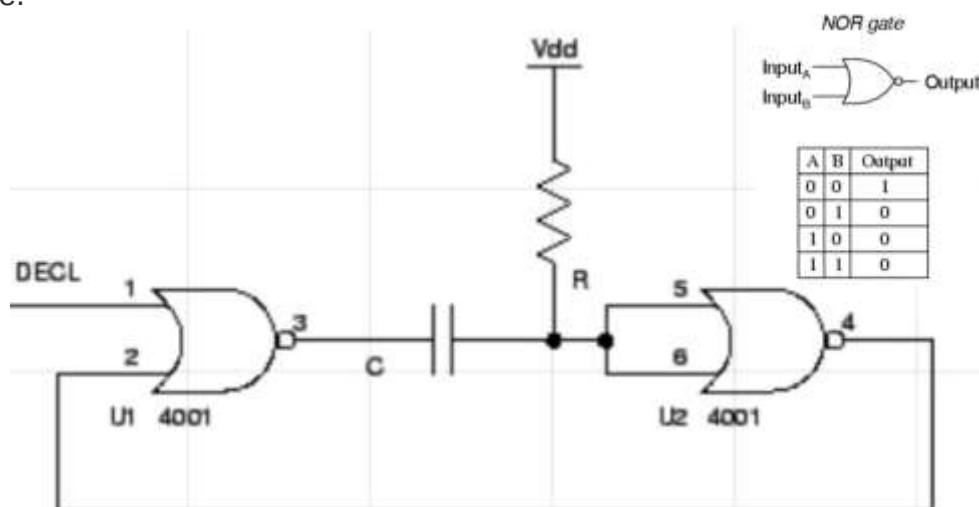
En logique séquentielle, le niveau de sortie dépend aussi du niveau passé des entrées.

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 temps de retour à l'état stable.

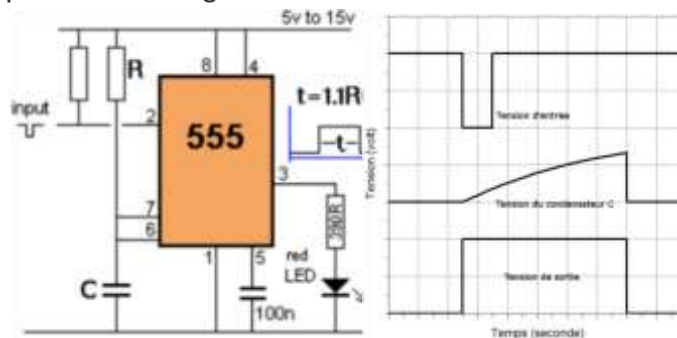
Au repos mettons l'entrée 1 à 0, et supposons 2 à 0, la sortie de cette porte NOR sera égale 1. L'entrée 5 est à 1 puisque reliée à Vdd par l'intermédiaire de R, 6 est à 1 aussi, la sortie est donc à 0. Si l'on ne change rien au niveau de l'entrée 1, rien ne bougera, l'état est stable.

Appliquons un niveau logique 1 sur l'entrée 1 durant un bref instant.

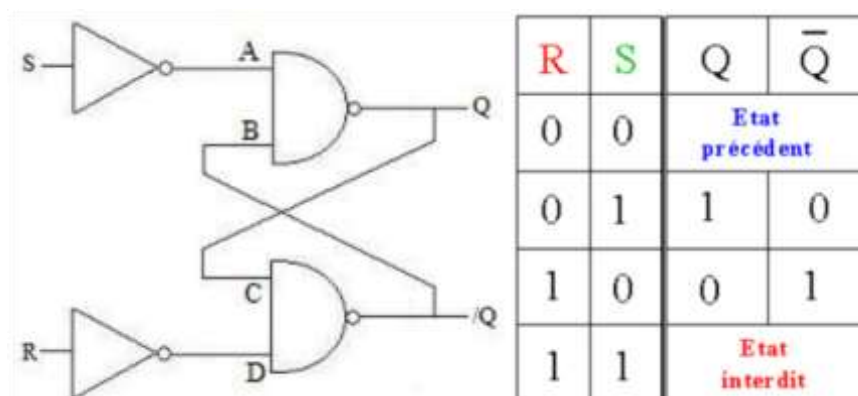
L'entrée 2 est toujours à 0, la sortie passe à 0 ce qui simultanément applique un 0 sur 5 et 6 et donc un niveau 1 à 4 et à 2. Donc une sortie 3 à 0.

Le condensateur se charge à travers R avec une constante de temps $= RC$.

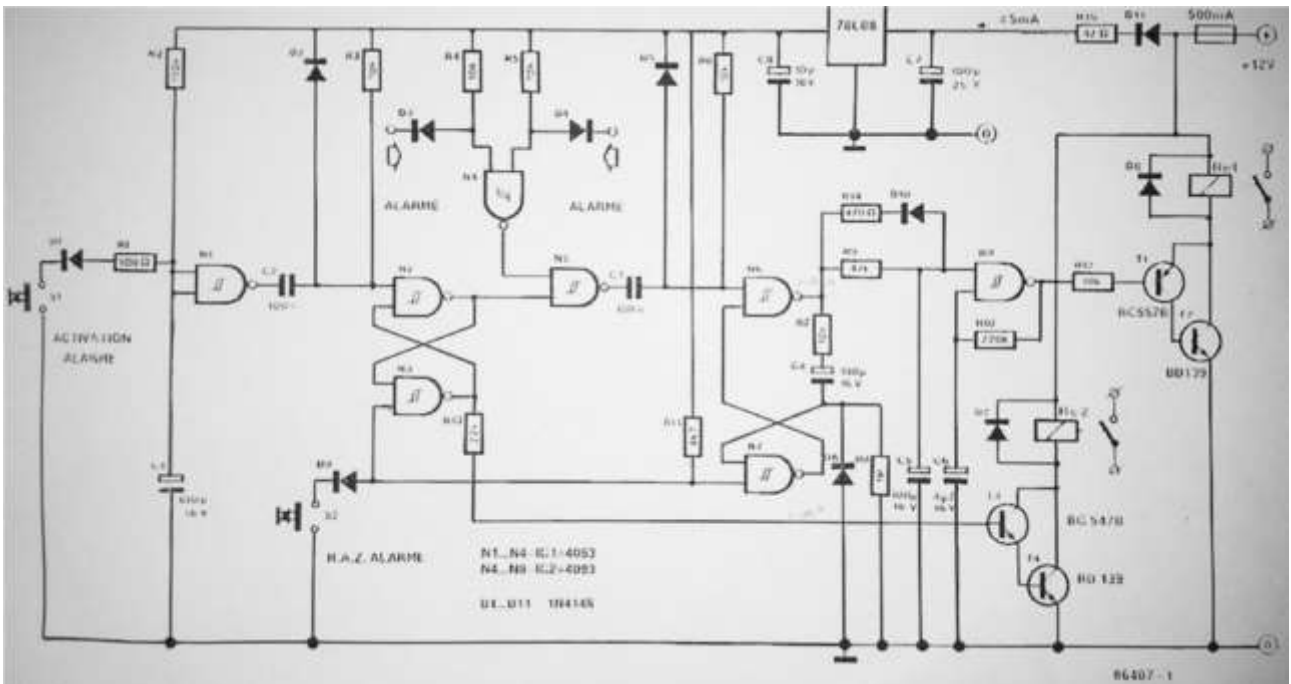
Le condensateur s'est chargé au bout d'un temps "t", les entrées 5 et 6 retrouvent un niveau 1 et provoque le basculement de 4 vers 0, suivi de 2, c'est l'état stable.



les bascules bistables :

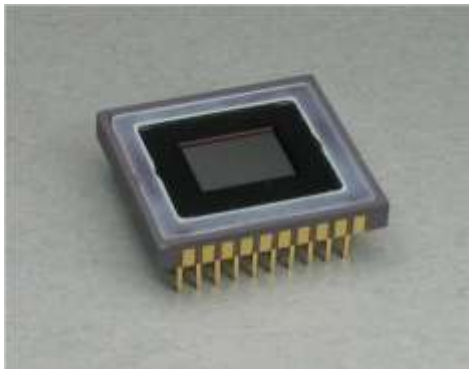


On les appelle bistables car ces bascules ont deux états stables ce qui signifie que s'il n'y a pas d'intervention sur la bascule, celle-ci est verrouillée dans son dernier état. Autrement dit, elle mémorise une information.

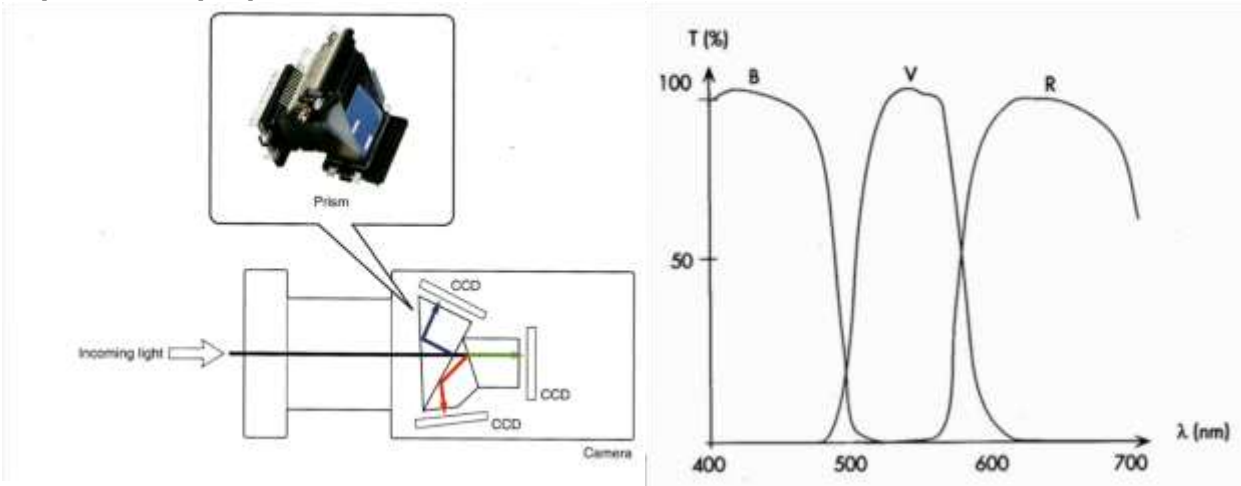


4) Les capteurs

Il existe deux grande familles de capteurs : les CCD et les CMOS, ils ont chacun leurs avantages et leurs inconvénients.
 Les tubes ont pratiquement disparus au profit des capteurs basés sur les semiconducteurs



Séparateur optique :



Un capteur 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 :

112,8 x 9,6mm (1 inch)

8,8 x 6,6mm (2/3 inch) 4,3 x 3,2mm (1/3 inch)

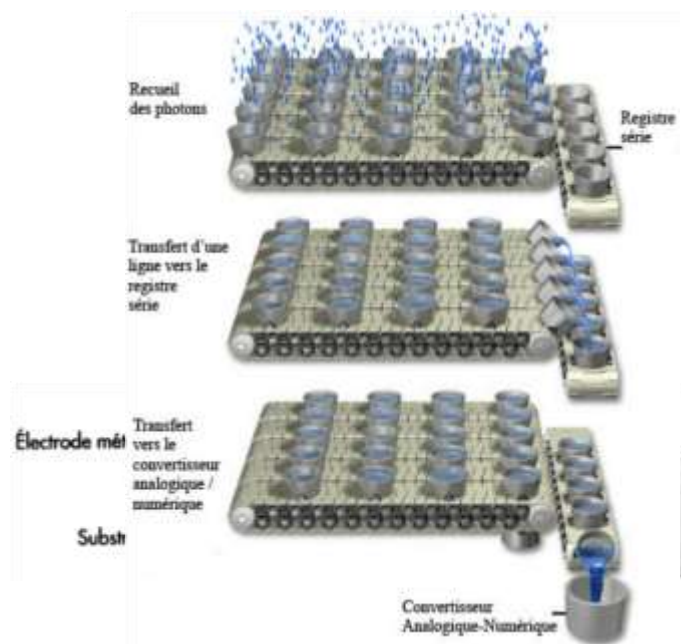
3,2 x 2,4mm (¼ inch) etc...



Les différents types de capteur

Le capteur CCD (Charge-Coupled Device) était le plus répandu il y a encore peu de temps, équipant presque tous les appareils photo et caméras vidéos. Le CCD est composé d'une matrice de cellules photosensibles qui transfère la charge vers un collecteur qui transfère à son tour l'ensemble des charges vers le convertisseur.

Le capteur CMOS (Complementary Metal Oxyde Semi-conductor) fonctionne sur le même principe, à quelques détails près : il se compose d'une matrice de cellules photosensibles également, mais au lieu de transférer la charge vers un collecteur, il la conserve et la transfère au convertisseur directement.



CCD :

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.

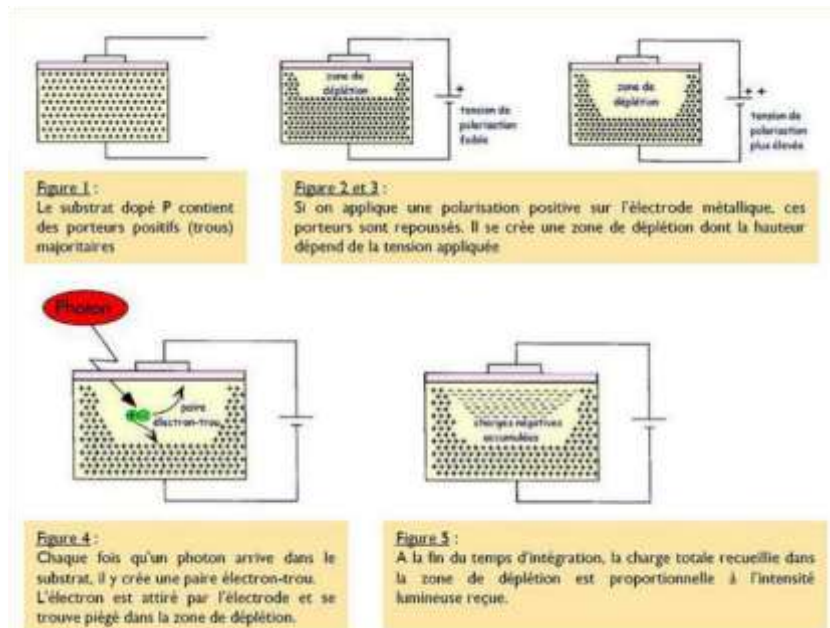
L'électron et le trou se séparent, du fait de la polarisation le trou est repoussé en profondeur et l'électron attiré par l'électrode (qu'il ne peut rejoindre du fait de la couche isolante) et reste dans la zone de déplétion.

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.

Le transfert des charges

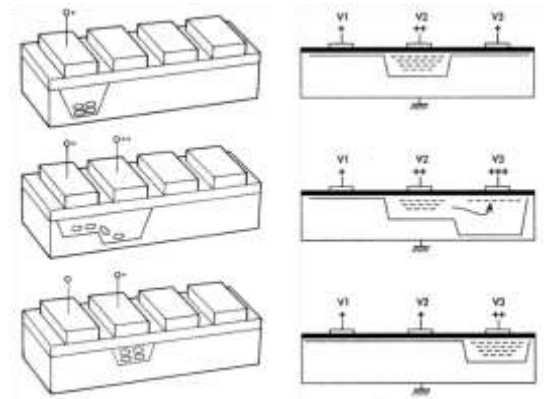
Il reste maintenant à récupérer ces charges et rendre la cellule de nouveau apte à capter les charges de la trame suivante.

Une analogie intéressante aux transfert de charge est le tapis roulant...



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.

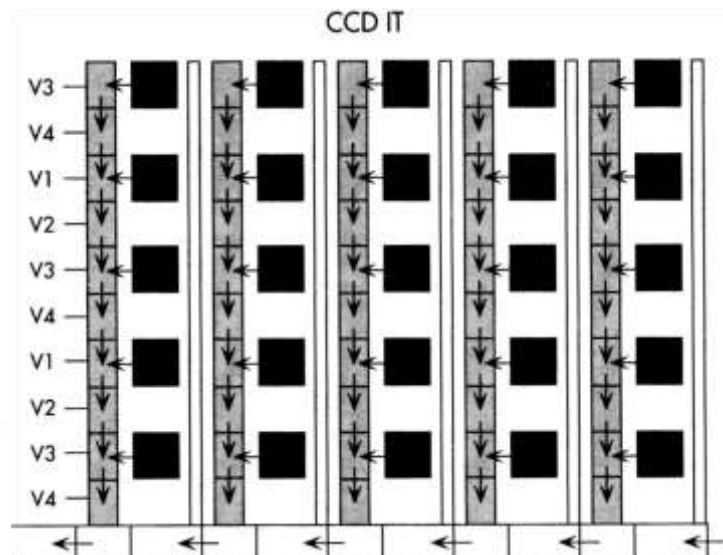
En MOS, nous avons les IT (Interline Transfert), les FT (Frame Transfert) et les FIT (Frame Interline Transfert).

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

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 photodétecteurs alternant avec des

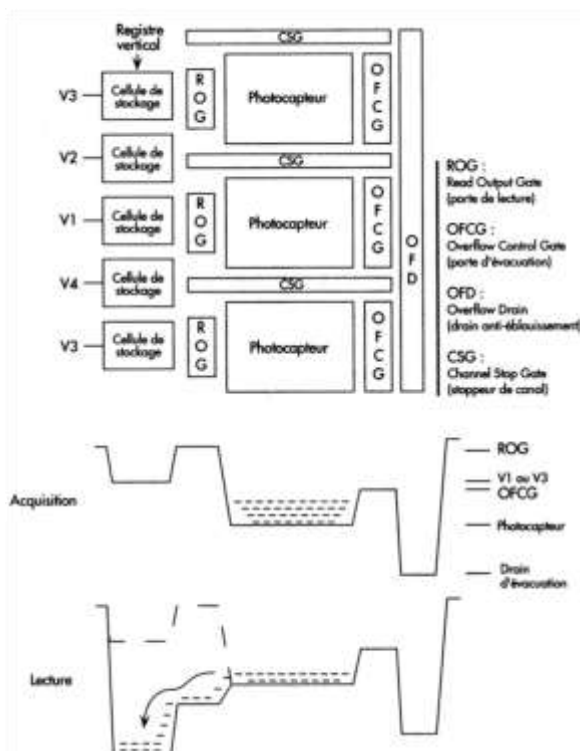


colonnes de cellules de stockage.

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 .

Le SMEAR

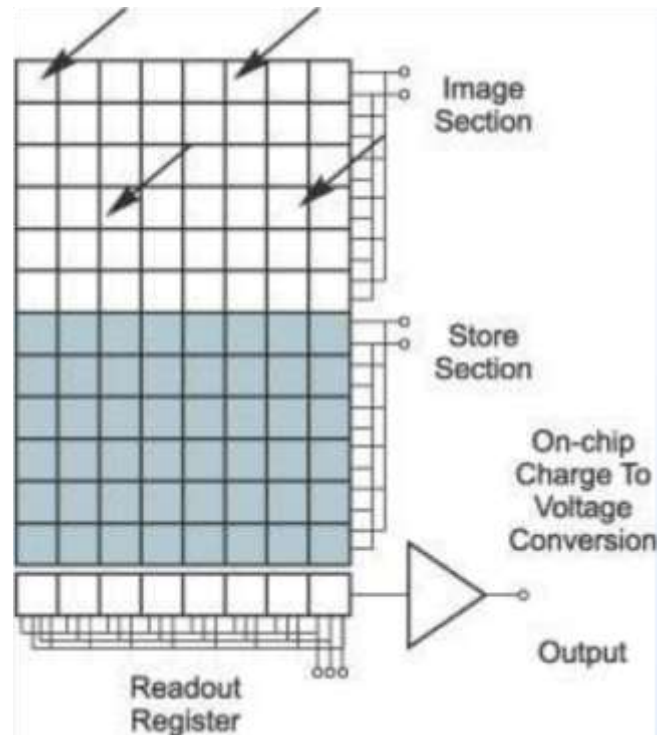
Le smear est caractéristique des capteurs CCD (surtout IT), il se traduit par une ligne verticale blanche ou rouge traversant une zone trop lumineuse (proje , soleil...).



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 (frame transfert)

La zone image d'un capteur FT n'est constituée que de photodétecteurs (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 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 (souvent 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.

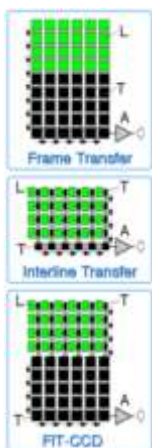
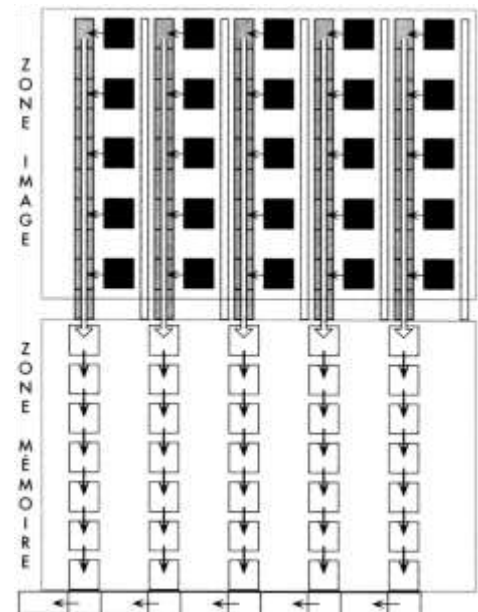


Le CCD FIT

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 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.

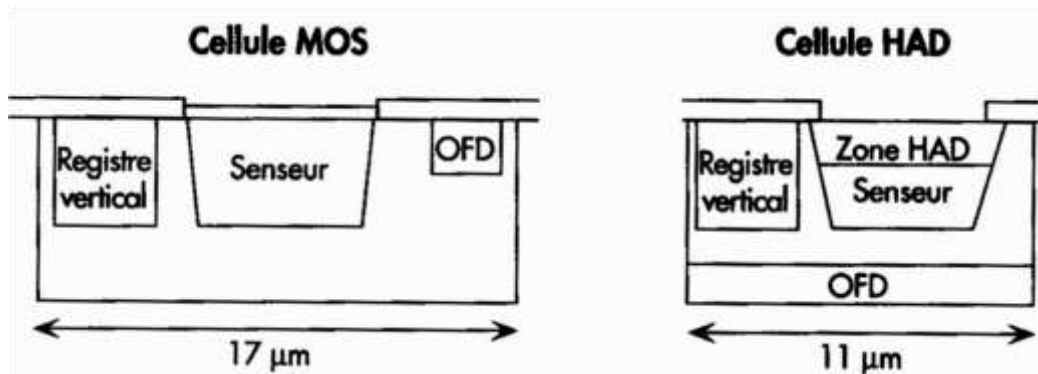


RÉSUMÉ

Les CCD HAD (amélioration du FIT)

Les CCD vu jusqu'à présent étaient du type MOS, il existe une gamme de CCD de type HAD (hole Accumulated Diode)

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.

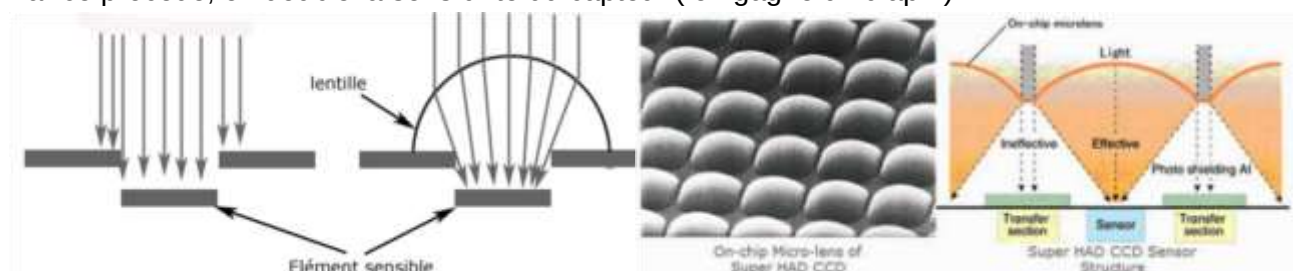
Dans une cellule Had, l'électrode ne se trouve plus en surface mais dans la profondeur, l'absence de cette électrode améliore la sensibilité du capteur qui est relativement faible dans les bleus. Le smear est en outre fortement réduit, en effet, les électrons provoqués par des rayons proche de l'infrarouge (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 photodétecteurs 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.

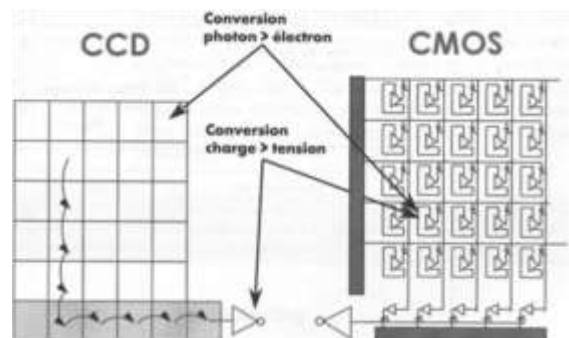
Par ce procédé, on double la sensibilité du capteur (on gagne un diaph).



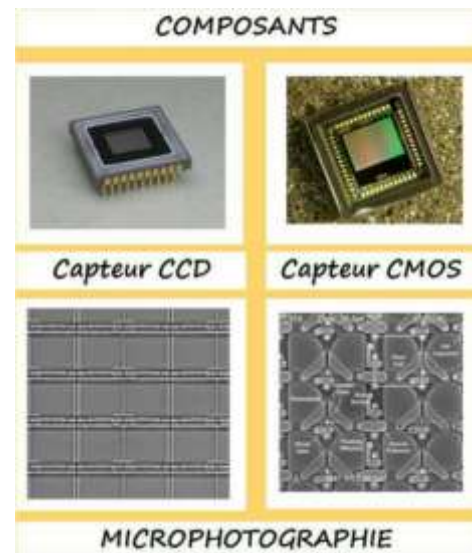
Les capteurs CMOS (Complementary Metal Oxyde Semiconductor)

Comme le CCD, le CMOS utilise le silicium dopé comme substrat et le même effet photoélectrique pour transformer les photons en charges.

La différence essentielle vient du fait que chaque cellule photosensible CMOS incorpore un réseau de transistors qui convertissent les charges en tension. Dans un capteur de type CCD, cette conversion est faite en un point unique, à la sortie du capteur après une série de registres à décalages.



Comme on peut le constater sur la figure ci-contre, les capteurs CCD et CMOS ne présentent pas de différence apparente à l'état de composants. Par contre au niveau de leur structure microscopique, la configuration est totalement différente comme le montre les microphotographies des composants. Alors que la surface de la cellule du capteur CCD est totalement réceptive au flux photonique, celle du capteur CMOS est en partie occupée par l'amplificateur et donc partiellement réceptive au flux photonique.



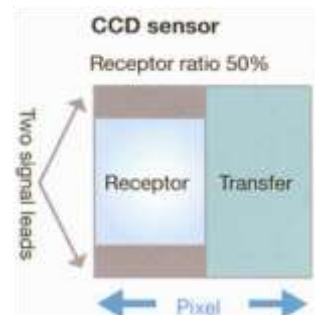
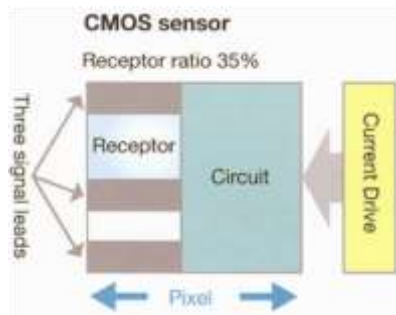
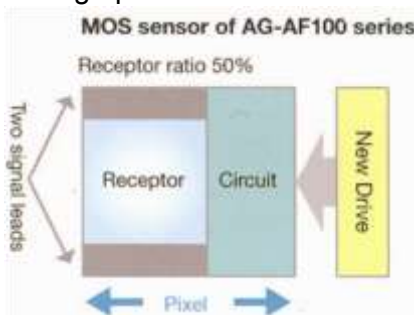
Avantages des CMOS :

- Ils reposent sur la même technologie que les microprocesseurs et les mémoires, les lignes de productions sont donc les mêmes, large diffusion = prix inférieurs.
- Consommation électrique moindre qu'un CCD

Inconvénients :

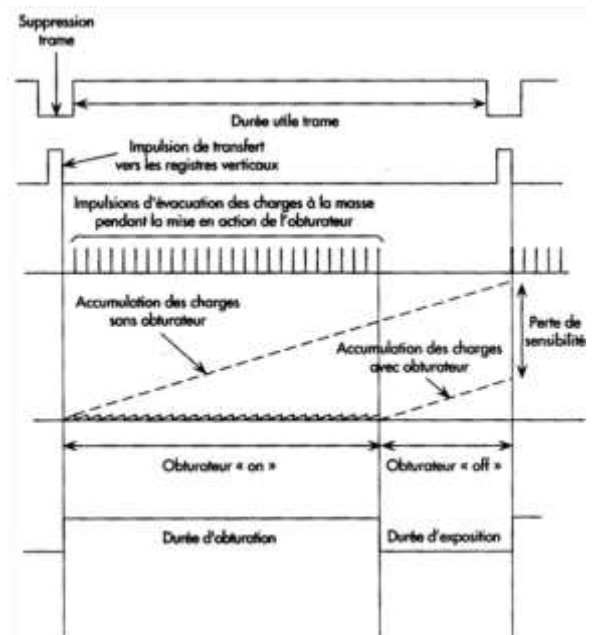
- Moins bonne linéarité puisque chaque pixel comporte ses propres circuits.
- Les nombreux circuits « obscurcissent » les zones sensibles. Le ccd reste donc plus sensible dans les basses lumières.

Dernièrement, plusieurs technologies basées autour du CMOS ont fait leur apparition. En tête de liste, Panasonic avec ses capteurs MOS (AG-AF100), un compromis permettant une qualité d'image proche du CDD avec la rapidité et la faible consommation du CMOS.



5) Caméra

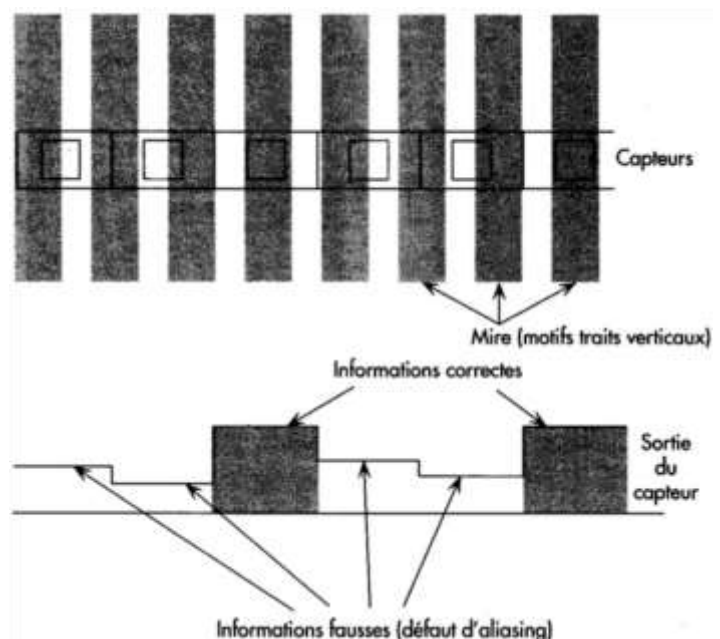
Le shutter (Obturateur électronique) La technologie des capteurs électronique permet de changer la durée d'impression de chaque trame. En fonctionnement normal, le temps de chargement de la zone d'acquisition vaut 1/50 de sec soit la durée d'une trame. Avec un shutter, on peut par exemple choisir de 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.



L'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).

$$F_{\text{éch}} = \frac{\text{Nombre de points par ligne}}{\text{Durée d'une ligne}}$$



Donc pour un capteur de 786 pts/ligne

Durée active d'une ligne = 52 μ s ●

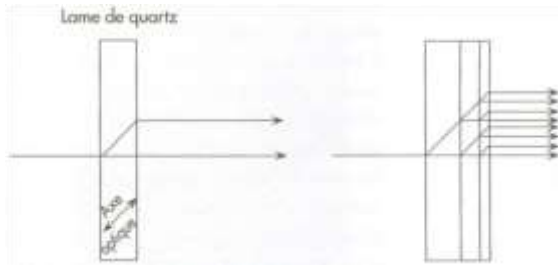
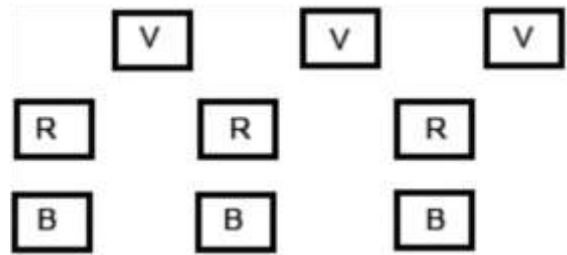
Fréq éch = 15 MHz

● il faut donc supprimer toute fréquence supérieure à 7.5 MHz

Solutions

Le décalage spatial :

Cette solution consiste à combler les zones aveugles du capteur en décalant horizontalement le capteur vert d'un demi pixel. 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.



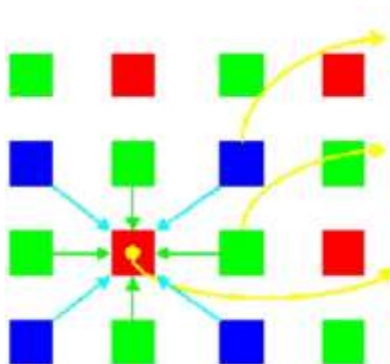
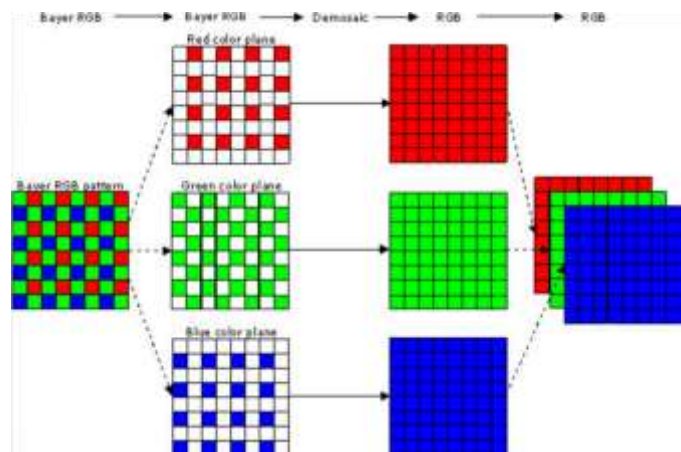
fréquences trop élevée d'atteindre le CCD.

Le filtre optique passe-bas :

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

Filtre de Bayer

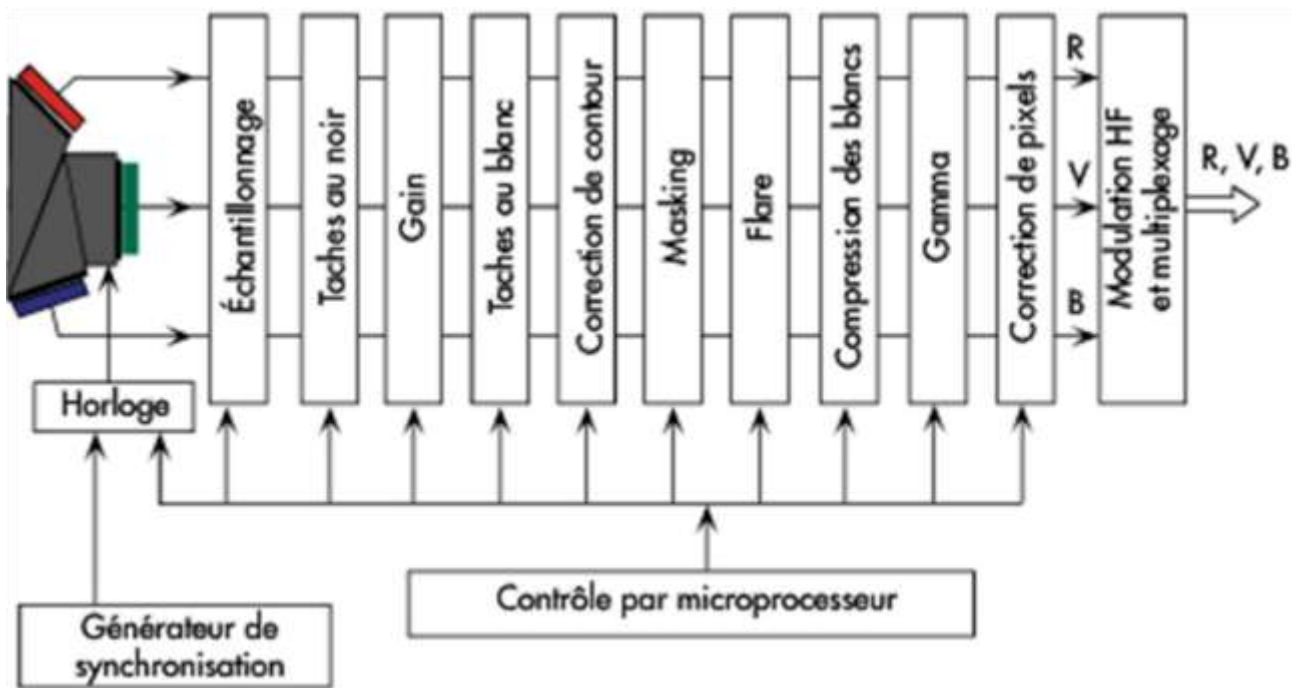
Si on ne souhaite utiliser qu'un seul capteur, il faudra surmonter chaque photosite d'un filtre pour l'affecter à une couleur donnée. L'introduction d'un filtre RVB de Bayer appelé aussi filtre en mosaïque conduira inévitablement à une perte de définition au niveau de chaque couleur. Le nombre de photosites sensibles au vert est deux fois plus élevé que ceux sensibles au bleu ou au rouge, ce qui correspond à la sensibilité de l'oeil. 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.



- le photosite mesure l'intensité du bleu
- le vert est calculé par moyenne sur les 4 verts voisins
- le rouge est calculé par moyenne sur les 4 rouges voisins
- le photosite mesure l'intensité du vert
- le bleu est calculé par moyenne sur les 2 bleus voisins
- le rouge est calculé par moyenne sur les 2 rouges voisins
- le photosite mesure l'intensité du rouge
- le bleu est calculé par moyenne sur les 4 bleus voisins
- le vert est calculé par moyenne sur les 4 verts voisins



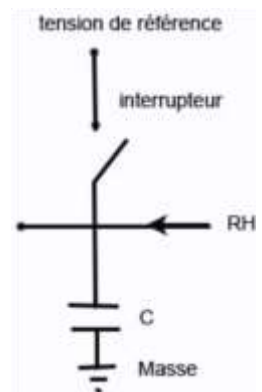
Traitement vidéo



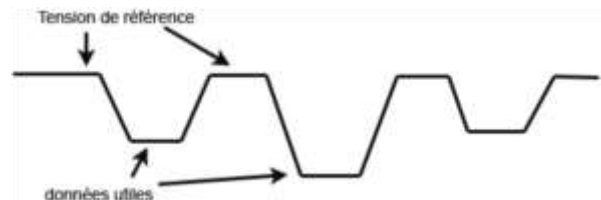
les charges sortant du registre horizontal sont inexploitable telles quelles, il faut les convertir en niveaux de tension, l'opération est réalisée par le circuit d'échantillonnage.

- 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 ou pixel.(environs 15MHz)



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

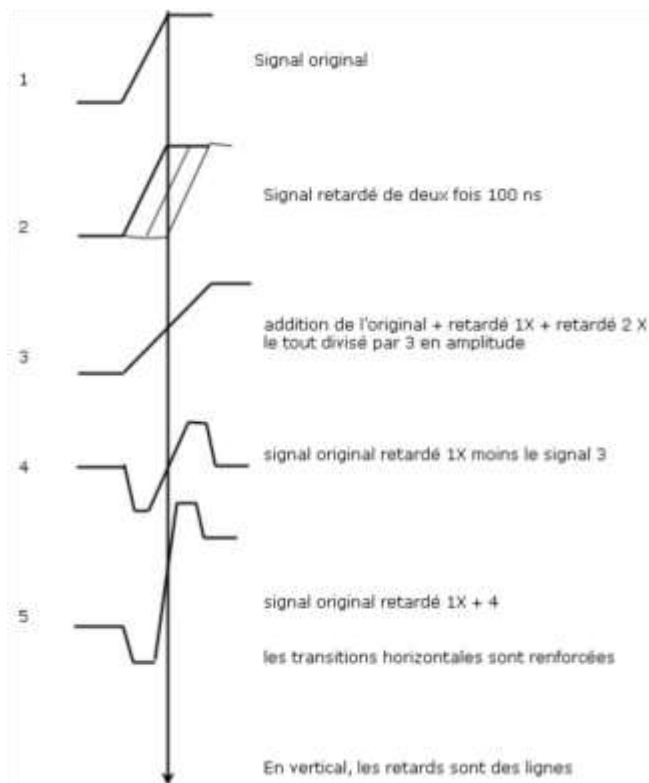
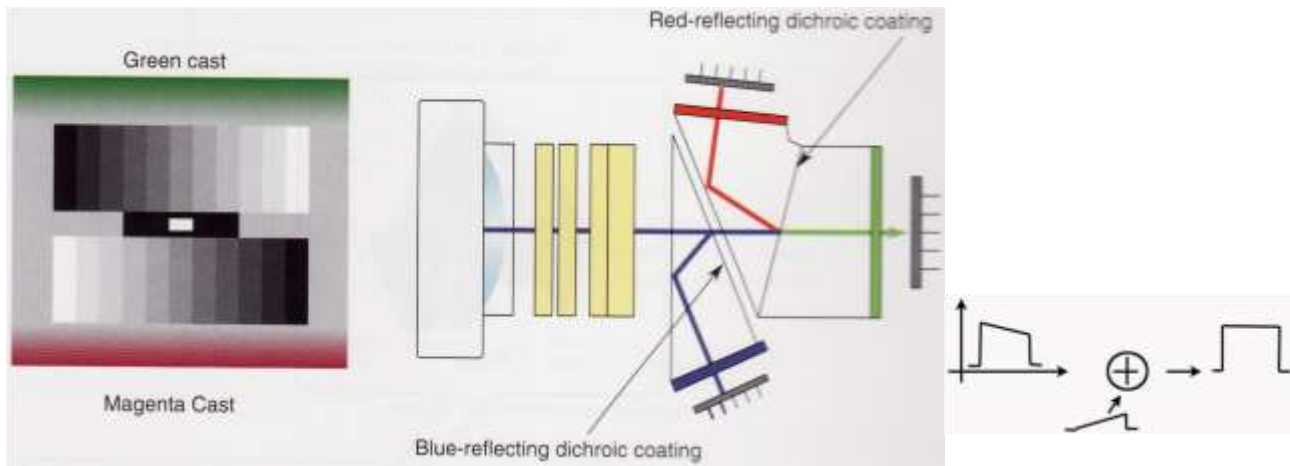


black shading (tache au noir) le black shading est provoqué par des variations de courant d'obscurité du à la température, cela donne des coloration en horizontal, vertical ou en parabole.

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

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.





correction de contour

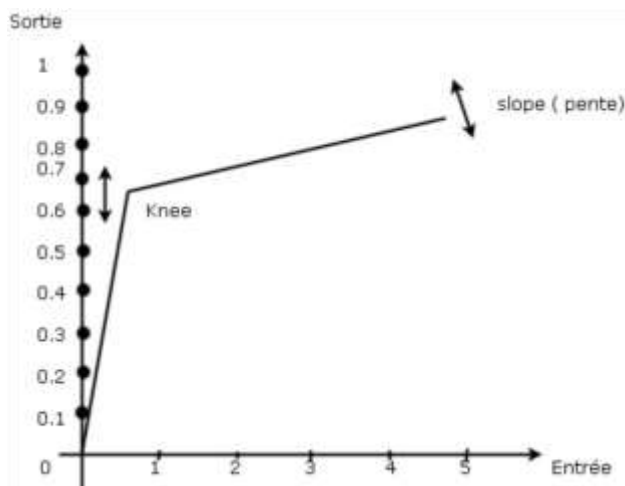
Le but de la correction de contour est de rajouter du piqué à une image rendue molle par le filtre passe-bas.

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

correction de flare

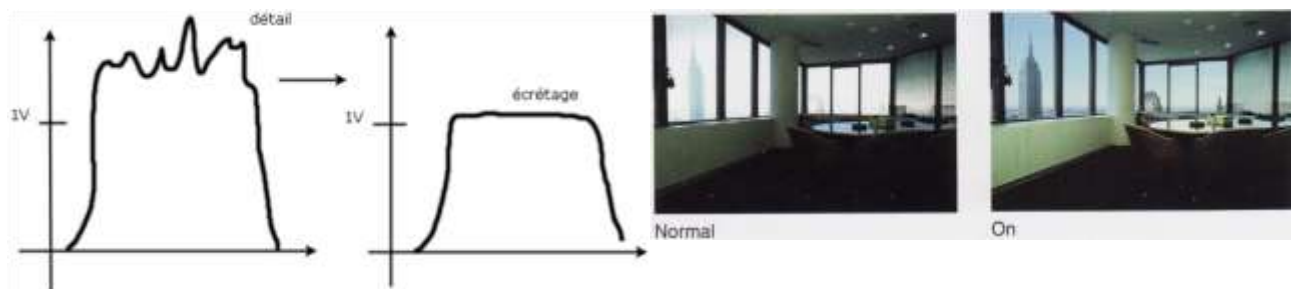
Causé par la diffusion parasite de lumière à l'intérieur de l'objectif, plus précisément quand l'ouverture ou la focale varie. En pratique on utilise une mire blanche avec un carré de velours noir au centre. - on zoome pour que le carré noir occupe 90% de l'image

- 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



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)



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

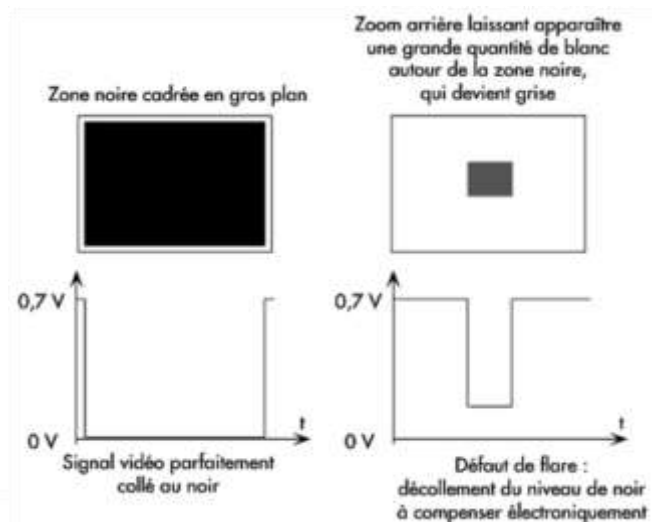
- mémorisation de leur position dans la tête de caméra

- 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.



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

- le knee ou seuil : niveau à partir

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

6) Le transistor à effet de champs

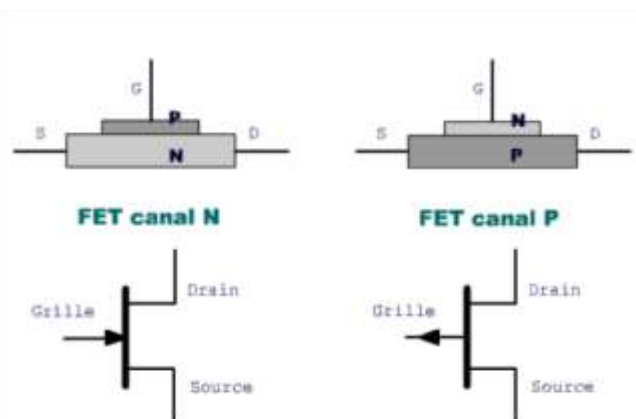
Intro

Dans la partie consacrée au transistor bipolaire, nous avons vu que le courant de sortie sur le collecteur est proportionnel au courant d'entrée sur la base.

Le transistor bipolaire est donc un dispositif piloté par un courant.

Le transistor à effet de champ (Field effect transistor ou FET) utilise une tension sur la borne d'entrée du transistor, appelée la grille afin de contrôler le courant qui le traverse. Cette dépendance se base sur l'effet du champ électrique généré par l'électrode de grille (d'où le nom de transistor à effet de champ).

Le transistor à effet de champ est donc un transistor commandé en tension.



Transistor bipolaire	Transistor à effet de champ
Emetteur - (E)	Source - (S)
Base - (B)	Grille - (G)
Collecteur - (C)	Drain - (D)

Comparaison entre les bornes du transistor bipolaire et du transistor à effet de champ.

Fonctionnement

Le FET "Field Effect Transistor" travaille d'une manière différente du transistor bipolaire. Ce dernier était commandé par un courant, le courant de base.

Dans la cas du FET la commande est faite par une tension appliquée à une "grille" ou "gate" en anglais.

Le courant traverse une mince canal de type N (ou P) surmonté d'une grille dopée en sens inverse.

Les deux électrodes situées de part et d'autre du canal sont appelées "Source" et "Drain".

Ces deux électrodes sont à priori équivalentes mais on appelle source celle qui fournit les porteurs majoritaires, les électrons dans un canal N ou les trous pour un FET à canal P.

La source est aussi le point par rapport auquel on mesure le potentiel de la grille V_{GS}

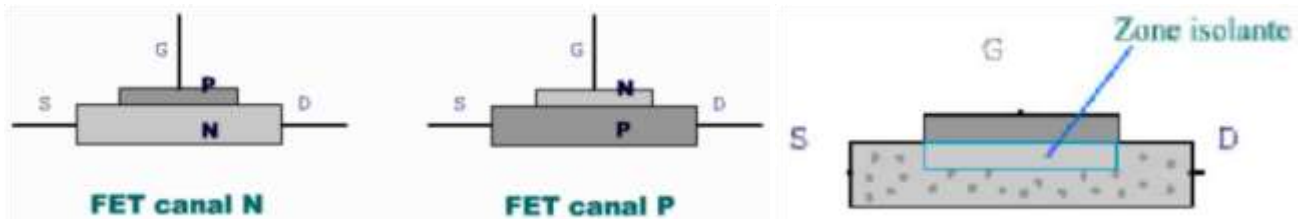
Le plus utilisé étant le FET à canal N, nous nous limiterons à celui là.

La grille étant dopée P et le canal en N, il se forme donc une zone de déplétion entre ces deux zones (comme une diode) : ions négatifs en haut, ions positifs en bas. Lorsque la tension de grille est de 0V ($V_{GS} = 0$) et qu'une petite tension (V_{DS}) est appliquée entre le drain et la source, la zone de déplétion est très fine.

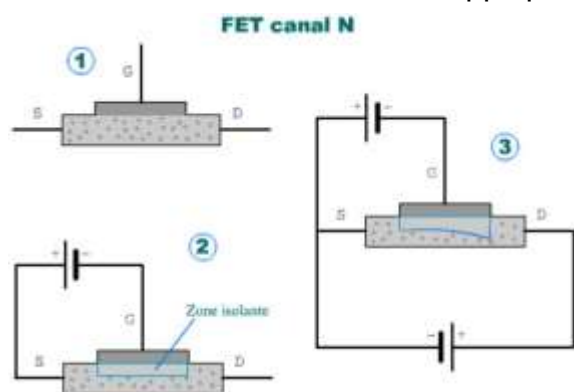
C'est là que le courant à travers le canal, I_D , est le plus grand. Ce courant s'appellera le courant maximum de saturation (I_{DSS}). Le FET est alors fortement conducteur. Comme la jonction PN grille-canal est polarisée en inverse, le courant qui va la traverser sera très faible et sera même fréquemment négligé.

Dans ce cas, le courant de source (I_S) sera égal au courant de drain (I_D).

$I_G = 0$ donc $I_D = I_S$



Le transistor se commande en appliquant entre la grille et la source une tension inverse.



Aucun courant ne passe donc dans la grille mais sous la grille le nombre de porteurs diminue comme dans une diode à jonction polarisée en sens inverse.

Le canal entre la source et le drain devient donc d'autant plus étroit que la tension entre la source et la grille est importante.

A partir d'une certaine valeur de la tension V_{GS} la résistance R_{DS} tend même vers l'infini car le canal ne comporte pratiquement plus de porteurs.

Le transistor MOSFET

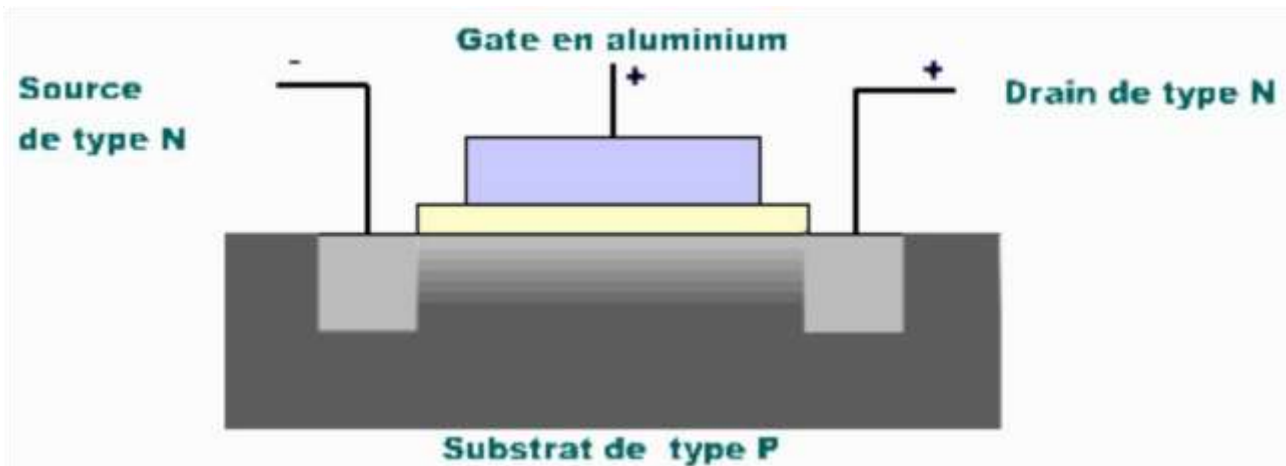
Le transistor FET est surtout utilisé comme composant analogique.

Les circuits intégrés numériques font plus souvent appel aux transistors MOSFET.

La grille en métal, généralement de l'aluminium, est séparée du substrat par un isolant (SiO_2).

Comme pour le FET, le MOSFET peut fonctionner avec un canal de type N ou P ; on parlera alors de NMOS ou de PMOS.

Prenons l'exemple d'un MOSFET à canal N.



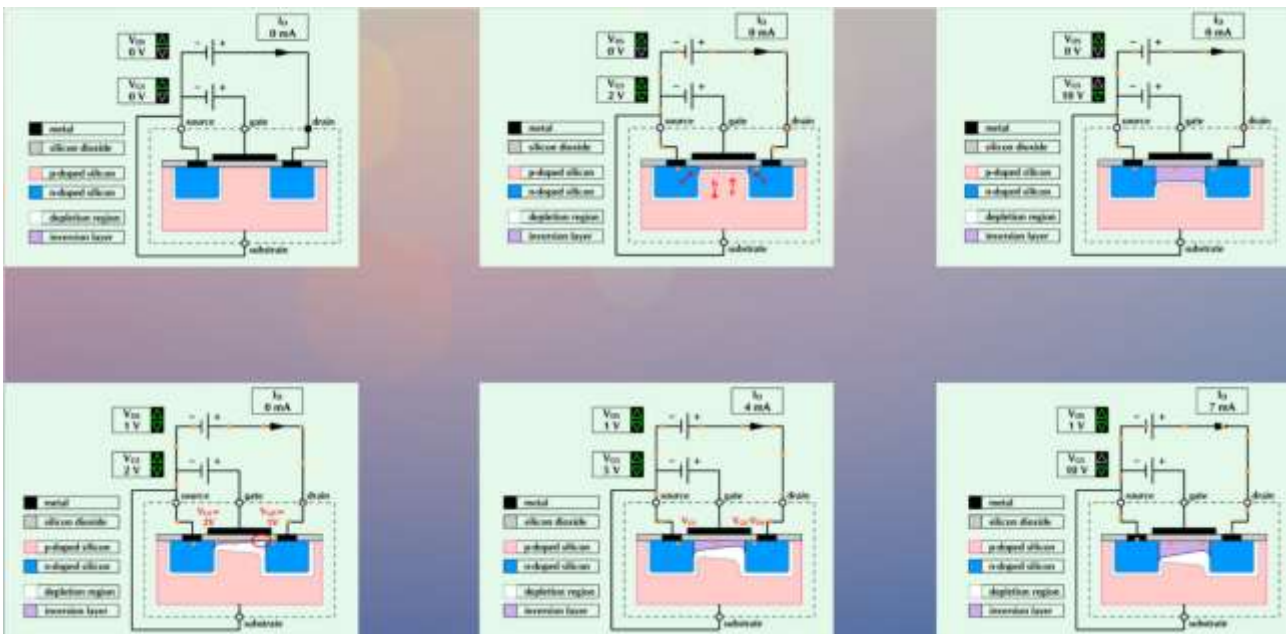
Le substrat est de type P. La source et le drain sont deux électrodes de type N insérées par diffusion dans le substrat.

Ces deux électrodes et la zone de type P qui les sépare, équivalent à deux diodes en tête bêche, le courant ne passe pas entre le drain et la source.

L'intervalle entre la source et le drain est recouvert d'une couche d'oxyde de silicium isolante puis d'une "grille" en aluminium .

Lorsque la grille est rendue positive par rapport au substrat, elle attire les électrons sous l'isolant pour former un canal entre la source et le drain qui ne contient plus de trous comme le reste du substrat mais un excédent d'électrons.

La tension appliquée à la grille fait varier la conductivité du canal. Elle module de ce fait le courant entre la source et le drain.



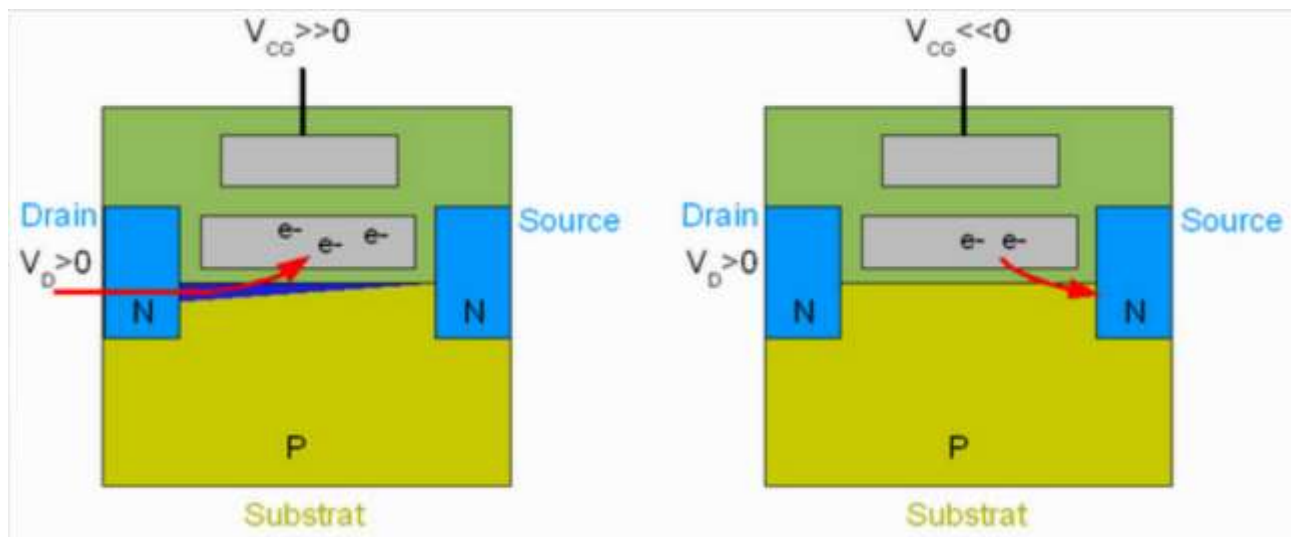
FGMOS – floating gate MOS transistor

On peut aussi utiliser les fet pour mémoriser une information (1 bit) , on utilise alors des fet à grille flottante.

Un transistor à grille flottante comporte deux grilles. La première, la grille de contrôle permet l'application d'une tension extérieure.

La seconde, la grille flottante est électriquement isolée du reste du dispositif. Lorsqu'une tension suffisante est appliquée à la grille de contrôle, des charges peuvent traverser l'oxyde de grille et se retrouvent piégées sur la grille flottante.

La tension de seuil du transistor est alors modifiée par la charge présente sur la grille.

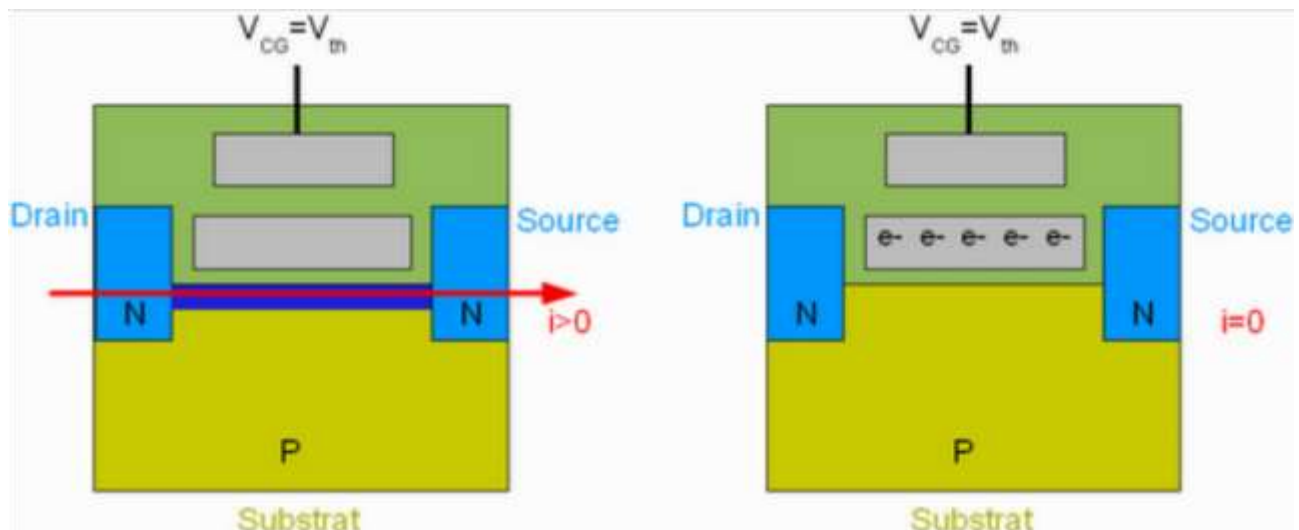


1/ écriture

L'application d'une tension importante sur la grille de contrôle, et d'une tension positive sur le drain, va "pincer" le canal et forcer l'injection d'électrons dans la grille flottante.

2/ pour vider la grille flottante

L'application d'une tension négative importante sur la grille de contrôle va fortement repousser les électrons, qui passeront par effet tunnel et seront évacués.



Pour lire la structure MOS, on applique une tension à la grille de contrôle, ce qui permet de faire circuler le courant dans le canal cela équivaut à lire un 1.

En revanche, si des électrons sont piégés dans la grille flottante, le comportement est différent.

Lors de la lecture, la même tension est appliquée à la grille de contrôle, mais cette fois cette tension est amoindrie, à cause des électrons présents dans la grille flottante.

Du coup, le canal est aminci, : la structure est isolante, ce qui équivaut à lire un 0.

L'effet tunnel pour les nuls.

Imaginez une balle que vous lanciez contre un mur. Soit elle est lancée assez fort, et elle passe au dessus du mur, soit elle n'est pas lancée assez fort, et elle rebondit.

La même chose existe, pour un électron essayant de sortir du métal qui le contient. Si on le lance assez fort, il franchit la barrière et retombe de l'autre côté (autrement dit, si on lui impose un champ électrique assez fort, il est capable de sortir du métal pour traverser le vide jusqu'à un autre métal ou matériau conducteur).

Mais là où une grosse différence intervient, c'est si vous ne lancez pas assez fort votre électron.

A la différence d'une balle, un électron est une sorte de nuage. Eh bien une partie de ce nuage peut passer le mur tandis que l'autre va rebondir. C'est la différence avec la balle. Confronté à une barrière, un électron a donc la possibilité de se scinder en deux : une partie franchit la barrière, et l'autre non.

Mais un tel état ne dure pas : un électron ne reste pas longtemps scindé, parce que les deux parties de l'électron interagissent avec le matériau dans lequel elles se trouvent. Et il se passe alors que l'une des deux parties disparaît, tandis que l'autre grossit : l'électron se retrouve alors entier d'un côté ou de l'autre. Il peut être passé ou pas, selon la partie qui grossit : en gros, la partie restée en arrière du mur a la possibilité d'être "téléportée" avec l'autre. Comme si il y avait eu un tunnel dans le mur par lequel elle serait passée. Alors qu'elle n'est passée nulle part !

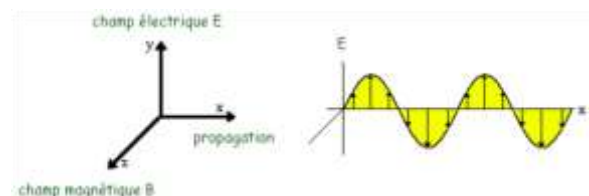
Si on lance des électrons contre une barrière, plus la barrière est petite, plus les électrons ont de chance de passer, par effet tunnel.

7) Écrans

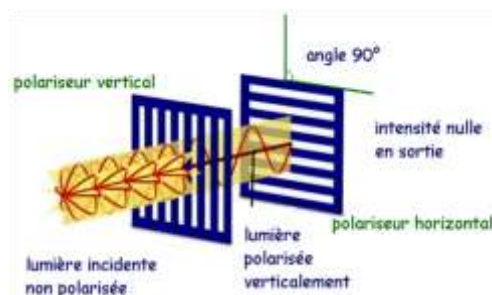
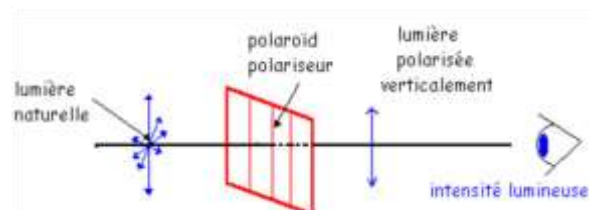
LCD et Plasma

Lumière naturelle et lumière polarisée La lumière est une onde électromagnétique composée d'un champ électrique E et d'un champ magnétique B . L'œil n'est sensible qu'à la composante électrique.

On distingue la lumière naturelle ou non polarisée, pour laquelle il n'y a pas de direction particulière. Et la lumière polarisée, pour laquelle tous les champs E sont parallèles.



La lumière peut être polarisée soit par un filtre appelé polaroïd (matière plastique formée de longues molécules parallèles) ou par simple réflexion sur un miroir.



Si on place un deuxième polaroïd sur le trajet de la lumière polarisée, l'intensité en sortie sera nulle quand ils sont croisés : c'est la loi de Malus.

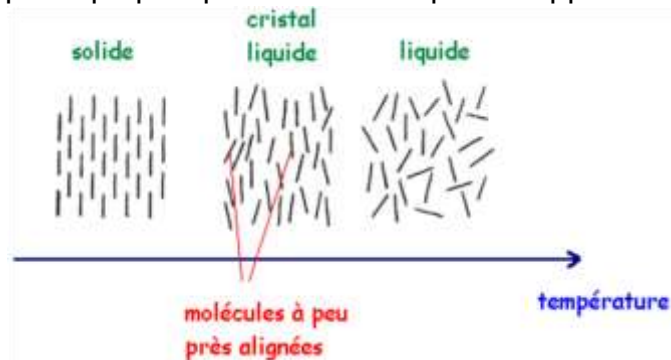
D'une façon générale, un afficheur LCD sera constitué :

- de deux polariseurs croisés qui ne laissent pas passer la lumière

- d'une substance placée entre ces deux polariseurs qui fait tourner la direction de polarisation et permet à la lumière de passer
- d'un dispositif de commande qui permettra ou non la rotation de polarisation

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 :

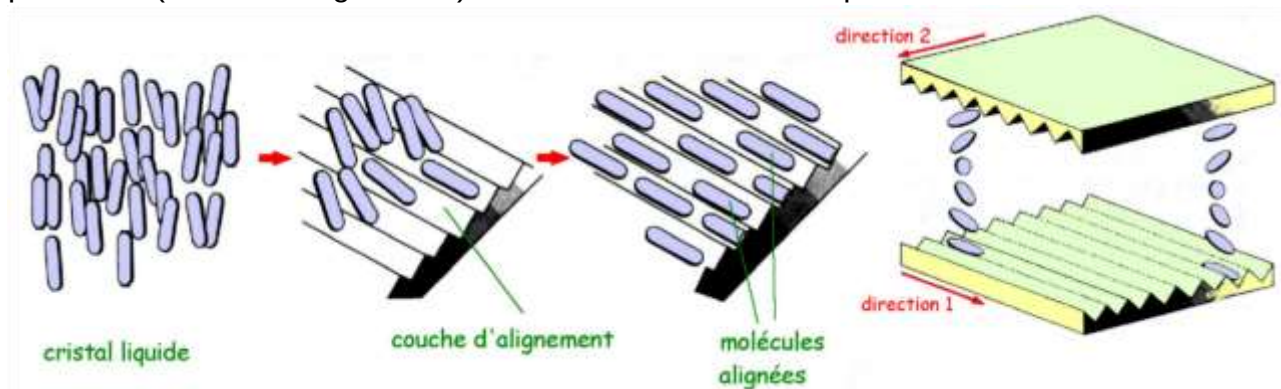
- la plupart des cristaux liquides sont des composés organiques avec des molécules allongées en forme de bâtonnets.
- ils ne présentent ni les arêtes vives, ni les facettes lisses et brillantes qui caractérisent ce qu'on appelle communément les cristaux.
- en revanche, ils possèdent des propriétés d'organisation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



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

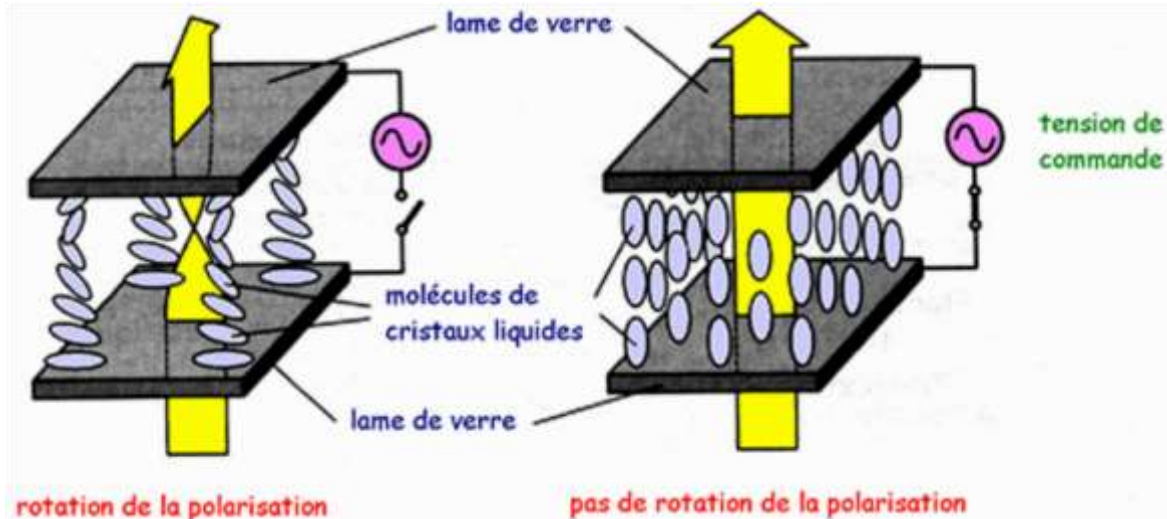


Lorsqu'on enferme une couche de cristaux liquides entre deux plaques gravées de sillons orientés dans deux directions différentes, l'orientation des molécules (à l'état de repos) passe progressivement de la direction (1) à la direction (2). Elle fait donc apparaître une torsion. C'est ce qu'on appelle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

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.

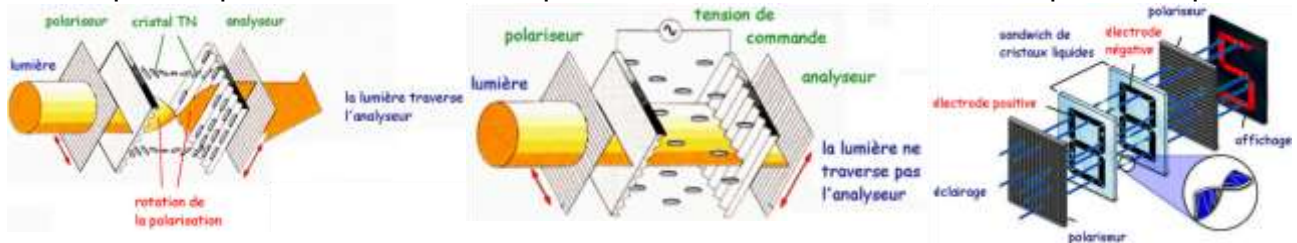


en jouant sur la valeur de la tension appliquée, il est possible d'obtenir des situations intermédiaires dans lesquelles seule une partie des rayons incidents voit sa direction de polarisation modifiée.

L'afficheur à deux couleurs

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



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

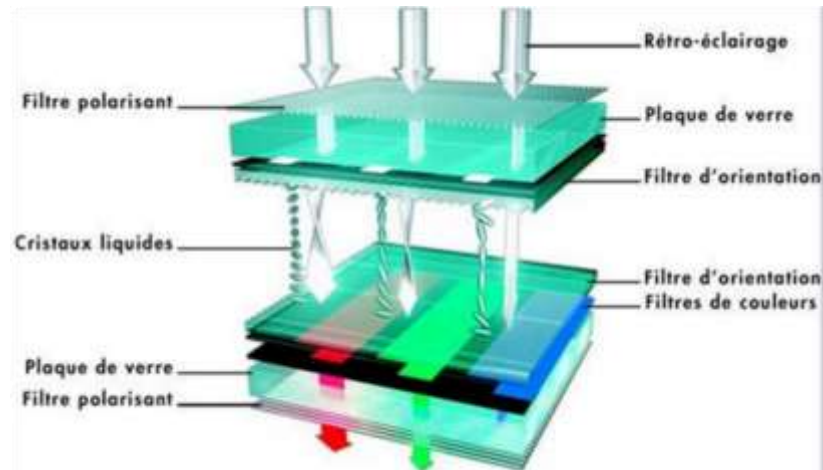
Les sources d'éclairage :

Les deux configurations des écrans LCD sont les écrans réfléchissants et ceux utilisant un rétro-éclairage:

- les écrans réfléchissants utilisent la lumière ambiante. Au repos, la lumière traverse le cristal avant de subir une réflexion dans un 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.

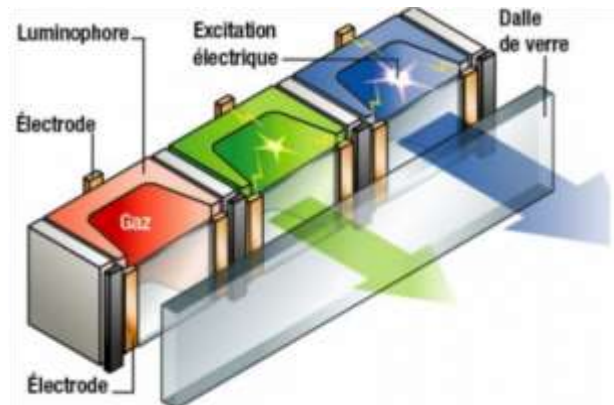
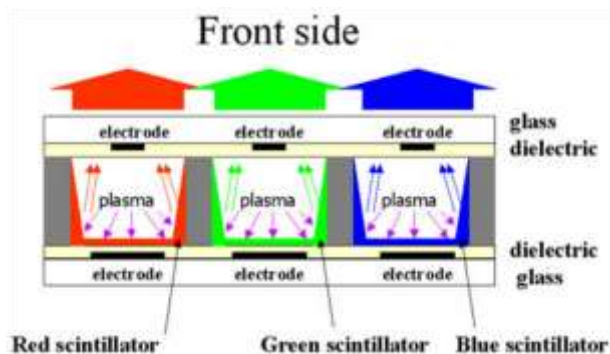


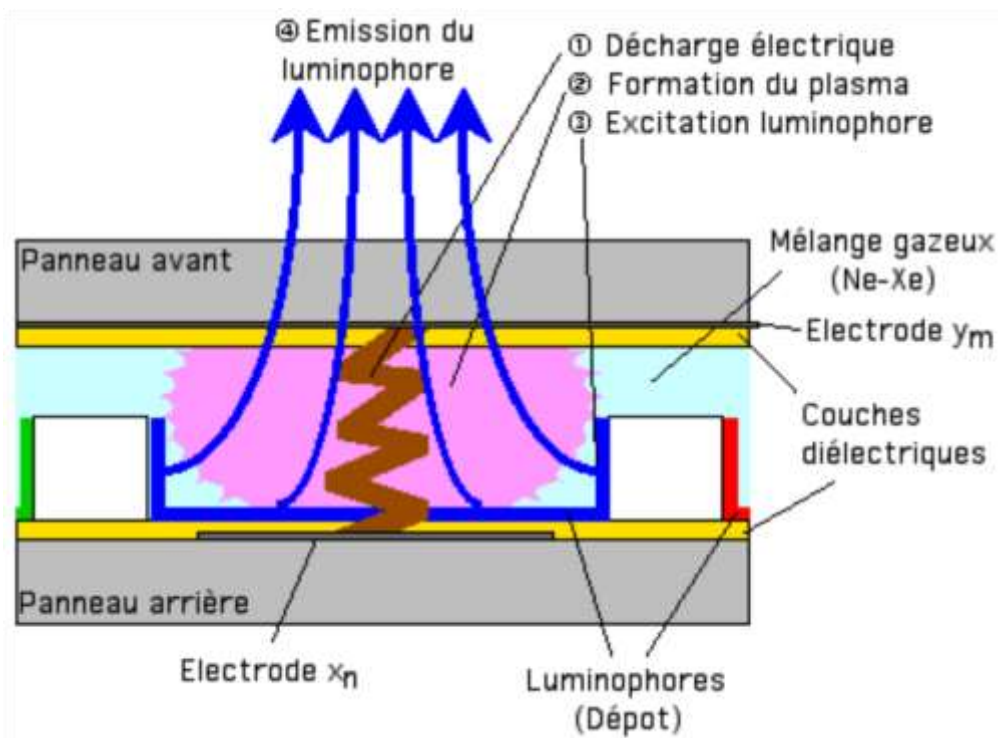
Les écrans à plasma

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 constitue 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, 3 par pixel, pas besoin de rétro éclairage.





Les écrans OLED

La technologie OLED est basée sur l'utilisation de diodes superposées qui, une fois alimentées par un courant électrique, émettent leur propre lumière, contrairement à d'autres types d'affichage tels que les écrans à cristaux liquides qui nécessitent un rétroéclairage.

Chaque pixel de l'écran est composé de trois diodes juxtaposées (bleue, rouge et verte) dont l'épaisseur ne dépasse pas un millimètre. chacune de ces diodes est constituée d'un semi-conducteur organique à base de carbone, oxygène, hydrogène et azote entouré par une cathode métallique (charge électrique positive) et une anode transparente (charge électrique négative).

L'ensemble repose sur un support en verre ou en plastique (notamment pour les applications des écrans souples).

