

COMPOSANTS NON LINÉAIRES À SEMI-CONDUCTEURS

PREMIÈRE ANNÉE DE L'ESPCI

JÉRÔME LUCAS

6 décembre 2018

Ce polycopié du cours d'électronique de l'ESPCI réalisé avec $\text{\LaTeX} 2_{\epsilon}$ est mis à votre disposition sous la licence ouverte conçue par Etalab.



LICENCE OUVERTE

OPEN LICENCE

www.etalab.gouv.fr/wp-content/uploads/2014/05/Licence_Ouverte.pdf

www.etalab.gouv.fr/wp-content/uploads/2014/05/Open_Licence.pdf

Table des matières

1	Semi-conducteurs	6
1.1	Bases élémentaires de la théorie des bandes	6
1.2	Principaux semi-conducteurs utilisés en électronique	6
1.2.1	Exemple du silicium.	6
1.2.2	Conduction par trous	7
1.2.3	Recombinaisons électrons-trous	8
1.3	Dopages	8
2	La jonction PN	8
2.1	Écrantage	10
2.2	La jonction PN comme composant électronique : La diode	10
2.3	Caractéristique statique courant/tension de la diode	11
3	Utilisation des diodes	12
3.1	Stratégies de calcul	14
3.2	Droite de charge	15
3.3	Montages de base	15
3.3.1	Redressement simple alternance	15
3.3.2	Redressement double alternance	16
3.3.3	Écrêtage (<i>clamping</i>)	16
3.3.4	Détection de crêtes	17
3.3.5	Logique à diodes	17
3.3.6	Diode de roue libre	17
3.3.7	Pompe à diode	18
4	Autres types de diodes	22
4.1	Diodes Zener	22
4.2	Diodes Électroluminescentes (LED)	23
4.3	Varicaps	24
4.4	Diodes PIN	24
4.5	Diodes Schottky	25
4.6	Diodes GUN (Hyperfréquences)	25
5	Interaction avec la lumière	26
5.1	Photopiles	26
5.2	Photodiodes	27
6	Critères de choix d'une diode.	30
7	Transistors	30
7.1	Qu'est ce qu'un transistor ?	30
7.2	Utilisation	30
8	Transistors à effet de champ	30
8.1	MOS canal N (NMOS)	31
8.2	MOS canal P (PMOS)	32
8.3	Caractéristiques statiques des transistors NMOS, modèle de Schichman et Hodges	34
8.3.1	Courant de grille	34
8.3.2	Courant drain-source	34
8.4	Validité du modèle	37
8.4.1	Paramètres d'influences	37
8.5	Stratégies de calcul : schémas équivalents	39

8.6	Caractéristique statique des transistors PMOS.	39
8.6.1	Mises en œuvre comparées des transistors NMOS et PMOS.	40
8.7	Transistors MOS à apauvrissement	41
9	Applications des transistors MOS	41
9.1	MOS en Commutation	41
9.1.1	NMOS en pull-down	42
9.1.2	PMOS en pull-up	42
9.1.3	Exemples d'application d'un NMOS en pull-down. : Pilotage d'une LED ou d'un relais par une porte logique.	42
9.2	Interrupteurs analogiques (<i>Analog switches</i>)	43
9.3	Portes logiques CMOS	46
9.3.1	L'inverseur MOS	46
9.3.2	Quelques autres portes CMOS	47
9.4	MOS en Amplification	48
9.4.1	Petit signal	49
9.4.2	Schéma équivalent petit signal	49
9.4.3	Méthodologie d'analyse du fonctionnement en petit signal d'un amplificateur à MOS.	51
10	Transistors bipolaires à jonction	53
10.1	Fonctionnement	53
10.2	Caractéristiques	54
10.3	Paramètres d'influences	56
10.3.1	Effet de la Température	56
10.3.2	Variation de β avec la fréquence	56
11	Applications des transistors à jonction	57
11.1	BJT en Commutation	57
11.2	BJT en Amplification	58
11.2.1	Schéma équivalent petit signal	58
11.2.2	Les trois montages de base des transistors bipolaires	59
11.2.3	Polarisation et emballement Thermique	61
11.2.4	Mise en œuvre d'un montage à émetteur commun, montage Cascode	62

1 Semi-conducteurs

De façon simplifiée, un conducteur est un matériau qui conduit l'électricité grâce à la présence de charges ou porteurs de charges mobiles en son sein. Dans les solides ces charges sont en général des électrons. Les isolants quant à eux sont des matériaux dans lesquels les charges ou porteurs de charges sont très rares ou inexistantes.

1.1 Bases élémentaires de la théorie des bandes

Dans un conducteur solide, à cause de l'interaction entre les atomes, il peut apparaître un niveau d'énergie quantifié au dessus du niveau d'énergie des orbitales de chaque atome dans lequel les électrons peuvent se déplacer. C'est la bande de conduction présentée figure 1. Cette bande est séparée de la bande de valence par une bande interdite. La bande de valence résulte de la mise en commun d'électrons périphériques dans les liaisons covalentes. L'existence et la taille de la bande interdite dépend du ou des matériaux constituant le solide, de sa structure (amorphe ou cristallin) et de sa température.

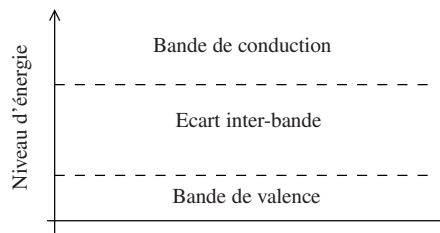


FIGURE 1 – Bandes d'énergie dans un solide

La bande de conduction existe toujours pour les conducteurs. Pour les isolants, cette bande de conduction n'existe pas.

Un semi-conducteur présente la même structure de bande qu'un conducteur, à ceci près qu'à basse température (0°K) la bande de conduction est vide. Elle se peuple avec la température. Un semi-conducteur est donc plus ou moins conducteur en fonction de la température.

1.2 Principaux semi-conducteurs utilisés en électronique

Les semi-conducteurs utilisés en électronique sont des mono-cristaux fabriqués artificiellement grâce à une technique de croissance appelée épitaxie. Ce sont principalement des mono-cristaux de :

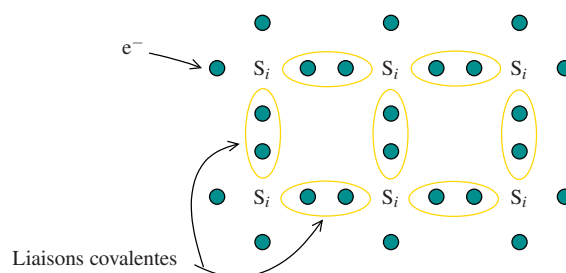
- Silicium, pour la très grande majorité des applications,
- Germanium, semi-conducteur historique dont l'usage a été abandonné,
- AsGa (Arsenure de galium), très utilisé pour les applications en micro ondes (fréquences de l'ordre de 1 à 100 GHz) .

1.2.1 Exemple du silicium.

Numéro atomique $Z = 14$. Les électrons sont donc répartis de la façon suivante :

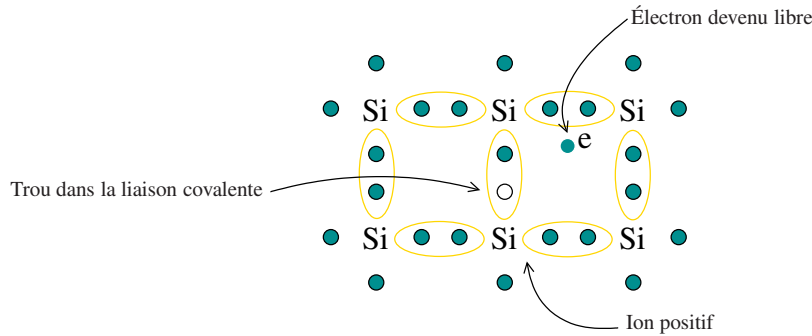
Couche	K	L	M
Nb électrons	2/2	8/8	4/8

Il y a donc 4 électrons dans la bande de valence. La figure suivante présente la structure en terme de liaisons covalentes d'un cristal de silicium à très basse température.



Cette figure montre la répartition des électrons dans la bande de valence. Il n'y a pas, à cette température, d'électron dans la bande de conduction.

Lorsque la température s'élève (température ambiante), certaines des liaisons covalentes sont cassées par l'agitation thermique et certains électrons passent dans la bande de conduction. On a alors, comme représenté dans la figure ci dessous, libération d'un électron. Ce dernier vient peupler la bande de conduction et il y a apparition d'un ion positif de Silicium et d'un trou dans la liaison covalente cassée.



Ce phénomène est proportionnellement nombre d'atomes peu important, Il est de l'ordre de 3 pour 10^{13} liaisons à 300 °K. Il est néanmoins suffisant pour que le silicium devienne conducteur, même si il reste un mauvais conducteur :

Silicium	$2,52 \cdot 10^{-4} \text{ S/m}$
Cuivre	$59,6 \cdot 10^6 \text{ S/m}$

Ce type de semi-conducteur naturel est dit semi-conducteur intrinsèque.

1.2.2 Conduction par trous

Il est très important de noter que le trou créé lors de la rupture de la liaison covalente participe lui aussi à la conduction. Les électrons encore liés aux atomes (les électrons dans la bande de valence) peuvent en effet changer d'atome comme présenté figure 2

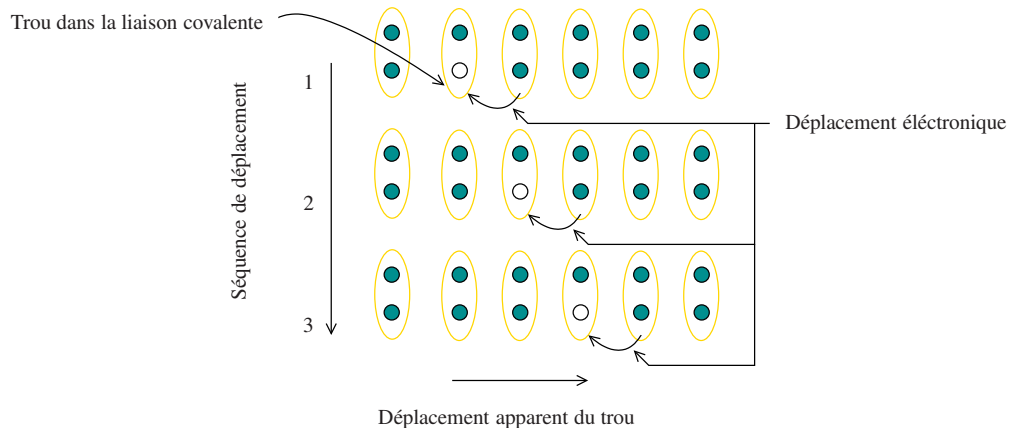


FIGURE 2 – Déplacement d'un trou.

Les électrons de valence sont cependant bien plus liés à leur atome que les électrons de la bande de conduction. En conséquence la mobilité des trous est bien plus faible que celle des électrons de conduction. L'équation 1 permet de définir la mobilité des porteurs de charge.

$$V_{(m/s)} = \mu E \quad (1)$$

Dans cette équation, E est le module du champ électrique vu par le porteur de charge, μ sa mobilité et V sa vitesse d'équilibre, conséquence du champ électrique. On voit ainsi que les mobilités s'expriment par exemple en $\text{m}^2\text{V}^{-1}\text{s}^{-1}$.

Les valeurs des mobilités des trous et des électrons sont les suivantes :

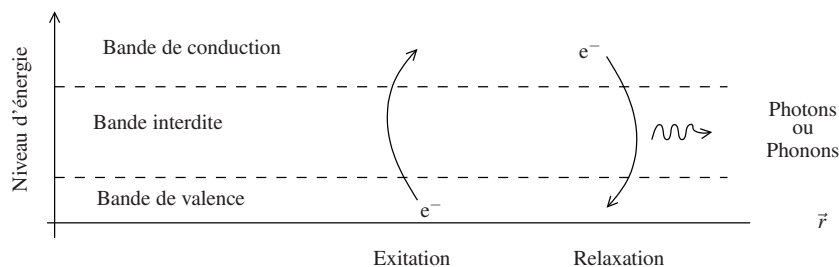
Électrons	$1500 \text{ cm}^2/\text{Vs}$
Trous	$475 \text{ cm}^2/\text{Vs}$

Finalement, on peut retenir que la mobilité des électrons est environ trois fois plus grande que celle des trous :

$$\mu_{e^-} \approx 3 \times \mu_{\text{Trous}} \quad (2)$$

1.2.3 Recombinaisons électrons-trous

Lorsqu'un électron libre rencontre un trou, il a évidemment tendance à reprendre la place libre qui correspond à un niveau d'énergie plus faible. Ce phénomène spontané est un phénomène de relaxation appelé recombinaison électron-trou. La différence d'énergie est convertie selon le semi-conducteur soit en phonon (vibrations du réseau cristallin) soit en photon comme présenté dans la figure ci-dessous. En ce qui concerne le silicium compte tenu des niveaux d'énergie mis en jeu, c'est en phonons. Il n'existe en effet pas de LED (Diodes Électroluminescence) en silicium.



La création des paires électron-trou et la recombinaison sont deux phénomènes concurrents qui s'équilibrent à une valeur qui dépend de la température.

1.3 Dopages

En considérant l'exemple du silicium, on voit qu'une liaison covalente cassée donne lieu à :

- $\nearrow$ un trou
- $\rightarrow$ un électron libre
- $\searrow$ un ion positif

Le dopage consiste à favoriser la création de trous ou d'électrons libres en insérant de façon contrôlée dans la maille cristalline des ions positifs ou négatifs qui constituent des impuretés vis à vis du cristal.

Pour le silicium :

Ions + (Bore, 3 e^- de valence)	$\Rightarrow$	création d'un trou	Dopage p
Ions - (Phosphore, 5 e^- de valence)	$\Rightarrow$	création d'un électron libre	Dopage n

Le taux d'impureté doit être contrôlé très précisément lors de l'épithaxie du cristal car la concentration de porteurs (trous ou électrons) en dépend. On crée ainsi des semi-conducteurs plus ou moins dopés : p, p^+ , p^{++} ou n, n^+ , n^{++} .

Les semi-conducteurs à dopage créés par impureté sont dit extrinsèques.

2 La jonction PN

On crée une jonction dite "métallurgique" entre un barreau de semi-conducteur N (dopé n) et un autre dopé p (semi-conducteur P). Jonction "métallurgique" (analogie avec la soudure ?) signifie qu'il ne s'agit pas d'un simple contact entre deux matériaux ce qui pourrait créer des barrières de potentiels, mais que l'on a un cristal continu dont le dopage change à la jonction.

À la jonction, les électrons majoritaires dans la zone N ont tendance à diffuser vers la zone P, et les trous de la zone P vers les zones N à cause du gradient de concentration de chacune des espèces de porteurs. Cela crée une zone de déplétion

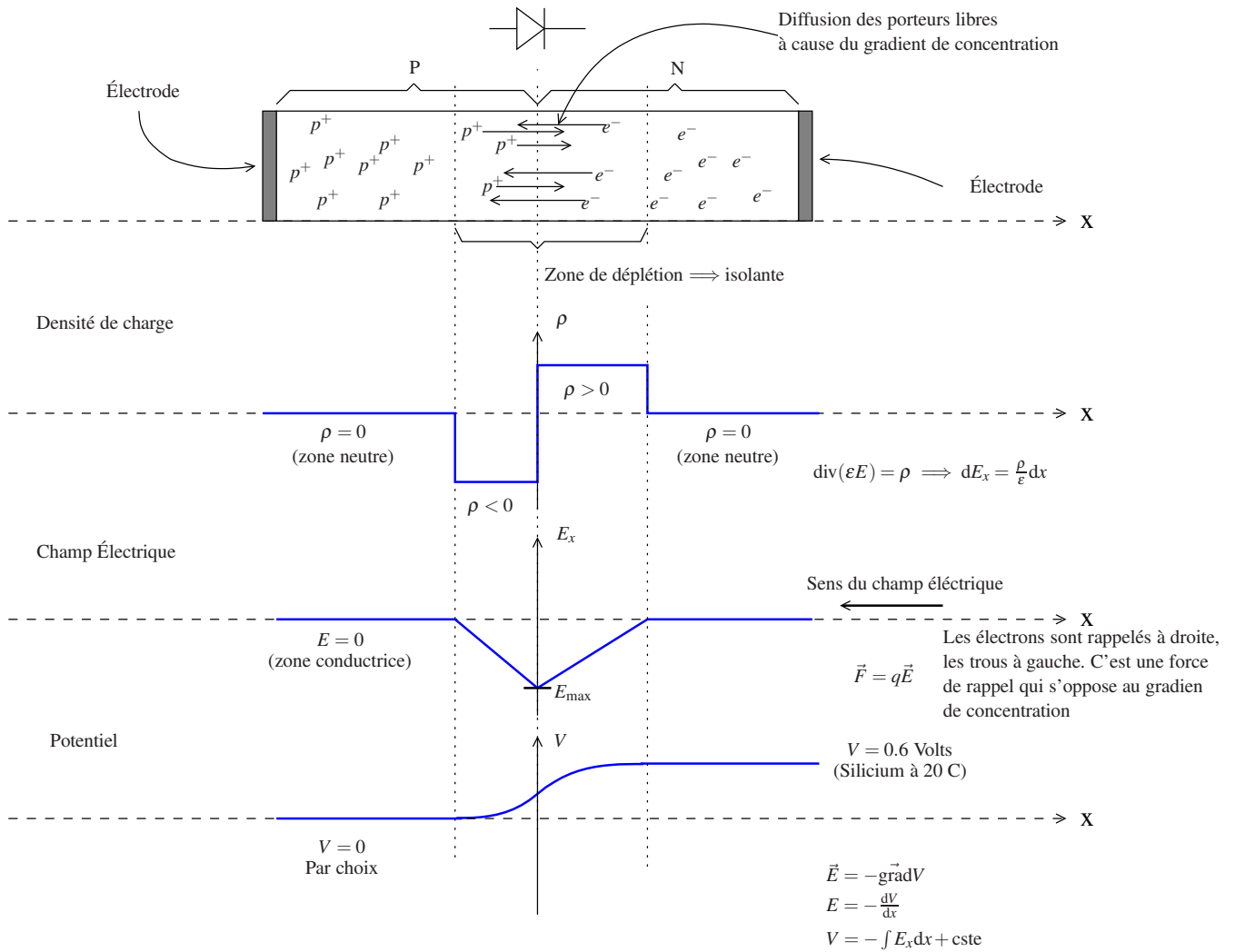


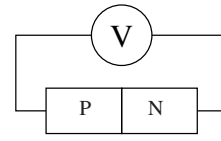
FIGURE 3 – Modèle simplifié de la jonction PN en circuit ouvert.

isolante comme présenté figure 3. On peut remarquer que les mobilités des trous et des électrons n'étant pas les mêmes, la zone de déplétion n'est pas symétrique.

La jonction présente alors un profil de densité de charge selon son axe x tel que celui représenté dans cette même figure. Il est alors possible de calculer le champ électrique résultant. En prenant le potentiel V nul pour l'électrode de la zone P, on en déduit la courbe de potentiel présentée. On voit qu'il apparaît dans la jonction une barrière de potentiel. Cette dernière vaut 0.6 Volts à 20°C pour le silicium avec les valeurs de dopage usuelles.

2.1 Écrantage

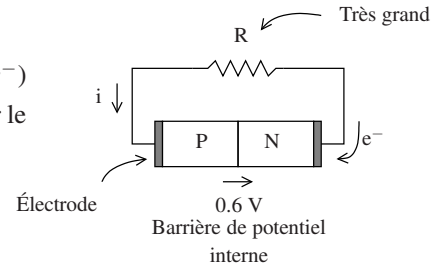
Nous venons de voir que la diode présente une barrière de potentiel interne. Pourtant, si l'on mesure le potentiel aux bornes d'une jonction avec le dispositif de droite par exemple :



On mesure un potentiel nul.

Ceci est dû à l'écrantage du potentiel par les charges libres (p^+ et e^-) présentes dans la jonction PN qui se déplacent pour venir rééquilibrer le potentiel.

En effet, le schéma équivalent du système de mesure est le suivant :



La résistance R est celle du voltmètre, mais c'est aussi celle de l'environnement (par exemple l'air ambiant) dès l'instant de la fabrication. Elle est grande, néanmoins les charges se déplacent tout de même pour venir équilibrer (écranter) le potentiel interne. Lorsque le régime statique est atteint, $i = 0$, et l'on mesure un potentiel nul aux bornes de la jonction.

ATTENTION : Cet écrantage n'annule pas les gradients de concentrations de porteurs internes. La densité de charge globale est nulle, mais les gradients ne le sont pas.

2.2 La jonction PN comme composant électronique : La diode

Le dipôle formé par une jonction PN est une diode. La figure 4 présente son symbole ainsi qu'un moyen mnémotechnique pour se souvenir de son orientation. Le trait de la cathode peut être complété pour former un N qui localise la position du dopage n.

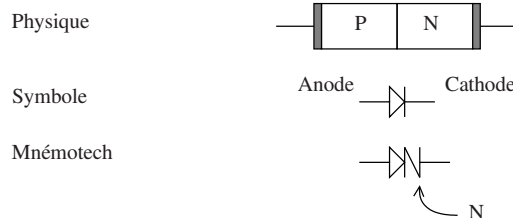


FIGURE 4 – La jonction PN réalise une diode.

À cause de la zone de déplétion à la jonction, la diode ne conduit pas le courant. Pour que des électrons puissent traverser la diode, il leur faut vaincre la barrière de potentiel interne. Le comportement de la diode va être différent selon qu'on la polarise en direct, c'est-à-dire avec une différence de potentiel qui s'oppose à la barrière de potentiel, ou en inverse, c'est-à-dire avec une tension qui renforce cette barrière. Ces deux modes de polarisation reviennent respectivement à diminuer jusqu'à annuler la zone de déplétion ou à agrandir cette même zone.

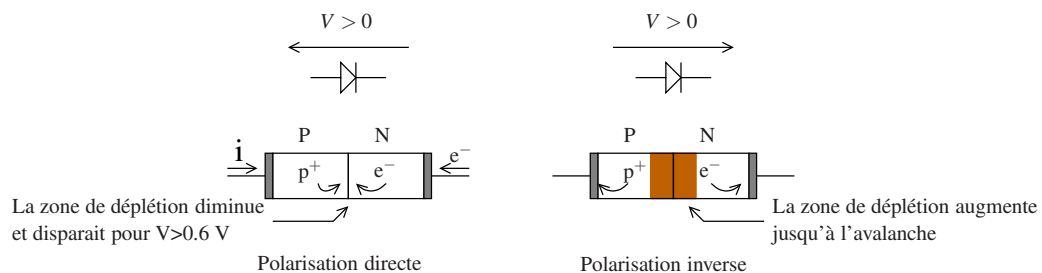


FIGURE 5 – Polarisation directe et indirecte de la diode.

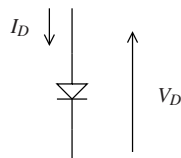
Remarque 1 : Même avec une zone de déplétion augmentée en polarisation inverse, la diode conduit un peu. Le courant vaut typiquement quelques dizaines de nA (25 nA pour la 1N4148). Cette conduction est due aux porteurs minoritaires qui restent présents dans la jonction : les électrons dans la zone P, les trous dans la zone N.

Remarque 2 : Lorsque la tension augmente en polarisation inverse, le champ électrique dans la zone de déplétion augmente. Les électrons minoritaires de la zone P prennent de la vitesse. Lorsque le potentiel aux bornes de la diode atteint une certaine valeur limite V_Z qui dépend du dopage, il y a avalanche. Les électrons franchissant la zone de déplétion ont assez d'énergie pour décrocher certains de ceux engagés dans les liaisons covalentes, qui eux mêmes en décrochent d'autres. La résistance de la jonction s'écroule.

Cette tension s'appelle la tension Zener. Elle est grande devant 0.6 V pour les diodes classiques (≈ 20 V pour la 1N4148). Elle est contrôlée et spécifiée pour les diodes dites Zeners.

2.3 Caractéristique statique courant/tension de la diode

La physique du solide permet, avec les considérations qualitatives précédentes, d'établir la caractéristique analytique de la diode. C'est l'équation 3 dite d'Ebers-Moll :



$$I_D = I_0(e^{\frac{V_D}{V_\phi}} - 1) \quad (3)$$

Dans cette équation I_D est le courant qui traverse la diode, V_D est la tension appliquée aux bornes de la diode, I_0 est le courant résiduel lorsqu'elle est polarisée en inverse, et V_ϕ la tension thermique. Cette équation ne décrit cependant pas le comportement en avalanche de la diode.

V_ϕ la tension thermique est donnée par l'équation 4

$$V_\phi = \frac{kT}{q} \quad (4)$$

Dans cette expression :

$$\begin{cases} k \text{ est la constante de Boltzman} & k = 1,38 \cdot 10^{-23} \text{ J/K} \\ T \text{ est la température en Kelvin} & T_K = T^\circ\text{C} + 273.15 \\ q \text{ est la charge de l'électron} & q = 1,6 \cdot 10^{-19} \text{ C} \end{cases}$$

Pour fixer les idées, le tableau suivant donne les valeurs de la tension thermique pour quelques valeurs de la température :

T °C	V_ϕ mV
20	25,3
25	25,7
30	26,1

La variation est, compte tenu de l'équation 4, linéaire avec la température avec une pente de $\frac{k}{q} = 0.086 \text{ mV/}^\circ\text{K}$.

La caractéristique courant tension pour la diode 1N4148 à 300°K est présentée figure 6.

La caractéristique de la diode dépend de la température au travers du coefficient de tension thermique V_ϕ . Comme V_ϕ augmente avec la température, I_D devrait diminuer. La caractéristique devrait se déplacer vers la droite avec la température. Cependant, le courant inverse I_0 dépend aussi de T comme présenté figure 5, où A est un coefficient qui dépend de la diode

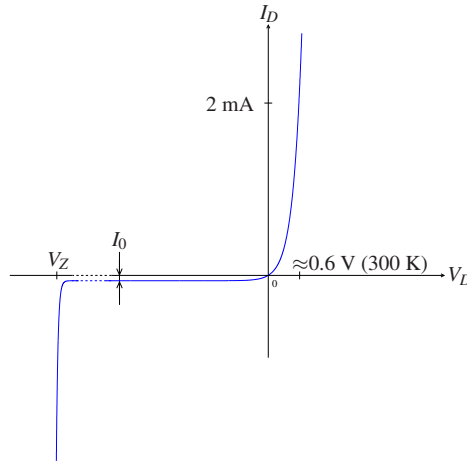


FIGURE 6 – Caractéristique courant tension de la diode 1N4148 à 300°K.

(matériau et dopage).

$$I_0 = AT^3 \quad (5)$$

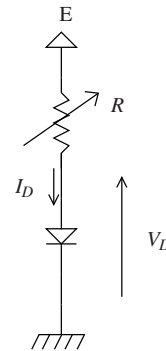
L'effet de la variation de I_0 avec la température sur I_D est prépondérant sur celui de la variation de V_ϕ avec la température ; I_D augmente avec T . La caractéristique devrait se déplacer vers la gauche avec la température.

Variation avec la température : Considérons le montage suivant :

On peut exprimer V_D : $V_D = E - RI_D$.

Cela permet de tracer les droites de charge (CF section 3.2) de la figure 7-gauche pour les variations de la caractéristique de la diode avec la température. Ces courbes ont été tracées pour $E = 5V$.

A partir de ce jeu de droites de charge on peut, pour chaque valeur de la résistance, en déduire la tension aux bornes de la diode en fonction de la température (figure 7-droite).



On remarque que toutes ces courbes présentent à peu près la même pente. On retrouve alors les valeurs classiquement admises dans la littérature :

$$\begin{aligned} V_D &\approx 0.6V \\ \frac{dV_D}{dT} &= -2\text{mV}/^\circ\text{C} \end{aligned}$$

$V_D = 0,6V$ est la valeur classique dans la littérature francophone pour les diodes en conduction. La littérature Anglo-saxonne utilise elle plutôt la valeur de $0,7V$. Il fait sans doute plus froid en Angleterre. Cette tension est souvent appelée tension de seuil de la diode.

3 Utilisation des diodes

La caractéristique réelle de la diode telle que présentée figure 6 est compliquée. On utilise, dès que possible, une des trois caractéristique simplifiée de la diode :

1. Cas des grands signaux : les tensions mises en œuvres sont grandes devant $0,6V$.

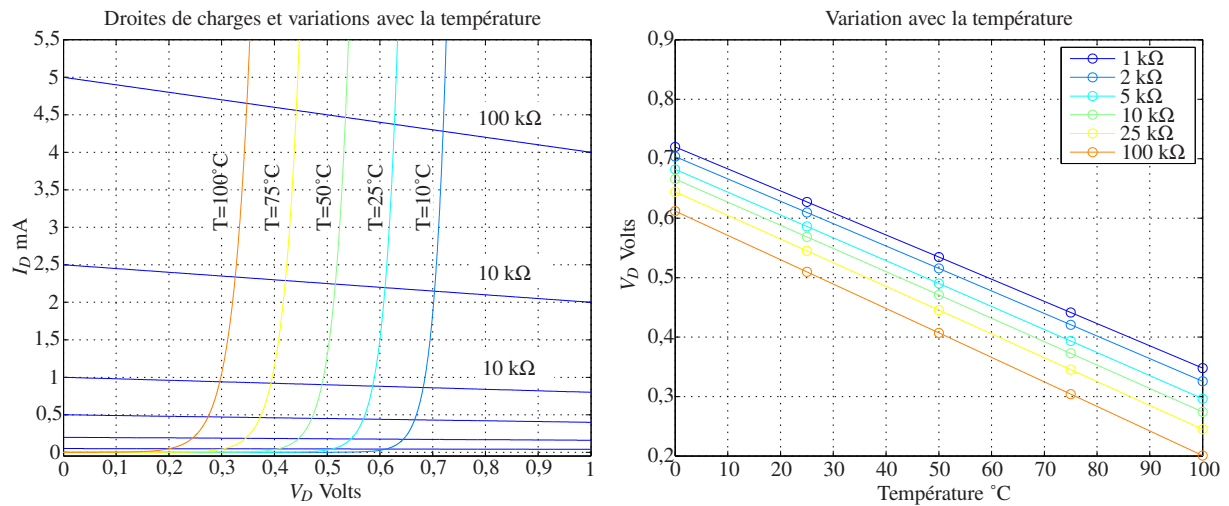
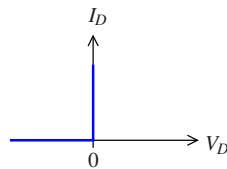
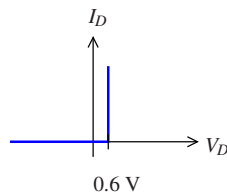


FIGURE 7 – Variation de la tension au borne d'une diode avec la température.

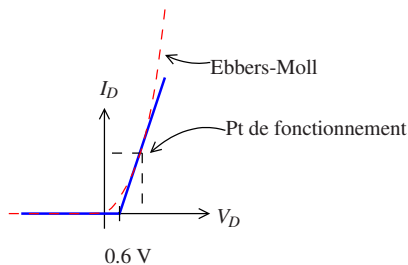
Dans ce cas on peut souvent négliger la tension de seuil, on utilise la caractéristique simplifiée suivante, dite de la diode idéale :



2. Signaux plus petits : on ne peut alors négliger la tension de seuil. On utilise le modèle de la diode idéale à seuil :



3. Cas des petites variations : si les variations de tensions sont faibles et situées dans la zone du coude de la diode, les modèles précédents ne donnent en général pas le bon résultat. Il est possible d'affiner en utilisant le modèle de la diode avec résistance différentielle. La valeur de la résistance différentielle est alors calculée en dérivant l'équation d'Ebers-Moll pour obtenir la pente ad-hoc autour du point de fonctionnement considéré.



4. Si aucun des modèles précédents ne convient, on utilise le modèle d'Ebers-Moll de l'équation 3. Dans tous les cas, on utilise le modèle le plus simple possible qui permet d'établir les bons résultats. Le bon résultat est celui qui correspond à ce que l'on peut réaliser et mesurer.

Retenons que :

ON UTILISE TOUJOURS LE MODÈLE LE PLUS SIMPLE POSSIBLE

3.1 Stratégies de calcul

La difficulté dans un circuit comportant des diodes consiste à trouver le bon angle d'attaque.

Il est souvent commode de raisonner en courant.

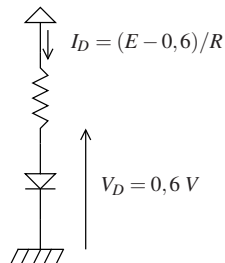
Avec le modèle de la diode idéale à seuil :

— Hypothèse 1 : $I_D = 0$; Diode Bloquée	
On calcule V_D	Si $V_D < 0,6 \text{ V}$ $\longrightarrow$ Hypothèse vérifiée
— Hypothèse 2 : $I_D \neq 0$; Diode passante $\iff V_D = 0,6 \text{ V}$	
On calcule I_D	Si $I_D > 0$ $\longrightarrow$ Hypothèse vérifiée

L'hypothèse diode bloquée ou passante est bien sûr guidée par le fait que le courant ne peut circuler que de l'anode vers la cathode dans la diode.

Exemple 1

$E \gg 0,6 \text{ V}$



L'hypothèse naturelle ici est diode passante. En effet avec $E > 0$ le courant ne peut circuler que du haut vers le bas.

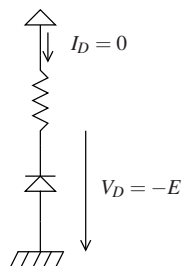
Cette hypothèse amène à $I_D = (E - 0,6)/R$. Comme $I_D > 0$ l'hypothèse est vérifiée.

L'hypothèse diode bloquée donnerait $V_D = E$, ce qui invaliderait l'hypothèse.

Sens direct

Exemple 2

$E \gg 0,6 \text{ V}$



L'hypothèse naturelle ici est diode bloquée. La diode ne peut ici laisser circuler le courant que du bas vers le haut.

Cette hypothèse amène à $V_D = -E$. Comme $V_D < 0,6 \text{ V}$, l'hypothèse est vérifiée.

L'hypothèse diode passante donnerait $I_D = -(E + 0,6)/R$. Dans ce cas $I_D < 0$ ce qui invaliderait l'hypothèse.

Sens inverse

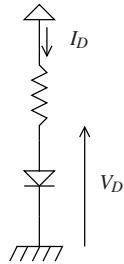
3.2 Droite de charge

Il est essentiel de savoir tracer les droites de charges pour prévoir le comportement de nombreux circuits non linéaires et diriger les raisonnements. Ce n'est pas un gadget, il faut savoir les tracer et penser à les utiliser.

Exemple : Tracé de la droite de charge correspondant au comportement en température de la section 2.3.

On calcule V_D

$$E \gg 0,6 \text{ V}$$

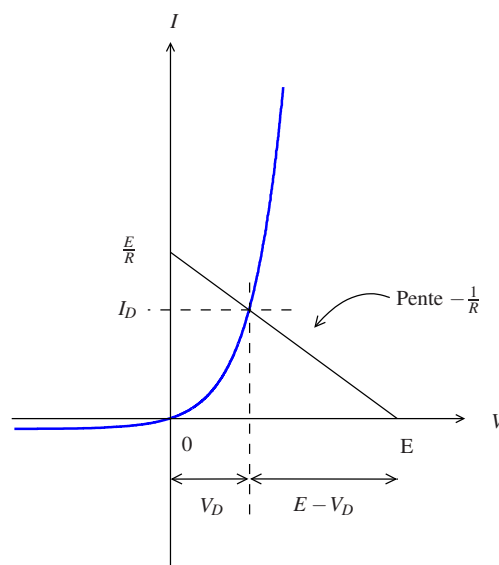


$$V_D = E - RI_D$$

On peut alors exprimer I_D :

$$I_D = -\frac{1}{R}V_D + \frac{E}{R}$$

On voit que c'est l'équation d'une droite de pente $-\frac{1}{R}$ qui intercepte l'axe des ordonnées en $\frac{E}{R}$, et l'axe des abscisses en E . On trace alors facilement la droite de charge comme présenté ci-dessous. Le résultat est obtenu ici avec le modèle d'Ebers-Moll comme le permet le tracé de la droite de charge, et non avec le modèle de la diode idéale à seuil comme dans la section précédente.

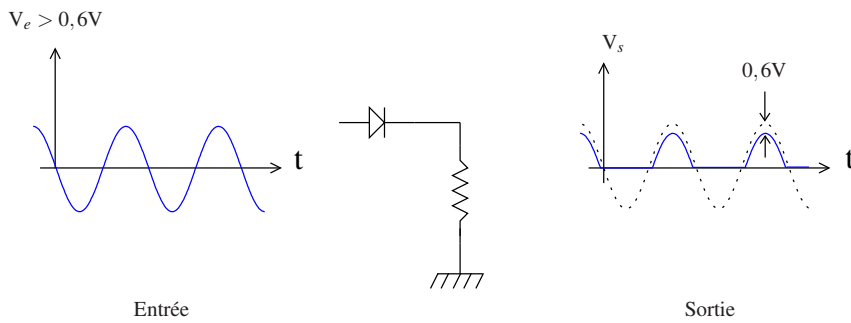


Conseil : Pensez à tracer les droites de charge pour comprendre ce qui se passe, avant de vous lancer dans les calculs.

3.3 Montages de base

Ici suivent des exemples choisis pour leur utilité pratique ou pour les raisonnements mis en œuvre pour leur compréhension.

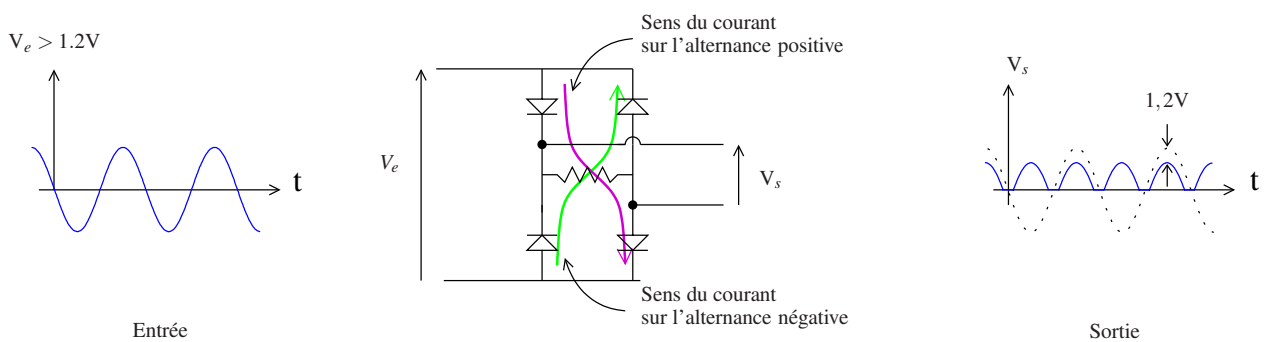
3.3.1 Redressement simple alternance



La diode coupe simplement l'alternance négative en ne laissant passer le courant que dans un sens. On peut remarquer sur la sortie l'effet de la tension de seuil. Le modèle utilisé ici est celui de la diode idéale à seuil.

3.3.2 Redressement double alternance

Le montage précédent ne permet pas de récupérer toute l'énergie du signal d'entrée, pour des applications d'alimentation par exemple. On lui adjoint une deuxième diode afin de permettre de récupérer les deux alternances. L'idée est de fournir au courant un chemin électrique qui lui permette de traverser la charge dans le même sens à chaque alternance. C'est le principe du pont en H.

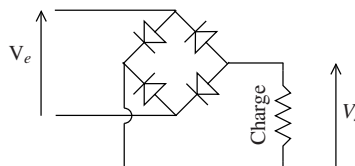


Comme on peut le voir dans le schéma ci-dessus, le courant passe toujours dans le même sens dans la charge quelque soit l'alternance. Ici encore, on peut remarquer sur la sortie l'effet de la tension de seuil. Le modèle utilisé ici est toujours celui de la diode idéale à seuil.

Le principe du pont en H est classique et est à retenir.

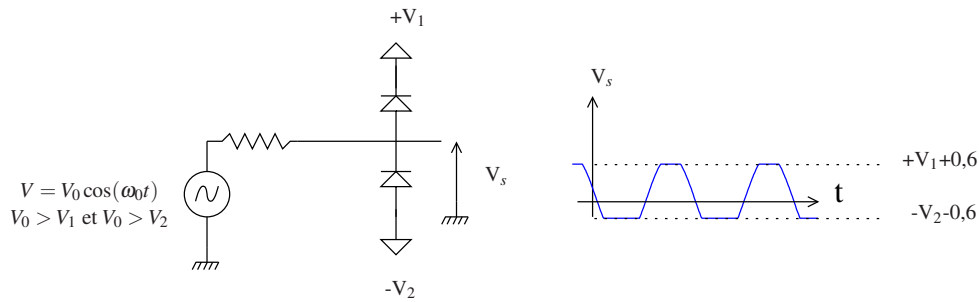
Remarque 1 : Ce montage ne fonctionne que chargé. En absence de la résistance, V_s n'est pas imposé par le circuit.

Remarque 2 : Il est souvent représenté comme suit et doit être reconnu au premier coup d'œil.



3.3.3 Écrêtage (*clamping*)

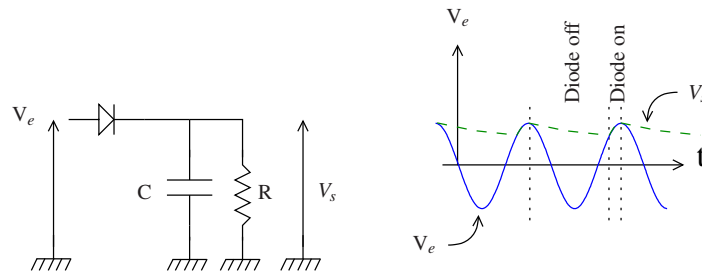
L'écrêtage est l'une des grandes utilisations des diodes. Il est souvent exploité pour protéger d'autres circuits en limitant leur tension d'entrée.



On peut avoir $V_1 = V_2 = 0$, dans ce cas l'excursion de sortie est limitée à $\pm 0,6$ V. C'est le montage type de protection contre les surtensions des amplificateurs d'instrumentation où les signaux à amplifier sont en général très petits devant 0,6 V.

3.3.4 Détection de crêtes

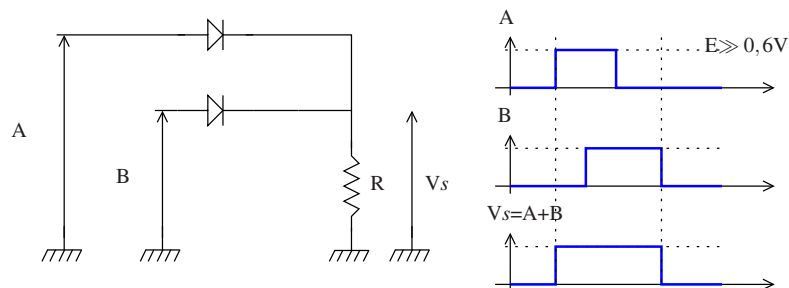
Le condensateur se charge sur les alternances positives au travers de la résistance différentielle de la diode qui est de l'ordre de la dizaine d'Ohm. Pendant les alternances négatives, la diode est bloquée et le condensateur se décharge dans la résistance avec la constante de temps $\tau = RC$.



En choisissant correctement τ , ce circuit peut être typiquement appliqué à la démodulation d'amplitude. Si τ est suffisamment grand, la tension de sortie peut être considérée comme constante et est égale à la valeur crête positive maximale atteinte.

3.3.5 Logique à diodes

Dès que l'une des diodes est passante, la tension aux bornes de R vaut E (diode idéale $E \gg 0,6$ V). La fonction logique réalisée est un OR (OU logique).



L'intérêt de ce montage n'est pas de remplacer une porte logique, mais réside dans ses niveaux de sortie. En fonction des diodes choisies, il permet de réaliser une porte OR supportant de grandes valeurs de courant et de tension.

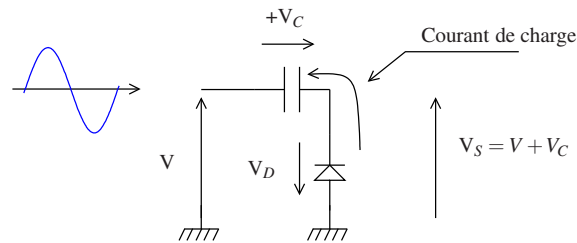
3.3.6 Diode de roue libre

Se reporter au chapitre alimentation à découpage.

3.3.7 Pompe à diode

Ce circuit est un circuit très simple, mais dont le fonctionnement est très astucieux et assez compliqué si l'on rentre vraiment dans les détails.

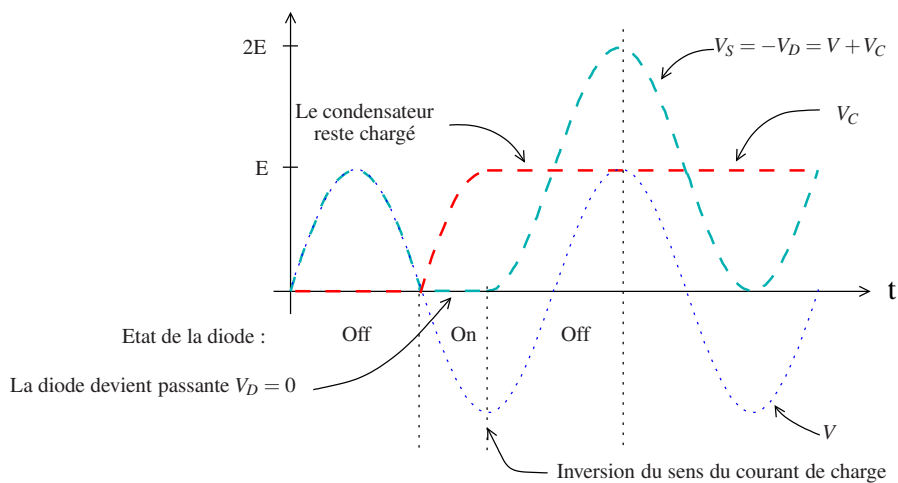
Premier étage : Pompage de la charge.



On utilise ici le modèle de la diode idéale. Le modèle de la diode à seuil pourrait être utilisé, c'est même nécessaire si l'amplitude du signal d'entrée V n'est pas beaucoup plus grande que 0,6 V. Les raisonnements sont les mêmes dans les deux cas. On se contentera du modèle le plus simple pour des raisons de clarté.

Le raisonnement en courant est ici, comme souvent avec les diodes, le plus efficace. Le condensateur ne peut se charger que si le courant est dans le sens indiqué sur la figure précédente. Pour cela, il faudra avoir $V_D > 0$, c'est-à-dire une alternance négative.

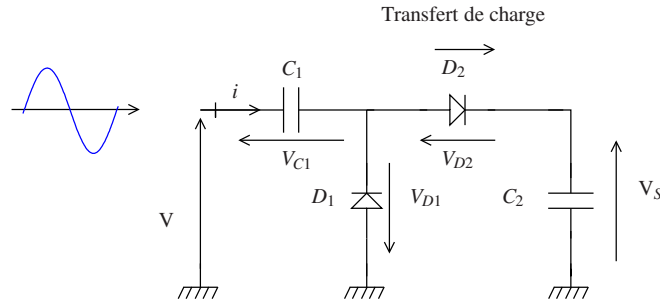
Avec le condensateur initialement déchargé, donc $V_C = 0$, V_D devient positif dès que $V < 0$. La diode devient passante et la tension de sortie est nulle. La diode reste passante jusqu'à ce que le sens du courant s'inverse, c'est-à-dire jusqu'à ce que la tension remonte ($i = \frac{dV}{dt}$). Ce fonctionnement est résumé dans la figure suivante.



On voit que l'on a bloqué la charge dans la diode ce qui conduit à un *offset* de la tension de sortie.

Retenons que la diode devient passante dès que la charge $+Q$ sur l'électrode de droite du condensateur est inférieure à CV .

Deuxième étage : on rajoute un étage de stockage de la charge. L'étage rajouté ressemble à un détecteur de crête. De cette façon, on cherche à obtenir en sortie une tension constante égale à $2E$. En pratique, c'est beaucoup moins simple car le fonctionnement de l'étage de pompage de charge est influencé par l'étage de transfert de charge. Néanmoins l'étage rajouté fonctionne sur le même principe que le détecteur de crête, c'est-à-dire en bloquant la charge dans le condensateur C_2 . Dans ce montage les capacités des condensateurs C_1 et C_2 sont les mêmes.



Les chronogrammes des différents signaux de la pompe à diode en fonctionnement sont représentés figure 8.

On suppose les condensateurs initialement déchargés. La tension d'entrée est $V = E \cos(\omega t)$. On pose $Q(t) = CV(t)$ la charge stockée dans l'un ou l'autre des condensateurs quand la différence de potentiel vaut $V(t)$ à leurs bornes et $Q_E = CE$ la valeur de cette charge lorsque $V(t) = E$.

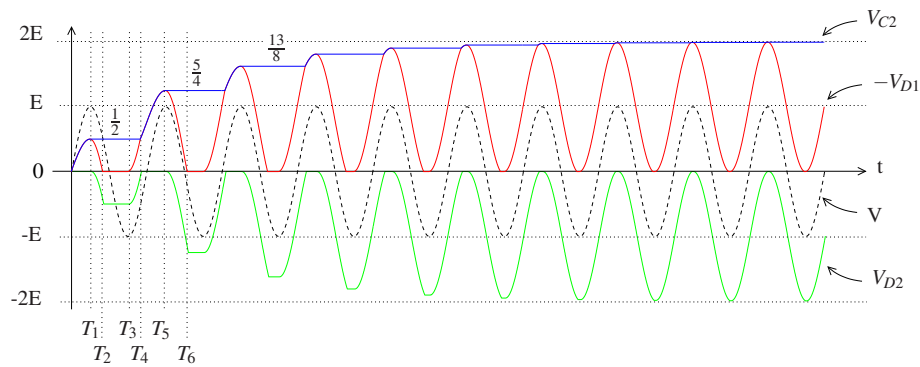


FIGURE 8 – Chronogrammes du doubleur de tension à pompe à diode

Analyse du fonctionnement :

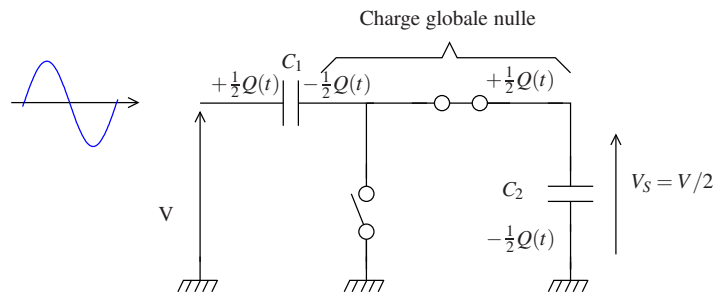
— $t = 0$:

$V = 0$, les deux condensateurs sont déchargés et $V_{C1} = V_{C2} = 0$. Le courant dans chacune des diodes est nul, elles sont donc bloquées et $V_S = 0$.

— $t = 0+$:

La tension d'entrée V devient positive, ce qui demande un courant d'entrée positif. D_1 reste bloquée et D_2 devient passante : $V_{D2} = 0$.

Le schéma équivalent du montage est donc :

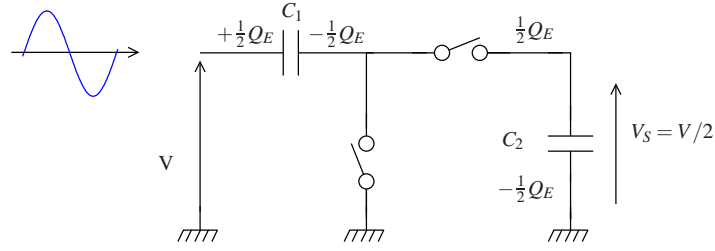


Comme D_1 est bloquée, la charge globale entre les deux condensateurs est nulle et se répartit comme dans la figure ci-dessus avec $V_S = V/2$ (pont diviseur capacitif).

De plus $V_{C1} = V - \frac{V}{2} = \frac{V}{2}$ et $V_{D1} = -V_S = -\frac{V}{2}$.

— $t = T_1$

Dès que la tension en entrée V diminue, la diode D_2 se bloque car le courant tend à s'inverser pour décharger C_2 au travers de D_2 pour suivre la tension $V(t)$. Les deux diodes sont bloquées et le schéma équivalent devient :

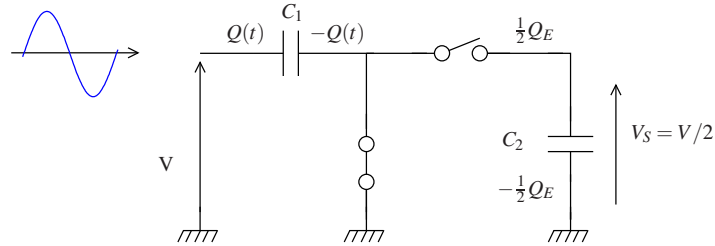


La charge piégée dans C_2 vaut $Q_E/2$ et la tension de sortie reste constante et vaut : $V_S = E/2$.

La charge piégée dans C_1 vaut $-Q_E/2$ et $V_{D1} = -(V - E/2) = E/2 - V$. La diode D_1 restera bloquée tant que $V_{D1} < 0$

— $t = T_2$

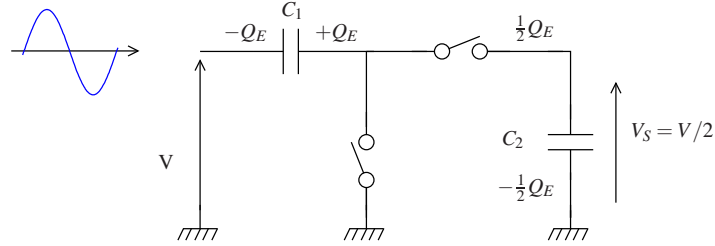
$V(T_2) = E/2 \Rightarrow V_{D1} = 0$. La diode D_1 redevient passante. Le schéma équivalent devient :



La charge dans C_1 recommence à suivre la tension $V(t)$. C_1 se charge jusqu'à ce que $V(t)$ atteigne sa valeur minimale $-E$, soit une charge $+Q_E$ sur l'électrode de droite de C_1 ($V(t) < 0$).

— $t = T_3$

D_1 se bloque car les charges ont maintenant tendance à quitter C_1 . D_1 et D_2 sont bloquées :

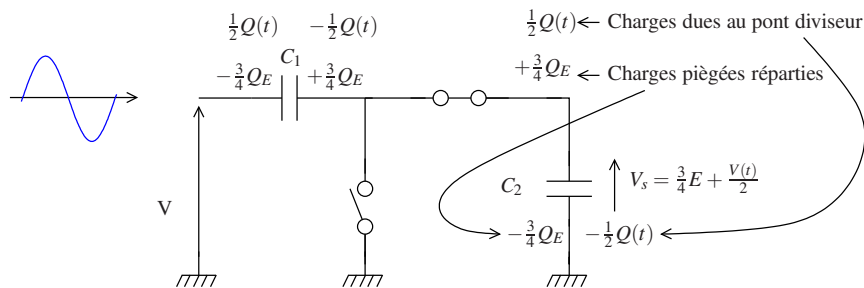


$V_{C1} = E$.

V_S reste constante.

— $t = T_4$

D_2 redevient passante dès que $V(t)$ est suffisamment grande pour recommencer à charger C_2 , c'est-à-dire lorsque $V(t) + E = \frac{E}{2}$, soit en $V(T_4) = -\frac{E}{2}$:

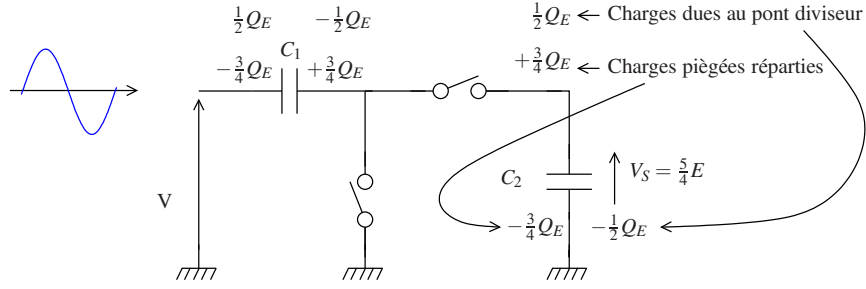


Les charges sur les électrodes internes des deux condensateurs (électrode de gauche de C_1 et électrode de droite de C_2) sont la somme des charges précédemment piégées, qui se répartissent équitablement $\left(+\frac{1}{2} \times \frac{3}{2} Q_E \right)$, et de la répartition de charges dues au pont diviseur capacitif qui est de nouveau actif $\left(\pm \frac{Q(t)}{2} \right)$. On a donc : $V_s = \frac{3}{4} E + \frac{V(t)}{2}$.

— $t = T_5$

D_2 se bloque comme dans l'alternance précédente lorsque $V(t)$ diminue pour prévenir la fuite des charges de C_2 .

Le schéma équivalent devient :



Compte tenu de la répartition des charges et de la valeur de $V(T_5)$: $V_s = \frac{3}{4}E + \frac{1}{2}E = \frac{5}{4}E$. V_s reste constante jusqu'à ce que D_1 se débloque. Comme la charge bloquée dans C_1 vaut $\frac{1}{4}QE$, la diode D_1 se débloquera dès que $V(t) + \frac{1}{4}E = 0$.

— $t = T_6$

$V(T_6) = -\frac{1}{4}E \implies D_1$ redevient passante. On a : $V_{D_1} = 0$.

Elle se re-bloquera comme précédemment au minimum de l'alternance de $V(t)$ et ainsi de suite.

Conclusion A chaque nouveau blocage de D_1 , la charge stockée dans C_2 s'ajoute à la charge $+QE$ pompée dans C_1 lors de la phase de conduction de D_1 . Cette charge se répartit et s'ajoute à celle du pont diviseur lorsque D_2 redevient passante. On aboutit à la récurrence suivante :

$$\begin{aligned}
 Q_0 &= 0 \\
 Q_1 &= \frac{1}{2}QE \\
 Q_2 &= (Q_1 + QE)\frac{1}{2} + \frac{1}{2}QE = \frac{1}{2}Q_1 + QE = \frac{5}{4}QE \\
 Q_3 &= \frac{1}{2}Q_2 + QE = \frac{13}{8}QE \\
 &\vdots \\
 Q_i &= \frac{1}{2}Q_{i-1} + QE
 \end{aligned} \tag{6}$$

Il s'agit d'une suite arithmético-géométrique que l'on peut résoudre :

$$Q_{n>0} = QE \left[\frac{2^{n+2} - 3}{2^{n+1}} \right] \tag{7}$$

Quant $n \rightarrow +\infty$ La charge tend vers $2QE$ et V_s vers $2E$ ce qui était le but recherché.

Généralisation En considérant la figure 9, on peut voir que le potentiel sur l'anode de D_2 (électrode de gauche) est de la forme $E + \cos(\omega t)$. L'idée vient naturellement de charger le doubleur par un autre doubleur et ainsi de suite. On obtient ainsi le montage de la figure 9.

Là encore l'analyse précise du fonctionnement n'est pas triviale. On peut vérifier cependant qu'une pompe à n étages permet de récupérer l'amplitude de la sinusoïde d'entrée multipliée par $2n$.

Les résultats obtenus avec le doubleur de tension montrent que tant que la tension finale n'est pas atteinte, les diodes D_1 et D_2 sont alternativement, en opposition de phase, bloquées et passantes.

Ce résultat se généralise aisément à une pompe à n étages en considérant le sens du courant demandé par le générateur. C'est ce fonctionnement qui vaut son nom à ce montage.

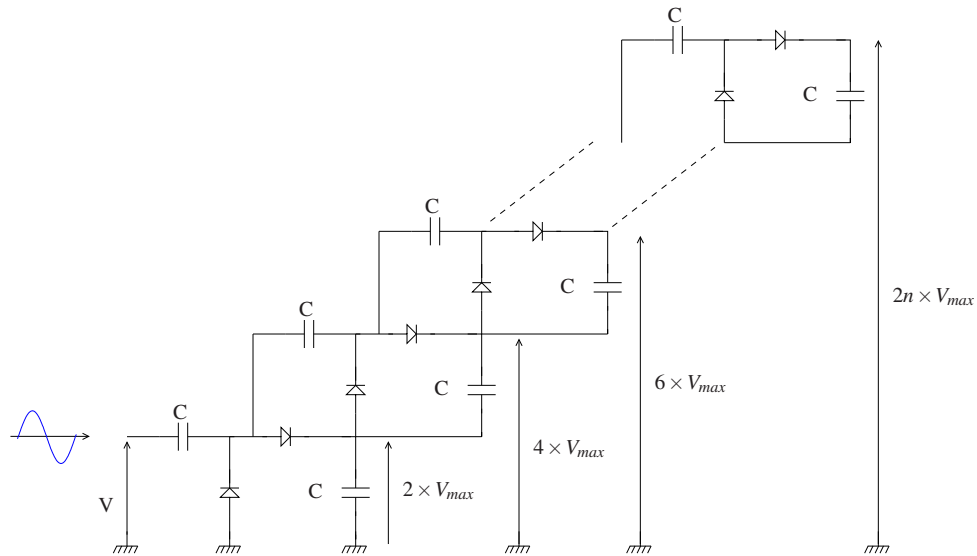


FIGURE 9 – Pompe à diode à n étages.

Applications La pompe à diode permet d’obtenir facilement de très hautes tensions (quelques centaines de kV). Le courant disponible en sortie est cependant très faible car il dépend de l’efficacité du pompage qui diminue avec la tension. En effet lorsque la tension de sortie est nulle on transfère une charge $Q_E/2$ à chaque alternance. Cependant, au niveau de charge correspondant à la n ème itération le transfert de charge n’est plus que de $\frac{3}{2^{n+1}}$.

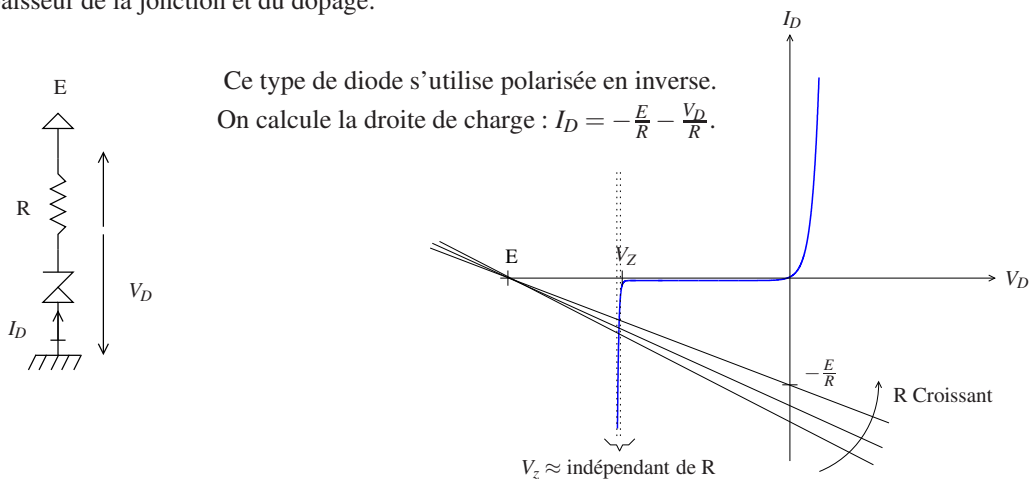
Ce montage reste cependant le montage de choix dans les applications qui demandent de très hautes tensions avec très peu de courant tel que les compteurs Geiger et pour la polarisation des diélectriques par exemple.

Il reste aussi un très bon exercice de compréhension du fonctionnement d’un montage à diodes.

4 Autres types de diodes

4.1 Diodes Zener

Ce type de diode tire parti de l’effet Zener. La tension d’avalanche est fixée précisément à la construction en jouant sur l’épaisseur de la jonction et du dopage.



De cette façon, comme la pente de la zone d’avalanche sur la caractéristique est très raide et peu dépendante de la température, on obtient facilement une référence de tension.

Attention : Bien que le choix de la résistance n’influe pas sur la tension délivrée, il permet de contrôler le courant

dans la diode : $I_D = \frac{E - V_Z}{R}$. Ce dernier doit être compatible avec ce que peut supporter la diode en terme de dissipation thermique.

Exemple d'application : référence de tension pilotable. le problème avec le montage précédent est que l'on ne peut pas mettre n'importe quelle charge en sortie. Si on tire trop de courant, le diode peut quitter le mode d'avalanche en faisant descendre la tension V_D en dessous de V_Z . Le lecteur pourra le vérifier.

Le montage de la figure 10 permet de contourner le problème car le courant I^+ reste toujours faible. C'est le transistor qui fournit le courant, la tension de sortie est simplement asservie sur la tension de la diode Zener en fonction de la valeur du pont diviseur. On obtient alors une référence de tension réglable avec une faible impédance de sortie : approximativement celle de l'alimentation en série avec la résistance passante du transistor.

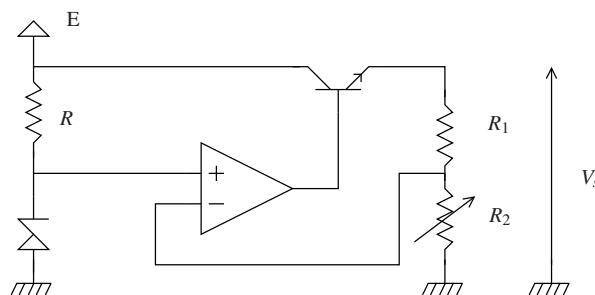
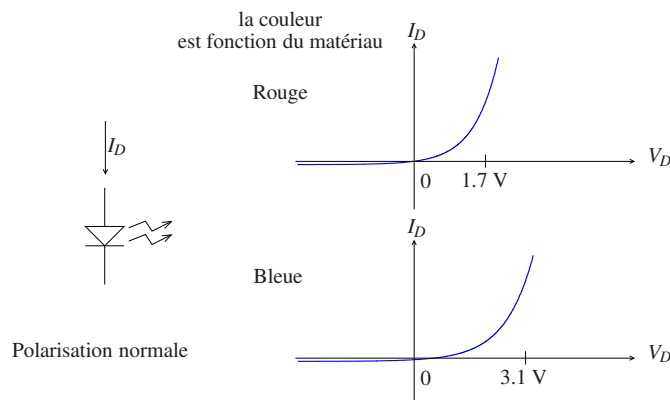


FIGURE 10 – Référence de tension réglable à diode Zener.

On vérifie aisément que $V_S = \frac{R_1 + R_2}{R_2} V_Z$.

4.2 Diodes Électroluminescentes (LED)

Les diodes électroluminescentes (Light Emitting Diodes) émettent de la lumière lorsqu'elles sont traversées par un courant. On les polarise donc en direct.



L'émission de lumière est due à la désexcitation d'un électron qui retourne dans la bande de valence. C'est en fait une recombinaison électron-trou comme présenté au paragraphe 1.2.3.

L'énergie, donc la fréquence de la lumière, dépend du matériau utilisé. C'est aussi le cas de la tension de seuil.

Remarque 1 : Pour le silicium, le gap d'énergie entre la bande de valence et la bande de conduction interdit la création d'un photon, un phonon est créé à la place (vibration du réseau cristallin). Il n'existe pas de LED au silicium.

Remarque 2 : Il est aussi possible grâce à cette émission de lumière de créer des diodes laser (CF cours d'optique deuxième année).

Un exemple intéressant : La figure 11 présente la mise en œuvre du pilotage d'une LED par une dynamo. Quel est le rôle de chacun des composants ?

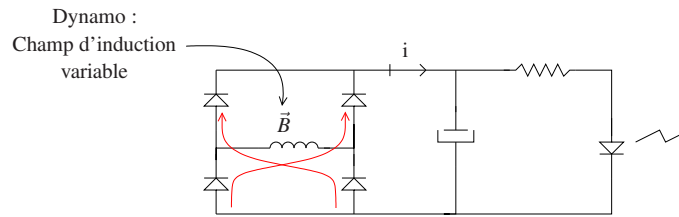
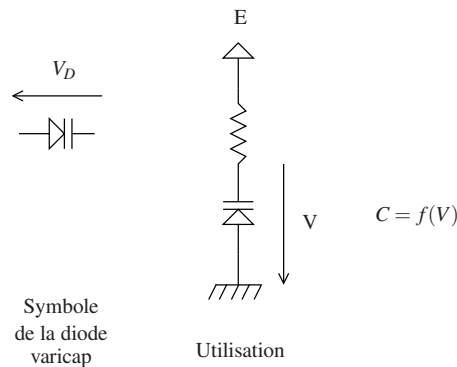


FIGURE 11 – Exemple d'application : alimentation d'une LED par une dynamo.

4.3 Varicaps

Comme présenté figure 5, la jonction PN présente une zone de déplétion lorsqu'elle est polarisée en inverse. Avec ses électrodes la zone de déplétion ainsi créée présente la structure d'un condensateur : Electrode-Isolant-Electrode. Comme l'épaisseur de la zone isolante (zone de déplétion) dépend de la tension aux bornes de la diode en polarisation inverse, on peut tirer parti de cet effet pour créer une capacité variable.



Certaines diodes comme la SOD 323 sont fabriquées pour exploiter cet effet. La figure 12 présente, à titre d'ordre de grandeur, les performances de cette dernière.

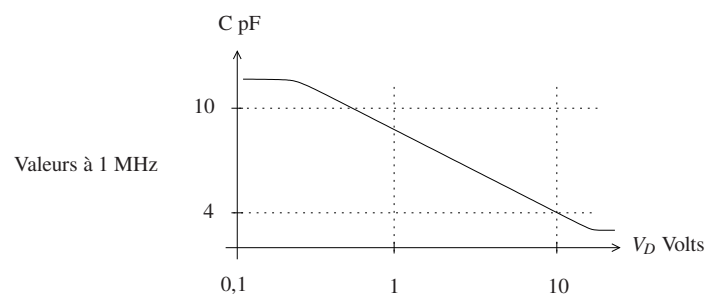
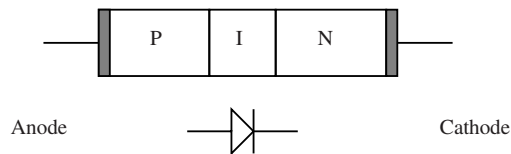


FIGURE 12 – Extrait des performances typiques SOD323.

Les varicaps sont typiquement utilisées dans le pilotage en tension d'oscillateurs.

4.4 Diodes PIN

Les diodes PIN sont constituées de deux jonctions entre un semi-conducteur intrinsèque I, et deux semi-conducteurs extrinsèques respectivement dopés P et N comme présenté ci-dessous.



Le fonctionnement détaillé de ce type de diodes ne sera pas abordé ici.

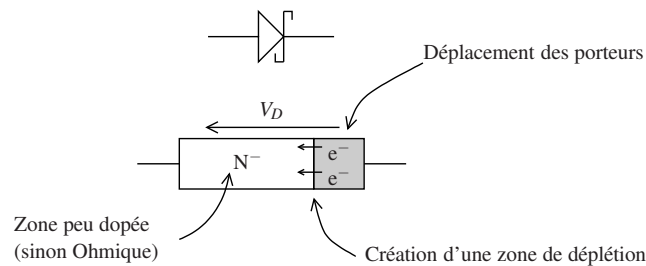
Principales différences par rapport à une diode classique PN :

Tension de seuil faible	$\approx 0,4 \text{ V}$
Tension de claquage élevée	$> 10 \text{ kV}$
Capacité propre faible et peu dépendante de la tension.	C'est tout l'inverse d'une varicap

Ces diodes sont très utilisées en hyperfréquence comme interrupteur. En BF on les utilise surtout pour leur faible tension de seuil ou leur tension de claquage élevée.

4.5 Diodes Schottky

Les diodes Schottky (*Schottky barrier diodes*) sont constituées d'une jonction métal semi-conducteur :



La création d'une zone de déplétion crée, comme dans la jonction PN, une barrière de potentiel.

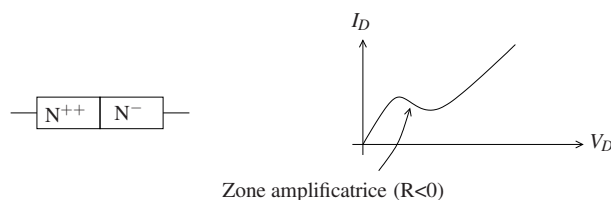
Principales différences par rapport à une diode classique PN :

Tension de seuil faible	$\approx 0,2 \text{ V}$
Commutation très rapide	Moins de porteurs à déplacer
Courant inverse élevé	$qq \mu A$

Elle sont très utilisées en protection contre les surtensions à cause de leur rapidité.

4.6 Diodes GUN (Hyperfréquences)

Leur usage est principalement limité aux hyperfréquences, où elles sont utilisées pour réaliser des oscillateurs en tirant parti de leur zone amplificatrice.



5 Interaction avec la lumière

Lorsqu'un photon provenant d'une source électromagnétique externe pénètre dans la zone de déplétion d'une jonction PN, il peut, si son énergie est suffisante, casser une liaison covalente et donc créer une paire électron-trou.

Pour le silicium, l'énergie minimale nécessaire pour le photon incident est de 1,1 eV. Cette condition peut s'exprimer comme :

$$h\nu > 1,1\text{eV} \quad (8)$$

h est la constante de Plank, ν est la fréquence du photon considéré.

On peut en déduire la condition équivalente sur la longueur d'onde du photon :

$$\lambda < 1,3 \text{ } \mu\text{m} \quad (9)$$

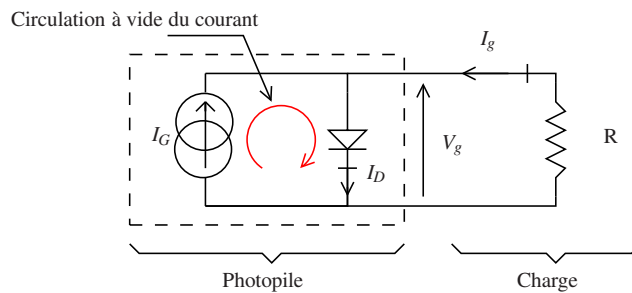
On se trouve donc dans l'infra-rouge lointain. Les jonctions PN au silicium sont donc sensibles aux sources lumineuses à partir de l'infra-rouge.

A cause du champ électrique qui règne dans la zone de déplétion en circuit ouvert (CF figure 3), ou des gradients de concentration de porteurs de charge si on est en circuit fermé l'électron et le trou ainsi créés se déplacent en sens inverse l'un de l'autre sans se recombiner et cela se traduit par l'apparition d'un courant.

Les photopiles (cellules photovoltaïques) et les photodiodes utilisent ce phénomène respectivement pour la production d'énergie et la détection des photons. Les deux types de composant sont des jonctions PN que la lumière peut atteindre. Ils ne sont cependant pas optimisés dans le même but. Dans le premier cas, le rendement énergétique est visé, dans le deuxième cas c'est la sensibilité. Les deux composants ne s'utilisent pas non plus de la même façon.

5.1 Photopiles

La photopile est une source de courant et s'utilise comme un générateur. Elle est simplement branchée sur une charge comme dans la figure ci-dessous. Son schéma équivalent comporte donc un générateur de courant. Il comporte aussi une diode qui modélise la caractéristique courant tension de la jonction PN. Comme la photopile peut être en circuit ouvert (non chargée) et que le courant généré circule, le générateur de courant et la diode du modèle doivent être orientés de façon ad-hoc.



Dans ce schéma, I_G est le courant généré par le flux Φ de photons incidents sur la jonction : $I_G = f(\Phi)$. La diode qui modélise la jonction est traversée par le courant I_D . La tension V_D est aussi la tension fournie par le générateur V_g . V_g et I_D sont donc reliés par l'équation d'Ebers-Moll (équation 3). Enfin I_g est le courant fourni à la charge.

On peut alors tracer la caractéristique courant-tension de la photopile chargée $I_g = f(V_g)$.

$$I_g = I_D - I_G \quad (10)$$

Dans cette équation, seul I_D dépend de V_g . La caractéristique de la photopile en circuit ouvert est donc celle d'une diode décalée vers le bas de I_G comme présenté figure 13.

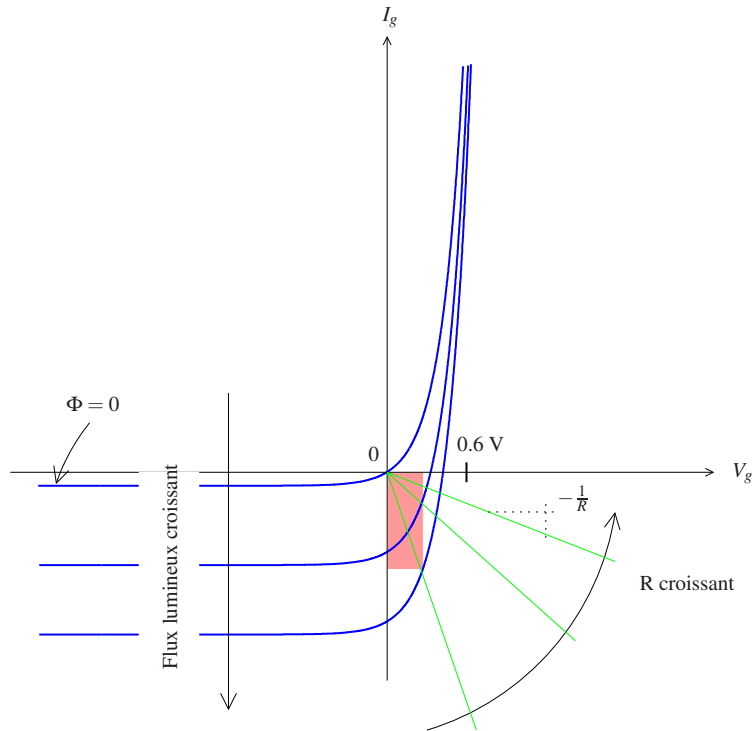


FIGURE 13 – Jeu de caractéristiques d’une photopile et droites de charges.

On peut alors tracer la droite de charge pour une charge résistive : $V_g = -RI_g$. On voit alors que compte tenu de la raideur de la caractéristique de la jonction PN, la tension aux bornes du générateur dépend peu de R et vaut approximativement 0,6 V pour une jonction au silicium. La meilleure valeur pour R est celle qui permet de récupérer le maximum de puissance, c’est-à-dire celle qui correspond à la plus grande surface du polygone (coloré) $V_D \times I_D$ (sa pointe droite en bas se trouve alors à peu près dans le coude de la caractéristique) dans la figure 13. Soit pour une valeur de R optimale qui dépend de Φ .

Il est possible, pour obtenir de plus grandes valeurs de tension, de cascader en série autant de photopiles que nécessaire. On obtient alors une tension de sortie multiple de 0,6 V.

5.2 Photodiodes

Alors que la photopile est utilisée comme un générateur et est polarisée en direct, la photodiode est utilisée comme un détecteur et est polarisée en inverse :

1. Chaque photon incident crée une paire électron trou.
2. Les électrons créés se déplacent sous l’action du champ électrique dû à la polarisation (E) et créent d’autres paires par collision avec ceux encore engagés dans des liaisons covalentes. Il en résulte une amplification du photocourant. Attention l’énergie est ici fournie par l’alimentation E .
3. La tension V , qui vaut E en l’absence de photons incidents, décroît rapidement.

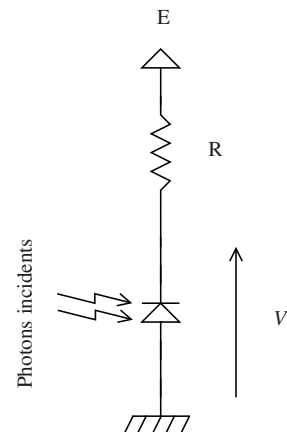
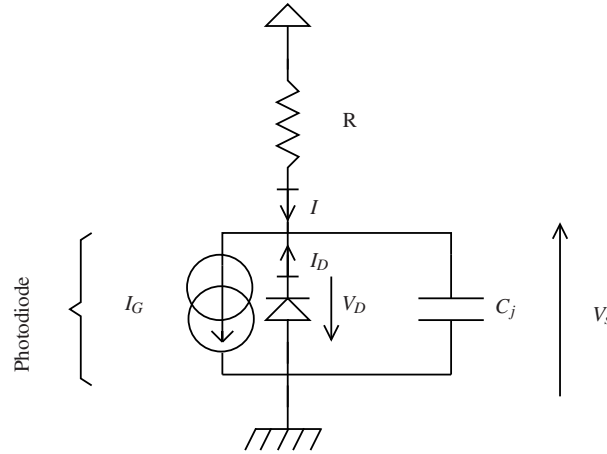


Schéma équivalent : le schéma équivalent est le même que celui de la photodiode. C’est en effet intrinsèquement le même composant. Cependant, comme on l’utilise en polarisation inverse, il faut tenir compte de la capacité de jonction C_j . Cette capacité a les mêmes origines que pour les varicaps (CF 4.3). On obtient donc le schéma équivalent suivant, dans lequel I_G est le photo courant multiplié.



Sensibilité statique : Ici encore pour obtenir la sensibilité statique, il faut tracer la droite de charge :

$$\begin{aligned} V = -V_D &= E - RI \\ I &= I_G - I_D \end{aligned} \quad (11)$$

On remarquera le signe moins devant la tension de la diode qui implique une symétrie de la caractéristique courant tension. Le tracé est présenté figure 14. Il permet de vérifier que la sensibilité augmente avec R . Il fait aussi apparaître que le système sature dès que le flux de photons devient trop grand à cause du coude de la diode. La réponse de la photodiode ne sera pas linéaire dans cette zone. Cela se traduit par une tension minimale V_{sat} en dessous de laquelle on ne peut descendre. Le bon choix pour R est donc un compromis entre saturation et sensibilité.

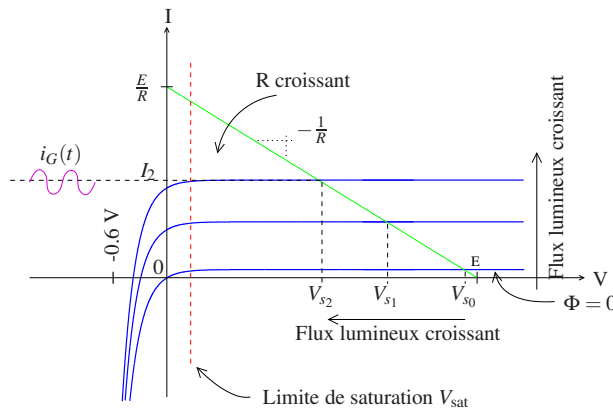
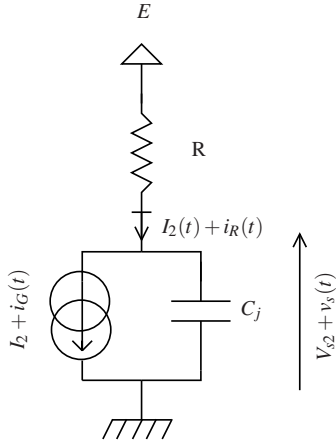


FIGURE 14 – Jeu de caractéristiques d’une photodiode et droite de charge.

Sensibilité dynamique, limitation de la bande passante par la capacité de jonction. La photodiode est un détecteur. Le flux de photon en entrée peut être très rapidement variable (dans les applications d’alimentation avec les photopiles, le problème ne se pose pas).

Considérons un flux de photons variable autour du flux Φ_2 : $\Phi(t) = \Phi_2 + \varphi(t)$ tel que on soit autour du point (V_{s2}, I_2) de la figure 14. Si de plus les variations du flux de photons sont telles que l’on ne s’approche pas de la zone de saturation (la caractéristique de la diode reste plate dans la zone de travail), alors les variations de courant dans la photodiode peuvent s’écrire : $I_G = I_2 + i_G(t)$ où $i_G(t)$ est proportionnel à $\varphi(t)$. Le comportement du montage à photodiode est alors décrit par le schéma suivant :



Dans cette figure, les tensions et courants ont été décomposés entre leur parties continues (en lettres capitales) et leur parties variables (en petites lettres).

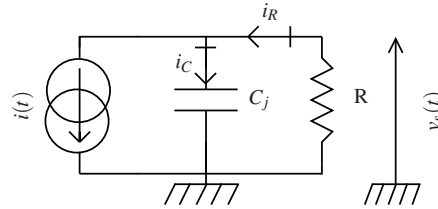
- La partie continue I_2 du courant de photon ne passe pas par le condensateur.
On a donc $I = I_R = I_2 + i_R(t)$, et $V_{s2} = E - RI_2$.
- La partie variable du courant passe dans la capacité de jonction et la résistance
d'où : $i_R(t) = i_G(t) + i_C(t)$.

D'après le schéma : $V_{s2} + v_s(t) = E - R(I_2 + i_R(t))$.

Finalement :

$$v_s(t) = -Ri_R(t) = -R.(i_G(t) + i_C(t)) \quad (12)$$

Ce résultat est celui que l'on obtiendrait aussi avec le schéma suivant qui est donc un schéma équivalent pour les petites variations du signal :



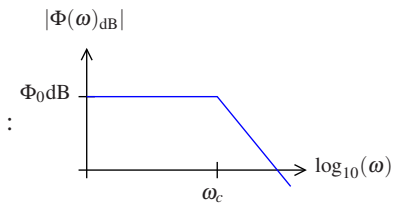
Notez le sens du courant dans le schéma équivalent qui permet de tenir compte du signe dans l'équation 12. Ce schéma permet de calculer facilement la réponse à la partie variable du signal grâce au pont diviseur de courant :

$$i_R(\omega) = i_G(\omega) \frac{1/jC\omega}{1/jC\omega + R} \quad (13)$$

Soit compte tenu de l'équation 12 ou le schéma équivalent pour les petites variations :

$$v_s(\omega) = -i_G(\omega) \frac{R}{1 + jRC\omega} \quad (14)$$

La capacité de jonction donne donc une fréquence de coupure à -3dB pour $\omega = \frac{1}{RC}$:



6 Critères de choix d'une diode.

Lorsque l'on doit choisir une diode, les paramètres principaux à considérer sont les suivants :

Nom	Description	Valeurs de référence (1N4148)
Courant moyen	Courant continu maximum que peut supporter la diode.	300 mA
Tension d'avalanche	V_Z tension Zener	75 V
Courant inverse	I_0	25 nA
Temps de recouvrement	Temps nécessaire pour passer de l'état bloqué à passant.	8 ns
Capacité de jonction	Valeur de la capacité en mode bloqué	4 pF

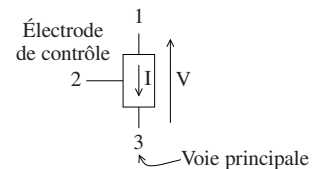
TABLE 1 – Principaux paramètres à considérer pour le choix d'une diode.

7 Transistors

7.1 Qu'est ce qu'un transistor ?

Il en existe plusieurs sortes. Dans tous les cas, ce sont des composants à trois électrodes (trois pattes). Deux des électrodes constituent la voie principale dans lequel le courant peut circuler si on leur applique une différence de potentiel. La circulation du courant dans la voie principale est contrôlée par le potentiel appliqué ou le courant que l'on fait circuler dans la troisième électrode qui constitue une électrode de pilotage.

Dans tous les cas, le courant dans l'électrode de contrôle est très petit devant le courant qui circule dans la voie principale.



7.2 Utilisation

Les transistors trouvent deux modes d'utilisation : la commutation et les applications linéaires (CF Systèmes linéaires), en particulier l'amplification. Dans le premier, cas le transistor est soit passant soit isolant et fonctionne comme un interrupteur commandé par le signal appliqué sur l'entrée de commande. Dans le deuxième, une tension constante dite de polarisation est appliquée en V. Cette tension peut être assez grande (quelques dizaines de Volts) et correspond potentiellement à un courant relativement grand devant l'entrée de commande. De cette façon, un petit signal (courant ou tension) de commande peut moduler un grand courant dans la voie principale du transistor, créant ainsi la fonction amplification par exemple.

Il existe deux familles de transistors :

1. Les transistors à effet de champ : FET (*Field Effect Transistors*). Ce sont ceux utilisés dans les composants logiques donc de fait les plus répandus.

Dans ces transistors, c'est le champ électrique interne qui est directement responsable du fonctionnement.

2. Les transistors bipolaires à jonctions : BJT (*Base Junction Transistors*). C'est le premier type de transistor découvert.

Leur fonctionnement est basé sur des effets de jonction comme pour les diodes.

8 Transistors à effet de champ

Parmi les transistors à effet de champ, on distingue les transistors à grille isolée et non isolée. Ce cours ne portera que sur les transistors à grille isolée que sont les MOSFET (*Metal Oxyde Silicium FET*). C'est en effet le composant roi de l'électronique logique.

Il en existe deux types, les transistors à canal N et les transistors à canal P selon le dopage du canal.

8.1 MOS canal N (NMOS)

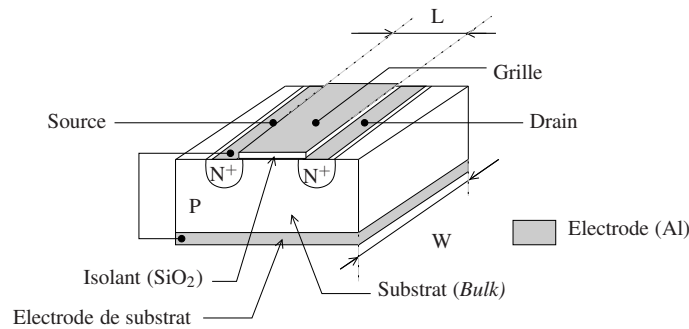
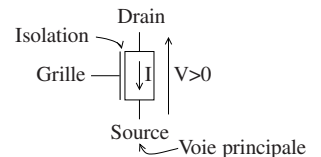


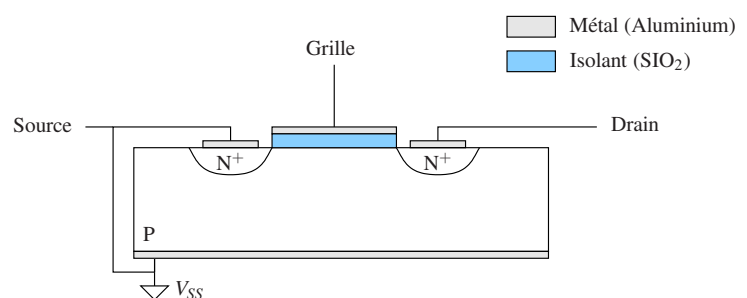
FIGURE 15 – Vue en perspective d'un MOSFET canal N.

Un transistor MOSFET, ou MOS, en abrégé est réalisé comme présenté figure 15. Il est constitué d'un substrat P déposé sur une électrode plane. Cet ensemble constitue le “bulk” ou “body” du transistor (*bulk* : partie massive en anglais). L'électrode de bulk est en général prise comme référence de potentiel. Au sommet du substrat on réalise deux caissons dopés N, recouverts chacun d'une électrode plane rectangulaire, qui constituent les électrodes de source et de drain. Entre les deux électrodes, on vient déposer une couche d'oxyde de silicium (SiO_2) isolante. Cette couche isolante est surmontée, elle aussi, d'une électrode rectangulaire de largeur W et de longueur L qui forme l'électrode de grille. Celle-ci est donc isolée du substrat, d'où le nom de transistor à grille isolée. Les grandeurs W et L sont caractéristiques du transistor réalisé et leur rapport W/L détermine certaines de ses propriétés. On brise la symétrie du système en connectant l'électrode de source au potentiel le plus bas (V_{SS}). C'est-à-dire à l'électrode de bulk dans la figure 15. Cette connexion permet d'imposer le sens du champ électrique qui sera généré par la polarisation de la grille.

Pour la petite histoire, les noms des électrodes : grille, drain, source sont issus de l'analogie de fonction de ces électrodes avec les composants électroniques à lampes tel que l'Audion de De Forest (1906). La grille est l'électrode de contrôle de ce type de transistor et les électrodes de drain et de source constituent la voie principale. Par définition le courant circule du drain vers la source. C'est cohérent avec la condition imposée : $V_D > V_S$, ou V_D et V_S sont respectivement les potentiels de Drain et de Source.

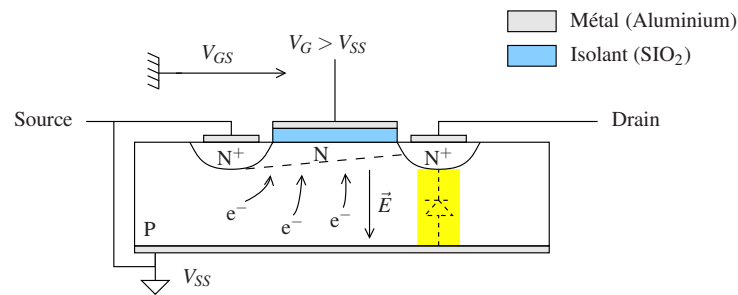


La représentation en perspective est intéressante pour donner une idée de la structure d'un MOS. Elle permet, en outre, de définir les dimensions W et L . Elle n'est cependant pas la plus pratique pour en étudier le fonctionnement. Dans la suite on lui préférera une représentation en coupe le long de sa dimension (L) comme ci contre :



Fonctionnement : Si on laisse la grille “en l'air”, c'est-à-dire non connectée, entre source et drain on voit une suite de semi-conducteurs NPN c'est-à-dire deux diodes tête-bêches et le courant ne peut pas passer : Source $\rightarrow$ Drain

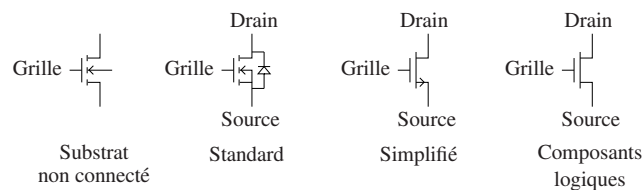
Si on applique un potentiel $V_G > V_{SS}$ sur la grille, alors on crée un champ électrique entre la grille et le substrat qui attire les électrons, porteurs minoritaires, présents dans le substrat P. Cela crée un canal dopé en électrons, donc N, entre les caissons de source et de drain comme représenté dans la figure ci dessous :



Le courant peut alors passer du drain vers la source, le transistor est passant. Comme le dopage du canal est de type N ce type de transistor est appelé NMOS.

Remarque : On notera dans cette situation la présence d'une diode entre drain et substrat. Avec l'électrode de substrat connectée à la masse, cette diode empêche de connecter le transistor à l'envers, c'est-à-dire avec $V_D < V_S$.

Symboles : On trouve de nombreux symboles pour le transistor NMOS dans la littérature et les documentations techniques. Les plus courants sont présentés ci dessous :



Tous les symboles ont en commun de montrer l'isolation de grille : l'électrode de grille n'est pas en contact dans les symboles avec le reste du composant.

Le premier symbole est très général, il représente les trois électrodes grille, drain et source. De plus, la connexion au substrat est représentée comme une électrode additionnelle au milieu. La flèche indique le sens de la diode substrat-canal. De cette façon, on pourra distinguer un MOS canal N d'un MOS canal P que l'on verra plus tard. Ce symbole représente un MOS dont le substrat n'est pas connecté à la source. Il n'est utilisé que dans le cadre de la micro électronique, par les fondeurs de silicium, seuls susceptibles de placer plusieurs MOS côte à côte sur un même substrat.

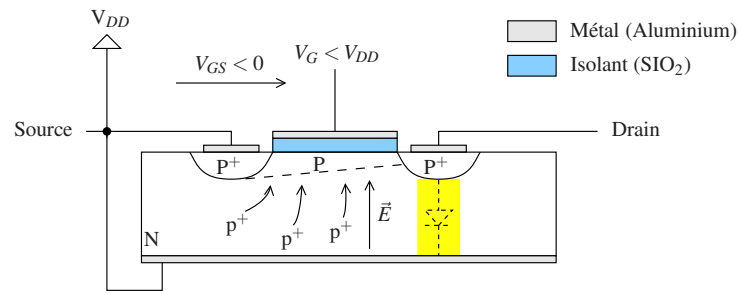
Le deuxième symbole est le symbole courant des NMOS comme composant unique dans les documentations des fabricants (exemple BS170) . Comme il n'y a qu'un seul MOS, le substrat est connecté à la source ce que fait apparaître le symbole. De plus la diode drain substrat est représentée. C'est une information redondante dans le symbole, mais qui peut aider à se souvenir du sens de circulation du courant, cette diode doit être bloquée sinon elle court-circuite le transistor.

Le troisième symbole, est un symbole simplifié très répandu, la flèche indique le sens d'écoulement normal du courant et est placée du côté source. Ce symbole est celui que nous utiliserons en général.

Le dernier symbole est celui souvent utilisé pour les schémas des composants logiques. Drain et source n'y sont pas distingués, mais dans ce contexte la source est nécessairement en bas, et le NMOS est passant pour un niveau haut sur sa grille . Ce symbole se comprend en considérant le symbole du MOS à canal P qui dans le même contexte présente un petit rond sur la grille pour signifier qu'il est passant pour un niveau bas sur sa grille (CF symboles PMOS).

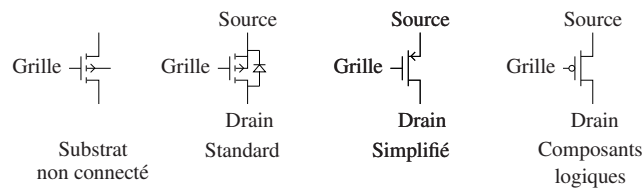
8.2 MOS canal P (PMOS)

Comme nous l'avons évoqué plus tôt, on peut aussi fabriquer des PMOS :



Dans la structure d'un PMOS, la source est cette fois-ci connectée au potentiel le plus élevé V_{DD} . Les caissons sont dopés P et le substrat N. Avec $V_{GS} < 0$ cette fois on attire les trous porteurs minoritaires du substrat entre les deux caissons de grille et de drain pour créer un canal P. Dans cette structure, la diode drain substrat est orientée cette fois ci du drain vers la source. En conséquence, dans un PMOS, le courant circule dans le canal de la source vers le drain.

Avec les mêmes principes que pour les symboles du transistor NMOS, les symboles du PMOS sont les suivants :



Dans ces schémas les sources sont placées en haut pour respecter la convention qui consiste à placer les potentiels les plus élevés vers le haut. On peut noter, dans les deux symboles du milieu, le changement de sens de la diode de substrat. Pour le symbole simplifié, la flèche qui indique le sens normal d'écoulement du courant est toujours placée sur la source. On notera que comme, dans un PMOS, le courant circule de la source vers le drain, la flèche pointe cette fois-ci vers le transistor. Le symbole utilisé pour les composants logiques présente comme mentionné plus haut un petit rond sur la grille qui rappelle qu'il est passant pour un niveau bas.

Intégration de plusieurs MOS sur un même substrat : Dans les coupes de transistors MOS présentées ci-dessus, la source était toujours connectée au substrat. Ce n'est évidemment pas toujours possible lorsque que l'on intègre plusieurs MOS sur le même substrat et qu'ils sont connectés en série. C'est encore moins possible si l'un des MOS est un NMOS et l'autre un PMOS.

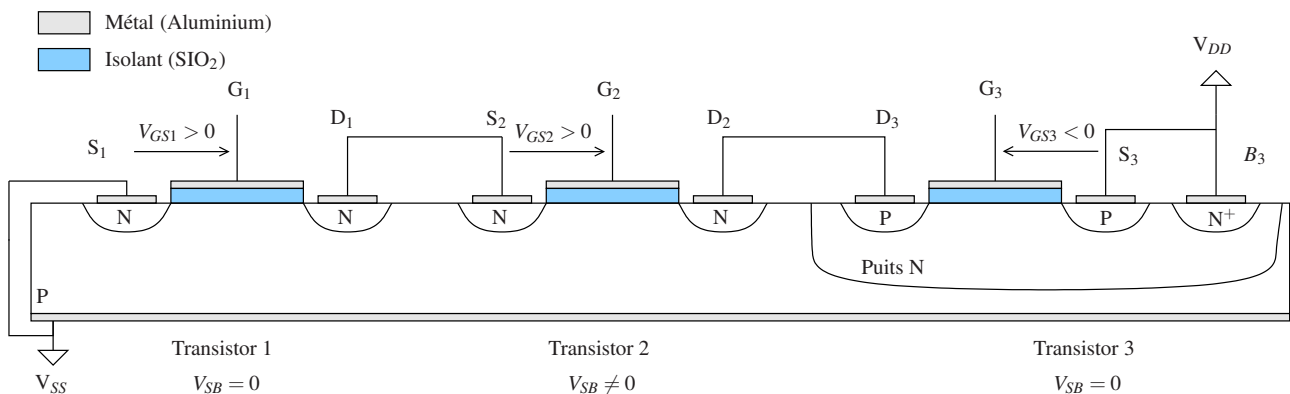


FIGURE 16 – Exemple d'intégration de trois MOSFET canal N et P sur un même substrat.

La figure ci-dessus montre, par exemple, comment il est possible de placer et de connecter en série trois transistors (N, N et P) sur le même substrat. Pour intégrer le transistor PMOS un puits P a été réalisé. Notez l'électrode supplémentaire B_3 qui est l'électrode de substrat du transistor PMOS. Elle permet de connecter au potentiel le plus élevé le substrat du PMOS. On vérifie que le sens possible pour l'écoulement du courant est bien, du drain vers la source, pour les transistors

NMOS et de la source vers le drain pour le PMOS. Enfin, on peut noter que la différence de potentiel grille-source du deuxième transistor n'est pas nulle. Cela influence bien sûr son fonctionnement.

Applications : Lorsqu'ils sont utilisés comme composants uniques et discrets, les MOS trouvent la plupart de leurs applications en commutation.

On peut citer les deux références suivantes de transistors MOS à tout faire : BS170 (NMOS), BS250 (PMOS).

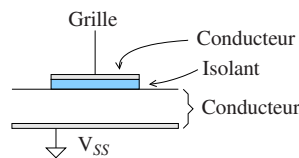
En électronique intégrée, le MOS est actuellement le composant roi. C'est le composant de tous les circuits logiques et numériques et des micro-processeurs. Il est aussi à la base des étages d'entrée des amplificateurs opérationnels modernes. En ce qui concerne les applications numériques, la technologie MOS est caractérisée par sa largeur de grille. En 2014, pour l'Intel Core I7 par exemple, elle est de 45 nm pour 10^9 transistors intégrés.

8.3 Caractéristiques statiques des transistors NMOS, modèle de Schichman et Hodges

En fonction de la tension appliquée sur la grille, le transistor est soit bloqué, soit passant si l'on a créé un canal. Ce modèle simpliste est suffisant pour expliquer ce qu'est un transistor et pour en comprendre le fonctionnement, mais il ne permet pas de calculer la valeur des composants qui vont entourer le MOS dans un circuit. Si le MOS est passant, qu'elle est sa résistance ? C'est une des premières questions que l'on peut se poser. La physique précise d'un transistor MOS est complexe, mais bien connue. Elle est hors du champ de ce cours et dans celui de physique du solide. Nous en utiliserons les résultats nécessaires sans les démontrer.

8.3.1 Courant de grille

Vue de la grille, un MOS est un condensateur dont la valeur dépend de W et L et de l'épaisseur de l'isolant (SiO_2) :



L'oxyde de silicium est un très bon isolant. On peut considérer pour la caractéristique statique (à fréquence nulle) que le courant de grille est nul :

$$I_G = 0 \quad (15)$$

L'ordre de grandeur de la résistance de fuite est en effet de l'ordre de grandeur de quelque TeraOhm ($10^{12} \Omega$).

8.3.2 Courant drain-source

Dans un premier temps, nous nous plaçons dans l'hypothèse $V_{SB} = 0$ pour établir la caractéristique du NMOS. Nous verrons comment la valeur de la tension source-substrat la modifie par la suite.

Transistor Bloqué : Comme nous l'avons vu dans la section précédente, on cherche à emmener des électrons porteurs minoritaires entre la source et la grille pour créer un canal de conduction. Ces électrons ne bougent pas librement dans un substrat P. Le canal ne peut se créer qu'à partir d'une certaine valeur de la tension V_{GS} que l'on appelle tension de seuil : V_T . Il est bloqué tant que :

$$V_{GS} < V_T \quad (16)$$

Le courant de drain est donc nul quelque soit la tension appliquée entre le drain et la source (zone bloquée figure 17).

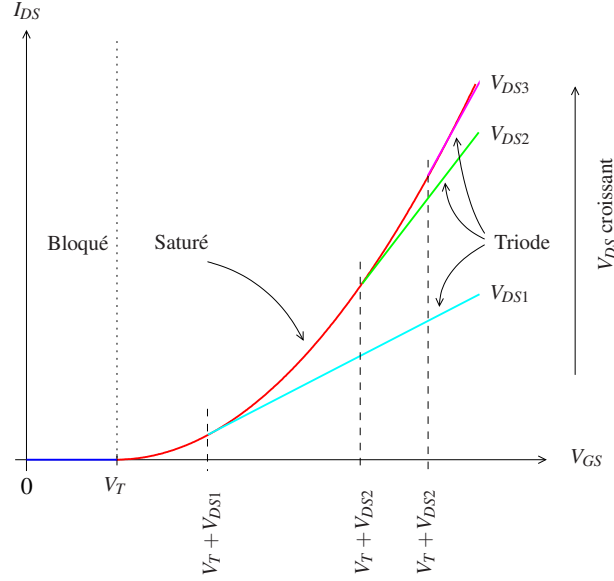
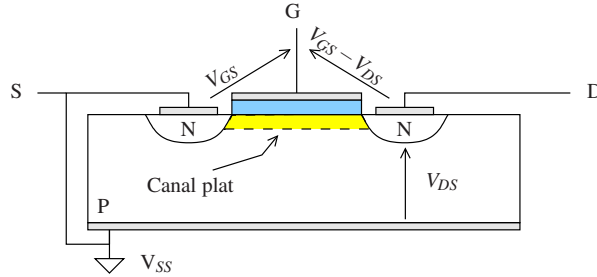


FIGURE 17 – Caractéristique NMOS : I_{DS} en fonction de V_{GS} .

Transistor passant : V_{GS} est maintenant supérieure à la tension de seuil, le canal existe. Il est alors possible d'appliquer une différence de potentiel V_{DS} pour y faire circuler du courant.

Mode Ohmique :

Si V_{DS} est suffisamment petite pour que $V_{GS} \approx V_{GS} - V_{DS}$ (c'est-à-dire $V_{GS} \approx V_{GD}$) alors le champ électrique qui règne dans la structure est homogène et le canal est plat, comme représenté ci-dessous.



L'épaisseur ep du canal est proportionnelle à $V_{GS} - V_T$: $ep = \alpha(V_{GS} - V_T)$, où α dépend de la mobilité des porteurs. Le courant I_{DS} circule donc dans une zone de longueur L et de section $S = \alpha(V_{GS} - V_T)W$ contenant des porteurs libres. Il se comporte donc comme une résistance de valeur $R = \rho L/S$, où ρ est la conductivité du canal. On en déduit avec la loi d'Ohm le courant drain-source : $\frac{\alpha}{\rho} \frac{W}{L} (V_{GS} - V_T) V_{DS}$. En posant $k = \frac{\alpha}{\rho}$, que l'on appelle le facteur de gain du transistor, on obtient :

Courant mode Ohmique

$$I_{DS} = k \frac{W}{L} (V_{GS} - V_T) V_{DS} \quad (17)$$

En remarquant que la condition $V_{GS} \approx V_{GS} - V_{DS}$ peut s'exprimer comme $V_{DS} \ll V_{GS}$, c'est-à-dire V_{DS} petit, on se trouve dans la zone ohmique de la caractéristique $I_{DS} = f(V_{GS})$ présentée figure 18.

k est généralement exprimé en $\mu A/V^2$. On peut montrer que :

$$k = \mu_{e^-} \cdot C_{ox} \quad (18)$$

Dans cette équation μ_{e^-} est la mobilité des électrons et $C_{ox} = \frac{\epsilon_{ox}}{d_{ox}}$ la capacité par unité de surface de la grille. d_{ox}

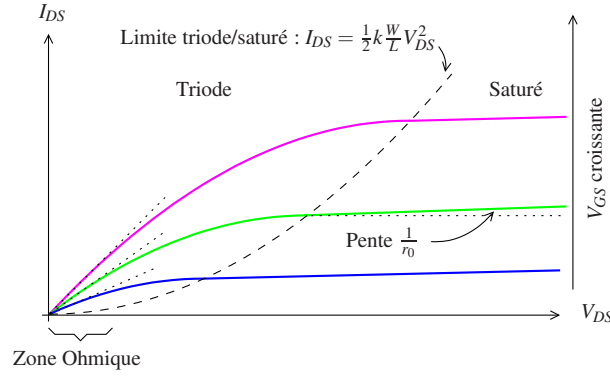


FIGURE 18 – Caractéristique NMOS : I_{DS} en fonction de V_{DS} .

est l'épaisseur de la couche isolante d'oxyde de Silicium et ϵ_{ox} sa permittivité (condensateur plan $C = \epsilon \frac{S}{d}$). Le lecteur vérifiera qu'avec la définition de la mobilité : $\vec{v}_{(m/s)} = \mu \vec{E}$, k est bien homogène à des Ampères par Volts carrés.

Mode Triode :

V_{DS} augmente. La condition $V_{GS} \approx V_{GD}$ n'est plus respectée. Le canal n'est plus symétrique (figure 19-c), il est plus épais du côté de la source où la différence de potentielle V_{GS} est plus grande que du côté du drain $V_{GD} = V_{GS} - V_{DS}$. Elle correspond en effet à un champ électrique plus intense. Le canal se pince (figure 19-d), c'est-à-dire que son épaisseur devient nulle lorsque V_{GD} devient égale à la tension de seuil : $V_{GD} = V_{GS} - V_{DS} = V_T$. On reste en régime triode tant que $V_{GD} > V_T$ soit tant que $V_{GS} - V_{DS} > V_T$. Finalement, on est en régime triode tant que : $V_{DS} < V_{GS} - V_T$.

Le modèle de Shichman et Hodges (S-H) donne le courant de drain pour le mode triode :

Mode triode

$$V_{DS} \leq V_{GS} - V_T \quad (19)$$

$$I_{DS} = k \frac{W}{L} \left[(V_{GS} - V_T) V_{DS} - \frac{V_{DS}^2}{2} \right] \quad (20)$$

Remarque : Lorsque V_{DS} est suffisamment petite, $V_{DS}^2 \ll (V_{GS} - V_T) V_{DS}$ et l'on retrouve l'expression du courant de mode Ohmique. Le mode Ohmique est un cas particulier du mode triode.

La figure 18, présente les modes ohmique et triode ainsi que la courbe qui sépare le mode triode du mode saturé qui suit. Cette courbe correspond à la limite $V_{DS} = V_{GS} - V_T$. L'équation de cette courbe s'obtient en remplaçant $(V_{GS} - V_T)$ par V_{DS} dans l'équation 20 :

Limite du mode triode

$$I_{DS} = \frac{1}{2} k \frac{W}{L} V_{DS}^2 \quad (21)$$

Mode saturé, modulation de la longueur de canal, effet Early :

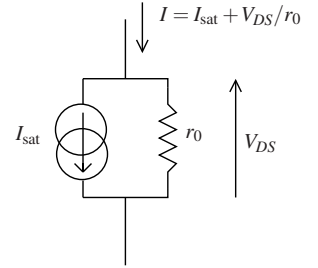
Lorsque le canal est juste pincé, le courant n'est contrôlé que par la zone de pincement. Le courant transistor est dit saturé car I_{DS} est alors quasi indépendant de V_{DS} . Il est donné par le modèle de S-H :

Courant de saturation

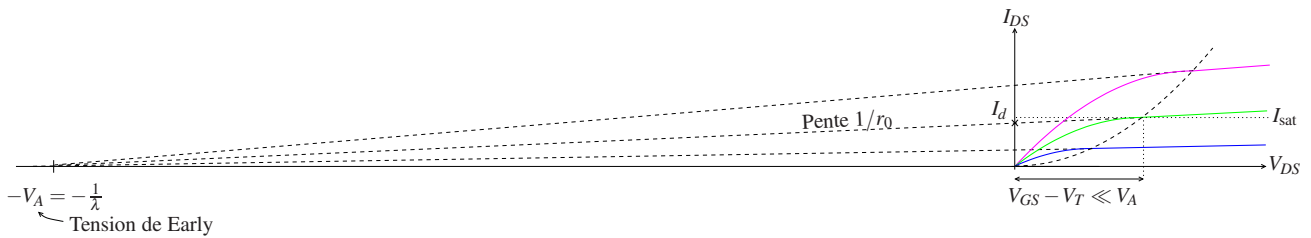
$$I_{DS} = I_{sat} = \frac{1}{2} k \frac{W}{L} (V_{GS} - V_T)^2 \quad (22)$$

I_{sat} ne dépend pas de V_{DS} . Le transistor se comporte comme une source de courant commandée par V_{GS} .

Bien sur, une source de courant idéale n'existe pas et I_{DS} dépend légèrement de V_{DS} comme présenté sur la courbe du milieu de la figure 18. Cet effet est simplement modélisable par une grande résistance en parallèle avec un générateur de courant fournissant I_{sat} . Cette résistance est en fait la résistance interne du générateur. Cet effet vient de la variation de la longueur du canal avec l'augmentation de V_{DS} comme présenté figure 19-e.



En effet, le canal existe car $V_{GS} > V_T$ et le courant est contrôlé par la zone de pincement. La zone pincée s'étend avec V_{DS} , ce qui se traduit par une légère augmentation du courant de saturation, qui est modélisable par une résistance en parallèle. Ce n'est là qu'un modèle de comportement. Le phénomène exact n'est pas simple.



Si l'on prolonge les droites correspondant aux zones saturées dans la caractéristique $I_{DS} = f(V_{DS})$ pour différentes valeurs de V_{GS} comme dans la figure ci-dessus, on constate qu'elles passent toutes par le point $(-V_A, 0)$. V_A est appelé tension de Early. Elle est proportionnelle à L la longueur du canal. $\lambda = \frac{1}{V_A}$ est appelé coefficient de modulation de la longueur du canal.

La résistance r_o dépend de V_{GS} : $r_o = \frac{\partial I_{DS}}{\partial V_{DS}}|_{V_{GS}=\text{cste}} = \frac{V_A}{I_d} = \frac{1}{\lambda I_d}$. Comme la valeur de r_o est très grande, $I_d \approx I_{sat}$ soit $r_o \approx \frac{1}{\lambda I_{sat}}$. Avec la même hypothèse on peut négliger $(V_{GS} - V_T) \frac{1}{r_o}$ devant $(I_{sat} - I_d)$, on écrit : $I_{DS} = I_{sat} + \frac{V_{DS}}{r_o}$. Finalement, on modélise l'effet de la variation de la longueur du canal par l'équation 23

Modélisation de l'effet Early

$$I_{DS} = \frac{1}{2} k \frac{W}{L} (V_{GS} - V_T)^2 (1 + \lambda V_{DS}) \quad (23)$$

Dans la plupart des cas, cet effet est négligeable (r_o très grande) et on n'en tiendra en général pas compte. En pratique on n'en tient compte que si on arrive à un résultat incohérent ou impossible sans en tenir compte.

ON UTILISE TOUJOURS LE MODÈLE LE PLUS SIMPLE POSSIBLE

8.4 Validité du modèle

Le modèle présenté est valable au premier ordre, il est donc relativement grossier. Il permet cependant les raisonnements et les calculs à la main à $\pm 10\%$ près. Lorsqu'une précision supérieure est nécessaire, ce qui n'est en général pas le cas pour l'électronique de tous les jours, il faut avoir recours aux logiciels de simulation.

8.4.1 Paramètres d'influences

Effet de substrat : Dans ce qui précède, le potentiel du substrat est supposé être le même que celui de la source : $V_{SB} = 0$. Or cela n'est parfois pas le cas, comme pour le transistor du milieu de la figure 16. Dans ce cas le fonctionnement du NMOS est légèrement modifié. Il ne s'agit cependant que d'un ajout au modèle précédent, ce dernier reste vrai. Le modèle corrigé est peu utile en pratique, et ne sera pas présenté ici.

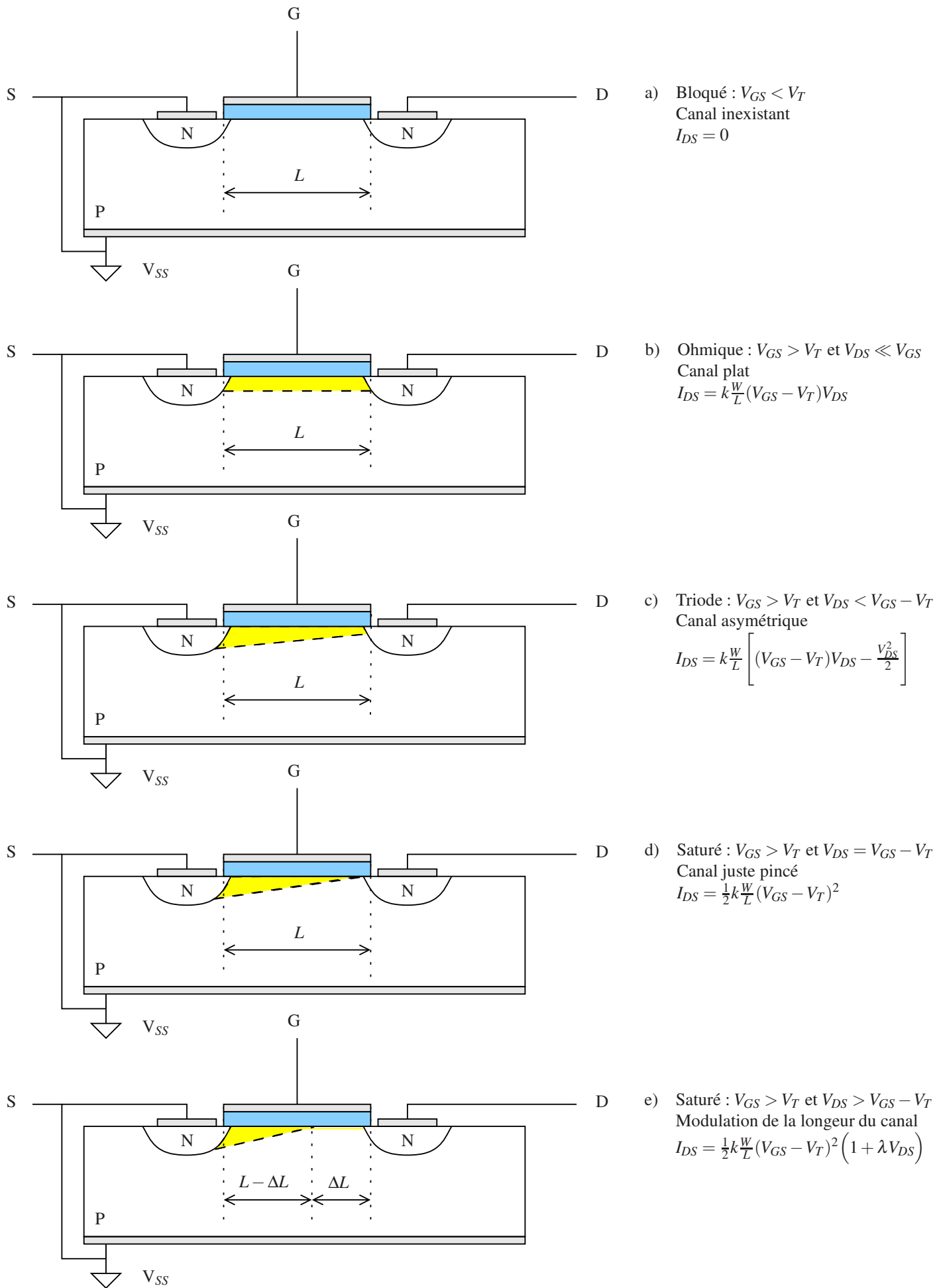


FIGURE 19 – Résumé du fonctionnement d'un transistor NMOS.

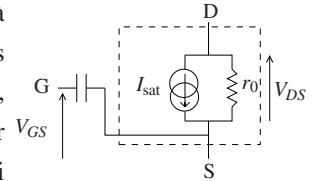
Effets de la température : Deux effets sont ici en compétition.

1. $\frac{\partial V_T}{\partial T} < 0$, la température aide à “décrocher” les porteurs minoritaires de réseau cristallin, ce qui fait augmenter I_{DS}
2. $\frac{\partial k}{\partial T} < 0$, la température augmente la résistivité du canal ce qui a tendance à diminuer I_{DS} .
3. Cet effet est prépondérant.

Finalement, I_{DS} diminue avec la température

8.5 Stratégies de calcul : schémas équivalents

Comme pour les diodes, il n’est pas toujours facile lorsque que l’on se trouve devant un circuit comprenant un transistor, de trouver le bon angle d’attaque. La bonne méthode consiste toujours à déterminer dans quel mode (Bloqué, Ohmique,..., ou saturé) se trouve le MOS. Pour les transistors MOS, cela se fait en déterminant V_{GS} puis, si $V_{GS} > V_T$, en comparant V_{GS} à $V_{DS} - V_T$. Ensuite, on remplacera le transistor par le schéma équivalent correspondant à son état. Pour un transistor saturé, on le remplace par le schéma ci-contre, qui fait apparaître un générateur de courant qui correspond au courant de saturation.



Le schéma équivalent doit faire apparaître toutes les électrodes d’entrée du MOS. Dans ce schéma le condensateur modélise la capacité grille-source. I_{DS} ne dépend pour un MOS saturé que de V_{GS} . La capacité grille-source dépend de la surface de l’électrode de grille $W \times L$. Elle est responsable des performances en terme de vitesse du transistor. Il faut en effet la charger ou la décharger pour modifier V_{GS} . L’amélioration des performances des MOS est actuellement liée à la diminution conjointe de la surface de l’électrode de grille et de la tension d’alimentation V_{DD} , ce qui permet de diminuer la quantité de charge nécessaire à la polarisation de la grille (CF table 2). Pour un NMOS en régime Ohmique, le schéma équivalent, serait celui d’une résistance pilotée par la valeur de V_{GS} . A l’aide du schéma équivalent, on peut commencer à “raisonner électronique”, c’est-à-dire à utiliser les lois d’Ohm, des nœuds, des mailles, et le théorème de superposition.

Année	W	V_{DD}
1974	3 μm	5 V
1995	0.35 μm	3.3 V
1999	0.18 μm	1.8 V
2007	65 nm	1.2-0.8 V

TABLE 2 – Évolution des tailles de gravures et tension d’alimentation en technologie MOS

8.6 Caractéristique statique des transistors PMOS.

Les transistor PMOS se comportent de la même façon que les NMOS au signe des tensions près ($V_{GS} < 0$ et $V_{DS} < 0$) et au sens du courant ($I_{DS} < 0$) puisqu’il circule de la source vers la grille comme présenté section 8.2. Pour créer le canal, il faut $V_{GS} < V_T$ avec $V_T < 0$. Un fois le canal créé, si l’on augmente $|V_{DS}|$ le canal est successivement plat, juste pincé, pincé avec modulation de la longueur. Le transistor est alors successivement en mode ohmique, triode, puis saturé comme le NMOS. Les équations du modèle de S-H permettent, là encore, de quantifier le courant. La figure 20 résume, en le comparant au NMOS, le fonctionnement des transistors PMOS. Notons que en général $|V_{Tp}| \neq |V_{Tn}|$. Ceci n’apparaît pas dans la figure 20 pour des raisons de clarté.

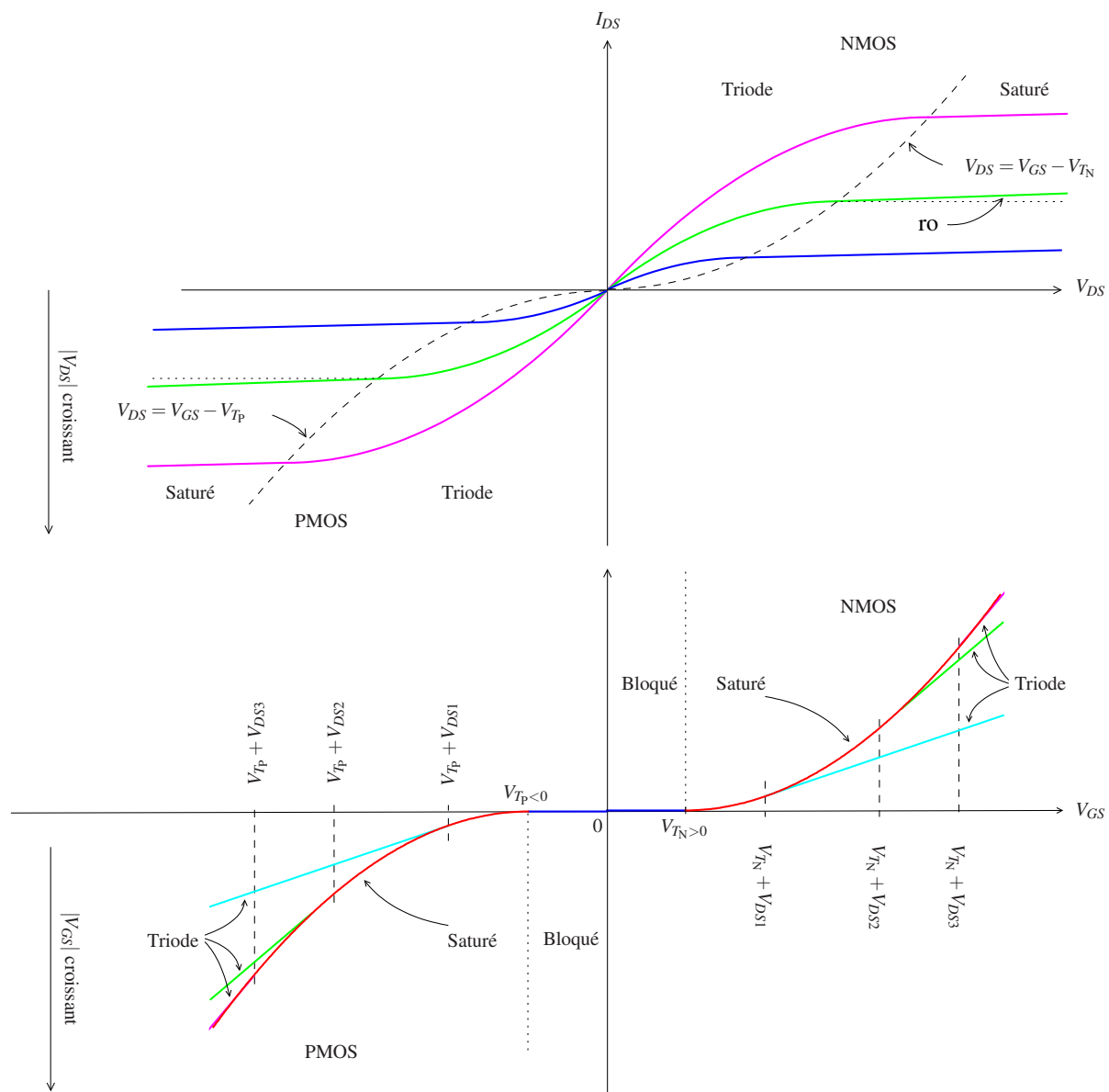
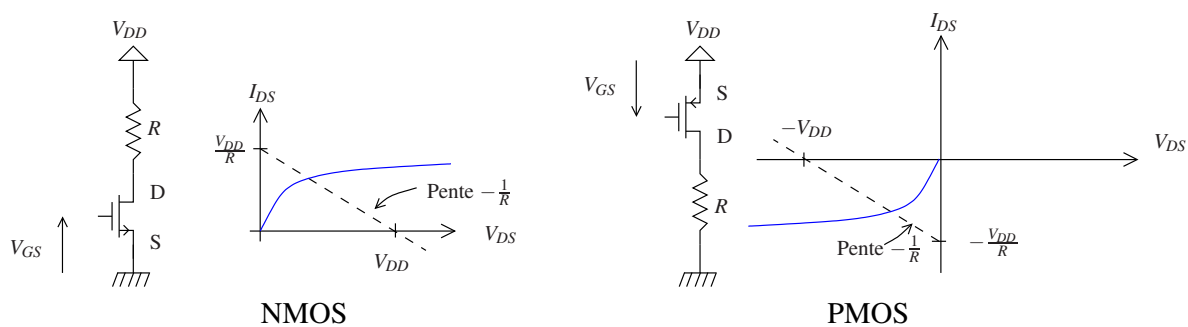


FIGURE 20 – Comparaisons des caractéristiques PMOS et NMOS.

8.6.1 Mises en œuvre comparées des transistors NMOS et PMOS.

Le fonctionnement des NMOS est contrôlé par la tension V_{GS} avec $V_G > V_S$, alors que pour les PMOS $V_G < V_S$. Ils doivent donc être mis en œuvre de telle sorte que ces conditions soient respectées. Cela se fait typiquement de la façon suivante :



la tension V_{GS} peut ainsi être directement imposée aux transistors pour contrôler leurs blocages respectifs. Notez que

R permet d'amener le transistor en mode ohmique, triode ou saturé, en fonction de sa valeur, en imposant la tension V_{DS} pour une tension V_{GS} donnée, comme le montrent les droites de charge.

8.7 Transistors MOS à appauvrissement

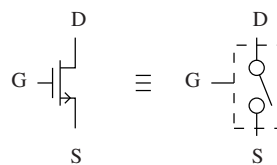
Pour être à peu près exhaustif sur les transistor MOS, il faut citer les transistors à appauvrissement. Les MOS P ou N que nous avons décrits jusqu'à présent sont des transistors dits à enrichissement car il faut créer le canal, c'est-à-dire l'enrichir en porteurs. Les transistors MOS à appauvrissement présentent un canal pré-dopé. Le transistor est donc passant lorsque aucune tension n'est appliquée sur la grille. Il faut venir chasser les porteurs du canal pour le bloquer. Ainsi pour un transistor à canal N il faut appliquer une tension $V_{GS} < 0$ pour chasser les porteurs majoritaires du canal pré-dopé et bloquer le transistor.

La caractéristique statique $I_{DS} = f(V_{GS})$ de ce type de transistors est grosso-modo la même que celle des transistors à enrichissement mais décalée, vers la gauche pour un NMOS ou vers la droite pour un PMOS, d'une tension constante qui dépend du taux de dopage et de la profondeur du canal à appauvrir.

9 Applications des transistors MOS

9.1 MOS en Commutation

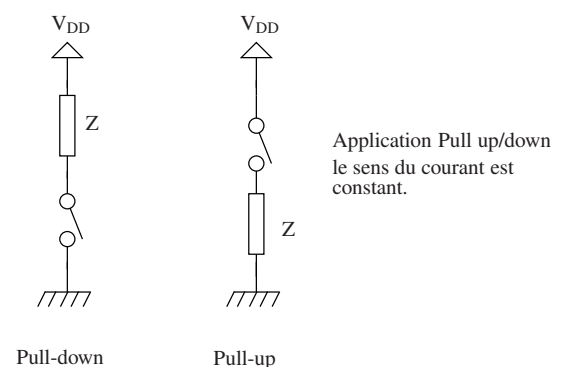
Lorsqu'il est utilisé en commutation, on cherche à le faire se comporter comme un interrupteur commandé :



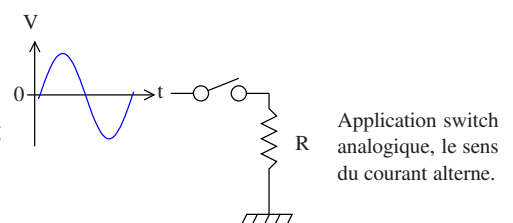
Idéalement il devrait présenter deux états : ON et OFF (passant et isolant). L'état OFF correspond simplement au blocage du transistor. L'état passant est plus compliqué. En effet pour un interrupteur l'état passant correspond à un état où il se comporte comme une résistance. Il faut donc pour cette application, que le MOS soit en mode Ohmique. De plus dans un interrupteur, le courant peut circuler dans les deux sens ce qui n'est pas le cas pour un MOS.

Il y a donc deux applications du MOS en commutation :

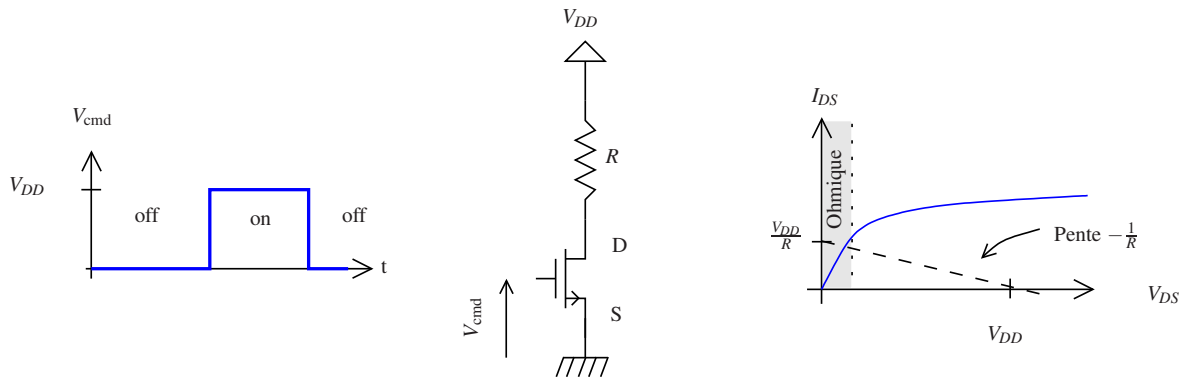
1. Le MOS dans l'état ON laisse toujours passer le courant dans le même sens. Ce sont les applications de "pull-up" ou de "pull-down" dans lesquelles on tire la tension vers le haut ou vers le bas. On utilise respectivement un NMOS et un PMOS :



2. Le MOS dans l'état ON doit laisser passer le courant dans les deux sens. Ce sont les applications d'interrupteurs analogiques (analog switches), et l'on est obligé d'utiliser deux MOS complémentaires.



9.1.1 NMOS en pull-down



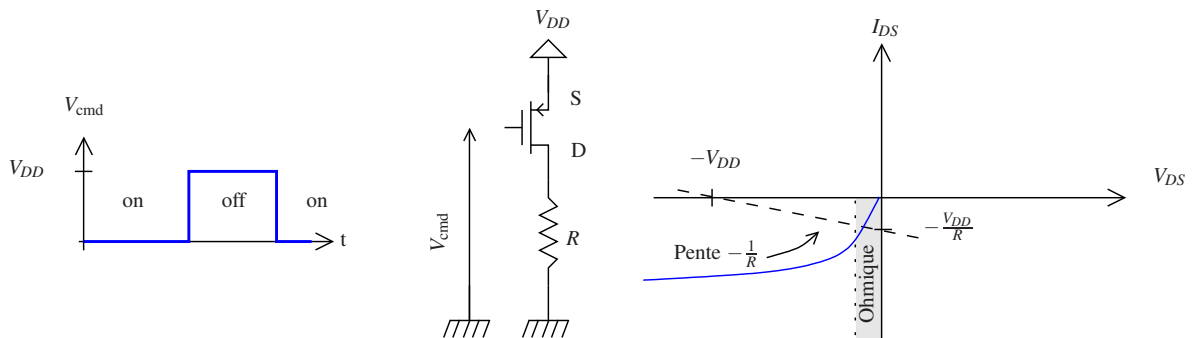
Pour être en mode ohmique, il faut que $V_{GS} \gg V_{DS}$. On commande donc le MOS avec la tension $V_G = V_{GS}$ la plus élevée possible : V_{DD} . De plus, comme on peut le voir dans la figure ci dessus, il faut que la résistance soit suffisamment grande pour rester en mode Ohmique. A ces conditions, la résistance passante du MOS vaut : $r_{on} = \frac{1}{k \frac{W}{L} (V_{DD} - V_T)}$.

Les documentations des composants (*Data Sheets*), précisent directement, l'ordre de grandeur de r_{on} , pour une valeur de V_{DD} et un courant donné. Ce courant est un courant maximum, ce qui revient à préciser la valeur minimale de R ($R > \frac{V_{DD}}{I_{max}}$).

Exemple : BS170 ; $r_{on} = 1,2 \Omega$ typique, (5Ω max) pour $V_{DD} = 10V$ et $I_{max} = 200 \text{ mA}$.

Remarque : On a en outre intérêt à prendre $R \gg r_{on}$ pour que la tension de drain soit la plus petit possible lorsque le MOS est passant. On se rapproche ainsi de l'interrupteur idéal.

9.1.2 PMOS en pull-up



Les principes sont les mêmes que pour le NMOS en pull_down.

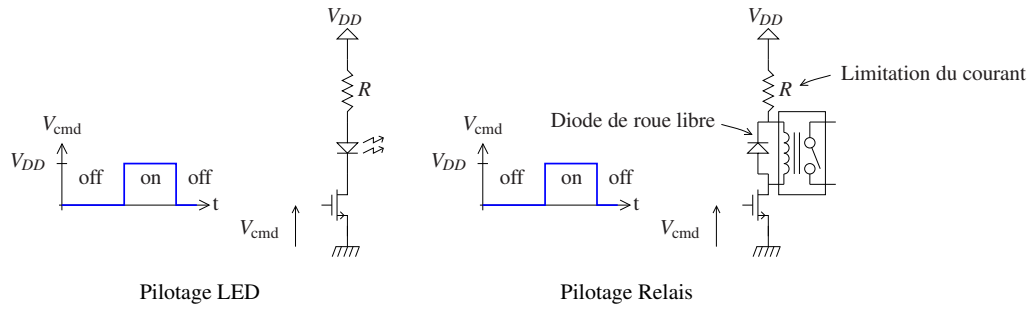
Attention cependant, ici $V_G = V_{DD} + V_{GS}$, le MOS est donc piloté de façon complémentaire par rapport à l'application pull-Down.

Exemple : BS250 ; $r_{on} = 8 \Omega$ typique, (14Ω max) pour $V_{DD} = 10V$ et $I_{max} = 200 \text{ mA}$.

On remarque que la valeur de r_{on} est supérieur pour le PMOS. Ce n'est pas général. A géométrie comparable, les résistances passantes des PMOS sont plus faibles que celles des NMOS. Cela est dû aux différences de facteur de gain des transistors de type N ou P. Il faut se souvenir que les porteurs majoritaires sont ici des trous moins mobiles que les électrons.

9.1.3 Exemples d'application d'un NMOS en pull-down. : Pilotage d'une LED ou d'un relais par une porte logique.

Le courant maximum que peuvent fournir les circuits logiques (portes, CPLD FPGA, micro_contrôleurs) est en général trop faible pour piloter une LED et encore moins un relais. On peut alors utiliser un NMOS par exemple :



Dans les deux cas on notera la présence de la résistance qui permet en le contrôlant de limiter le courant : $I \approx \frac{V_{DD}}{R}$, ainsi que la présence de la diode de roue libre dans le cas du relais qui permet d'éviter les claquages à la rupture du courant lors du passage en mode "off".

9.2 Interrupteurs analogiques (*Analog switches*)

Comme nous l'avons déjà évoqué plus haut, il n'est pas possible avec un seul transistor de type N avec la source connectée au bulk et à V_{SS} , ou de type P avec la source connectée à V_{DD} de réaliser un switch analogique. Cependant, si l'électrode de Bulk n'est pas connectée à la source comme présenté figure 21, les transistors NMOS et PMOS sont des composants symétriques. Le courant peut passer dans les deux sens. La source et le drain du transistor ne sont alors déterminés que par le potentiel appliqué sur chacune des électrodes.

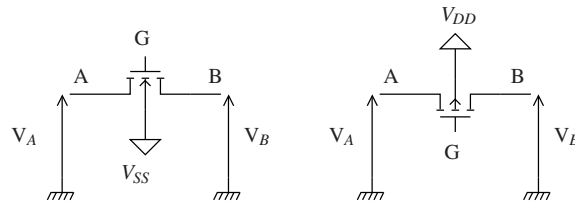


FIGURE 21 – Transistors NMOS et PMOS utilisés en interrupteurs analogiques

Dans les conditions normales de mise en œuvre : $V_{DD} > V_A > V_{SS}$ et $V_{DD} > V_B > V_{SS}$ et le potentiel de référence (la masse) est le point milieu de l'alimentation $\frac{V_{DD}+V_{SS}}{2}$. En général on prend $V_{SS} = -V_{DD}$.

Pour le NMOS :

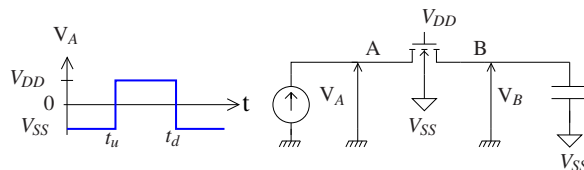
Si $V_A > V_B$, le courant circule de A vers B, l'électrode A joue le rôle du drain et l'électrode B joue le rôle de la source et inversement.

Le NMOS est bloqué si $V_G - V_{SS} < V_T$ le canal n'existe pas. En effet comme ni l'électrode de droite ni l'électrode de gauche ne sont connectées à la masse, c'est la différence de potentiel $V_G - V_{SS}$ qui crée le champ électrique responsable de l'apparition du canal. Toutes les tensions doivent être comparées à V_{SS} pour déterminer le fonctionnement du NMOS en utilisant le modèle de S-H. On admettra ici que l'on peut l'utiliser sans trop faire d'erreurs.

Blocage : Pour créer un canal le plus grand possible, on fixe $V_G = V_{DD}$. Le NMOS est passant.

Il est alors possible d'étudier le comportement du NMOS selon que $V_A > V_B$ et inversement.

Pour cela on étudie la réponse du NMOS chargé par un condensateur pour voir l'évolution de la tension de sortie dans le temps.



On suppose le condensateur initialement déchargé

1. $t = t_u^-$, $V_{AB} = 0$ aucun courant ne circule dans le MOS.
2. $t = t_u^+$; $V_A > V_B \rightarrow$ la source est du coté droit.

$$\begin{aligned} V_{GS} - V_T &= (V_{DD} - V_{SS}) - V_T \implies V_{DS} > V_{GS} - V_T, \text{ Le transistor est saturé, le condensateur de sortie se} \\ V_{DS} &= V_{DD} \end{aligned}$$
charge à courant constant $V_B(t) = \frac{1}{C} I_{sat} t$.
3.
$$\begin{aligned} V_{GS} - V_T &= (V_{DD} - V_B(t)) - V_T \\ V_{DS} &= V_{DD} - V_B(t) \end{aligned}$$
, V_{DS} reste toujours supérieur à $V_{GS} - V_T$ le transistor ne passe jamais en mode ohmique. La résistance passante apparente (Pente de la caractéristique $I_{DS} = f(V_{DS})$) est très grande.
4. Lorsque $V_B(t)$ atteint $V_{DD} - V_T$, le transistor se bloque. On n'a jamais atteint le mode ohmique !
 $V_{GS} = V_{DD} - V_B$. Lorsque $V_B = V_{DD} - V_T$, le transistor reste bloqué comme présenté figure 22.
5. $t = t_d^+$, $V_A = V_{SS}$, $V_B = V_{DD} - V_T$, $V_B > V_A$. La source est maintenant à gauche.

$$\begin{aligned} V_{GS} - V_T &= (V_{DD} - V_{SS}) - V_T \\ V_{DS} &= V_{DD} - V_T - V_{SS} \end{aligned}$$
Le MOS est à la limite de la saturation. Le condensateur commence à se décharger. $V_{GS} = \text{cste}$ et V_{DS} diminue. Le MOS passe en mode triode puis Ohmique. La résistance apparente diminue fortement. Elle devient celle du canal en mode ohmique

Ce comportement est résumé figure 22.

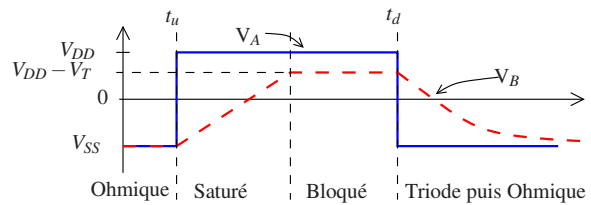


FIGURE 22 – Réponse à un pulse d'un NMOS monté en *switch*

On en déduit que le NMOS se comporte comme une résistance pour $V_A - V_B < 0$. Il ne conduit pas bien le courant dans l'autre sens.

— Pour le PMOS

Si $V_A > V_B$, le courant circule de A vers B, l'électrode A joue le rôle de la source et l'électrode B joue le rôle du drain et inversement.

Son comportement est symétrique de celui du NMOS. Le lecteur pourra vérifier que le PMOS se comporte comme une petite résistance pour $V_A - V_B > 0$, et qu'il ne conduit pas bien le courant dans l'autre sens.

On peut résumer ces comportements en terme de résistance passante apparente r_{on} et l'on obtient l'allure présentée figure 23. Le NMOS et le PMOS présentent une faible résistance passante chacun dans leurs domaines respectifs de tension ; On les associe donc en parallèle pour obtenir une résistance passante faible et plus constante avec $V_A - V_B$.

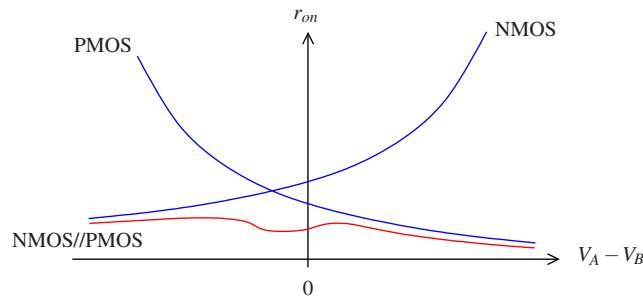
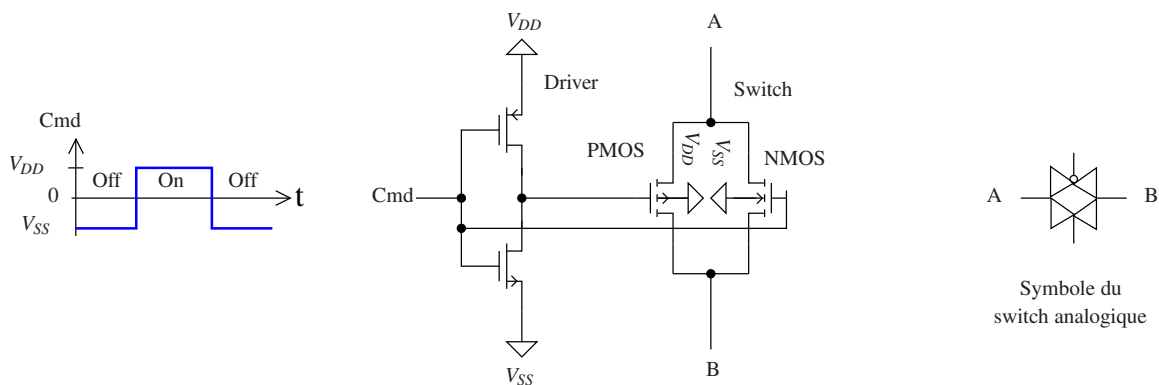


FIGURE 23 – Résistance passante des MOSFET en fonction du signe de la tension à leur bornes.

Finalement on arrive à la structure suivante pour l'interrupteur analogique :



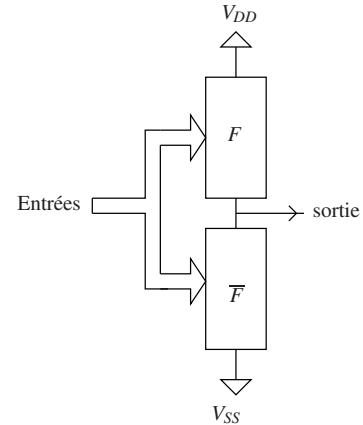
L'interrupteur comprend deux sous parties, le switch proprement dit que nous venons d'étudier et la partie driver. Celle ci inclue deux MOS complémentaires qui travaillent respectivement en Pull up et pull down. Elle permet d'assurer le pilotage en tensions complémentaires des transistors du switch. Le lecteur pourra vérifier que les deux transistors du switch sont bien respectivement bloqués et passants ensemble.

Utilisations Les switches analogiques permettent de réaliser en logique des multiplexeurs et peuvent être utilisés pour certaines fonctions logiques ainsi que pour des sorties trois états (0, 1, haute impédance). Ils peuvent remplacer avantageusement parfois les relais, en particulier pour des applications de mesure. Ils commutent en effet beaucoup plus rapidement que ces derniers et ne présentent pas d'usure mécanique. Le contact réalisé est cependant beaucoup moins bon ($r_{on} = \text{qq } \Omega$), et ils peuvent réinjecter des charges dans le circuit commuté (dans les électrodes A et B) lors de leur fermeture. Ces charges qui sont celles qui formaient le canal sont en faible quantité (quelques pico-Coulomb). Cela peut cependant être rédhibitoire dans certaines applications. Il faut garder à l'esprit que 10 pC dans un condensateur de 10 pF cela fait 1 Volt aux bornes du condensateur.

9.3 Portes logiques CMOS

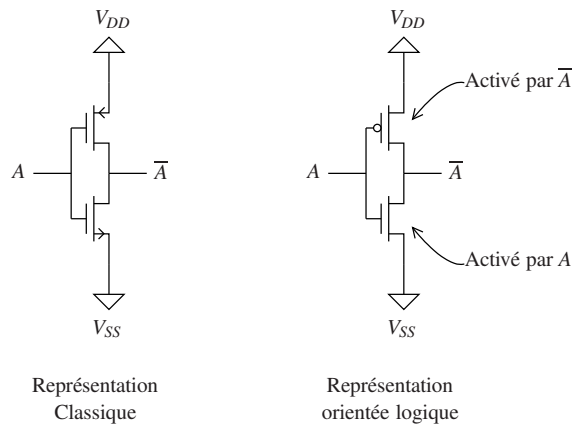
Les MOS sont les composants roi de l'électronique logique. Ils sont à la base de la technologie CMOS (Complementary MOS). Cette technologie est caractérisée par une consommation statique nulle. Ils ne consomment que lorsqu'ils commutent.

Pour réaliser une fonction logique F en technologie CMOS, on réalise un circuit de pull up qui est passant quant $F = 1$, bloqué sinon, et un circuit de pull down qui est passant quant $\bar{F} = 1$. De cette façon V_{DD} et V_{SS} sont isolés, et il n'y a pas de consommation statique. Quant $F = 1$, la sortie est tirée vers le haut, et la sortie vaut 1. Quant $F = 0$, la sortie est tirée vers le bas, et la sortie vaut 0.



9.3.1 L'inverseur MOS

C'est le plus simple des composants, les fonctions F et $\bar{F}$ sont directement réalisées par des transistors complémentaires :

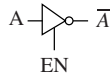


La figure ci-dessus présente l'inverseur en utilisant (à gauche) les symboles classiques pour les NMOS et PMOS. On constate que les deux transistors sont montés en pull up et pull down. Le NMOS est bloqué lorsque l'entrée A est au niveau bas ($V_{GS} = 0$; $A = V_{SS} = 0$ logique) alors que le PMOS est passant ($V_{GS} = V_{DD} - V_{SS} < V_T < 0$; $\bar{A} = V_{DD} = 1$ logique). Si $A = V_{DD} = 1$ logique, la situation s'inverse et le PMOS se bloque alors que le NMOS devient passant. Le PMOS est donc actif pour une entrée logique égale à zéro, alors que le NMOS est actif pour une entrée logique égale à 1. Ils travaillent respectivement en logique négative et positive. Pour résumer, les transistors PMOS de pull up inversent les entrées et les transistors de pull down prennent les entrées directes. C'est pourquoi on utilise souvent la représentation orientée logique de droite pour les schémas internes de composants logique à CMOS.

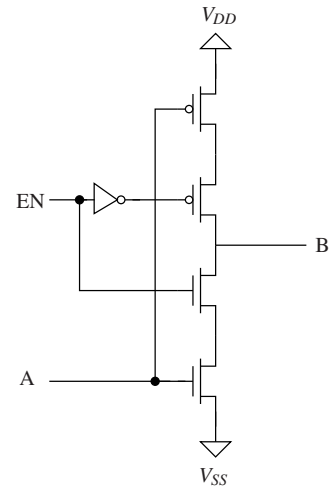
Inverseur à sortie haute impédance (HZ)

Il est parfois utile pour les circuits logiques de présenter une sortie qui peut passer en haute impédance selon la table de vérité suivante par exemple :

EN	A	B
0	X	HZ
1	1	0
1	0	1



On modifie alors la conception comme présenté dans la figure de droite. Tels qu'ils sont commandés les deux transistors centraux sont tous les deux passants ou bloqués en même temps, ce qui permet d'isoler la sortie.



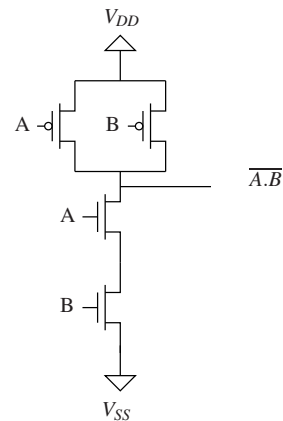
NOT trois états CMOS

9.3.2 Quelques autres portes CMOS

il n'est pas dans la portée de ce cours de présenter la réalisation de toutes les portes logiques possibles. On présentera cependant quelques unes des implantations des portes de base.

Fonction NAND (NON ET)

Cette fonction est une des plus simple à réaliser (plus simple que AND par exemple). La fonction à réaliser est $\overline{A.B}$. Soit $F = \overline{A} + \overline{B}$ (Th de Morgan) et $\overline{F} = A.B$. Les fonctions OR (OU) sont réalisées par des circuits en parallèle. En effet dans ce cas le courant peut circuler dans une branche OU bien dans l'autre. De la même façon les fonctions AND (ET) sont réalisées par des circuits en série. Le courant doit passer dans l'une ET l'autre des branches. On réalise donc la fonction $F = \overline{A} + \overline{B}$ avec deux transistors de pull up en parallèle qui inversent directement les entrées. $\overline{F}$ est réalisé avec deux transistors de pull down montés en série qui sont actifs sur les entrées directes.



NAND CMOS

Fonction AND (ET)

La fonction AND est simplement réalisée avec la fonction NAND suivie de l'inverseur précédent. Il faut remarquer ici qu'il n'est pas simple de réaliser directement la fonction AND pour laquelle on aurait une fonction de pull down qui prendrait en entrée $\overline{A}$ et $\overline{B}$ alors que les transistors de pull down prennent naturellement A et B en entrée. Les entrées naturelles de pull up sont les entrées inversées alors que les entrées naturelles de pull down sont les entrées directes.

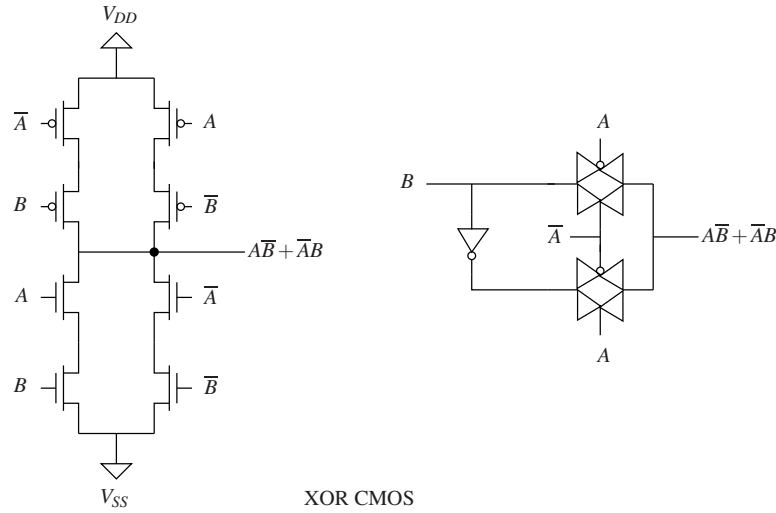
Fonctions plus complexes, exemple du XOR.

Compte-tenu de ce qui précède, on réalise facilement la fonction NOR (NON OU), puis le OU à partir de cette dernière.

Les fonctions logiques sont plus ou moins faciles à réaliser selon les principes énoncés ci-dessus. En effet, dans certains cas F et $\overline{F}$ prennent à la fois des variables directes et inversées en entrée. Dans ces cas là, on peut simplement inverser les entrées pour suivre la méthodologie précédente, mais ce n'est pas forcément optimal en nombre de transistors.

C'est le cas par exemple du XOR (Ou exclusif). Dans ce cas on a : $F = A\overline{B} + \overline{A}B$ qui dépend à la fois d'entrées inversées et non inversées. On calcule $\overline{F} = AB + \overline{A}\overline{B}$.

On obtient alors le schéma de gauche. Compte-tenu des inverseurs, ce schéma nécessite 12 transistors.



XOR CMOS

Il est aussi possible d'optimiser la conception en remarquant que $A = 1 \Rightarrow \text{XOR} = \bar{B}$, et que $A = 0 \Rightarrow \text{XOR} = B$. En fonction de A , il faut donc sortir B ou $\bar{B}$. On aboutit alors au schéma de droite qui ne nécessite que huit transistors. La réalisation de droite est présentée à titre d'exemple. Il en existe plusieurs autres. Le choix entre les diverses réalisations se fait principalement en fonction du nombre de transistors et de la vitesse recherchée pour la porte.

9.4 MOS en Amplification

Reprenons le montage de base pour un transistor NMOS, et traçons la caractéristique $V_s = f(V_e)$ comme dans la figure 24. Compte-tenu du câblage, on a $V_e = V_{GS}$ et $V_s = V_{DS}$. Lorsque V_e augmente en partant de zéro Volts, le transistor est bloqué jusqu'à ce que $V_e > V_T$. On peut voir à l'aide de la droite de charge, qu'il est ensuite saturé puis qu'il passe en mode triode pour finir en mode ohmique.

Remarque 1 : V_s n'atteint jamais zéro Volts. En effet le schéma équivalent final est un pont diviseur de tension entre R et la résistance de mode Ohmique finale $r_{on} = \frac{1}{k \frac{W}{L} (V_{DD} - V_T)}$.

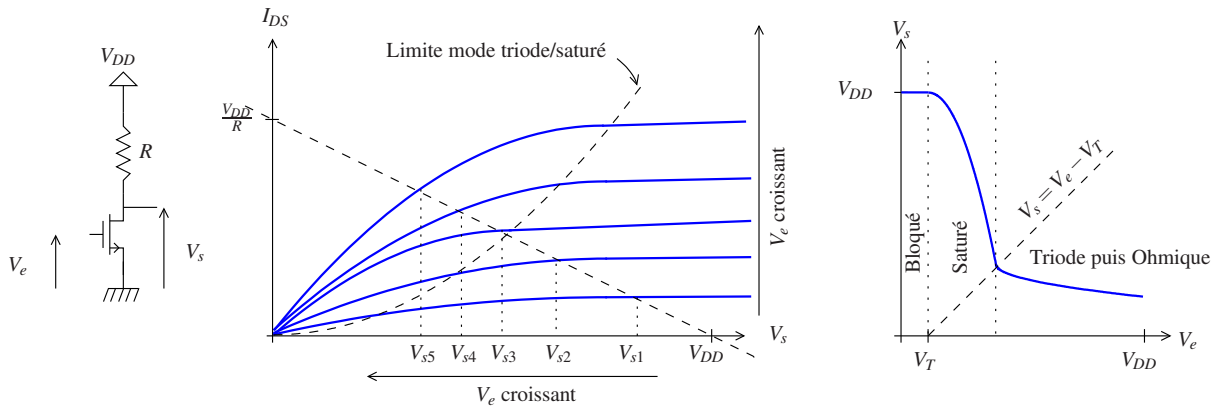


FIGURE 24 – Caractéristique $V_s = f(V_e)$ pour le montage de base NMOS

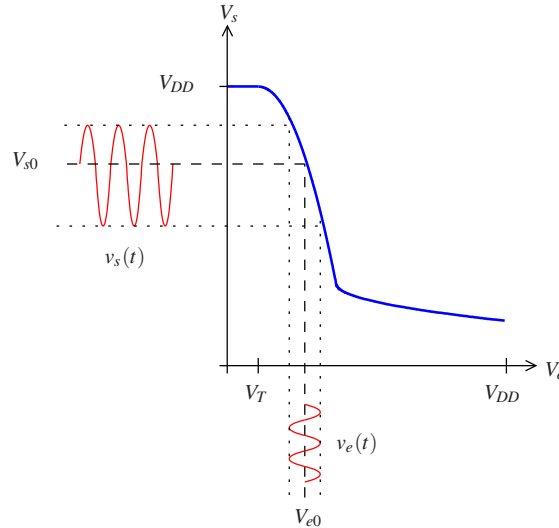
Remarque 2 : Les résultats de la figure 24 sont aussi valables pour le NMOS utilisé en pull down. On constate que dans ce mode, le NMOS passe par le mode saturé avant de finir en mode ohmique. On peut vérifier, en chargeant le NMOS par un condensateur, de l'ordre de grandeur de sa capacité de grille, pour modéliser le chargement par une autre porte logique et permettre le calcul d'un temps de réponse, que le temps passé dans le mode saturé est négligeable devant le temps total de commutation. C'est pourquoi en mode pull up et pull down, on ne considère en général que les états bloqués et passant ohmique (r_{on}) des MOS.

9.4.1 Petit signal

Comment réaliser un amplificateur avec un NMOS, et dans quel mode doit-on l'utiliser ?

Un amplificateur idéal réalise la fonction $V_s = A \times V_e$, avec A le plus grand possible. Ce n'est manifestement pas le cas du NMOS de la figure 24 si on considère la caractéristique $V_s = f(V_e)$.

C'est cependant possible si l'on prend pour V_e , un signal d'amplitude limitée $v_e(t)$, que l'on appellera petit signal, variant autour d'une valeur fixe V_{e0} que l'on appelle la polarisation comme présenté ci-dessous pour un petit signal sinusoïdal :



La réponse du NMOS est un signal variable $v_s(t)$ variant autour d'une valeur constante V_{s0} : $V_s(t) = V_0 + v_s(t)$. On constate que plus la pente de la caractéristique $V_s = f(V_e)$ est grande, plus le module du signal de sortie ($|v_s(t)|$) est grand et donc plus $|A| = \frac{|v_s(t)|}{|v_e(t)|}$ est grand. On a donc intérêt à travailler dans le domaine de tension où le NMOS est saturé. Comme la pente de la caractéristique est négative, on voit aussi que A est négatif. $A = -\frac{|v_s(t)|}{|v_e(t)|}$. Enfin $v_s(t)$ n'est manifestement pas symétrique autour de V_{s0} , contrairement au petit signal d'entrée. Cet effet est la distorsion due à la courbure de la caractéristique qui est, en mode saturé, un morceau de parabole. Pour annuler cet effet, il faudrait que la caractéristique soit rectiligne. Il est possible d'approcher cette situation en prenant une amplitude suffisamment faible pour $v_e(t)$ afin que la parabole soit localement assimilable à sa tangente avec le moins d'erreur possible (les termes de deuxième ordres doivent être négligeables). Nous venons ainsi d'introduire la notion de petit signal.

Remarque : l'amplitude du signal de sortie est forcément limitée par V_{DD} . Par conséquent : $|v_e(t)| < \frac{V_{DD}}{A}$. Pour que $|A|$ soit grand, la condition petit signal est assez naturellement respectée.

Résumé : Pour utiliser un MOS (N ou P) efficacement en amplification il faut :

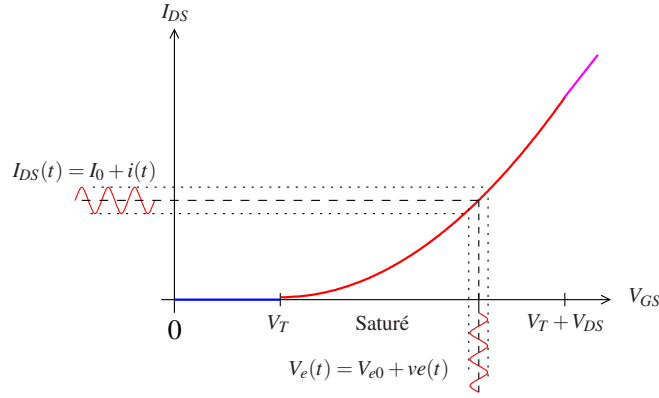
1. L'amener autour d'un point de fonctionnement où il sera en mode saturé. On parle de polarisation.
2. L'exciter avec un signal variant autour de ce point de polarisation. L'amplitude de ce signal doit être suffisamment faible pour que la réponse du MOS y soit proportionnelle (linéarité). On parle d'excitation en petit signal.

9.4.2 Schéma équivalent petit signal

Considérons à nouveau le montage de base de la figure 24. Pour une excitation petit signal v_e autour de V_{e0} :

$$V_e(t) = V_{e0} + v_e(t) \quad (24)$$

Le courant qui traverse le MOS s'écrit $I_{DS}(t) = I_{DS0} + i(t)$ comme présenté dans la figure suivante :



Si R a été correctement choisie pour que le MOS soit saturé, avec la transconductance g (variation de courant à la sortie/variation de tension à l'entrée) :

$$g(V_{e0}) = \frac{\partial I_{\text{sat}}(t)}{\partial V_e} = k \frac{W}{L} (V_{e0} - V_T) \quad (25)$$

et le courant de polarisation I_0 :

$$I_o = \frac{1}{2} k \frac{W}{L} (V_{e0} - V_T)^2 \quad (26)$$

Le courant qui traverse le MOS peut être réécrit comme :

$$I_{\text{sat}}(t, V_e) = I_0(V_{e0}) + g(V_{e0})v_e(t) \quad (27)$$

En séparant la partie polarisation (signal constant) du petit signal, on peut réécrire l'équation 24 comme

$$V_{s0} + v_s(t) = \underbrace{V_{DD} - RI_0}_{\text{Réponse du circuit à la polarisation}} - \underbrace{Rg(V_{e0})v_e(t)}_{\text{Réponse au petit signal}} \quad (28)$$

Ce qui donne pour le petit signal uniquement :

$$v_s(t) = -Rg(V_{e0})v_e(t) \quad (29)$$

Dans ce qui suit et en particulier dans les schémas, on notera $g(V_{e0}) = g$ pour des raisons de clarté.

L'équation 29 donne la réponse au petit signal du circuit. Le montage ci-dessous présente le circuit électrique qui donne le même résultat :

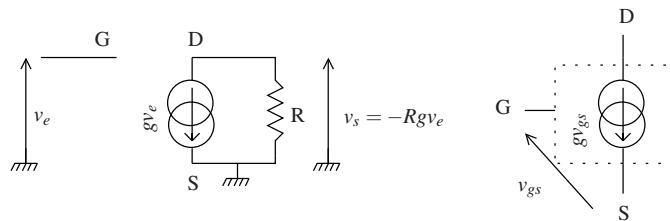


Schéma équivalent petit signal du montage de base

Schéma équivalent petit signal du MOS

Le MOS en petit signal se comporte comme un générateur de courant commandé en tension par v_e . On en déduit un

premier schéma équivalent petit signal pour le MOS. La grille est présente dans le schéma équivalent, mais n'est connectée à rien car le courant d'entrée est nul.

Amélioration du modèle :

Dans le modèle ci-dessus, nous n'avons tenu compte ni de la capacité Grille-Source, ni de l'effet Early. Il existe bien sûr d'autres effets qui pourraient aussi être pris en compte. Les autres capacités parasites grille-drain et drain-source pourraient être ajoutées au modèle. Cependant, elles sont de moindre valeurs et leurs effets ne sont pas notables aux fréquences qui nous intéressent. On ne rajoutera pas non plus, l'effet de la différence de potentiel source-bulk dans le cas où source et bulk ne sont pas connectés. La prise en compte de cet effet est anecdotique dans le cas de l'utilisation usuelle des transistors MOS.

Si on considère l'équation 23, l'effet Early se traduit par un courant supplémentaire au courant de saturation proportionnel à V_{DS} . Il apparaîtra donc comme une résistance en parallèle au générateur de tension commandé sur le schéma petit signal. La valeur de cette résistance r_o est donnée par l'équation 23 : $r_o = \frac{V_A}{I_0}$. Ou $I_0 = \frac{1}{2}k\frac{W}{L}(V_{GS0} - V_T)^2$ et où V_A est la tension de Early.

Le modèle complet est présenté figure 25.

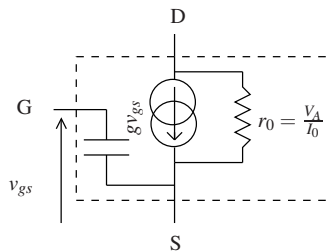
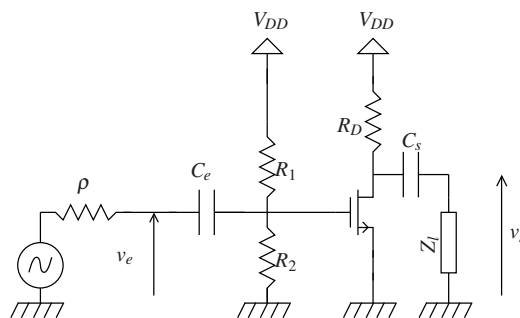


FIGURE 25 – Schéma équivalent petit signal d'un transistor NMOS

9.4.3 Méthodologie d'analyse du fonctionnement en petit signal d'un amplificateur à MOS.

Pour fixer les idées, nous prendrons comme exemple le circuit suivant :



La méthode présentée ici est valable pour tout type de circuit en petit signal. Elle comporte deux volets

1. Étude de la réponse aux petits signaux.
2. Étude de la polarisation

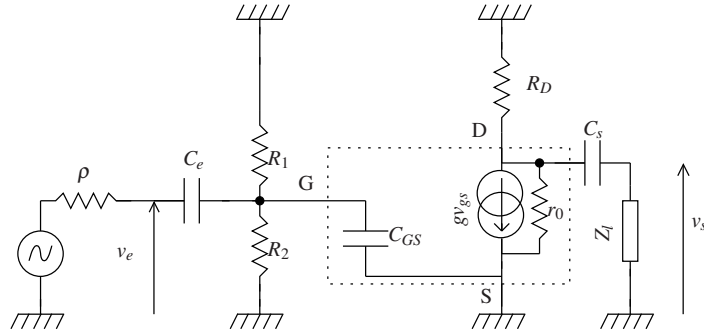
En général, on s'intéresse au comportement petit signal. Quant on parle d'amplificateur, c'est le petit signal qui est amplifié pas la polarisation. On cherche à avoir un gain donné. Par exemple $|H(\omega)| \frac{|v_s(\omega)|}{|v_e(\omega)|} = 10$.

L'étude petit signal, va nous permettre d'exprimer le gain de l'amplificateur en fonction des paramètres du schéma équivalent petit signal. Ceux-ci et en particulier la transconductance g dépendent de la polarisation. L'étude de la polarisation permettra de fixer g .

À la fin de l'étude, on pourra obtenir la réponse complète du circuit avec le théorème de superposition, en additionnant le signal de polarisation au comportement petit signal.

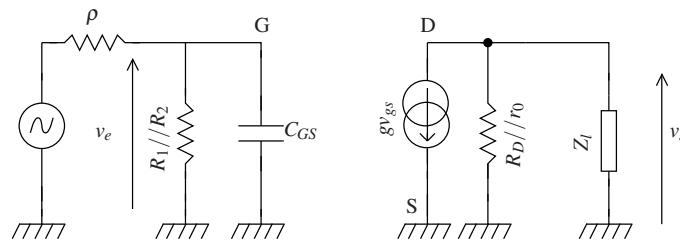
Étude petit signal :

Si on considère l'équation 28, on voit que la polarisation est la partie constante dans le temps et le petit signal la partie variable dans le temps. Par conséquent, pour le petit signal, tout ce qui est constant est nul. Pour obtenir le schéma équivalent petit signal il faut donc remplacer dans le schéma tous les potentiels constants par la masse et le transistor par son schéma équivalent petit signal :



Dans le même ordre d'idée, les condensateurs en entrée et sortie dans le montage C_e et C_s ont pour rôle de séparer la polarisation du continu. Ils coupent naturellement le continu, et doivent être choisis assez grands pour laisser passer le petit signal dans la bande passante du circuit. On les assimile donc souvent dans un premier temps à des fils tant que l'on ne s'intéresse pas à calculer la valeur qu'il faut leur donner pour que cela soit vrai.

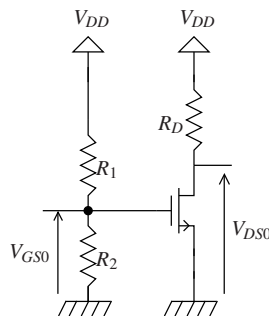
Le schéma équivalent peut donc être simplifié et réorganisé de la façon suivante :



Avec ce schéma on calcule aisément les caractéristiques de l'amplificateur ainsi réalisé :

Réponse fréquentielle à vide ($Z_l = \infty$)	$H(\omega) = -(R_D // r_0)g \approx -R_D g$
Impédance d'entrée	$Z_{in} = R1 // R2 // C_{GS}$
Impédance de sortie	$Z_{out} = R_D // r_0 \approx R_D$

Étude de la Polarisation : La polarisation consiste à amener le circuit autour du point de fonctionnement qui nous intéresse. Les tensions et courants de polarisation sont des signaux continus. On commence par dessiner le schéma du circuit en ne tenant compte que du continu. Les condensateurs sont pour les signaux continus des circuits ouverts. Ils isolent la partie centrale du circuit du générateur et de la charge :



L'étude de la polarisation est essentielle car c'est elle qui donne sa validité au schéma équivalent petit-signal. En particulier, il faut s'assurer que le transistor soit saturé.

Le courant de grille est nul en continu donc : $V_{GS0} = \frac{R_2}{R_1+R_2}V_{DD}$.

On peut alors calculer le courant de polarisation si le transistor est saturé : $I_o = \frac{1}{2}k\frac{W}{L}(V_{GS0} - V_T)^2$. On ne tient pas compte de l'effet Early pour la polarisation en général car λ est suffisamment faible pour que le surcroît de précision soit ridicule devant la précision des résistances.

Le transistor est saturé si $V_{DS0} > V_{GS0} - V_T$. Soit :

$$R_D < \frac{(V_{DD} + V_T - V_{GS0})}{\frac{1}{2}k\frac{W}{L}(V_{GS0} - V_T)^2} \quad (30)$$

Remarque 1 Cette condition est de la forme " $R_D < \text{quelque chose}$ ". On pouvait s'y attendre car si R_D est trop grand, le courant traversant le MOS devient petit, et il passe en mode triode voir ohmique.

Remarque 2 Il ne faut pas perdre de vue que les paramètres sont liés entre eux. Les résistances R_1 et R_2 permettent de fixer V_{GS0} , mais elles interviennent aussi dans la valeur de l'impédance d'entrée. De même R_D est lié à la condition de saturation, mais aussi à la valeur de l'impédance de sortie. I_o quant à lui fixe bien sur la valeur de g . Le travail de conception consistera bien souvent à trouver des compromis entre ce dont on rêve ($Z_{in} = \infty$, $Z_{out} = 0$, ...) et ce qui est réalisable.

10 Transistors bipolaires à jonction

C'est le transistor historique découvert en 1947. Sa forme actuelle est cependant différente. Il en existe deux versions :

1. transistor NPN



2. transistor PNP

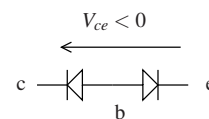


Dans les deux cas, l'épaisseur de la jonction est faible et l'émetteur est plus fortement dopé que la base ou le collecteur. Les porteurs majoritaires y sont en grand nombre.

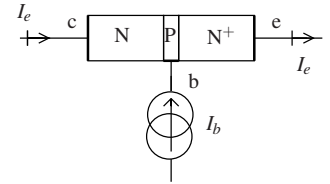
10.1 Fonctionnement

Nous allons décrire le fonctionnement du NPN qui est le transistor le plus classique. Tout est transposable au PNP en changeant la nature des porteurs dans le discours : $e^- \longleftrightarrow p^+$.

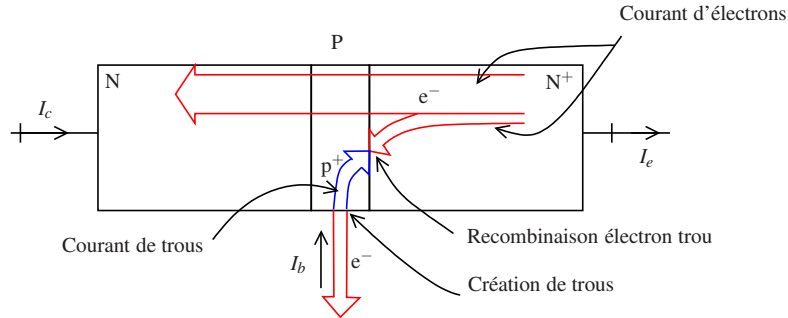
Si la base n'est pas connectée, pour traverser le transistor entre collecteur et émetteur un courant va "voir" une suite de deux jonctions, NP puis PN. Aucun courant ne peut donc circuler.



Si on rend la jonction base-émetteur passante en la polarisant correctement, c'est à dire à $V_{be} \approx 0.6 \text{ V}$, c'est-à-dire en y injectant du courant (cf Figure de droite), alors, les électrons très nombreux dans l'émetteur rentrent dans la base avec de l'élan, comme celle-ci est très fine, ils ont une probabilité non nulle de se retrouver directement dans le collecteur.



En fait seule la petite fraction des électrons entrant dans la base par l'émetteur qui correspond au fonctionnement normal de la jonction base-émetteur se recombinaient avec les trous, les autres continuent leur chemin dans le collecteur :



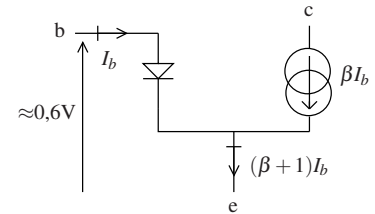
Il s'avère, et c'est l'effet transistor, que la quantité d'électrons collectés par le collecteur est contrôlée par la quantité de trous créée au niveau de l'électrode de base par le flux d'électrons sortant. Elle est même proportionnel à ce dernière :

$$I_c = \beta I_b \quad (31)$$

Comme il y a conservation du nombre de porteurs, la loi des nœuds s'applique :

$$I_e = I_c + I_b \quad (32)$$

Le fonctionnement du transistor peut donc être schématisé comme présenté figure de droite. La jonction base-émetteur se comporte comme une diode normale, sa polarisation permet de contrôler le courant I_b dont dépendent les courants d'émetteur et de collecteur. Le transistor est un amplificateur de courant. Remarquons que si la diode base-émetteur est bloquée, le transistor est bloqué.

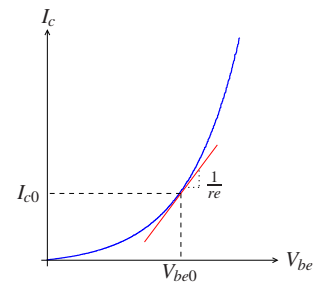


10.2 Caractéristiques

$$I_c = f(V_{be})$$

Cette caractéristique suit à β près la caractéristique d'une diode selon la relation d'Ebers-Moll. Comme on n'utilise pas la diode base-émetteur en inverse, en mode actif normal (non bloqué), on a $I_c = \beta I_b$. En mode actif on approxime cette relation en négligeant le terme -1 devant l'exponentielle dans l'équation d'Ebers-Moll (équation 3) par :

$$I_c \approx \beta I_0 e^{\frac{V_{be}}{V_\phi}} \quad (33)$$



L'expression de la pente autour d'un point de polarisation (V_{be0} , I_{c0}) s'exprime en fonction de V_ϕ :

$$\frac{\partial I_c}{\partial V_{be}}(V_{be0}) = \beta I_0 \frac{1}{V_\phi} e^{\frac{V_{be0}}{V_\phi}} = \frac{I_{c0}}{V_\phi} \quad (34)$$

Cette grandeur est aussi l'inverse de la résistance dynamique r_e (cf figure précédente), soit :

$$r_e = \frac{V_\phi}{I_{c0}} \quad (35)$$

La relation 35 permet de calculer facilement la résistance dynamique d'émetteur en connaissant I_{c0} .

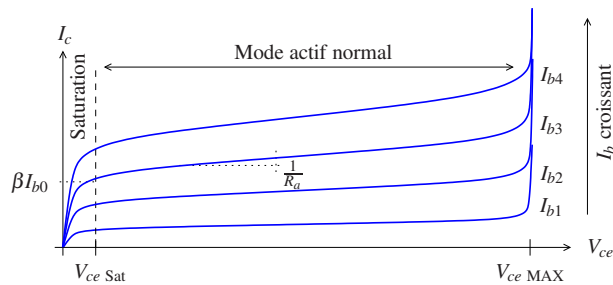
On prend en général :

$V_\phi = 26 \text{ mV à } 300^\circ\text{K}$

(36)

$$\underline{I_c = f(V_{ce})}$$

Compte-tenu de la relation 31, la caractéristique $I_c = f(V_{ce})$ est obtenue à partir de la relation d'Ebers-Moll. Il y a cependant une valeur maximale de la tension V_{ce} à partir de laquelle le courant traverse simplement par effet d'avalanche le transistor. Cette zone est dangereuse pour le transistor car elle correspond à une puissance $P = V_{ce\text{MAX}} \times I_c$ susceptible de l'endommager.



En mode actif normal, c'est-à-dire dans la zone de la caractéristique dans laquelle on cherche à utiliser le transistor, le courant y est quasi constant et contrôlé par le courant qui circule dans la base. Dans cette zone le transistor se comporte comme un générateur de courant piloté. Il existe cependant une valeur de la tension V_{ce} minimale en-dessous de laquelle le transistor est dit saturé et qui correspond en fait au coude de la caractéristique de la diode base-émetteur.

Il existe comme pour les transistors MOS une tension de Early V_A qui permet d'écrire en négligeant $V_{ce\text{Sat}}$ la résistance du générateur en mode actif normal comme :

$$R_a = \frac{V_A}{\beta I_{b0}} \quad (37)$$

Schéma équivalent

Compte-tenu de ce qui précède, on améliore le schéma équivalent :

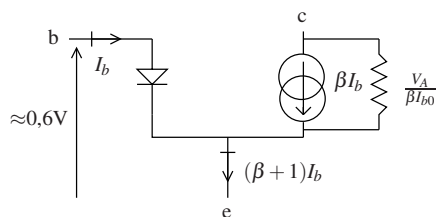


FIGURE 26 – Schéma équivalent transistor NPN.

Transistor idéal

Un transistor NPN idéal, aurait pour paramètres :

$V_{ce\text{Sat}} = 0$
$V_{ce\text{MAX}} = \infty$
$R_a = \infty$

En appliquant les mêmes raisonnements au transistor PNP, on aboutit au schéma équivalent de la figure 30. Les valeurs des paramètres pour le transistors idéal sont les mêmes que pour le transistor NPN. Notez que seule l'orientation de la diode et le sens du courant change.

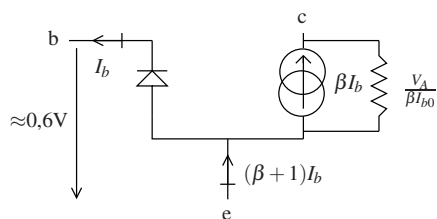


FIGURE 27 – Schéma équivalent transistor PNP.

10.3 Paramètres d'influences

10.3.1 Effet de la Température

La température influence le comportement d'un transistor bipolaire de deux façons :

1. Par l'intermédiaire de la variation du courant de la diode base-émetteur avec la température : $\frac{\partial V_{be}}{\partial T} = -2\text{mV}/^\circ\text{C}$
2. Par le gain en courant β , lui même **très** dépendant de la température.

La figure 28 illustre cette dépendance pour le transistor NPN 2N222 qui est un transistor à tout faire, standard.

On peut voir que le gain varie au minimum de 100% entre 25°C et 125°C.

Toute conception de circuit basée sur un transistor bipolaire peut donc se baser sur l'hypothèse β grand, typiquement $(\beta + 1) \approx \beta$, mais pas sur une valeur précise de β . On n'utilisera β pour le calcul des circuits qu'à titre d'ordre de grandeur.

10.3.2 Variation de β avec la fréquence

Il existe une fréquence de coupure au-delà de laquelle la valeur du gain en courant diminue fortement. Les variations de β peuvent être modélisées comme un phénomène du premier ordre.

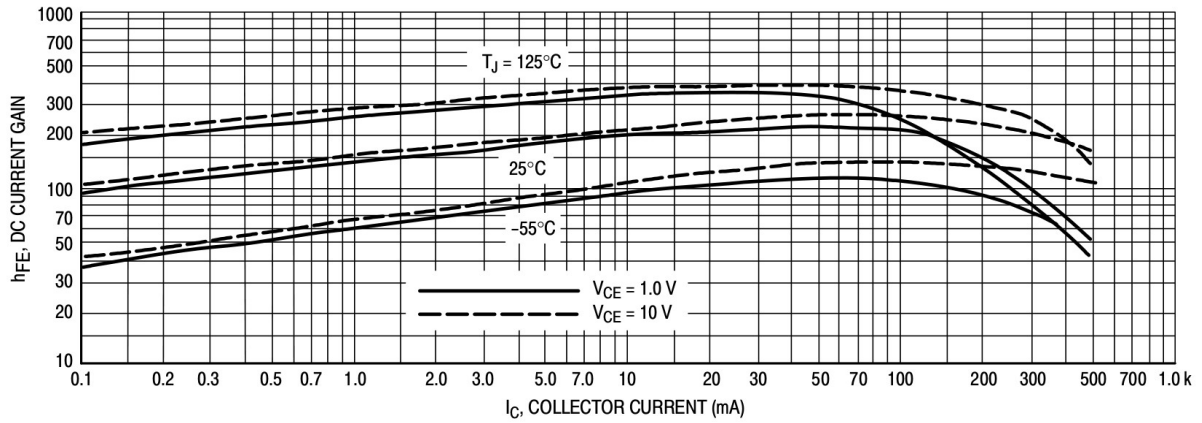
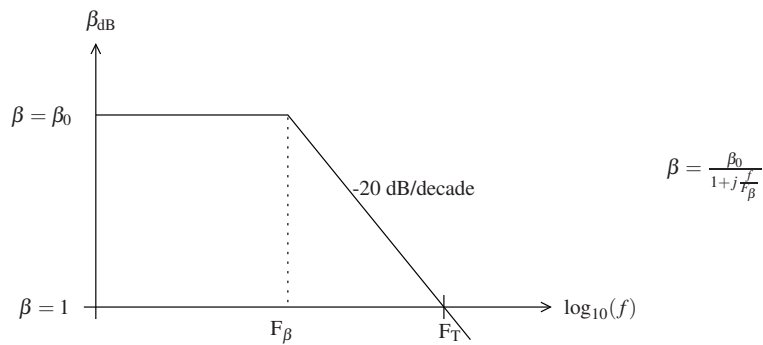


FIGURE 28 – Variation du gain en courant β avec le courant de collecteur et la température.



Le paramètres F_β et F_T sont en général mentionnés dans les documentations des constructeurs.

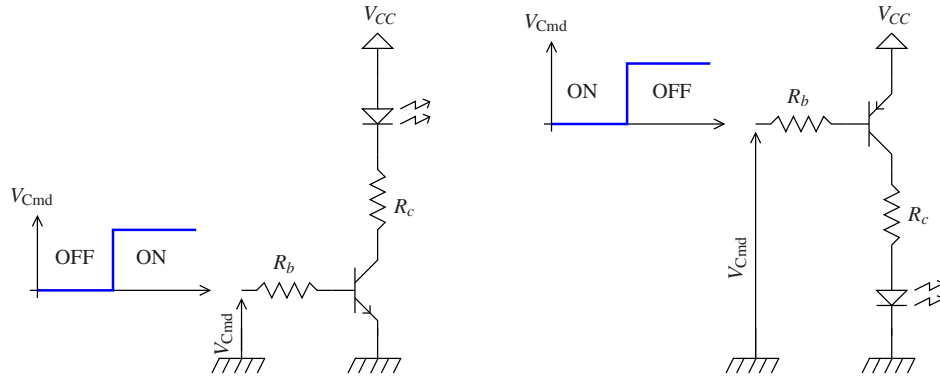
11 Applications des transistors à jonction

Comme pour les MOS, les transistors bipolaires à jonctions sont utilisés pour des applications de commutation ou d'amplification. Dans le premier cas on les utilise en mode bloqué ($I_c = I_e = I_b = 0$)/saturé ($V_{ce} = 0$), dans l'autre cas en mode actif normal (ni bloqué ni saturé, $I_c = \beta I_b$).

11.1 BJT en Commutation

On cherche en commutation à faire fonctionner le transistor en mode bloqué-saturé. Le transistor passant doit être saturé pour avoir $V_{ce} \approx 0$.

Exemple : Allumage d'une diode. Le transistor se pilote en contrôlant la tension aux bornes de la diode base-émetteur. Le pilotage se fait donc avec des tensions complémentaires selon que l'on utilise un NPN ou un PNP.

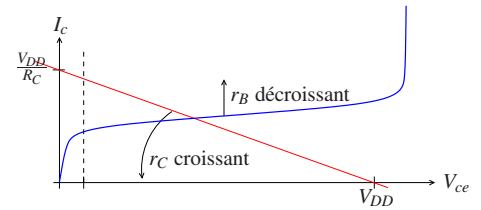


Analyse et dimensionnement du montage à transistor NPN

- Si $V_{cmd} = 0$, alors la diode base-émetteur est bloquée et le transistor aussi. La LED est éteinte.
- Si $V_{cmd} = V_{CC}$, on suppose ici $V_{CC} \gg 0,6 \text{ V}$, alors la diode base-émetteur est passante et $V_{be} \approx 0,6 \text{ V}$.

La droite de charge ci-contre montre qu'il faut que R_c soit suffisamment grande et/ou R_b suffisamment petite pour que le transistor sature.

En pratique, on calcule R_c pour obtenir le courant désiré dans le collecteur, celui qui convient à la LED dans notre exemple. Ensuite on détermine R_b pour que le transistor soit saturé sans toutefois dépasser le courant maximal admissible pour la diode base-émetteur.



Montage à transistor PNP

Le fonctionnement du montage à transistor PNP est analogue sauf que $V_{be} \approx V_{cmd} - V_{DD}$. Le transistor est donc bloqué pour $V_{cmd} = V_{DD}$ et passant pour $V_{cmd} = 0$.

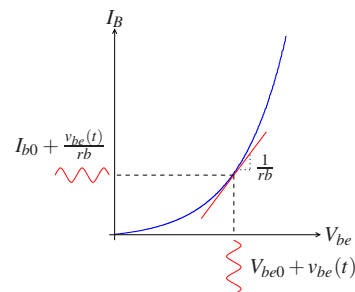
Circuits logiques à transistors bipolaires : Il est possible avec des transistors bipolaires à jonction de faire des composants logiques. C'est le cas des familles de composants logiques TTL (Transistor Transistor Logique) ou ECL (Emitter Coupled Logic). Elles sont cependant maintenant obsolètes car elles consomment beaucoup plus que les familles logiques à transistors MOS. En outre, ces derniers commutent maintenant aussi vite que les familles logiques à transistors à jonction, grâce à la réduction de la taille de la gravure qui a permis de diminuer les capacités de grille.

11.2 BJT en Amplification

Comme pour les transistors MOS, on amène le transistor autour d'un point de fonctionnement (la polarisation) puis on considère les variations (le petit signal) autour de ce point.

11.2.1 Schéma équivalent petit signal

Reprenons le schéma équivalent de la figure 26. C'est pour le moment le schéma d'un composant non linéaire à cause de la diode base-émetteur. En condition de petit signal, il est possible de le linéariser. Considérons la caractéristique de la diode base-émetteur pour des petits signaux lorsque l'on polarise le transistor autour du point de fonctionnement (I_{B0}, V_{be0}).



Au premier ordre, on a : $I_b = I_{b0} + i_b = I_{b0} + \frac{v_{be}}{r_b}$ ou r_b est la résistance dynamique de la diode base-émetteur $r_b = \frac{\partial V_{be}}{\partial I_b}$ au point de polarisation.

Comme $I_b = \frac{1}{\beta} I_c$, on a :

$$r_b = \beta r_e \quad (38)$$

où r_e est la résistance dynamique de l'émetteur tel que définie équations 34 et 35.

Finalement pour le petit signal, on remplace la diode base-émetteur par la résistance dynamique de base dans le schéma équivalent de la figure 26, et l'on obtient le schéma équivalent petit signal de transistor NPN :

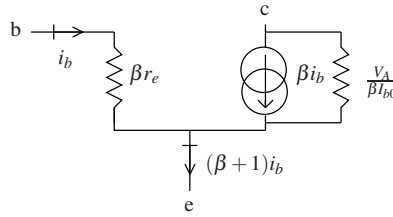


FIGURE 29 – Schéma équivalent petit signal transistor NPN.

Compte tenu des équations 35 et 36, on utilise en général l'expression suivante pour calculer r_e :

$$r_e = \frac{26 \text{ mV}}{I_{c0} \text{ mA}} \quad (39)$$

Dans cette expression, I_{c0} est le courant de polarisation de collecteur.

Les raisonnements précédents sont transposables au transistor PNP, et l'on obtient le schéma équivalent de la figure 30 pour le transistor PNP. La valeur de r_e et de la résistance de Early se calcule de la même façon que pour le transistor NPN.

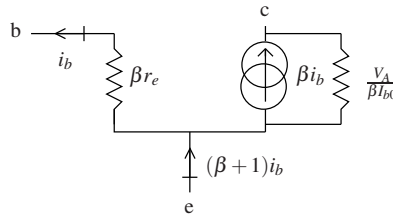


FIGURE 30 – Schéma équivalent petit signal transistor PNP.

11.2.2 Les trois montages de base des transistors bipolaires

La figure 31 présente les trois montages de base des transistors bipolaires en amplificateurs petits signaux. Ils sont présentés pour des transistors NPN, mais ils peuvent être transposés à des transistors PNP. Dans cette figure, on n'a pas représenté les circuits de polarisation. C'est pourquoi par exemple l'alimentation V_{CC} apparaît non connectée pour le montage base commune. Les noms des montages font référence à l'électrode qui est connectée à la masse pour le petit signal. Ainsi pour le montage émetteur commun, l'émetteur est à la masse pour le petit signal. Il en est de même pour les montages collecteur commun et base commune.

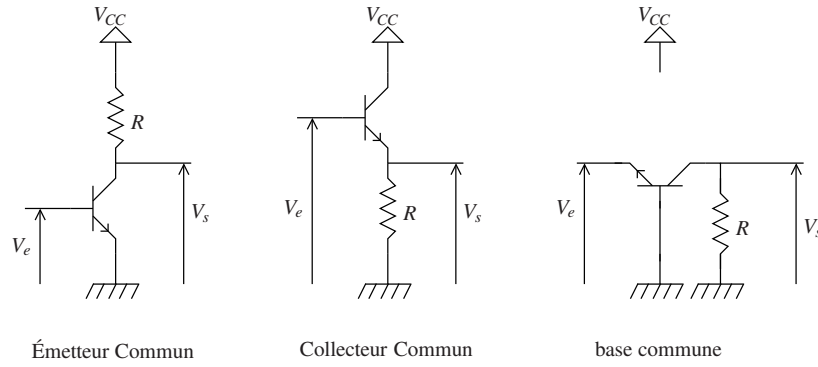
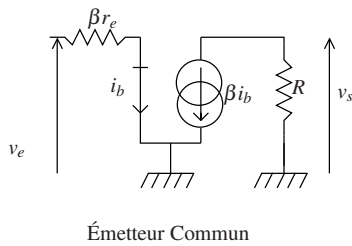


FIGURE 31 – Les trois montages de bases des transistors bipolaires à jonction.

Ces trois montages ont chacun des caractéristiques différentes (gain, impédances d'entrées et de sorties, etc..) qui président à leurs usages. Le lecteur pourra pour s'entraîner recalculer les résultats présentés ci-dessous.

Émetteur commun C'est le montage à utiliser en première intention. Avec le schéma équivalent petit signal de la figure 29, en négligeant la résistance de Early pour simplifier (elle n'apporte pas grand chose à la discussion et complexifie grandement l'étude). On obtient le schéma équivalent petit signal dont on déduit les caractéristiques petit signal du montage.

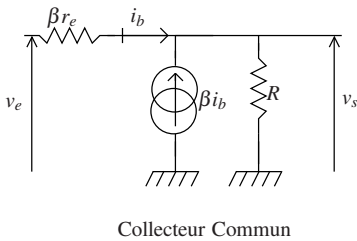


$$\begin{aligned}\frac{v_s}{v_e} &= -\frac{R}{r_e} \\ Z_{in} &= \beta r_e \\ Z_{out} &= R\end{aligned}$$

Ordres de grandeur typiques : $r_e = 5\Omega$, $R = 500\Omega$.

C'est un amplificateur inverseur. L'impédance d'entrée n'est pas trop petite, et l'impédance de sortie pas trop grande même pour un gain typique de -100.

Collecteur commun Ce montage est appelé *emitter follower* dans la littérature anglo-saxonne car c'est un suiveur. C'est évident si on considère la figure 31. L'entrée et la sortie sont séparées par la tension de diode qui reste à peu près constante $\approx 0.6V$ donc nulle pour le petit signal. On a donc $v_s \approx v_e$.



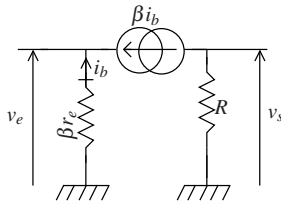
$$\begin{aligned}\frac{v_s}{v_e} &= \frac{R}{R + \frac{\beta}{\beta+1}r_e} \approx 1 \\ Z_{in} &= \beta \left(\frac{\beta+1}{\beta}R + r_e \right) \approx \beta R \\ Z_{out} &= \frac{\beta+1}{\beta}r_e // R \approx r_e\end{aligned}$$

Ordres de grandeur typiques : $r_e = 5\Omega$, $R = 500\Omega$.

L'impédance d'entrée est plutôt grande et l'impédance de sortie plutôt petite, ce qui fait de ce montage un bon étage d'adaptation d'impédance.

Remarque : Le résultat $Z_{out} \approx r_e$ peut se voir directement dans l'expression du gain : c'est la valeur qu'il faut donner à R dans cette expression pour obtenir un gain de $\frac{1}{2}$. Avec la méthode de la demi-charge, comme $R \gg r_e$ si l'on charge le suiveur par r_e , on divise la tension de Thévenin par deux.

Base commune Ce montage n'est utilisé qu'en haute fréquence en particulier dans le montage cascode (CF précepteurat) car il est peu sensible à l'effet Miller.



Base commune

$$\begin{aligned}\frac{v_s}{v_e} &= \frac{R}{r_e} \\ Z_{in} &= \frac{\beta}{\beta+1} r_e \approx r_e \\ Z_{out} &= R\end{aligned}$$

Ordres de grandeur typiques : $r_e = 5\Omega$, $R = 500\Omega$.

Son gain et son impédance de sortie sont les mêmes que pour l'émetteur commun. Par contre son impédance d'entrée est beaucoup moins bonne.

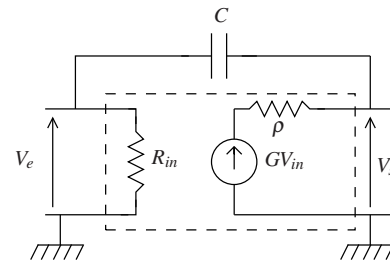
Effet Miller

L'effet Miller décrit la capacité vue en entrée d'un amplificateur lorsque sa sortie est rebouclée sur l'entrée par un condensateur comme présenté à droite. En général ce condensateur n'est pas placé la volontairement, c'est un condensateur parasite.

Le courant d'entrée dans l'amplificateur bouclé est la somme du courant circulant dans le condensateur et de celui circulant dans R_{in} , soit :

$I_{in} = V_e \frac{1}{R_{in}} + V_e(1-G) \frac{1}{\rho + Z_c}$. On en déduit facilement que l'impédance d'entrée de l'amplificateur bouclé est :

$$Z_{in} = R_{in} // \frac{\rho + Z_c}{1-G}$$



A basses fréquences, $|Z_c|$ est très grand devant R_{in} . L'impédance d'entrée du système bouclé est R_{in} .

A moyennes fréquences, c'est-à-dire lorsque $|Z_c| \gg \rho$ et que $|Z_c|$ est de l'ordre de grandeur de R_{in} , l'impédance d'entrée de l'amplificateur bouclé est $Z_{in} = R_{in} // \frac{Z_c}{1-G}$. Cela signifie que l'on voit en entrée, en parallèle avec R_{in} un condensateur de valeur $(1-G)C$. Ceci qui vient limiter la bande passante de l'amplificateur en atténuant le signal d'entrée.

A hautes fréquences enfin, l'impédance d'entrée tend vers $Z_{in} = \frac{\rho}{(1-G)}$. Comme en général $|G| \gg 1$ et que ρ est petite, Z_{in} devient petite et l'amplificateur n'est plus utilisable normalement.

Remarque : $(1-G)$ doit être une quantité positive sinon on est en réaction et le système n'est pas stable. L'effet Miller rend instables les amplificateurs à gain positif et diminue le gain des amplificateurs à gain négatif.

Nous n'avons pas, dans le schéma équivalent des transistors bipolaires, introduit les capacités parasites : capacités base-émetteur, base-collecteur, émetteur-collecteur. On doit les prendre en compte pour comprendre l'intérêt du montage base commune. Parmi ces trois capacités parasites, la capacité émetteur-collecteur est beaucoup plus faible que les autres. Le montage base commune est donc beaucoup moins sensible que le montage émetteur commun à l'effet Miller.

11.2.3 Polarisation et emballement Thermique

Dans ce paragraphe, nous allons présenter les circuits de polarisation des montages amplificateurs que nous n'avons pas décrit dans ce qui précède.

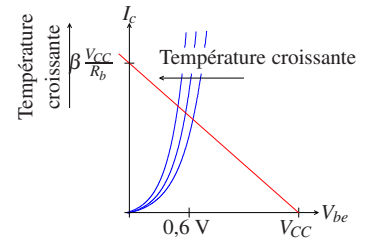
Pour le montage de première intention : l'émetteur commun, le circuit de polarisation le plus simple est le montage a) de la figure 32. Dans ce montage, on utilise comme pour les amplificateurs à MOS des condensateurs de découplage pour séparer le petit signal de la polarisation. Ce montage, qui a le mérite d'être très simple, est malheureusement réputé pour ne pas être stable thermiquement. En effet dès la mise en route de l'amplificateur, le transistor dissipe de la puissance par effet Joule, principalement entre émetteur et collecteur $P = V_{ce} \times I_{ce}$, ce qui fait monter sa température. Or, le gain en

courant β augmente avec la température et la tension au bornes de la diode base-émetteur diminue elle aussi de 2 mV par degré Celsius.

Traçons la droite de charge pour le montage a :

$$V_{be} = V_{CC} - R_b I_b = V_{CC} - R_b \frac{I_c}{\beta}$$

On voit qu'avec l'augmentation de la température, les effets sur β et sur la caractéristique se cumulent pour augmenter le courant ce qui fait à son tour monter la température. Le système s'emballe et comme $V_{ce} \rightarrow 0$ le transistor finit par saturer.



Notons que parmi ces deux effets, celui de β est prépondérant.

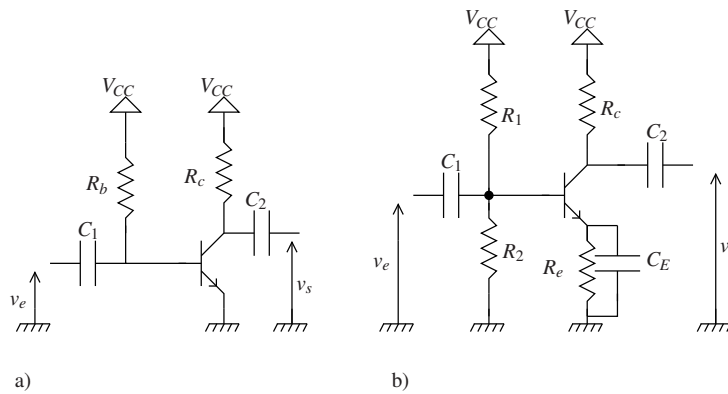


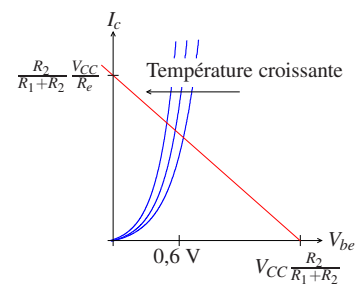
FIGURE 32 – Montages émetteur commun avec circuits de polarisation.

Le montage b) de la figure 32, présente la méthode de polarisation classique qui permet de limiter l'emballement thermique. Dans ce montage l'émetteur n'est pas directement connecté à la masse. On rajoute une résistance d'émetteur R_e et l'on impose le potentiel de la base à l'aide d'un pont diviseur. Bien sûr, comme on désire toujours réaliser un montage émetteur commun on rajoute en parallèle à cette résistance un condensateur de découplage qui permet au petit signal de continuer à voir un montage émetteur commun.

La droite de charge pour le montage (b) est donnée par :

$$V_{be} = V_{CC} \frac{R_2}{R_1 + R_2} - R_e I_e \approx V_{CC} \frac{R_2}{R_1 + R_2} - R_e I_c$$

La droite de charge est maintenant indépendante de β . L'augmentation de la température n'a d'effet que sur la caractéristique. Cet effet est suffisamment faible pour que l'on atteigne un équilibre thermique avec l'air ambiant. L'emballement est contrôlé.



11.2.4 Mise en œuvre d'un montage à émetteur commun, montage Cascode

Cette étude est traitée en préceptorat et n'est pas abordée ici.