

CHAPITRE III

MATERIAUX SEMI-CONDUCTEURS

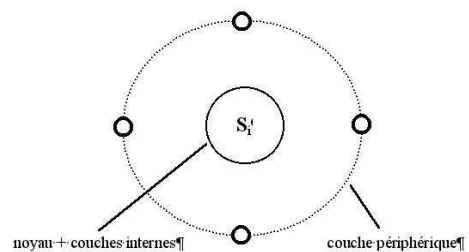
Les matériaux semi-conducteurs sont utilisés dans les dispositifs électroniques servant dans les circuits électroniques et dans la branche de l'électrotechnique connue sous le nom d'électronique de puissance. Ils sont aussi utilisés dans les cellules photovoltaïques servant à convertir l'énergie solaire en électricité. Les matériaux les plus utilisés sont le Silicium (silicon Si), le Germanium (Ge) et des composés tel l'arseniure de Galium (Galium Arsenide GaAs)

I –SEMI-CONDUCTEUR PUR

Il est facile de simplifier au maximum la représentation de Niels Bohr de l'atome en ne laissant apparaître que le noyau avec les couches électroniques internes et la dernière couche électronique, appelée couche périphérique ou couche de valence.

Un matériau semi-conducteur (silicium ou germanium) a la particularité de posséder 4 électrons périphériques, soit exactement la moitié d'une couche complètement saturée.

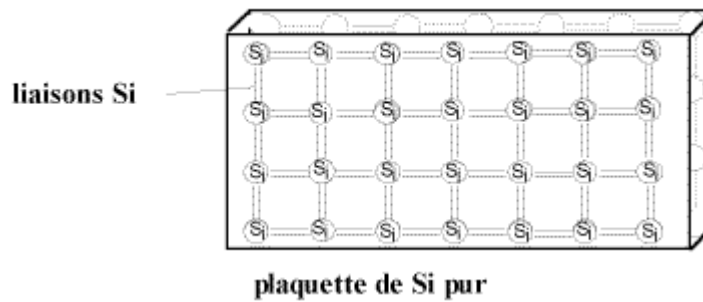
Cette particularité va lui donner un comportement particulier en ce qui concerne les phénomènes électriques, entre autres.



Le matériau semi-conducteur le plus utilisé actuellement est le SILICIUM.

Toutefois, pour utiliser du silicium dans les applications électriques, il faut obtenir des plaquettes d'une pureté extraordinaire. La pureté est de l'ordre de un atome impur pour un million d'atomes de silicium. Le tout reste totalement stable si la température est très basse.

Pour illustrer non seulement un seul atome de silicium mais une plaquette entière, nous simplifions la représentation en ne faisant apparaître que les noyaux avec les couches atomiques intérieures par les cercles comme ci-dessus et avec des traits doubles pour illustrer les électrons périphériques entre chaque atome.

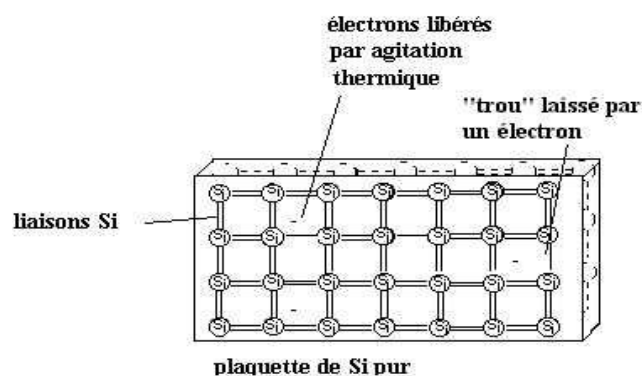


De plus, les atomes du silicium purifié s'organisent entre eux de manière très régulière, suite aux traitements subis, ce qui nous amène à parler d'un cristal semi-conducteur, ou d'une structure cristalline. Cette organisation donne des propriétés électriques particulières au silicium électronique.

Grâce à l'organisation cristalline, chaque atome est entouré de quatre atomes voisins qui vont combiner ensemble leurs électrons de valence de fait que chaque atome se trouve entouré de huit électrons périphériques. Ce qui donne la propriété d'un isolant parfait:

A TRES BASSE TEMPERATURE, AU VOISINAGE DU ZERO ABSOLU (0 KELVIN) LE SILICIUM PUR EST UN ISOLANT PARFAIT.

Dès que la température augmente, l'agitation des atomes entre eux va bousculer cet ordre établi et des électrons périphériques peuvent se retrouver arrachés à la liaison cristalline des atomes. Ces électrons se retrouvent à une distance des noyaux qui leur permet de se déplacer dans la plaquette de silicium.



Les électrons ainsi libérés ont chacun rompu une liaison cristalline du silicium. Ils ont donc laissé derrière eux un emplacement vide, nous parlons d'un "trou". Ces électrons vont se déplacer librement dans la plaquette jusqu'au moment où ils rencontrent un "trou" et se fixer à nouveau dans le réseau.

Ce déplacement aléatoire d'électrons (dans n'importe quel sens) correspond à un courant électrique aléatoire qui représente ce que nous appelons du souffle électronique. Toutefois, ce courant est très faible et nous parlons de conduction intrinsèque.

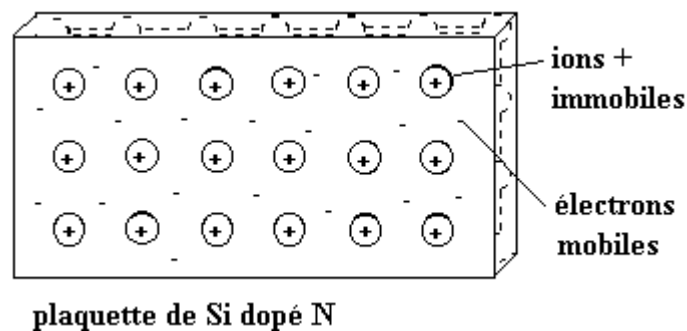
Cette conduction intrinsèque est pratiquement non mesurable et est souvent indésirable, les courants résultants sont de l'ordre du nanoampère et appelés courants de fuites. Même non mesurable, ils existent néanmoins et deviennent trop importants si la température n'est pas contrôlée.

UN SEMI-CONDUCTEUR EST DONC TRES SENSIBLE A LA TEMPERATURE ET NECESSITERA DES MOYENS EXTERNES DE STABILISATION. SANS QUOI UN EMBALLEMENT THERMIQUE ENTRAÎNE TRES VITE LA DESTRUCTION DU SEMI-CONDUCTEUR.

II –DOPAGE D’UN SEMI-CONDUCTEUR

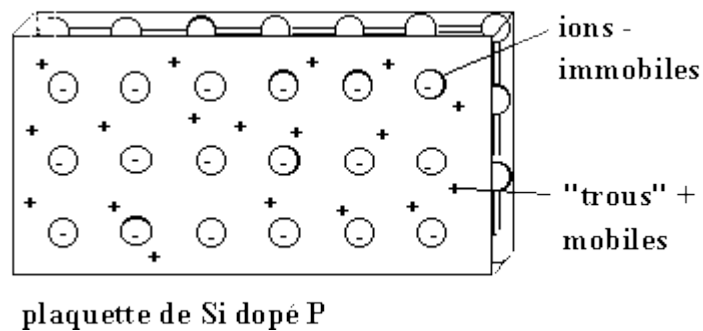
Afin d'améliorer la conduction d'un semi-conducteur, les fabricants injectent dans une plaquette semi-conductrice des matériaux étrangers, ou impuretés, qui possèdent un nombre d'électrons périphériques juste inférieur ou juste supérieur aux 4 électrons du semi-conducteur.

Le dopage N consiste à ajouter au semi-conducteur des atomes possédants 5 électrons périphériques. Quatre de ces électrons vont participer à la structure cristalline, et un électron supplémentaire va se retrouver libre et pouvoir se déplacer dans le cristal. Nous parlons de porteurs de charges mobiles.



Les ions + sont fixes car ils font partie de la structure atomique cristalline de la plaquette de silicium. Rappelons que les ions comprennent le noyau des atomes et qu'ils sont gros, lourds et solides par rapports aux porteurs de charges mobiles. Un électron est environ 2000x plus petit qu'un seul proton.

Le dopage P consiste à ajouter au semi-conducteur des atomes possédants 3 électrons périphériques. Ces trois électrons participent à la structure cristalline, mais un "trou" est créé par chaque atome étranger puisqu'il lui manque un électron périphérique.

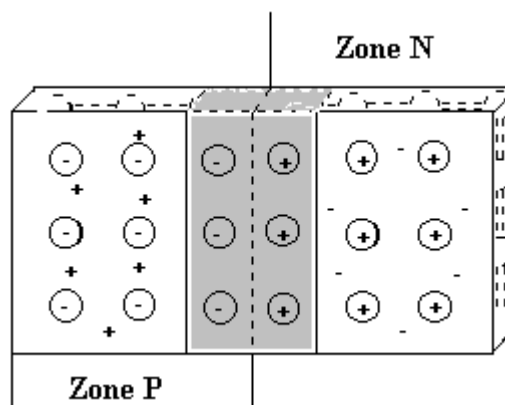


Les "porteurs de charges électriques" mobiles sont responsables de la conduction d'une plaquette de silicium dopée. Si la proportion de dopage est de l'ordre de dix atomes de dopant P pour 100 atomes de silicium, la conductibilité du semi-conducteur est améliorée dans la même proportion, soit de 10%.

Il est donc possible de "régler" la conduction d'un semi-conducteur en choisissant la quantité de dopage. A l'intérieur d'un circuit intégré, il est aisé d'imaginer des zones plus ou moins dopées de manière à obtenir des résistances électriques.

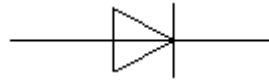
III – LA JONCTION P – N (Effet Diode)

Il est aisé de s'imaginer ce qui se passe lorsque deux zones de dopage P et N sont réalisées sur une même plaquette. A la jonction (réunion) des deux zones, les porteurs de charges mobiles se combinent et il apparaît une zone d'espace vide de porteurs de charges mobiles, donc une zone isolante.



En observant attentivement la polarité résultante des charges électriques de la zone isolante, on s'aperçoit qu'elle est inverse à la polarité de la région de dopage: charge positive dans la zone N et charge négative dans la zone P.

Cela signifie qu'une jonction PN non alimentée est à l'image d'un condensateur, à savoir deux zones conductrices séparées par une zone isolante. Le symbole de la diode représente directement ces deux zones:



L'anode représente la zone P

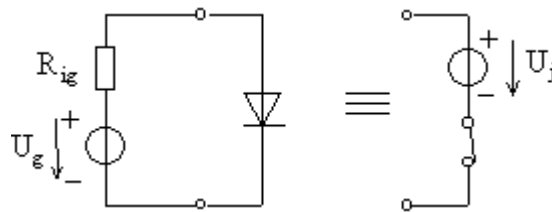
zone N

La cathode représente la

IV – MECANISME DE CONDUCTION D'UNE DIODE

Lorsque l'on alimente une diode, donc une jonction PN, l'effet change selon la polarité de la tension appliquée. Une diode ne laisse passer le courant que dans un seul sens.

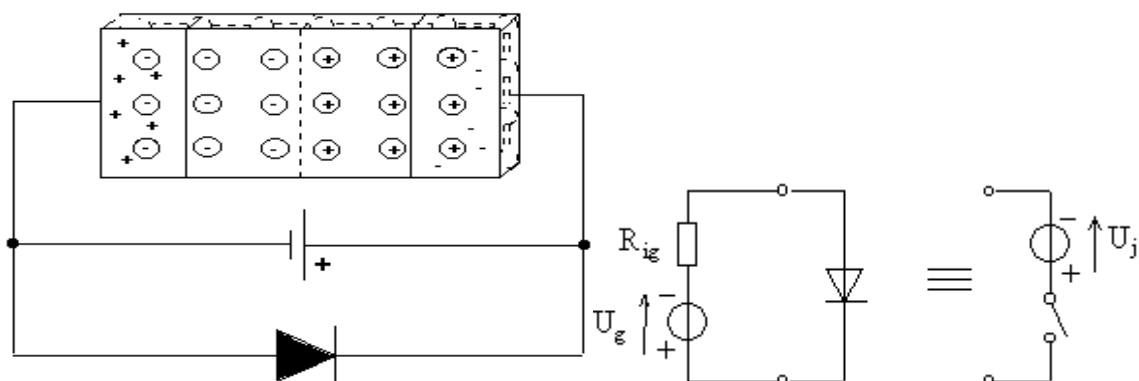
Polarisation directe



$$U_j = 0,6 - 0,7V$$

Nous constatons la circulation du courant électrique. La jonction est conductrice en présentant une différence de potentiel de 0,6 - 0,7 volts à ses bornes. La diode conduit et nous pouvons idéaliser ce fonctionnement en remplaçant la diode par un générateur DC et un interrupteur fermé. Nous parlons de: polarisation dans le sens passant, ou sens direct; courant direct; en anglais "forward".

Polarisation inverse



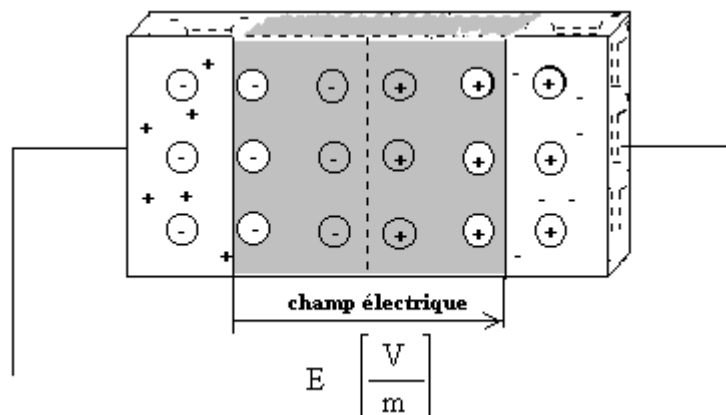
$$U_j = U_g$$

Les porteurs de charges mobiles sont attirés vers les connexions extérieures par la présence des charges électriques de l'alimentation. Nous constatons l'élargissement de la zone vide de porteurs de charges. La jonction est isolante. La diode est bloquée et nous pouvons représenter cet état par un interrupteur ouvert. Nous parlons ici de: polarisation dans le sens bloquant, ou dans le sens inverse; en anglais "reverse"

V – EFFET ZENER

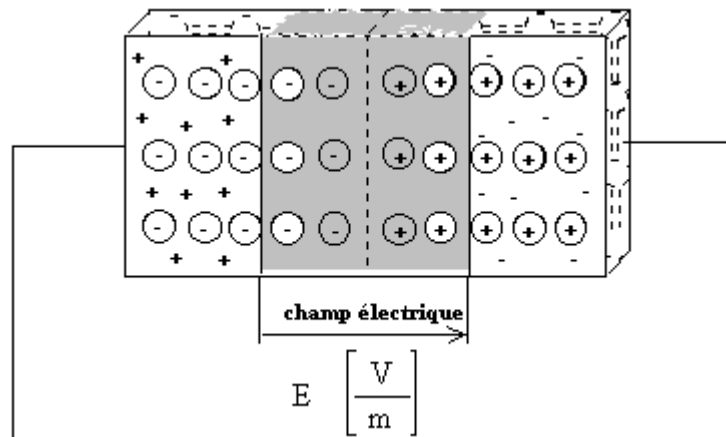
Nous savons qu'une jonction P-N ne laisse passer le courant que dans un seul sens, en polarisation directe. Par contre, en polarisation inverse, il ne circule pratiquement aucun courant tant que la tension inverse ne dépasse pas une valeur limite, limite appelée tension de claquage $U_{INV \text{ MAX}}$.

La tension inverse donne naissance à un champ électrique à l'intérieur de la plaquette de Silicium, et plus précisément dans la zone isolante de la jonction.



La quantité faible des éléments de dopage fait que la jonction PN d'une diode conventionnelle supporte des champs électriques intenses correspondant à des tensions inverses allant de plusieurs centaines de volts à plusieurs milliers de volts. Si cette tension $U_{INV \text{ MAX}}$ est dépassée, l'emballement thermique détruit la diode dans la plupart des cas.

Par contre, les diodes Zener ont été spécialement conçues et fabriquées de manière à pouvoir être utilisées en polarisation inverse, dans la zone de claquage, notamment en modifiant les dimensions de la jonction et surtout la quantité de dopage des zones P et N.

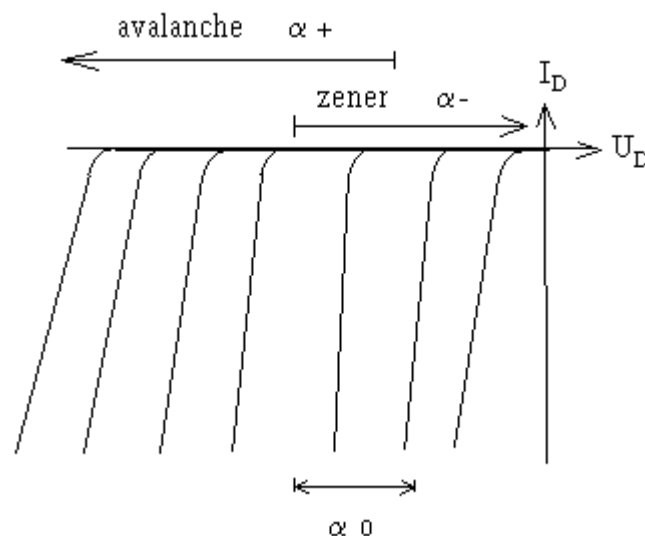


Plus fortement dopée que les diodes conventionnelles, un champ électrique relativement faible devient déjà suffisamment intense pour que les liaisons de covalence s'affaiblissent et se rompent. Les porteurs de charges (des éléments de dopage) ainsi libérés sont assez nombreux pour que le courant augmente brutalement et pour que la tension aux bornes de la diode ne varie pratiquement plus. C'est ce qui est appelé l'effet Zener.

Pour d'autres diodes Zener, il est possible que sous l'action du champ électrique interne, les porteurs de charges minoritaires (du silicium) de la zone isolante acquièrent une énergie telle qu'il puisse y avoir ionisation par choc, et, par effet d'avalanche, le courant croît extrêmement vite. La tension aux bornes de la diode ne varie pratiquement pas non plus. C'est ce qui est appelé effet d'avalanche.

En pratique, pour les diodes dont la tension zener dépasse 10V, seul l'effet d'avalanche est possible. Ce qui a pour conséquence que la caractéristique de la diode est moins franche (la pente est plus grande), et le coefficient de température est positif.

Les diodes dont la tension Zener est inférieure à 5V ont une jonction très mince et seul l'effet Zener peut avoir lieu, ce qui entraîne que la caractéristique de la diode est très raide et, de plus, ces diodes ont un coefficient de température négatif.



Entre 5V et 10V, les deux effets peuvent se combiner, et la caractéristique est la plus raide ainsi que le coefficient de température qui peut être proche de zéro. Ce qui signifie que les diodes Zener prévues pour un fonctionnement inverse compris entre 5V et 10V seront utilisées pour un fonctionnement très stable.

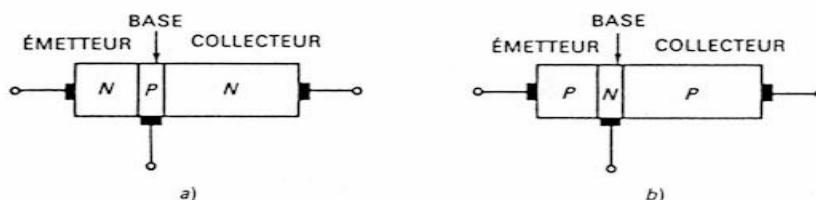
V – EFFET TRANSISTOR

Avant 1950, tout équipement électronique comportait des tubes à vide. Ces "lampes" lumineuses, dont le filament à lui seul consommait déjà quelques watts, nécessitaient une alimentation volumineuse et un dispositif de circulation d'air afin de dissiper la chaleur émise.

Vers 1951, un dispositif à semi-conducteur à deux jonctions appelé transistor (de l'expression anglaise TRANSfer resISTOR) a été inventé par Bardeen, Brattain et Schockley. L'impact de cette découverte sur l'électronique fut considérable et donna naissance à toute une industrie ainsi qu'au développement de composants électroniques allant jusqu'au microprocesseur actuel présent dans tout micro-ordinateur.

Un transistor est réalisé à partir d'une plaquette semi-conductrice, germanium ou silicium, dans laquelle trois zones de dopage P ou N, telles que nous l'avons déjà décrit pour les diodes, sont créées, ce qui représente la réalisation de deux jonctions P-N. Il est donc possible d'obtenir des transistors PNP ou NPN qui sont complémentaires. Les courants et les tensions d'un transistor PNP sont opposés aux courants et tensions d'un transistor NPN, mais le principe de fonctionnement est exactement le même pour les deux types.

Vu à l'aide d'un microscope, l'intérieur d'un transistor peut se visualiser (en simplifié, bien sûr) selon le dessin ci-dessous.

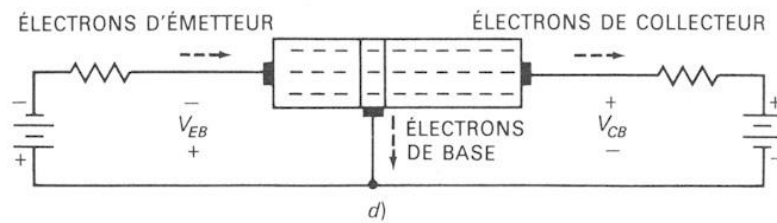


Les trois régions des transistors. a) Transistor NPN. b) Transistor PNP.

Il est utile de savoir que la zone d'émetteur est fortement dopée, la zone de base est très étroite et faiblement dopée, alors que la zone collecteur, dopée moyennement, est par contre la plus large des trois zones et dissipe le plus de chaleur.

L'effet transistor peut se résumer (et se simplifier) de la manière suivante:

Les électrons venant de l'émetteur (cas du NPN) arrivent dans la région de la base. Deux trajets sont offerts, l'un vers la zone de la base, et l'autre à travers la jonction collecteur et ensuite dans la région du collecteur.



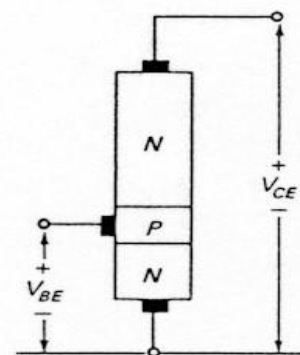
La base étant très mince et remplie d'électrons provenant de l'alimentation de l'émetteur, ceux-ci provoquent un champ électrique sur la jonction collecteur, renforcé par la polarisation inverse de l'alimentation de collecteur, qui entraîne une conduction par effet d'avalanche, à l'image d'une diode Zener.

L'effet transistor n'existe donc que si la diode émetteur est polarisée dans le sens direct et la diode collecteur dans le sens inverse.

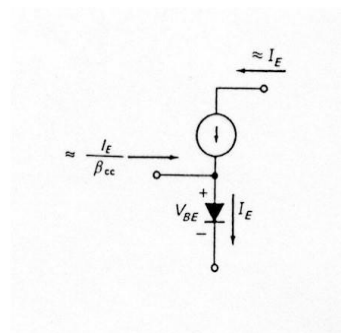
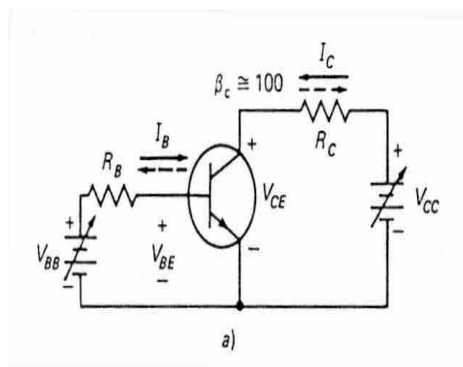
En résumé, pour que l'effet transistor existe, il faut :

- que la diode collecteur soit polarisée en inverse
- que la diode émetteur soit polarisée dans le sens direct

et dans ce cas



- le transistor devient une source de courant commandée : le faible courant de base commande un fort courant d'émetteur (et donc de collecteur)

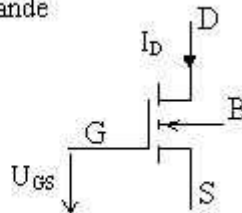


Le transistor peut se représenter à l'aide d'un schéma équivalent composé d'une source de courant (commandée) et d'une diode.

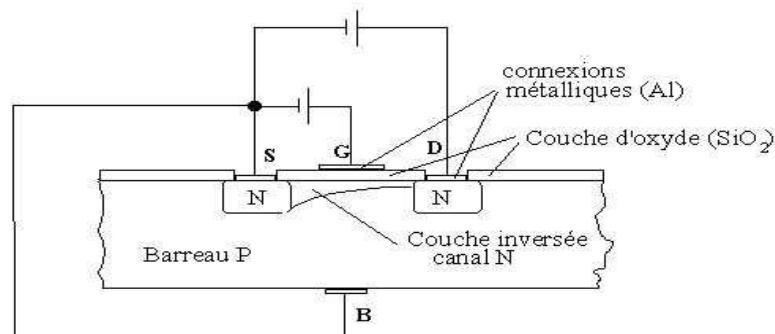
V –EFFET DE CHAMP

Dans l'effet Zener, nous avons remarqué que les porteurs de charges mobiles de l'élément de dopage pouvaient se déplacer sous l'influence d'un champ électrique. Par exemple, une jonction PN se vide lors d'une polarisation inverse, ce qui amène l'existence d'un champ à l'intérieur même de la jonction.

La tension U_{GS} commande le courant de drain I_D



Ce phénomène est utilisé par les transistors à effet de champ que nous pouvons résumer de la manière suivante, un canal conducteur (P ou N) est rendu plus ou moins conducteur grâce à une tension de commande qui le vide tout ou en partie de ses porteurs de charges mobiles.



Sur un barreau de silicium P, deux zones N sont diffusées pour former le drain et la source. Le barreau P forme également un condensateur avec la grille dont le diélectrique est la couche d'oxyde.

Lorsque la grille est rendue positive par rapport à la source, les électrons du barreau sont attirés dans la zone située entre le drain et la source. Par cet artifice, un canal de type N est créé entre la source et le drain. Si une tension est appliquée entre le drain et la source, un courant de drain I_D circulera.

En variant la tension de commande U_{GS} , la densité des électrons dans le canal change. Ce qui signifie que le courant de drain varie ou que la résistance de passage du drain est modifiée, ce qui revient au même.

L'avantage de cette commande, dite en tension, est qu'il n'y a aucun courant de grille, donc aucun courant de commande, ce qui diminue d'autant la puissance de commande nécessaire.

De plus, comme il n'y a pas d'électrons à déplacer dans la grille, la rapidité de commande est accrue ce qui permet à ces transistors de "travailler" plus haut en fréquence.

En technique audio, ils sont parfois préféré aux transistors bipolaires (PNP ou NPN) car le résultat de cette commande en tension est d'obtenir un son plus "chaud", à l'image des tubes électroniques sous vide