

CHAPITRE 4

LES ÉMETTEURS RADAR

Rédigé avec la collaboration de Messieurs Jacques DEBRAND et André BERGES

Première partie : Etat solide et tubes hyperfréquences

1	COMPOSANTS HYPERFRÉQUENCES DE PUISSANCE	3
1.1	PRINCIPES DE BASE	3
1.2	CLASSIFICATION DES TRANSISTORS DE PUISSANCE	5
1.3	MMICs (Microwaves Monolithic Integrated Circuits)	8
1.4	DIODES IMPATT	8
1.5	DIODES GUNN	8
2	ÉTAGES DE PUISSANCE	8
2.1	AMPLIFICATEURS DE PUISSANCE À TRANSISTORS	8
2.2	ÉTAGES DE PUISSANCE À DIODES IMPATT	13
2.3	PROBLÈMES DE DISSIPATION THERMIQUE	14
3	LE KLYSTRON	14
3.1	GÉNÉRALITÉS	14
3.2	PRINCIPE DE FONCTIONNEMENT	15
3.3	EXPRESSION DU MOUVEMENT DES électrons	16
3.4	APPLICATION AU KLYSTRON AMPLIFICATEUR	17
3.5	LE KLYSTRON OSCILLATEUR	20
4	LE TUBE À ONDE PROGRESSIVE	23
4.1	GÉNÉRALITÉS	23
4.2	TUBE À ONDE PROGRESSIVE DE TYPE « O »	24
4.3	PRINCIPE DE FONCTIONNEMENT DU T.P.O.	24
4.4	GAIN DU TUBE À ONDE PROGRESSIVE	25
4.5	DIFFÉRENTS TYPES DE STRUCTURES UTILISÉS	26
4.6	TUBE À ONDE PROGRESSIVE DE TYPE « M »	29
4.7	PERFORMANCES DES TUBES À ONDES PROGRESSIVES	30
5	LE CARCINOTRON	31
5.1	GÉNÉRALITÉS	31
5.2	CONSTITUTION DU CARCINOTRON	31
5.3	PRINCIPE DE FONCTIONNEMENT	32
5.4	LE CARCINOTRON « M »	32
5.5	PERFORMANCES DES CARCINOTRONS	33
6	LE MAGNÉTRON	33
6.1	GÉNÉRALITÉS	33
6.2	DESCRIPTION DU MAGNÉTRON	33
6.3	PRINCIPE DE FONCTIONNEMENT DU MAGNÉTRON	34
6.4	DIFFÉRENTS TYPES DE STRUCTURES ANODIQUES	36
6.5	MAGNÉTRONS COAXIAUX-MAGNÉTRONS AGILES	37
6.6	MAGNÉTRONS SYNCHRONISÉS	38
6.7	PERFORMANCES DES MAGNÉTRONS	39
7	LES TUBES À CHAMPS CROISÉS ET À FAISCEAU RÉENTRANT	40
7.1	PRINCIPES GÉNÉRAUX	40
7.2	AMPLIFICATEURS À CHAMPS CROISÉS	41
7.3	CARACTÉRISTIQUES COMMUNES	42

INTRODUCTION GÉNÉRALE

L'émetteur est l'élément qui fournit sa puissance au radar. C'est un organe fondamental et son choix conditionne en grande partie la conception générale de l'équipement, tout particulièrement en ce qui concerne les possibilités futures de traitement du signal.

Aussi, l'émetteur est souvent une des premières options à définir lors de l'élaboration d'un projet de radar.

On peut classer les émetteurs de radar en deux grandes catégories :

- *les émetteurs oscillateurs de puissance*, dans lesquels l'énergie est directement obtenue par mise en oscillation de l'élément de puissance ;
- *les émetteurs à chaîne d'amplification*, dans lesquels le signal à émettre est progressivement porté à la puissance convenable, par amplification dans *un* ou *plusieurs* étages.

Les émetteurs oscillateurs de puissance sont les plus anciens et les plus rustiques.

La recherche d'une grande souplesse d'utilisation et de grandes stabilités plaide en faveur des chaînes pilotées dont l'emploi se généralise dans les radars modernes. Dans certains cas cependant, l'emploi d'oscillateurs de puissance synchronisés peut s'avérer intéressant.

Dans le présent chapitre seront examinés successivement :

- *les étages de puissance état solide*, dont l'emploi se généralise soit comme étages intermédiaires de chaînes d'amplification, soit pour réaliser des émetteurs complets ;
- *les tubes hyperfréquences*, qui conservent l'exclusivité des très hautes puissances ;
- *les modulateurs*, éléments de commande (et de régulation) du tube de puissance, transposant le signal de synchronisation en un signal de commande accessible au tube hyperfréquence, et lui délivrant l'énergie nécessaire.

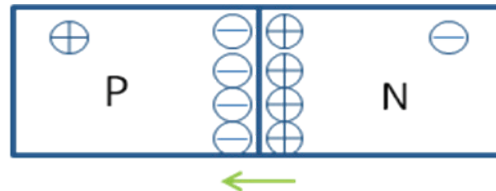
1 COMPOSANTS HYPERFRÉQUENCES DE PUISSANCE

1.1 PRINCIPES DE BASE

1.1.1 Diode à jonction PN

La jonction PN, comporte :

- une région dopée « P » (déficit d'électrons ou trous),
- une région dopée N (excès d'électrons).

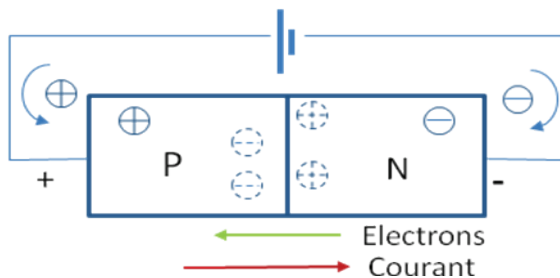


Zone de déplétion

⊖ N porteur majoritaire électron

⊕ P porteur majoritaire trou

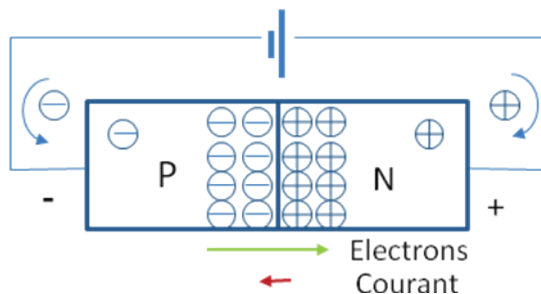
A la jonction des trous diffusent dans la région N et des électrons dans la région P. Il en résulte une charge d'espace et un champ électrique interne qui s'oppose à la diffusion des porteurs majoritaires.



En polarisation directe

(P+, N-)

l'afflux d'électrons dans la zone N et de trous dans la zone P vient détruire la barrière le courant peut circuler



En polarisation inverse

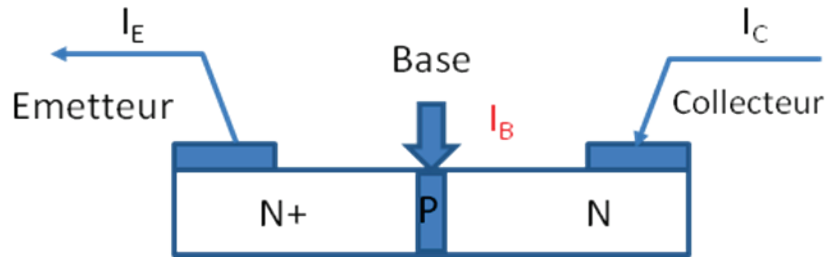
(P-, N+)

l'afflux d'électrons dans la zone P et de trous dans la zone N viennent renforcer la barrière le courant est bloqué

1.1.2 Transistor bipolaire NPN

Dans le Transistor bipolaire NPN, une base mince dopée « P » (déficit d'électrons), de faible épaisseur, agit sur la communication entre un émetteur fortement dopé N (excès d'électrons) et un collecteur faiblement dopé N.

Un faible courant communiqué à la base se propage traverse la jonction base/émetteur (polarisation directe) ce qui engendre un afflux d'électrons dans la base qui devient conductrice, d'où l'apparition d'un courant fort collecteur-émetteur.



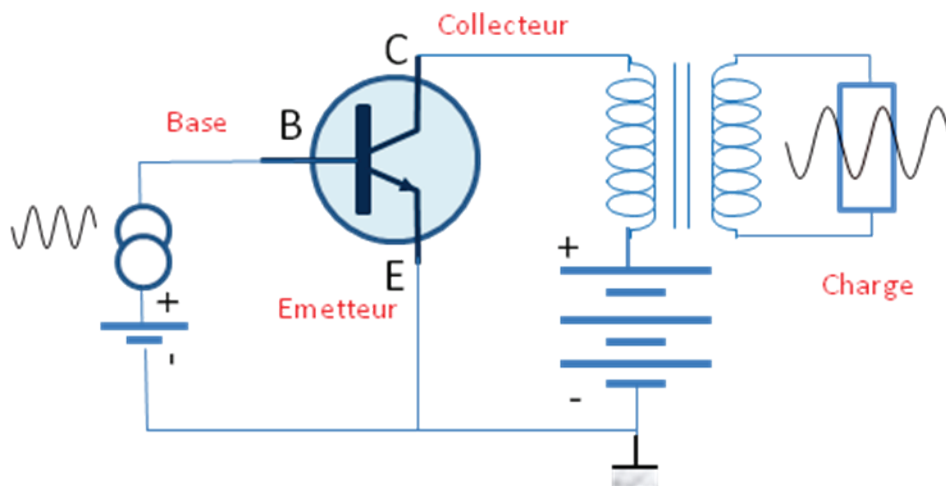
Ce transistor est donc un amplificateur de courant. Il répond aux relations :

$$I_E = I_B + I_C$$

$$I_C = \beta \cdot I_B$$

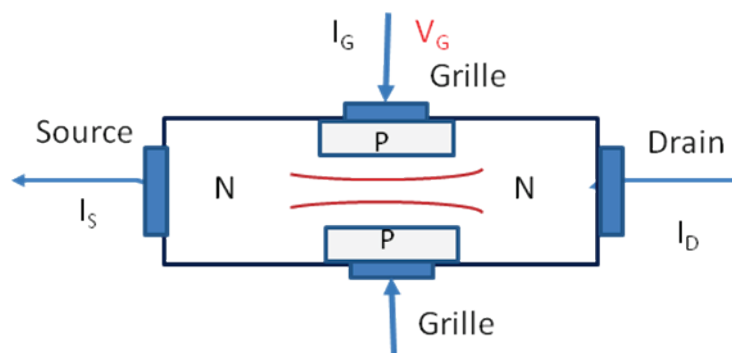
$$50 < \beta < 100$$

Il en va de même pour une variation sinusoïdale du courant de base, qui entraîne une variation sinusoïdale du courant de collecteur dont l'énergie peut être communiquée à une charge, selon le schéma type suivant.



1.1.3 Transistor FET à jonction

Dans le transistor à effet de champ (FET) à jonction. Le canal de conduction entre le drain et la source, correspond à la région N encadrée par deux régions P connectées à l'électrode de grille. Cette grille sert à polariser la jonction PN en inverse, introduisant des trous dans le canal N de façon à moduler la largeur de ce canal.



Ce transistor est donc un modulateur de courant par la tension de grille, le courant de grille étant négligeable. Il répond aux relations :

$$I_S = I_G + I_D \approx I_D$$

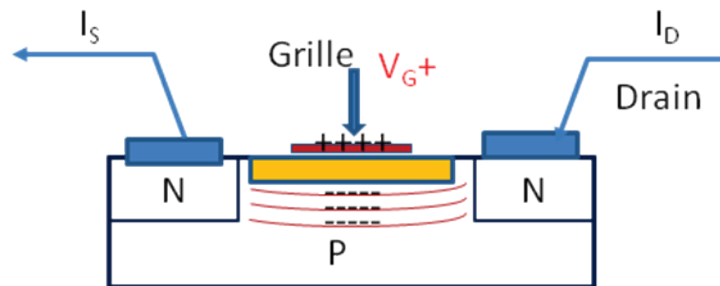
$$I_D = f(V_G)$$

1.1.4 Transistor à effet de champ MOSFET

Dans le transistor effet de champ sur un barreau dopé P, deux zones N sont diffusées pour former le drain et la source.

Une grille métallique, soumise à une tension positive, engendre une charge dans la capacité formée avec l'isolant, attirant les électrons du barreau dans la zone située entre le drain et la source et venant ainsi former un canal de conduction modulable entre la source et le drain.

Une faible tension communiquée à la grille engendre un courant fort entre la source et le drain.

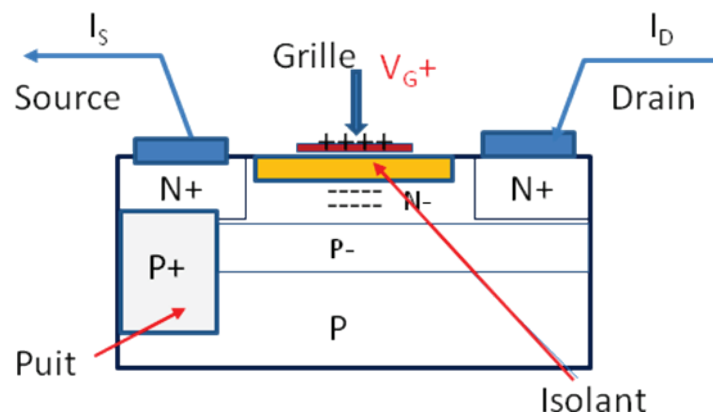


Ce transistor est également un modulateur de courant par la tension de grille, le courant de grille étant négligeable. Il répond aux relations :

$$I_S = I_G + I_D \approx I_D$$

$$I_D = f(V_G)$$

Sur ce principe général des variantes comme le LDMOS sont élaborées, pour améliorer la tenue en puissance.



D'autres variantes plus complexes, utilisant l'effet de champ, comme les HEMTs et PHEMTs, sur substrat AsGa ou encore GaN permettent des niveaux de puissance encore supérieurs.

1.2 CLASSIFICATION DES TRANSISTORS DE PUISSANCE

Les progrès des composants hyperfréquences de puissance état solide sont tels qu'ils équipent actuellement tous les étages de puissance intermédiaires et de plus en plus (radars, autoguides, fusées de proximité, sondes altimétriques...) leur étage final.

Nous distinguerons :

Les amplificateurs à transistor dans lesquels peuvent être utilisés :

- des transistors à jonction bipolaires :
 - Si BJTs (Silicon Bipolar Junction Transistor) ;
 - HBTs (Heterojonction Bipolar Transistor) sur substrat GaAs ;
- des transistors FET à jonction JFET (Junction Field Effect Transistor)
- des transistors FET à grille métallique, agissant par la création d'électrons venant former un canal de conduction modulable
 - le MOSFET (Metal Oxyde Semi-conductor Field Effect Transistor) : grille métallique isolée de la couche active par un oxyde isolant. les transistors MOS classiques sont limités en puissance du fait de leur faible tension de claquage
 - le LDMOS (Laterally diffused MOS) Il se distingue du MOSFET par un puits dopé p+, jouant le rôle de masse RF entre la source et la face arrière du composant. On atteint 1W/mm à 3.2 GHz pour une polarisation de 50 V et 2W/mm à 1 GHz pour une polarisation de 70 V.
 - le MESFET (Metal Semi-conductor Field Effect Transistor) : grille métallique à barrière Schottky.
 - les transistors MISFET.InP (Metal Insulator Semiconductor FET), sur substrat en phosphore d'Indium ;
- des transistors FET à hétérojonction, agissant par la création et le contrôle d'un gaz d'électrons dans un matériau faiblement dopé où les électrons peuvent se déplacer plus rapidement.
 - les transistors HFETs (Hétérostructure Field Effect Transistors), sur substrat en arseniure de gallium (GaAs) ; ou nitrure de gallium (GaN)
- et aussi :
 - le TEGFET (Two dimensionnal Electron Gas Field Effect Transistor) ;
 - le HFET (Heterostructure Field Effect Transistor) ;
 - le MODFET (MODulation Doped Field Effect Transistor) ;
 - le PHEMT (Pseudomorphic Hight Electron Mobility Transistor) ;
 - le PMHFET (PseudoMorphic Heterostructure Field Effect Transistor) ;
- des transistors de puissance grand gap : semi-conducteurs à grande bande interdite qui permettent d'atteindre de fortes puissances, du fait de leur fort champ électrique de claquage :
 - les transistors MESFETs SiC : MESFETs en Carbure de Silicium (SiC)
 - les transistors HEMTs GaN : HEMTs en Nitrure de Gallium.

Les amplificateurs à diode, libres ou synchronisés pour lesquels à base de diodes IMPATT (IMPact Avalanche and Transit Time), qui apportent un supplément de puissance au-delà de 50 GHz.

Les performances des principaux composants utilisés sont résumées ci après.

1.2.1 TRANSISTORS BIPOLAIRES SILICIUM SI.BJTS

Ces transistors sont des transistors à jonction « N.P.N » de technologie épitaxiale, les dopages étant effectués par diffusion d'arsenic, phosphore et bore dans le substrat silicium.

Les transistors silicium sont couramment utilisés à des fréquences allant jusqu'à 4 GHz. Des

puissances importantes (typique : 600 W en bande L, 200 W en bande S) sont obtenues en régime pulsé sur plusieurs centaines de microsecondes pour un facteur de forme de 10 à 20 %.

Des puissances de l'ordre de 1 kW crête ont été obtenues à 1 GHz pour un facteur de forme de 1 %. Les rendements obtenus varient de 45 % à 1 GHz à 30 % à 3 GHz.

1.2.2 TRANSISTORS GAAS BIPOLAIRES a HeTeROJONCTION : GAAS.HBTS

Une autre technique pour les transistors bipolaires consiste à l'utilisation d'hétérojonctions favorisant la mobilité des porteurs.

Des puissances intéressantes peuvent ainsi être obtenues jusqu'à 20 GHz (typiquement 200 W crête à 3 GHz, 10 W crête à 20 GHz) avec des rendements de l'ordre de 50 %.

1.2.3 FETs À HÉTÉROJONCTIONS – H.FETS

L'utilisation de couches stratifiées de matériaux (hétérojonctions) ou dopages (dopage planar) différents permet de constituer des canaux de « gaz d'électrons » de mobilité plus élevée, donc de conductibilité plus faibles, améliorant les performances hyperfréquences du transistor.

On obtient ainsi, jusqu'à 1 W/mm de grille à 50 MHz, par l'utilisation de doubles hétérojonctions pseudo - morphiques (Al GaAs – GaInAs – GaAs).

1.2.4 transistors « MESFET_sA_sG_a »,

Les transistors FET en arséniure de gallium sont considérés comme le composant hyperfréquence de base de 5 GHz à 30 GHz. Ils délivrent des niveaux de puissance de : 30 W à 5 GHz, 10 W à 10 GHz, 1 W à 30 GHz.

1.2.5 MISFETs

Par rapport à l'arséniure de gallium, le phosphore d'indium présente des avantages de tension de claquage et de vitesse électroniques favorables aux grandes puissances. Jusqu'à 4 W/mm de grille ont été obtenues sur des jonctions S₁O₂ – I_nP.

1.2.6 LES TRANSISTORS DE PUISSANCE GRAND GAP

Les transistors MESFETs SiC : MESFETs en Carbure de Silicium (SiC) qui peuvent atteindre 25 W à 50 W en régime pulsé (200 micro secondes facteur de forme 10 %, 5 à 10 W/mm) ; 10 W en continu, sous des tensions d'alimentation de 100 volts

Les transistors HEMTs GaN : HEMTs en Nitrure de Gallium (GaN) qui peuvent atteindre 50 à 100 W en régime pulsé (30W/mm) sous des tensions d'alimentation de 120 volts

1.2.7 TRANSISTORS HFETS SUR NITRURE DE GALLIUM

Les transistors sur nitrure de gallium, GaN - HFETs et HEMT, sont capables de niveaux de puissance supérieurs ceux qui sont possibles avec le silicium ou l'arséniure de gallium. La technologie avancée en nitrure de Gallium sur Carbure de silicium, a une conductivité thermique excellente favorisant son emploi pour des puissances élevées. Elle conduit à des impédances plus hautes plus faciles à accorder dans une large bande de fréquence.

Les puissances typiques obtenues en onde continues sont de 500 W en bande L et S, 100 W en bande C et 50 W en bande Ku (14GHz).

Les puissances crêtes obtenues en régime pulsé ($\tau = 300 \mu s$ facteur de forme 10%) sont de 800 W en bande L et S, 150 W en bande C.

1.3 MMICS (MICROWAVES MONOLITIC INTEGRATED CIRCUITS)

Le principe des MMICs est décrit en annexe au chapitre 2.

Des performances légèrement inférieures (-3 dB) aux MESFETs, les MMICs présentent cependant les avantages des circuits intégrés : faible volume, faible masse, faible coût, reproductibilité, capacité de fabrication en grandes séries.

D'un usage courant en communications (notamment spatiales), ils sont bien adaptés à l'emploi sur des antennes actives à balayage électronique et utilisés également en guerre électronique.

1.4 DIODES IMPATT

Les diodes IMPATT sont des diodes à avalanche, polarisées aux environs de leur tension de claquage, dans lesquelles l'effet d'avalanche combiné au temps de transit des électrons dans la jonction crée un courant en opposition de phase avec la tension incidente, d'où un effet de résistance négative.

Les meilleures performances sont obtenues sur des diodes GaAs en dessous de 50 GHz (10 W CW à 10 GHz) et sur des diodes silicium au-dessus de 50 GHz (1 W CW à 100 GHz).

Au-delà de 100 GHz, les puissances décroissent comme $1/f^3$. Des puissances 5 à 20 fois plus élevées peuvent être obtenues en régime pulsé pour des facteurs de forme faibles (qq%).

1.5 DIODES GUNN

Les diodes Gunn, GaAs ou InP utilisent un effet de résistance négative liée à la chute de la vitesse des électrons au-delà d'un point de polarisation. Couramment utilisées dans les oscillateurs, leurs performances de puissance inférieures de 10 dB environ à celles des diodes IMPATT les écartent aujourd'hui des oscillateurs de puissance.

2 ÉTAGES DE PUISSANCE

2.1 AMPLIFICATEURS DE PUISSANCE À TRANSISTORS

Les amplificateurs de puissance à transistors ne diffèrent pas fondamentalement dans leur structure des amplificateurs hyperfréquences décrits au chapitre 3. Devront cependant être pris en compte dans leur conception :

- la mise en parallèle de plusieurs cellules élémentaires souvent nécessaire pour obtenir la puissance désirée et qui se traduit par :
 - une diminution des impédances d'entrée et de sortie, par mise en parallèle des capacités et des résistances,
 - un problème d'équilibrage électrique (dispersion entre les cellules d'un même circuit) sur un composant dont la largeur totale peut atteindre plusieurs millimètres,
 - une complexité accrue des interconnexions sur le composant ;
- le problème de l'évacuation thermique d'une puissance importante dans un volume réduit ;
- la non linéarité du dispositif liée à un fonctionnement en signal fort, en particulier une saturation du composant se traduisant par une compression du gain par rapport au régime linéaire. Les gains aux fortes puissances sont de 7 à 13 dB par étage.

Le bilan de puissance d'un transistor se traduit par l'égalité :

$$P_e + P_{ac} = P_s + P_d$$

où :

- P_e : puissance hyperfréquence d'entrée
- P_s : puissance hyperfréquence de sortie
- P_{ac} : puissance de l'alimentation continue
- P_d : puissance dissipée sous forme thermique

Par définition, le gain est le rapport :

$$G = \frac{P_s}{P_e}$$

On définit sur un transistor de puissance un *rendement en puissance ajoutée* R_{PAE} , tel que :

$$R_{PAE} = \frac{P_s - P_e}{P_{ac}} = 1 - \frac{P_d}{P_{ac}}$$

directement caractéristique de la puissance dissipée par le transistor.

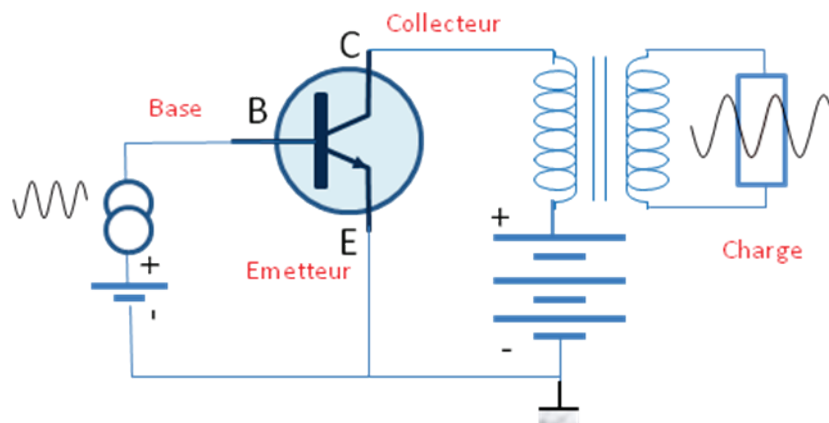
Ce rendement est inférieur au rendement de drain, $R_D = P_s/P_{ac}$ auquel il est relié par la relation :

$$R_{PAE} = R_D \left(1 - \frac{1}{G} \right)$$

Le rendement de drain dépend du mode de fonctionnement du transistor.

En classe A (point de polarisation centré dans la caractéristique de conduction du transistor) ce rendement atteint au maximum 50 %.

Un schéma type de mode de fonctionnement utilisant un transistor bipolaire, est le suivant :



Si R est la résistance de charge vue du primaire du transformateur (ici $R = 50 \Omega$) et V_0 la tension d'alimentation, la tension réellement appliquée au transistor est :

$$V_{CE} = V_0 - R \cdot I_C$$

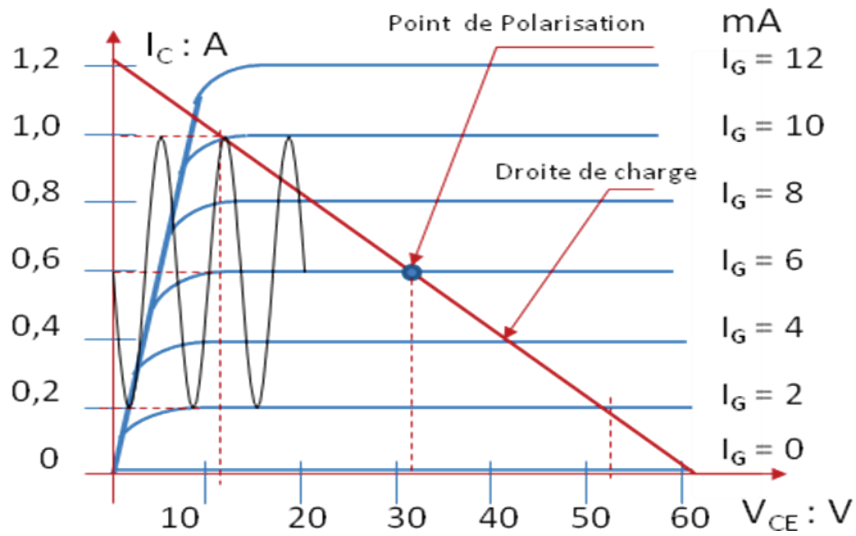
ce qui définit la droite de charge par 2 points

$$I_C = 0 ; V_{CE} = V_0 \text{ et } I_C = V_0/R ; V_{CE} = 0$$

Le point de polarisation est choisi de telle manière que la tension de grille restera toujours

négative, dans l'exemple traité :

$$I_G = 8 \text{ mA}$$



Au point de fonctionnement :

$$I_{C0} = 0,6 \text{ A} ; V_{CE0} = 32 \text{ Volts.}$$

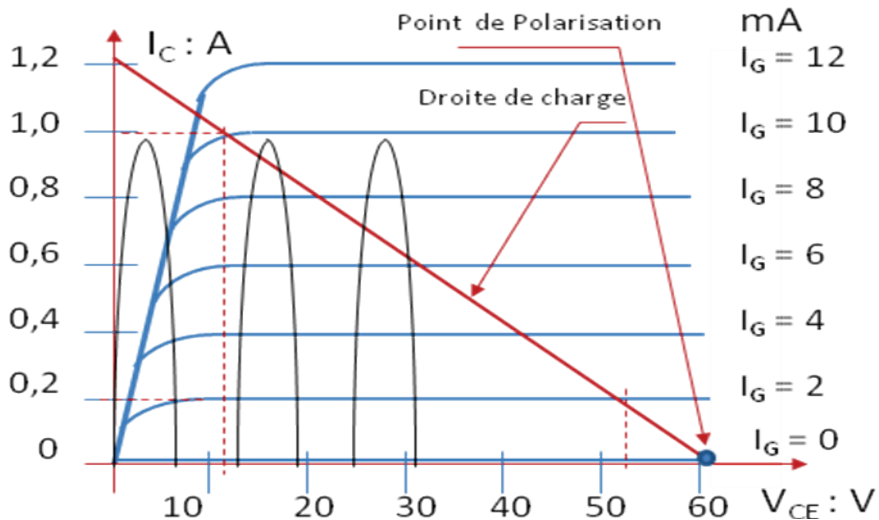
« I_G » varie sinusoïdalement de $\pm 4 \text{ mA}$

« I_C » varie de $\pm 0,4 \text{ A}$ « V_{CE} » de $\pm 20 \text{ volts}$

La puissance efficace fournie est

$$W = R \cdot I_C^2 / 2 = 4 \text{ W}$$

En classe B (point de polarisation à courant nul) le fonctionnement est le suivant :

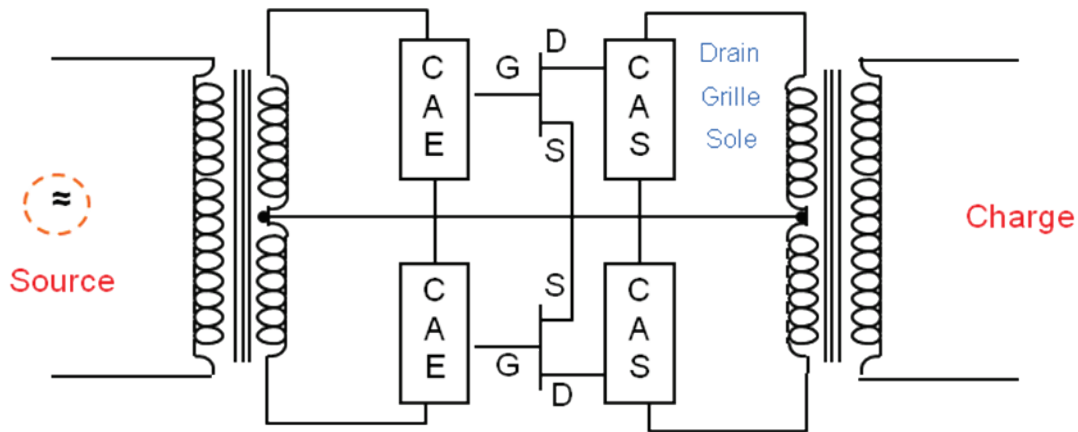


« I_C » varie de 0 à 1 A « V_{CE} » de 10 à 60 volts soit une variation de 50 volts

La puissance crête fournie est

$$W = R \cdot I_C^2 = 50 \text{ W}$$

La classe B est souvent associée à un montage « Push Pull » dont un schéma de principe est donné ci après :

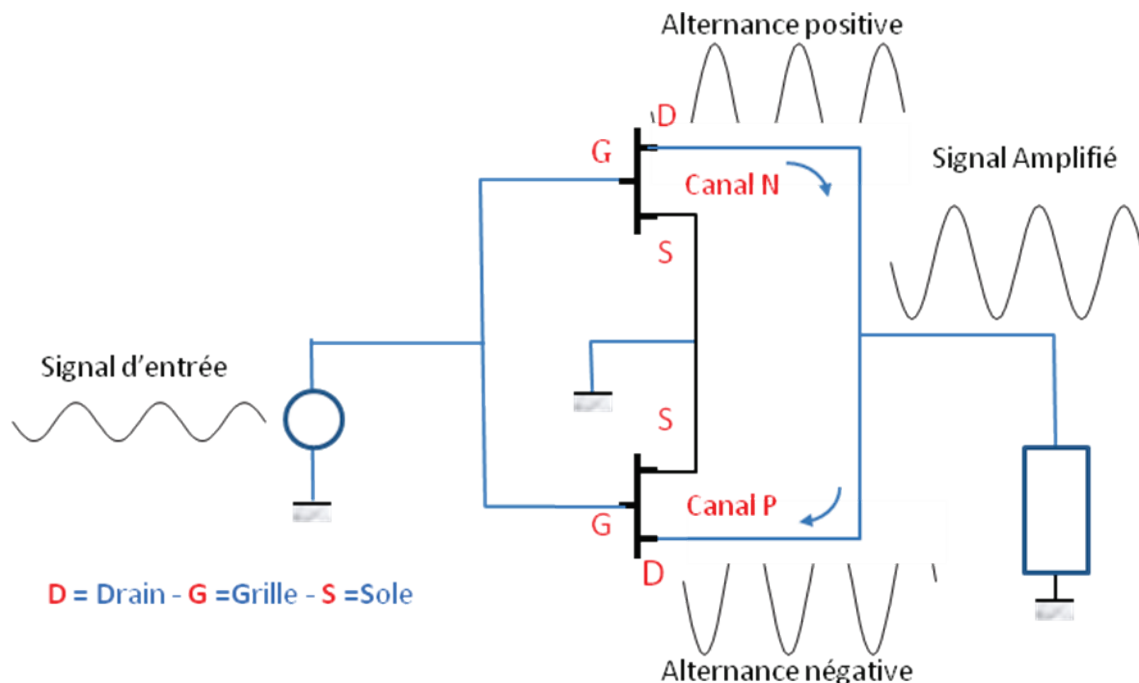


CAE : Circuit d'adaptation entrée - CAS : Circuit d'adaptation de sortie

Le rôle des transformateurs est double :

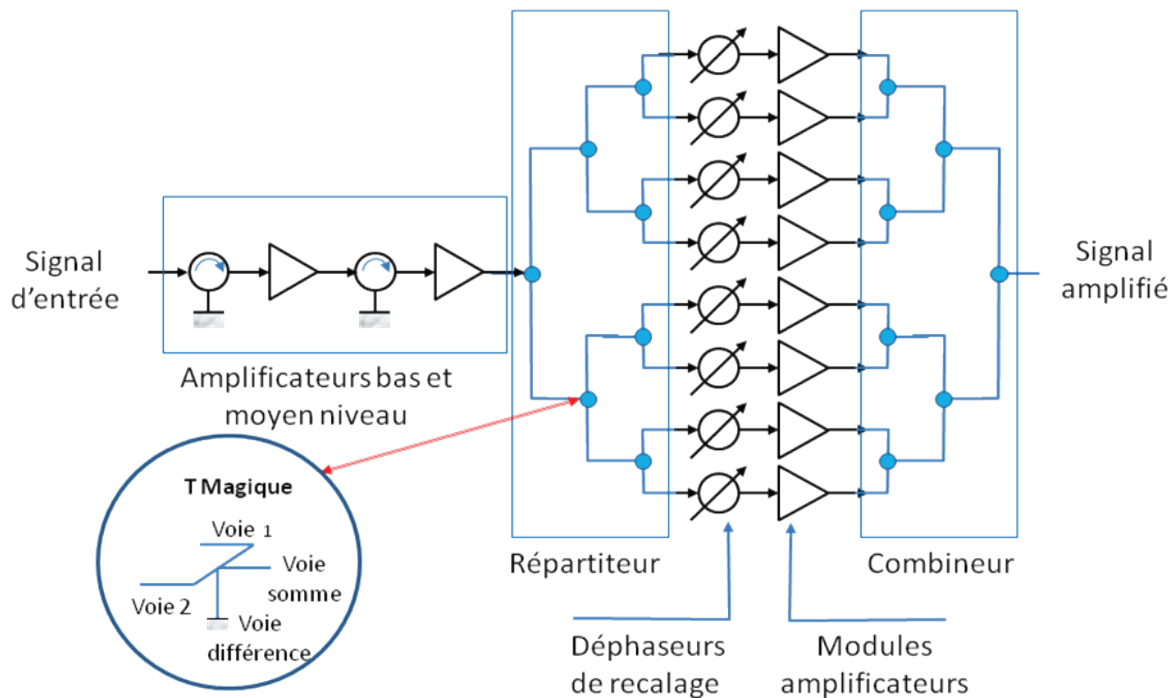
- présenter à chaque transistor des signaux positifs en inversant le sens des courants sur l'une des deux voies, qui verra positivement les alternances négatives ;
- adapter l'impédance d'entrée et de sortie du dispositif (par exemple à 50 Ω)

Il existe d'autres types de montage, utilisant des transistors complémentaires, l'un conduisant naturellement les alternances positives, l'autre les alternances négatives, sans faire emploi à des transformateurs. Un schéma de principe est donné ci après :



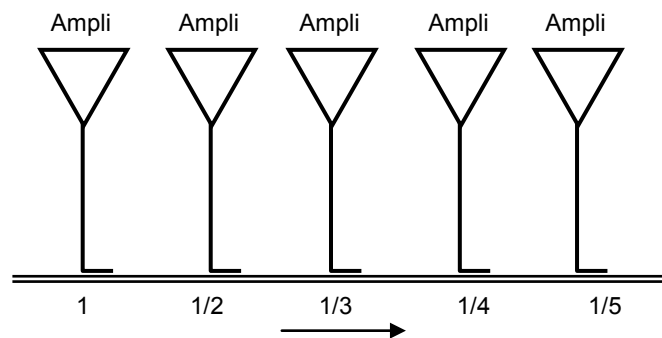
La décomposition en série de Fourier du signal démontre un rendement maximum de $\pi / 4$. Avec le transistor de l'exemple précédent, la puissance efficace obtenue serait de 25 W.

L'augmentation des puissances peut conduire à paralléliser des amplificateurs à l'aide de différentes structures : structures en arbre utilisant des coupleurs 3 dB ou des T magiques en anneau :



Combineurs N voies de différentes structures : Planar, Wilkinson, Radial, Coaxial... permettant l'addition d'un grand nombre de voies,

Structures en lignes ou chaque étage amplificateur de rang N est relié à la ligne par un coupleur de coefficient $1/N$:



Sur des modules préalablement adaptés, ces recombinaisons de puissance se font avec un bon rendement (0,8 à 0,9) permettant ainsi de réaliser des émetteurs état solide modulaires de puissances importantes, présentant en outre des caractéristiques intéressantes de fiabilité et de maintenabilité, les modules pouvant être changés sans arrêter le fonctionnement de l'ensemble.

A titre d'exemple, on peut citer en bande L, un émetteur 10 kW crête 1 kW moyen réalisés par l'association de 20 modules de 600 W crête dans un combineur à 20 voies, utilisant des transistors bipolaires silicium comme élément de puissance.

Ainsi que d'autres réalisations (Estimations) :

- Emetteur bande S état solide 16 modules 800 W moyen ;
- Emetteur bande L état solide 32 modules 2 KW moyen ;
- Utilisation de modules état solide dans les antennes actives :
 - Antenne radar sol bande S : environ 200 modules (10 à 20 KW moyen ?) ;
 - RBE2 bande X : environ 900 modules émission – réception (1KW moyen ?).

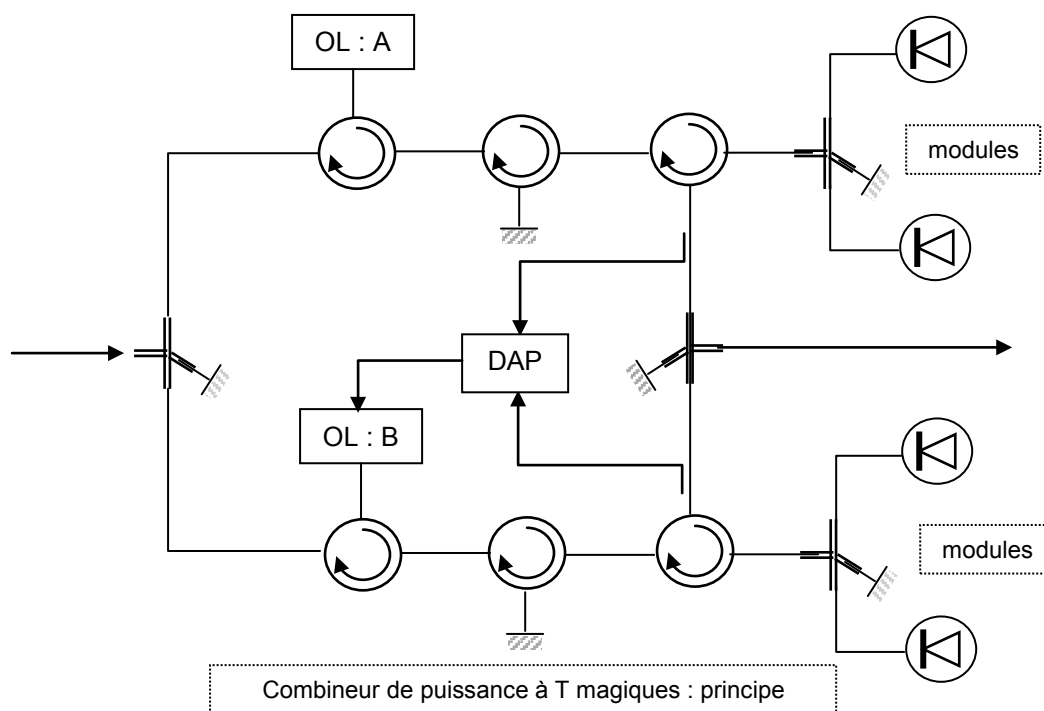
2.2 ÉTAGES DE PUISSANCE À DIODES IMPATT

Les étages de puissance à diodes IMPATT sont utilisés à des niveaux du watt à quelques dizaines de watt, dans des applications particulières en ondes centimétriques (autodirecteurs) et pour des émissions cohérentes en ondes millimétriques. L'élément de base de ces étages est un oscillateur de puissance synchronisé, les gains par étage étant de l'ordre de 10 dB.

Les puissances disponibles par diodes étant limitées on sera amené à combiner ces puissances, soit par addition de modules à travers une structure adaptée, soit par combinaison de diodes dans un même module.

Dans le premier cas, chacun des modules devra être adapté à l'impédance de charge de manière à obtenir la meilleure addition des puissances. L'emploi des diodes en boîtiers pré accordés facilite ces problèmes d'adaptation. Les structures de recombinaison doivent rattraper les défauts résiduels d'adaptation entre modules, des combinaisons de T magiques et circulateurs peuvent être utilisées à cette fin.

Dans l'exemple ci après quatre modules de puissance sont combinés deux à deux à travers des T magiques chaque oscillateur ainsi constitué étant piloté par des oscillateurs intermédiaires (A et B) eux-mêmes synchronisés par une référence unique.



Les deux chaînes ainsi constituées sont ensuite combinées, leur équilibrage étant assuré par une boucle de phase. Les circulateurs intermédiaires servent d'isolateurs entre étages.

Les associations de ce type donnent de très bons résultats de rendement et stabilité de phase ; elles deviennent rapidement très complexes dès que le nombre d'éléments à associer devient important.

Une autre méthode consiste à associer directement un certain nombre de diodes dans un combineur, formé d'une cavité résonnante, couplée à sa périphérie à des cavités périphériques recevant les diodes oscillatrices. Le volume résonnant étant par nature surdimensionné, des adaptateurs complémentaires, filtre de mode, cavité auxiliaire sont utilisés pour piéger les modes d'oscillations parasites et ainsi favoriser le mode d'oscillation principal qui est couplé au circuit de sortie.

Un grand nombre de diodes peuvent être aussi associées simultanément dans un oscillateur unique avec cependant un rendement de combinaison plus faible du fait des problèmes d'adaptation et de parité entre les diodes rendus délicats du fait de leur association directe.

À titre d'exemple, on peut citer en bande Ka un combineur à 12 diodes IMPATT qui délivre une puissance de 25 W moyens en impulsions avec un facteur de forme de 25 %.

2.3 PROBLÈMES DE DISSIPATION THERMIQUE

Un problème majeur des amplificateurs de puissance état solide est la nécessité d'évacuer la chaleur engendrée à l'intérieur du composant, plus précisément au niveau de la jonction, donc dans un très petit volume.

Indépendamment des procédés propres de fabrication des puces pour favoriser l'évacuation des calories, celles-ci devront être intégrées à leurs boîtiers ou à leur circuit d'utilisation en prévoyant des drains thermiques propres à les relier à des structures aptes à évacuer la chaleur. Aux faibles puissances, un refroidissement par convection naturelle peut être suffisante sur des composants munis de radiateurs individuels.

Une convection forcée devra être envisagée pour des puissances supérieures (à partir de 50 W). Pour des systèmes plus compacts ou plus complexes, des dissipateurs relieront les composants à une structure elle-même refroidie par convection forcée ou par liquide.

Le problème du refroidissement est un problème majeur et conduit à rechercher des rendements très élevés (> 50 %) notamment dans le cas des antennes actives.

3 LE KLYSTRON

3.1 GÉNÉRALITÉS

Le *klystron* est un tube à modulation de vitesse du faisceau électronique, il peut être soit amplificateur, soit oscillateur (*Klystron reflex*). Son principe est basé sur l'interaction d'un faisceau d'électrons, modulé en vitesse ou en densité, avec un ensemble de cavités.

Le tube se compose principalement :

- d'une cathode en général portée à un potentiel négatif et de son dispositif de chauffage ;
- d'une anode à la masse ;
- éventuellement d'une grille de modulation en impulsions du tube (cette modulation peut s'effectuer également en agissant directement sur le potentiel relatif cathode-anode) ;
- l'ensemble cathode, anode, grille constitue le canon du tube, duquel est issu un faisceau d'électrons cylindrique et homogène ;
- d'un dispositif de focalisation qui empêche le faisceau issu du canon de diverger sous l'action de la force répulsive des électrons entre eux (*charge d'espace*) : cette focalisation est réalisée à l'aide d'un champ magnétique axial obtenu à partir d'un électro-aimant ou d'un aimant permanent ;
- de deux cavités séparées par un espace de glissement ;
- d'un collecteur qui recueille le faisceau électronique à l'extrémité du tube ;
- d'un dispositif de refroidissement du tube, à air ou à eau, suivant le niveau de puissance.

La dimension et la forme des cavités sont déterminées afin que : d'une part elles soient accordées sur la fréquence de l'onde à amplifier, et d'autre part, elles produisent un champ électrique intense au voisinage de leur axe, c'est-à-dire, dans la zone où le faisceau électronique les traverse.

3.2 PRINCIPE DE FONCTIONNEMENT

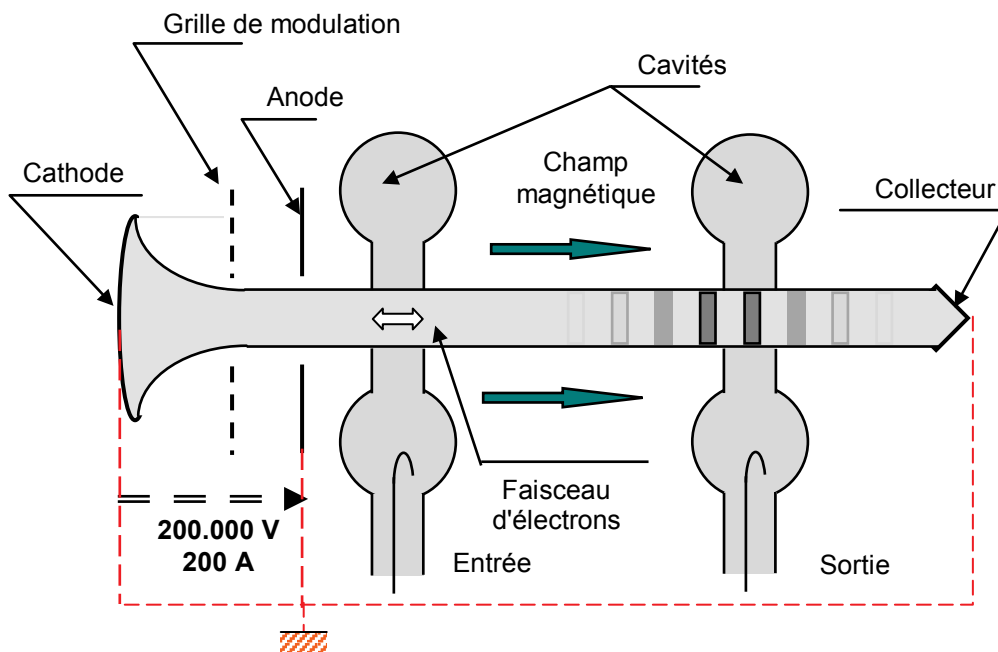
A l'arrivée dans la cavité d'entrée, le faisceau électronique issu du canon est monocinétique et sa vitesse « v_0 » est déterminée par la différence de potentiel V_0 entre anode et cathode. Le signal à amplifier qui arrive dans la cavité d'entrée produit un champ électrique sinusoïdal important dans la zone de traversée du faisceau électronique.

Nous supposons que le temps de transit des électrons dans les cavités est négligeable, c'est-à-dire que la variation du champ électrique auquel ils sont soumis lors de leur passage dans la cavité est faible.

Pendant une demi-période du signal à amplifier, les électrons du faisceau rencontreront au droit de la cavité d'entrée un champ électrique de même sens que leur vitesse, ils seront donc soumis à une force qui aura tendance à les ralentir car ils sont de charge négative et :

$$\vec{F} = e \cdot \vec{E}$$

Pendant la demi-période suivante, le champ électrique deviendra de sens opposé à la vitesse des électrons qui seront alors accélérés.

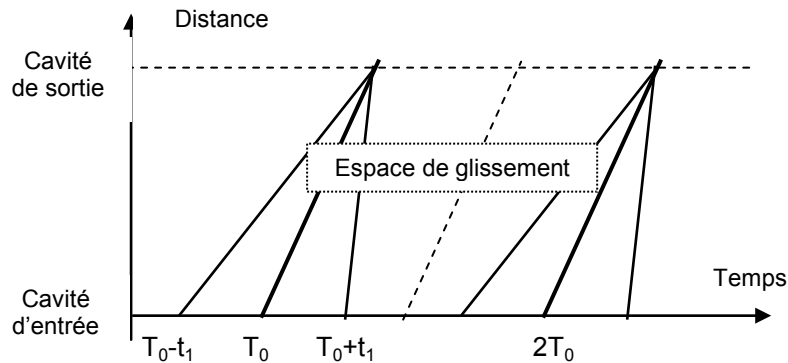


Suivant leur instant d'arrivée dans la cavité d'entrée, les électrons du faisceau sont donc *accélérés* ou *freinés*, le faisceau monocinétique à l'origine est modulé en vitesse à la sortie de la cavité. Cette modulation de vitesse provoque, au cours du trajet des électrons dans l'espace de glissement, une modulation de densité ; les électrons ont tendance à se grouper en paquets. Ce phénomène est illustré à l'aide du diagramme d'*Applegate* représenté ci-après :

Soit un électron arrivant dans la cavité d'entrée à l'instant T_0 où le champ électrique est nul, et décroissant, sa vitesse n'est pas modifiée et reste égale à v_0 (*pente de la droite dans le diagramme*) ; cet électron sera pris comme électron de référence.

Un électron arrivant dans la cavité d'entrée à l'instant T'_0 soit un peu avant l'électron de

référence, rencontre un champ retardateur et voit sa vitesse diminuer, la pente de la droite est donc plus faible ; cet électron aura tendance à être rattrapé par l'électron de référence.



De même un électron arrivant à l'instant T₀ soit un peu après l'électron de référence, rencontre un champ accélérateur et voit sa vitesse augmenter, la pente de la droite est donc plus grande ; cet électron a tendance à rattraper l'électron de référence. Le phénomène se reproduit d'une manière identique à la période suivante.

On remarque sur le diagramme que des électrons ayant quitté la cavité d'entrée à des instants différents se retrouvent à une distance d de cette cavité au même instant.

Les électrons ont donc tendance à se grouper en paquets autour de l'électron de référence à une certaine distance d de la cavité d'entrée.

Si à cette distance d, on place la cavité de sortie, celle-ci va être traversée par des paquets d'électrons à des instants séparés de la période de l'onde à amplifier. Ces paquets d'électrons induisent dans la cavité de sortie, accordée sur la pulsation du signal à amplifier, un champ hyperfréquence à cette même pulsation.

3.3 EXPRESSION DU MOUVEMENT DES ELECTRONS

Les électrons accélérés sous la tension V₀ ont acquis, avant d'atteindre la cavité d'entrée, une vitesse v₀ donnée par :

$$v_0 = a \sqrt{V_0} \quad \text{avec :} \quad a = \sqrt{-\frac{2e}{m}} \approx 6 \cdot 10^5$$

(e charge de l'électron et m masse de l'électron, V₀ tension d'accélération du faisceau).

La tension dans la cavité d'entrée est de la forme V₁.sin ωt.

Cette tension agit sur les électrons de telle sorte qu'à la sortie de la cavité, ils ont une vitesse :

$$v = a \sqrt{V_0 + V_1 \cdot \sin \omega t} = a \cdot \sqrt{V_0} \sqrt{1 + \frac{V_1}{V_0} \cdot \sin \omega t}$$

En posant α = V₁/V₀ et en supposant V₁ << V₀, ce qui est pratiquement, toujours le cas, on obtient en faisant un développement limité :

$$v = v_0 \left(1 + \frac{\alpha}{2} \cdot \sin \omega t \right)$$

Si on désigne par t₁ l'instant où l'électron passe dans la cavité d'entrée et, par t₂ l'instant où il a parcouru une distance d par rapport à cette cavité, on peut écrire :

$$t_2 - t_1 = \frac{d}{v} = \frac{d}{v_0 \left(1 + \frac{\alpha}{2} \cdot \sin \omega t_1 \right)}$$

équivalent à pour α petit :

$$t_2 - t_1 = \frac{d}{v_0} \left(1 - \frac{\alpha}{2} \cdot \sin \omega t_1 \right)$$

$d/v_0 = t_0$, représente le temps mis par l'électron de référence pour parcourir la distance d .

L'instant d'arrivée t_2 des électrons à la distance d de la cavité d'entrée, en fonction de l'instant de départ t_1 , est donné par :

$$t_2 = t_1 + t_0 \left(1 - \frac{\alpha}{2} \cdot \sin \omega t_1 \right)$$

En multipliant par ω et en posant $\omega t = \tau$ *angle de transit*, l'expression devient :

$$\tau_2 = \tau_1 + \tau_0 \left(1 - \frac{\alpha}{2} \cdot \sin \tau_1 \right) = \tau_1 + \tau_0 - \frac{\alpha \tau_0}{2} \cdot \sin \tau_1$$

soit en posant $\alpha \cdot \tau_0 / 2 = k$:

$$\tau_2 = \tau_1 + \tau_0 - k \cdot \sin \tau_1$$

3.4 APPLICATION AU KLYSTRON AMPLIFICATEUR

3.4.1 Courant dans la cavité de sortie

Les paquets d'électrons traversant la cavité de sortie y induisent un champ hyperfréquence, la cavité étant accordée sur la pulsation ω , seule la composante « ω » de l'onde induite est à considérer. Cette composante sera obtenue en cherchant le terme à la pulsation ω de la décomposition en série de Fourier du courant I_2 traversant la cavité de sortie, soit :

$$I_2 = a_0 + \sum_1^{\infty} a_n \cos n \omega t_2$$

donc à la pulsation ω :

$$I_2(\omega) \cong a_1 \cdot \cos \omega t_2$$

avec :

$$a_1 = \frac{\omega}{\pi} \int_{-\frac{\pi}{\omega}}^{\frac{\pi}{\omega}} I_2 \cdot \cos \omega t_2 \cdot dt_2$$

ou :

$$a_1 = \frac{1}{\pi} \int_{-\pi}^{\pi} I_2 \cdot \cos \tau_2 \cdot d\tau_2$$

Or, I_0 étant le courant qui traverse la cavité d'entrée, on peut écrire la conservation de la charge :

$$I_0 \cdot dt_1 = I_2 \cdot dt_2$$

ou en multipliant par ω :

$$I_o \cdot d\tau_1 = I_2 \cdot dt_2$$

Donc :

$$a_1 = \frac{1}{\pi} \int_{-\pi}^{\pi} I_o \cdot \cos \tau_2 \cdot d\tau_1$$

et en remplaçant τ_2 par sa valeur en fonction de τ_1 :

$$a_1 = \frac{1}{\pi} \int_{-\pi}^{\pi} I_o \cdot \cos (\tau_1 + \tau_o - k \cdot \sin \tau_1) \cdot d\tau_1$$

ou en prenant comme origine τ_o , ce qui conduira à écrire : $I_2(\omega) \cong a_1 \cdot \cos (\omega t_2 - \tau_o)$:

$$a_1 = \frac{I_o}{\pi} \int_{-\pi}^{\pi} \cos (\tau_1 - k \cdot \sin \tau_1) \cdot d\tau_1$$

$$a_1 = \frac{2 I_o}{\pi} \cdot \int_0^{\pi} \cos (\tau_1 - k \cdot \sin \tau_1) \cdot d\tau_1$$

Or, on démontre que :

$$\frac{1}{\pi} \int_0^{\pi} \cos (\tau_1 - k \cdot \sin \tau_1) \cdot d\tau_1 = J_1(k)$$

qui est la fonction de *Bessel* d'ordre 1.

Donc :

$$a_1 = 2 I_o \cdot J_1(k)$$

et :

$$I_2(\omega) = 2 I_o \cdot J_1(k) \cdot \cos (\omega t - \tau_o)$$

Le maximum de $J_1(k)$ est obtenu pour $k = 1,84$ et vaut 0,58. Le maximum du courant dans la cavité de sortie est donc :

$$I_2(\omega) = 1,16 I_o \cdot \cos (\omega t - \tau_o)$$

avec :

$$k = \frac{\alpha \tau_o}{2} = \frac{V_1}{2 V_o} \cdot \frac{\omega \cdot d}{v_o} = 1,84$$

ainsi pour V_1 , ω , V_o donnés, il existe une valeur optimum de la longueur de l'espace de glissement : d .

3.4.2 Rendement du klystron amplificateur

Le rendement η est le rapport de la puissance P_2 disponible (*en alternatif*) dans la cavité de sortie, à la puissance consommée P_o (*en continu*) :

$$\eta = \frac{P_2}{P_0} = \frac{\frac{V_2 \cdot I_2}{2}}{V_0 \cdot I_0} = \frac{V_2 \cdot I_2}{2 V_0 \cdot I_0}$$

Le rendement maximum est obtenu pour :

- I_2 maximum égale $1,16 I_0$,
- V_2 maximum égale V_0 car si V_2 dépasse V_0 les électrons seront renvoyés vers la cathode.

$$\eta_{\max} = \frac{1,16 \cdot V_0 \cdot I_0}{2 V_0 \cdot I_0} = 0,58$$

Le rendement maximum théorique du klystron amplificateur à deux cavités serait donc 58 % .

En réalité, on atteint seulement des rendements de l'ordre de 25 % avec des klystrons à deux cavités, ceci est dû :

- à ce que V_2 est inférieur à V_0 ,
- à l'influence du temps de transit des électrons dans les cavités,
- au phénomène de charge d'espace qui a tendance à défavoriser le groupement des électrons,
- aux rendements des couplages *entrée* et *sortie*.

Pour améliorer le rendement du klystron, on peut placer entre la cavité d'entrée et celle de sortie une ou plusieurs cavités intermédiaires non chargées. Ces cavités ont pour but d'améliorer le groupement des électrons qui fournissent alors une puissance plus importante à la cavité de sortie.

Le faisceau électronique issu de la cavité d'entrée est faiblement modulé en densité lorsqu'il atteint une cavité intermédiaire, il y induit un champ plus important que celui dû au signal d'entrée, cette cavité non chargée ayant un grand coefficient de surtension. Ce champ important réagit sur le faisceau en provoquant un renforcement du groupement des électrons, groupement qui s'améliore au passage d'une deuxième cavité intermédiaire et ainsi de suite.

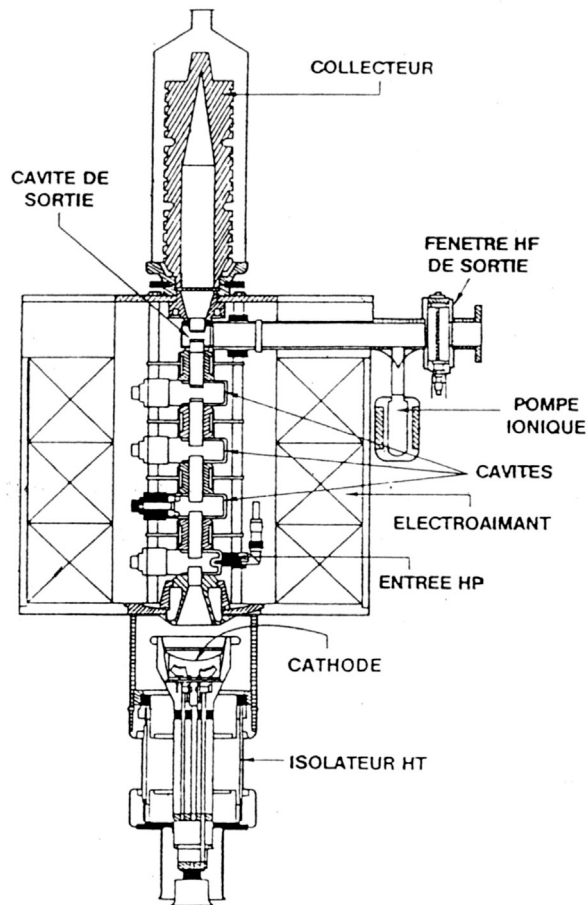
En fait, le nombre de cavités intermédiaires est limité par des considérations pratiques, liées aux pertes propres introduites par chaque cavité supplémentaire. On n'utilise guère plus de 2 à 3 cavités intermédiaires sauf dans des cas particuliers de klystron à très large bande (jusqu'à 8 cavités réalisées). La figure page suivante présente la coupe d'un klystron à 5 cavités.

3.4.3 Performances des klystrons amplificateurs

Les caractéristiques principales des klystrons sont, en bandes P à C :

- un gain important de l'ordre de 45 à 50 dB,
- une bande passante comprise entre 5 et 12 % voire 15 %,
- des rendements de 35 à 55 % en impulsions, jusqu'à 70 % en CW,
- des puissances moyennes allant jusqu'au MW en CW, typiquement 10 à 200 kW à la large bande,
- des puissances crête liées aux tensions d'alimentation : 30 MW avec 300 kV, 10 MW avec 150 kV, 1 MW avec 90 kV,
- des rotations de phase importantes de 5 à 10° par % de tension nécessitant une régulation de leurs alimentations.

Les klystrons sont utilisés de 400 MHz à 30 GHz.



klystron à 5 cavités

3.5 LE KLYSTRON OSCILLATEUR

Pour réaliser un oscillateur avec un klystron à deux cavités, on peut prélever une partie de l'énergie dans la deuxième cavité et la ramener dans la première cavité. Lors du couplage ainsi réalisé entre ces deux cavités, si l'on respecte certaines conditions d'amplitude et de phase, le tube peut osciller.

La sortie s'effectue alors par une boucle de couplage dans la première cavité.

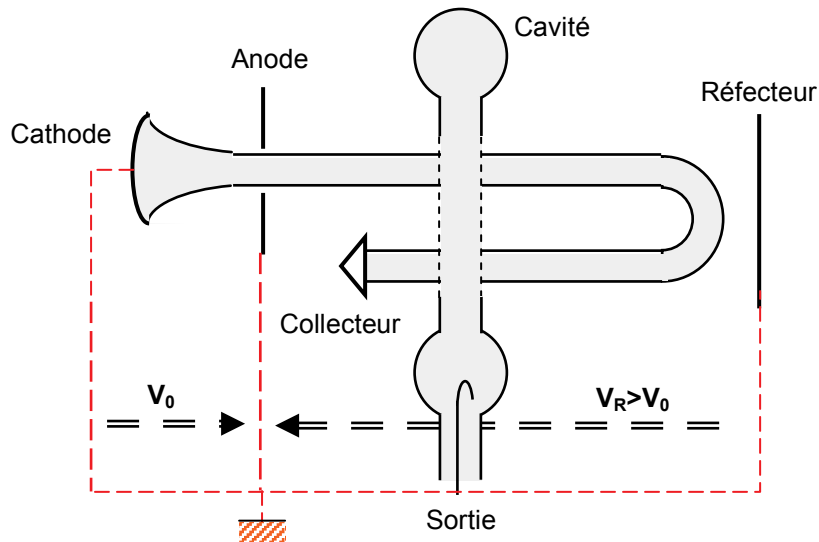
Cependant, on utilise plus fréquemment un klystron à une seule cavité, dit *klystron reflex*, dont nous allons étudier le fonctionnement.

3.5.1 Principe du klystron reflex

Le klystron reflex est constitué :

- d'un canon identique au klystron amplificateur,
- d'une cavité,
- d'un réflecteur porté à un potentiel plus négatif que la cathode.

A la sortie de la cavité, le faisceau entre dans une zone où il est soumis à une différence de potentiel négative, il est donc freiné et comme $|V_R| > |V_o|$, à une certaine distance de la cavité, sa vitesse va s'annuler, le faisceau retourne alors vers la cavité où il y passe une deuxième fois mais en sens inverse.



Soit x l'abscisse d'un électron par rapport à la sortie de la cavité, λ la distance cavité réflecteur et v_0 la vitesse du faisceau à la sortie de la cavité, le mouvement du faisceau électronique dans l'espace cavité réflecteur est donné par (« d » : la distance anode réflecteur) :

$$x = \frac{1}{2} \cdot \gamma t^2 + v_0 t$$

$$\gamma = \frac{F}{m} = \frac{e \cdot E}{m} = \frac{-e \cdot V_R / d}{m}$$

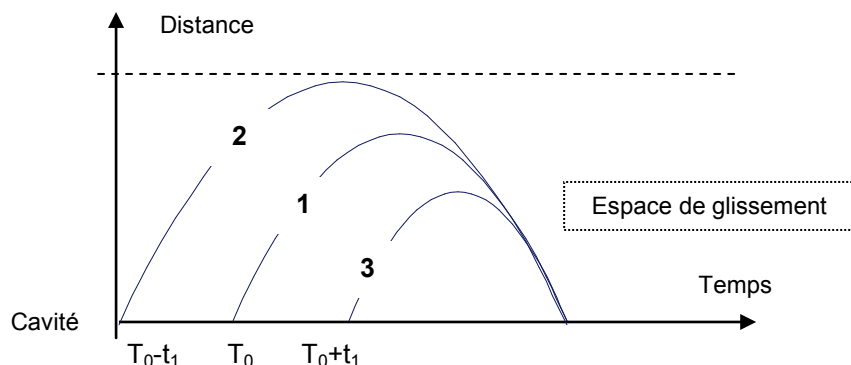
$$x = -\frac{1}{2} \cdot \frac{e}{m} \cdot \frac{V_R}{d} \cdot t^2 + v_0 t$$

Le temps que met l'électron pour effectuer l'aller et retour est donc :

$$t_0 = \frac{2m}{e} \cdot \frac{d \cdot v_0}{V_R}$$

Le mouvement des électrons dans l'espace cavité-réflecteur est une fonction parabolique du temps ; l'ordonnée du sommet de la parabole dépend de la vitesse initiale v_0 et de la tension V_R . Si nous supposons que des oscillations existent dans la cavité (comme dans tout oscillateur, il faut qu'à un instant donné une perturbation donne naissance à des oscillations qui s'amplifient jusqu'à atteindre un niveau de stabilité), le champ créé au droit du faisceau va provoquer une modulation de vitesse du faisceau électronique comme précédemment.

On peut représenter le mouvement des électrons sur un diagramme identique au diagramme d'Applegate, les trajectoires étant des paraboles au lieu de droites.



L'électron quittant la cavité à l'instant T_0 où le champ est nul et croissant conserve sa vitesse initiale v_0 et décrit la parabole 1 (*électron de référence*).

- L'électron quittant la cavité un peu avant l'électron de référence et rencontrant un champ accélérateur, voit sa vitesse augmenter, il s'approche plus près du réflecteur et met donc plus longtemps pour revenir vers la cavité ; cet électron sera rattrapé par l'électron de référence ; il décrit la parabole 2.
- Un électron quittant la cavité un peu après l'électron de référence à un instant où le champ est retardateur, voit sa vitesse diminuer, il s'approche moins près du réflecteur et met donc un temps moins long pour revenir vers la cavité, cet électron a tendance à rattraper l'électron de référence, il décrit la parabole 3.

Nous voyons donc que des électrons quittant la cavité à des instants différents y reviennent au même instant, il y a aussi groupement autour de l'électron de référence qui dans ce cas, contrairement au précédent, se trouve à l'instant où le champ passe d'accélérateur à retardateur. Ce sont, par conséquent, des paquets d'électrons espacés d'une période qui traversent la cavité à leur retour.

Pour que ces paquets d'électrons cèdent de l'énergie à la cavité, il faut qu'ils soient freinés et donc qu'ils traversent cette cavité à un instant où le champ électrique est retardateur, c'est-à-dire, accélérateur pour le faisceau qui entre dans la cavité.

Ainsi, l'optimum sera obtenu lorsqu'un électron parti à un instant où le champ nul passe d'accélérateur à retardateur, y revient lorsque le champ est accélérateur maximum, il se sera donc écoulé un temps égal au $3/4$ de la période à un multiple de la période près.

Soit :

$$t_0 = \frac{3}{4}T + nT = \left(\frac{3}{4} + n\right)T$$

ou en multipliant par ω :

$$\tau_0 = 2\pi \left(\frac{3}{4} + n\right)$$

Nous voyons apparaître ici, la notion de modes d'oscillations possibles du klystron, chaque mode correspond à une valeur de n (0, 1, 2, etc.). Or :

$$t_0 = \frac{2m}{e} \cdot \frac{d \cdot v_0}{V_R}$$

est inversement proportionnel à V_R .

Les différents modes apparaîtront lorsque l'on agira sur V_R pour certaines plages de variation de V_R .

On démontre que le rendement théorique optimal du klystron reflex (pour une tension de modulation proche de la tension de faisceau) s'écrit :

$$\eta_{\max} \approx - \frac{2,4 \sin \tau_0}{\tau_0}$$

soit théoriquement 50 % pour $\tau_0 = 3\pi/2$. Dans la réalité, les rendements des klystrons reflex ne dépassent guère quelques %.

La puissance de sortie et la plage de variations de V_R décroissent lorsque l'ordre du mode croît. Par contre la fréquence d'oscillation et l'excursion de fréquence à l'intérieur d'un mode croissent avec l'ordre de ce mode.

3.5.2 Performances des klystrons reflex

La gamme d'utilisation des klystrons reflex se situe entre 1 000 MHz et 100 000 MHz, pour des puissances de sortie variant de quelques watts aux fréquences basses, à quelques centaines de milliwatts aux fréquences élevées. Les tensions de fonctionnement utilisées vont de quelques centaines de volts à 1 kV.

Autrefois utilisés comme oscillateurs locaux ou sources de modulable en fréquence, le klystron réflex se voit aujourd'hui détrôné par les composants état solide.

4 LE TUBE À ONDE PROGRESSIVE

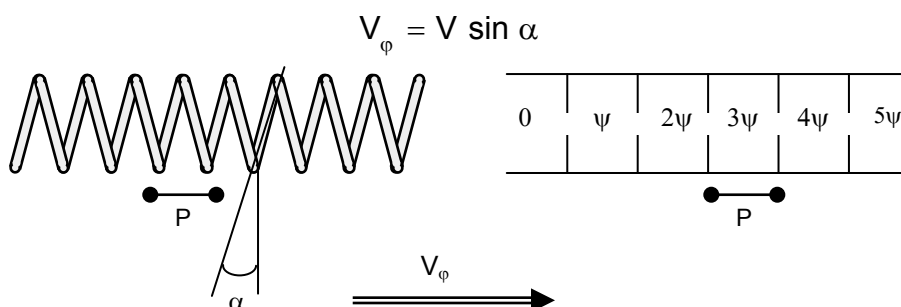
4.1 GÉNÉRALITÉS

Le fonctionnement du tube à onde progressive, en abrégiation TOP ou TPO (en anglais TWT ; Travelling - Wave - Tube), se rapproche de celui du klystron en ce sens que l'on utilise une modulation de vitesse du faisceau électronique par interaction avec une onde hyperfréquence. Cependant, cette interaction, qui dans le klystron est limitée au niveau des cavités, s'effectue sur une grande longueur dans le tube à onde progressive.

On aura donc à faire à une interaction continue d'une onde hyperfréquence et d'un faisceau électronique se propageant à la même vitesse le long de l'axe du tube.

La structure doit également être telle que le champ électrique soit sensiblement axial de manière à ce qu'il agisse sur le faisceau électronique.

La structure la plus simple est l'hélice, dans ce cas si V est la vitesse de l'onde sur l'hélice, sa vitesse axiale est :



Des *structures périodiques* sont également utilisées. Dans ce cas, si ψ est le déphasage entre deux cellules de la ligne, p le pas de la ligne, la vitesse de phase V_ϕ de l'onde est telle que, pour une onde de pulsation ω :

$$\frac{P}{V_\phi} \cdot \omega = \psi + 2n\pi \quad \text{soit : } V_\phi = \frac{\omega}{\beta} = \frac{\omega \cdot P}{\psi + 2n\pi}$$

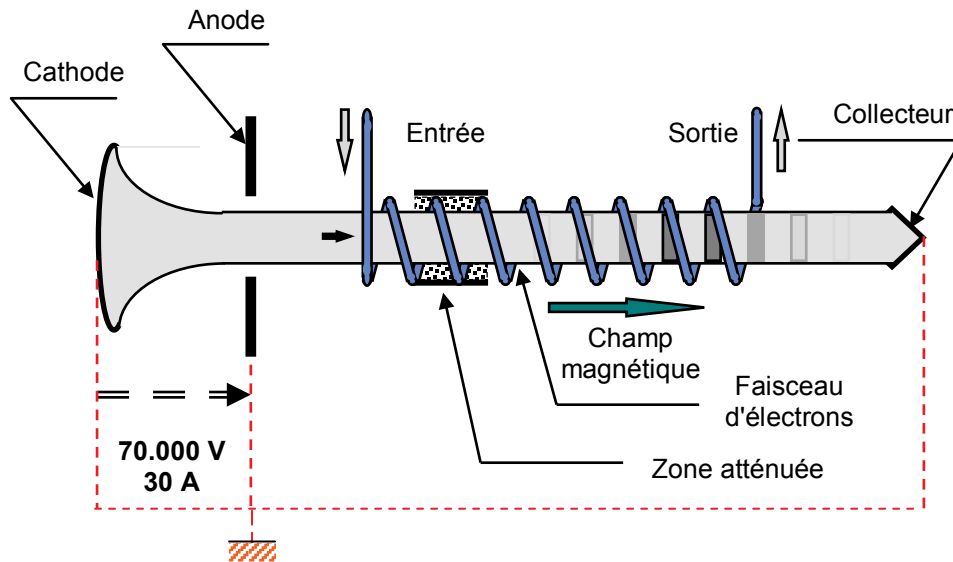
La vitesse de phase V_ϕ peut donc être choisie de telle manière qu'elle soit proche de la vitesse v_0 des électrons dans le faisceau *onde et faisceau d'électrons se déplaçant en synchronisme*. Généralement c'est le mode fondamental de propagation dans la ligne qui est utilisée ($n = 0$).

Ce type d'architecture moins « surtendue » que celle du klystron permet d'obtenir une bande passante instantanée plus élevée.

4.2 TUBE À ONDE PROGRESSIVE DE TYPE « O »

Le tube à onde progressive se compose :

- d'un canon analogue à celui du klystron,
- d'une structure destinée à guider l'onde à amplifier ; cette structure comporte une partie atténuatrice près de son entrée pour éviter que des ondes se dirigeant vers le canon puissent se propager et provoquer des phénomènes oscillatoires,
- d'un collecteur refroidi par air ou par eau,
- d'un focalisateur magnétique dans le sens du faisceau d'électron.

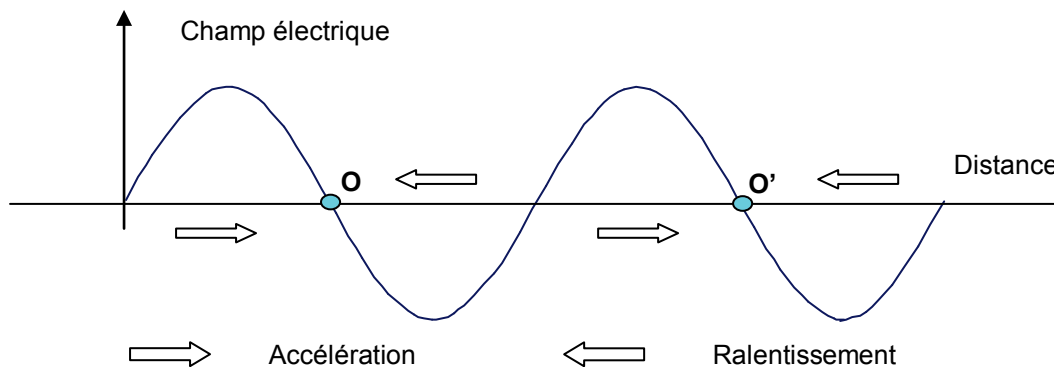


4.3 PRINCIPE DE FONCTIONNEMENT DU T.P.O.

L'onde à amplifier se propageant le long de la structure crée le long de l'axe une répartition de champ électrique sinusoïdale de pulsation ω . Le faisceau électronique qui se déplace à la même vitesse que l'onde est donc soumis à cette répartition de champ électrique, certaines zones du faisceau sont soumises à un champ électrique accélérateur, d'autres à un champ électrique retardateur.

Les électrons tels que O et O' sont soumis à un champ électrique nul, leur vitesse n'est donc pas modifiée (*électrons de référence*). Un électron situé un peu avant (*côté canon*) l'électron de référence est soumis à un champ accélérateur, sa vitesse augmente et il a tendance à rattraper l'électron de référence.

Un électron situé un peu plus loin (*côté collecteur*) que l'électron de référence est soumis à un champ électrique retardateur, sa vitesse diminue et il a tendance à être rattrapé par l'électron de référence.



Il y a donc groupement des électrons autour de l'électron de référence, et ce groupement se poursuit d'une façon continue puisque l'onde et le faisceau se déplacent sensiblement à la même vitesse.

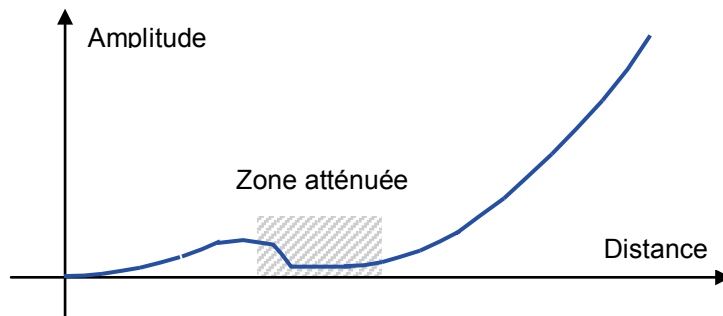
Si l'on fait en sorte que la vitesse du faisceau électronique soit légèrement supérieure à celle de l'onde, les paquets d'électrons vont avoir tendance à se déplacer par rapport à l'onde pour se trouver en permanence dans une zone où le champ électrique est retardateur. Ils seront alors freinés et céderont ainsi de l'énergie à l'onde hyperfréquence.

Cette interaction excite *deux types d'ondes* en chaque point de l'hélice, l'une (directe) se dirigeant dans le sens de l'onde excitatrice, l'autre en sens inverse. On suppose en outre que la structure ne permet que la propagation d'ondes *progressives*, c'est-à-dire d'ondes dont la vitesse de groupe et la vitesse de phase sont dans le même sens.

La synchronisation est alors possible pour les ondes (*directes*) se propageant dans le sens du faisceau dont des composantes en phase s'ajoutent, alors que les ondes (*inverses*) se propageant en sens inverse ont des phases relatives quelconques et n'interfèrent que très peu avec le faisceau d'électrons.

Ce phénomène favorise le groupement des électrons, qui cèdent alors de plus en plus d'énergie à la structure, d'où une croissance exponentielle de l'onde en fonction de la distance parcourue, comme sur le schéma ci-après :

La présence d'une zone atténuée à proximité de l'entrée de la structure ramène à une valeur très faible l'amplitude de l'onde utile et annule ainsi les ondes inverses issues de l'interaction ou d'une réflexion en bout de ligne.



Cependant, le groupement des électrons du faisceau ayant déjà été amorcé, le phénomène d'amplification décrit ci-dessus reprend après l'atténuation et se poursuit d'une manière exponentielle jusqu'à l'extrémité de la structure.

4.4 GAIN DU TUBE À ONDE PROGRESSIVE

Soit V_1 l'amplitude de l'onde à l'entrée et V_2 celle de la sortie ; si l'on ne tient pas compte de la zone atténuée, la relation liant V_2 à V_1 est de la forme :

$$V_2 = \frac{V_1}{3} \exp \left[\frac{\sqrt{3}}{2} \cdot C \cdot \beta_o \cdot L \right] \text{ expression dans laquelle :}$$

- L est la longueur de la structure ;
- β_o la constante de propagation axiale, $\beta_o = \omega / v_o$ (ω pulsation du signal à amplifier, v vitesse de propagation axiale de l'onde ou vitesse du faisceau) ;
- C est une constante, pour un tube donné, que l'on nomme *facteur de gain de Pierce* ;

Cette constante est liée à l'impédance caractéristique de la structure Z_c et à l'impédance du faisceau Z_o par la relation :

$$C^3 = \frac{Z_c}{4 Z_o}$$

Pour les tubes classiques, C est compris entre 0,01 et 0,1.

Le gain en tension du tube a donc pour expression :

$$\frac{V_2}{V_1} = \frac{1}{3} \exp \left[\frac{\sqrt{3}}{2} \cdot C \cdot \beta_o \cdot L \right]$$

En général, on pose :

$$\beta_o \cdot L = 2 \pi \cdot N$$

où N est la longueur de la structure exprimée en longueurs d'ondes λ_o le long de l'axe du tube ($\lambda_o = 2 \pi / \beta_o$), et on exprime le gain en puissance en décibels :

$$G_{(dB)} = 20 \log_{10} (V_2 / V_1)$$

d'où l'expression théorique du gain :

$$G_{(dB)} = 47,3 C \cdot N - 10$$

En fait, si l'on tient compte de la zone atténuée, des pertes dans la structure et de l'adaptation des couplages d'entrée et de sortie, la formule pratique du gain s'écarte quelque peu de la théorie tout en restant de la forme :

$$G_{(dB)} = K_1 \cdot C \cdot N - K_2$$

Pour augmenter le gain du tube à onde progressive on pourra soit :

- *Augmenter le facteur C* , ce qui conduit à utiliser une structure d'impédance caractéristique Z_c élevée et un canon donnant une impédance de faisceau Z_o faible.
- *Augmenter N* , c'est-à-dire la longueur du tube. On pourrait ainsi obtenir un gain aussi grand que souhaité, cependant, outre les problèmes techniques posés par la réalisation d'un tube très long.

Une limitation provient du fait que les électrons du faisceau constamment freinés cèdent de moins en moins d'énergie à la structure (*la vitesse des paquets d'électrons devient égale et même inférieure à celle de l'onde*) lorsque l'on se rapproche de la sortie du tube.

On peut compenser en partie cette diminution de vitesse du faisceau en adoptant une structure produisant également un ralentissement de l'onde, par exemple, en faisant varier le pas de l'hélice.

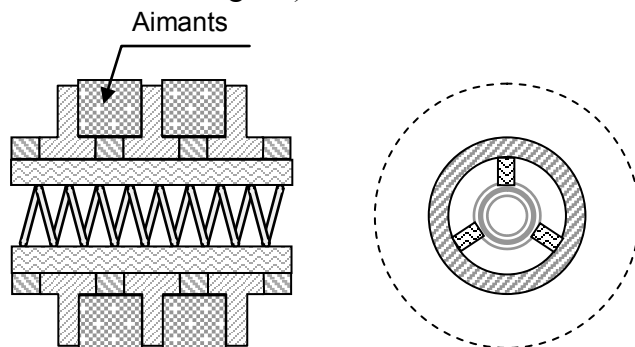
Les gains des T.P.O. couramment utilisés sont de 20 à 30 dB, mais on peut réaliser des tubes atteignant 60 dB de gain.

4.5 DIFFÉRENTS TYPES DE STRUCTURES UTILISÉS

4.5.1 L'hélice

L'hélice est la structure la plus simple et fut la première utilisée, elle est en outre peu sélective en fréquence. Pour augmenter sa rigidité, elle est fixée à l'enveloppe du tube par trois tiges à 120°.

C'est par nature, une structure permettant une très large bande d'utilisation (jusqu'à plusieurs octaves). La focalisation est fréquemment assurée par des aimants permanents à polarité alternée (PPM Periodic Permanent Magnets).



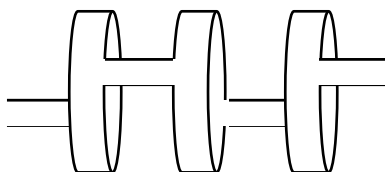
Les hélices sont, soit serrées, soit brasées. Les performances typiques de ces TOP sont les suivantes :

- hélices serrées de 500 W CW de 3 à 10 GHz, 100 W CW à 30 GHz,
- hélices brasées de 4 kW CW de 3 à 10 GHz, 200 W CW à 30 GHz,
- hélices fortes puissances crête : 30 kW crête de 3 à 5 GHz, 20 kW crête à 10 GHz, 1 kW crête à 30 GHz (facteur de forme 5 à 7 %).

Des progrès sur la technologie des hélices serrées (Brevet THOMSON-CSF) permettent d'obtenir également de fortes puissances avec des hélices serrées.

4.5.2 Structure « ring and bar » et « ring and loop »

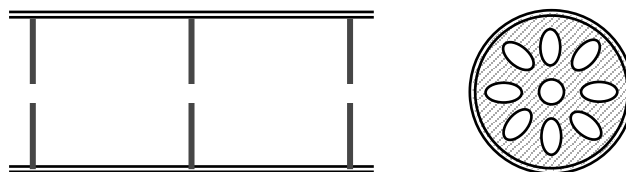
Cette structure dérivée de l'hélice possède à peu près les mêmes propriétés, toutefois sa bande passante est plus faible.



Elle se compose d'anneaux reliés par des bandes, droites (ring and bar) ou couchées (ring and loop). Son utilisation se situe à des niveaux de puissances de quelques dizaines de kilowatts crête et quelques centaines de watts moyens.

4.5.3 Guide circulaire chargé par des iris structure : clover-leaf et centipède

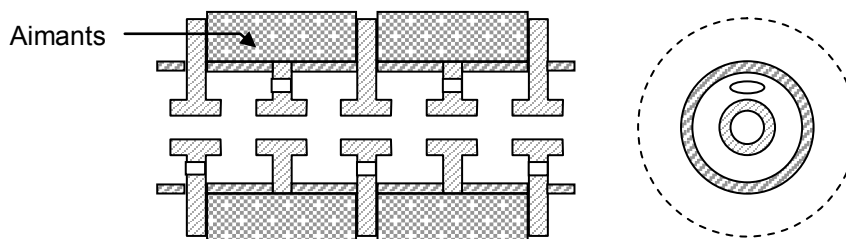
C'est un guide circulaire dans lequel on place des iris régulièrement espacés, de telle sorte que l'ensemble soit accordé sur la fréquence à amplifier.



Des structures plus ou moins complexes ont ainsi été créées, selon l'épaisseur des disques, leur découpe, leurs positions dans le guide. Les plus connues sont les structures CLOVER-LEAF (alternance de disques minces et épais) et CENTIPÈDE (disques à découpes complexes). Ces structures à bandes étroites (5 %) sont compatibles de fortes puissances : plusieurs MW crête 10 kW moyens en bande S, avec des focalisations par solénoïdes.

4.5.4 Cavités couplées à focalisation PPM

Développés pour des utilisations aéroportées ces tubes à cavités couplées par des iris ont une structure proche des précédents et permettent également dans une bande limitée, des puissances crêtes et moyennes importantes.

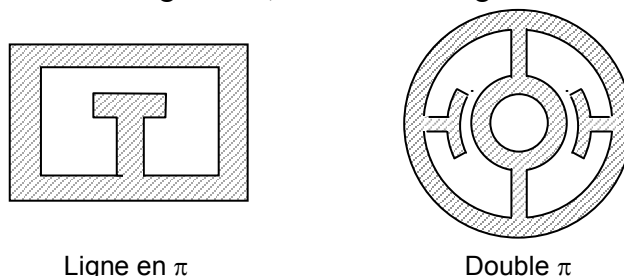


Utilisés en bande X et C leurs puissances moyennes sont de 1 à 2,5 kW avec des facteurs de forme allant de quelques % à 20 % pour des bandes de 5 à 10 %, certaines réalisations possèdent deux régimes de fonctionnement à deux puissances crêtes différentes (rapport 4 à 8).

4.5.5 Structures à barreaux, en échelles, lignes en π

En guide circulaire, les iris peuvent être dessinés, à partir de barreaux disposés radialement, de forme et de disposition diverses, la structure ainsi constituée est à large bande, tout en restant compatibles de puissances crêtes élevées.

Citons en particulier les TOP à ligne en π , dérivée de la ligne en π des tubes de type M :



Ligne en π

Double π

Des TOP de ce type, focalisés par solénoïdes permettent d'atteindre des puissances crêtes proches du MW et des puissances moyennes de plusieurs dizaines de kW en bande S dans des bandes dépassant 10 %.

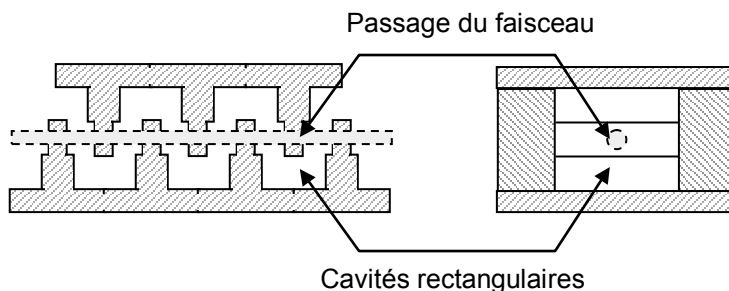
On peut utiliser également des structures telles que des guides d'onde à nervures, des lignes interdigitales et des lignes en méandres, et ceci plus fréquemment dans les tubes de types « M ».

4.5.6 TOP millimétriques

La montée en fréquence des radars (détection des câbles, suivi de terrain, autoguides air-sol...) a amené à rechercher des puissances à des fréquences élevées, tout en conservant la cohérence des signaux émis. Les TOP millimétriques répondent à cette question.

Les circuits réalisés consistent en un empilage de cavités rectangulaires reliées par des fenêtres (ou fentes) de couplage alternées. Cette technologie très simple permet de réaliser la ligne par l'usinage de deux pièces qui seront ensuite réunies par deux parois latérales.

Les puissances obtenues sont de l'ordre de 3 kW crête à 18 GHz, 500 W crête à 100 GHz avec des facteurs de forme de 2 à 10 % dans des bandes passantes de quelques %.



4.6 TUBE À ONDE PROGRESSIVE DE TYPE « M »

Les tubes examinés jusqu'ici étaient du type « O », c'est-à-dire, que le champ électrique, le champ magnétique et la vitesse du faisceau sont colinéaires ; on peut également réaliser des tubes fonctionnant sur le même principe mais, cette fois, dans le domaine d'interaction, les champs électriques, magnétiques et la vitesse du faisceau sont tous trois perpendiculaires entre eux.

Un tel tube est dit de type « M » et appelé T.P.O.M. Il fait partie de la famille des CFA (Cros Field Amplifier).

Un T.P.O.M. se compose principalement :

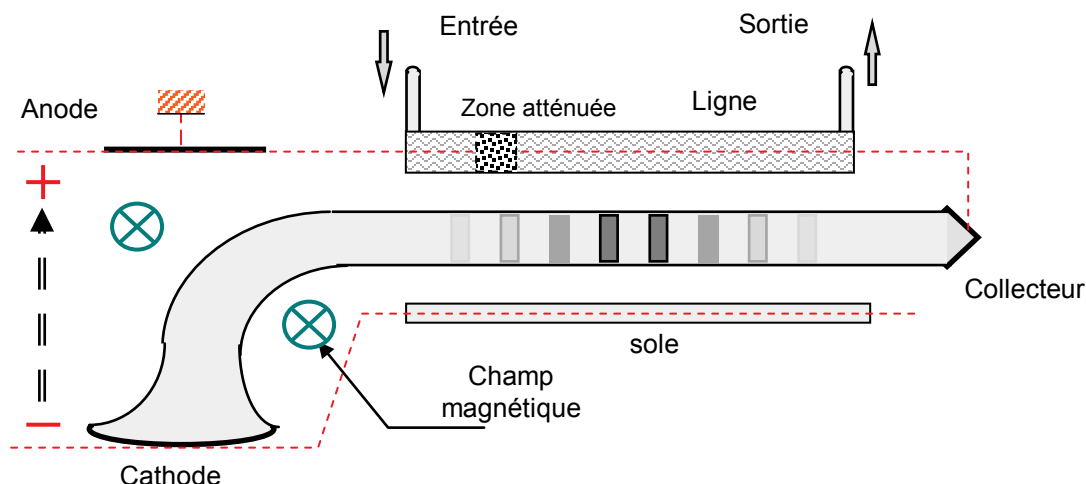
- d'un canon (cathode – anode),
- d'une structure portée à un potentiel positif par rapport à la cathode,
- d'une sole portée à un potentiel négatif par rapport à la cathode,
- d'un collecteur.

Les trajectoires des électrons issus de la cathode sont courbées sous l'effet de la force due au champ magnétique B_0 jusqu'à l'entrée dans l'espace structure-sole où la force due au champ électrique et celle due au champ magnétique s'équilibrent.

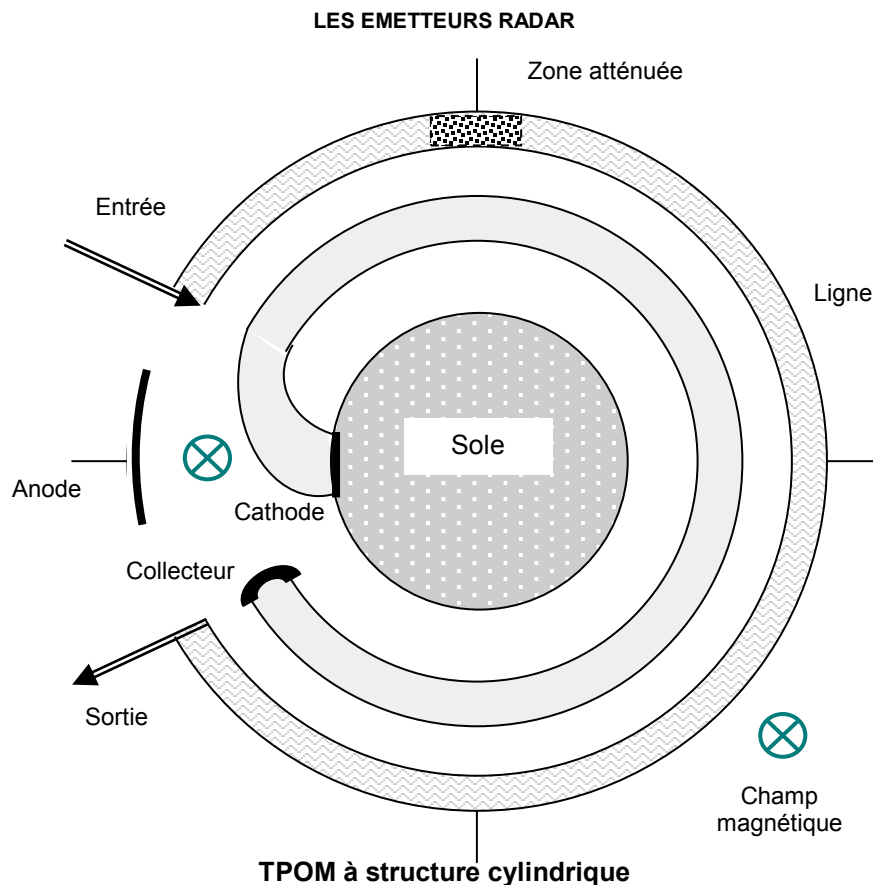
Le faisceau suit une trajectoire rectiligne au voisinage de la structure, de vitesse : $V = E_0/B_0$ (Cf. §6.3)

Les T.P.O.M. peuvent également être réalisés sous forme cylindrique, ce qui présente l'avantage de les rendre plus compacts et de faciliter l'application du champ magnétique.

Les figures page suivante présentent les deux structures de TPOM :



TPOM à structure linéaire



Dans la famille des *amplificateurs à champ croisé* (voir paragraphe 7), les T.P.O.M. font partie des tubes *à faisceaux collectés*.

Les avantages des tubes de type « M » par rapport à ceux de type « O » sont :

- leur meilleur rendement > 50 %,
- un fonctionnement nécessitant des hautes tensions plus faibles,
- des puissances de sortie plus élevées (jusqu'à 5 MW en bande L).

Leurs inconvénients sont de présenter un gain faible < 20 dB et une stabilité bien inférieure.

4.7 PERFORMANCES DES TUBES À ONDES PROGRESSIVES

Sans parler des tubes de type M qui seront repris au paragraphe 7, les TOP sont les tubes qui présentent la plus large gamme d'utilisation :

- fréquences émises de 1 GHz à 100 GHz,
- bande passante de quelques % à 30 %, certaines réalisations à très large bande atteignant même plusieurs octaves,
- gain de 20 dB (TOP miniatures) à 60 dB,
- tensions d'alimentation du kV (top à hélice de faible puissance) à 70 kV (TOP de très grandes puissances crête : typique 1 MW), voire 100 kV,
- gamme de puissances et facteurs de forme dépendant des structures utilisées.

Leurs rendements sont moyens, 30 à 40 % et exigent l'emploi de collecteurs déprimés, le rendement de ligne étant de l'ordre de 20 %.

Par ailleurs, le TOP quoique de bruit très faible (comparable à celui du Klystron) reste assez sensible aux modulations de phase par la tension d'alimentation (jusqu'à 15° par % de tension).

Les TOP, dans leur version TOP miniatures (faible gain, puissance limitée) pourraient être des éléments pour des antennes actives de très grande puissance.

5 LE CARCINOTRON

5.1 GÉNÉRALITÉS

Le carcinotron est un oscillateur dont la constitution et le principe de fonctionnement sont très voisins de ceux du tube à onde progressive. Sa caractéristique essentielle est sa très large bande de fonctionnement, sa fréquence d'oscillation étant une fonction continue de sa tension d'anode.

La différence principale réside dans le fait que le TOP amplificateur utilise la propagation d'une onde progressive dont la vitesse de phase et la vitesse de groupe sont toutes deux dirigées vers le collecteur, alors que le carcinotron utilise la propagation d'une onde régressive (c'est d'ailleurs l'origine de son nom tiré du grec « karcinos » qui veut dire écrevisse) dont la vitesse de phase est dirigée vers le collecteur et la vitesse de groupe vers le canon.

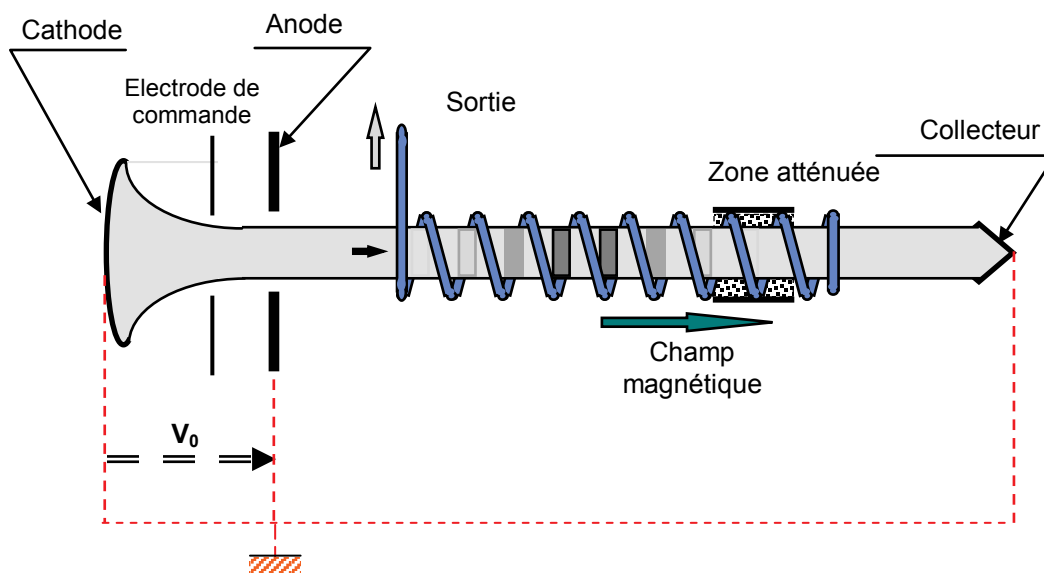
La vitesse de groupe caractérise le déplacement de l'énergie portée par une onde, ainsi dans le carcinotron, l'énergie se déplace en sens inverse du faisceau électronique et la sortie se fera donc du côté canon sur la structure.

5.2 CONSTITUTION DU CARCINOTRON

On remarque que sa constitution est très voisine du tube à onde progressive, à la différence près que la zone atténuée se trouve cette fois, du côté collecteur de façon que les ondes directes que peut propager la structure se trouvent affaiblies et que seule subsiste l'onde inverse qui est la seule utile.

De plus, le canon comporte une électrode de commande qui fait varier l'intensité du faisceau en conservant la tension d'anode constante, tension qui détermine la vitesse du faisceau et donc la fréquence d'oscillation du tube. La modulation en amplitude ou en impulsion s'effectuera par l'électrode de commande, et la modulation de fréquence ou la commande de fréquence s'effectue par l'anode.

Le carcinotron se présente schématiquement sous la forme suivante :



5.3 PRINCIPE DE FONCTIONNEMENT

Le raisonnement fait pour le tube à onde progressive reste valable, à condition de considérer que la structure ne permet de propager que des ondes dont la vitesse de groupe et la vitesse de phase sont de sens opposés.

On suppose en outre, *au départ* (comme pour tout oscillateur), l'existence sur la structure d'une onde pouvant provoquer l'oscillation ; soit dans le cas présent, d'une onde dont la vitesse de phase est voisine de celle des électrons dans le faisceau.

Comme dans le T.P.O., cette onde favorise le groupement des électrons en paquets, et si leur vitesse de déplacement est telle qu'ils soient constamment freinés, il y a cession d'énergie des électrons à la structure.

Celle-ci se manifeste sous la forme d'excitation d'une onde inverse (*vitesse de groupe vers le canon et vitesse de phase dans le sens du faisceau*) et d'une onde de sens opposé (*vitesse de groupe vers le collecteur et vitesse de phase en sens inverse du faisceau*).

On montre alors que, seule l'onde inverse possède des composantes pouvant se sommer en phase et réagir sur le faisceau en favorisant le groupement des électrons, d'où une amplification de l'onde initiale jusqu'à arrivée à un état d'équilibre.

Par contre, les ondes de sens opposé ne se somment pas en phase, et n'interfèrent que très peu avec le faisceau d'électrons. Elles sont en outre, affaiblies grâce à la présence de la zone atténuée de la structure côté collecteur.

La vitesse de phase de l'onde inverse excitée est en outre, fonction de la fréquence du signal, et doit rester voisine de la vitesse v_0 du faisceau, de manière à conserver le synchronisme nécessaire à la formation des paquets d'électrons.

Il en résulte que la fréquence d'oscillation du tube sera fonction de la vitesse v_0 du faisceau électronique, donc de la tension anode-cathode V_0 et, comme v_0 est proportionnel à $(V_0)^{1/2}$, la fréquence d'oscillation variera comme $(V_0)^{1/2}$.

On commandera la fréquence d'oscillation du carcinotron à partir de la tension d'accélération du faisceau V_0 et ceci sur une très large bande de fréquence, donc, une très grande plage de variation de V_0 .

Par ailleurs, si le courant du faisceau électronique est constant, l'énergie que possède le faisceau est d'autant plus grande que V_0 est grand et il en est de même de l'énergie cédée à la structure, donc de la puissance délivrée par le tube. De ce fait, la puissance de sortie du tube croît lorsque la fréquence augmente et inversement.

On peut corriger ce défaut à l'aide de l'électrode de commande, ce qui compensera la diminution de la vitesse du faisceau par une augmentation du courant de ce faisceau.

5.4 LE CARCINOTRON « M »

Le carcinotron « M » est au carcinotron « O », ce que le T.P.O.M. est au T.P.O., d'où, une analogie de propriétés.

Schématiquement, il se présente sous une forme analogue à celle du T.P.O.M., la zone atténuée de la structure étant côté collecteur et la sortie du tube côté canon.

Les structures les plus utilisées dans les carcinotrons « M » sont les lignes interdigitales qui se prêtent bien à la propagation d'une onde inverse.

Dans le carcinotron « M » la fréquence d'oscillation est déterminée par la différence de

potentiel V entre la structure et la sole, car la vitesse de faisceau V_o est donnée par :

$$v_o = \frac{E_o}{B_o} = \frac{V_o}{d \cdot B_o}$$

(d distance entre la structure et la sole). Nous voyons cette fois que la fréquence d'oscillation variera proportionnellement à la tension V_o , par contre, les conditions liées à une bonne formation du faisceau d'électron viennent limiter les variations de V_o .

Les avantages du carcinotron « M » résident dans :

- son rendement élevé pouvant dépasser 60 %,
- sa puissance de sortie élevée pouvant atteindre 1 MW crête et 1 kW moyen.

Par contre, il présente une bande de fonctionnement deux fois plus faible environ que le carcinotron « O » et son émission s'accompagne d'un bruit plus important.

5.5 PERFORMANCES DES CARCINOTRONS

Les carcinotrons fonctionnent depuis la bande L jusqu'à des fréquences proches de 1000 GHz. Ce sont les tubes oscillateurs qui ont permis d'obtenir les fréquences les plus élevées. Leur bande de fonctionnement peut dépasser l'octave. Les puissances de sortie sont comprises entre quelques dizaines de milliwatts et quelques centaines de milliwatts pour les carcinotrons « O », les carcinotrons « M » permettant d'atteindre des puissances beaucoup plus élevées (1 MW crête, 1 kW moyen).

Jusqu'à 30 GHz les carcinotrons O ont été remplacés par des dispositifs à état solide, mais ils gardent leur intérêt dans les domaines millimétrique et décimillimétrique.

Les carcinotrons M continuent à être utilisés comme tubes de brouilleurs.

6 LE MAGNÉTRON

6.1 GÉNÉRALITÉS

Le magnétron est un tube oscillateur à champs croisés, donc de type « M ». C'est le tube hyperfréquence le plus anciennement utilisé, il a équipé pratiquement tous les émetteurs des premiers radars décimétriques et centimétriques.

Ce type de tube est capable de fournir des puissances élevées avec un très bon rendement, mais sa bande de fonctionnement instantanée est très faible et son bruit à l'émission est relativement important.

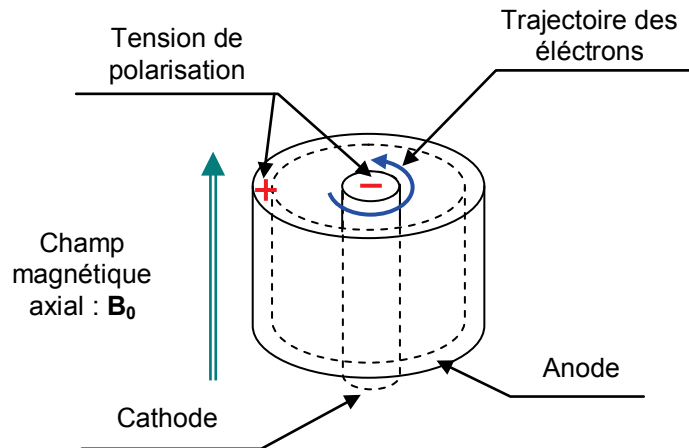
Le magnétron est donc un tube simple, bon marché, très compact et de mise en œuvre aisée. Il est particulièrement indiqué chaque fois que l'on a besoin d'un émetteur robuste et peu complexe, aussi, il a longtemps équipé un très grand nombre d'émetteurs de radars et d'autodirecteurs de missiles.

Par contre, il se prête très mal au calcul et aucune théorie complète de son fonctionnement n'existe à l'heure actuelle. On se contente d'expliquer physiquement le mécanisme de production des oscillations et tous les tubes sont réalisés à partir des résultats obtenus sur les tubes précédents.

6.2 DESCRIPTION DU MAGNÉTRON

Le magnétron se compose :

- d'une cathode de forme cylindrique,
- d'une anode en forme de couronne entourant la cathode et portée à un potentiel positif par rapport à celle-ci,
- d'un certain nombre de cavités creusées tout le tour de l'anode, régulièrement espacées et couplées à l'espace anode cathode,
- d'un dispositif de couplage dans l'une des cavités destiné à prélever l'énergie de sortie.



Le tube est fermé par deux disques assurant l'étanchéité, et baigné dans un champ magnétique axial B_0 produit le plus souvent par un aimant permanent. L'anode est refroidie par air ou par circulation d'eau.

6.3 PRINCIPE DE FONCTIONNEMENT DU MAGNÉTRON

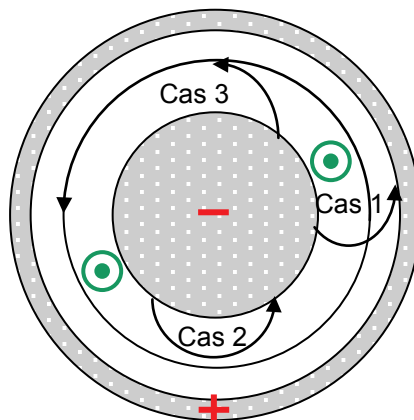
En l'absence de champ hyperfréquence, les électrons sont soumis à une force dont l'expression générale est :

$$\vec{F} = e \cdot (\vec{E}_0 + \vec{v} \wedge \vec{B}_0)$$

avec :

- e : charge de l'électron ($e = -1,6 \cdot 10^{-19}$ coulomb)
- E_0 : champ électrique statique ;
- B_0 : induction magnétique statique ;
- v : vitesse de l'électron.

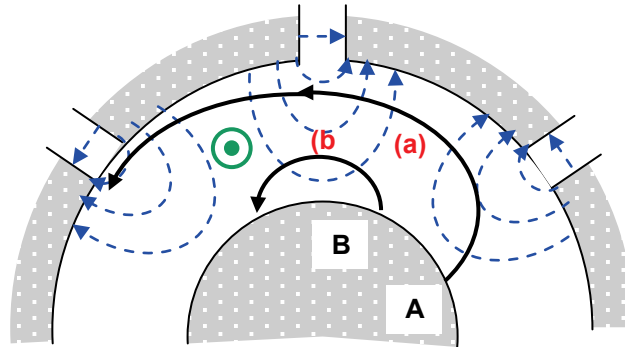
Suivant les valeurs de B_0 et de la tension d'anode V_0 , les trajectoires d'électrons peuvent avoir les allures suivantes :



- Cas 1 : les électrons atteignent l'anode.
- Cas 2 : les électrons reviennent sur la cathode.
- Cas 3 : les électrons tournent autour de la cathode.

L'anode est percée de cavités. Supposons qu'il apparaisse un champ hyperfréquence dans les cavités et que celles-ci soient réalisées de telle sorte qu'elles fonctionnent sur le « mode π », c'est-à-dire que le déphasage entre les champs dans deux cavités successives est égal à π .

Examinons l'allure des trajectoires des électrons en présence de la répartition du champ hyperfréquence issu des cavités.



V_0 et B_0 sont choisis de telle manière, qu'en absence de champ hyperfréquence, les trajectoires des électrons soient de type « 3 ». Sur la figure sont en outre présentées les lignes de champ hyperfréquence figées à un instant donné.

Un électron émis à cet instant par une zone située en A sur la cathode rencontre un champ électrique retardateur et voit sa vitesse diminuer. La force due au champ magnétique va également diminuer ainsi que la courbure de la trajectoire ; l'électron aura donc tendance à se rapprocher de l'anode. Pendant tout son parcours au voisinage de la cavité, l'électron qui est freiné cède de l'énergie à l'onde hyperfréquence existant dans la cavité.

De plus, si la tension d'anode et le champ magnétique sont choisis de telle sorte que le temps que met l'électron pour passer du voisinage de la première cavité au voisinage de la seconde est voisin de la demi-période, lorsque cet électron atteint la zone d'influence de la deuxième cavité, le champ électrique a changé de sens et il se trouve encore en présence d'un champ retardateur, sa vitesse continue à diminuer ainsi que la courbure de sa trajectoire.

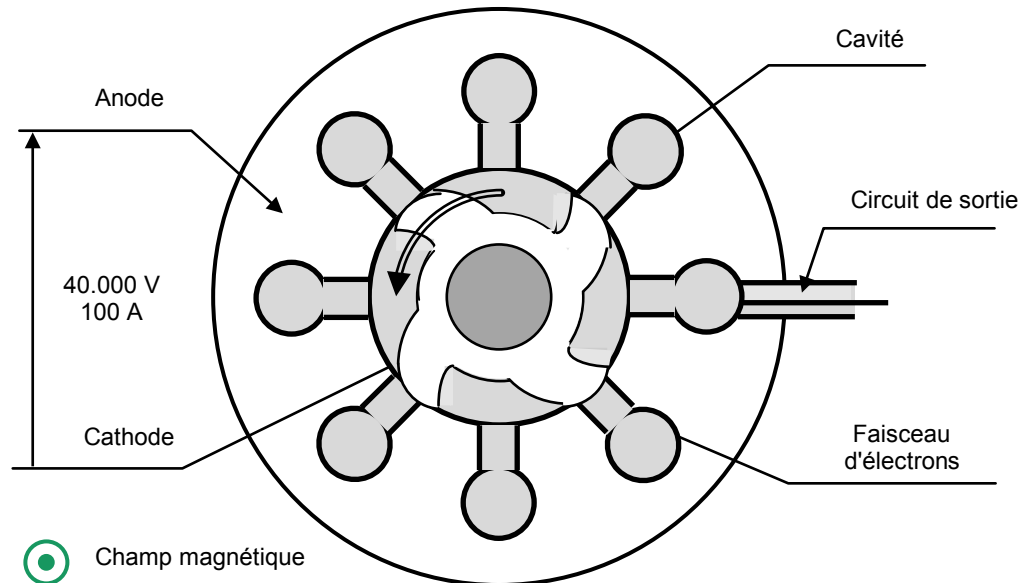
Le même phénomène peut se produire au niveau de la cavité suivante et ceci jusqu'à ce que l'électron soit capté par l'anode. Il décrit ainsi une trajectoire passant au voisinage de plusieurs cavités auxquelles il cède de l'énergie produisant l'onde hyperfréquence utile.

Considérons maintenant un électron émis toujours au même instant mais provenant d'une zone B sur la cathode, cet électron rencontre un champ accélérateur, sa vitesse augmente, la force due au champ magnétique va augmenter également ainsi que la courbure de sa trajectoire. L'électron a donc tendance à revenir très rapidement sur la cathode. Cet électron, ayant été accéléré, a prélevé de l'énergie dans la cavité correspondante.

Tous les électrons issus de la zone A pendant une demi-période où la configuration reste la même, décriront des trajectoires de type (a). Par contre, tous les électrons issus de la zone B pendant la même demi-période suivront des trajectoires du type (b) et reviendront frapper la cathode.

A la demi-période suivante, la configuration de champ s'inverse et ce sont les électrons issus de la zone B qui décrivent des trajectoires de type (a) pendant que les électrons issus de la zone A reviennent sur la cathode.

Ainsi pour un magnétron à N cavités, on a pendant une demi-période $N/2$ faisceau d'électrons ayant des trajectoires de type (a). A la demi-période suivante, la figure s'est décalée de $1/N$ tours et ainsi de suite.



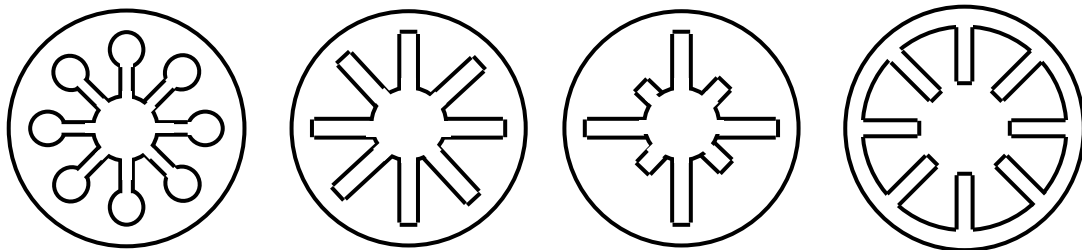
Cet ensemble de faisceaux électroniques cède une énergie importante aux cavités entretenant ainsi l'oscillation et fournissant la puissance utile.

Dans le même temps, les électrons qui retournent sur la cathode consomment de l'énergie, mais en quantité infiniment plus faible que celle fournie précédemment. En effet, ils restent dans une zone où le champ électrique hyperfréquence est faible et ce, pendant un temps très court.

Le bombardement électronique auquel est soumis la cathode devient très intense et provoque un échauffement important, de sorte que dès que le magnétron a démarré, on peut couper le chauffage de cathode.

6.4 DIFFÉRENTS TYPES DE STRUCTURES ANODIQUES

L'anode d'un magnétron est constituée par une couronne en cuivre dans laquelle sont taillées des cavités pouvant avoir des formes différentes ; quatre exemples sont représentés sur la figure ci-après :



Chaque cavité constitue un circuit oscillant accordé sur une fréquence donnée, l'ensemble de ces circuits étant mis en série par le couplage entre cavités et l'espace inter - électrodes. Si on considère que le régime hyperfréquence est équivalent à la propagation d'une onde le long de l'anode, cette propagation aboutira à un régime résonnant si le déphasage de l'onde sur un tour est égal à $2k \cdot \pi$, de manière à ce que la sommation des diverses composantes du signal hyperfréquence se fasse en phase.

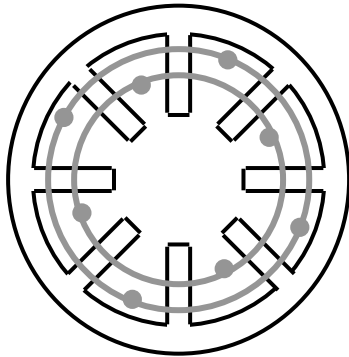
D'où la condition générale de résonance rapportée au déphasage entre *deux cavités successives* :

$$\varphi = 2k \cdot \pi / N$$

- φ : déphasage entre deux cavités voisines ;
- N : nombre de cavités.

On aura donc autant de modes d'oscillation que de cavités. Par exemple, pour un magnétron à 6 cavités, les 6 valeurs de φ correspondant aux 6 modes sont :

$$\varphi = \pm \pi/3, \pm 2\pi/3 \text{ et } \pm \pi$$



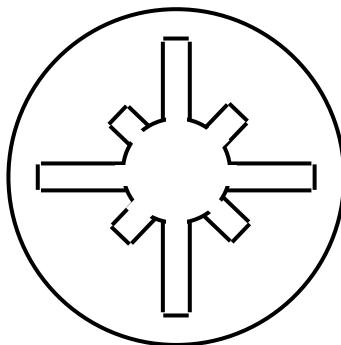
La fréquence de résonance de l'ensemble sera différente suivant le mode d'oscillation sélectionné. Aussi, pour que la fréquence d'oscillation du magnétron soit toujours la même, il faut qu'il démarre toujours sur le même mode.

Nous constatons que le mode π existe quel que soit le nombre (*pair*) de cavités. Par ailleurs, il est possible d'ajouter au magnétron un dispositif simple favorisant le démarrage sur ce mode π , ce dispositif est appelé « *strapping* ».

Le « *strapping* » consiste à relier, par un fil conducteur, de deux en deux, les parties de l'anode séparant deux cavités successives, obligeant ainsi les cavités à se retrouver en phase de deux en deux, de sorte que les parties reliées soient toujours au même potentiel. Le mode π qui est le seul respectant la condition précédente peut se propager sans opposition.

Pour les autres modes que le mode π , il apparaît une différence de potentiel, entre les parties reliées, créant un courant dans les fils. Ce courant a tendance à amortir ces modes lorsqu'ils commencent à prendre naissance.

Dans le cas où l'on ne veut pas avoir recours au « *strapping* », on peut utiliser une structure d'anode du type « *rising-sun* ». Cette structure a pour propriété de favoriser l'établissement du mode π ; elle se compose d'un entrelacement de cavités de dimensions différentes et se comporte comme deux ensembles résonnants, l'un, constitué par les grandes cavités, l'autre, par les petites.

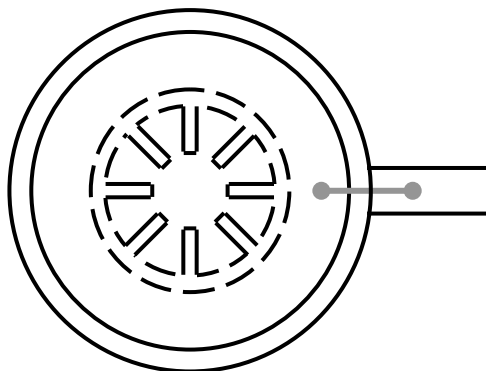


La structure « *rising-sun* » est surtout utilisée aux fréquences élevées (au-delà de la bande K_u), pour lesquelles les dimensions de l'anode deviennent si petites que la réalisation du « *strapping* » serait très délicate.

6.5 MAGNÉTRONS COAXIAUX-MAGNÉTRONS AGILES

Un magnétron fonctionnant en mode π peut être couplé à une cavité extérieure fonctionnant sur le mode TE_{011} , mode dont les lignes de champ sont des cercles, par des fentes faites à la périphérie des cavités, une cavité sur deux et donc rayonnant en phase.

L'ensemble peut être accordé mécaniquement en jouant sur la fréquence de résonance de la cavité coaxiale à l'aide d'un piston jouant sur la hauteur de la cavité. Schématiquement, la coupe d'un magnétron coaxial est la suivante :



Les magnétrons coaxiaux offrent une très bonne stabilité de fréquence due à la cavité externe, et sont compatibles d'une large gamme de durées d'impulsions. Aux fréquences élevées, bandes K_a , K_u , ils permettent d'obtenir des puissances crêtes et moyennes plus élevées.

La bande couverte par accord mécanique peut atteindre 12 %. Un tel dispositif peut être utilisé pour réaliser des magnétrons à fréquence agile, en actionnant le piston d'accord à l'aide d'un asservissement utilisant un moteur linéaire comme élément de puissance.

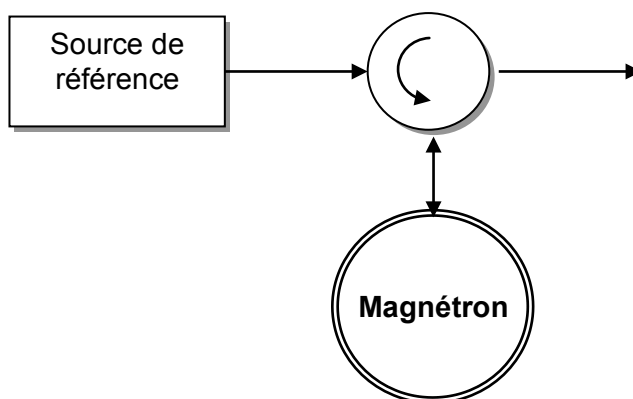
De tels magnétrons ont été réalisés, notamment dans les bandes X pour des puissances crêtes de 100 à 200 kW avec des rapidités de réponse de quelques millièmes de secondes entre deux positionnement de fréquence (précision ≈ 1 MHz).

6.6 MAGNÉTRONS SYNCHRONISÉS

Le magnétron est par nature un oscillateur incohérent dont les oscillations démarrent sur le bruit engendré par le tube. Cependant dans chaque pulse, les magnétrons peuvent posséder un haut niveau de cohérence, lié au coefficient de surtension de leur structure électronique.

De ce fait, il est possible de rendre le train d'impulsions émis par un magnétron complètement cohérent, si on parvient à contrôler la phase de démarrage des oscillations.

Cette cohérence peut être introduite en amorçant les oscillations du magnétron en lui « injectant » un signal de synchronisation couvrant le bruit naturel du tube de manière à ce que les oscillations prennent naissance avec une phase déterminée par le signal de synchronisation, selon le schéma de principe ci-après.



Le signal de référence peut alors avoir un niveau très faible, -30 à -50 dB sous le signal émis. Une cohérence plus serrée encore, peut être obtenue pour des niveaux plus élevés du signal de référence, qui peut alors non seulement piloter la phase de démarrage de l'oscillation, mais l'asservir pendant toute la durée de l'émission.

Les niveaux de synchronisation requis sont alors beaucoup plus élevés (-10 à -20 dB).

Les performances obtenues sont alors du même ordre que celles d'une boucle de phase (cf. chapitre 3 paragraphe 5.7 et 5.8), dans une bande de synchronisation définie par la relation :

$$\Delta F = \frac{2 F_o}{Q_{\text{ext}}} \sqrt{\frac{P_{\text{entrée}}}{P_{\text{sortie}}}}$$

Dans le cas d'un fonctionnement en impulsions deux effets doivent être pris en compte :

- l'effet de boucle de phase pendant la durée de l'impulsion, la compression de bruit (près de la porteuse) étant d'autant plus efficace que Q est faible, et le bruit propre du magnétron (loin de la porteuse) d'autant plus faible que Q est élevé,
- l'accrochage de cette boucle de phase pendant le front de montée (et de descente) de l'impulsion, avec un transitoire d'autant plus court que Q est faible.

En impulsions courtes, c'est le second phénomène qui domine. Le bilan global est donc favorable à l'emploi de magnétrons à Q relativement faibles. Les bruits additifs des magnétrons sont alors de l'ordre de -100 dB/Hz, les gains en puissance de l'ordre de 10 dB.

En outre, pour que cette synchronisation puisse être efficace, il faut que la dérive thermique du magnétron soit inférieure à sa bande de synchronisation, notamment au démarrage de l'émetteur (cas des autodirecteurs), la technologie et la métallurgie du tube doivent être spécialement étudiées à cet effet.

6.7 PERFORMANCES DES MAGNÉTRONS

Les magnétrons sont utilisés dans une bande de 400 à $100\,000$ MHz pour des puissances crêtes variant de 5 MW pour le bas de gamme à 1 kW pour le haut. Leur rendement de l'ordre de 50% en bas de gamme décroît également jusqu'à 15% avec la fréquence.

Les performances typiques des magnétrons sont données ci-après pour des facteurs de forme de $1/1000$ et des durées d'impulsions de 10 à $0,1$ μs .

Bande	L	S	C	X	Ka	Ku	W
F : GHz	1,3	3	5	10	16	35	94
P _c : MW	5	2,5	1,5	0,8	0,1	0,05	0,001
η : %	50	50	50	50	40	30	15
V _A : kV	50	40	35	30	15	13	12

Les magnétrons peuvent également être utilisés en radar à des facteurs de forme élevés ; citons par exemple en bande K_a des magnétrons synchronisés 300 W crête, 75 W moyens dont le facteur de forme est de 25% .

Les performances de stabilité naturelle des magnétrons sont très moyennes, elles peuvent devenir bien meilleures en mode synchronisé. Les qualités propres des magnétrons : robustesse, compacité, faible coût, faibles tensions d'alimentation, bon rendement, bonne tenue dans les environnements sévères, font qu'ils resteront encore longtemps compétitifs.

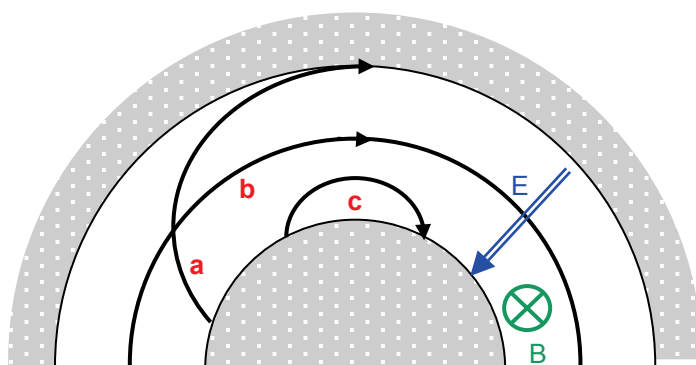
7 LES TUBES À CHAMPS CROISÉS ET À FAISCEAU RÉENTRANT

7.1 PRINCIPES GÉNÉRAUX

Dans ces tubes comme pour le magnétron les électrons se déplacent entre deux électrodes parallèles planes ou circulaires dans un champ électrique E créé par une différence de potentiel entre ces deux électrodes et un champ magnétique transversal B . Il en résulte un mouvement à accélération transversale constante défini par la relation suivante, exprimant l'équilibre entre les forces électrique, magnétique et centrifuge :

$$\frac{m v^2}{R} = e \cdot (\vec{E} + \vec{v} \wedge \vec{B})$$

La stabilité de la trajectoire est obtenue pour des vitesses faibles sept fois inférieures à la vitesse de la lumière environ. L'énergie cinétique des électrons est plus faible que dans les klystrons ou TOP mais par contre les électrons évoluent dans un milieu de potentiel variable.



En absence de champ hyperfréquence, les trajectoires d'équilibre sont des trajectoires de type **b**, décrites sur la figure ci – avant. Il en résulte qu'en régime statique, les électrons évoluent dans une couche limite parallèle aux électrodes formant un « faisceau d'électrons type M ».

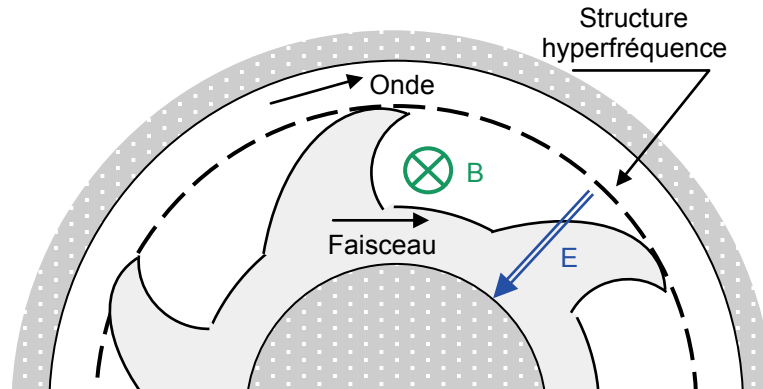
En présence d'un champ électrique complémentaire le mécanisme d'action peut s'expliquer comme suit :

- les électrons qui rencontrent un champ longitudinal retardateur voient leur vitesse diminuer. La force due au champ magnétique va également diminuer ainsi que la courbure de la trajectoire (type **a**). Les électrons vont alors se rapprocher de l'anode, dans ce mouvement leur vitesse totale, donc leur énergie cinétique varie peu, mais ils passent d'un potentiel faible au potentiel de l'anode, plus élevé, et peuvent ainsi céder au milieu de l'énergie potentielle ;
- l'effet inverse se produit pour un champ longitudinal accélérateur, les électrons type **C** voient leur courbure augmenter et rejoignent le potentiel zéro en consommant de l'énergie.

Quoique le mécanisme de fonctionnement soit assez complexe, on peut retenir que :

- les électrons ralentis (type **a**) qui cèdent leur énergie à l'onde guidée par la structure hyperfréquence rejoignent l'anode en fin de course. Le bombardement de la structure hyperfréquence est donc inhérent au principe du tube type M ;
- la composante transversale du champ électrique joue un rôle important en plaçant en bonne position la plupart des électrons ;
- ceux qui ne sont pas mis en phase bombardent la cathode en prélevant de l'énergie à l'onde.

Dans les tubes à champs croisés, le champ électrique complémentaire est fourni par la structure hyperfréquence qui fait partie de l'anode. Si la vitesse de phase de l'onde propagée par la structure hyperfréquence est en synchronisme avec la vitesse d'entraînement du faisceau, un couplage de l'onde au faisceau peut se produire comme schématisé sur la figure suivante, qui rappelle celle déjà vue pour le magnétron.

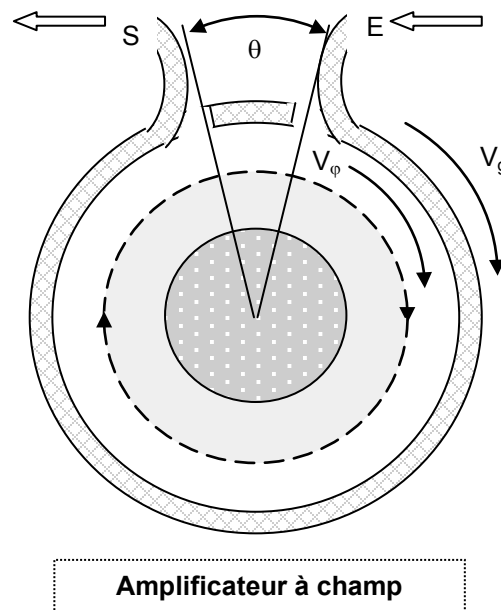


Le rendement du tube à champs croisés sera d'autant meilleur que l'énergie cinétique dissipée dans le bombardement de la ligne est faible, donc pour des faibles vitesses de faisceau.

7.2 AMPLIFICATEURS À CHAMPS CROISÉS

Ces tubes sont de forme circulaire et comportent une cathode annulaire analogue à celle des magnétrons, qui peut être chaude ou froide. Dans le second cas, c'est le champ HF qui vient amorcer l'émission des électrons en créant un bombardement électronique. Les ondes propagées dans la structure peuvent être directes (vitesse de groupe et de phase de même sens) ou inverses (vitesse de groupe et de phase de sens contraire).

Les tubes en onde directe sont dénommés ACC (amplificateurs à champs croisés ou CFA en anglais) et les tubes en onde inverse AMPLITRONS. Leur structure est légèrement différente.



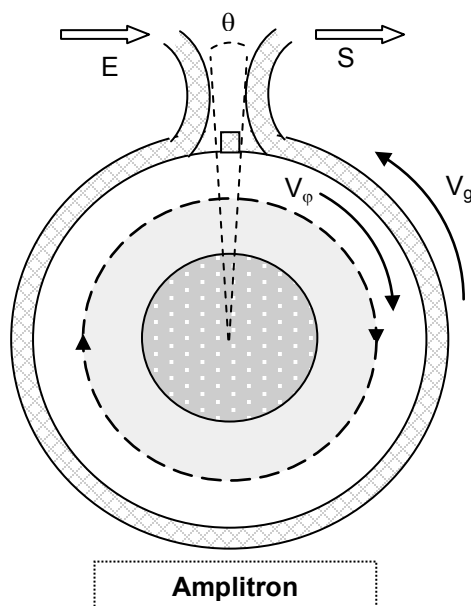
Amplificateur à champ

Dans les ACC la ligne est une ligne à propagation directe : hélice sur pieds par exemple d'une longueur électrique assez faible (14 à 18λ). L'onde interagit progressivement avec le faisceau et leurs énergies de modulation croissent ensemble jusqu'à la sortie hyperfréquence.

Le faisceau traverse ensuite un espace de glissement assez long où en absence de modulation hyperfréquence il retourne à l'état « statique » avant d'interférer à nouveau avec l'onde incidente.

Ce type de structure est adapté à des puissances de quelques dizaines, à quelques centaines de kilowatts crête en bande S et C dans une bande de 10 % avec un facteur d'utilisation de quelques % et un rendement de l'ordre de 50 %. Leur gain est de 15 dB environ.

Dans les AMPLITRONS la ligne est une ligne à propagation inverse, ligne en échelle chargée par exemple, qui permet de grandes puissances moyennes. L'espace de glissement est limité à une longueur d'onde, de telle manière que le faisceau encore formé interfère en phase avec les champs HF très puissants qu'il rencontre en pénétrant dans la structure hyperfréquence.



On peut atteindre ainsi des rendements exceptionnels dépassant 70 %, toute l'énergie HF du faisceau étant pratiquement récupérée. Par contre, leur gain (10 dB) et leur bande passante (7 %) est en moyenne plus faible que celle des AAC, à des niveaux de puissance il est vrai supérieurs. Les valeurs typiques de puissance crête envisageables sont :

BANDE	S	C	X	Ku
P_c	600 kw à 6 MW	400 kw à 4 MW	100 kw à 1 MW	50 kw à 500 kw

pour un facteur d'utilisation de quelques %.

7.3 CARACTÉRISTIQUES COMMUNES

Les tubes à champs croisés sont défavorisés par leur faible gain et leur bruit de phase supérieur d'environ 20 dB à celui des TOP. Par contre ils ont pour avantage :

- d'être peu sensibles en rotation de phase ($0,1^\circ$ par % de tension),
- d'avoir de très bons rendements,
- d'être très compacts,
- d'utiliser des THT peu élevées (10 à 40 KV suivant la puissance),
- d'être transparents quand ils ne sont pas alimentés.

Leur créneau se situe donc comme « boosters » dans des matériels relativement rustiques.
