

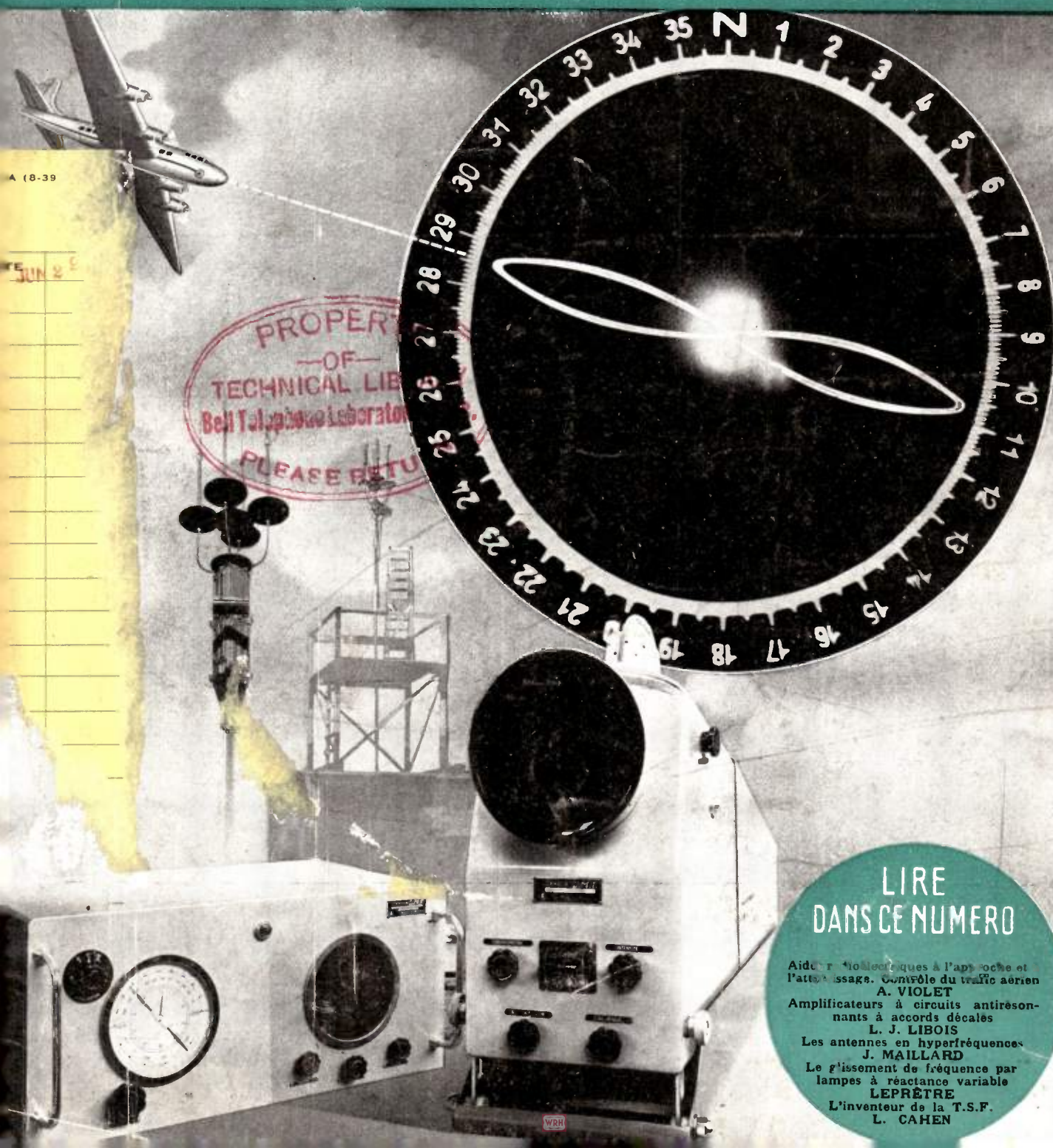
L'ONDE ÉLECTRIQUE

29^e ANNÉE N° 264

MARS 1949

PRIX : 160 FRANCS

BULLETIN DE LA SOCIÉTÉ DES RADIOÉLECTRICIENS
ÉTIENNE CHIRON, ÉDITEUR, 40, RUE DE SEINE, PARIS-6^e



LIRE DANS CE NUMERO

Aides radioélectriques à l'approche et
l'atterrissage. Contrôle du trafic aérien
A. VIOLET
Amplificateurs à circuits antiréson-
nants à accords décalés
L. J. LIBOIS
Les antennes en hyperfréquences
J. MAILLARD
Le glissement de fréquence par
lampes à réactance variable
LEPRÊTRE
L'inventeur de la T.S.F.
L. CAHEN

CABLES ISOLÉS AU *styroflex...*



Coaxiaux en paires symétriques en quartes souples, rigides pour H. F. et M. F.
Les meilleures caractéristiques de transmission Qualité et régularité.
Pour téléphonie multiple, Radiodiffusion, Radiogoniométrie, Télévision, etc.

DIRECTION SERVICES COMMERCIAUX
89, RUE DE LA FAISANDERIE - PARIS (16^e)



SERVICES D'INSTALLATION - LABORATOIRES
TÉLÉPHONE : TROCADERO 12-71 et 12-85

PURITEC DOMENACH

MAZDA



ÉMISSION
RÉCEPTION
TUBES À RAYONS
CATHODIQUES

3T1000 A
TRIODE de 1000 Watts
Un des nouveaux Tubes
de la série A émission

P.54

COMPAGNIE DES LAMPES

29, RUE DE LISBONNE - PARIS-VIII^e

CHAMPROSAV

Régisseur exclusif de la Publicité de l'Onde Electrique : Agence Publiciter-Domenach 21, Rue des Jeuneurs PARIS-2^e, Tél. CENTral 97-63

L'ONDE ÉLECTRIQUE

Revue Mensuelle publiée par la Société des Radioélectriciens
avec le concours du Centre National de la Recherche Scientifique

ABONNEMENT, D'UN AN

FRANCE. 1650 »

ETRANGER. 1950 »

ÉDITIONS

Etienne **CHIRON**

40, Rue de Seine — PARIS (6^e)

C. C. P. PARIS 53-35

Prix du Numéro :

160 francs

Vol. **XXIX**

MARS 1949

Numéro 264

SOMMAIRE

	Pages
Louis Lumière.....	89
Aides radioélectriques à l'approche et à l'atterrissage. Contrôle du trafic aérien.....	A. VIOLET 91
Les antennes en hyperfréquences.....	J. MAILLARD 110
Amplificateurs à circuits antirésonnants à accords décalés.....	L.J. LIBOIS 124
Le glissement de fréquence par lampes à réactance variable....	LEPRÊTRE 130
L'inventeur de la T.S.F.	L. CAHEN 137

Sur la couverture :

Radiogoniométrie T.H.F. à lecture oscillographique et lever de doute automatique, type R.G.O. 22 B de la Société « Le Matériel Téléphonique ».

SOCIÉTÉ DES RADIOÉLECTRICIENS

FONDATEURS

† Général **FERRIÉ**, Membre de l'Institut

† H. **ABRAHAM**, Professeur à la Sorbonne.

† A. **BLONDEL**, Membre de l'Institut.

P. **BRENOT**, Directeur à la Cie Générale de T. S. F.

J. **CORNU**, Chef de bataillon du Génie e. r.

† A. **PÉROT**, Professeur à l'Ecole Polytechnique.

† J. **PARAF**, Directeur de la Sté des Forces Motrices de la Vienne
La Société des Ingénieurs Coloniaux.

BUTS ET AVANTAGES OFFERTS

La Société des Radioélectriciens, fondée en 1921 sous le titre « Société des Amis de la T. S. F. », a pour buts (art. 1 des statuts) :

1^o De contribuer à l'avancement de la Radiotélégraphie théorique et appliquée, ainsi qu'à celui des Sciences et Industries qui s'y rattachent;

2^o D'établir et d'entretenir entre ses membres des relations suivies et des liens de solidarité.

Les avantages qu'elle offre à ses membres sont les suivants

1^o Service gratuit de la revue mensuelle *L'Onde Électrique*;

2^o Réunions mensuelles, avec conférences, discussions et expériences sur tous les sujets d'actualité technique ;

3^o Visites de diverses installations radio-électriques : stations d'émission et de réception, postes de navires et d'avions, laboratoires, radiophares expositions, studios, etc. ;

4^o Renseignements de tous ordres (joindre un timbre pour la réponse).

COTISATIONS

Elles sont ainsi fixées

1^o Membres titulaires, particuliers 1.000 fr.
sociétés ou collectivités 5.000 fr.

2^o Membres titulaires, âgés de moins de vingt-cinq ans, en cours d'études 500 fr.

Les membres de ces deux catégories, résidant à l'étranger, doivent verser en plus un supplément pour frais postaux de 300 fr.

3^o Membres à vie :

Les particuliers, membres titulaires, peuvent racheter leur cotisation annuelle par un versement unique égal à quinze fois le montant de cette cotisation, soit 15.000 fr.

4^o Membres donateurs :

Seront inscrits en qualité de donateurs, les membres qui auront fait don à la Société, en plus de leur cotisation, d'une somme égale au moins à 3.000 fr.

5^o Membres bienfaiteurs :

Auront droit au titre de « Bienfaiteurs », les membres qui auront pris l'engagement de verser pendant cinq années consécutives, pour favoriser les études ou publications techniques ou scientifiques de la Société une subvention d'au moins 10.000 fr.

Adresser la correspondance administrative et technique, et effectuer le versement des cotisations à l'adresse suivante

SOCIÉTÉ DES RADIOÉLECTRICIENS

14, avenue Pierre-Larousse, Malakoff (Seine)

Tél. : **ALÉSIA 04-16** — Compte de chèques postaux n^o 697-38

Les correspondants sont priés de rappeler chaque fois le numéro d'inscription porté sur leur carte.

CHANGEMENTS D'ADRESSE : Joindre 20 francs à toute demande.

SOCIÉTÉ DES RADIOÉLECTRICIENS

PRÉSIDENTS D'HONNEUR

R. MESNY (1935) — † H. ABRAHAM (1940).

ANCIENS PRÉSIDENTS DE LA SOCIÉTÉ

MM.

- 1922 M. DE BROGLIE, Membre de l'Institut.
 1923 H. BOUSQUET, Prés. du Cons. d'Adm. de la Cie Gle de T. S. F.
 1924 R. DE VALBREUSE, Ingénieur.
 1925 † J.-B. POMEY, Inspecteur général des P. T. T.
 1926 E. BRYLINSKI, Ingénieur.
 1927 † Ch. LALLEMAND, Membre de l'Institut.
 1928 Ch. MAURAIN, Doyen de la Faculté des Sciences de Paris
 1929 † L. LUMIÈRE, Membre de l'Institut.
 1930 Ed. BELIN, Ingénieur.
 1931 C. GUTTON, Membre de l'Institut.
 1932 P. CAILLAUX, Conseiller d'Etat.
 1933 L. BRÉGUET, Ingénieur.
 1934 Ed. PICAULT, Directeur du Service de la T. S. F.
 1935 R. MESNY, Professeur à l'Ecole Supérieure d'Electricité.
 1936 R. JOUAUST, Directeur du Laboratoire Central d'Electricité.
 1937 F. BEDEAU, Agrégé de l'Université, Docteur ès sciences.
 1938 P. FRANCK, Ingénieur en chef de l'Aéronautique.
 1939 † J. BETHENOD, Membre de l'Institut.
 1940 † H. ABRAHAM, Professeur à la Sorbonne.
 1945 L. BOUTHILLON, Ingénieur en chef des Télégraphes.
 1946 R.P. P. LEJAY, Membre de l'Institut.
 1947 R. BUREAU, Directeur du Laboratoire National de Radio-
 électricité.
 1948 Le Prince Louis de BROGLIE.

BUREAU DE LA SOCIÉTÉ

Président :

M. M. PONTE, Directeur général adjoint de la Cie Gle de T. S. F.

Vice-Présidents :

MM. P. ABADIE, Ingénieur en chef au L. N. R.
 G. LEHMANN, Ingénieur-Conseil.
 De MARE, Ingénieur.

Secrétaire général :

M. R. RIGAL, Inspecteur général adjoint des P. T. T.

Trésorier :

M. R. CABESSA, Ingénieur au L. C. T.

Secrétaires :

MM. L. J. LIBOIS, Ingénieur des P. T. T.
 M. PIRON, Ingénieur du Génie Maritime.
 J. DOCKES, Ingénieur des P. T. T.

SECTIONS D'ÉTUDES

N°	Dénomination	Présidents	Secrétaires
1	Etudes générales.	M. DE MARE.	M. FROMY.
2	Matériel radioélectr.	M. AUBERT.	M. ADAM.
3	Electro-acoustique.	M. BEDEAU.	M. POINCELOT.
4	Télévision	M. MALLEIN.	M. ANGEL.
5	Hyperfréquences.	M. GOUDET.	M. GUÉNARD.
6	Electronique.	M. LÉAUTÉ.	M. BRACHET.
7	Documentation.	M. VILLENEUVE.	M. CHARLET.

Les adhésions pour participation aux travaux des sections doivent être adressées au Secrétariat de la Société des Radioélectriciens, 10 avenue Pierre-Larousse, à Malakoff (Seine).

INFORMATIONS

SOCIÉTÉ FRANÇAISE DES INGÉNIEURS COLONIAUX

La Société Française des Ingénieurs Coloniaux (11, rue Tronchet, Paris — 8^e — Anjou 14-65) organise en mai 1949 à Paris un « Congrès International d'Ingénieurs pour le développement des pays d'outre-mer ». Le programme du Congrès comporte l'étude des problèmes d'équipement et d'industrialisation et la création d'un organisme international permanent d'étude et de liaison entre les groupes d'ingénieurs adhérents.

La Société Française des Ingénieurs Coloniaux invite tous les ingénieurs (civils et militaires) les Sociétés, les Entreprises que leur activité a mis ou met en contact ou en relations avec les Pays Neufs, à exposer, par des rapports ou communications, leur point de vue, leurs conceptions, leurs connaissances sur des sujets de leur choix entrant dans le cadre du congrès ; elle les prie de se faire connaître d'urgence en écrivant à son Siège, 11, rue Tronchet (8^e) pour préciser les sujets ou questions qu'ils désirent traiter.

D'autre part, à l'occasion de ce congrès la Société des Ingénieurs Coloniaux organise également une rétrospective des réalisations techniques françaises dans le monde et une exposition de matériel et machines destinés aux pays d'outre-mer.

DEMANDE D'EMPLOI

0.38. — Compagnie Générale de T.S.F. 23, rue du Maroc, PARIS (19^e) demande un Physicien Grande Ecole pour travaux de recherches, domaine des hyperfréquences. Ecrire à M. LE POLLES avec curriculum vitae.

SOCIÉTÉ FRANÇAISE DES ÉLECTRICIENS

La Société Française des Electriciens a décidé l'organisation chaque année d'une « Semaine de Documentation technique » destinée à mettre au courant les ingénieurs des derniers progrès de la technique.

La première semaine a eu pour objet l'électronique et ses diverses applications industrielles. Elle a comporté une série de conférences, des séances de travaux pratiques et des visites d'usines, avec remise aux auditeurs, en plus du texte des conférences, d'une importante documentation.

Deux Sessions de cette Semaine ont eu lieu en novembre 1948 et janvier 1949. Une troisième Session aura lieu, avec le même programme, du 21 au 26 mars prochain, au siège de l'Ecole Supérieure d'Electricité 10 avenue Pierre Larousse à Malakoff (Seine).

AVIS

Le service de l'Onde Electrique à chacun de ses membres constitue une lourde charge pour la Société des Radioélectriciens.

Pour que le fonctionnement de notre Société ne soit pas rendu difficile, il est indispensable que les cotisations soient versées dès le début de l'année.

Le Bureau de la Société se verra donc, à partir du 1^{er} avril, dans l'obligation de percevoir par recouvrement, les cotisations qui ne lui seraient pas parvenues à cette date, cotisations augmentées des frais ainsi entraînés.

Ci-après les nouveaux taux des cotisations :

1.000 frs pour les membres titulaires ordinaires
 (plus 300 fr pour les membres résidant à l'Etranger)

et ramenée à 500 frs pour les membres titulaires âgés de moins de 25 ans en cours d'études.

5.000 frs pour les Sociétés et Collectivités.

LOUIS LUMIÈRE (1864-1948)

Louis Lumière est né le 5 octobre 1864 à Besançon. Il était le fils d'Antoine Lumière, d'abord peintre d'enseignes, puis photographe dans cette ville. Sa famille s'étant établie à Lyon, Louis Lumière fit ses études à l'Ecole de la Martinière, la célèbre école technique qui fut fréquentée par beaucoup d'industriels de la région lyonnaise, et dont il sortit 1^{er} en 1880.

Son père ayant entrepris la fabrication des plaques sèches au gélatinobromure d'argent appelé à remplacer le collodion humide que chaque opérateur devait préparer lui-même quelques instants avant son utilisation, Louis Lumière n'avait pas vingt ans lorsqu'il mit au point dans la petite usine de Monplaisir, la fabrication de la célèbre plaque « Etiquette bleue » qui devait acquies rapidement par la suite une renommée mondiale.

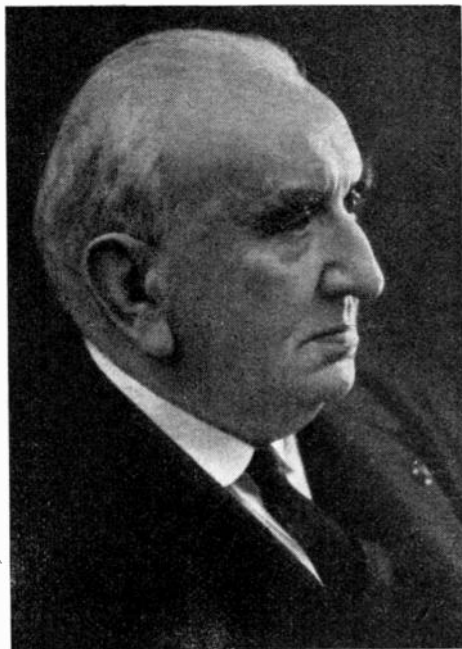
Malgré le peu de loisirs que lui laissaient les travaux qu'il poursuivait avec son frère Auguste pour développer la fabrication des plaques et papiers sensibles et créer de toutes pièces l'organisation qui devait assurer à la France, pour de longues années, une situation privilégiée dans le domaine de l'Industrie Photographique, Louis Lumière se passionnait pour le problème de la photographie animée. C'est vers la fin de l'année 1894, au cours d'une nuit d'insomnie, que fut imaginé par lui le mécanisme qui devait permettre la synthèse du mouvement. L'appareil qu'il réalisa sous le nom de « Cinématographe » était mécaniquement parfait et son principe en a été conservé malgré tous les perfectionnements apportés par la suite à cette géniale invention.

Le Cinématographe Lumière, breveté le 13 février 1895, fut présenté en séance privée, à la Société d'Encouragement à l'Industrie Nationale, le 22 Mars de la même année. Le film projeté représentait la sortie des ouvriers de l'usine de Monplaisir. Une nouvelle présentation eut lieu à Lyon, le 10 juin 1895 à l'occasion du Congrès des Sociétés Photographiques de France qui se tenait dans la salle du

Palais de la Bourse. On pouvait voir sur l'écran le célèbre astronome Janssen, Président du Congrès, s'entretenir avec Victor Lagrange, Conseiller Général du Rhône. La première séance destinée au grand public eut lieu à Paris, le 28 décembre 1895, dans le sous-sol du « Grand Café » 14, Boulevard des Capucines. Cette dernière présentation eut un succès considérable et a mérité le qualificatif de « Grande première historique » que lui ont donné ceux qui assistèrent à cette séance mémorable. Une plaque commémorative apposée sur la façade de l'immeuble construit à cet emplacement rappelle cet événement.

Mais, si le cinématographe est la plus populaire des inventions de Louis Lumière, il en est une autre qui permit à la photographie de faire un immense progrès dans le domaine de la couleur. Depuis longtemps déjà de nombreux chercheurs s'étaient attaqué au problème de l'enregistrement des couleurs. Au premier rang de ceux-ci il faut citer Gabriel Lippmann et son ingénieuse méthode interférentielle dont l'application pratique présentait de telles difficultés qu'elle est demeurée une très belle expérience de laboratoire permettant de vérifier l'exactitude des théories admises sur la nature intime des phénomènes lumineux. Venant après Lippmann et reprenant le principe de l'analyse et de la synthèse trichromes

imaginés simultanément en 1869 par deux français : Louis Ducos du Hauron et Charles Cros, Louis Lumière s'orienta vers une nouvelle solution consistant à fixer sur un vernis isolant des grains microscopiques de fécule colorée et à les recouvrir d'une mince couche d'émulsion sensible. Présenté à l'Académie des Sciences le 30 mai 1904, le procédé « Autochrome » permit dès 1907 de réaliser automatiquement, sans aucune modification préalable des appareils, la photographie en couleurs naturelles. La mise au point de ce procédé a nécessité de la part de son auteur la réalisation de véritables « tours de force » industriels, mais grâce à lui des milliers de photographes ont pu enregistrer avec une surprenante fidélité, les teintes



Les plus riches et les nuances les plus subtiles que peut offrir la nature dans toutes ses manifestations et permettre notamment la reproduction exacte des chefs-d'œuvre de la peinture.

Après avoir donné à la photographie le moyen de restituer le mouvement et la couleur, Louis Lumière entreprit l'étude du relief. La « Photostéréosynthèse » fût la solution originale qu'il présenta à l'Académie des Sciences le 8 novembre 1920. Elle permettait de reconstituer l'espace réel au moyen d'une série d'images planes. En outre, un procédé de cinéma en relief, utilisant le principe des anaglyphes fût, après avoir fait l'objet d'un Compte-Rendu à la séance du 27 février 1935 de l'Académie des Sciences, présenté au public d'une salle parisienne au mois de mai 1936.

En dehors de ces travaux qui suffiraient amplement à immortaliser le nom de leur auteur, Louis Lumière contribua largement, en collaboration avec son frère Auguste et le chimiste Seyewetz, par ses recherches et par ses nombreuses communications dans les milieux scientifiques et sociétés savantes au perfectionnement des procédés photographiques,

notamment par d'importantes études sur la chimie du développement.

L'activité de l'illustre inventeur s'exerça encore dans bien d'autres domaines : Réchauffement des nacelles d'avions par combustion catalytique de l'essence, étude des récepteurs acoustiques et de la restitution des sons, organes de prothèse pour les mutilés, etc ...

Membre de plusieurs académies étrangères et Docteur honoris causa de l'Université de Berne, Louis Lumière présida d'importantes Sociétés scientifiques et fût élu Membre de l'Académie des Sciences en 1919.

La Ville de Paris lui avait décerné la grande médaille d'or et son Jubilé scientifique avait été célébré en Sorbonne en présence du Président de la République et de nombreuses délégations étrangères, le 6 novembre 1935. Il fût enfin élevé à la dignité de Grand-Croix de la Légion d'Honneur.

Retiré pendant les dernières années de sa vie sur la côte méditerranéenne à Bandol, dont il était citoyen d'honneur, il y est décédé le 6 juin 1948.



AIDES RADIOÉLECTRIQUES A L'APPROCHE ET A L'ATTERRISSAGE. — CONTRÔLE DU TRAFIC AÉRIEN ⁽¹⁾

PAR

A. VIOLET

*Agrégé de l'Université
Ingénieur au Matériel Téléphonique*

L'objet de cet exposé est de rappeler brièvement les moyens actuellement disponibles en vue de faciliter par des procédés radioélectriques le guidage à petite distance et l'atterrissage des avions, d'examiner les données du problème à résoudre à ce point de vue et de décrire le matériel proposé, étudié et partiellement réalisé jusqu'à ce jour aux U. S. A. par les « Federal Telecommunication Laboratories ».

La première partie au moins se rapporte à des dispositifs déjà cités ou décrits devant les membres de la Société des Radioélectriciens par le Colonel PENIN et par M. FAGOT (O. E. Janvier, Février, Mars 1948). Elle est donc aussi succincte que possible et doit être considérée comme un rappel préliminaire.

La seconde partie montre le chemin qui reste à parcourir avant que l'aviation possède, en vue du radioguidage à petite distance, de l'approche des terrains et de l'atterrissage, les moyens radioélectriques satisfaisants. L'énoncé des buts à atteindre et des dispositifs qu'on peut envisager pour y parvenir s'inspirent largement des idées directrices contenues dans le rapport final de la R. T. C. A. Américaine (Radio Technical Commission for Aeronautics).

Quant à la troisième partie elle est empruntée aux renseignements et spécifications des « Federal Telecommunication Laboratories ».

Généralités.

L'approche immédiate de l'aéroport et l'opération de l'atterrissage ont toujours constitué la partie la plus difficile et la plus dangereuse des voyages aériens usuels.

L'atterrissage lui-même exige une précision qui est rarement nécessaire au cours du vol. De plus, la convergence d'un grand nombre d'appareils vers le très petit espace que constitue la piste multiplie à son voisinage les risques de collision.

Quand les conditions météorologiques sont défavorables, les méthodes de radionavigation doivent donc, non seulement permettre l'atterrissage lui-même, mais, auparavant, donner individuellement aux arrivants des directives précises en vue de les amener dans la zone de fonctionnement du dispositif d'atterrissage sans visibilité, sans qu'il en résulte aucun risque de collision.

Il est inutile de doter un terrain d'un équipement d'A. S. V. si par ailleurs les aides radioélectriques n'y sont pas suffisamment développées pour qu'il soit possible :

1° d'entre en liaison individuellement avec les appareils situés au voisinage.

2° en même temps, de repérer individuellement leurs positions, et de les suivre dans leurs évolutions.

3° de donner à chaque appareil sa position par rapport au terrain et par rapport aux autres appareils.

A ces conditions de temps de paix s'en ajoute

une autre, inspirée par la nécessité de ne pas être pris au dépourvu en temps de guerre :

Pouvoir repérer instantanément tout appareil ennemi ou tout au moins suspect, c'est-à-dire, inversement, reconnaître de façon sûre tout avion ami.

Cette quatrième condition doit être en réalité appliquée dès le temps de paix. Toute attaque brusquée, sans avertissement préalable, et effectuée peut-être par un seul appareil ennemi porteur d'un engin de destruction très puissant doit pouvoir être repérée avant que cet appareil ait atteint son objectif. Naturellement, c'est non seulement au voisinage des aéroports que le dispositif permettant de remplir la quatrième condition devra être établi, mais en toutes sortes de points du territoire.

Pour l'instant, si l'on s'en tient à l'approche d'un terrain, on s'aperçoit que la solution complète du problème se heurte surtout à des difficultés suscitées par l'accroissement du trafic ; c'est ainsi que la tour de contrôle se trouve actuellement souvent contrainte de donner aux appareils environnant le terrain l'ordre d'évoluer en cercles superposés, afin de prendre leur tour avant l'atterrissage. Cette façon d'opérer est empirique et malencontreuse. C'est une méthode de moindre mal, mais l'effet psychologique produit sur les passagers d'un avion par l'attente est à coup sûr désastreux. Il n'est pas très agréable, ni très réconfortant lorsqu'on pense au niveau d'essence, de passer une heure dans la « crasse » au-dessus de l'aérodrome.

Les finances des compagnies de navigation s'en ressentent aussi. Une étude faite par l'« Air Transport Association » des U. S. A. a montré que pour 1946 la raréfaction du trafic due au manque d'aides-radio pour le guidage sans visibilité et le temps perdu par suite de la congestion des aéroports s'étaient traduits par une perte de quarante millions de dollars sur l'ensemble des lignes aériennes américaines.

(1) Communication présentée le 20 novembre 1948, devant les membres de la Société des Radioélectriciens.

PREMIÈRE PARTIE

AIDES ACTUELLES A L'APPROCHE ET A L'ATTERRISSAGE

A l'heure actuelle, les aides à l'approche consistent dans le radioguidage des appareils à partir de leur arrivée à portée visuelle du terrain (60 à 150 km) jusqu'à leur entrée dans la zone contrôlée par le dispositif d'atterrissage sans visibilité. L'atterrissage lui-même est en effet, conformément aux recommandations de l'O. A. C. I. (1) effectué généralement à l'aide d'un équipement distinct.

On peut donc passer en revue séparément les dispositifs actuellement utilisés pour les deux fonctions, mais pour la première au moins il n'est pas toujours très facile de faire un choix entre les équipements démodés encore en service, les systèmes modernes adoptés avec certitude, et les systèmes nouveaux en cours d'installation ou d'essais pour lesquels les recommandations de l'O. A. C. I. elles-mêmes sont sujettes à révision.

I. — Aides à l'approche.

1° La liaison V. H. F.

Emetteurs-récepteurs de bord communiquant, dans la gamme des fréquences de 118 à 132 Mc/s, avec les émetteurs-récepteurs des tours de contrôle.

Les appareils pratiquement utilisés comportent généralement un assez petit nombre de canaux. Il en résulte trop souvent un embouteillage des conversations.

Il semble que la puissance d'émission de l'appareil de bord puisse être limitée à 5 ou 10 watts, celle de l'émetteur au sol étant de 50 à 100 watts.

La tendance actuelle est d'augmenter considérablement le nombre des canaux disponibles. On parle, par exemple, de projets à plus de 500 canaux.

2° Le radiogoniomètre V. H. F. au sol.

Même gamme que les appareils précédents.

Cet équipement joue un peu le même rôle que le radar d'approche, mais ne permet évidemment le relèvement que si l'avion émet. Il a alors sur le radar d'approche ordinaire (radar primaire) l'avantage d'identifier l'avion avec qui l'opérateur communique. On peut donc estimer qu'il serait très intéressant de combiner les deux dispositifs, et de projeter sur l'écran du radar d'approche le relèvement donné par le radiogoniomètre, en vue de l'identification du correspondant.

Par ailleurs, s'il est utilisé seul, il ne donne que l'azimut et non la distance. Il est donc nécessaire d'employer deux ou mieux trois radiogoniomètres pour obtenir le point de l'avion par triangulation. Leurs indications doivent alors être centralisées à grande distance, et leur manœuvre doit logiquement être télécommandée.

La plupart des radiogoniomètres V. H. F. sont automatiques, avec possibilité de passage en « manuel ». L'indicateur est habituellement un tube cathodique, et le relèvement est donné soit directement sans indétermination de sens, soit avec lever de doute additionnel, automatique ou non.

Dans le cas où le radiogoniomètre est unique, au voisinage d'un terrain, il est souvent situé dans l'axe de la piste, de sorte que les avions qui effectuent leur atterrissage le survolent à faible altitude. Le retournement rapide de l'indication du tube cathodique peut alors être très utile au sol pour indiquer ce passage.

La précision pratique de ces appareils est de l'ordre de 2 degrés en général.

3° Les radiophares au sol (Ranges).

Les radiophares au sol, associés avec des récepteurs de bord spéciaux — qui peuvent d'ailleurs comporter des parties communes importantes avec les récepteurs de bord de radiotéléphonie — sont de types très divers.

Le plus moderne est le radiophare omnidirectionnel ou « omnirange ». C'est un radiophare V. H. F. à polarisation horizontale donnant à l'avion son azimut par comparaison de phase entre une modulation à 30 pps, dont la phase est proportionnelle à l'angle d'azimut, et une modulation à 30 pps portée par une modulation intermédiaire à 10.000 pps et dont la phase coïncide avec la modulation précédente dans la direction du Nord.

Cet équipement est en cours d'installation aux U. S. A. Il doit, en principe, constituer la base de tout le système de radionavigation américain. Il est facile à interpréter et à appliquer au pilotage automatique, mais sa précision n'est pas très grande et, quoiqu'il ait été recommandé par l'O. A. C. I. et standardisé comme système international, son adoption semble se heurter à des réserves de la part de plusieurs états-membres. Certains techniciens estiment d'ailleurs que c'est un palier inutile avant l'installation progressive d'un système cohérent et perfectionné de radioguidage à petite distance tel que celui que propose la R. T. C. A. (voir deuxième partie).

4° Les radars au sol.

Ceux-ci sont soit des radars d'approche, pour repérage des avions en vol à des distances variant entre 10 et 100 km, soit des radars de terrain, pour surveillance de la piste. On cherche à les améliorer surtout en ce qui concerne la suppression des échos fixes et l'identification des avions par l'emploi de répondeurs de bord. La figure 1 donne la photographie d'un écran de P. P. I. ordinaire (Plan position indicator).

(1) Organisation de l'Aviation Civile Internationale.

Tels sont les principaux procédés de radioguidage à petite distance. Il faut signaler que des systèmes

de navigation à moyenne distance comme le Gee ou le Decca présentent une précision suffisante

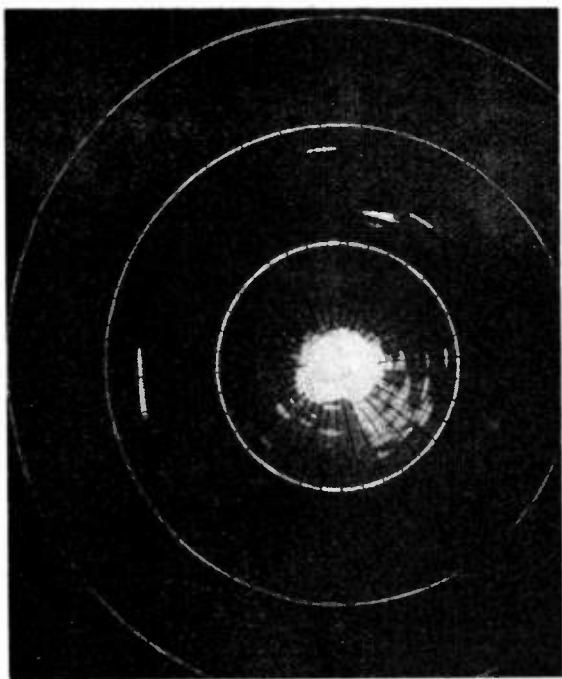


FIG. 1. — Ecran de P. P. I.

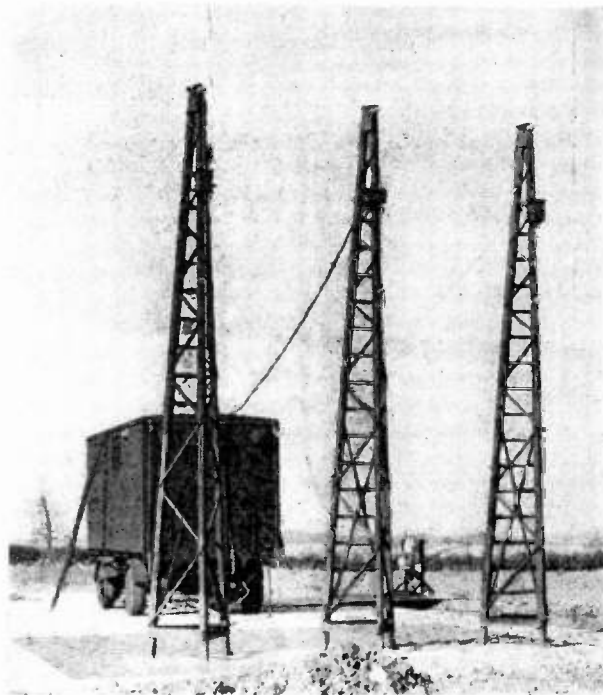


FIG. 2. — Emetteur S. B. A.

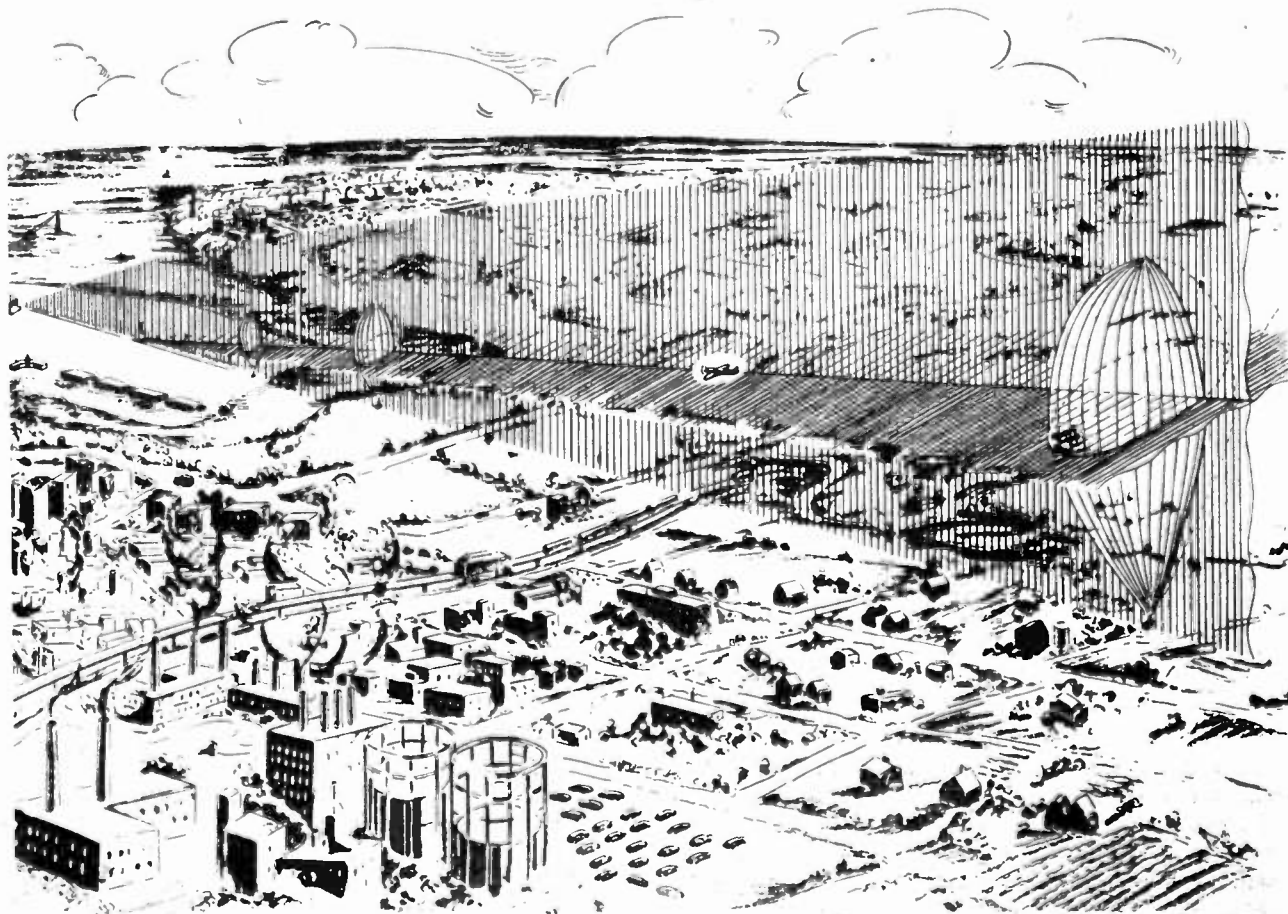


FIG. 3. — Représentation d'un atterrissage à l'aide d'un équipement I. L. S.

pour pouvoir être utilisés au radioguidage même dans la zone d'approche de terrains éloignés de centaines de kilomètres des émetteurs.

II. — Aides à l'atterrissage.

Deux sortes de méthodes sont actuellement en usage pour les atterrissages sans visibilité, celle de l'atterrissage à l'aide des instruments de bord, par référence à des axes ou des plans définis une fois pour toutes par des émetteurs fixes, et celle de l'atterrissage commandé du sol par un opérateur qui suit l'avion dans ses évolutions et le guide en « phonie » jusqu'à la piste.

1° La première méthode est appliquée dans :

a) le S. B. A. anglais, encore en usage sur les lignes britanniques (figure 2). C'est une modification du « Lorenz ». L'émetteur est situé à quelques centaines de mètres au-delà de la piste. Il travaille avec une puissance de 500 watts sur 34,8 ou 36,6 Mc/s, et est modulé électroniquement à 1150 pps. L'aérien est constitué par une antenne centrale d'émission flanquée de deux brins non excités à une distance d'environ $\frac{\lambda}{4}$.

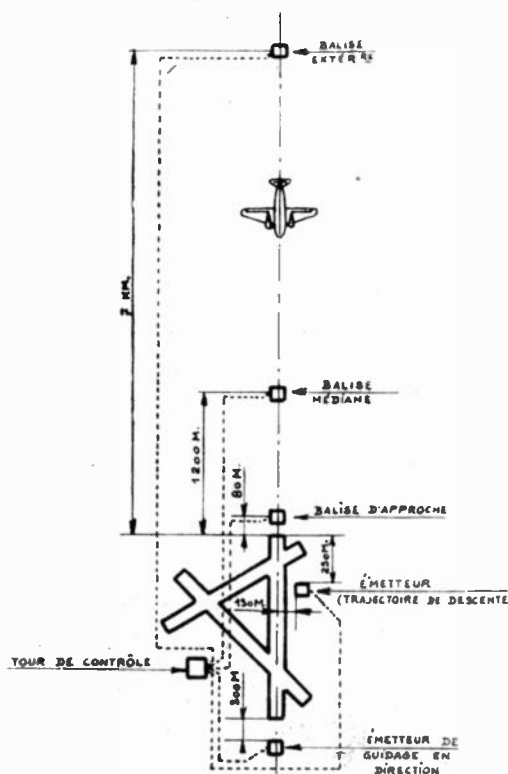


FIG. 4. — Disposition des émetteurs d'un équipement I. L. S.

Ces brins sont manipulés mécaniquement : chacun d'eux est coupé en deux parties par un interrupteur actionné par une came, ce qui modifie considérablement le diagramme d'émission de l'ensemble constitué par l'antenne centrale et chaque réflecteur. Les interruptions se produisent de façon complémentaire, le rythme de manipulation de l'un des réflecteurs correspondant à des A et celui de l'autre à des N. Lorsqu'un récepteur d'avion est

dans le plan médiateur des aériens, les A et N enchevêtrés produisent une modulation continue. De part ou d'autre, c'est l'un des signaux qui l'emporte. On définit donc ainsi un plan de guidage.

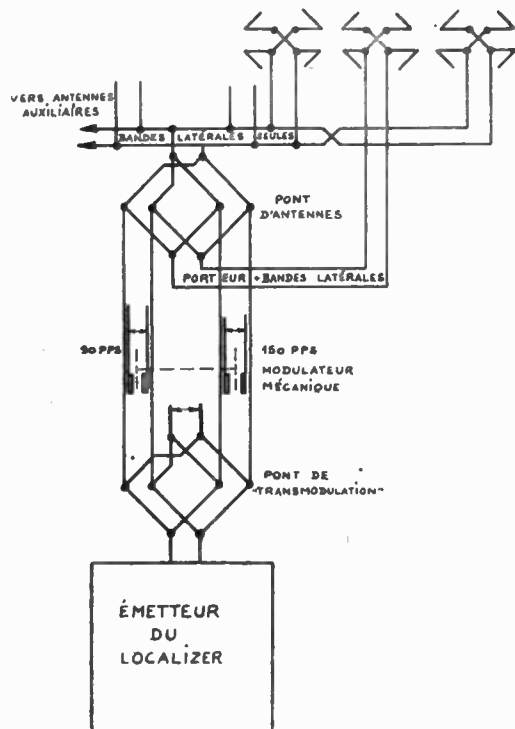


FIG. 5. — Schéma du modulateur mécanique et des ponts du Localizer (SCS51).

L'équipement ne comporte pas de trajectoire de descente, la méthode à champ constant autrefois utilisée ayant été abandonnée ; mais il est complété par deux balises à 38 Mc/s, de 2 watts de

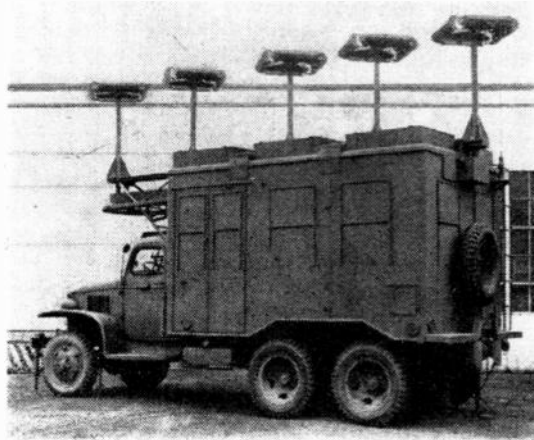


FIG. 6. — Emetteur de Localizer du S. C. S. 51.

puissance, situées l'une à 150 m de l'extrémité d'approche de la piste (modulation à 1700 pps, manipulation à 6 points par seconde) et l'autre à environ 3200 m au-delà (modulation à 700 pps manipulation à deux traits par seconde).

b) les S. C. S. 51, FTR. 51 et ILS. 2.

Ces appareils comportent un localizer à 110 Mc/s

environ, qui définit le plan d'approche, un glide-path à 330 Mc/s environ, qui définit la courbe de descente, et trois balises à 75 Mc/s qui permettent le recoupement des indications précédentes et donnent aussi la distance à l'aérodrome (figures

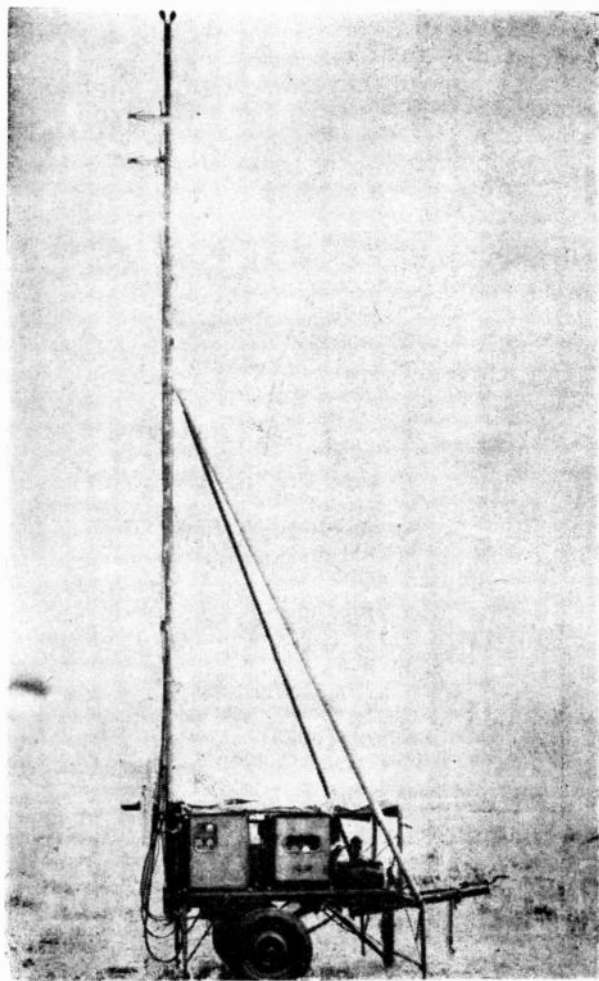


FIG. 7. — Émetteur de Glide-path du S. C. S. 51

3 et 4). Les particularités techniques les plus intéressantes de ces appareils résident dans la modulation mécanique à 90 pps et 150 pps qui a lieu à la sortie de l'émetteur, les ponts à lignes qui évitent la transmodulation et répartissent le porteur et les bandes latérales dans les aériens, et ces aériens eux-mêmes, qui sont pour le « Localizer » des antennes cadres à polarisation horizontale, et, pour le glide-path, des éléments rayonnants à polarisation horizontale avec réflecteurs (figures 5, 6, 7).

Le diagramme de directivité horizontale du localizer est donné par deux figures symétriques avec de petits lobes secondaires au voisinage du centre, l'une des figures correspondant au contour d'égale intensité de champ de l'onde modulée à 90 pps, et l'autre à celui de l'onde modulée à 150 pps (figure 8).

Pour le glide-path, la trajectoire de descente est donnée par l'intersection des lobes d'égale intensité de champ des ondes modulées à 90 pps et à 150 pps. Ces lobes sont eux-mêmes déterminés par

l'interférence des ondes directes et des ondes réfléchies sur le terrain (figure 9).

L'indicateur de bord est constitué de deux aiguilles croisées, une pour le localizer, l'autre pour le glide-path (fig. 10).

Le SCS 51 est l'équipement mobile de l'aviation militaire américaine. Son localizer comporte 5 an-

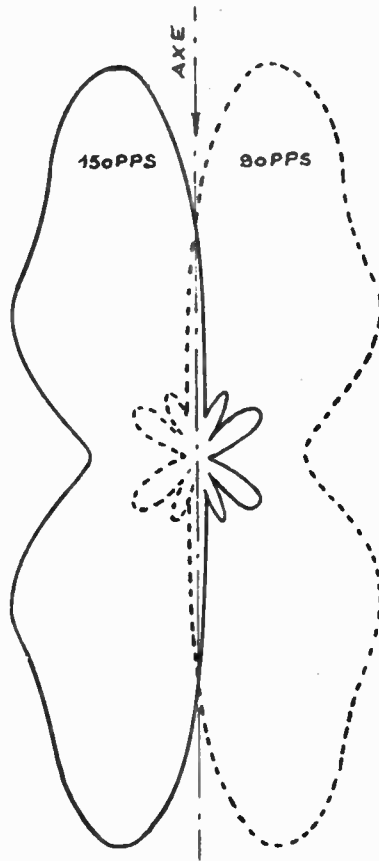


FIG. 8. — Diagramme de directivité du Localizer (S. C. S. 51)

tennes-cadres sur le camion du « Localizer » : trois antennes principales et deux antennes secondaires. Le glide-path est porté par une remorque. La liaison entre les appareils est généralement effectuée par des émetteurs-récepteurs mobiles.

Le FTR. 51 est l'équipement fixe correspondant. La directivité et la symétrie du diagramme du localizer sont améliorées par l'emploi de 7 antennes.

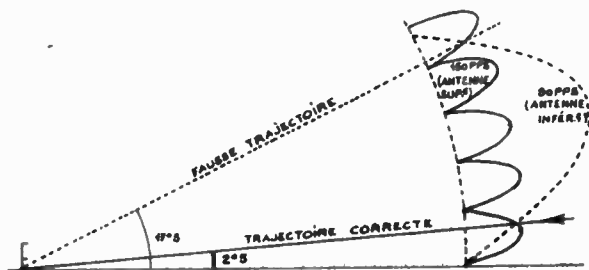


FIG. 9. — Trajectoire de descente (glide-path) du S. C. S. 51.

Il peut être télécommandé avec indication de l'alarme à distance. Quant à l'ensemble ILS. 2 (figures 11 et 12), il présente par rapport au pré-

cèdent un certain nombre d'améliorations : puissance accrue, possibilité d'émission simultanée du signal de guidage et d'une modulation vocale, à

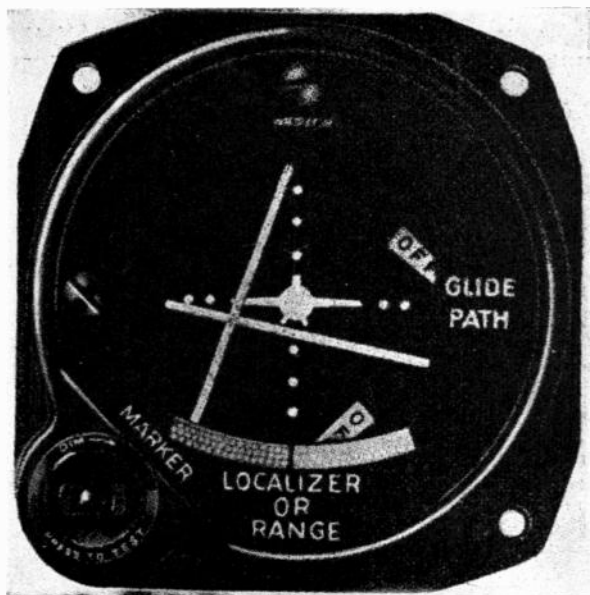


FIG. 10. — Indicateur de bord d'atterrissage sans visibilité.

l'aide d'un pont supplémentaire, transformation éventuelle pour guidage par comparaison de phase.

En principe en effet, ces équipements qui sont actuellement les équipements « standard » adoptés par l'O. A. C. I. doivent être modifiés à partir de 1951 pour fonctionner par comparaison de phase suivant le même principe que le radiophare omnidirectionnel. Toutefois la Grande-Bretagne a fait savoir récemment qu'elle n'envisageait pas cette transformation.

2° la méthode du guidage en « phonie » de l'avion dont les évolutions sont suivies à l'aide de radars est appliquée dans le G.C.A. (Ground Control Approach) qui est adopté par l'Armée et la Marine de Guerre Américaines. Dans le G. C. A. un premier équipement radar sur 3000 Mc/s explore la zone environnant le terrain et repère tous les avions

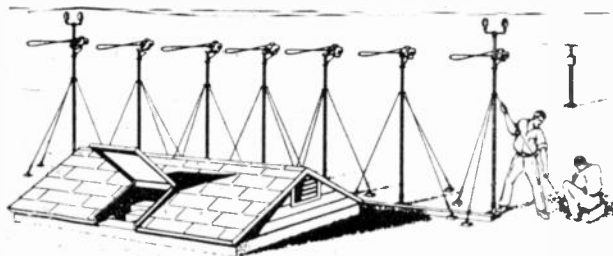


FIG. 11. — Installation du Localizer I. L. S. 2.

qui s'y présentent. Par radiotéléphonie (liaison V.H.F. usuelle), un contrôleur au sol leur donne toutes les indications nécessaires et, au fur et à mesure que l'atterrissage est possible, les amène dans le champ du deuxième équipement. Celui-ci, dans la gamme des 10.000 Mc/s, comporte deux faisceaux d'explo-

ration se balançant, l'un horizontalement, l'autre verticalement. Les positions des images de l'avion sur deux écrans cathodiques fournissent ses coordonnées au contrôleur d'atterrissage, et il donne en conséquence au pilote par radiotéléphonie, et de façon permanente, toutes les indications nécessaires à l'atterrissage.

Ce qu'on peut appeler la querelle de l'I.L.S. et du G.C.A. n'est pas close. En réalité les deux équipements se complètent, et les améliorations qu'on y apporte peu à peu font que chacun d'eux a ten-

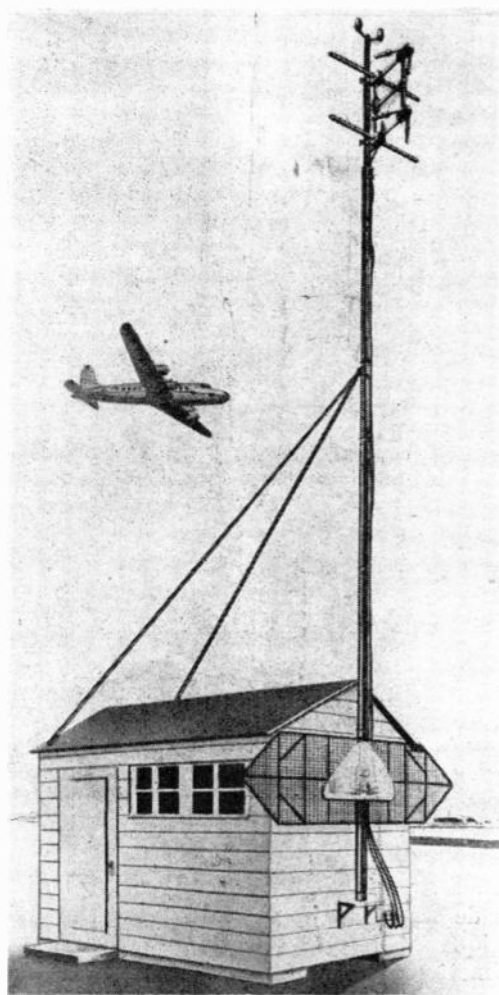


FIG. 12. — Emetteur de « Glide-path » de l'I. L. S. 2.

dance à acquérir les qualités de l'autre (communication « phonie » dans l'ILS 2, par laquelle un contrôleur au sol peut transmettre au pilote divers renseignements fournis par un radar d'approche et un radar de surveillance du terrain autonomes ; d'autre part, recherche d'un fonctionnement automatique du G.C.A., éliminant les défaillances et les erreurs du personnel).

Sans revenir sur les divers points de comparaison habituellement développés à ce sujet, il faut dire en passant que le G.C.A. a pour principal avantage de permettre l'atterrissage des petits avions équipés d'un simple émetteur-récepteur V.H.F. Or le G.C.A. aurait pu récemment être rendu complètement auto-

matique, avec transmission à l'avion des directives nécessaires par le moyen de dispositifs de signalisation conventionnels et possibilité de commande du pilote automatique. Les difficultés linguistiques reprochées au G.C.A. disparaissent donc, mais il en résulte la nécessité d'un équipement de bord complexe, et l'avantage qu'il présentait pour les petits avions disparaît aussi.

Tels sont les dispositifs radioélectriques actuels

d'aides à l'atterrissage. Il serait injuste de passer sous silence les moyens non radioélectriques tels que le balisage du terrain par des éclairages spéciaux de grande brillance, doublant le balisage nocturne habituel, ou encore le système « Fido » qui, par combustion d'huile lourde au voisinage de la piste, dans des brûleurs rapprochés, est capable de dissiper certains brouillards et d'élever artificiellement le plafond.

DEUXIÈME PARTIE.

PRINCIPES D'UN FUTUR SYSTÈME DE RADIONAVIGATION A PETITE DISTANCE

En fait le terme d'aides à l'approche doit s'entendre dans un sens très large, englobant à la fois, pour le sol, le contrôle du trafic, et, pour l'avion, les moyens d'atteindre sûrement sa destination sans risques de collision ni d'attente.

et qu'il doit avoir aussi durant son vol observé un horaire soigneusement établi. C'est donc l'ensemble du vol qui dépend en fait des conditions d'approche et d'atterrissage.

Ce problème se présente avec une particulière acuité dans certaines régions des Etats-Unis. Aussi en 1946 le Congrès Américain, saisi de la question, ordonna-t-il une enquête sur les principes des solutions à adopter. Mais il obtint durant l'hiver 1946-47 des réponses tellement confuses et contradictoires que la Commission Technique Radio de l'Aéronautique (RTCA) désigna un groupe d'études spécial pour élaborer des recommandations à l'usage des constructeurs.

(La RTCA est la seule organisation des U. S. A. où se trouvent représentés l'« Air-Force », la Marine, l'Administration de l'Aéronautique Civile, les Compagnies privées de Transport et l'Industrie Radio).

Les réunions eurent lieu en 1947 et aboutirent à un rapport final d'où l'on peut extraire un certain nombre d'idées directrices intéressantes.

1° Le rapport insiste d'abord sur le caractère incohérent des solutions actuelles :

Jusqu'à présent, chaque fois qu'une aide nouvelle a été nécessaire, on a réalisé un nouvel appareil pour l'assurer. Cette absence de méthode et de vue d'ensemble risque de conduire à des poids et des complexités prohibitifs aussi bien pour le matériel de bord que pour le matériel au sol. On estime que, si l'on continue ainsi, l'avion devra transporter quatorze appareils différents pour accomplir les opérations de navigation et de contrôle d'une façon satisfaisante.

De plus, le pilote devient alors incapable de suivre les indications de tous ces appareils (la figure 13 représente l'équipement de la cabine de pilotage d'un DC. 6).

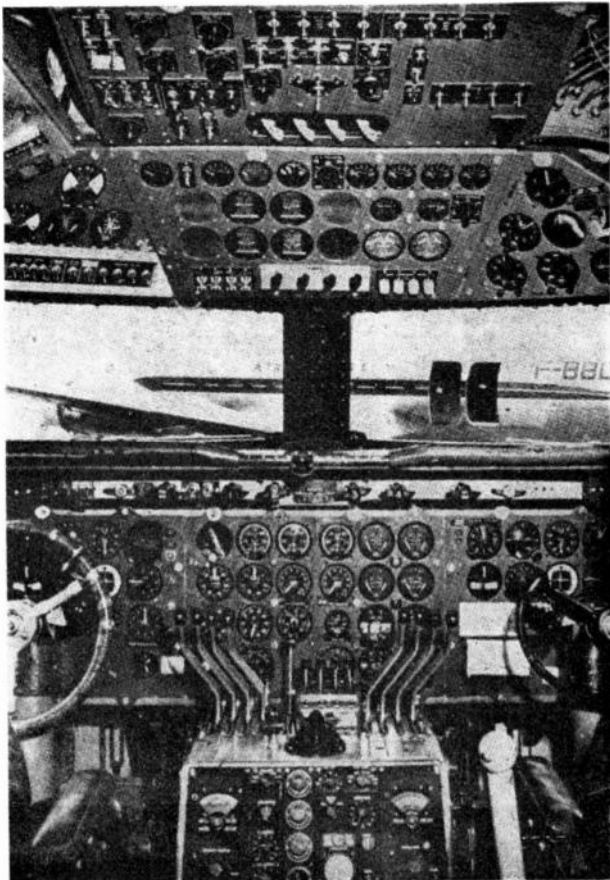


FIG. 13. — Cabine de pilotage d'un D. C. 6.

En effet, pour ne pas avoir à « stationner » au-dessus de l'aéroport d'arrivée, l'avion doit s'y présenter à une heure convenable qui lui est réservée. Cela signifie qu'il doit être parti au moment prévu,

2° Le groupe d'études de la RTCA pose ensuite un certain nombre de principes fondamentaux auxquels doit obéir le dispositif de contrôle du trafic. Certains de ces principes n'appellent aucun commentaire (sécurité de fonctionnement, possibilités d'emploi sur tous les types d'avions, minimum d'entraînement indispensable, absence de difficultés linguistiques, possibilités d'identification des appareils, etc...); mais certains points méritent une mention particulière :

Automatisme maximum : Le pilote ne doit avoir à intervenir que de façon exceptionnelle. C'est l'équipement qui travaille et qui joue tout seul un rôle de radio-navigateur.

Evolution progressive : Le système adopté permettra une évolution continue à partir des systèmes actuels vers la future solution d'ensemble.

Capacité de trafic : Ce n'est pas le système lui-même qui doit limiter la densité du trafic, mais uniquement des considérations de sécurité ou de capacité d'écoulement des aéroports.

Partage des responsabilités : Le contrôle du trafic proprement dit sera laissé à un organisme au sol, mais le pilote aura la responsabilité, tout en se conformant aux indications données par cet organisme, d'éviter les collisions.

Répartition des équipements : Les parties les plus lourdes, volumineuses et compliquées devront faire partie du matériel au sol.

Sécurité nationale : Quoique le rapport ne soit pas très explicite on peut penser qu'il s'agit des conditions suivantes :

- l'équipement ne doit pas pouvoir être utilisé par les avions ennemis pour leur propre navigation.
- l'équipement doit pouvoir distinguer les avions amis des avions ennemis.

3° Après avoir posé ces principes, la RTCA examine les dispositifs actuels, et leur accorde, au mieux, une note de 4 sur 10. Quant aux équipements proposés par diverses compagnies ou par divers organismes, il lui a été impossible de leur donner une note supérieure à 7.

Aussi la RTCA en vient-elle à conseiller le développement d'un nouveau système étudié suivant des idées directrices d'ensemble, et dont l'application constituerait un plan de 15 ans.

4° Entrant un peu plus dans les détails, la RTCA propose alors un système comportant deux équipements de bord, plusieurs équipements au sol pour chaque zone de navigation, et un équipement régional de contrôle du trafic (La figure 14 donne le bloc schématique de l'ensemble).

Les deux équipements de bord sont un équipement de navigation et un équipement de contrôle.

L'équipement de navigation est un équipement à nombreux canaux (environ 60), dans la gamme de 960 à 1215 Mc/s. Il a les fonctions suivantes :

Mesure de la distance de l'avion à la station au sol (balise répondeuse) choisie,

Mesure de son azimut par rapport à cette station.

Détermination de la position de l'avion par rapport à un localiser d'ASV. et par rapport à un glide-path,

Indication du passage au-dessus de certains points (« marker beacons »),

Représentation sur écran de la position de l'avion et des appareils voisins, transmise depuis le sol,

Communication en « phonie » sur deux voies.

Naturellement cet équipement est complété par des indicateurs qui sont des instruments de mesure, des voyants, et un écran de tube cathodique, ainsi que par l'appareillage de commande du pilote automatique.

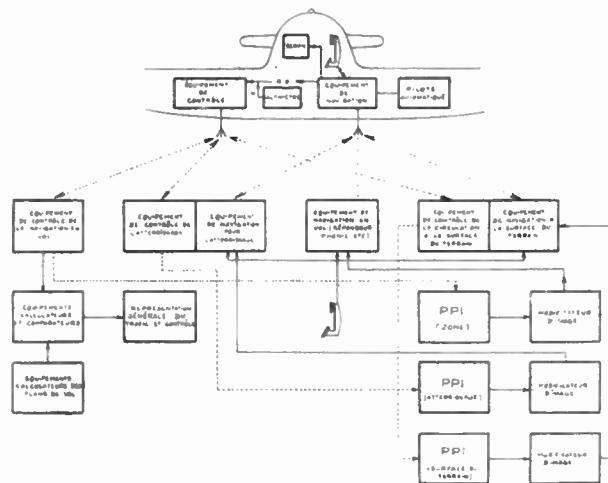


FIG. 14. — Bloc schématique du système de navigation et de contrôle du trafic recommandé par la R. T. C. A.

L'ensemble constitue donc une sorte d'unité autonome pour chaque avion, en liaison soit avec l'équipement au sol de navigation en vol proprement dite, soit avec l'équipement d'atterrissage, soit avec l'équipement de navigation à la surface du terrain, suivant le cas.

Les renseignements obtenus en vol, joints à l'altitude et à l'identité de l'avion, servent à coder l'émission du deuxième équipement, ou équipement de contrôle. Ce deuxième équipement de bord doit fonctionner entre 1365 Mc/s et 1660 Mc/s. Il constitue un répondeur pour l'équipement au sol de contrôle du trafic aérien, ou pour l'équipement au sol de contrôle de l'atterrissage, ou pour l'équipement au sol de contrôle du mouvement à la surface du terrain, suivant le cas. Un codage spécial lui permet aussi de transmettre les demandes éventuelles de modifications du programme de vol formulées par le pilote, ainsi que le signal d'accusé de réception de certains signaux de sécurité transmis par les équipements au sol (changement de route, attente, attention, etc...) ou de certains avertissements (vous êtes trop haut, trop bas, déporté, en avance, en retard, etc...).

Le contrôle du trafic d'un aéroport se fait au bureau de l'aérodrome à l'aide, en particulier, de trois indicateurs de position « PPI » relatifs respectivement à la zone d'approche, à la zone de l'aéroport (atterrissage et envol) et à la surface même du terrain. Les indications utiles sont envoyées au

pilote soit sous forme d'indications symboliques (allumage de voyants marqués : montez, descendez, plus à droite, etc...) soit sous forme de modificateurs d'images (flèches tracées électroniquement, etc., par exemple pour indiquer le trajet à suivre à la surface du terrain) qui apparaissent à bord sur l'écran représentant la zone considérée.

Enfin l'organisation générale du trafic comporte des équipements susceptibles de calculer automatiquement les horaires, de vérifier l'absence de risques de collision, etc... ainsi que de comparer les renseignements émis par l'avion avec les valeurs prévues lors de l'établissement du plan de vol et de signaler les désaccords au pilote. Ces équipements calculateurs sont centralisés dans un bureau de contrôle général du trafic, et la liaison avec les équipements émetteurs-récepteurs des diverses zones du territoire est effectuée par lignes hertziennes.

Dans ces conditions, un vol se présente de la façon suivante :

Le pilote soumet d'abord un projet de vol au bureau local de contrôle du trafic, qui le transmet par téléphone au bureau central. Celui-ci le transmet à son tour au bureau du lieu de destination pour avis, modifications éventuelles et acceptation. Le moment d'arrivée de l'appareil est réservé, retransmis au demandeur et par ailleurs inscrit au programme d'arrivée du terrain. Si l'avion arrive au terrain à l'heure dite, il pourra atterrir sans retard.

Au bureau de contrôle du départ, l'heure d'arrivée permet de calculer l'heure du départ, qui est à son tour inscrite au programme du terrain de départ.

Le programme total du vol est alors établi, et comparé dans un calculateur automatique avec tous les autres programmes déjà approuvés pour la même route. S'il n'y a pas de difficulté, il est automatiquement approuvé et enregistré. S'il est incompatible avec les autres vols, le calculateur automatique calcule un autre plan de vol.

Le pilote est finalement tenu au courant de la décision prise.

L'avion prend le départ en utilisant son équipement de navigation, d'abord sur la position « terrain », ensuite sur la position « décollage-atterrissage » et enfin sur la position « navigation en vol ». L'équipement de contrôle fonctionne par contre en permanence sans commutation.

S'il le désire, le pilote branche son équipement « navigation » sur le pilote automatique.

Pendant ce temps, au sol, un calculateur compare les positions suivantes de l'avion : Position obtenue au sol, position obtenue à bord, position prévue au programme de vol. Si les deux premières indications ne coïncident pas, le pilote en est averti, et doit prendre garde. Si des écarts importants se manifestent entre la position prévue au programme et la position réelle, un nouveau programme de vol est établi et transmis au pilote.

Le pilote peut aussi demander des informations au contrôle de trafic en appuyant sur des boutons associés au tableau indicateur de l'équipement de

contrôle. Dans certains cas la réponse est transmise automatiquement sur le canal de l'avion, parfois elle est dûe au personnel du bureau de contrôle de trafic.

Tandis que l'avion s'approche du terrain d'atterrissage, l'opérateur au sol, relatif à ce terrain est tenu au courant de sa progression par l'équipement de contrôle du trafic, qui donne sur un tableau l'identité de l'avion, son heure probable d'arrivée et sa position actuelle. Quand l'avion entre dans la zone finale d'approche, il apparaît aussi sur le PPI. L'image ainsi réalisée est modifiée par un « modificateur d'image » qui donne le plan de la région, des détails sur le temps et les routes d'approche préférables. C'est cette image modifiée qui est envoyée à l'avion par l'équipement de navigation, tandis que l'équipement de contrôle reçoit des instructions de guidage.

Le moment venu, le pilote commute son équipement de navigation sur « atterrissage » et commence sa descente. L'opération est suivie par un observateur au sol qui regarde l'image obtenue sur le PPI d'atterrissage. S'il faut envoyer à l'avion une information pendant l'atterrissage, elle peut être transmise par l'équipement de contrôle, ou par le canal « parole » de l'équipement de navigation d'ASV, ou encore elle peut être indiquée directement par le modificateur d'image d'atterrissage.

Aussitôt que l'avion atteint la piste, il est visible sur l'indicateur PPI de terrain, et l'opérateur du terrain lui donne les instructions nécessaires pour dégager la piste.

.....
Tel serait un vol effectué en utilisant les équipements prévus par la RTCA.

Ce programme, qui ne laisse rien au hasard (ni aux conditions météorologiques) peut paraître au moins ambitieux (pour ne pas dire utopique). Il est probable pourtant que c'est bien une solution de ce genre qu'il faudra adopter finalement, au moins sur les parcours aériens très fréquentés.

La circulation ferroviaire a connu une évolution analogue. Dans les déserts, on peut se contenter de « voies sans disques », mais sur les lignes à grand trafic, le bloc automatique s'est finalement imposé.

Les conclusions de la RTCA mènent pour l'aviation à une conception semblable : il ne doit plus, si on les accepte, subsister d'aides à la navigation à moyenne distance, mais seulement des aides à grande distance (Consol, Loran, Navaglobe, etc...) et des aides à courte distance telles que celles qu'on vient d'envisager. Les premières conviennent aux traversées océaniques ou désertiques, les autres aux vols continentaux ordinaires. Il devient par ailleurs impossible, dans un système cohérent de navigation à courte distance, de dissocier les problèmes d'atterrissage, d'approche et de navigation. On ne pourrait pas demander un rendement efficace au poste d'aiguillage d'une gare parisienne si les trains de banlieue n'y convergeaient pas avec un horaire précis à quelques dizaines de secondes près. Il ne viendrait à l'idée de personne qu'on puisse lancer des trains au hasard, et à aviser ensuite au milieu d'un inextricable embouteillage. C'est

pourtant ce qui se produit parfois à l'heure actuelle à l'entrée des aéroports internationaux.

Ils ont à faire face journellement à des centaines d'atterrissages, et l'on doit bien convenir que, dans ces conditions, il n'y a plus place pour l'empirisme dans la navigation aérienne.

Aussi les aides à courte distance, avec contrôle du trafic, sont elles appelées à se développer de plus en plus, et à jalonner les lignes.

De même que le train est signalé d'une gare à l'autre avec prévision du retard éventuel, l'avion sera suivi dans sa course, piloté comme sur des rails, et son heure probable d'arrivée indiquée continuellement à l'aéroport.

Il existe un autre problème, c'est celui des conséquences de l'adoption d'un plan aussi rigide sur l'essor de l'avion privé individuel. Sera-t-il possible de le munir, pour un prix acceptable, des équipements de bord nouveaux ? Pourra-t-il trouver place sur des routes aériennes devenues obligatoires ? Ou bien au contraire, de même que l'on

ne voit pas de particuliers se promener dans des autorails privés sur les lignes de chemin de fer, ne devra-t-on pas le proscrire de façon définitive pour les itinéraires et les aéroports très fréquentés ?

En ce qui concerne les aéroports, la réponse est apparemment affirmative. Il est nécessaire de prévoir des terrains spéciaux pour les avions privés. Pour la navigation elle-même, sans doute faudra-t-il un peu de souplesse pour trouver un « modus vivendi », quoique des raisons de défense nationale militent en faveur de la suppression pure et simple de ce mode de transport s'il ne se plie pas aux exigences précédemment exposées.

Quoi qu'il en soit, l'intérêt de l'enquête menée par la RTCA est d'obliger les radioélectriciens à réviser leurs idées sur le problème du radioguidage, et à le « repenser » dans son ensemble. Il ne s'agit plus pour eux de réaliser des appareils au fur et à mesure des besoins, mais au contraire de prendre en mains l'ensemble du problème de la navigation aérienne, de le résoudre et d'en faire triompher la solution.

TROISIÈME PARTIE

ÉQUIPEMENTS DES "FEDERAL TELECOMMUNICATION LABORATORIES" :

NAVAR, NAVAGLIDE ET NAVASCREEN

Les principes énoncés par la RTCA sont récents, mais, depuis plusieurs années, les laboratoires de diverses grandes firmes américaines s'étaient déjà attaqué au problème. Il n'est pas douteux d'ailleurs que c'est en tenant compte des recherches auxquelles ils se sont livrés que la RTCA a pu finalement préciser ce que les Américains attendent du système futur de radionavigation et de contrôle du trafic.

Parmi ces firmes, il faut tout particulièrement citer Hazeltine (Lanac), RCA (Teloran) et Federal (Navar et Navaglide) pour les aides à la navigation et à l'atterrissage, ainsi que la Sperry pour l'atterrissage grâce à des ondes centimétriques (voir conférences de MM. Penin et Fagot).

On trouvera ci-dessous quelques précisions sur les travaux théoriques et pratiques effectués par les Federal Telecommunication Laboratories.

I. --- Le « Navar »

Le point de départ de l'étude du Navar fut le même que celui qui présida aux travaux de la commission de la RTCA ; nécessité d'un système cohérent susceptible d'adoption progressive. D'où divers stades d'études, dont seules les plus urgentes ont été pratiquement réalisées.

Les premiers équipements à prévoir sont les radars de surveillance au sol, les balises répondeuses au sol et les indicateurs de distance et d'azimut à bord des avions.

Ces appareils ont été construits et certains ont fait l'objet de commandes du gouvernement américain.

Les études portent ensuite sur la transmission au sol d'informations supplémentaires relatives aux avions, comme l'identité et l'altitude ; l'élaboration d'un équipement automatique réunissant toutes les informations sur un seul tableau ; la transmission automatique de ces informations à distance, et la possibilité de les combiner, classer, coordonner, etc... et d'effectuer sur elles diverses opérations automatiques. On envisage enfin la transmission à l'avion d'informations supplémentaires (positions des autres avions, vent, trajectoire à suivre, etc...) sous forme d'images schématiques.

Toutes les caractéristiques du Navar ne répondent pas absolument à l'avance aux directives de la RTCA, mais elles les ont, sans doute, assez souvent inspirées et il semble intéressant de donner sur elles quelques précisions techniques.

A. --- Mesure de la distance (D. M. E.).

La position de l'appareil est déterminée en coordonnées polaires (distance et azimut). La distance est donnée par l'équipement DME (distance measuring equipment) dont il existe depuis plusieurs années des réalisations aussi bien pour l'appareil de bord que pour la balise répondeuse au sol.

Le principe du DME Federal est classique ; il consiste dans la mesure du temps qui sépare le

départ d'impulsions émises par l'avion (reçues ensuite par la balise répondeuse et réémises sur une fréquence légèrement différente) de la réception du signal de retour, compte-tenu du temps pris par la balise répondeuse pour effectuer son travail (figure 15).

Les fréquences utilisées sont comprises entre 960 et 1215 Mc/s. L'émetteur de bord envoie environ 30 impulsions par seconde avec une puissance de crête de 2 Kw, et la portée utilisable est de l'ordre de 150 km, sous réserve évidemment que l'avion et la balise se trouvent en visibilité directe.

La balise répondeuse est réalisée de façon à ré-

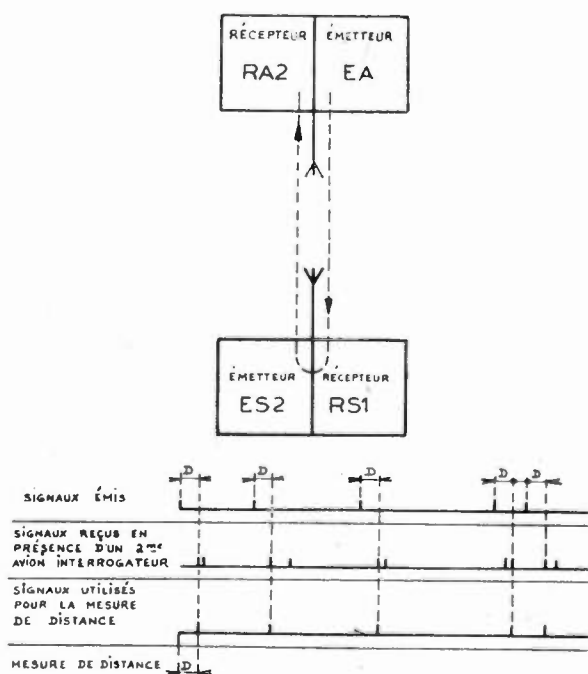


FIG. 15. — Principe de l'équipement Navar « D. M. E. »

pondre indifféremment à un grand nombre d'interrogateurs de bord, à condition qu'ils travaillent sur la fréquence caractéristique de réception de cette balise. Par conséquent le récepteur de bord doit être capable de choisir parmi toutes les impulsions qui lui parviennent celles qui correspondent à son interrogation, et elles seules. A cet effet, la vitesse de répétition des impulsions de chaque avion est caractéristique de cet appareil. Elle est voisine de 30 impulsions par seconde, mais varie erratiquement autour de cette valeur. A la réception, des circuits spéciaux constituant le « Strobe » examinent les impulsions reçues par un procédé stroboscopique, et permettent de choisir seulement celles qui se succèdent exactement au même rythme que les impulsions émises.

Les conditions primordiales de fonctionnement sont d'abord une bonne séparation des canaux, pour éviter des réponses indésirables dues à d'autres balises que celle qu'on interroge, et l'absence d'erreurs dans les réponses simultanées d'une seule balise à un grand nombre d'avions.

Par suite, en s'attaquant au problème du DME, les Federal Telecommunication Laboratories ont

posé en principe que les points les plus importants sont pratiquement :

- 1^o Précision de la mesure
- 2^o Temps nécessaire à la mesure,
- 3^o Nombre d'avions desservis par une balise,
- 4^o Nombre de doubles canaux disponibles pour les différentes balises,
- 5^o Puissance moyenne d'émission nécessaire à bord et au sol.

A propos du nombre de canaux et du nombre d'avions desservis, différents travaux de firmes américaines tendaient à adopter un petit nombre de canaux de stabilité peu poussée, avec utilisation de différents codages d'impulsions en vue de différencier les signaux lors de leur réception par les balises. Federal, par contre, voulant réserver le codage à d'autres usages s'orienta vers la technique de nombreux canaux (51 par exemple) relativement proches stabilisés par quartz. Le développement promettait d'être beaucoup plus long mais plus sûr pour l'avenir. Effectivement, les 51 doubles canaux espacés de 2,375 Mc/s prévus par Federal dans la bande de 960 à 1215 Mc/s ont reçu la consécration de la RTCA.

La figure 16 donne le bloc schématique de l'appareil de bord.

L'oscillateur stabilisé par quartz utilisé à la réception sert aussi à la stabilisation de l'émission

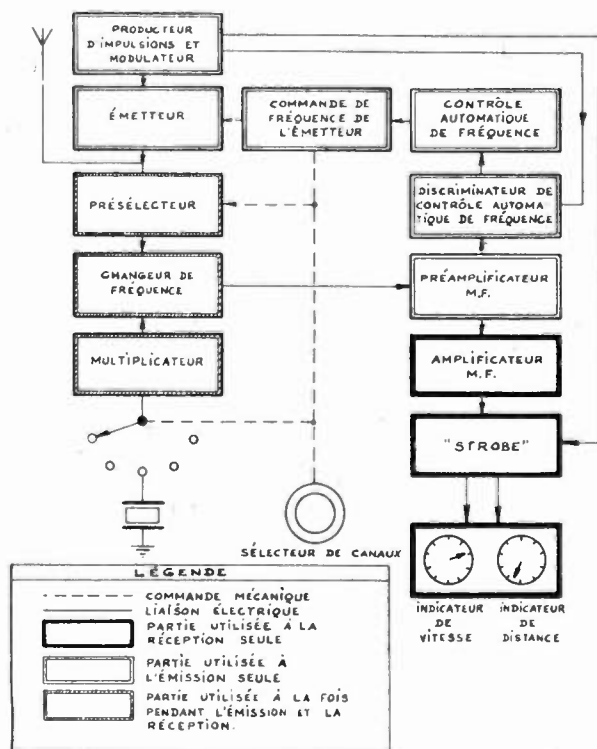


FIG. 16. — Bloc schématique de l'équipement de bord D. M. E.

à l'aide d'un circuit de correction automatique de fréquence commandant un moteur qui agit sur le réglage de la cavité résonante associée à l'oscillateur à impulsions. Lors d'un changement de canal, la commutation des quartz s'accompagne d'une modification mécanique préréglée de la cavité ré-

sonante, le réglage fin étant effectué automatiquement comme on vient de le voir.

La partie fondamentale de l'équipement est constituée par le « Strobe » qui traduit sur un appareil de mesure gradué en « miles » l'intervalle de temps entre le départ de l'impulsion et l'arrivée de la réponse. Il donne également la vitesse radiale de



FIG. 17. — Équipement de bord D. M. E.

l'appareil par rapport à la balise. Lorsque le récepteur perd la réponse de la balise, une lampe s'allume, mais l'indicateur de distance continue à fonctionner pendant quelques secondes en faisant la supposition que la vitesse radiale reste constante.

Si, alors les impulsions sont définitivement perdues, le dispositif stroboscopique de recherche se met en route et ne se fixe que sur les impulsions dont le rythme correspond à celles de l'avion et qui constituent les signaux désirés. L'appareil comporte deux gammes de fonctionnement : de 0 à 120 miles et de 0 à 12 miles. Le poids actuel de l'équipement de bord est d'environ 20 kgs (fig. 17).

On trouvera en appendice quelques détails supplémentaires sur le D. M. E.

B. — Mesure de l'azimut.

L'indicateur d'azimut utilise le même émetteur au sol et le même récepteur à bord que le D. M. E., mais ces équipements sont complétés par un deuxième récepteur à bord, et, au sol, par un émetteur que l'on dénommera ES1 avec un aérien tournant très directif. Cette deuxième émission a lieu sur une fréquence voisine de 3000 Mc/s.

Le principe de fonctionnement est le suivant (figure 18).

Un faisceau très directif d'impulsions rapprochées produites par l'émetteur au sol ES1 balaye l'horizon à une vitesse de rotation d'environ un tour par seconde. Il est reçu à bord par un récepteur spécial que l'on dénommera RA1. Par ailleurs, la balise au sol du D.M.E. (ES2) émet, chaque fois que le faisceau passe par le Nord, une impulsion puissante non directionnelle que reçoit le récepteur de bord du D.M.E. (RA2).

L'azimut est donné par le temps écoulé entre

l'arrivée du signal Nord non directionnel et l'interception du faisceau directif. Comme le récepteur RA2 effectue déjà la réception des signaux de retour de l'appareil de mesure de distance, ces impulsions sont différenciées des signaux du DME par leur longueur. Naturellement, il ne peut recevoir que les impulsions Nord provenant de la balise interrogée par son DME, les autres balises voisines étant caractérisées par des canaux différents.

Il n'en est pas de même pour l'émission des signaux directifs qui, elle, a lieu sur une fréquence fixe, commune à toutes les stations au sol, le récepteur RA1 étant accordé en permanence sur cette

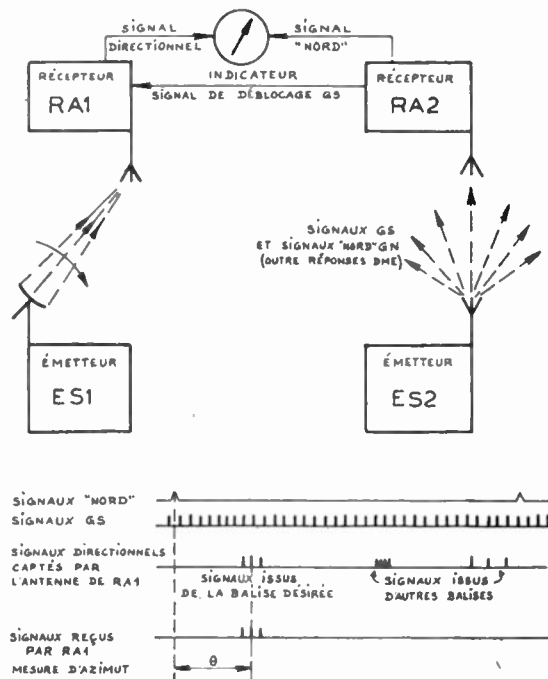


FIG. 18. — Principe de la mesure d'azimut.

même fréquence. Dès lors l'avion reçoit par son récepteur RA1 non seulement les impulsions de la balise de référence, mais les impulsions émises par les balises voisines. Il faut donc les reconnaître et les différencier.

A cet effet, l'émetteur non directionnel ES2 (répondeur du DME) envoie des impulsions non directionnelles au même rythme exactement que les impulsions directionnelles de l'émetteur ES1. Ces impulsions synchrones non directionnelles G5, amplifiées à bord par le récepteur RA2 accordé sur la balise servent à débloquent le récepteur RA1, si bien que c'est seulement lorsqu'elles arrivent de l'émetteur non directionnel que les impulsions directionnelles correspondantes peuvent être reçues par le récepteur RA1.

Naturellement les impulsions synchrones G5 diffèrent par leur longueur et leur rythme des impulsions non directionnelles « Nord », ce qui permet de les séparer dans le récepteur RA2. Elles diffèrent aussi évidemment des impulsions du DME. La traduction du temps écoulé entre l'arrivée des impulsions « Nord », et celle des impulsions directionnelles d'azimut s'effectue à l'aide d'une aiguille sur un cadran gradué en degrés (figure 19).

Cet indicateur d'azimut peut, comme l'indicateur de DME, être reproduit en divers points de l'avion. Il est susceptible d'actionner le pilote automatique en vue du parcours de trajectoires rectilignes radiales. L'emploi d'un calculateur électronique doit permettre aussi de suivre des trajectoires rectilignes quelconques.



FIG. 19. — Indicateur de bord combiné (Azimut et distance).

Les équipements de « Navar » ci-dessus décrits ont été développés et expérimentés.

Le deuxième stade de l'étude comprend des appareils dont les caractéristiques de base sont indiquées ci-dessous et qui sont actuellement en cours d'étude.

C. — « Navamander »

La fonction de l'appareillage « Navamander » est d'envoyer à l'avion des instructions de vol conventionnelles, telles que « Plus haut », « Plus à droite », « Moins vite », etc., qui sont automatiquement transcrites sur le tableau de bord à l'aide, par exemple, de voyants lumineux. Cette méthode évite les difficultés de langues et la dispersion de l'attention du pilote qui, en une fraction de seconde, reçoit et interprète l'instruction.

De plus, l'avion peut répondre à la station au sol, soit que le pilote fasse savoir, en appuyant sur un bouton, que le message a été reçu et compris, soit que l'équipement renvoie automatiquement la réponse à une question posée. Dans ce dernier cas, non seulement la réponse est automatique, mais il n'est même pas nécessaire que le pilote sache que cette liaison a lieu.

L'équipement Navamander (figure 20) comporte au sol un nouvel émetteur ES3 accordé sur la même fréquence d'émission que le répondeur du DME (ES2) et attaquant un aérien très directif. Par des dispositifs de télécommande, cette antenne suit le

curseur d'azimut que l'observateur au sol déplace sur l'écran du « Navaspector » qu'on décrira plus loin.

Des impulsions codées, dites impulsions d'interrogation, sont transmises par cet émetteur qui travaille, comme on vient de le voir, sur la même

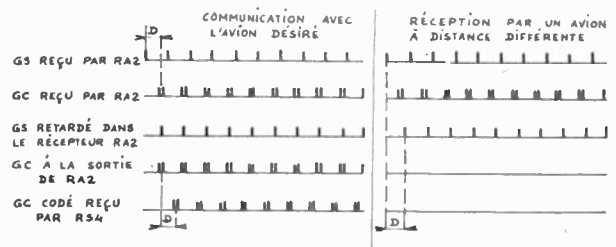
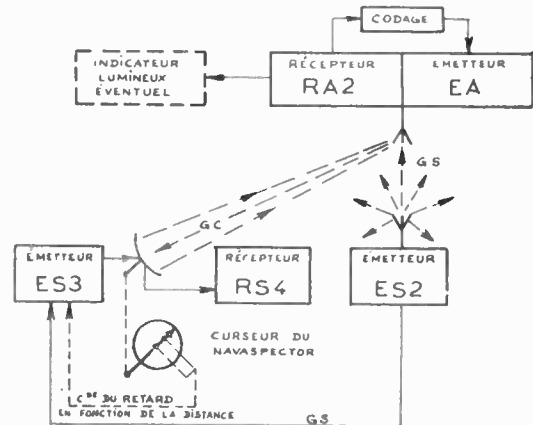


FIG. 20. — Principe du « NAVAMANDER ».

fréquence que l'émetteur non directionnel d'azimut, qui est aussi le répondeur du DME. Ces impulsions d'interrogation atteignent donc tous les avions dont le récepteur RA2 est accordé sur la station au sol et qui se trouvent dans le faisceau.

Il faut maintenant y adjoindre un moyen de sélection, afin de ne permettre la réception que sur l'avion désiré, c'est-à-dire sur celui qui se trouve à une distance donnée de l'émetteur. Ce but est atteint par un décalage proportionnel à la distance entre les impulsions de synchronisation Gs de l'émetteur non directionnel ES2 et les impulsions d'interrogation transmises par l'émetteur ES3. Ce décalage est obtenu automatiquement au sol à l'aide de l'indicateur de distance associé à l'écran du Radar Navaspector (voir plus loin) ; le déplacement d'un curseur télécommande le décalage entre les impulsions Gs et les impulsions du Navamander Gc. Or, à bord, un décalage analogue est fourni à la sortie du récepteur du DME, si bien que seules les impulsions Gc séparées de Gs par un délai caractéristique de l'appareil intéressé traversent un « filtre de distance » (dispositif de déblocage) dans le récepteur RA2 de l'appareil. Tous les autres avions restent insensibles à ces impulsions. Chaque impulsion d'interrogation Gc consiste en deux impulsions brèves rapprochées dont l'intervalle constitue le codage et est caractérisé-

cathodique et marquent les positions de tous les avions situés dans le champ de la balise interrogée. La base de temps du tube est déclenchée par les impulsions de synchronisation non directionnelles émises par la station au sol en même temps que ses propres impulsions d'interrogation. Le temps

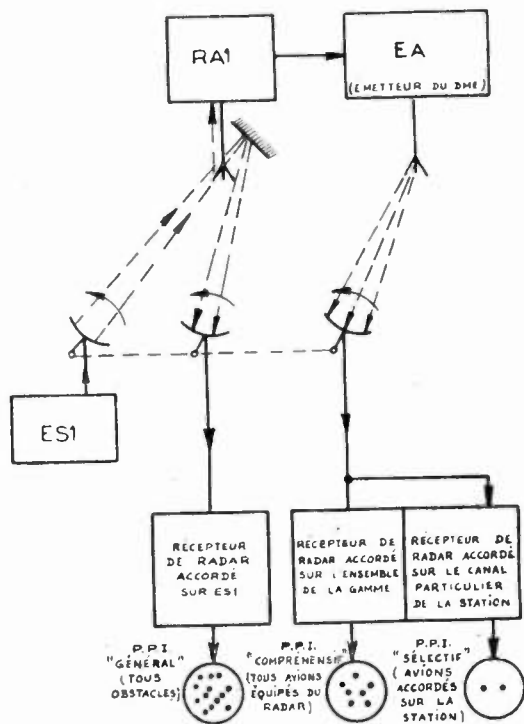


FIG. 22. — Principe du « NAVASCOPE ».

qui sépare l'impulsion d'interrogation de l'impulsion réémise représente le double du temps nécessaire au trajet balise-avion (1). C'est aussi celui qui sépare la réception de l'impulsion de synchronisation de la réception de l'impulsion de retour.

Pour obtenir une représentation correcte en azimut, il suffit d'utiliser les signaux « Nord », transmis par l'émetteur au sol non directionnel ES1 en vue de la mesure de l'azimut.

Self-identification

Parmi tous les avions ou obstacles représentés sur le tube indicateur de l'avion, il est nécessaire que le pilote reconnaisse son propre appareil. Cet effet est produit par l'application au Wehnelt du tube d'une tension très forte en même temps que l'impulsion caractéristique de l'appareil.

Cette tension est obtenue grâce à un « marqueur de distance » qui la produit seulement un certain temps après l'arrivée de chaque impulsion de synchronisation du navascope, ce temps étant tout simplement donné par le DME. De même on s'arrange (« marqueur d'azimut ») pour que cette tension ne soit appliquée que lorsque les impulsions d'azimut déclenchent l'indicateur d'azimut de l'avion. Dès lors, elle n'agit que lorsque les valeurs de distance et d'azimut correspondent à l'appareil. Lui

seul est signalé par un renforcement très marqué de son image sur le tube cathodique.

Représentation sélective de la zone d'altitude où évolue l'avion :

Il est très intéressant pour le pilote de pouvoir voir seulement, et sélectionner, la trace des avions dont l'altitude est peu différente de la sienne. On peut admettre par exemple que seuls lui importent ceux dont l'altitude diffère de la sienne de moins de 200 mètres.

A cet effet, chaque impulsion utilisée pour le navascope et effectuant les 3 trajets Sol-Avion, Avion-Sol, puis Sol-Avion, subit lors de sa réémission par l'avion un codage sous forme du dédoublement en deux impulsions dont l'écart est proportionnel à l'altitude barométrique. La station au sol réémet ce signal sans modification du codage.

Or, elle réémet aussi tous les signaux des avions au voisinage, chacun avec sa caractéristique altimétrique propre. Dès lors, à bord de l'avion considéré, on reçoit les impulsions codées en provenance de tous les avions voisins. Un système de filtrage commandé aussi par le codeur barométrique élimine celles dont la caractéristique de codage est nettement différente de celle de l'avion, de telle sorte que seules celles qui en diffèrent peu sont traduites sur le tube cathodique de bord.

Remarque. — Le fonctionnement correct du navascope exige que tous les avions soient munis de l'équipement Navar.

Autres caractéristiques :

On envisage aussi les adjonctions suivantes au Navascope :

1^o Transmission de signaux donnant électroniquement sur le tube cathodique de bord une flèche lumineuse rectiligne indiquant la direction du vent observé au sol. La linéarité de la flèche sert en même temps à vérifier le bon fonctionnement des circuits de balayage.

2^o Superposition mécanique sur l'écran de bord d'un disque transparent, portant des flèches indicatrices de la direction de l'avion, télécommandé par le compas gyroscopique.

3^o Possibilité de superposition de disques transparents colorés indiquant, pour chaque station au sol, la conformation du terrain.

En résumé le Navar présente, outre les avantages dus au système de coordonnées polaires, l'intérêt de remplir un grand nombre de fonctions avec seulement 2 récepteurs et 1 émetteur de bord sans aériens tournants ni directionnels, (figure 23). Les différentes fonctions sont séparées en général par des filtres d'impulsions dont la technique était déjà connue pendant la guerre.

Son élaboration constitue l'essentiel des recherches de radionavigation des Federal Telecommunication Laboratories. Il est complété pour l'atterrissage et le contrôle général du trafic par le Navaglide et le Navascreen.

2^o Le Navaglide

Le Navar est l'équipement de radionavigation

(1) Abstraction faite du temps passé dans les équipements, qui est fixe, et dont on tient compte.

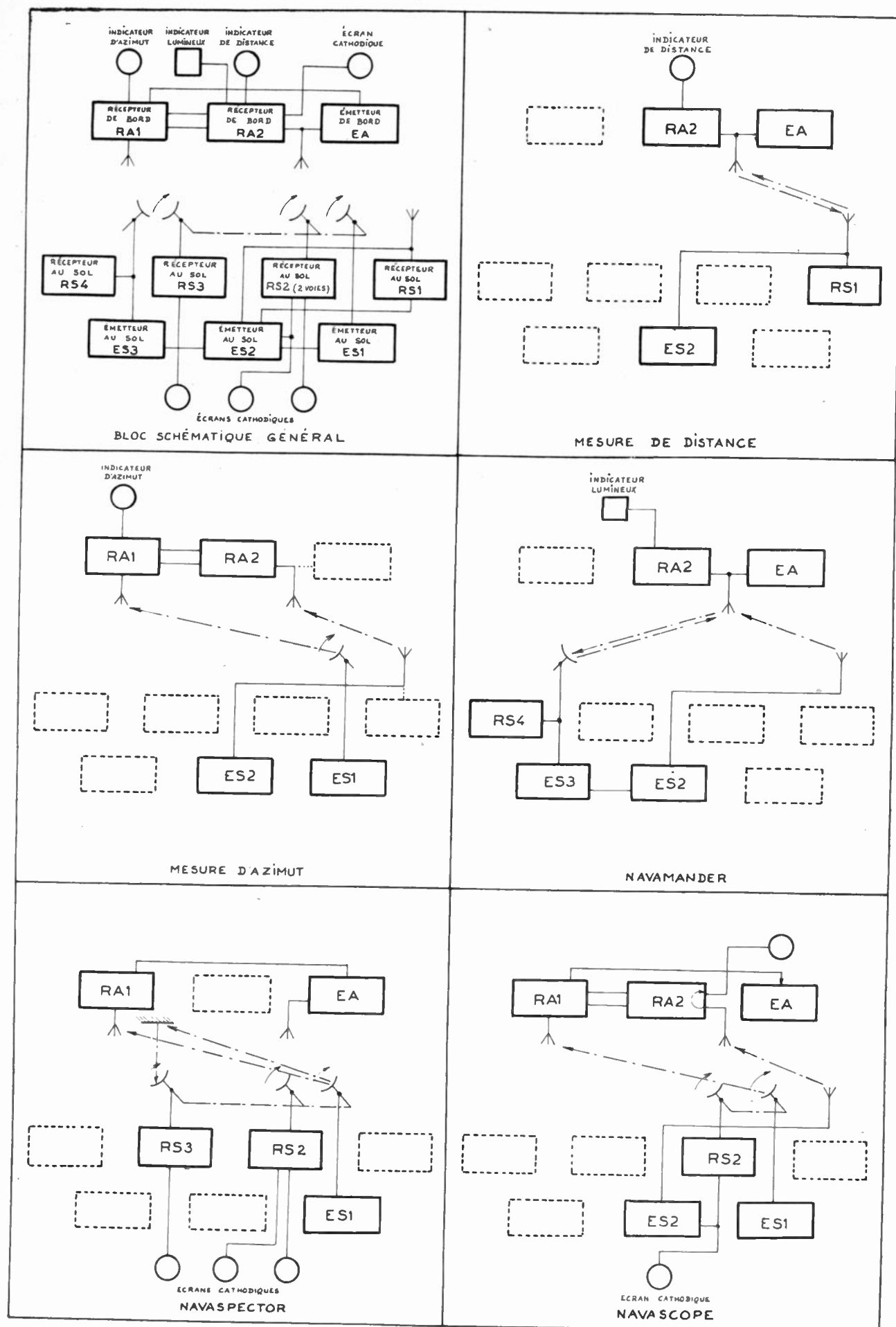


FIG. 23. — Ensemble du système « NAVAR » et diverses fonctions

et de contrôle du trafic à petite distance. Le Navaglide est l'équipement de radioguidage lors de l'atterrissage.

Le développement du Navaglide n'est pas considéré comme très urgent puisque les équipements d'atterrissage sans visibilité actuels sont, pense-t-on généralement, aptes à remplir tous les besoins pendant plusieurs années.

Toutefois, on peut se rendre compte qu'il serait avantageux d'utiliser des ondes ultra-courtes qui permettent l'emploi d'aériens très directifs, évitant les réflexions indésirables. Il serait d'autre part souhaitable de donner à l'avion, à tout instant, sa distance au point d'atterrissage. On peut y parvenir par l'emploi d'une balise répondeuse tout à fait analogue à une balise de Navar (DME), fonctionnant dans la même gamme, avec le même récepteur de bord.

La détermination de la trajectoire de descente elle-même s'effectue dans le Navaglide (figure 24) avec l'indicateur habituel à aiguilles croisées, mais la différence essentielle de principe entre le Nava-

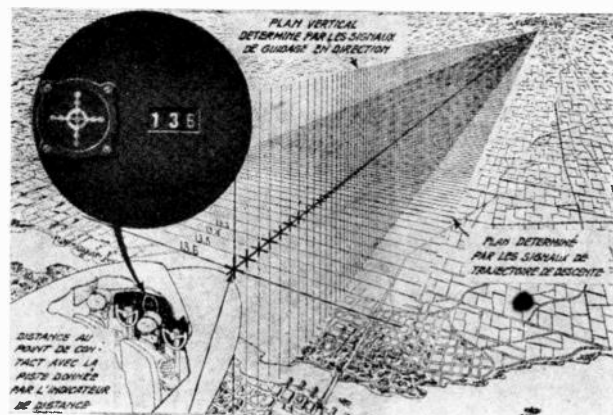


FIG. 24. — NAVAGLIDE.

glide et les équipements actuels est que les faisceaux de détermination du plan vertical de l'axe et de la trajectoire de descente sont émis sur la même fréquence, et successivement. Ils sont reçus à bord sur un seul appareil qui les distingue par des modulations à fréquence très élevée. Le rythme de commutation est suffisamment rapide pour donner à la sortie du récepteur l'impression d'une émission continue de 4 signaux simultanés.

3° Le « Navascreen »

Le Navascreen a pour but de réunir en un seul point pour chaque aéroport ou chaque centre de contrôle du trafic tous les renseignements connus sur les appareils évoluant au voisinage.

Il réalise sur un grand écran muni d'une carte du terrain la projection d'indicatifs lumineux donnant chacun la position et l'identité de l'appareil, ainsi que l'altitude (par la couleur de cette projection) et la direction (par une flèche).

D'autre part, un dispositif électromécanique fait mouvoir automatiquement chaque indicatif en fonction du déplacement de l'image sur le PPI correspondant. De même la couleur change avec la couche d'altitude utilisée.

Inversement, l'image obtenue sur le grand écran est reprojétée sur chaque tube cathodique PPI de façon que les observateurs puissent parer à toute erreur et corriger toute discordance entre la position relevée sur le PPI et la position calculée et projetée sur l'écran du Navascreen.

Le système présente aussi l'intérêt de pouvoir donner à l'avance, sur interrogation, la situation probable du trafic dans un délai déterminé.

Les informations rassemblées sur l'écran du Navascreen sont recueillies par des observateurs qui les envoient aux autres aéroports ou au centre général de contrôle du trafic de la région, où elles peuvent à nouveau être figurées sous forme de tableaux lumineux.

Conclusion

Les systèmes Navar, Navaglide et Navascreen ne sont pas en tous points conformes aux recommandations de la RTCA américaine, mais ils procèdent de la même inspiration. En fait la RTCA semble être allée plus loin encore que le projet initial des Federal Telecommunication Laboratories, vieux maintenant de plusieurs années.

Il est agréable de mentionner pour finir que cet essai de vue d'ensemble d'un problème général et complexe est pour une part l'œuvre d'ingénieurs français expatriés. Souhaitons que sur le sol métropolitain aussi de telles études soient engagées avec le souci d'aboutir dans un délai raisonnable et de donner à l'aéronautique des aides radioélectriques complètes, économiques en poids, en prix et en personnel. Il ne s'agit pas tant en effet de réaliser d'importants progrès techniques que d'imaginer des mises en œuvre rationnelles de dispositifs existants, en gardant une connaissance lucide de leurs avantages et surtout de leurs imperfections et de leurs limites, de façon à en tirer le meilleur parti possible.

APPENDICE

Détails techniques sur l'équipement DME Federal

1) Émetteur-récepteur de bord :

- 51 canaux prééglés,
- Consommation : 250 watts ;
- Puissance de crête : 2 kw ;
- Poids : 20 kgs ;
- Réjection des canaux adjacents : 70 db ;
- Réjection de la moyenne fréquence : 80 db ;
- Réjection de la fréquence image : 70 db ;
- Sensibilité effective du récepteur : 15 microvolts.

Antenne unique à large bande, verticale, omnidirectionnelle dans le plan horizontal.

L'émission s'effectue sur une fréquence inférieure à celle de la réception (émission entre 960 et 1085 Mc/s et réception entre 1090 et 1215 Mc/s), et la différence des deux fréquences est toujours de 125,5 Mc/s. Ce changement de fréquence est effectué par la balise répondeuse.

À l'émission un oscillateur à impulsions constitué

d'un tube «light-house» 2 C 39, dont la fréquence est déterminée par deux conducteurs coaxiaux cylindriques, fournit l'oscillation à l'antenne. Des impulsions de 2500 volts sont appliquées à sa plaque, et assurent alors une puissance de crête de sortie de 2 KW. Ces impulsions sont fournies elles-mêmes par un thyatron attaqué par un oscillateur blocking ; elles ont une forme approximativement trapézoïdale, leur durée moyenne est de 1,8 microseconde, le temps de croissance étant de 0,25 microseconde et le temps de décroissance de 0,4 microseconde, avec une fréquence de répétition de l'ordre de 30 par seconde.

Le récepteur comporte un oscillateur à quartz sur 45 Mc/s environ, suivi de trois étages classe C (un tripleur et deux doubleurs) à large bande. Les signaux reçus par l'antenne, commune à l'émission et à la réception, passent dans un présélecteur à cavités résonnantes couplées et subissent un changement de fréquence dans deux cristaux IN 28 excités en parallèle par le signal et en push-pull par la sortie du multiplicateur, ce qui est équivalent à un push-push pour l'harmonique 2 de ce multiplicateur. C'est finalement entre cet harmonique et le signal incident que se fait le changement de fréquence.

La fréquence ainsi obtenue a une valeur nominale de 64 Mc/s pour le porteur. Elle est appliquée à un amplificateur MF à large bande (8 Mc/s) puis à un amplificateur MF principal d'environ 1,5 Mc/s de largeur de bande. Les impulsions attaquent alors le Strobe (voir plus loin).

La commutation des quartz de l'oscillateur local du récepteur s'accompagne du réglage mécanique des cavités d'accord du récepteur et de la cavité qui détermine la fréquence d'émission correspondante. La stabilisation exacte de la fréquence d'émission est réalisée de la façon suivante :

Une petite fraction de la tension transmise à l'antenne subit, pendant l'émission, le même changement de fréquence que le signal de retour, mais, maintenant elle se trouve à 61,5 Mc/s au-dessous de la fréquence de l'oscillateur local. Le signal MF de fréquence nominale 61,5 Mc/s ainsi obtenu est amplifié dans le préamplificateur MF à bande large, puis actionne un dispositif de correction automatique de fréquence qui vient corriger mécaniquement le réglage de la cavité oscillatrice de façon à ramener la fréquence porteuse à 61,5 Mc/s exactement au-dessous de la fréquence de référence (24 fois la fréquence de l'oscillateur local). Ce dispositif de correction est tout à fait classique. Son originalité réside dans son application à une cavité résonnante.

Pratiquement, après amplification des signaux MF, un discriminateur de type usuel fournit des impulsions dont la polarité dépend du sens du désaccord. Elles sont amplifiées et, suivant leur signe, actionnent l'un ou l'autre de deux relais qui peuvent faire tourner, dans un sens ou dans l'autre, le moteur d'ajustage précis de la cavité. Comme le réglage principal, cet ajustage précis est obtenu en faisant varier la longueur du conducteur interne au moyen

d'une partie tubulaire télescopique. La construction mécanique est très soignée, en vue d'éviter tout jeu nuisible. Pendant l'émission, une fraction de la tension de sortie sert à débloquent le dispositif de correction automatique de fréquence qui est inopérant dans les intervalles entre les impulsions d'émission.

Lors de la réception, les signaux sont amplifiés dans le préamplificateur MF à bande large, puis dans l'amplificateur MF à bande étroite, sur 64 Mc/s.

La détection s'effectue dans un discriminateur spécial qui assure une très bonne protection contre les signaux situés dans les canaux adjacents (70 db). Il combine en opposition la tension aux bornes du primaire d'un transformateur accordé sur 64 Mc/s, laquelle présente deux bosses sur les 2 canaux adjacents, et la tension aux bornes du secondaire, laquelle ne présente qu'une bosse centrée sur le canal désiré. Pour un mélange convenable des 2 tensions, les impulsions qui se trouvent sur le canal désiré ont une polarité, et celles des canaux adjacents la polarité inverse. Il est alors facile de les éliminer.

Après cette détection, les impulsions sont amplifiées dans deux étages et attaquent le « Strobe » qui est la partie fondamentale de l'appareil.

Un châssis dénommé « Sanaphant » fournit, à partir de l'impulsion d'émission une impulsion retardée qui est différenciée et transformée en deux impulsions identiques.

L'équipement est construit de façon à maintenir automatiquement l'impulsion de retour « à cheval » sur ces deux impulsions retardées. La constante de temps qui détermine le retard est elle-même commandée par une tension continue qui est par ailleurs envoyée à l'indicateur de distance. Chaque fois que les deux impulsions retardées sont en avance ou en retard sur l'impulsion de réponse, une détection différentielle donne une tension positive ou négative qui corrige indirectement la tension continue de commande du retard. Le dispositif rappelle donc un dispositif correcteur automatique de fréquence, mais c'est un dispositif de correction automatique de temps de retard.

En même temps la vitesse de variation de ce temps de retard mesure la vitesse radiale de l'avion ce qui, par l'introduction d'une constante de temps appropriée, permet de maintenir le fonctionnement du DMF pendant quelques secondes même si ces signaux disparaissent, en supposant constante cette vitesse radiale (opération de « prédiction »).

Si, au bout de ce temps, les réponses sont définitivement perdues, il se déclenche un dispositif de balayage du temps de retard relativement lent quoique beaucoup plus rapide que celui qui correspondrait au déplacement de n'importe quel avion. Parmi toutes les impulsions qui sont les réponses de la balise interrogée à tous les avions interrogateurs, le dispositif stroboscopique ne rencontre qu'un groupe d'impulsions dont le rythme puisse correspondre aux impulsions émises par l'avion. C'est sur lui que se stabilise immédiatement, sui-

vant le procédé précédemment indiqué, le dispositif de commande automatique du temps de retard, tandis que le système de recherche stroboscopique cesse d'être alimenté.

En résumé, le pilote sélectionne sa fréquence d'émission, et, du même coup, sa fréquence de réception, pour interroger la balise qu'il désire. La balise répond indifféremment à tous les avions qui l'interrogent.

L'appareil reçoit toutes les réponses et choisit celles qui correspondent à son rythme d'émission ; c'est après avoir fait ce choix qu'il mesure sa distance à la balise et en suit les variations.

L'équipement peut évidemment être relié à un pilote automatique en vue de décrire des trajectoires circulaires, de rayon donné, autour de la balise.

2) Balise répondeuse au sol.

L'aérien de la balise répondeuse, commun à la réception et à l'émission est constitué d'éléments rayonnants coniques empilés donnant une directivité marquée dans le plan horizontal.

Lors de la réception, les signaux sont filtrés dans des cavités résonnantes, puis subissent un changement de fréquence à l'aide d'un oscillateur à cristal suivi d'étages multiplicateurs. L'amplificateur à

moyenne fréquence, accordé sur 31,25 Mc/s a une largeur de bande de 1,5 Mc/s (7 étages).

Il est suivi d'un détecteur analogue à celui de l'équipement de bord et de deux étages d'amplification. Le gain est tel qu'un signal d'entrée de 5 microvolts permet la manipulation correcte de l'émetteur. Les impulsions subissent auparavant un retard réglable, de telle sorte que le retard total introduit par la balise répondeuse soit de 8 microsecondes, puis elles modulent l'émetteur. Lors de cette réémission, sur une fréquence qui diffère de la fréquence de réception de 125,5 Mc/s, la puissance de crête est de 20 KW. Un dispositif de correction automatique de fréquence est nécessaire. A cet effet, un oscillateur stabilisé par quartz, suivi d'étages multiplicateurs de fréquence, produit un battement au voisinage de 5 Mc/s avec le signal émis. Après amplification à large bande autour de 5 Mc/s, les signaux résultants sont injectés dans deux filtres, l'un passe-haut, l'autre passe-bas, suivis d'une détection différentielle. La tension continue sert à actionner par l'intermédiaire de thyatron et de relais un moteur qui tourne dans le sens convenable pour commander la fréquence d'accord de la cavité résonnante.

La réjection du signal adjacent est de 70 décibels. La portée utilisable est de l'ordre de 150 à 200 km.

LES ANTENNES EN HYPERFRÉQUENCES

PAR

Le Capitaine de Corvette J. MAILLARD

du Centre National d'Études des Télécommunications

Division « Tubes et Hyperfréquences »

I. — Généralités sur les antennes radioélectriques ordinaires.

Les antennes permettent d'envoyer à travers l'espace des ondes électromagnétiques ou de capter de telles ondes.

Faisons en premier lieu une remarque fondamentale : une antenne ne rayonne pas par elle-même, à la façon dont rayonnent une corde de piano qui vibre ou les parois d'un caisson musical. L'énergie rayonnante électromagnétique n'est pas emmagasinée par l'antenne puis rejetée par elle dans l'espace. Toutes les antennes radioélectriques ne sont pas autre chose que des guides d'ondes. Elles permettent aux ondes qui arrivent canalisées par

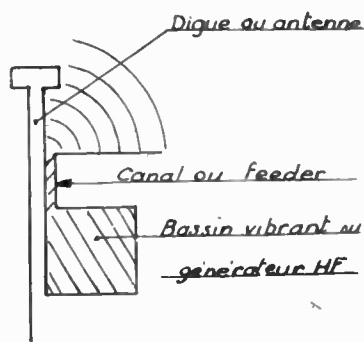


FIG. 1.

les feeders soit de s'épanouir progressivement, soit de trouver un volume résonnant ayant une large ouverture et ainsi d'attaquer l'espace dans de bonnes conditions. Les conducteurs qui constituent l'antenne ne distribuent pas mais absorbent (d'autant moins qu'ils sont plus conducteurs), de l'énergie électromagnétique qui pénètre par leur surface extérieure.

L'antenne se conduit comme une digue de dimensions telles qu'elle transmet au large (fig. 1) des vibrations de l'eau produites dans un bassin. Sans la digue les vibrations du bassin se repercuteraient beaucoup moins au large mais il est clair que la digue ne vibre pas et absorbe une partie de l'énergie rayonnante.

Les caractéristiques d'une antenne expriment d'une part le champ rayonné dans l'espace c'est-à-dire ses qualités pour la communication envisagée et d'autre part l'impédance qu'elle présente aux bornes du feeder qui l'alimente c'est-à-dire ses qualités pour extraire du générateur haute fréquence le maximum d'énergie.

Étant donné une antenne de forme quelconque alimentée en deux points, il est théoriquement possible, en intégrant les équations de Maxwell, de connaître pour chaque fréquence l'impédance de cette antenne et le champ qu'elle rayonne. Cependant, même pour les formes les plus simples (ellipsoïde, deux cônes...) les calculs sont très compliqués et ils deviennent inextricables pour des structures plus complexes.

En pratique on procède par approximations en s'arrêtant généralement à la première. Pour ce faire on suppose a priori la distribution des courants dans l'antenne. Le champ électrique rayonné à grande distance découle de la formule :

$$(1) \quad E = -j \frac{k}{c} \int \varphi (\vec{I})_N dl \quad \text{en unités de Gauss} \quad (1)$$

le courant est supposé de la forme $e^{j(\omega t + \theta)}$ en tous les points du conducteur ; $j = \sqrt{-1}$; ω est la pulsation $\omega = 2\pi / \lambda$ étant la fréquence ; c est la vitesse de la lumière ; $k = \frac{\omega}{c} = \frac{2\pi}{\lambda}$, λ étant la longueur d'onde ; $\vec{I}$ est le vecteur intensité de courant s'écoulant suivant le conducteur dont dl est un élément de longueur. $(\vec{I})_N$ est la composante de ce vecteur perpendiculaire à la direction du point où est calculé le champ électrique. Enfin $\varphi = \frac{e^{-jkr}}{r}$

r étant la distance séparant l'élément dl du point où E est calculé.

La façon dont $\vec{E}$ est calculé montre que le vec-

(1) Nota : Les unités de Gauss correspondent à $\mu_0 = 1$ $\epsilon_0 = 1$ dans les relations fondamentales de l'électrostatique et du magnétisme. E est exprimé en unités électrostatiques (U. E. S.) H en Gauss.

teur obtenu est perpendiculaire à la direction de propagation. Nous savons que le champ magnétique $\vec{H}$ est aussi dans un plan perpendiculaire à cette direction, est en phase avec $\vec{E}$ et est perpendiculaire à $\vec{E}$ de telle sorte que $\vec{E}$, $\vec{H}$, $\vec{P}$ ($\vec{P}$ étant un vecteur ayant même direction et même sens que l'onde qui se propage) forment un trièdre de sens direct (fig. 2). On a en grandeur 11^{Gauss}

$$= P^{\text{ues}} = \frac{E^{\text{Volt/cm}}}{300}$$

La connaissance des champs fournit la puissance rayonnée par l'antenne. On sait que le flux d'énergie rayonnante est donné par le vecteur de Poynting $\vec{P} = \frac{c}{4\pi} \vec{E} \times \vec{H}$ ou en unités pratiques pour

$$\text{une onde plane (2) } P^{\text{Watts/cm}^2} = \frac{10}{4\pi} E^{\text{Volt/cm}} \times 11^{\text{Gauss}}$$

$$= \frac{E^{\text{Volt/cm}}}{120\pi}$$

A travers une surface ds il s'écoule une puissance instantanée égale à $P ds \cos a$, a étant l'angle que fait $\vec{P}$ avec la normale à cette surface (fig. 3).

De la connaissance des courants et de la puis-

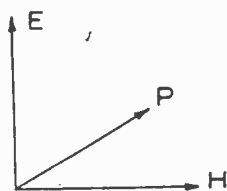


Fig. 2

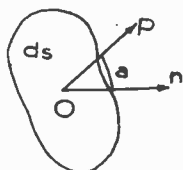


Fig. 3

sance rayonnée dans l'espace on déduit aisément la résistance de rayonnement de l'antenne et plus précisément la résistance présentée par l'antenne au feeder qui l'alimente. Il est très légitime d'opérer ainsi pour les raisons suivantes :

1° Le long des conducteurs filaires les courants principaux pour les fréquences envisagées sont dirigés suivant les dimensions longitudinales des fils ; les dimensions transversales des fils sont très petites devant les longueurs d'onde.

2° Les antennes étudiées se comportent soit comme un circuit résonnant ayant un bon coefficient de surtension, soit comme une ligne parcourue par des ondes progressives. En ces deux cas extrêmes la répartition des courants dépend très peu de la valeur de la résistance de rayonnement calculée d'après cette répartition supposée a priori.

3° Il est plus facile et plus rapide pour avoir une antenne répondant à certaines conditions de se servir d'approximations aisément concevables et calculables que de vouloir résoudre rigoureusement le problème à partir des conditions fixées.

Cependant, en opérant ainsi, on n'obtient pas la valeur de la réactance de l'antenne, c'est-à-dire la vraie valeur de l'impédance de l'antenne et sa

variation en fonction de la fréquence. Cette variation conditionne la largeur de bande de fréquences dans laquelle l'antenne est utilisable. Si l'antenne est parcourue par des ondes progressives, sa largeur de bande est très grande et généralement plus que suffisante. Si elle est le siège d'ondes stationnaires, elle présente pour certaines fréquences (d'autant moins éloignées de la fréquence fondamentale que son coefficient de surtension est plus élevé) une réactance telle que peu de puissance est alors extraite du feeder. Comme il a été dit ci-dessus, le calcul de cette réactance est trop laborieux ou même pratiquement impossible. On la mesure expérimentalement sur l'antenne construite avec les données principales : diagramme de rayonnement, résistance pour la fréquence fondamentale. Pour augmenter la largeur de bande il faut diminuer le coefficient de surtension. Pour cela il suffit d'augmenter le diamètre des fils ou de les remplacer par des cônes. En augmentant le diamètre des fils on diminue l'énergie réactive qui se trouve concentrée très près des conducteurs ; en disposant des cônes on oblige les ondes à s'épanouir plus progressivement qu'avec les conducteurs cylindriques ce qui diminue les réflexions d'ondes, donc l'énergie réactive. La largeur de bande se définit habituellement par la différence des deux fréquences pour lesquelles la puissance transmise est la moitié de la puissance transmise sur la fréquence fondamentale.

Les considérations ci-dessus paraissent s'appliquer uniquement aux antennes d'émission. En fait il existe entre l'émission et la réception une réciprocité telle que les caractéristiques d'une antenne peuvent être considérées comme indépendantes du sens dans lequel l'antenne est utilisée. J. R. Carson a démontré que si une force électromotrice U appliquée en un point a_1 d'une antenne A_1 produit un courant I en un autre point a_2 d'une antenne A_2 , réciproquement la même force électromotrice U appliquée en a_2 produit en a_1 le même courant I . De ceci découle que la directivité d'une antenne, sa largeur de bande, son adaptation optima sont les mêmes à l'émission et à la réception et que le rapport de la puissance maxima émise par une antenne à la puissance maxima captée par une seconde antenne est indépendant du sens de la transmission.

On peut concevoir physiquement cette réciprocité en gardant les hypothèses justifiées ci-dessus, à savoir qu'il existe pour l'antenne étudiée un état vibratoire des grandeurs électriques bien déterminé auquel correspond une répartition précise des courants. On ne fait de différence entre l'émission et la réception qu'en adoptant pour celle-ci une distribution de courant où les phases relatives θ de l'émission sont changées en $-\theta$ si bien que le diagramme de rayonnement de l'antenne isolée est symétrique par rapport au centre de l'antenne, du diagramme à l'émission. Ceci revient pour une onde progressive parcourant l'antenne à changer le sens de propagation de cette onde et le sens de transfert de l'énergie. Or il n'y a interaction énergétique entre les ondes planes que pour celles qui ont même fréquence et cheminent suivant la même direction et le même sens. Deux ondes planes qui se croisent interfèrent

mais la puissance moyenne diffusée par ces deux ondes à travers un plan indéfini est la somme des puissances moyennes évaluées indépendamment pour chacune d'elles. Il n'en est pas de même pour deux ondes cheminant suivant le même sens qui peuvent rayonner une puissance nulle si leurs amplitudes se détruisent. Il en résulte que l'antenne de réception placée dans le champ d'une onde plane rayonne de l'énergie suivant son propre diagramme sans aucune perturbation due à l'onde incidente. Seulement, dans la région où les ondes émises par

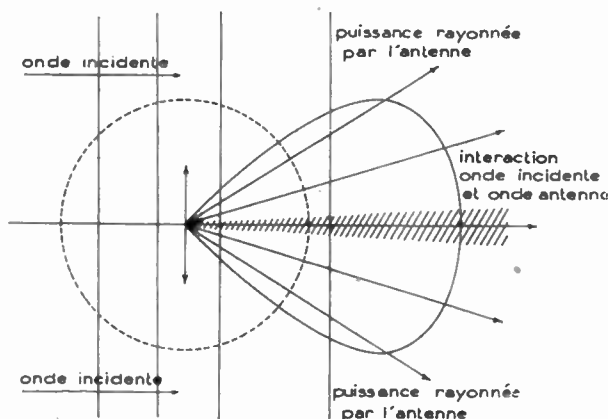


Fig. 4

l'antenne et l'onde incidente ont même sens de propagation il y a une réaction énergétique si bien que le flux moyen du vecteur de Poynting à travers une surface fermée entourant l'antenne dépend

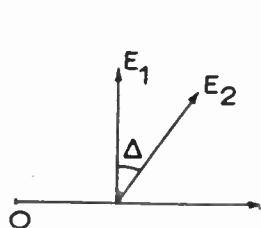


Fig. 5

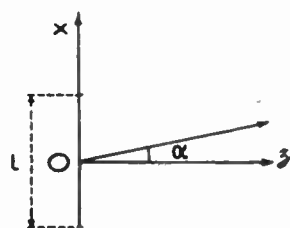


Fig. 6

de l'onde incidente et de l'onde rayonnée suivant cette direction. (fig. 4) Le calcul montre que la puissance maxima qu'une antenne peut extraire d'une onde plane incidente vaut.

$$(3) \quad W_s = c \frac{\lambda^2}{(4\pi)^2} \cos^2 \Delta \cdot g_s \cdot E_1^2$$

E_1 étant la grandeur du champ de l'onde incidente Δ l'angle que fait avec E_1 le Champ E_s émis par l'antenne dans la direction vers laquelle se propage l'onde incidente (fig. 5) et g_s le gain de l'antenne suivant cette direction. Ce gain se définit, le champ dû à l'antenne étant seul considéré, comme le rapport de la puissance rayonnée dw dans un angle solide élémentaire $d\Omega$ suivant cette direction à la puissance totale rayonnée W_T divisée par $\frac{4\pi}{d\Omega} g_0 = \frac{4\pi}{d\Omega} W_T$. Ce gain mesure la directivité de l'antenne suivant

la direction considérée. Pour une antenne omnidirectionnelle il est égal à 1 dans toutes les directions.

Si le champ électrique E_1 provient d'un émetteur ayant un gain g_1 dans la direction du récepteur nous aurons.

$$(4) \quad E_1^2 = g_1 \frac{W_1}{c R^2} \text{ U. E. S.}$$

W_1 étant la puissance totale rayonnée par l'émetteur et R la distance séparant l'émetteur du récepteur. Des formules (3) et (4) se déduit la formule fondamentale des communications par ondes.

$$(5) \quad \frac{W_s}{W_1} = \frac{g_1 g_s}{R^2} \frac{\lambda^2}{(4\pi)^2} \cos^2 \Delta$$

La symétrie de cette formule exprime bien la réciprocité entre l'émission et la réception.

Cette formule montre tout l'intérêt qu'il y a à augmenter le gain pour améliorer une communication radioélectrique et pour obtenir un gain élevé il faut que l'antenne soit fortement directive ce qui conduit à avoir des aériens dont les dimensions sont grandes devant la longueur d'onde utilisée.

Supposons pour prendre l'exemple le plus simple que nous ayons une base de longueur l sur laquelle

les vecteurs élémentaires ($\vec{I} dl$) ont la même direction. La formule (1) s'écrira

$$(6) \quad \vec{E} = -j \frac{k}{c} (\vec{I}_0)_N \varphi_0 \int_{-\frac{l}{2}}^{+\frac{l}{2}} e^{j(kx \sin \alpha + \theta)} h(x) dx$$

I_0 et φ_0 représentent les valeurs de I et de φ pour le milieu de la base. $h(x) dx$ indique les variations d'amplitude du vecteur $I dl$ le long de la base et θ ses variations de phase. α est le complément de l'angle que fait avec la base la direction suivant laquelle E est estimée (fig. 6) La directivité de la base est donnée par le facteur D qui dans la formule (6) traduit l'influence de l'angle α

$$(7) \quad D(\alpha) = \frac{1}{l} \int_{-\frac{l}{2}}^{+\frac{l}{2}} h(x) e^{j(kx \sin \alpha + \theta)} dx$$

Soit par exemple $\theta = 0$. $h(x) = 1$ (tous les éléments sont en phase et ont même amplitude). La directivité dans le plan xoy est donnée par le facteur.

$$(8) \quad D_1 = \frac{\sin \left(\frac{kl}{2} \sin \alpha \right)}{\frac{kl}{2} \sin \alpha}$$

ce qui se traduit par la courbe cartésienne de la fig. 7 ou le diagramme polaire de la fig. 8. On voit qu'une telle distribution de courants est directive perpendiculairement à la base. Le demi-angle d'ouverture du lobe principal vaut $\frac{\lambda}{l}$ mais il existe des lobes latéraux.

On peut réduire ces lobes latéraux en prenant une fonction $h(x)$ convenable. La formule (7) montre que la fonction $D_1(\alpha)$ n'est autre que la transformée de Fourier de la fonction $h(x)$. En remplaçant la

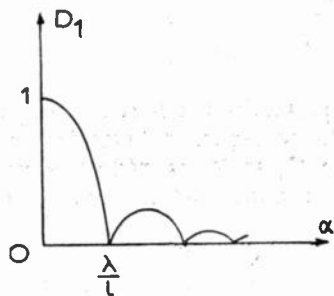


FIG. 7

forme rectangulaire de la fonction $h(x)$ par une forme en cloche ou même simplement par une forme triangulaire ce qui dans tous les cas revient à diminuer l'amplitude du courant sur les bords de la base, nous réduirons les lobes latéraux, mais

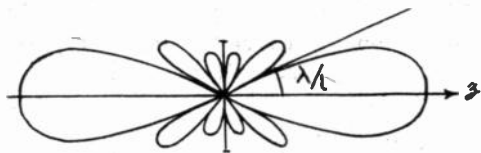


FIG. 8

augmenterons la largeur du lobe principal comme l'indique la figure 9. Il y a donc un compromis entre d'une part la directivité et le gain et d'autre part l'existence de lobes secondaires.

La directivité exprimée par la formule (8) n'est

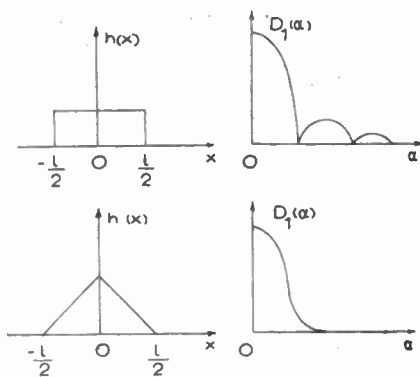


FIG. 9

valable que dans un plan. Pour avoir une directivité dans deux plans rectangulaires il convient que la base soit une véritable surface. En ce cas la directivité est donnée par le facteur.

$$(9) \quad D_2(\alpha, \beta) = \frac{\sin\left(\frac{ka}{2} \sin \alpha\right)}{\frac{ka}{2} \sin \alpha} \cdot \frac{\sin\left(\frac{kb}{2} \sin \beta\right)}{\frac{kb}{2} \sin \beta}$$

a, b étant les côtés de la base (fig. 10) et α, β les

angles de la direction OP avec les plans yoz et xoz. Bien entendu une telle base rayonnerait symétriquement suivant la direction Oz. Pour que le rayonnement n'existe que d'un côté il conviendrait de placer un écran réflecteur.

Supposons maintenant que sur la base de longueur l on ait $h(x) = 1$ et $\theta = -q k x$ c'est-à-dire que tous les éléments ont même amplitude mais

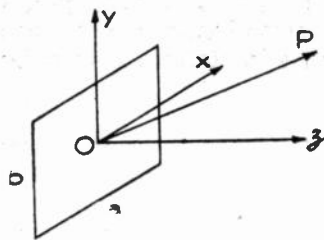


FIG. 10

qu'il existe le long de Ox une différence de phase correspondant à une onde progressive se propageant avec la vitesse de phase $v = \frac{c}{q}$

La directivité est alors donnée par la formule :

$$(10) \quad D_s = \frac{\sin\left[\frac{kl}{2}(\sin \alpha - q)\right]}{\frac{kl}{2}(\sin \alpha - q)}$$

Si $q < 1$ cette expression est analogue à la formule 8 sauf que la directivité principale n'a plus lieu pour $\alpha = 0$ mais pour un angle α_0 tel que $\sin \alpha_0 = q$ (fig. 11). Pour $q = 1$, c'est-à-dire lorsque la vitesse de phase est égale à la vitesse de propagation des

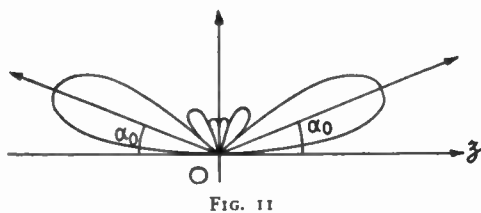


FIG. 11

ondes dans l'espace, $\alpha_0 = \frac{\pi}{2}$. En ce cas une simple base de longueur 1 donne un pinceau directif situé dans un cône dont le demi angle d'ouverture est donné par la relation $\frac{kl}{2}(\sin \alpha - 1) = -\pi$ d'où

la valeur approximative $\sqrt{\frac{2\lambda}{l}}$ pour cet angle. L'axe du pinceau coïncide avec la direction de la base et son sens est celui de la propagation des ondes sur la base (fig. 12). Si q est légèrement supérieur à 1 (mais tel que $(q - 1) \frac{kl}{2} < \pi$) c'est à-dire si la vitesse v est légèrement inférieure à la vitesse c , nous avons encore un pinceau dirigé suivant Ox et plus étroit que pour $q = 1$. Cependant la valeur relative des lobes secondaires a augmenté. Nous retrouvons ici sous une autre forme le compromis déjà signalé

ci-dessus entre l'étroitesse du pinceau et les lobes secondaires.

Quoiqu'établies pour des bases très simples et supposées continues les formules (7) (8) (9) (10) suffisent pour concevoir les structures des aériens directifs. En particulier, elles s'appliquent conve-

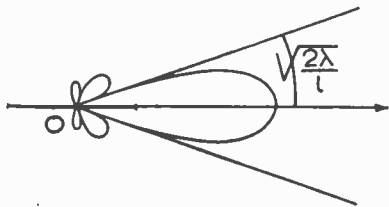


FIG. 12

nablement aux structures discontinues des antennes réelles si les divers éléments sont espacés de moins d'une fraction de longueur d'onde. Si cette condition-ci n'est pas remplie il existe plusieurs azimuts dans lesquels tous les éléments rayonnent en phase.

Enfin pour clore ces généralités sur les antennes, il convient de mentionner la dualité qui existe

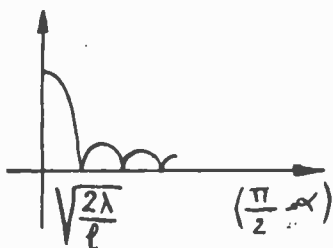


FIG. 12 bis

entre les champs magnétiques et électriques. A tout diagramme de rayonnement correspond un diagramme de rayonnement conjugué où les champs électriques et magnétiques sont permutés. A ces diagrammes correspondent aussi des antennes conjuguées dont l'exemple le plus simple réside dans le doublet et une boucle élémentaire. Le diagramme du rayonnement de la boucle est le même que celui de doublet où l'on a permuté champ électrique et champ magnétique. Une telle considération permet de transposer facilement des résultats déjà obtenus et de réaliser, sans autres calculs, de nouveaux modèles d'antennes.

II. — Particularités des antennes en hyperfréquences.

Toutes les propriétés énoncées ci-dessus pour les antennes s'appliquent évidemment aux antennes pour hyperfréquences. Celles-ci peuvent être constituées par simple transposition des antennes ordinaires et comme les dimensions géométriques varient proportionnellement aux longueurs d'onde il en résulte soit des aériens peu encombrants, soit des possibilités de directivité qui sont pratiquement interdites aux ondes plus longues.

Cependant alors que pour les ondes longues et même pour les ondes métriques les notions de différence de potentiel entre les conducteurs ou de courants suivant leurs longueurs sont toutes naturelles,

elles perdent toute signification en hyperfréquences lorsque les dimensions transversales des conducteurs sont du même ordre de grandeur que les longueurs d'onde utilisées λ . En outre, pour réduire les pertes, l'alimentation des antennes se fait par des guides d'ondes semblables aux tuyaux acoustiques qui canalisent les ondes non entre deux conducteurs comme cela se passe aux fréquences ordinaires mais à l'intérieur d'une enveloppe métallique continue sur laquelle la répartition des courants est complexe. Aussi préfère-t-on en hyperfréquences effectuer tous les calculs en considérant non les courants sur les conducteurs mais les champs électriques et magnétiques que ces conducteurs canalisent. Dans les calculs les conducteurs interviennent comme des surfaces sur lesquelles les champs doivent satisfaire à certaines conditions, par exemple pour un conducteur parfait la composante tangentielle du champ électrique doit y être nulle.

Pour effectuer ces calculs on fait en hyperfréquences les mêmes hypothèses que dans les ondes radio-électriques ordinaires en substituant la notion de champ à celle de courant : on suppose qu'à chaque antenne correspond un état vibratoire bien déterminé et que cet état se traduit par une répartition des champs bien définie. Ces hypothèses restent légitimes car les guides d'onde qui alimentent les antennes et ainsi engendrent l'état vibratoire supposé sont généralement conçus pour ne laisser passer qu'un seul type d'onde c'est-à-dire pour n'admettre qu'une seule répartition des champs.

Pour calculer les champs à grande distance on ne peut utiliser la formule (1) qui fait intervenir les courants. On se sert des formules de Kottler qui permettent de connaître le champ en tout point extérieur à une surface enfermant les antennes d'émission si l'on connaît la distribution des champs sur cette surface. A grande distance le champ est donné par la formule suivante :

$$(11) \quad \vec{E} = \frac{jk}{4\pi} \iint_S \varphi [\vec{u} \times (\vec{E}_1 \times \vec{n}) + (\vec{H}_1 \times \vec{n})\vec{N}] dS$$

en unités de gauss, où S désigne la surface sur laquelle est prise l'intégrale. En pratique cette surface

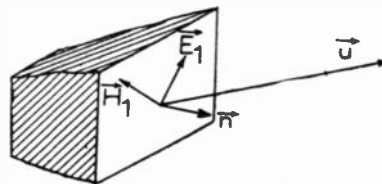


FIG. 13

se compose de surfaces de conducteurs sur lesquelles les champs sont nuls et d'une surface sur laquelle on a supposé une certaine répartition des champs (fig. 13).

E_1 et H_1 expriment les champs électriques et magnétiques sur la surface S. $\vec{n}$ désigne le vecteur unitaire de la normale à la surface S dirigée vers l'extérieur. $\vec{u}$ désigne le vecteur unitaire de la direc-

tion dans laquelle est calculé le champ $\vec{E}$; $(\vec{H}_1 \times \vec{n})_N$ indique que seule est prise la composante du vecteur $(\vec{H}_1 \times \vec{n})$ perpendiculaire à la direction $\vec{u}$. Les autres termes ont la signification déjà donnée pour la formule (1).

Notons que pour le calcul de $\vec{E}$ seules interviennent les composantes de $\vec{E}_1$ et $\vec{H}_1$ tangentes à la surface S .

Il apparaît clairement que le terme :

$[\vec{u} \times (\vec{E}_1 \times \vec{n}) + (\vec{H}_1 \times \vec{n})_N] dS$ de la formule (11) remplace le terme $(I)_N dl$ de la formule (1). Nous pouvons donc transposer très facilement les résultats du paragraphe précédent sur la directivité des réseaux d'antennes.

Soit une surface rayonnante rectangulaire et plane, de côtés a , b . Si sur cette surface $\vec{E}_1$ et $\vec{H}_1$ sont constants la directivité du diagramme de rayonnement sera donnée par la formule (10) et nous aurons la possibilité en réduisant l'amplitude des champs sur les bords de la surface de rechercher un compromis entre le gain maximum et les lobes secondaires.

Nous devons cependant souligner une différence entre les deux diagrammes de directivité. Le diagramme obtenu plus haut pour une surface d'éléments vibrant tous en phase était symétrique par rapport au plan de la surface. Or dans la formule (11) le terme $\vec{u} \times (\vec{E}_1 \times \vec{n})$ change de sens avec $\vec{u}$ tandis que le terme $(\vec{H}_1 \times \vec{n})_N$ est indépendant de ce sens. En particulier si les composantes tangentielles de $\vec{E}_1$ et $\vec{H}_1$ sont égales, en phase et font entre elles un angle de 90° comme c'est le cas pour une onde plane qui se propage normalement à la surface, le champ sera nul dans la direction opposée à celle correspondant au gain maximum.

Calculons dans une telle hypothèse qui correspond au gain maximum possible de la surface rayonnante, le champ dans la direction optima. La formule (11) donne :

$$(12) \quad E = \frac{S}{\lambda R} E_1$$

S étant la grandeur de la surface rayonnante et R la distance à laquelle le champ est évalué. Or la puissance totale rayonnée est alors égale à

$\frac{c}{4\pi} E_1^2 S$. On en conclut que le gain maximum possible d'une surface plane rayonnante est :

$$(13) \quad g_m = \frac{4\pi S}{\lambda^2}$$

formule qu'on aurait pu tirer directement de la formule (5) en supposant l'antenne réceptrice et captant toute l'énergie incidente sur la surface S . Cette formule traduit la diffraction des ondes.

Si la surface S est un rectangle très allongé tel qu'une de ses dimensions soit inférieure à une longueur d'onde tandis que l'autre dimension vaut plu-

sieurs longueurs d'onde les formules (7) (8) et (10) du paragraphe précédent lui seront applicables.

De même que pour les ondes longues on a imaginé pour obtenir des antennes d'ouvrir vers l'espace les lignes transportant les courants il est naturel en hyperfréquences de créer des surfaces rayonnantes en dilatant progressivement le guide d'onde par où arrive l'énergie électromagnétique. On réalise ainsi des cornets qui rayonnent perpendiculairement à leur ouverture.

Pour avoir un ruban rayonnant auquel puissent être appliquées les formules (7) (8) et (10) il est naturel encore de réaliser un guide qui laisse fuir de l'énergie tout au long de son axe. Une façon d'obtenir ceci est de canaliser les ondes dans un tube (plein ou creux) constitué par un diélectrique. Les dimensions de ce guide diélectrique qui constitue ce qu'on appelle couramment une antenne diélectrique sont choisies de telle sorte que les ondes s'y propagent avec une vitesse légèrement inférieure à la vitesse c . Si le guide diélectrique est plein cette condition signifie que les ondes cheminant à l'intérieur de l'antenne subissent, sur la surface diélectrique-air une réflexion voisine de la réflexion totale. Si le guide diélectrique est creux, l'épaisseur des parois conditionne la fuite d'énergie recherchée.

Un deuxième procédé est de pratiquer des ouvertures rayonnantes tout le long d'un guide. Les phases respectives de ces fentes rayonnantes sont facilement réglées par leur espacement. De telles ouvertures perturbent évidemment la répartition des courants sur les parois du guide. Il est commode pour étudier et réaliser des fentes de se servir de la

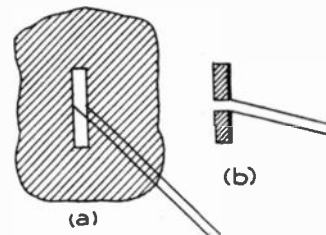


FIG. 14

dualité signalée au paragraphe précédent entre les champs électriques et magnétiques. Cette dualité se traduit en hyperfréquences par des diagrammes de rayonnement conjugués (identiques en permutant champ électrique et champ magnétique) pour une fente et pour l'écran complémentaire, par exemple, comme indique dans la figure 14 la fente (a) entourée d'un plan indéfini, et le dipôle (b). Afin de présenter au guide une charge non réactive et ainsi d'avoir le plus faible coefficient de réflexion ces fentes sont faites résonnantes pour la fréquence utilisée.

Tous les procédés indiqués ci-dessus ne sont qu'une transposition aux hyperfréquences des moyens utilisés aux fréquences ordinaires. Situées à mi-chemin entre ces fréquences-ci et les ondes lumineuses les

hyperfréquences se prêtent aussi à la transposition des procédés de l'optique. C'est ainsi que l'on se sert couramment de projecteurs paraboliques identiques aux projecteurs lumineux et calculés en première approximation comme ceux-ci en faisant intervenir la notion de rayons émis depuis une source supposée ponctuelle. Cependant, contrairement aux projecteurs lumineux où chaque élément de surface représentant une faible partie de la surface totale, a des dimensions géométriques grandes devant la longueur d'onde, la surface totale d'un projecteur pour hyperfréquences conditionne l'ouverture du faisceau rayonné conformément à la formule (13). Cette petitesse relative de la surface du projecteur rend le diagramme de rayonnement très sensible à l'amplitude, aux phases et à la polarisation du champ aux divers points de cette surface. Suffisante pour l'optique géométrique l'approximation qui remplace les ondes par des rayons, ne l'est donc pas tout à fait pour construire des projecteurs hyperfréquences qui réclament ou outre l'analyse des champs. En agissant sur ceux-ci on peut d'ailleurs trouver un compromis entre les feuilles latérales et la directivité comme indiqué ci-dessus pour les cornets.

Une deuxième transposition des procédés de l'optique réside dans l'utilisation de lentilles. Des lames de diélectriques convenablement taillées donnent de bons résultats mais elles présentent le défaut d'être bien lourdes. Or un diélectrique peut être considéré comme un réseau agissant sur les vitesses des phases des ondes. Il était naturel en hyperfréquences de songer à réaliser une lentille par assemblage de guides. Dans ceux-ci la vitesse de phase des ondes est plus grande que la lumière. Ils peuvent donc constituer un milieu dont l'indice est inférieur à l'unité. Ainsi des lames métalliques planes distantes entre elles de la valeur a représentent un milieu dont l'indice est donné par la formule :

$$(14) \quad n = \sqrt{1 - \left(\frac{\lambda}{2a}\right)^2}$$

La figure 15 représente une lentille réalisée avec de telles lames dont les plans sont perpendiculaires au plan de la figure (a doit être compris entre $\frac{\lambda}{2}$ et λ pour que les ondes soient guidées et pour qu'il n'y ait qu'un seul type d'onde). Les mêmes remarques formulées pour les miroirs s'appliquent évidemment aux lentilles : il faut que les champs aient une amplitude et une polarisation convenable sur toute la surface d'attaque de la lentille.

Nous avons vu dans le paragraphe précédent qu'une caractéristique importante de l'antenne est la largeur de bande dans laquelle elle est utilisable. Une variation de la fréquence peut affecter soit l'alimentation de l'antenne, soit la distribution des champs sur la surface rayonnante.

Comme dans les ondes ordinaires l'antenne est alimentée par un feeder que l'on s'efforce de terminer sur son impédance caractéristique afin d'éviter

des ondes stationnaires et une charge réactive à l'émetteur. Les fréquences pour lesquelles la charge présentée par l'antenne au feeder provoque une réflexion importante des ondes vers l'émetteur déterminent la largeur de bande. En hyperfréquences il arrive couramment que le feeder ait une longueur représentant un nombre élevé de longueurs d'ondes.

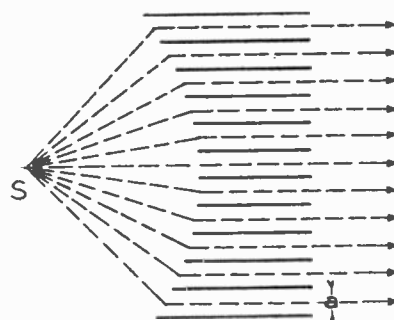


FIG. 15

En ce cas et si le taux d'ondes stationnaires est assez élevé une très faible variation de fréquence entraîne une variation importante de la charge (active ou réactive) de l'émetteur. Un inconvénient grave peut alors surgir : la plupart des émetteurs hyperfréquences notamment les radars sont des auto-oscillateurs ; de la grande sensibilité de la charge à la fréquence peut résulter une instabilité provoquant un changement brutal de l'onde émise, phénomène courant chez de tels émetteurs ; la réception est alors nulle tandis que l'émetteur paraît fonctionner convenablement.

Ces considérations montrent toute la nécessité de réaliser une bonne adaptation du feeder dans une large gamme, la largeur de bande étant conditionnée davantage par un taux maximum d'ondes stationnaires compatible avec la longueur du feeder plutôt que par la considération unique de la composante réactive de la charge présentée par l'antenne du feeder. Dans toute la gamme désirée cette charge doit s'écarter très peu de l'impédance caractéristique du feeder. Les cornets et antennes diélectriques réalisent généralement une bonne adaptation au feeder. Si la charge du feeder est un petit doublet rayonnant vers un projecteur on améliore la largeur de bande par les procédés utilisés dans les ondes plus longues. Si l'organe rayonnant est une ouverture dans un guide on lui donne une force convenable ou on la munit de flasques soit analogues à celles d'un cornet, soit constituant un volume résonnant à large bande (fig. 16) Lorsque l'antenne élémentaire rayonnant vers le projecteur reçoit une partie des ondes réfléchies la charge présentée par cette antenne au feeder n'est pas la même que si l'antenne était isolée dans l'espace. Or le trajet aller et retour des ondes étant encore de plusieurs longueurs d'ondes une faible variation de fréquence modifie considérablement la charge. Afin d'éviter cet inconvénient on fait en sorte que l'influence de l'onde réfléchie soit très faible sur la charge. Un moyen radical d'y arriver est évidemment de placer le foyer rayonnant en dehors des ondes réfléchies.

Lorsque la variation de fréquence affecte la distribution des champs sur la surface émissive le résultat pratique en est soit une diminution du gain et l'apparition de lobes secondaires plus marqués, soit un changement dans la direction de l'émission. Dans ces deux éventualités la communication

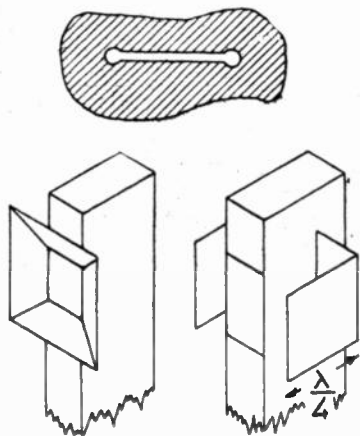


Fig. 16

devient plus mauvaise. Une telle perturbation provient du fait que les longueurs radioélectriques des chemins parcourus dans l'antenne par les divers rayons de l'onde ne conservent pas les valeurs relatives fixées pour la fréquence fondamentale. Il en est ainsi par exemple dans les lentilles et dans les guides rayonnant par des fentes. Il est naturel d'atténuer ce défaut de la même façon que les lentilles optiques sont rendues achromatiques : un double trajet est imposé aux rayons ; pour chaque rayon une variation de fréquence augmente la longueur d'un trajet et diminue celle de l'autre de telle sorte que la longueur radioélectrique totale reste constante.

III. — Réalisation et performances des antennes hyperfréquences.

Il convient de rappeler que si la compréhension des théories est nécessaire elle n'est nullement suffisante pour réaliser de bonnes antennes hyperfréquences. L'expérimentation est indispensable. Pendant la guerre les anglais et les américains ont multiplié leurs expériences sur les antennes. Afin d'opérer rapidement ils réalisèrent des modèles en bois revêtus d'une peinture conductrice ou de feuilles métalliques.

Les cornets furent beaucoup employés dans les débuts des hyperfréquences. Pour que le rayonnement soit convenable il faut que les ondes soient sensiblement en phase dans l'embouchure. On admet que le guide crée des ondes sphériques et qu'il suffit pour que ces ondes rayonnent comme une onde plane dans l'embouchure que la flèche de la calotte sphérique limitée par les bords du cornet soit inférieure à $\frac{\lambda}{2}$. Ceci entraîne la relation :

$$(15) \quad l \geq \frac{a^2}{2\lambda}$$

a étant la plus grande dimension transversale et l la longueur du cornet (fig. 17). Le gain théorique donné par la formule (13) s'écrit en pratique pour toutes les surfaces rayonnant perpendiculairement à leur plan (cornets, miroirs, lentilles) :

$$(16) \quad g = A \frac{S}{\lambda^2}$$

A étant un coefficient généralement compris entre 3 et 10 et qualifiant la réalisation de l'antenne

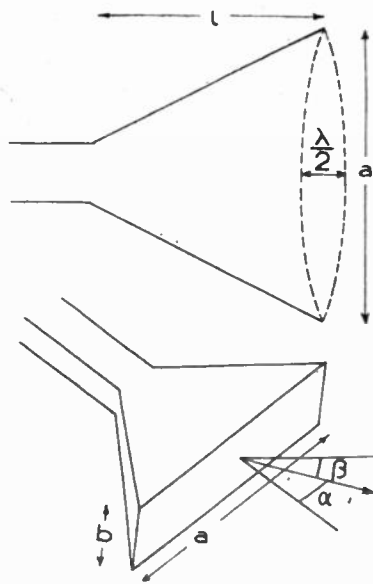


Fig. 17

En première approximation les demi-angles d'ouverture du pinceau rayonné valent $\alpha = \frac{\lambda}{a} \beta = \frac{\lambda}{b}$

En pratique on mesure généralement non l'angle 2α entre les deux directions, difficiles à estimer, où le rayonnement devient nul, mais l'angle $2\alpha_1$ entre les deux directions où la puissance rayonnée est la moitié de la puissance rayonnée dans la di-

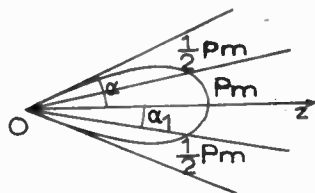


Fig. 18

rection optima. (En ces directions le gain est moitié du gain maximum et l'amplitude du champ vaut 0,7 fois le champ dans la direction optima (fig. 18)

Cet angle α_1 , est sensiblement égal à $\frac{\alpha}{2}$ dans le cas du diagramme à gain maximum ainsi que l'indique la formule (8).

Permettant à l'onde issue d'un guide de s'épanouir progressivement les cornets réalisent une bonne adaptation au guide et présentent l'avantage d'une uti-

lisation dans une large bande de fréquences mais ainsi qu'il ressort de la formule (15) ils sont fort encombrants.

Les antennes diélectriques ont un rayonnement longitudinal correspondant à la formule (10) et la (fig. 12). Pour que ce rayonnement existe la formule (11) indique qu'il est nécessaire que l'onde cheminant dans l'antenne ait des composantes tangentielles du champ électromagnétique dirigées suivant cette

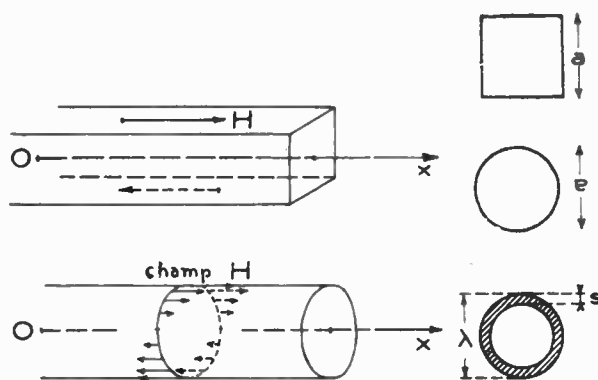


Fig. 19

longueur. (fig. 19) Le fait que le champ rayonné est polarisé suffit à indiquer que l'onde ne peut être une onde ayant une symétrie de révolution par rapport à l'axe longitudinal. En pratique, cette onde est le type d'onde normal dans les guides rectangulaires, ce qui présente l'avantage de n'avoir dans

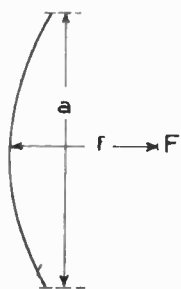


Fig. 20

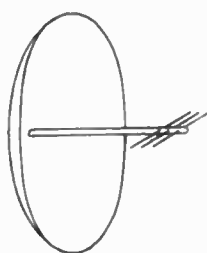


Fig. 21

l'antenne aucun type d'onde parasite, les dimensions étant telles que les autres types d'onde se trouvent éliminés par leur fréquence de coupure. Les champs tangentiels longitudinaux sont des champs magnétiques disposés comme l'indique la (fig. 19 a). La longueur d'onde λg dans un guide diélectrique plein est donné par la formule :

$$(17) \quad \lambda g = \frac{\lambda}{\sqrt{\epsilon - \left(\frac{\lambda}{\lambda c}\right)^2}}$$

ϵ étant la constante diélectrique et λc la longueur d'onde de coupure. Faire cheminer les ondes à une vitesse voisine de c signifie que $\lambda g = \lambda$ d'où la

$$\text{formule : (18)} \quad \lambda c = \frac{\lambda}{\sqrt{\epsilon - 1}}$$

Si l'antenne est rectangulaire de grand côté a on a : $\lambda c = 2a$ (fig. 19 b). Si l'antenne est cylindrique de

diamètre a on a : $\lambda c = \frac{\pi a}{1,84}$ (fig. 19 c)

d'où approximativement le diamètre suivant pour les antennes diélectriques :

$$a : \square \text{ antenne rectangulaire } a \approx \frac{\lambda}{2} \frac{1}{\sqrt{\epsilon - 1}} \quad (19)$$

$$a : \bigcirc \text{ antenne cylindrique } a \approx \frac{\lambda}{1,7} \frac{1}{\sqrt{\epsilon - 1}} \quad (19)$$

Ces diamètres correspondent à une réflexion totale des ondes guidées sur la surface air-diélectrique. Aussi sont-ils pratiquement un maximum et il convient pour que l'énergie diffusée reste constante d'amincir l'antenne dans le sens de la propagation. Les antennes diélectriques cylindriques sont généralement tronconiques et d'un diamètre variant entre $0,6 \frac{\lambda}{\sqrt{\epsilon - 1}}$ et $0,4 \frac{\lambda}{\sqrt{\epsilon - 1}}$

Si le guide diélectrique est creux l'expérience a montré qu'il convenait d'adopter un diamètre voisin de la longueur d'onde et une épaisseur voisine

de $s = \frac{\lambda}{20} \frac{1}{\sqrt{\epsilon - 1}}$ (fig. 19 d) Le pinceau rayonné par

une antenne diélectrique est plus étroit que celui donné par la formule (10) où l'on fait $q = 1$ parce que d'une part les ondes cheminent à une vitesse légèrement inférieure à celle de la lumière et que d'autre part les surfaces supérieures et inférieures du guide sont distantes d'une valeur non négligeable devant la longueur d'onde de telle sorte qu'il en résulte un effet directif. On peut adopter en première approximation la formule suivante pour le demi-angle d'ouverture.

$$(20) \quad \alpha = \sqrt{\frac{\lambda}{2l}}$$

et pour le gain

$$(21) \quad g = \Lambda \frac{l}{\lambda}$$

Ces antennes présentent les avantages d'un faible encombrement et d'une largeur de bande assez grande mais leurs diagrammes de rayonnement ont des lobes secondaires très marqués. Leur directivité est améliorée par leur groupement en réseaux.

Les projecteurs sont constitués par les surfaces de paraboloïdes ou de cylindres paraboliques. Leur distance focale f est généralement comprise entre $0,25 a$ et $0,6 a$, a étant l'ouverture du miroir. Leur surface est constituée par une surface pleine ou un grillage métallique dont les ouvertures ont des dimensions inférieures à $\frac{\lambda}{2}$ dans le sens perpendiculaire au champ électrique. Il convient qu'il n'y ait pas entre les rayons provenant de la source des différences de marche supérieures à $\frac{\lambda}{8}$ ce qui

revient à exiger une précision de $\frac{\lambda}{16} \left(\pm \frac{\lambda}{32} \right)$ dans

la construction de la surface du miroir. Pour diminuer les lobes secondaires le foyer rayonnant est tel que l'amplitude du champ sur les bords du projecteur est environ le dixième de l'amplitude au centre (fig. 20).

Les projecteurs ont de bons diagrammes de rayonnement tout en présentant un encombrement nettement inférieur à celui des cornets. Nous avons vu que leur largeur de bande ne dépend pas d'eux mais de leur alimentation.

Lorsque le projecteur est un paraboloïde l'alimentation est fournie par une source ponctuelle. Celle-ci peut être un doublet accompagné d'une surface réfléchissante ou de plusieurs autres doublets obligeant la source à ne rayonner que vers le miroir. (fig. 21).

Plus généralement cette source est l'ouverture d'un guide qui présente sur un dipôle l'avantage de mieux définir la polarisation du champ. Si la source est placée au milieu des ondes réfléchies on peut éviter que le feeder détériore le diagramme

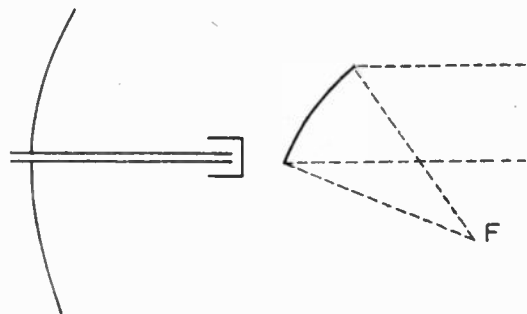


Fig. 22

Fig. 23

de rayonnement en le faisant arriver à la source au travers du paraboloïde parallèlement aux rayons réfléchis comme indiqué sur la fig. 21 pour un coaxial ou la fig. 22 pour un guide.

Si du fait des ondes qui retournent sur la source la largeur de bande se trouve trop étroite on utilise une surface de paraboloïde telle que les ondes ne rencontrent plus le foyer (fig. 23) mais cette

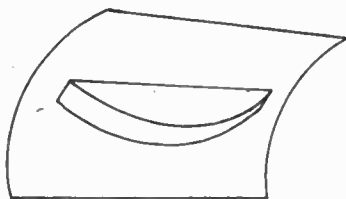


Fig. 24

solution augmente l'encombrement de l'antenne. Pour éviter les sauts d'onde de l'émetteur on admet que pour un feeder d'une longueur de $60 \lambda_g$, alimenté par un magnétron, le taux d'ondes stationnaires $\left(\frac{\text{champ maximum}}{\text{champ minimum}} \right)$ dans le guide ne doit pas dépasser 1,5. La largeur de bande ainsi définie est plus étroite que celle définie comme ci-dessus

par la considération de la puissance d'émission, d'où l'avantage d'utiliser des feeders courts.

Lorsque le miroir est un cylindre parabolique destiné à créer un pinceau étroit, il faut l'alimenter par une ligne qui peut elle-même être la surface rayonnante d'un autre élément de cylindre parabolique (fig. 24). Dans les radars anglais d'aviation cette ligne est constituée par des fentes résonantes

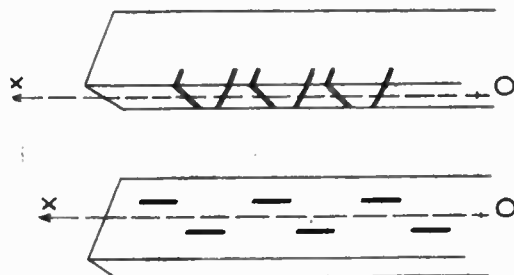


Fig. 25 et 26

découpées dans un guide. Elles se situent soit sur le grand côté, soit sur le petit côté. Dans le premier cas (fig. 25) elles sont longitudinales, leur couplage est d'autant plus élevé qu'elles sont écartées de la ligne médiane Ox et leur phase varie de 180° suivant le côté de cet axe où elles se situent. Dans le deuxième cas (fig. 26) elles sont transversales (pour être résonnantes il faut alors qu'elles débordent sur le grand côté), leur couplage est d'autant plus élevé qu'elles se rapprochent de l'axe Ox et leur phase varie de 180° suivant le sens de leur inclinaison. Le champ rayonné par une de ces lignes est évidemment polarisé à 90° du champ rayonné par l'autre.

Pour que l'énergie rayonnée par la ligne soit constante sur sa longueur il faut augmenter le couplage des fentes à mesure que l'on se rapproche de l'extrémité du guide. On ne peut augmenter trop ce couplage sans provoquer des réflexions ; aussi l'é-

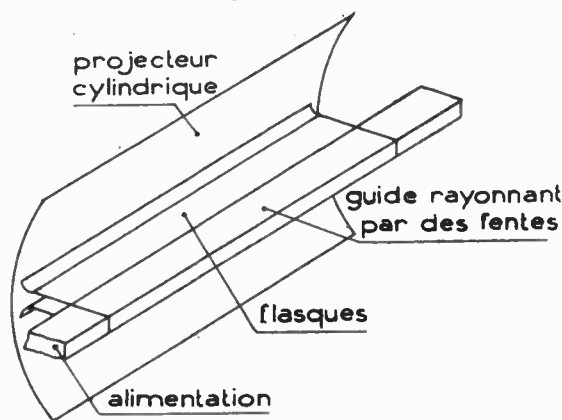


Fig. 27

nergie disponible à l'extrémité du guide n'est pas nulle mais environ le dixième de l'énergie totale rayonnée. Elle est absorbée dans une charge dissipative.

Pour avoir un faisceau rayonné perpendiculaire à la ligne il faut que toutes les fentes rayonnent

en phase ce qui conduit avec la disposition alternée des figures 25 et 26 à les espacer de $\frac{\lambda g}{2}$. Un tel es-

pacement a le grave défaut d'ajouter les réflexions inévitables que produisent toutes les fentes. Aussi en pratique on préfère mettre entre elles une différence de phase de 200° au lieu de 180° . Ceci incline d'environ 3° par rapport à la normale le faisceau émis par la ligne et le rend conique au lieu de cylindrique. Il faut donc incliner la ligne par rapport aux génératrices du miroir. On la munit en outre de flasques planes destinées à rendre le faisceau cylindrique et de flasques évasées assurant une distribution correcte de l'amplitude du champ (fig. 27).

Ces flasques peuvent être rendues telles que le cheminement des ondes y soit d'une sensibilité à la fréquence corrigeant la perturbation qui résulte de la variation avec la fréquence de la longueur d'onde dans le guide λg . La largeur de bande est ainsi augmentée par un procédé rappelant ceux de l'optique.

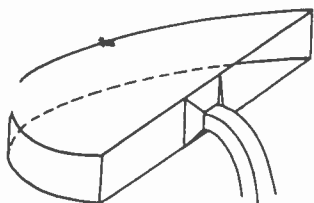


Fig. 28

Les miroirs cylindriques sont particulièrement employés lorsque le diagramme de rayonnement désiré est fort différent dans deux plans perpendiculaires. Lorsque dans un de ces plans on ne désire qu'une faible directivité l'aérien peut alors être constitué d'un élément de cylindre parabolique compris entre deux plans. Une telle antenne est appelée antenne « cheese » par les Anglais en raison de sa forme (fig. 28).

Les lentilles constituées par des guides d'onde sont de forme concave puisque l'indice est inférieur

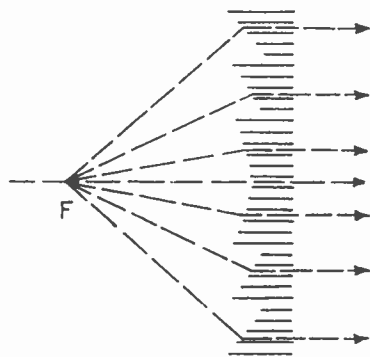


Fig. 29

à l'unité. Le choix de l'indice résulte d'un compromis. n voisin de 1 signifie des lentilles trop épaisses, n trop faible engendre des réflexions des ondes sur la surface de la lentille et impose de sévères

conditions pour la précision de réalisation. En pratique on adopte $n = 0,5$ ($a = 0,58 \lambda$) ou $n = 0,6$ ($a = 0,63 \lambda$). Pour réduire l'épaisseur de la lentille, comme l'a déjà fait Fresnel pour les lentilles de phare, la surface concave est une surface à échelons (fig. 29) où chaque discontinuité correspond à une différence de marche de λ pour les rayons. En ce cas aucune épaisseur de la lentille n'est plus

grande que $\frac{\lambda}{1-n}$. On réalise aussi des lentilles à épaisseur constante mais à lames d'épaisseur variable.

La distance focale habituellement employée vaut 0,5 à 1 fois la plus grande dimension d'ouverture. Une lentille est donc beaucoup moins encombrante qu'un cornet. Avec $n = 0,5$ il faut que l'espace-

ment des plaques soit constant à $\pm \frac{\lambda}{75}$ (pour $n = 0,6$, $\pm \frac{\lambda}{50}$). l'épaisseur de la lentille est moins

critique et doit être tenue à $\pm \frac{\lambda}{16(1-n)}$ pour

maintenir la phase à $\pm \frac{\lambda}{16}$. Ces tolérances sont

moins sévères que celles imposées aux miroirs car il est plus facile de respecter les conditions qui s'adressent à des parties métalliques étroitement

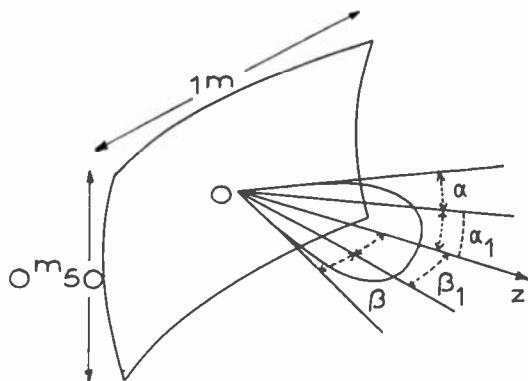


Fig. 30

solidaires et relativement petites que celles qui concernent une grande surface de par sa nature plus exposée à des défauts résultant de son utilisation.

Les formules ci-dessus permettent de chiffrer approximativement les performances des antennes hyperfréquences. Les projecteurs de radars de bord fonctionnant sur $\lambda = 10$ cm et ayant environ 1 m sur 0 m 50 ont un gain moyen de 300 (25 décibels) et des faisceaux ayant des angles d'ouverture.

$$\begin{array}{ll} 2 \alpha = 24^\circ & 2 \beta = 12^\circ \text{ (figure 30).} \\ 2 \alpha_1 = 12^\circ & 2 \beta_1 = 6^\circ \end{array}$$

Les antennes de radar sur $\lambda = 3$ cm pour conduite de tir anti-aérien ont des faisceaux ayant des angles d'ouverture de l'ordre du degré et des gains de plusieurs milliers (environ 35 décibels).

IV — Utilisation.

Les appareils radars de divers types fonctionnant sur des longueurs d'onde de 10 cm, 3 cm, 1 cm représentent actuellement la principale application des hyperfréquences. Rappelons que l'équation du radar s'écrit :

$$(22) \quad \frac{W_2}{W_1} = g_1 g_2 \frac{\lambda^2}{(4\pi)^2} \frac{\Sigma}{4\pi} \frac{1}{R^4}$$

où les termes ont la signification donnée pour la formule (5) et où Σ représente la surface équivalente de l'objectif qui capte une énergie $P\Sigma$ (P étant le vecteur de Poynting au lieu où se trouve l'objectif) et la diffuse dans tout l'espace. Pour un avion cette surface est de quelques mètres carrés.

Les antennes les plus simples sont celles des appareils radars de veille auxquels on demande une bonne précision azimutale et une ouverture en site suffisante pour détecter à la fois les bâtiments et les avions. De tels appareils ont des antennes genre « cheese » ou des miroirs paraboliques dont la dimension horizontale est environ le double de la dimension verticale.

Pratiquement les avions ont un plafond et il est

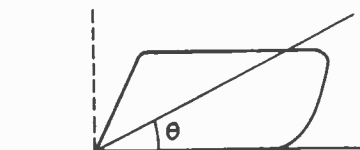


Fig. 31

inutile d'envoyer à des sites élevés une énergie telle qu'elle permettrait de détecter des appareils volant à 50 ou 100 kilomètres d'altitude. Il est préférable pour un radar de veille de diminuer l'énergie rayonnée aux sites élevés et de renforcer celle émise dans les sites faibles de telle sorte que les avions volant à la même altitude soient détectés pareillement quelle que soit leur distance. Ceci exige un diagramme théorique dans le plan vertical

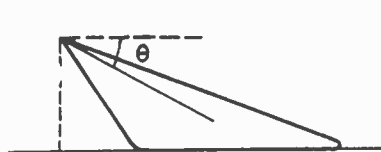


Fig. 32



Fig. 33

tel que l'amplitude soit donnée par la formule :

$$(23) \quad E(\theta) = \frac{K}{\sin \theta} = K \operatorname{cosec} \theta$$

θ étant l'angle de la direction suivant laquelle est estimée E , avec le plan horizontal. Avec un tel diagramme le gain de l'antenne pour le site θ est proportionnel à $\frac{1}{\sin^2 \theta}$ et la formule (22) indique que la puissance reçue par le récepteur ne dépend que de l'altitude de l'objectif.

De telles antennes appelées par les anglo-saxons

« cosecant antennas » ont été réalisées. Elles servent soit à la veille soit à la télévision radar du sol. Dans ce cas-ci il convient que les points du sol donnent le même écho quelle que soit leur distance. Leur diagramme vérifie la formule (23) entre deux angles θ_1 et θ_2 . Pour un radar de veille il est semblable à celui de la figure 31, pour un radar de télévision à celui de la figure 32.

Un projecteur parabolique ayant une courbure spéciale dans le plan vertical peut donner un diagramme cosecant mais les meilleurs résultats ont été obtenus par des miroirs cylindriques ayant un tracé spécial (fig. 33) alimenté par une ligne. Les

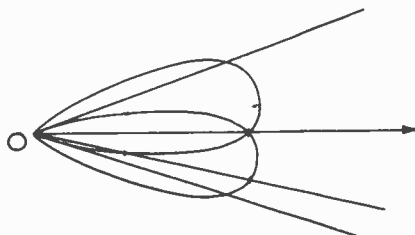


Fig. 34

deux définitions verticale et horizontale étant bien séparées ils ne présentent pas les défauts de polarisation inévitables avec un parabolicoïde.

Les appareils de conduite de tir doivent donner une précision en direction de l'ordre de quelques minutes. On ne peut évidemment songer à créer des faisceaux de cet ordre. La précision nécessaire est obtenue par le procédé du double lobe, procédé depuis longtemps en usage dans la marine dans les chenaux optiques REP. L'amplitude d'un écho varie rapidement avec la direction de l'antenne lorsque cette direction est voisine d'un des côtés du lobe du diagramme polaire. Soient deux antennes dont les directions principales font entre elles un angle approximativement égal à l'angle $2\alpha_1$ (fig. 34). La précision de la mesure de la direction pour laquelle les deux antennes donnent des échos de même amplitude est autrement plus grande que la précision avec laquelle on apprécie sur une seule antenne le maximum d'écho. En pratique avec deux lobes de quelques degrés on arrive à une précision suffisante pour la conduite de tir. Pour une direction de tir contre avions il est indispensable d'avoir deux lobes dans un plan horizontal et deux lobes dans un plan vertical. Outre la précision dans la direction ce procédé fournit les éléments nécessaires pour une conduite automatique du tir puisque les grandeurs relatives des deux échos indiquent quelle est la direction de l'antenne par rapport à celle du but.

Dans certains radars, particulièrement les anciens radars sur ondes métriques il y a autant d'antennes que de lobes désirés. En hyperfréquences les lobes sont généralement obtenus non simultanément mais successivement avec une rapidité telle que l'œil ne s'aperçoive pas de cette différence et c'est la même antenne qui crée tous les lobes. De telles antennes sont généralement constituées par des projecteurs. Les lobes sont obtenus soit par un déplacement

de la source rayonnante dans le plan focal, soit par l'alimentation successive de plusieurs sources espacées dans le plan focal, soit encore lorsque le miroir est cylindrique et alimenté par une ligne, par une variation des phases des éléments radiants de cette ligne. S'il est inutile grâce au procédé du double lobe d'avoir des faisceaux très étroits pour obtenir une précision suffisante pour le tir, notons cependant qu'une grande directivité de chaque lobe permet seule de différencier les objectifs situés dans des directions voisines.

Certains radars de conduite de tir ou de bombardement sont conçus pour rechercher l'objectif. Il est alors nécessaire qu'ils puissent balayer continuellement une certaine région de l'espace et avec une rapidité telle que la distance de l'objectif n'ait sensiblement pas varié entre deux balayages successifs. Des résultats convenables n'ont pu, jusqu'ici, être obtenus en déplaçant devant le projecteur la source rayonnante, comme il est fait pour avoir des lobes voisins ; lorsqu'on écarte notablement la source du foyer le diagramme de rayonnement devient très mauvais. Par contre, il existe

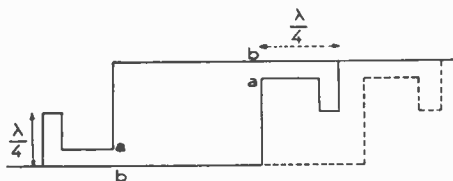


Fig. 35

Les replis en $\frac{\lambda}{4}$ font que le contact entre *a* et *b* est électriquement parfait

des appareils où le balayage est obtenu en faisant varier la phase des éléments radiants d'une ligne. Un procédé employé par les Américains consiste à déplacer l'un par rapport à l'autre les deux petits côtés d'un guide, comme indiqué dans la fig. 35.

La phase des éléments radiants peut encore être changée en agissant sur chacun d'eux individuellement comme il est fait en ondes métriques. Pour que le système ne soit pas trop complexe il convient alors que ces éléments soient en nombre limité. Or, d'une part la longueur de la base détermine l'angle du faisceau, et un espacement des éléments radiants supérieurs à une longueur d'onde crée des faisceaux parasites. Le problème est résolu par l'emploi d'éléments radiants ayant chacun un effet directif marqué, à savoir des antennes diélectriques.

Mais le procédé le plus simple et actuellement le plus employé consiste en un balancement mécanique de l'antenne. Pour un balayage alternatif rapide une telle opération nécessite un équilibrage dynamique de l'ensemble particulièrement soigné. Cependant vu les difficultés rencontrées pour l'obtention de bons balayages par les procédés ci-dessus indiqués la solution entièrement mécanique est à adopter chaque fois qu'elle se révèle possible.

En France la Compagnie Générale de Télégraphie sans fil a essayé de résoudre le problème de balayage rapide pour une antenne de veille sur tout l'horizon. Une telle solution rendrait inutile l'orientation gyroscopique de l'écran exigée par la ré-

manence, elle-même nécessaire pour un balayage lent, d'où son emploi possible pour les bâtiments non munis de gyroscopes. Le dispositif réalisé (fig. 36 et 36 bis) consiste en un réseau fixe de cornets disposés circulairement. Ces cornets sont alimentés par un petit cornet C dont la faible inertie mécanique permet de lui inculquer une grande vitesse de rotation (de l'ordre de 10 tours par seconde). Tous les problèmes évoqués ci-dessus montrent combien la réalisation des antennes hyperfréquences, comme d'ailleurs celle des guides ou tubes hyperfréquences, pose des problèmes mécaniques qui exigent des constructeurs une grande compétence, une grande habileté et aussi des machines-outils de précision.

L'utilisation des antennes hyperfréquences nécessite parfois (notamment dans l'aviation) que ces antennes soient entourées d'un dôme. Ces dômes ne doivent ni absorber ni réfléchir les ondes émises. La réflexion est encore plus gênante que l'absorption car elle peut provoquer un taux d'ondes stationnaires propice aux sauts d'onde de l'émetteur. Pour les ondes qui atteignent normalement sa surface les réflexions peuvent être atténuées si le dôme

a une épaisseur égale à $\frac{\lambda}{2}$, les ondes réfléchies sur

la surface extérieure annulant les ondes réfléchies sur la surface intérieure. Un autre procédé consiste à rechercher des matériaux présentant aux ondes une impédance voisine de celle de l'air. Pour les avions de très grande vitesse, l'emploi des fentes rayonnantes paraît être la solution désirable.

La principale utilisation des hyperfréquences autre que les divers appareils radar (veille, navigation, conduite de tir...) consiste dans les communications (téléphonie, télégraphie), par câbles hertziens. Ainsi que nous l'avons vu les ondes peuvent être canalisées dans un faisceau spatial très étroit et la même

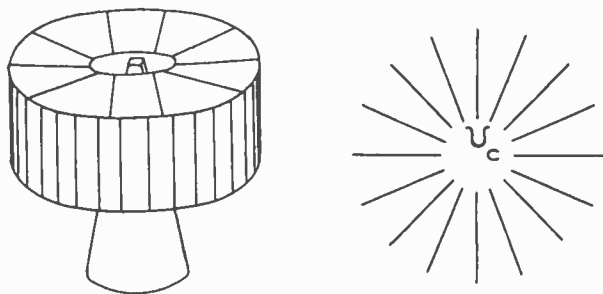


Fig. 36 et 36 bis

directivité s'exprime à la réception comme à l'émission. Contrairement aux émissions radioélectriques ordinaires qui encombrant l'espace de telle sorte qu'il est nécessaire de prévoir des attributions de fréquences aux divers utilisateurs, les émissions par hyperfréquences peuvent se faire beaucoup plus discrètement sans crainte de gêner ni d'être gênés par les autres émissions analogues d'où la dénomination justifiée de câbles hertziens par comparaison avec les câbles ordinaires qui canalisent tout près d'eux les ondes radioélectriques des transmissions filaires. Soulignons que pour de telles

communications il est avantageux d'avoir des faisceaux étroits non seulement pour la direction et le gain mais encore par la possibilité d'éviter ainsi des réflexions des ondes sur le sol ; ceci évite les interférences entre l'onde directe et l'onde réfléchie d'où résulte parfois une baisse de la qualité de la communication. C'est pour l'utilisation en câbles hertziens que les Américains ont mis au point des lentilles ayant un gain de l'ordre de 10.000 (40 décibels) et un faisceau d'ouverture de 2°. Il est à prévoir dans les années à venir un grand développement de ces câbles hertziens qui sont plus économiques à établir que les câbles filaires.

BIBLIOGRAPHIE

DAVID. — Cours de radiotechnique de l'Ecole des Officiers de transmissions. — Chapitre II.

RIGAL. — Cours de radioélectricité générale de l'école nationale supérieure des télécommunications.

L. DE BROGLIE. — Problèmes de propagations guidées des ondes électromagnétiques. — Gauthier-Villars.

J. A. RATCLIFFE. — Aerials for radar equipment.
Radiolocation Convention 26-29 Mars 1946.
The Institution of electrical engineers.

Les antennes diélectriques. — *Onde Electrique*. — Octobre 1946.

H. GUTTON. — Les projecteurs d'ondes centimétriques. — *Onde Electrique*. — Décembre 1946.

G. GOUDET — Dispositifs rayonnants pour ondes centimétriques.
Annales des Télécommunications - Mars 1948.

H. T. FRUS. — W. D. LEWIS : Radar antennas. — *The Bell System Technical Journal*. — Avril 1947.

G. GOUDET. — Une formule de rayonnement électromagnétique.
Onde Electrique. — Août 1947 p. 313.

AMPLIFICATEURS A CIRCUITS ANTIRÉSONNANTS A ACCORDS DÉCALÉS

PAR

L. J. LIBOIS

Ingénieur des P. T. T.

Il est maintenant courant d'utiliser pour la réalisation des amplificateurs moyenne et haute fréquence la technique dite des « circuits décalés ». On peut être amené à déterminer les caractéristiques de l'amplificateur de façon que son temps de transmission soit dans la bande passante aussi constant que possible ; mais le plus souvent on cherche à obtenir, avec un minimum de tubes électroniques, une amplification donnée et dans la bande passante la courbe de réponse en amplitude la plus plate possible. Ce sont les formules relatives à ce dernier cas que nous nous proposons d'examiner ici. Nous voudrions montrer en particulier que ces formules sont très faciles à obtenir mathématiquement et surtout qu'elles sont d'un emploi très pratique : elles permettent, sans abaques ou calculs plus ou moins longs, de déterminer immédiatement tous les éléments de l'amplificateur quelque soit le nombre d'étages nécessaire.

La technique des amplificateurs moyenne et haute fréquences à circuits antirésonnants décalés (stagger-tuned amplifiers) a déjà fait l'objet de plusieurs études théoriques et a donné lieu à de nombreuses réalisations pratiques. Ces réalisations concernent le plus souvent des amplificateurs moyenne fréquence constitués par des groupes de circuits décalés (groupes de deux, trois, quatre circuits, etc...) l'amplificateur comprenant plusieurs groupes à la suite. Un groupe souvent utilisé est celui de trois circuits. Il n'est pas plus difficile d'ailleurs de calculer l'amplificateur avec un seul groupe de p circuits tous décalés les uns par rapport aux autres.

Le problème général à résoudre est le suivant : obtenir pour l'amplificateur le gain désiré et choisir les différents paramètres dont on dispose de manière à ce que la courbe de réponse de cet amplificateur satisfasse à un certain nombre de conditions fixées à l'avance.

Quel genre de courbe de réponse (amplitude et phase) va-t-on chercher à réaliser ? Cela dépend évidemment des conditions de transmission que

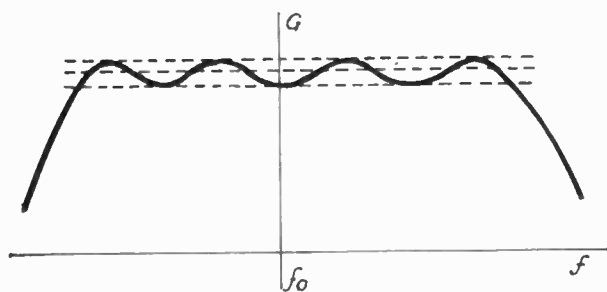


Fig. 1

l'on s'impose. On peut par exemple ne pas avoir à tenir compte de la courbe de phase et chercher à obtenir dans une bande de fréquences donnée un gain aussi élevé que possible, ce gain pouvant osciller autour de sa valeur moyenne dans des limites que l'on se fixera. La courbe de réponse en amplitude aura par exemple l'allure de celle représentée fig. 1.

Ce cas a été étudié théoriquement par Baum pour un amplificateur à pentodes (*Journal of applied*

physics, juin septembre 1946) Baum a montré que l'on pouvait manipuler mathématiquement une telle courbe de réponse en faisant appel aux polynômes de Tchebycheff et arriver ainsi à calculer complètement les paramètres des circuits.

Le cas que nous nous proposons d'examiner et qui a aussi fait l'objet d'études de Baum est plus

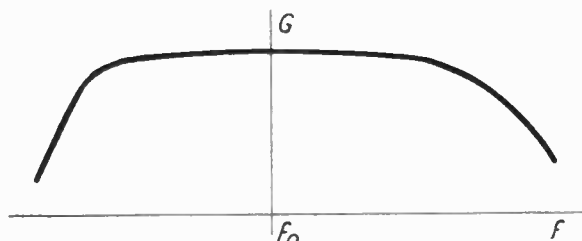


Fig. 2

simple mais n'en présente pas moins un grand intérêt pratique : c'est celui des amplificateurs à gain plat, c'est-à-dire dont la courbe de réponse en amplitude est dans une bande passante donnée aussi constante que possible et cela sans oscillations (fig. 2). Il est bien entendu que dans ce cas la caractéristique du temps de transmission ne sera pas dans cette même bande « la plus plate possible ».

Nous voudrions montrer sur cet exemple important comment on peut très simplement aboutir aux formules exactes qui permettent de calculer tous les éléments des circuits, et comment on peut s'en servir pratiquement. Nous rappellerons également quelques relations générales notamment celle qui relie le gain à la bande passante. Nous indiquerons aussi l'allure que présente dans le cas étudié la courbe du temps de transmission en fonction de la fréquence.

I. — Forme mathématique de la courbe de réponse cherchée.

Soit un amplificateur à pentodes de p étages : nous admettons que les pentodes sont parfaites, c'est-à-dire que leur résistance interne est infinie et que leur capacité parasite grille-plaque est nulle. En réalité comme nous le verrons plus loin, cette capacité ne peut pas toujours être négligée, surtout

si les fréquences utilisées sont élevées et les surtensions des circuits importantes.

Dans l'hypothèse d'une pentode idéale de pente S l'amplification a pour expression :

$$A = S |Z|$$

L'impédance Z sera celle d'un circuit antirésonnant L, C, R (fig. 3).

$$\frac{1}{Z} = \frac{1}{L\omega j} + C\omega j + \frac{1}{R}$$

$$Z = \frac{R}{1 + R \left(\frac{1}{L\omega j} + C\omega j \right)}$$

$$Z = \frac{R}{1 + jR \sqrt{\frac{C}{L}} \left(\omega \sqrt{LC} - \frac{1}{\omega \sqrt{LC}} \right)}$$

d'où finalement :

$$Z = \frac{R}{1 + jQ \left(\frac{f}{f_0} - \frac{f_0}{f} \right)}$$

— Q : surtension effective du circuit (l'amortissement est produit par la résistance mise en parallèle

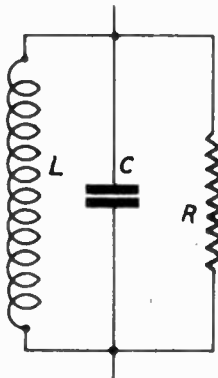


Fig. 3

sur le circuit L, C , les pertes de la self, la résistance d'entrée du tube électronique etc...)

— f_0 : fréquence de résonance du circuit.

L'amplification est donc donnée par :

$$A = \frac{S R}{\sqrt{1 + Q^2 \left(\frac{f}{f_0} - \frac{f_0}{f} \right)^2}}$$

Sur cette expression on remarque aisément que l'amplification est la même pour deux valeurs de f symétriques géométriquement par rapport à f_0 (c'est-à-dire dont le produit est égal à f_0^2).

L'amplification des p étages sera donnée par la formule :

$$A = \prod_{n=1}^{n=p} \frac{S_n R_n}{\sqrt{1 + Q_n^2 \left(\frac{f}{f_{0n}} - \frac{f_{0n}}{f} \right)^2}}$$

Soient alors f_1 et f_2 les fréquences limites qui définiront la bande passante, il est naturel de choisir

comme fréquence centrale de référence f_0 la moyenne géométrique de f_1 et de f_2 .

$$f_0^2 = f_1 f_2$$

Si l'amplificateur était composé d'un seul circuit c'est cette fréquence qui s'introduirait immédiatement comme nous venons de le voir, et l'on utiliserait comme paramètre $\left(\frac{f}{f_0} - \frac{f_0}{f} \right)$.

Essayons de généraliser au cas d'un amplificateur à plusieurs étages. Posons donc :

$$\left\{ \begin{array}{l} x = \frac{f}{f_0} - \frac{f_0}{f} \\ x_n = \frac{f}{f_{0n}} - \frac{f_{0n}}{f} \end{array} \right.$$

Pour simplifier les écritures, considérons d'autre part, au lieu de l'expression complète de l'amplification, seulement l'expression sous le radical puisque c'est elle qui caractérise la courbe de réponse de l'amplification.

$$P = \prod_{n=1}^{n=p} \left(1 + Q_n^2 x_n^2 \right)$$

Il est naturel d'envisager alors un ensemble de circuits groupés en symétrie géométrique autour de la fréquence centrale f_0 , les deux circuits d'une même paire ayant même coefficient de surtension.

Dans ces conditions P est fonction des $(x_n^2 + x_m^2)$ et des $(x_n^2 x_m^2)$ expressions qui, on s'en rend compte aisément, s'expriment uniquement en fonction de la variable x . P est donc le produit de facteurs de la forme $\alpha x^4 + \beta x^2 + \gamma$ correspondant à chaque groupe de circuits symétriques par rapport à f_0 . Finalement, et c'est là le résultat intéressant, P est un polynôme en x de degré $2p$ c'est-à-dire une fonction mathématique qui pourra être manipulée assez simplement.

Comment obtenir alors, d'une part dans la bande passante $f_1 f_2$, que nous supposons, pour simplifier, définie à 3 db d'affaiblissement, une courbe de réponse aussi plate que possible, et en dehors de la bande passante un affaiblissement croissant rapidement ? Il suffit d'identifier $P(x)$ à un polynôme $Q(x)$ de même degré et qui réponde aux exigences précédentes. Ce polynôme $Q(x)$ a une expression simple :

$$Q(x) = K^2 \left[1 + \left(\frac{x}{x_L} \right)^{2p} \right]$$

$|x_L|$ valeur de $|x|$ pour $f = f_1$ ou $f = f_2$

$$|x_L| = |x_1| = |x_2| = \frac{|x_1| + |x_2|}{2} = \frac{f_1 - f_2}{f_0}$$

$$x_L = \frac{B}{f_0}$$

B étant la bande passante f_1, f_2 définie comme nous l'avons dit à 3 db. K est un coefficient d'identification qui représente en somme la perte de gain dû

aux désaccords des circuits : en effet, à la fréquence f_0 l'amplification aura pour valeur :

$$A_{f_0} = \frac{\prod_{n=1}^{n=p} S_n R_n}{K}$$

$\prod_{n=1}^{n=p} S_n R_n$ est l'amplification que l'on aurait si tous les circuits étaient accordés sur f_0 .

Nous appellerons K le *facteur de réduction*.

Remarquons que plus le nombre p de circuits sera grand, plus la courbe de réponse de l'amplificateur prendra une forme rectangulaire : A restera sensiblement constant tant que $\left| \frac{x}{x_L} \right|$ ne sera pas trop voisin de 1, puis décroîtra brutalement pour $|x| > |x_L|$.

II. — Calcul des paramètres des circuits.

Il s'agit maintenant d'identifier les deux polynômes $P(x)$ et $Q(x)$; pour ce faire, la méthode classique consiste à identifier les racines complexes des équations :

$$\begin{cases} P(x) = 0 \\ Q(x) = 0 \end{cases}$$

Racines de $Q(x) = 0$

$$1 + \left(\frac{x}{x_L} \right)^{2p} = 0$$

Les racines de cette équation sont :

$$(1) \quad \begin{cases} x = x_L e^{j(2n-1) \frac{\pi}{2p}} \\ n = 1, 2, \dots, p \end{cases}$$

Racines de $P(x) = 0$

$$\prod_{n=1}^{n=p} (1 + Q_n^2 x_n^2) = 0$$

Les racines de cette équation sont évidemment :

$$(2) \quad \begin{cases} x_n = \pm \frac{Q_n}{j} \\ n = 1, 2, \dots, p \end{cases}$$

x et x_n sont des fonctions de f : prenons donc f comme variable et posons :

$$f = \varphi + j\psi$$

Les équations (1) s'écrivent alors :

$$\frac{\varphi + j\psi}{f_{0n}} - \frac{f_{0n}}{\varphi + j\psi} = \pm \frac{j}{Q_n}$$

En égalant les parties réelles et les parties imaginaires :

$$\begin{cases} \varphi^2 + \psi^2 = f_{0n}^2 \\ \frac{\psi}{f_{0n}} = \pm \frac{1}{2Q_n} \end{cases}$$

D'autre part les équations (2) s'écrivent :

$$\frac{\varphi + j\psi}{f_0} - \frac{f_0}{\varphi + j\psi} = x_L e^{j(2n-1) \frac{\pi}{2p}}$$

On en déduit immédiatement en éliminant φ et ψ

$$(F) \quad \begin{cases} \left(\frac{f_{0n}}{f_0} + \frac{f_0}{f_{0n}} \right) \times \frac{1}{2Q_n} = \frac{B}{f_0} \sin(2n-1) \frac{\pi}{2p} \\ \left| \frac{f_{0n}}{f_0} - \frac{f_0}{f_{0n}} \right| \sqrt{1 - \left(\frac{1}{2Q_n} \right)^2} = \frac{B}{f_0} \cos(2n-1) \frac{\pi}{2p} \end{cases}$$

Ces formules permettent de calculer rapidement avec la précision nécessaire un amplificateur à circuits décalés. Il est d'ailleurs dans la pratique une précision qu'il est absolument inutile de dépasser. Le plus simple est de partir des formules approchées puis de corriger ensuite les valeurs trouvées pour f_0 et Q , en utilisant les formules précédentes.

Dans la très grande majorité des cas ($B < 0,4 f_0$ environ) on pourra négliger le terme $\sqrt{1 - \left(\frac{1}{2Q_n} \right)^2}$. En faisant une approximation de plus on écrira en posant :

$$\Delta f_{0n} = |f_{0n} - f_0|$$

$$\begin{cases} \frac{1}{Q_n} = \frac{B}{f_0} \sin(2n-1) \frac{\pi}{2p} \\ \Delta f_{0n} = 0,5 B \cos(2n-1) \frac{\pi}{2p} \end{cases}$$

Si le nombre des étages est impair on voit d'après ces formules qu'un circuit sera accordé sur la fréquence centrale f_0 et que sa surtension sera $Q = \frac{f_0}{B}$ c'est-à-dire sa bande passante B .

III. — Relations générales. Détermination du nombre d'étages.

Les relations précédentes permettent de calculer les éléments des circuits (Q_n, f_{0n}) lorsque l'on s'est fixé la bande passante f_1, f_2 et le nombre de circuits. Comment déterminer alors le nombre d'étages nécessaires pour obtenir l'amplification désirée ? Les relations que nous allons maintenant examiner fourniront une réponse immédiate à cette question.

Facteur de réduction K .

L'amplification a comme nous l'avons vu pour expression :

$$A = \prod_{n=1}^{n=p} \frac{S_n R_n}{\sqrt{1 + Q_n^2 x_n^2}} = \frac{\prod_{n=1}^{n=p} S_n R_n}{K \sqrt{1 + \left(\frac{x}{x_L} \right)^{2p}}}$$

K est le coefficient que nous avons appelé facteur de réduction, c'est le coefficient par lequel il faut diviser le produit des amplifications maxima de chaque étage pour avoir l'amplification totale à la fréquence centrale.

$$K^2 = \prod_{n=1}^{n=p} (1 + Q_n^2 x_{0n}^2)$$

$$x_{0n} = \frac{f_{0n}}{f_0} - \frac{f_0}{f_{0n}}$$

Or d'après les formules générales (F) il apparaît aisément en élevant les deux équations au carré et en les additionnant que l'on a :

$$x_L^2 = \left(\frac{1}{Q_n}\right)^2 + x_{on}^2$$

d'où l'expression de K :

$$K = (x_L)^p \prod_{n=1}^{n=p} Q_n$$

Si l'on prend comme référence la surtension $Q = \frac{1}{x_L}$ (dans le cas d'un nombre impair de circuits Q est le coefficient de surtension du circuit central) il vient :

$$K = \prod_{n=1}^{n=p} \frac{Q_n}{Q}$$

$$Q = \frac{1}{x_L} = \frac{f_0}{B}$$

Valeur approchée de K : D'après les formules approchées on a vu que :

$$Q_n \approx Q \times \frac{1}{\sin(2n-1) \frac{\pi}{2p}}$$

d'où :

$$K \approx \prod_{n=1}^{n=p} \frac{1}{\sin(2n-1) \frac{\pi}{2p}}$$

Pour calculer cette expression, on écrit :

$$\sin \frac{(2n-1)\pi}{2p} = \frac{e^{j(2n-1)\frac{\pi}{2p}} - e^{-j(2n-1)\frac{\pi}{2p}}}{2j} = \frac{x_n^2 - \frac{1}{x_n}}{2j}$$

$$\prod_{n=1}^{n=p} \sin(2n-1) \frac{\pi}{2p} = \left(\frac{1}{2}\right)^p \times \left(\frac{1}{j}\right)^p \prod_{n=1}^{n=p} \frac{x_n^2 - \frac{1}{x_n}}{2j}$$

Les x_n étant les racines de l'équation :

$$1 + (x^2)^p = 0$$

toutes les fonctions des racines de cette équation sont nulles sauf leur produit :

$$x_1^2 x_2^2 \dots x_p^2 = (-1)^p$$

d'où :

$$\prod_{n=1}^{n=p} \sin(2n-1) \frac{\pi}{2p} = \frac{1}{2^{p-1}}$$

$$K \approx 2^{p-1}$$

Cela veut dire en somme que le facteur de réduction croît de 6 dbs quand augmente d'une unité le nombre des étages de l'amplificateur.

Gain et bande passante.

Nous venons de voir que le gain total à la fréquence centrale pouvait se mettre sous la forme :

$$G_0 = \prod_{n=1}^{n=p} \frac{S_n R_n}{Q_n} \quad \left(Q = \frac{f_0}{B}\right)$$

Admettons que chaque étage à même coefficient de mérite $\frac{S}{2\pi C}$

$$Q_n = R_n C_n \omega_{on}$$

$$G_0 = \left(\frac{S}{2\pi C}\right)^p \times (Q)^p \times \prod_{n=1}^{n=p} \frac{1}{f_{on}}$$

Comme les circuits sont disposés symétriquement par rapport à f_0 (symétrie géométrique)

$$\prod_{n=1}^{n=p} f_{on} = f_0^p$$

D'où en introduisant le gain moyen par étage :

$$G_m = (G_0)^{\frac{1}{p}}$$

La formule fondamentale :

$$G_m \times B = \frac{S}{2\pi C}$$

Le produit de la bande passante de l'amplificateur par le gain moyen par étage est égal au coefficient de mérite des étages.

Autrement dit : si l'on veut un gain G_0 et une bande passante B le nombre d'étages nécessaire sera p tel que :

$$(G_m)^p = G_0$$

G_m étant donné par la relation fondamentale.

Cas de plusieurs groupes de circuits.

Au lieu d'utiliser un seul groupe de circuits tous décalés les uns par rapport aux autres, on peut envisager plusieurs groupes de circuits comportant chacun moins d'étages, ce qui peut être quelquefois plus pratique pour les réglages.

Supposons que nous ayons q groupes de p circuits chacun. L'amplification est alors :

$$A = \frac{A_{f_0}}{\left[1 + \left(\frac{x}{x_L}\right)^{2p}\right]^{q/2}}$$

au lieu de :

$$A = \frac{A_{f_0}}{\sqrt{1 + \left(\frac{x}{x_L}\right)^{2pq}}}$$

dans le cas d'un seul groupe de p.q circuits.

Calculons la bande passante à 3 dbs d'affaiblissement :

$$\left[1 + \left(\frac{x}{x_L}\right)^{2p}\right]^q = 2$$

$$\left(\frac{x}{x_L}\right) = \left(2^{\frac{1}{q}} - 1\right)^{\frac{1}{2p}}$$

D'où une certaine réduction de la bande passante : la formule précédente permet de déterminer la valeur de x_L qui correspond à la bande cherchée.

Affaiblissement en dehors de la bande passante.

Dans certains cas on est obligé de tenir compte de l'affaiblissement en dehors de la bande passante si l'on veut avoir par exemple un véritable effet de filtrage permettant de séparer différents canaux.

Soit N l'affaiblissement en népers que l'on désire

obtenir à des fréquences telles que $x/x_L = k$, $\left(x = \frac{kB}{f_0}\right)$
on écrira :

$$1 + \left(\frac{x}{x_L}\right)^{2p} = e^{2N}$$

On pourra négliger 1 devant e^{2N} : on pourra en déduire alors le nombre d'étages qu'il faudra utiliser pour atteindre cet affaiblissement :

$$P_{Nk} = \frac{N}{\log k}$$

Si l'on employait plusieurs groupes de circuits on serait amené à un nombre total de circuits plus élevé.

IV. — Variations du temps de transmission.

Examinons rapidement l'allure que présente la courbe du temps de transmission en fonction de la fréquence dans le cas d'un amplificateur à gain plat du type précédent. Pour cela on peut utiliser les relations liant l'amplitude à la phase. (réseaux à minimum de phase) : d'une manière plus élémentaire on peut partir des formules mêmes donnant dans le cas considéré le temps de transmission.

L'amplification d'un étage a pour expression :

$$\Lambda = - \frac{\Lambda_0}{1 + j Q_n x_n}$$

$$\operatorname{tg} \varphi_n = + Q_n x_n$$

$$\tau_n = \frac{d\varphi_n}{df} = \frac{Q_n \frac{dx_n}{df}}{1 + Q_n^2 x_n^2}$$

$$\theta = \frac{1}{2\pi} \tau = \frac{1}{2\pi} \sum_{n=1}^{n=p} \frac{d\varphi_n}{df} \quad (\theta \text{ temps de transmission}).$$

τ_n a pour expression :

$$\tau_n = \frac{d\varphi_n}{df} = \frac{\frac{Q_n}{f} \left(\frac{f_{on}}{f} + \frac{f}{f_{on}} \right)}{1 + Q_n^2 \left(\frac{f}{f_{on}} - \frac{f_{on}}{f} \right)^2}$$

Nous supposons pour simplifier que B/f_0 est relativement faible :

$$\frac{1}{Q_n} \approx \frac{B}{f_0} \sin \alpha_n$$

$$\left| \frac{f_0}{f_{on}} - \frac{f_{on}}{f_0} \right| \approx \frac{B}{f_0} \cos \alpha_n$$

$$\alpha_n = (2n - 1) \frac{\pi}{2p} \quad n = 1, 2, \dots, p$$

On posera d'autre part :

$$f = f_0 + \varepsilon$$

$$\text{et } \frac{2\varepsilon}{B} = \beta$$

Dans ces conditions on peut se rendre compte que τ_n se met sous la forme :

$$\tau_n \approx \frac{2}{B} \frac{\sin \alpha_n}{1 + 2\beta \cos \alpha_n + \beta^2}$$

$$\tau \approx \frac{2}{B} \sum_{n=1}^{n=p} \frac{\sin \alpha_n}{1 + 2\beta \cos \alpha_n + \beta^2}$$

Cherchons comment varie τ en fonction de β .

a) On remarque d'abord que si l'on change β en $-\beta$, τ ne change pas (les circuits sont groupés par paire correspondant à α_n et à $\pi - \alpha_n$ donc symétrie par rapport à la fréquence centrale, avec les hypothèses faites, bien entendu).

b) La dérivée $\tau'(\beta)$ a pour valeur :

$$\tau'(\beta) \approx \frac{2}{B} \sum_{n=1}^{n=p} \frac{-2 \sin \alpha_n (\beta + \cos \alpha_n)}{(1 + 2\beta \cos \alpha_n + \beta^2)^2}$$

le signe de $\tau(\beta)$ pour $\beta \ll 1$ est donné par :

$$\tau'(\beta) \approx \frac{2}{B} \sum_{n=1}^{n=p} -2 \sin \alpha_n (\beta + \cos \alpha_n) (1 - 4\beta \cos \alpha_n)$$

$$\tau'(\beta) \approx \frac{4}{B} \beta \sum_{n=1}^{n=p} (-4 \sin^3 \alpha_n + 3 \sin \alpha_n)$$

$$\tau'(\beta) \approx \frac{4}{B} \beta \sum_{n=1}^{n=p} \sin 3 \alpha_n$$

Si n est différent de 1, $\sum \sin 3 \alpha_n$ est positif, donc $\tau'(\beta)$ est du signe de β : τ présente un *minimum* pour $f = f_0$ fréquence centrale.

Dans le cas d'un seul circuit :

$$\tau \approx \frac{2}{B} \times \frac{1}{1 + Q^2 \left(\frac{f}{f_0} - \frac{f_0}{f} \right)^2}$$

présente un maximum pour $f = f_0$.

Aux fréquences limites de la bande $\beta = \pm 1$.
Pour $\beta = +1$ par exemple :

$$\tau'(\beta) \approx - \frac{2}{B} \sum_{n=1}^{n=p} \operatorname{tg} \frac{\alpha_n}{2}$$

Comme $\frac{\alpha_n}{2} < \frac{\pi}{2}$, $\tau'(\beta)$ est négatif puisque c'est

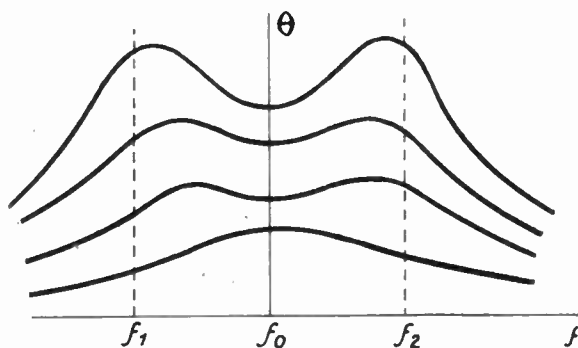


Fig. 4

une somme de termes tous négatifs. Le temps de transmission décroît.

Remarquons d'ailleurs que l'on a alors :

$$\tau \approx \frac{1}{B} \sum_{n=1}^{n=p} \operatorname{tg} \frac{\alpha_n}{2}$$

Quand le nombre d'étages augmente τ et τ' augmentent de plus en plus puisque :

$$\frac{x_p}{2} \longrightarrow \frac{\pi}{2}$$

Finalement les courbes du temps de transmission ont des allures telles que celles représentées fig. 4. Les variations de θ au voisinage de f_0 sont évidemment plus rapides que celles de A puisque θ varie comme le carré de $(f - f_0)$.

V. — Réalisation d'un amplificateur à circuits décalés.

Nous considérons comme précédemment le cas d'un amplificateur à gain plat dans une certaine bande de fréquences. La première question que l'on peut se poser est de se demander s'il est préférable pratiquement de réaliser l'amplificateur avec un seul groupe de circuits tous décalés les uns par rapport aux autres, ou bien avec plusieurs groupes de circuits. Théoriquement, dans le cas particulier envisagé, nous venons de voir qu'un seul groupe est plus intéressant : en réalité, un inconvénient du groupe unique est parfois de conduire à des circuits en bout de bande présentant des surtensions trop élevées. Ces coefficients de surtension élevés sont en très haute fréquence difficiles à obtenir ; les résistances d'entrée des lampes pouvant être relativement faibles même avec des tubes HF. D'autre part, la présence d'impédances de charges élevées devient en haute fréquence très rapidement gênante à cause de la réaction grille-plaque. Si cette réaction n'est pas négligeable, la détermination a priori des caractéristiques de l'amplificateur se complique beaucoup.

A titre d'indication examinons rapidement l'impédance ramenée à l'entrée d'un tube électronique par une charge que l'on supposera assimilable à un circuit antirésonnant. Si l'impédance interne du tube est grande vis-à-vis de Z l'impédance d'entrée a pour expression :

$$Z_e = \frac{1 + Z \gamma p}{\gamma p (1 + SZ)}$$

On peut négliger $Z \gamma p$ devant l'unité :

$$Z_e = \frac{1}{\gamma p (1 + SZ)}$$

Soit donc :

$$\frac{1}{Z_e} = \frac{1}{R} + Cp + \frac{1}{Lp}$$

Z_e s'écrit alors :

$$Z_e = \frac{R (1 - LC \omega^2) + L \omega j}{\gamma \omega j [R (1 - LC \omega^2) + (1 + SR) L \omega j]}$$

Si l'on suppose SR assez grand ce qui est précisément le cas gênant :

$$Z_e = \frac{1}{SR \gamma \omega j} + \frac{C}{S \gamma} \left(1 - \frac{\omega_0^2}{\omega^2} \right)$$

L'impédance ramenée est équivalente à une capacité $SR \gamma$ en série avec une résistance variable

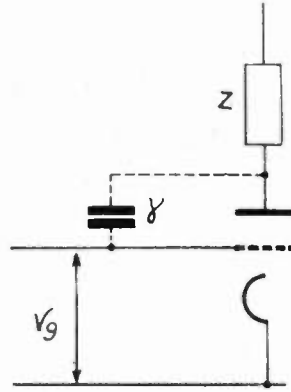


Fig. 5

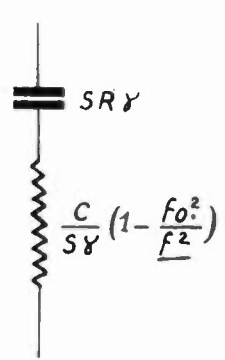


Fig. 6

$\frac{C}{S \gamma} \left(1 - \frac{f_0^2}{f^2} \right)$. En dessous de la fréquence de résonance, la résistance est donc négative et au-dessus positive.

Avec une 6 AK 5 par exemple, on aura γ de l'ordre de $5/100 \text{ pF}$.

Si $R = 3\,000 \text{ ohms}$, la capacité est de l'ordre de $0,7 \text{ à } 0,8 \text{ pF}$; avec une capacité entre étages de $10 \text{ à } 11 \text{ pF}$ et des résistances parallèles de $3\,000 \text{ à } 4\,000 \text{ ohms}$, il n'est plus possible de négliger cette impédance ramenée à l'entrée de la lampe.

Pour réaliser pratiquement un amplificateur dans lequel la réaction grille-plaque n'est pas tout à fait négligeable il est intéressant de procéder à l'accord de chaque circuit en maintenant chargé le tube qui suit ; pour cela, il est nécessaire d'utiliser un volt-

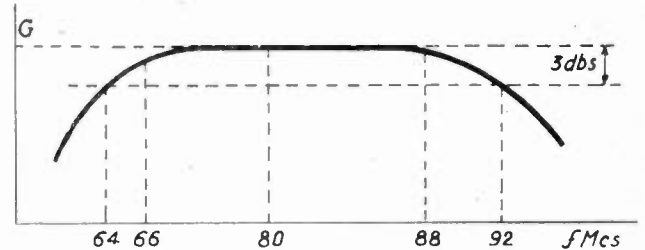


Fig. 7

mètre présentant une grande impédance d'entrée (capacité faible, résistance élevée) ; on peut employer par exemple une tête de voltmètre constituée par un cristal à très faible capacité, cette tête pouvant se mettre facilement en parallèle sur un circuit de charge sans en changer les caractéristiques.

Voici à titre d'exemple, la courbe de réponse (fig. 7), d'un amplificateur à gain plat à très large bande de 10 étages. Les 10 étages sont constitués par 2 groupes de 5 circuits : les tubes utilisés sont des 6 AK 5.

Le coefficient de mérite par étage est d'environ 70 Mc/s .

$$S = 5 \text{ mA/volt}$$

$$C = 10,5 \text{ à } 11 \text{ pF}$$

Le gain total est de 70 dbs .

LE GLISSEMENT DE FRÉQUENCE PAR LAMPES A RÉACTANCE VARIABLE⁽¹⁾

PAR

R. LEPRÊTRE

ancien élève de l'Ecole Polytechnique

ing. E.S.E., chef de la division « Aviation civile » du C.N.E.T.

Les lampes à réactance variable sont destinées à produire par variations des tensions appliquées sur une de leur grille, un changement de la fréquence des circuits oscillants aux bornes desquels elles sont montées.

L'espace cathode plaque est alors assimilable à un élément *réactif*, venant se mettre en dérivation avec celui du circuit oscillant. Le type du montage donne ainsi à une même lampe des caractères capacitifs ou selfiques selon les besoins.

Les lampes à réactance variable sont principalement utilisées dans certains cas de téléphonie par modulation de fréquence. Elles peuvent aussi être utilisées pour obtenir dans un récepteur un balayage en fréquence autour d'une fréquence préétablie.

Nous allons étudier le mécanisme du « glissement de fréquence » dans de tels montages et calculer les glissements et les amortissements correspondants à chacun d'eux.

Il y a lieu souvent dans les cas où l'amplification n'est pas surabondante de choisir très convenablement les éléments du montage de façon à ne pas avoir un amortissement des circuits oscillants trop considérable et c'est en cela que les résultats suivants, obtenus de la façon la plus rigoureuse possible, peuvent être utiles. D'autre part, nous indiquerons les avantages respectifs des différents montages en fonction de la fréquence.

I. MONTAGE DE LAMPE A RÉACTANCE SELFIQUE

Soit le schéma de la figure 1. Le circuit résonnant est le circuit accordé d'une lampe autooscillatrice non représentée. A ses bornes est monté le dispositif représenté ci-dessus où T est une lampe triode ou pentode, r une résistance et C_1 une capacité dont nous préciserons les valeurs ci-dessous. C_2 est une simple capacité de liaison.

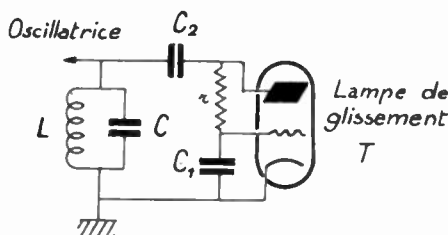


Fig. 1

Nous allons montrer que l'impédance équivalente à la lampe T et à son circuit r, C_1 est selfique et nous allons calculer le glissement et l'amortissement correspondant, de la manière la plus rigoureuse possible.

La tension $\mathcal{E}_p$ est la somme de la tension $\mathcal{E}_r$ et de la tension $\mathcal{E}_g$ aux bornes de la capacité.

On a par conséquent le diagramme des tensions représenté, figure 2 où $\mathcal{E}_g$ est en quadrature arrière de $\mathcal{E}_r$ puisque $\mathcal{E}_g$ est une différence de potentiel aux bornes d'une capacité.

La résistance r sera prise grande devant l'impédance $\frac{1}{C_1 \omega} \left(\frac{\mathcal{E}_g}{\mathcal{E}_r} \leq 1 \right)$ ce qui est la condition pour que l'impédance équivalente à la triode soit presque uniquement selfique. Ceci peut être vu sur le diagramme de la manière

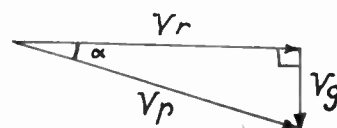


Fig. 2

suivante : le courant plaque de la lampe est la somme de deux composantes : l'une, la plus importante en phase avec $\mathcal{E}_g$, l'autre en phase avec $\mathcal{E}_p$. Si l'angle α est petit, le courant plaque sera en quadrature arrière avec $\mathcal{E}_p$ et la lampe équivalant à une self.

Appelons $\mathcal{Z}$ l'impédance du circuit équivalent à l'ensemble des circuits LC et r, C_1 .

Nous avons :

$$\frac{1}{\mathcal{Z}} = \frac{1}{R} - \frac{j}{S} + \frac{1}{r - \frac{j}{C_1 \omega}}$$

(1) Bibliographie : « La modulation de fréquence » par Mr Besson (Chiron).

Wide Deviation Reactance Modulator. (Electronics Avril 1948).

$\frac{1}{R}$ et $\frac{1}{S}$ étant l'admittance ohmique et l'admittance réactive du circuit LC au voisinage de l'accord.

$$\left(\frac{1}{S} = \frac{1}{L\omega} - C\omega \right)$$

Introduisons dans les calculs la quantité α représentant l'angle entre $\mathcal{V}_r$ et $\mathcal{V}_p$. α qui est petit comme nous l'avons vu est approximativement égal à

$$\frac{\mathcal{V}_g}{\mathcal{V}_r} = \frac{1}{r C_1 \omega} \# \frac{\mathcal{V}_g}{\mathcal{V}_p}$$

d'où :

$$\frac{1}{\mathcal{Z}'} = \frac{1}{R} - \frac{j}{S} + \frac{1/r}{1 - \alpha j}$$

$$\frac{1}{\mathcal{Z}'} = \frac{1}{R} - \frac{j}{S} + \frac{1}{1 + \alpha^2} \frac{1 + \alpha j}{r}$$

$$\frac{1}{\mathcal{Z}'} = \frac{1}{R} + \frac{1/r}{1 + \alpha^2} - \frac{j}{S} + \frac{\alpha j/r}{1 + \alpha^2}$$

$$\frac{1}{\mathcal{Z}'} = \frac{1}{R} + \frac{1}{r} \left(\frac{1}{1 + \alpha^2} \right) + j \left(\frac{\alpha}{1 + \alpha^2} \times \frac{1}{r} - \frac{1}{S} \right)$$

α^2 étant faible devant 1

$$\frac{1}{\mathcal{Z}'} = \frac{1}{R} + \frac{1}{r} + j \left(\frac{\alpha}{r} - \frac{1}{S} \right)$$

Appliquons la formule fondamentale de la lampe en notation complexe.

$$\rho \mathcal{V}_p = k \mathcal{V}_g + \mathcal{V}_p$$

avec $\mathcal{V}_p = -z' \mathcal{V}_p$

On en déduit que :

$$\mathcal{V}_p = \frac{k \mathcal{V}_g}{\rho + \mathcal{Z}'} = \frac{k \mathcal{V}_g}{\mathcal{Z}'} \frac{1}{\frac{\rho}{\mathcal{Z}'} + 1}$$

d'où en remplaçant $\frac{1}{\mathcal{Z}'}$ par sa valeur ;

et en multipliant haut et bas par la quantité imaginaire conjuguée.

$$\mathcal{V}_p = k \mathcal{V}_g \frac{\left(\frac{1}{R} + \frac{1}{r} \right) \left[1 + \rho \left(\frac{1}{R} + \frac{1}{r} \right) \right] + \rho \left(\frac{\alpha}{r} - \frac{1}{S} \right)^2 + j \left(\frac{\alpha}{r} - \frac{1}{S} \right)}{\left[\rho \left(\frac{1}{R} + \frac{1}{r} \right) + 1 \right]^2 + \rho^2 \left(\frac{\alpha}{r} - \frac{1}{S} \right)^2}$$

$$\text{soit } \mathcal{V}_p = k \mathcal{V}_g \frac{\Re}{\Omega}$$

Calculons maintenant $\mathcal{Z}$ l'impédance équivalente de la lampe. On a, $\frac{1}{\mathcal{R}}$ et $\frac{1}{\mathcal{S}}$ étant ses admittances ohmique et réactive équivalentes.

$$\frac{1}{\mathcal{Z}} = \frac{1}{\mathcal{R}} + \frac{1}{j\mathcal{S}}$$

or $\frac{1}{\mathcal{Z}} = \frac{\mathcal{V}_p}{\mathcal{V}_g}$ ou en remplaçant $\mathcal{V}_p$ par sa valeur $\frac{k \mathcal{V}_g \Re}{\Omega}$

précédemment calculée $\frac{1}{\mathcal{Z}} = k \frac{\mathcal{V}_g}{\mathcal{V}_p} \times \frac{\Re}{\Omega}$

Déterminons alors le rapport $\frac{\mathcal{V}_g}{\mathcal{V}_p}$ en amplitude et en phase au moyen du diagramme de Fresnel

On a vu précédemment que $\alpha \# \frac{\mathcal{V}_g}{\mathcal{V}_p}$ d'où :

$$\frac{\mathcal{V}_g}{\mathcal{V}_p} = \frac{V_g}{V_p} (\sin \alpha - j \cos \alpha) \# \alpha (\alpha - j)$$

$$\text{d'où : } \frac{1}{\mathcal{Z}} = k \alpha (\alpha - j) \left(\frac{\Re}{\Omega} \right)$$

Posons pour simplifier les calculs.

$$\frac{1}{R} + \frac{1}{r} = \frac{1}{R'} \quad (1)$$

$$\left(\frac{\alpha}{r} - \frac{1}{S} \right) = -\frac{1}{S'} \quad (2)$$

Il vient :

$$\frac{1}{\mathcal{Z}} = k \alpha (\alpha - j) \left[\frac{1}{R'} \left(1 + \frac{\rho}{R'} \right) + \frac{\rho}{S'^2} - \frac{j}{S'} \right]$$

d'où :

$$\frac{1}{\mathcal{R}} = k \alpha \frac{\frac{\rho \alpha}{S'^2} - \frac{1}{S'} + \frac{\alpha}{R'} \left(1 + \frac{\rho}{R'} \right)}{\left(\frac{\rho}{R'} + 1 \right)^2 + \frac{\rho^2}{S'^2}} \quad (3)$$

$$\frac{1}{\mathcal{S}} = k \alpha \frac{\frac{\rho}{S'^2} + \frac{\alpha}{S'} + \frac{1}{R'} \left(1 + \frac{\rho}{R'} \right)}{\left(\frac{\rho}{R'} + 1 \right)^2 + \frac{\rho^2}{S'^2}} \quad (4)$$

Calcul de l'admittance réactive x en parallèle sur le circuit oscillant :

L'admittance réactive cherchée équivalente à la lampe T et au circuit $r C_1$ s'obtient en ajoutant à l'admittance réactive $\frac{1}{\mathcal{S}}$ de la lampe T seule, l'admittance réactive du circuit $r C_1$ soit : $-\frac{\alpha}{r}$

x étant l'admittance réactive totale ou a $x = \frac{1}{\mathcal{S}} - \frac{\alpha}{r}$

Or la formule (4) précédemment établie ne donne $\frac{1}{\mathcal{S}}$ qu'en fonction de $\frac{1}{S'}$ c'est-à-dire de $\frac{1}{S}$ (formule (2)).

Or il existe entre ces 3 grandeurs $\mathcal{S}$, S et S' des relations qui vont nous permettre de calculer intrinséquement x

$$\text{Nous avons vu que } \frac{1}{S} = -C\omega + \frac{1}{L\omega}$$

où ω est la nouvelle pulsation d'accord définie par l'équation.

$$\left(x + \frac{1}{L\omega} \right) = C\omega$$

$$\text{d'où } x = C\omega - \frac{1}{L\omega} = -\frac{1}{S} = -\frac{1}{S'} - \frac{\alpha}{r} \text{ (d'après (2))}$$

c'est-à-dire $\frac{1}{S'} = -\left(x + \frac{\alpha}{r}\right)$ (5)

et $\frac{1}{S'} = \frac{\alpha}{r} - \frac{1}{S} = \left(x + \frac{\alpha}{r}\right)$ (6)

L'équation (4) explicitée à l'aide des équations (5) et (6) devient, en posant $x + \frac{\alpha}{r} = y$ (7)

(remarquons que y n'est autre que $\frac{1}{S'}$ admittance réactive de la lampe seule) et en faisant apparaître la pente de la lampe :

$$y = p \alpha \frac{\rho^2 y^2 - \alpha \rho y + \frac{\rho}{R'} \left(1 + \frac{\rho}{R'}\right)}{\left(\frac{\rho}{R'} + 1\right)^2 + \rho^2 y^2} \quad (8)$$

Cette équation (8) nous donne y et par suite d'après (7) l'admittance équivalente x en fonction de la pente p et de la résistance intérieure de la lampe, de $\frac{1}{R'} = \frac{1}{R} + \frac{1}{r}$ et de $\alpha = \frac{1}{rc_1 \omega}$ dont on peut avoir une valeur approchée en prenant pour ω la valeur de la pulsation d'origine.

Bien que l'équation (8) ne soit pas explicitée en y , on peut démontrer que x est toujours positif donc que l'admittance de l'ensemble de la lampe T et de son circuit rc_1 est selfique

Posons en effet :

$$t = \frac{\rho^2 y^2 - \alpha \rho y + \frac{\rho}{R'} \left(1 + \frac{\rho}{R'}\right)}{\left(\frac{\rho}{R'} + 1\right)^2 + \rho^2 y^2} \quad (9)$$

Il vient :

$$y = p \alpha t \quad (10)$$

Or trouver la solution y de l'équation (8) revient à chercher l'intersection des courbes définies par les équations (9) et (10) et il faut démontrer que cette solution est toujours telle que x soit positif (figure 3).

La fonction t est toujours positive quel que soit y : car l'équation $\rho^2 y^2 - \alpha \rho y + \frac{\rho}{R'} \left(1 + \frac{\rho}{R'}\right) = 0$ a toujours des racines imaginaires puisque l'expression

$$\alpha^2 \rho^3 - \frac{4 \rho^3}{R'} \left(1 + \frac{\rho}{R'}\right) = \rho^3 \left[\alpha^2 - \frac{4 \rho}{R'} \left(1 + \frac{\rho}{R'}\right) \right]$$

est négative.

D'un autre côté cette fonction a une dérivée qui est égale à :

$$t' = \frac{\alpha \rho^3 y^2 + 2 \rho^2 y \left(\frac{\rho}{R'} + 1\right) - \alpha \rho \left(\frac{\rho}{R'} + 1\right)}{\left[\left(\frac{\rho}{R'} + 1\right)^2 + \rho^2 y^2\right]^2}$$

t' s'annule pour :

$$y = -\rho^2 \left(\frac{\rho}{R'} + 1\right) \pm \sqrt{\rho^4 \left(\frac{\rho}{R'} + 1\right)^2 + \alpha^2 \rho^4 \left(\frac{\rho}{R'} + 1\right)^2}$$

soit en négligeant α^4 devant 1

$$y = \frac{\alpha}{2} \left(\frac{1}{R'} + \frac{1}{\rho}\right) \neq 0$$

$$y = \left(\frac{1}{R'} + \frac{1}{\rho}\right) \left(\frac{2}{\alpha} + \frac{\alpha}{2}\right)$$

t' est négatif entre les racines, positif à l'extérieur, la fonction t croît donc depuis $y = -\infty$ passe par

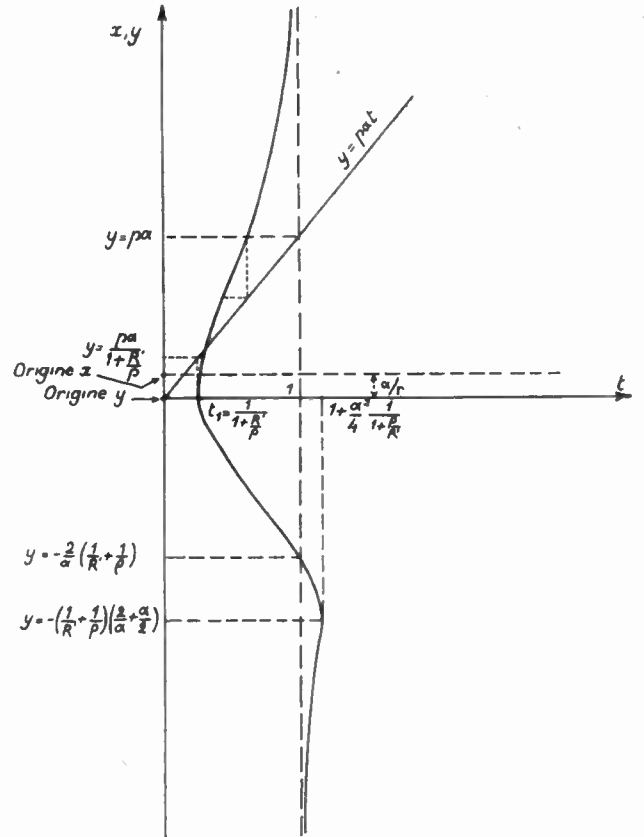


Fig. 3

un maximum, décroît, passe par un minimum pour y voisin de 0 puis croît jusqu'à $y = +\infty$ (voir figure 3)

Elle est asymptote à la valeur $t = +1$ pour $y = \pm \infty$.

La valeur de t minimum est : $t_1 = \frac{1}{1 + \frac{\rho}{R'}}$ (11)

D'après l'équation (7), l'origine des x est au dessus de l'axe des t et distante de cet axe de — la quantité $\frac{\alpha}{r}$. Le point d'intersection des courbes (9) et (10) c'est-à-dire de la courbe du 3^e degré que nous venons de tracer et de la droite ($y = p \alpha t$) a une ordonnée x supérieure à celle du point de cette dernière droite d'abscisse t_1 (équation (11)).

Or cette ordonnée est égale à :

$$x_1 = y_1 - \frac{\alpha}{r} = p \alpha t_1 - \frac{\alpha}{r} = p \alpha \frac{1}{1 + \frac{R'}{\rho}} - \frac{\alpha}{r}$$

$$\text{soit } Xq = \alpha \left(\frac{\rho}{1 + \frac{R'}{\rho}} - \frac{1}{r} \right) \quad (11 \text{ bis})$$

or on constate que la quantité $\left[\frac{\rho}{1 + \frac{R'}{\rho}} - \frac{1}{r} \right]$ est

positive, pour les valeurs courantes de montage (voir l'exemple plus loin) ; ce qui revient à dire que l'admittance capacitive du circuit $r c_1$ est inférieure à l'admittance selfique de la lampe T seule.

La réactance du montage (T, r, c_1) est donc selfique.

La self équivalente est alors égale à :

$$L = \frac{1}{x \omega}$$

Une valeur très approchée est obtenue en confondant x avec y d'où $L \approx \frac{1}{y \omega}$

La pulsation ω des formules est la pulsation de l'oscillatrice montée avec le circuit T r C_1 , la lampe T ayant ses caractéristiques internes normales.

Balayage ou modulation de fréquence — Calcul du glissement.

De part et d'autre de la fréquence f_1 correspondant à la pulsation ω d'équilibre des oscillations relative à des caractéristiques moyennes de la lampe T, il est possible de moduler en fréquence l'oscillatrice en faisant varier les caractéristiques internes de la lampe T par variations de la tension appliquée à une grille.

L'excursion de la fréquence en fonction de correspondant donné par l'équation (8) est obtenue à partir des équations

$$y - y_1 = \Delta y = c (\omega - \omega_1) - \left(\frac{1}{L \omega} - \frac{1}{L \omega_1} \right) = C \Delta \omega + \frac{\Delta \omega}{L \omega \omega_1}$$

$$\text{d'où } \Delta \omega = \frac{\Delta y}{C + \frac{1}{L \omega \omega_1}}$$

Si l'on suppose que ω et ω_1 sont voisins de ω_0 pulsation d'accord du circuit oscillant monté seul,

on peut confondre $\frac{1}{L \omega \omega_1}$ avec C et écrire

$$\Delta \omega = \frac{\Delta y}{2C} \text{ c'est à dire : } \Delta f = \frac{\Delta y}{4 \pi C} \quad (12)$$

Equation approchée de l'admittance réactive dans le cas général.

Nous pouvons voir d'après l'examen des courbes représentatives des équations (9) et (10) que dans l'équation (9) le terme $-\alpha \rho y$ est petit devant $\rho^2 y^2$ puis-

que ρy au moins égal à $\frac{\alpha k}{1 + \frac{R'}{\rho}}$ est grand devant α

($k > 1 + \frac{R'}{\rho}$) même pour les triodes.

On peut donc écrire :

$$t = \frac{\rho^2 y^2 + \frac{\rho}{R'} \left(1 + \frac{\rho}{R'} \right)}{\left(\frac{\rho}{R'} + 1 \right)^2 + \rho^2 y^2} \quad (13)$$

Nous pouvons également voir sur la figure, que le point de fonctionnement est compris entre des valeurs de y limites définissant la relation :

$$\frac{\rho \alpha}{1 + \frac{R'}{\rho}} \leq y < \rho \alpha$$

Cet intervalle sera du reste d'autant plus resserré que ρ sera grand devant R' (cas des pentodes à grande résistance interne).

Nous pouvons prendre des limites plus rapprochées de y en menant la parallèle à l'axe des t depuis les points de la droite $y = p \alpha t$ correspondant aux valeurs précédentes jusqu'à la courbe, en les rappe-
lant verticalement jusqu'à la droite puis en le reportant de nouveau horizontalement sur la courbe.

Pour cela il suffit de reporter dans l'équation (13) les deux valeurs limites précédentes, d'où deux valeurs de t et de les reporter dans l'équation (10) donnant deux nouvelles valeurs de y , également plus rapprochées. On procédera de même jusqu'à ce que Δy soit inférieure à une limite acceptable et on prendra pour valeur de y une valeur moyenne des limites du dernier intervalle.

Nous trouvons ainsi :

Pour la limite inférieure :

$$t = \frac{\frac{k^2 \alpha^2}{(1 + R'/\rho^2)^2} + \frac{\rho}{R'} \left(1 + \frac{\rho}{R'} \right)}{\left(\frac{\rho}{R'} + 1 \right)^2 + \frac{k^2 \alpha^2}{\left(1 + \frac{R'}{\rho} \right)^2}}$$

et y est égal à :

$$y_1 = p \alpha t = p \alpha \left[1 - \frac{1 + \rho/R'}{(\rho/R' + 1)^2 + \frac{k^2 \alpha^2}{\left(1 + \frac{R'}{\rho} \right)^2}} \right]$$

Pour la valeur supérieure :

$$y_2 = p \alpha \left[1 - \frac{1 + \rho/R'}{(\rho/R' + 1)^2 + k^2 \alpha^2} \right]$$

La différence entre les deux limites est égale à :

$$y_2 - y_1 = p \alpha (1 + \rho/R') \frac{k^2 \alpha^2 \left(1 - \frac{1}{(1 + R'/\rho)} \right)}{\left[(\rho/R' + 1)^2 + \left(\frac{k \alpha}{1 + R'/\rho} \right)^2 \right] \left[(\rho/R' + 1)^2 \right]}$$

C'est pour les triodes à faible résistance interne que cet intervalle est le plus grand.

Pour ces tubes les plus petites valeurs sont telles que

$$k \propto \# 1$$

$$\rho/R' = 1/2.$$

$$y_s - y_i = \# 0,1 p \alpha$$

On prendra pour y une valeur intermédiaire entre les deux limites, D'où la valeur suivante de l'admittance réactive :

$$x = -\frac{\alpha}{r} + p \alpha \left[\frac{1 + \rho/R'}{(\rho/R' + 1)^2 + \frac{k^2 \alpha^2}{2}} \right] \quad (14)$$

approchée à moins de $\pm 0,05 p \alpha$ près.

Remarque : L'erreur commise sur l'évaluation de y diminue en même temps que ρ/R' augmente, les deux limites se resserrant davantage.

Equation approchée de l'admittance réactive et du glissement dans le cas d'une pentode.

Dans ce cas ρ étant grand devant R' on peut négliger 1 devant ρ/R' .

L'équation (13) s'écrit :

$$1 = \frac{\rho^2 y^2 + \rho^2/R'^2}{\rho^2 y^2 + \rho^2/R'^2} = 1$$

Il vient alors :

$$y = p \alpha \quad (15)$$

La courbe représentative se confond avec son asymptote $l = 1$ et le point de fonctionnement est bien l'intersection de celle-ci avec la droite $y = p \alpha l$.

L'expression du glissement qui était dans le cas général $\Delta f = \frac{\Delta x}{4\pi C}$ peut s'écrire

$$\Delta f = \frac{\alpha \Delta p}{4\pi C}$$

et en remplaçant α et C par sa valeur en fonction de r, c_1, L et f on écrit :

$$\frac{\Delta f}{f} = \frac{L}{2r_1 C} \Delta p \quad (16)$$

Calcul de la conductance en parallèle sur le circuit oscillant.

Remplaçons dans l'équation (3) $\frac{1}{S}$, par $-\left(x + \frac{\alpha}{r}\right)$ (équation 5) ou par $-y$ (équation 7) où y est sensiblement égal à x , admittance du montage donnée par les formules (8) ou (13).

Il vient :

$$\frac{1}{\mathcal{R}} = p \alpha \frac{\alpha \rho^2 y^2 + \rho y + \alpha \frac{\rho}{R'} \left(1 + \frac{\rho}{R'}\right)}{\rho^2 y^2 + \left(\frac{\rho}{R'} + 1\right)^2}$$

y étant, comme nous l'avons vu, toujours positif,

$\frac{1}{\mathcal{R}}$ sera toujours positif car l'expression

$$\alpha \rho^2 y^2 + \rho y + \alpha \frac{\rho}{R'} \left(1 + \frac{\rho}{R'}\right)$$

dont les racines sont 0 et $-\frac{2}{\alpha \rho}$ (en négligeant α' devant 1) est positive pour y positif.

Nous allons maintenant calculer une valeur inférieure de la valeur de $1/\mathcal{R}$.

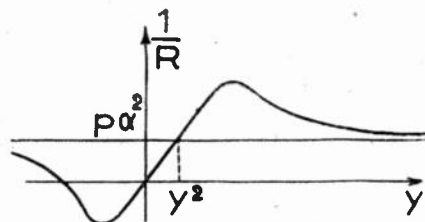


Fig. 4

L'expression précédente peut se mettre sous la forme :

$$\frac{1}{\mathcal{R}} = p \alpha^2 \left[1 + \frac{\frac{\rho y}{\alpha} - \left(1 + \frac{\rho}{R'}\right)}{\rho^2 y^2 + \left(\frac{\rho}{R'} + 1\right)^2} \right] \quad (17)$$

Le second terme entre parenthèses s'annule pour $\frac{\rho y}{\alpha} = 1 + \frac{\rho}{R'}$, c'est-à-dire pour $y_2 = \frac{\alpha}{\rho} \left(1 + \frac{\rho}{R'}\right)$ et est toujours positif pour des valeurs de y plus grandes. (voir courbe fig. 4).

Or y_2 est toujours plus petit que la valeur approchée de l'admittance réactive $y_1 = \frac{p \alpha}{1 + \frac{\rho}{R'}}$, car

$$\frac{\alpha}{\rho} \left(1 + \frac{\rho}{R'}\right) < \frac{p \alpha}{(1 + \rho/R')},$$

On constate en effet que l'on a en général (voir l'exemple plus loin)

$$\frac{\rho}{R'} + \frac{R'}{\rho} < k - 2 \quad (17 \text{ bis})$$

D'autre part y_1 est lui-même plus faible que l'admittance y , le terme complémentaire de l'équation (14) est positif et une limite inférieure très approchée de $1/\mathcal{R}$ est par suite obtenue en prenant pour y la valeur y_1 .

Une limite supérieure très approchée peut être obtenue en prenant pour y la valeur $p \alpha$ qui est supérieure à l'admittance y du point de fonctionnement et inférieure à la valeur de y correspondant au maximum de $\frac{1}{\mathcal{R}}$ ($p \alpha < \frac{1}{\rho} + \frac{1}{R'}$, c'est à dire $\alpha < \frac{1}{k} \left(1 + \frac{\rho}{R'}\right)$: inégalité toujours satisfaite, sauf dans le cas des pentodes à fort coefficient d'amplification, mais dans ce cas la conductance est quasiment égale à la limite inférieure).

D'où les deux limites rapprochées inférieure et supérieure de la conductance totale $\left(\frac{1}{r} + \frac{1}{R}\right)$ en parallèle sur le circuit oscillant.

$$\frac{1}{r} + p \alpha^2 \left[1 + \frac{\frac{k}{1 + R'/\rho} - \left(1 + \frac{\rho}{R'}\right)}{\left(1 + \frac{R'}{\rho}\right)^2 + \left(\frac{\rho}{R'} + 1\right)^2} \right] \quad (18)$$

$$\frac{1}{r} + p \alpha^2 \left[1 + \frac{k - \left(\frac{\rho}{R'} + 1\right)}{k^2 \alpha^2 + \left(\frac{\rho}{R'} + 1\right)^2} \right]$$

limites toutes deux supérieures à $\frac{1}{r} + p \alpha^2$

Formule de la conductance dans le cas de la pentode.

La conductance équivalente à l'ensemble T, r, c₁ calculée à partir de l'équation (17) où l'on remplace y par p α ou à partir de (18), devient égale à :

$$\frac{1}{r} + p \alpha^2 \left[1 + \frac{k - \frac{\rho}{R'}}{k^2 \alpha^2 + \frac{\rho^2}{R'^2}} \right] \quad (19)$$

Remarque : Dans la plupart des cas $\frac{\rho}{R'}$ est négligeable devant k et même k α.

Discussion des formules du glissement et de l'amortissement.

Dans le cas où la lampe à réactance variable T est employée pour moduler en fréquence l'oscillateur d'un émetteur on voit qu'il y a conjointement modulation de fréquence et modulation d'amplitude puisque l'amortissement et par suite, la surtension du circuit oscillant, varient avec la modulation.

Dans le cas où la lampe T est employée pour balayer en fréquence dans les circuits H. F. d'un récepteur, l'amortissement se traduit par une variation de la sensibilité du récepteur.

Choix de p et de α.

On remarque que le glissement Δ / (équation (12) et (14) est proportionnel grosso modo à p α et que la conductance en parallèle sur le circuit oscillant l'est grosso modo à p α² (formules 18 et 19). Nous verrons en effet ultérieurement que l'on prend r très grand.

Il s'en suit donc, suivant l'usage que l'on veut faire de la lampe à réactance, un compromis dans l'adoption des valeurs de p et de α, tout en cherchant à accroître le plus possible la valeur de la pente et à diminuer celle de α (I).

(1) Ce qui concerne notre choix, α petit devant 1, fait au début de cet exposé.

Choix de r et de C₁.

Pour une fréquence donnée, c'est-à-dire le rapport r c₁, étant choisi il est avantageux au point de vue de l'amortissement de prendre r le plus grand possible et c₁ le plus petit possible. Les formules (16 et 17) contiennent en effet le terme $\frac{1}{r}$.

Choix de ρ.

Les formules (14) et (18) montrent qu'on a avantage au point de vue du glissement et de l'amortissement à prendre une lampe ayant une très grande résistance interne. D'où l'avantage des pentodes.

Choix de L et C.

L'équation (12) montre que pour un Δ y donné, le glissement de fréquence sera d'autant plus grand que C sera faible et donc que L sera grand. On a donc avantage à prendre un circuit oscillant ayant la plus grande self possible. Par contre au point de vue de l'amortissement si L est trop grand pour une valeur donnée de la résistance équivalente à la lampe de glissement, l'amortissement sera excessif. Il y a donc lieu de prendre pour L une valeur moyenne.

Influence de la fréquence.

Supposons que l'on ait adopté, quelle que soit la fréquence, les mêmes valeurs de p α et C₁ (C₁ le plus faible possible déterminé par la valeur de la capacité grille cathode de la lampe).

$\frac{1}{r}$ qui est égal à α C₁ ω, augmente avec la fréquence. La résistance en parallèle sur le circuit oscillant diminue donc quand la fréquence croît.

Pratiquement pour avoir le même amortissement en H F qu'en BF on devra prendre un coefficient α plus petit en HF qu'en BF d'où un glissement relatif plus petit.

En effet la formule (16) s'écrit :

$$\frac{\Delta f}{f} = \frac{L}{2 r c_1} \Delta p = \frac{L \omega}{2} \alpha \Delta p$$

et L ω peut être considéré comme constant.

Les basses fréquences sont donc beaucoup plus avantageuses.

Remarque : Si l'on veut étudier par contre l'influence de la fréquence pour un montage déterminé (p, L, r C₁ donnés ; C et α variables) la formule (16) montre que le glissement relatif Δ / est indépendant de la fréquence. Quant à l'amortissement il varie peu avec la fréquence. En effet la conductance dont une valeur approchée est

$$(1/r + p \alpha^2) = 1/r + p/r^2 c_1^2 \omega^2$$

diminue quand la fréquence augmente et la résistance à l'anti résonance du circuit oscillant déterminé croît avec la fréquence, il y a par suite grosso modo compensation.

Exemple de Montage.

Supposons que nous voulions obtenir des glissements importants sans trop d'amortissement du circuit oscillant.

Comme nous l'avons dit nous prendrons une penthode à forte pente, la 6 A C 7 possède à ce point de vue des caractéristiques intéressantes.

La pente est de 9 mA/V avec une tension plaque de 300 V et une tension écran de 150 V et une polarisation grille de -3 V. Cette pente varie de 16 mA/V à 1 mA/V lorsque la polarisation varie de 0 à 5 V environ.

La résistance interne est pour la valeur moyenne :

$$\rho = 750.000 \text{ Ohms.}$$

Le coefficient d'amplification statique est $k = 6.750$.

I.) Supposons que l'on fonctionne sur une fréquence $f = 10 \text{ Mégacycles/seconde}$.

Nous prendrons comme self.

$$L = 8 \mu \text{ H.}$$

La valeur de la capacité C d'accord est par suite.

$C = 31 \text{ picofarads}$, valeur acceptable compte tenu des capacités de sortie et des connexions.

La résistance à l'antirésonance pour un coefficient de surtension de 200 sera égale à :

$$R = 100.000 \text{ Ohms.}$$

Nous prendrons r grand comme nous l'avons dit, soit :

$$r = 100.000 \text{ ohms}$$

$$\text{d'où } \frac{1}{R'} = \frac{1}{r} + \frac{1}{R} = \frac{1}{50.000} \text{ ohms}$$

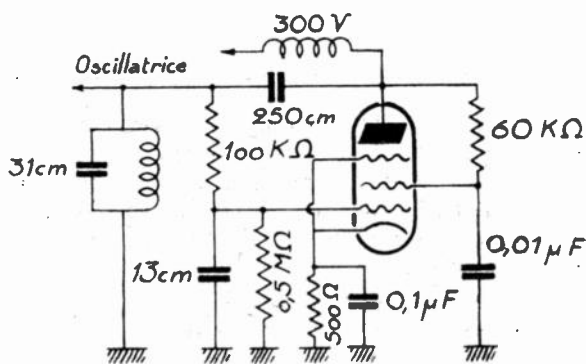


Fig. 5

Nous prendrons α très petit devant 1 soit :

$$\alpha = \frac{1}{80} \text{ d'où } C_1 = \frac{1}{r \alpha \omega} = 13 \text{ picofarads.}$$

Valeur compatible avec la capacité d'entrée de la 6 A C 7.

La résistance d'amortissement correspondant à la pente la plus forte ($p = 16 \text{ mA/V}$) ($\rho = 420.000 \text{ Ohms}$) est égale à : (voir formule 19)

$$\frac{1}{\frac{1}{r} + p \alpha^2 \left(1 + \frac{k - \frac{\rho}{R'}}{k^2 \alpha^2 + \frac{\rho^2}{R'^2}} \right)} = 66.000 \text{ ohms}$$

$$\frac{1}{r} + p \alpha^2 \left(1 + \frac{k - \frac{\rho}{R'}}{k^2 \alpha^2 + \frac{\rho^2}{R'^2}} \right)$$

ce qui représente un amortissement de 2,5.

Le glissement relatif est égal pour une variation

de pente de 15 mA/V à : (voir formule 16)

$$\frac{\Delta f}{f} = \frac{L}{2 r c_1} \Delta p = 4,75 \% , \text{ soit un balayage en fréquence de } 475 \text{ Kcs.}$$

La figure 5 donne le schéma de montage avec les valeurs que nous venons d'adopter.

Remarque : Nous avons été amené à déclarer dans le précédent exposé que certaines inégalités étaient en pratique toujours exactes ? Nous pouvons le vérifier dans l'exemple ci-dessus :

$$\text{— C'est ainsi que la quantité } \left(\frac{p}{1 + \frac{\rho}{R'}} - \frac{1}{r} \right)$$

de l'équation 11 bis s'annule pour des valeurs de p égales à

$\left(1 + \frac{R'}{\rho} \right) \times \frac{1}{r}$ soit à $\frac{1}{50} \text{ mA/V}$. La quantité ci-dessus est donc toujours positive. Ce qui démontre le caractère selfique de la lampe T.

— D'autre part l'inégalité (17)

$\frac{\rho}{R'} + \frac{R'}{\rho} < k - 2$ se vérifie dans l'exemple ci-dessus :

$$\frac{420.000}{50.000} + \frac{50.000}{420.000} < 6.748$$

2^e Exemple.

Avec la même 6 A C7 prenons une fréquence moyenne de 200 Kc/s et calculons le glissement et l'amortissement :

$$f = 200 \text{ Kc/s.}$$

$$L = 400 \mu \text{ H}$$

$$C = \frac{1}{L \omega^2} = 1560 \text{ pf.}$$

$$R = Q L \omega = 100.000 \Omega$$

Soit :

$$C_1 = 13 \text{ pf}$$

$$r = 1 \text{ M}\Omega.$$

$$\alpha = \frac{1}{r C_1 \omega} = \frac{6}{100}$$

$$\frac{1}{R'} = \frac{1}{R} + \frac{1}{r} = \frac{1}{100.000} + \frac{1}{1.000.000}$$

$$R' = 90.000 \Omega$$

Calculons l'amortissement d'après (19) on trouve $R = 620.000 \Omega$.

Calculons le glissement d'après (16). On trouve

$$\frac{\Delta f}{f} = 24,5 \%$$

$$\Delta f = 49 \text{ Kc/s.}$$

On a par suite en comparant à l'exemple sur la fréquence de 10 Mcs, à la fois un glissement relatif plus important et un amortissement plus faible.

(A suivre)

L'INVENTEUR DE LA T. S. F.

PAR

L. CAHEN

Ancien Ingénieur en Chef des P. T. T.

Il est peu d'inventions importantes dont la paternité n'ait donné lieu à des revendications de priorité. Mais, pour beaucoup, l'accord s'est fait au bout d'un certain temps et n'a plus ensuite été sérieusement remis en question ; pour d'autres, au contraire, surtout celles mettant en jeu des amours-propres nationaux, les polémiques n'ont jamais cessé et, à la façon des fièvres endémiques, elles reprennent de temps en temps de la virulence.

La radiotélégraphie, a soulevé dès le début des contestations acharnées et continue à les soulever. Dans un article des « Annales télégraphiques » de mars-avril 1898, sur lequel nous reviendrons à plusieurs reprises, M. Voisenat, Ingénieur en chef des télégraphes français, le constatait en ces termes :

« La presse technique a montré comment les dispositions utilisées (par Marconi en 1896) provenaient des travaux de Hertz ; elle a enregistré les revendications présentées par divers savants qui ont découvert les propriétés utilisées et réalisé les divers appareils employés par M. Marconi ; elle a enfin reproduit les arguments présentés par celui-ci pour apprécier l'importance de son travail personnel.

« D'après certains auteurs, M. Marconi n'aurait rien inventé et son mérite, en la circonstance, aurait seulement consisté à lancer dans le grand public, au moyen d'une habile réclame et de hauts patronages des procédés entièrement connus, décrits dans des publications remontant à plusieurs années et déjà considérées comme classiques dans l'enseignement des Universités. Peut-être la polémique n'eût-elle pas atteint le caractère aigu auquel elle est arrivée si, connaissant davantage les antériorités de la question, M. Marconi n'eût revendiqué comme lui étant propres les procédés et les appareils dont il est question et s'il se fût borné à faire breveter leur application à la transmission des signaux ».

Depuis lors, on sait la campagne ardente menée en France en faveur d'Auguste Branly qui, pour beaucoup de français, est « le père de la T. S. F. ». Enfin, depuis 1925, les savants et les services publics russes ont exalté l'œuvre de leur compatriote Popov, qui est désormais officiellement, en Russie soviétique, l'inventeur de la T. S. F. En 1945, le « cinquantenaire » de l'invention de la T. S. F. par

Popov » a donné lieu à de grandioses manifestations avec présentation de nombreux témoignages de savants ayant collaboré avec Popov ou l'ayant connu. Ces attestations ont été recueillies et commentées non seulement dans la presse scientifique ou technique, mais dans celle d'informations quotidiennes ou périodiques, notamment en Angleterre et en France ; étant donné l'extrême sensibilité de l'opinion mondiale, on pense bien que l'impartialité et la sérénité n'ont pas toujours guidé les affirmations des auteurs des articles consacrés à la question.

Nous voulons examiner cette question avec une objectivité aussi complète que possible, en tirant surtout parti des faits recueillis dans deux articles du Wireless Engineer et en les confrontant avec ceux qu'apporte l'article déjà cité de M. Voisenat, qui a l'intérêt de montrer ce qui, à cette époque, était déjà venu à la connaissance des techniciens ; nous ajouterons que M. Voisenat, que nous avons connu personnellement, était un homme d'une probité intellectuelle scrupuleuse qui n'aurait jamais, sous aucun prétexte, sollicité ou déformé les informations qu'il possédait ; d'ailleurs, patriote ardent, écrivant dans les beaux jours de la lune de miel franco-russe, il aurait certainement été heureux de mettre en lumière les mérites de Popov, mais il n'aurait pourtant pas pour cela employé d'arguments ou d'informations contestables. (1)

* * *

Les articles auxquels nous nous référons signalent tous que de multiples essais de télégraphie sans fil ne faisant pas intervenir les ondes électromagnétiques avaient eu lieu avant Hertz. La télégraphie par le sol (T.P.S.) dont Lindsay avait émis l'idée dès 1831 consiste à établir entre deux points du sol une différence de potentiel qui produit des courants

(1) L'article de M. Voisenat comprend une étude des précurseurs de Marconi et le début de l'examen des appareils de ce dernier. Il prévoit une suite qui n'a jamais paru. Les Annales télégraphiques ont cessé leur publication en 1900, c'est-à-dire près de deux ans plus tard, ce qui ne peut pas expliquer cette interruption. L'Administration française d'alors a-t-elle trouvé imprudent de laisser, dans un journal dépendant d'elle, émettre des opinions de nature à entraîner des polémiques ? Il serait intéressant de le savoir et s'il n'y a pas eu intervention quelconque.

dans les différents filons conducteurs et on peut ainsi détecter entre deux autres points du sol à une certaine distance des premiers des impulsions reproduisant celles de l'origine. Boursoult, en 1871, essaya par ce moyen d'établir des relations entre Paris assiégé et des personnes de l'autre côté de l'armée assiégeante. On a relaté des essais faits par l'A. E. G. en Allemagne en 1890. La T. P. S. (avec récepteur téléphonique) a reçu une utilisation limitée dans la guerre de 1914-1918 et peut-être dans la dernière.

Les inductions électrostatiques et électromagnétiques fournissent un moyen de faire passer des impulsions de courant entre deux circuits sans liaison conductrice. Pour l'induction électrostatique, on peut citer des essais d'Edison (1885) avec des ballons à surface métallisée. Il prit également un brevet pour communications télégraphiques entre terre et bateau, les stations côtières ayant des antennes et des prises de terre, les stations sur bateau des antennes en forme de L. Le courant émetteur était interrompu à cadence rapide par un interrupteur rotatif. Le récepteur était un téléphone spécial « électromotographe » d'Edison. L'induction électromagnétique a reçu des applications plus étendues. Entre 1892 et 1895, le Post Office anglais, sous l'inspiration de son ingénieur en chef Preece, essaya des communications de ce genre pour suppléer de façon provisoire au service de câbles interrompus, à des distances variant de 3 à 7 milles anglais.

Aussi, quand les essais de Marconi eurent été connus, comme il ne s'agissait au début que de distances de même ordre, Preece et un certain nombre de techniciens soutinrent-ils que le nouveau système s'expliquait par l'induction sans intervention d'ondes. (2).

On peut encore citer comme exemple de communications de ce genre les essais d'Edison pour communications avec des trains en marche, idée qui a été depuis lors reprise fréquemment.

Il faut mentionner spécialement les expériences de Joseph Henry en 1842. Ce savant voulu se rendre compte à quelle distance on pouvait transmettre une impulsion par induction. Il plaça un circuit en fils à un étage supérieur d'un bâtiment et l'excitait par une étincelle longue d'un pouce environ émanant du conducteur primaire d'une machine électrique : le circuit secondaire ou récepteur était placé dans la cave située à 30 pieds en dessous avec intercalation de deux planchers et deux plafonds. Le courant produit était suffisamment intense pour magnétiser des aiguilles placées dans une spirale de cuivre. Il inséra aussi une spirale semblable dans un fil placé sur la toiture de son cabinet de travail et constata la magnétisation des aiguilles par des décharges atmosphériques provenant d'éclairs éloignés. Il est

vraisemblable que le fil était mis à la terre à son extrémité inférieure. Ces expériences ont donc, comme l'écrit le *Wireless Engineer*, un « parfum hertzien ». D'ailleurs, en 1851, rappelant ces expériences, Henry, bien avant les publications de Maxwell, écrivait : « Comme il s'agit de courants dont le sens alterne, ils doivent produire dans l'espace, environnant des mouvements d'un sens et de l'autre analogues, sinon identiques, à des ondulations ».

On cite enfin des essais de Hughes, en 1879, qui ne firent l'objet d'aucune publication, mais les expériences eurent lieu en présence de W. H. Preece et de Sir William Crookes qui en témoignèrent. On transmit des signaux « radioélectriques » (?) à une distance de 60 pieds. L'appareil récepteur était dans la poche et un écouteur téléphonique à l'oreille. On put recevoir des signaux jusqu'à environ 1/2 mille de distance. Une nouvelle démonstration eut lieu devant les professeurs Stokes et Huxley. On raconte que ceux-ci estimèrent que ces résultats pouvaient s'expliquer par l'induction électromagnétique normale et ne semblaient pas offrir de vastes possibilités. Hughes, découragé, renonça à poursuivre ses essais. A Fahia qui, en 1899, lui demandait de faire valoir les titres qu'il s'était acquis, il répondit : « Les expériences de Hertz étaient beaucoup plus concluantes que les miennes, quoiqu'il ait employé un récepteur beaucoup moins efficace que le mien ou le cohéreur. Je pensai alors qu'il était maintenant trop tard pour mettre en avant mes propres expériences et j'ai été forcé de laisser d'autres faire la découverte que j'avais précédemment faite en ce qui concerne la sensibilité du contact microphonique et son emploi dans la réception des ondes électriques aériennes ».



Il semble donc bien qu'Edison, et surtout Hughes, aient produit et détecté des ondes électromagnétiques mais tout au moins le premier, sans en concevoir la vraie nature. On sait que Maxwell avait signalé comme conséquence de sa théorie électromagnétique de la lumière, la possibilité de produire de telles ondes et de les propager dans les milieux diélectriques. C'est en 1888 qu'Henri Hertz exécuta la série des expériences mémorables qui établirent l'existence de telles ondes et vérifièrent qu'elles possédaient tous les caractères essentiels de l'onde lumineuse. Les longueurs d'ondes alors obtenues étaient de l'ordre de quelques mètres, la puissance de l'émetteur et la sensibilité du récepteur étaient également faibles, et Hertz ne pensa pas qu'on pût arriver par le moyen de ces ondes à envoyer des signaux à une distance pratiquement intéressante.

Les expériences de Hertz eurent naturellement beaucoup de retentissement dans les milieux savants d'Europe et d'Amérique, et on chercha dans les principaux laboratoires à les répéter et à en faciliter la démonstration, en augmentant la puissance émise, la sensibilité et la visibilité de la réception. Pour la réception, intervint en 1890 la découverte par Branly de la sensibilité aux ondes électriques des semi-conducteurs à contacts imparfaits, tels que limaille métallique, dont les variations

(2) M. Preece, ingénieur d'ailleurs actif et « efficace », se fit remarquer pendant toute sa carrière par son conservatisme technique intrinsèque et obstiné. Vers 1885, alors que les abonnés au téléphone en Amérique se comptaient déjà par dizaines de milliers, il affirmait que ce succès était dû à des circonstances spéciales aux Etats-Unis et qu'en Europe le téléphone ne ferait jamais une concurrence sérieuse au télégraphe. Un peu plus tard, il s'éleva avec violence contre les théories d'Heaviside en faveur de l'augmentation de l'inductance des lignes téléphoniques.

de conductance sous diverses actions avaient déjà été signalées par Calzecchi.

Les années 1891 et 1892 furent marquées par les expériences et les démonstrations faites devant les Sociétés d'Electriciens d'Amérique, d'Angleterre et de France par Nicolas Tesla et qui frappèrent vivement l'imagination des techniciens. A vrai dire, Tesla anticipant sur un avenir lointain, mettait surtout l'accent sur la transmission de l'énergie par les ondes plutôt que sur celle des signaux : il a montré surtout l'allumage de lampes et n'a pas insisté sur les modes spéciaux de détection télégraphique. Mais il a nettement prévu l'avenir de la T. S. F., affirmant la possibilité de liaisons transatlantiques. Tesla étudiait surtout la génération industrielle des ondes, soit par un alternateur spécial (de 10 000 à 30 000 c/s) soit par un éclateur produisant des oscillations à très haute fréquence (de l'ordre du megacycle/seconde), pour lequel il inventait son fameux transformateur appelé bobine Tesla ou simplement « Tesla » qui fut d'un usage courant dans les débuts de la T. S. F.

Tesla indiquait de plus qu'en reliant son système générateur à une plaque métallique *rayonnant* dans l'air (le mot est de lui) on augmentait l'énergie transmise et qu'on améliorait la réception en reliant l'appareil détecteur à une plaque semblable.

Malgré l'importance accordée à ces communications, il est curieux de constater qu'aucun des articles sur lesquels nous nous appuyons, l'un presque contemporain des faits, les autres tout récents, ne parlent du rôle joué par Tesla dans la naissance de la T. S. F.

Peut-être faut-il en chercher l'explication dans le fait que la mise au point des systèmes de détection, sur lequel allait se concentrer l'effort des chercheurs, ne suivit pas la marche indiquée par Tesla ; de même les plaques rayonnantes, si elles contiennent à nos yeux l'essentiel du principe de l'antenne, en diffèrent dans leur réalisation matérielle et ne semblent pas avoir contribué effectivement, à leur naissance. La conception des ondes entretenues par générateur était trop en avance sur ce temps pour avoir été retenue. Tesla a donc été un grand précurseur, mais un peu « en marge ». Toutefois on peut dire qu'après lui, la question de la T. S. F. était posée.

En tous cas, le professeur Lodge en Angleterre, entreprit, à partir de 1893, d'étudier les ondes au moyen de l'appareil Branly auquel il donna le nom de « cohéreur » en vertu d'une théorie qui n'obtint d'ailleurs pas l'assentiment de Branly. Il étudia aussi les procédés de « décoherence » c'est-à-dire de retour de l'appareil à sa faible conductibilité initiale, sans lequel toute signalisation continue eût été impossible. L'augmentation de passage du courant due aux ondes était décelée par le fonctionnement d'une sonnerie (ou un toc téléphonique).

Et c'est ici que se pose le problème Popov que nous allons maintenant étudier en détail.

Il n'est pas contesté que l'objet premier des recherches de Popov fut de déceler les phénomènes d'électricité atmosphérique, en particulier les orages. Dès 1888, Lodge avait émis l'idée que les décharges

de la foudre devaient avoir un caractère oscillatoire (à haute fréquence). Elles devaient donc relever des systèmes de détection des ondes hertziennes et le montage de Lodge pouvait leur être appliqué. Pour permettre la réception, Popov remplaça la sonnerie de Lodge par un relais fermant le circuit d'un appareil Morse.

Puisqu'il s'agissait de recevoir les décharges de la foudre, il était naturel de chercher à les collecter par l'emploi d'un paratonnerre. C'est ce que fit en effet Popov. Puis, quand il ne pouvait disposer de paratonnerre, il y substituait un fil, tiré d'abord dans le laboratoire, puis s'élevant dans l'air et soutenu par des mâts isolants : l'appareil récepteur était relié à l'extrémité inférieure de ce fil et la terre reliée à l'autre borne du système. C'est bien déjà l'antenne de Marconi que celui-ci revendique dans son brevet initial. Il est vrai que le dispositif d'Edison signalé plus haut avait déjà les caractères principaux d'une antenne. D'autre part, on signalé l'utilisation antérieure, également pour des recherches météorologiques, de dispositifs ayant quelque ressemblance, mais sans liaison avec le sol. Enfin nous avons parlé des plaques rayonnantes à l'émission et à la réception de Tesla. Néanmoins, le mérite de l'invention de l'antenne, tout au moins pour la détection des orages, ne semble pas contesté à Popov, et comme, dans les circuits de réception, le dispositif pouvait être transporté tout entier sans modification, les principes de la jurisprudence et ceux de l'équité sont d'accord pour attribuer à Popov l'antenne de T. S. F., ce qui suffit à lui donner une place éminente parmi les metteurs au point de la nouvelle technique.

Mais les défenseurs de Popov vont plus loin et affirment que Popov a non seulement prévu mais réalisé l'application de son appareil à la transmission des signaux à distance. C'est ce point sur lequel s'est concentrée la polémique et sur lequel il nous faut à présent insister.

Signalons d'abord que si Hertz avait considéré comme chimérique la transmission de signaux par ondes hertziennes, certains savants après lui n'avaient pas été de son avis. Nous avons rappelé les prédictions de Tesla qui semblent avoir été vite oubliées à cause sans doute, de leur apparente exagération et aussi parce qu'elles étudiaient peu les détecteurs qui furent ensuite adoptés. Mais en 1892 Sir William Crookes écrivait, en se référant au tube de Branly « Ainsi s'est révélée la possibilité de télégraphier sans fils, poteaux ou câbles ou tout autre ensemble de nos coûteuses applications actuelles, par des signaux du code Morse... Les opérateurs doivent accorder leurs appareils à une longueur d'onde déterminée, p. ex. 50 yards ».

D'autre part, notons que les seules publications relatives à Popov parues avant le dépôt des demandes de Marconi sont les suivantes :

Popov présenta, en mai 1895, son appareil à la Société de physique de St-Petersbourg en même temps qu'un mémoire sur l'action des oscillations électriques sur les poudres métalliques. Un résumé de sa conférence a paru dans l'*Elektrichestvo* de St-Petersbourg de 1896. Le récepteur a été décrit

dans le Journal de janvier 1896 sous le titre « Appareil pour détecter et enregistrer les oscillations électriques de l'atmosphère ». Quoique servant seulement comme indicateur d'orage, Popov exprima l'espoir qu'avec certains perfectionnements il pourrait être employé pour la réception de signaux quand on aurait découvert une source d'oscillations suffisamment forte. Il aménagea le frappeur du cohéreur de façon à lui faire actionner une sonnerie et connecta également avec le circuit un enregistreur électromagnétique à cylindre tournant sur lequel s'inscrivaient les décharges de la foudre. L'Electrician du 10 décembre 1897 publia une lettre du Popov datée du 26 novembre précédant avec la traduction d'extraits du Mémoire de Popov de janvier 1896. On lit entre autres choses dans la lettre : « De juillet 1895 jusqu'à novembre, mon appareil fonctionna dans de bonnes conditions. J'ai pu détecter des ondes électromagnétiques à une distance de 1 kilomètre en travaillant avec le vibreur de Hertz comportant des sphères de 30 cm et le relais Siemens et Halske. Avec le vibreur de Bjoerknes de 90 cm de diamètre et un relais plus sensible, j'obtins un bon fonctionnement à 5 km. Des remarques précédentes, on peut conclure que le récepteur Marconi est la reproduction de mon enregistreur de décharges de foudre ». De ces dernières phrases peut-on conclure, comme le fit récemment le professeur Ashby, que Popov a effectivement transmis des messages à distance et accuser avec lui Fleming d'avoir, dans son ouvrage classique, en rendant compte des travaux et des mémoires de Popov, oublié sciemment ces titres à l'invention de la T. S. F. ? L'article du *Wireless Engineer* de janvier 1948 proteste contre cette accusation ; il pense que Popov a entrepris seulement dans ces essais de 1896, de produire des ondes simulants des décharges de foudre pour voir à quelle distance on pouvait recevoir celles-ci.

Dans la même lettre, écrite après publication des essais de Marconi, Popov exprime d'ailleurs l'espoir que son appareil pourra servir à la signalisation à longue distance aussitôt qu'on aura trouvé un générateur plus puissant. Notons à ce sujet, qu'à part les indications sur le vibreur, Popov ne donne aucun renseignement sur son dispositif d'émission et notamment sur l'emploi d'une antenne.

Voici ce que dit à ce sujet M. Voisenat, écrivant avant février 1898 « M. Popov a d'ailleurs procédé à cette époque, au moyen de son dispositif, à des essais de correspondance télégraphique qui lui ont donné des résultats encourageants. Les conclusions du mémoire font ressortir les services que peut rendre son appareil pour l'étude des orages et pour la transmission des signaux. Ils indiquent cependant la nécessité d'y apporter quelques perfectionnements et de recourir à un générateur d'oscillations électriques d'une puissance suffisante ».

Ainsi, M. Voisenat adopte, à cette date, la thèse actuelle des « propopovistes ». Il dit même que le compte-rendu d'expériences de transmission existe dans les mémoires et non dans la lettre postérieure à l'Electrician. Il n'a sans doute pas lu les Mémoires russes, mais seulement des extraits de traduction, ce qui diminue la valeur de son affirmation. Mais

il n'en reste pas moins qu'à cette date un homme impartial, et ayant cherché à se renseigner, croit à la réalité des expériences de transmission de Popov.

En sens inverse, il faut admettre, avec le *Wireless Engineer*, que Marconi, cherchant à prendre un brevet, n'a pu publier aucun renseignement sur ses appareils et ses essais avant le dépôt de ce brevet (mai 1896), tandis que Popov, professeur public et n'ayant aucune arrière-pensée commerciale, pouvait faire connaître sans délai ses idées et ses réalisations.

Telles sont les informations relatives à Popov, qui résultent de publications antérieures ou de peu postérieures aux brevets de Marconi. Mais, beaucoup plus tard, les partisans de la revendication « popoviste » en ont apporté de nouvelles. En effet, dans une publication russe de 1945 (l'Invention de la radio par Popov) on reproduisait trois lettres écrites respectivement en décembre 1925 et janvier 1926 par trois professeurs russes qui avaient assisté aux démonstrations faites par Popov le 12 mars 1896. Ces lettres, en même temps qu'une brochure de Rycklin, publiée en 1945, soutenaient que Popov avait été empêché de publier le résultat de ses expériences par les autorités navales russes. En mars 1895, Popov fit des démonstrations de la réception d'ondes émises par un oscillateur de Hertz au moyen d'un relais et d'une frappe automatique. En mars 1896, on soutient qu'il a donné une nouvelle démonstration ; cette fois, l'émetteur était dans une autre pièce de l'Université, à environ 250 m de distance et pourvu d'une clef d'émission. A l'appareil récepteur, le tambour à rotation lente employé pour l'enregistrement des décharges avait été remplacé, en septembre 1895, par un récepteur Morse. Après une introduction explicative, l'expérience commença et quand les lettres apparurent sur la bande mobile, on les inscrivit sur un tableau noir où les spectateurs lurent d'abord les mots « Heinrich Hertz ». Rybkin, qui était assistant de Popov, avertit Popov, avant la démonstration, qu'il eût à faire très attention à ses paroles à cause de l'importance militaire de la question et, en conséquence, Popov, dans son compte-rendu, présenta seulement l'expérience qu'on vient de décrire comme un moyen de vérifier l'existence des ondes hertziennes. Lebedinsky, dans sa lettre, dit avoir conservé la bande avec les mots « Heinrich Hertz » jusqu'en 1918-1919, où elle disparut avec le reste de sa bibliothèque à Riga.

Popov était alors attaché à l'Ecole des torpilles et placé sous les ordres des autorités navales militaires. Il est cependant étonnant que — tout au moins dans les analyses qui nous en ont été transmises — on ne trouve aucune description de ces importantes expériences dans le numéro jubilaire spécial de 1925 de l'*Electrichesvo*, auquel le professeur Lebedinski a donné un article ne faisant pas mention de la transmission des mots « Heinrich Hertz » et de leur enregistrement sur bande. Dans cet article, il dit que Popov améliora l'appareil de Lodge en rendant plus régulière l'action du frappeur décohérent et en utilisant un récepteur Morse, plus pratique que le galvanomètre employé par Lodge. Il est certain que, dans une action en justice