

L'ONDE ÉLECTRIQUE

35^e ANNÉE. N° 334

JANVIER 1955

PRIX: 250 FRANCS

Library
National Laboratories
National Bureau of Standards
Boulder, Colorado

REVUE DE LA SOCIÉTÉ DES RADIOÉLECTRICIENS
ÉDITIONS CHIRON. 40, RUE DE SEINE. PARIS - 6^e



GENERA-
TEUR DE SIGNAUX
MARCONI T.F. 867 avec son
ATTENUATEUR DE SORTIE
(détail au 1^{er} plan)
Importateur : LELAND-RADIO

LIRE DANS CE NUMERO

Procédé pour la télévision en couleurs,
BOUTRY, BILLARD, LE BLANC — Télévision
en couleurs, N.T.S.C., HIRSCH — Télé-
vision aux U.S.A., MAYER — Faisceaux
hertziens, DUCOT — Signal isolé sur
circuits résonnants, J. MARIQUE —
Pletrage en reproduction sonore,
SACERDOTE, CACIOTTI,
SACERDOTE — Semi-
conducteurs, B. PIS-
TOULET...

Library
Boulder Laboratories
National Bureau of Standards
Boulder, Colorado
JUN 13 1956

6 ANNÉES D'EXPERIENCE ET DE RÉALISATIONS...



...DANS LE DOMAINE DU
CALCUL ANALOGIQUE
LINÉAIRE & NON LINÉAIRE

L'INSTRUMENT DE TRAVAIL
DE L'INGÉNIEUR MODERNE

OME L2

CALCULATEUR ANALOGIQUE
USAGE GÉNÉRAL ET DE SIMULATION



SOCIÉTÉ D'ELECTRONIQUE & D'AUTOMATISME
138, BOULEVARD DE VERDUN - COURBEVOIE (Seine) * DEF. 41-20

SIMULATEURS * CALCULATRICES NUMÉRIQUES UNIVERSELLES * APPLICATIONS INDUSTRIELLES

Régisseur exclusif de la Publicité de l'Onde Électrique : Agence-Publéditec-Domenach, 161, boulevard Saint-Germain, PARIS - 6^e. Tél. LIT. 79-53.

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	ÉDITIONS CHIRON	Prix du
FRANCE 2500 F	40, Rue de Seine — PARIS (6 ^e)	numéro :
ÉTRANGER 2800 F	C. C. P. PARIS 53-35	250 francs

Vol. XXXV JANVIER 1955 Numéro 334

SOMMAIRE

	Pages
Un procédé pour la transmission simultanée de deux programmes de télévision ou d'un spectacle en couleurs.....	G. BOUTRY P. BILLARD L. LE BLAN 5
Principes du système de télévision en couleur simultanée N.T.S.C.....	C. HIRSCH..... 22
Télévision en couleurs aux U.S.A.	C.G. MAYER 33
Procédé technique pour l'amélioration des performances des faisceaux hertziens en télé-phonie	C. DUCOT..... 41
Effet d'un signal radioélectrique isolé sur des circuits résonnants	J. MARIQUE 55
Le pleurage dans les systèmes de reproduction sonore.....	C.B. SACERDOTE M. CACIOTTI G. SACERDOTE 62
Mesure des propriétés électriques des semi-conducteurs.....	B. PISTOULET 71
Vie de la Société	77

Sur la couverture :

Générateur de signaux MARCONI T. F. 867 avec son atténuateur de sortie (détail au premier plan).
Importateur : LELAND-RADIO

Les opinions émises dans les articles ou comptes rendus publiés dans L'Onde Electrique n'engagent que leurs auteurs.

SOCIÉTÉ DES RADIOÉLECTRICIENS

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

2° 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° Service gratuit de la revue mensuelle *L'Onde Electrique* ;

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

3° Visites de diverses installations radioélectriques : stations d'émission et de réception, postes de navires et d'avions, laboratoires, radiophares, expositions, studios, etc ;

4° Renseignements divers (joindre un timbre pour la réponse).

COTISATIONS

1° Membres titulaires, particuliers	1 600 F.
— sociétés ou collectivités 10 000 F.	au gré
25 000 F.	de la société
ou 50 000 F.	ou collectivité
2° Membres titulaires, agés de moins de vingt-cinq ans, en cours d'études en France.....	800 F.
Les membres de ces deux catégories, résidant à l'étranger, doivent verser en plus un supplément pour frais postaux de 400 F.	
3° Membres à vie :	
Les particuliers, membres titulaires peuvent racheter leur cotisation annuelle par un versement unique égal à dix fois le montant de cette cotisation soit	
16 000 F.	
4° Membres donateurs :	
Seront inscrits en qualité de donateurs, les membres qui feraient don à la Société, en plus de leur cotisation, d'une somme d'au moins 5 000 F.	
5° 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	
15 000 F	

Adresser la correspondance administrative et technique, et effectuer le versement des cotisations au Secrétariat de la Société des Radioélectriciens

SOCIÉTÉ DES RADIOÉLECTRICIENS

10, Avenue Pierre-Larousse, Malakoff (Seine)

Tél. ALÉSIA 04-16 — Compte de chèques postaux Paris 697-88

CHANGEMENTS D'ADRESSE : Joindre 20 francs à toute demande

SOCIÉTÉ DES RADIOÉLECTRICIENS

PRÉSIDENTS D'HONNEUR

† R. MESNY (1947) — † H. ABRAHAM (1947)

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 VALBREUZE, 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 es-Sciences.
 1938 P. FRANCK, Ingénieur général de l'Air.
 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, Secrétaire perpétuel de l'Académie des Sciences.
 1949 M. PONTE, Directeur Général Adjoint de la Cie Gle de T.S.F.
 1950 P. BESSON, Ingénieur en Chef des Ponts et Chaussées.
 1951 Général LESCHI, Directeur des Services Techniques de la Radiodiffusion — Télévision Française.
 1952 J. de MARE, Ingénieur Conseil.
 1953 P. DAVID, Ingénieur en chef à la Marine.
 1954 G. RABUTEAU, Directeur Général de la Sté « Le Matériel Téléphonique ».

BUREAU DE LA SOCIÉTÉ

Président (1955)

M.H. PARODI, Membre de l'Institut, Professeur au Conservatoire National des Arts et Métiers.

Président désigné pour 1956 :

M.R. RIGAL, Ingénieur Général des Télécommunications.

Vice-Présidents :

MM. E. FROMY, Directeur de la Division Radioélectricité du L.C.I.E.
 A. ANGOT, Ingénieur militaire en Chef, Directeur de la Section d'Etudes et de Fabrications des Télécommunications.
 C. BEURTHERET, Ingénieur en Chef à la C. F. T. H.

Secrétaire Général :

M.J. MATRAS, Ingénieur Général des Télécommunications.

Trésorier :

M.R. CABESSA, Ingénieur à la Société L.M.T.

Secrétaires :

MM. R. CHARLET, Ingénieur des Télécommunications.
 J.M. MOULON, Ingénieur des Télécommunications
 P. DEMAN, Ingénieur des Télécommunications.

SECTIONS D'ÉTUDES

N°	Dénomination	Présidents	Secrétaires
1	Etudes générales.	Colonel ANGOT	M. TOUTAN.
2	Matériel radioélectr.	M. LIZON	M. GAMET.
3	Electro-acoustique.	M. CHAVASSE.	M. POINCELOT
4	Télévision.	M. MALLEIN	M. ANGEL
5	Hyperfréquences.	M. WARNECKE	M. GUÉNARD
6	Electronique.	M. CAZALAS.	M. PICQUENDAR
7	Documentation.	M. CAHEN.	Mme ANGEL.
8	Electronique appliq.	M. RAYMOND.	M. LARGUIER.

GROUPE DE GRENOBLE

Président. — M.-J. BENOIT, Professeur à la Faculté des Sciences de Grenoble, Directeur de la Section de Haute Fréquence à l'Institut Polytechnique de Grenoble.

Secrétaire. — M. J. MOUSSIEGT, Chef de Travaux à la Faculté des Sciences de Grenoble.

GROUPE D'ALGER

Président. — M. A. BLANC-LAPIERRE, Professeur à la Faculté des Sciences d'Alger.

Secrétaire. — M. J. SAVORNIN, Professeur à la Faculté des Sciences d'Alger.

GROUPE DE L'EST

Président. — G. GOUDET, Directeur de l'Ecole Nationale Supérieure d'Electricité et de Mécanique de Nancy.

Secrétaire. — E. GUDEFIN, Assistant à l'E. N. S. E. M.

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).

CONSEIL DE LA SOCIÉTÉ DES RADIOÉLECTRICIENS

- J. BOULIN, Ingénieur des Télécommunications à la Direction des Services Radioélectriques.
 F. CARBENAY, Ingénieur en Chef au Laboratoire National de Radio-électricité.
 G. CHEDEVILLE, Ingénieur Général des Télécommunications.
 R. FREYMAN, Professeur à la Faculté des Sciences de Rennes.
 J. MARIQUE, Secrétaire Général du C.C.R.M. à Bruxelles.
 F.H. RAYMOND, Directeur de la Société d'Electronique et d'Automatisme.
 J.L. STEINBERG, Maître de Recherches au C.N.R.S.
 L. DE VALROGER, Directeur du Département Radar-Hyperfréquences de la Cie Française Thomson-Houston.
 J. ICOLE, Ingénieur en chef des Télécommunications, Chef du Département Faisceaux-Hertziens, Direction des Lignes Souterraines à Grande Distance.
 J. LOCHARD, Lieutenant Colonel, Chef des Services Techniques du Groupe de Contrôle Radioélectrique.
 N'GUYEN THIEN CHI, Chef de Département à la Cie Gle de T.S.F., Ingénieur-Conseil Cie Industrielle des Métaux électroniques.

- MM. G. POTIER, Ingénieur à la Société « Le Matériel Téléphonique ».
 P. RIVIÈRE, Chef du Service « Multiplex » de la Sté Française Radioélectrique.
 M. SOLLIMA, Directeur du Groupe Electronique de la Cie Française Thomson-Houston.
 H. TESTEMALE, Ingénieur des Télécommunications.
 A. VIOLET, Chef de Groupe à la Sté « Le Matériel Téléphonique ».
 l'Ingénieur J. BLOESMMA, Ingénieur Radio E. S. E.
 A. CHARLES, Général du C. R., Conseiller technique à la Société Kodak-Pathé.
 A. DIDIER, Professeur au C. N. A. M.
 G. GOUDET, Directeur de l'Ecole Nationale Supérieure d'Electricité et de Mécanique de Nancy.
 M. D. INDJOUJIAN, Ingénieur des Télécommunications, chargé du Département Télécommande du C. N. E. T.
 P. MANDEL, Ingénieur en Chef à la Société Nouvelle de l'outillage R.B.V. et de La Radio-Industrie.
 A. PAGES, Ingénieur à la Compagnie Générale de T. S. F.
 H. TANTER, Ingénieur au L. C. T.

RÉSUMÉ DES ARTICLES

UN PROCÉDÉ POUR LA TRANSMISSION SIMULTANÉE DE DEUX PROGRAMMES OU D'UN SPECTACLE EN COULEURS, par M. BOUTRY, BILLARD et LE BLAN, *Laboratoires d'Electronique et de Physique appliquée*. *Onde Electrique* de Janvier 1955 (pages 5 à 21).

L'exposé débute par un rappel succinct du mode de transmission de deux informations au moyen d'un multiplex à impulsions modulées en amplitude, étendu au cas où la bande de chaque information couvre la bande totale offerte à la transmission. On obtient ainsi à la station réceptrice une reproduction séquentielle des informations, qui peut être employée en télévision pour la transmission de deux images reproduites par « séquences de points ».

La suite de l'exposé traite de l'application directe au doublement de la capacité d'une liaison vidéo par faisceau hertzien.

Enfin, le signal de multiplex servant d'agent de transmission peut être traité, à la réception, d'une manière simple conduisant à la séparation des informations par l'emploi de diodes. La simplicité du procédé conduit, si l'on ne prend pas de précautions, à certains défauts, et l'on présente une méthode générale susceptible d'y remédier.

TÉLÉVISION EN COULEUR N.T.S.C. par M. HIRSCH. *Onde Electrique* de Janvier 1955 (pages 22 à 32).

Le système de télévision en couleur N.T.S.C. produit les images en couleur par transmission simultanée d'une porteuse transmettant la brillance et d'une sous-porteuse transmettant la couleur. Le signal de brillance produit l'image habituelle en noir et blanc. Le signal de couleur (chrominance) est utilisé à la manière d'un pinceau pour colorier l'image « noir et blanc ». Le signal de couleur ne demande qu'une bande de fréquence étroite (0,5 — 1,3 Mc/s) parce que l'œil n'est pas sensible à la couleur des détails fins. Autrement dit, le pinceau utilisé peut être large. Les teintes sont produites par les angles de phase correspondants de la porteuse de couleur. La pureté de la couleur est produite par les changements dans l'amplitude de cette porteuse, le blanc correspondant à l'amplitude zéro de la sous-porteuse de couleur.

Dans les parties non colorisées de l'image, le signal complet ne diffère donc en rien du signal en noir et blanc parce que le signal complémentaire de couleur disparaît alors.

Le signal de couleur emprunte la même bande que le signal en noir et blanc, son spectre de fréquences venant s'imbriquer dans les portions vides du spectre du signal noir et blanc.

TÉLÉVISION EN COULEUR AUX ÉTATS-UNIS, par C. G. MAYER. *Onde Electrique* de Janvier 1955 (pages 33 à 40).

L'auteur examine rapidement le problème technique fondamental, qui consistait à établir des caractéristiques de transmission s'adaptant à la largeur de bande actuelle de 6 Mc/s des canaux de télévision et satisfaisant aux exigences de compatibilité et d'utilisation optimale du spectre de fréquence. On montre que la solution adoptée a été rendue possible par l'utilisation des propriétés particulières de la vue en ce qui concerne la couleur et par la séparation des signaux issus de la caméra de couleur, de telle façon que les composantes de couleur et de brillance soient transmises simultanément à l'intérieur de la bande de 6 Mc/s. Enfin, l'article décrit les organes d'extrémité dont le développement était indispensable pour que la télévision en couleur devienne une réalité.

PROCÉDÉ TECHNIQUE POUR L'AMÉLIORATION DES PERFORMANCES DES FAISCEAUX HERTZIENS EN TÉLÉPHONIE, par C. DUCOT, *(Laboratoires d'Electronique et de Physique appliquées ; Paris)*. *Onde Electrique* de Janvier 1955 (pages 41 à 54).

On sait que l'utilisation de la modulation de fréquence est le seul moyen de transmettre sur les faisceaux hertziens avec une vitesse suffisante l'information des signaux téléphoniques multiplex à nombre moyen ou élevé de voies et des signaux de télévision. Néanmoins ce procédé présente à l'heure actuelle des limitations, notamment en téléphonie, en raison des bruits de fluctuation et d'intermodulation.

L'auteur commence par analyser à ce point de vue les recommandations du C.C.I.F. Il souligne la nécessité de considérer solidairement les effets de ces deux sortes de bruits pour définir la qualité d'un système.

Il envisage ensuite l'emploi d'un système à double modulation de fréquence, dont les caractéristiques sont étudiées. Parmi ces caractéristiques, la surpuissance à l'émission, rendue possible par l'emploi d'un tube oscillateur à réflexions multiples, s'avère intéressante.

Puis l'auteur procède à une comparaison quantitative entre les propriétés des systèmes à simple et à double modulation de fréquence.

Il indique quelques résultats expérimentaux à l'appui des considérations précédentes.

EFFET D'UN SIGNAL RADIOÉLECTRIQUE ISOLÉ SUR DES CIRCUITS RÉSONNANTS, par J. MARIQUE, Ingénieur A.I. Br et Radio E.S.E., Secrétaire général du C.C.R.M. *Onde Electrique* de Janvier 1955 (pages 55 à 61).

Recherche des amplitudes maxima du courant dû à un signal de haute fréquence dans un circuit résonnant simple, et dans un système amplificateur comportant trois circuits identiques. Influence de la durée du signal et de l'écart entre la fréquence du signal et la fréquence de résonance du système. La notion de durée d'établissement du courant au sens du C.C.I.R. ne se raccorde facilement à la constante de temps du système que quand la fréquence du signal et la fréquence de résonance sont égales.

LE PLEURAGE DANS LES SYSTÈMES D'ENREGISTREMENT SONORE, par Ing. Cesarina BORDONE SACERDOTE (*Istituto Elettrotecnico Naz. Torino*) et Ing. Mario CACIOTTI (*Radiotelevisione Italiana*) et Prof. Gino SACERDOTE (*Istituto Elettrotecnico Nazionale, Torino*). *Onde Electrique* de Janvier 1955 (pages 62 à 70).

On rappelle les méthodes de mesure du pleurage dans l'enregistrement sonore et parmi ses causes on étudie particulièrement l'élasticité des bandes magnétiques. On examine le pleurage dans les réenregistrements répétés au point de vue théorique et expérimental. Par des mesures subjectives on recherche les limites de perception du pleurage dans la reproduction de la musique.

MESURE DES PROPRIÉTÉS ÉLECTRIQUES DES SEMI-CONDUCTEURS, par B. PISTOULET, Docteur ès-Sciences. *Onde Electrique* de Janvier 1955 (pages 82 à 87).

Les principales mesures effectuées sur les matériaux semi-conducteurs concernent la résistivité, la constante de Hall, et la durée de vie des porteurs minoritaires. Pour les deux premières on étudie plus particulièrement la méthode dite des pointes, rapide, précise, et qui évite le découpage des échantillons. On décrit l'appareillage réalisé.

PAPERS SUMMARIES

TECHNICAL ARRANGEMENTS FOR IMPROVING THE PERFORMANCE OF RADIO LINKS, by C. DUCOT (*Laboratoires d'Electronique et de Physique appliquées ; Paris*). *Onde Electrique*, January 1955 (pages 41 to 54).

It is known that frequency modulation provides the only means on radio link transmission to give sufficient speed for conveying the information in multiplex telephone systems of moderate or large number of channels, or for television transmission. Nevertheless at the present time these systems are limited especially for telephony, by resistance noise and intermodulation.

The author commences by analysing the C.C.I.F. recommendations from this point of view. The necessity for examining the effects of both these types of noise is emphasised in defining the quality of a system.

A double frequency modulation system is then considered and the characteristics studied. Among these characteristics it is interesting to note the increase of transmitted power, made possible by the use of multiple reflex oscillator tube.

A quantitative comparison of the properties of simple and double frequency modulation systems is then given.

Several experimental results arising from the above considerations are included.

EFFECT OF A TRANSIENT RADIO SIGNAL ON RESONANT CIRCUITS, by J. MARIQUE, *Ingénieur A.I. Br et Radio E.S.E., Secrétaire Général du C.C.R.M.* *Onde Electrique*, January 1955 (pages 55 to 61).

This is a study of the maximum current amplitudes resulting from a high frequency signal in a simple resonant circuit and in an amplifier system comprising three identical circuits. The influence of signal duration and of the difference between signal frequency and the resonant frequency of the system is shown. The idea of time of rise of current, in the C.C.I.R. sense, is easily reconcilable with the time constant of the system only when the signal frequency is the same as the resonant frequency.

PRINTING EFFECTS IN SOUND RECORDING SYSTEMS, by Ing. CESARINA BORDONE SACERDOTE, (*Istituto Elettrotecnico Naz. Torino*) and Ing. MARIO CACIOTTI (*Radiotelevisione Italiana*) and Prof. GINO SACERDOTE (*Istituto Elettrotecnico Nazionale, Torino*). *Onde Electrique*, January 1955 (pages 62 to 70).

Methods of measuring printing in sound recording systems are recalled and among the causes the elasticity of magnetic tape is considered. Printing on repeated recordings (dubbings) is examined from a theoretical and experimental standpoint. By subjective measurements the limits of audibility of printing in music reproduction are established.

MEASUREMENT OF ELECTRICAL PROPERTIES OF SEMI-CONDUCTORS, by B. PISTOULET, *Docteur ès-Sciences*. *Onde Electrique*, January 1955 (pages 82 to 87).

The basic measurements made on semi-conductor materials refer to resistivity, Hall constant and lifetime of minority carriers. For the first two ones, the so-called « four probes method » is more extensively described. This method is quick and accurate, and avoids slicing into samples. The purposely designed measuring equipment is described.

A SYSTEM FOR THE TRANSMISSION OF EITHER TWO PROGRAMMES SIMULTANEOUSLY OR A COLOUR PROGRAMME, by M. BOUTRY, BILLARD and LE BLAN, *Laboratoires d'Electronique et de Physique appliquées*. *Onde Electrique*, January 1955 (pages 5 to 21).

The article commences with a short recapitulation of the methods of transmitting two pieces of information by means of multiplexing a pulse amplitude signal, extended to the case where the frequency band of each signal covers the whole band transmitted. In this method the receiving station receives the two signals sequentially, and the method may be applied to television for the transmission of two programme signals produced by a dual sequential process.

Following on this the article deals with the direct application of doubling the capacity of a television radio link.

Finally the multiplex signal which carries this transmission can be dealt with at the receiver by a simple method using diodes for the separation of the signals. The simplicity of the method leads to certain faults if precautions are not taken, and a general method of remedying such faults is given.

N.T.S.C. COLOUR TELEVISION SYSTEM, by M. HIRSCH. *Onde Electrique*, January 1955 (pages 22 to 32).

The N.T.S.C. colour system produces a colour signal by the simultaneous transmission of a luminance carrier and a colour sub-carrier. The luminance signal corresponds to the normal black and white. The colour (chrominance) signal is used like a brush to add colour to the black and white picture. The colour signal requires only a narrow band of frequencies (0.5 — 1.3 Mc/s) as the eye is not sensitive to fine detail in colour. In other words a coarse brush is used. The tones are reproduced by the phase angles of the colour carrier. The degree of saturation of the colour is reproduced by the changes in magnitude of this carrier, white corresponding to zero amplitude of the colour sub-carrier.

In colourless parts of the picture, the complete signal is no different from the black and white signal because the colour signal disappears in this case.

The colour signal is included in the same frequency band as the black and white signal, as its spectrum is interleaved with the black and white signal spectrum.

COLOUR TELEVISION IN THE UNITED STATES OF AMERICA, by C.G. MAYER. *Onde Electrique*, January 1955 (pages 33 to 40).

The author considers briefly the fundamental technical problem, which is to establish the transmission characteristics applicable to a television channel of effective bandwidth of 6 Mc/s, and which satisfy the requirements of compatibility and efficient use of frequency band. It is shown that the solution adopted has been made possible by exploiting the special features of colour vision, and by separating the signals produced by the colour camera in such a way that the colour and luminance components are transmitted simultaneously within the 6 Mc/s band. Finally the article describes the receiving apparatus upon whose development the reality of colour television is dependent.

LE COMPORTEMENT DE L'ŒIL ET LA TRANSMISSION RATIONNELLE D'IMAGES SIMULTANÉES

III. — UN PROCÉDÉ POUR LA TRANSMISSION SIMULTANÉE DE DEUX PROGRAMMES OU D'UN SPECTACLE EN COULEURS

PAR

MM. G.-A. BOUTRY, P. BILLARD, L. LE BLAN

*Laboratoires d'Electronique
et de Physique Appliquées*

I. — INTRODUCTION ET GÉNÉRALITÉS.

A. — Type de reproduction envisagée et mode opératoire.

Cette partie de l'exposé est consacrée à la description d'un procédé qui s'appuie sur la possibilité de reproduire une image télévisée par « séquences de points ». Nous ne reviendrons pas ici sur le principe d'une telle reproduction, dont le lecteur pourra trouver une description, dans le cas de transmission d'un programme unique, dans la littérature technique (1).

Comme il ressort de la deuxième partie de cet exposé, le fait qu'une reproduction à « séquences de points », conduise, pendant une trame, à transmettre seulement la moitié de l'information contenue dans cette trame, permet d'envisager d'écouler, dans la bande video normale du noir et blanc, deux programmes distincts, au même standard que le noir et blanc conventionnel.

On peut montrer (2) que la reproduction par séquences de points, telle qu'elle a déjà été considérée pour un programme unique (1) peut être définie comme une transmission séquentielle binaire de vitesse maximum compatible avec la bande passante. On peut montrer également que les « signaux de commutation » utilisés dans les transmissions séquentielles ordinaires (trames ou lignes) pour ouvrir ou fermer la ligne unique de transmission à chacun des deux signaux successivement écoulés, doivent avoir ici une forme très précise, qui est analogue à celle des fonctions de sondage d'un multiplex binaire à impulsions modulées en amplitude.

Soit donc $f_1(t)$ et $f_2(t)$ les deux informations à transmettre. Comme nous supposons qu'il s'agit de signaux video, nous nous permettrons de limiter leurs variations à l'intervalle (0,1), 0 correspondant au noir et 1 au blanc. Ces deux informations, correspondant en général au même standard de balayage, couvrent la même bande $(0, \omega_D)$.

Soit une fréquence ω_s un peu supérieure à ω_D , les « fonctions de commutation », que nous appellerons « fonctions de sondage » par analogie avec la technique des multiplex, s'écrivent :

$$(1) \quad \begin{cases} g_1 = 1 + 2 \cos \omega_s t \\ g_2 = 1 - 2 \cos \omega_s t \end{cases}$$

On fabrique le signal :

$$(2) \quad \Sigma = [f_1(t) g_1(t) \pm f_2(t) g_2(t)]_{\omega_s}$$

l'indice ω_s indiquant que la bande de Σ est limitée à ω_s par un filtre passe-bas présentant un gain 1/2 à cette fréquence.

Ce signal est utilisé pour assurer la transmission. A la station réceptrice, on fabrique les quantités suivantes, que nous appelons produits de réception.

$$(3) \quad \pi_1 = [g_1 \Sigma]_{\omega_s} \quad \text{et} \quad \pi_2 = [g_2 \Sigma]_{\omega_s}$$

π_1 redonne $f_1(t)$ seule, mais affectée d'une « structure de points » qui correspond au caractère séquentiel de la transmission. De même π_2 donne $f_2(t)$ seule, et affectée également d'une « structure de points ».

Nous allons montrer au paragraphe suivant comment ces résultats sont obtenus.

(1) Wilson BOOTHROYD, *Electronics* 22 N° 12, Déc. 1949, p. 88.

(2) Travaux non publiés actuellement.

B. — Détails de la transmission.

Désignons par F_1 et F_2 les spectres des signaux à transmettre $f_1(t)$ et $f_2(t)$ et soit, par exemple, $f_1 g_1$:

$$(1) \quad f_1 g_1 = f_1 + 2 f_1 \cos \omega_s t$$

La quantité $2 f_1 \cos \omega_s t$ présente un spectre symétrique par rapport à la fréquence ω_s , et nous désignerons par $M(F_1)$ la bande latérale inférieure de ce spectre, qui s'étend de $(\omega_s - \omega_p)$ à ω_s . M peut être considéré ici comme un opérateur.

Pour une composante sinusoïdale de f_1 , $\cos pt$ par exemple, l'opération « M » fournit $\cos(\omega_s - p)t$. Nous aurons besoin par la suite de connaître le résultat de l'opération « $M \times M$ ». Il est clair que si l'on fait subir une nouvelle fois à $\cos pt$ l'opération, on obtient $\cos pt$. On a donc la relation :

$$(5) \quad MM(F) = F$$

On peut donc écrire :

$$(6) \quad \begin{cases} [f_1 g_1]_{\omega_s} = F_1 + M(F_1) \\ [f_2 g_2]_{\omega_s} = F_2 - M(F_2) \end{cases}$$

et, en se bornant pour l'instant dans (2) au signe + :

$$(7) \quad \Sigma = F_1 + F_2 + M(F_1) - M(F_2)$$

La répartition spectrale de Σ est représentée figure 1. En trait plein on montre la limite des

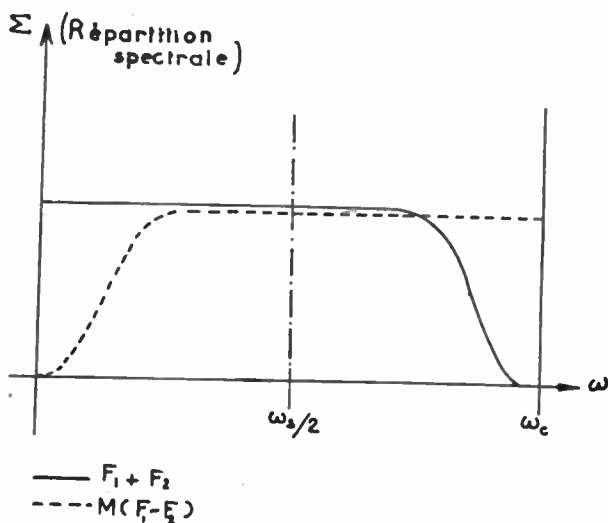


FIG. 1.

spectres non transposés F_1 et F_2 , et en pointillé la limite des spectres transposés $M(F_1)$ et $M(F_2)$.

On pourra noter que la fréquence $\omega_s/2$ constitue un axe de symétrie pour la représentation spectrale de Σ , symétrie sans changement de signe pour $[f_1 g_1]_{\omega_s}$ et avec changement de signe pour $[f_2 g_2]_{\omega_s}$.

A partir de (7), on peut calculer les produits de réception π et π_2 . On trouve ainsi facilement :

$$(8) \quad \begin{aligned} \pi_1 &= [F_1 + F_2 + M(F_1) - M(F_2)] + [M(F_1) \\ &\quad + M(F_2) + MM(F_1) - MM(F_2)] \\ &\longrightarrow = 2[F_1 + M(F_1)] \end{aligned}$$

On aurait de même :

$$(9) \quad \pi_2 = 2[F_2 - M(F_2)]$$

On voit ainsi que π_1 ne dépend que de $f_1(t)$, et π_2 seulement de $f_2(t)$. Cependant, π_1 , par exemple, n'est pas seulement fonction de F_1 , mais aussi de $M(F_1)$. Nous verrons au cours du Chapitre II que $M(F_1)$ dans π_1 , et $M(F_2)$ dans π_2 , sont les termes qui amènent les « structures de point » sur les signaux disponibles pour l'attaque des kinescopes reproducteurs. Nous allons montrer dans le paragraphe suivant que ces termes disparaissent si l'on considère une image complète.

C. — Reproduction correcte en quatre trames.

Considérons la trame N° n . Dans le mode habituel de reproduction pour la télévision en noir et blanc, les trames N° n et $n + 1$ par exemple fournissent une image complète. Si ω_l est la fréquence de trames, on peut alors définir une fréquence d'images :

$$(10) \quad \omega_l = \frac{\omega_l}{2}$$

Pour le type de reproduction étudié ici, reproduction séquentielle binaire, la moitié seulement de l'information de chaque trame est transmise pendant la durée de balayage de celle-ci. Il en résulte que le complément de l'information transmise sur une voie pendant le balayage de la trame N° n devra être transmis pendant le balayage de la trame N° $n + 2$, dont les lignes se superposent à celles de la trame N° n .

Si nous prenons comme origine des temps, pour chaque trame, le début du balayage de cette trame, le résultat précédent sera obtenu si, exprimées en fonction de cette origine, les fonctions de sondage pour la trame N° n et la trame N° $n + 2$ sont « entrelacées », ou, en d'autres termes, si l'on a :

a) Trame N° n

$$(11) \quad \begin{cases} g_1 = 1 + 2 \cos \omega_s t \\ g_2 = 1 - 2 \cos \omega_s t \end{cases}$$

b) Trame N° $n + 2$

$$(12) \quad \begin{cases} g_1 = 1 - 2 \cos \omega_s t \\ g_2 = 1 + 2 \cos \omega_s t \end{cases}$$

Si l'on considère, par exemple, la suite des maxima de g_1 suivant (11), et cette suite suivant (12), on constate que les deux suites sont exactement entrelacées.

Pour obtenir les relations précédentes, on peut facilement démontrer qu'il suffit d'avoir :

$$(13) \quad \omega_s = \frac{2k+1}{2} \omega_i$$

En pratique, et pour des raisons de commodité, on s'impose une condition plus restrictive :

$$(14) \quad \omega_s = \frac{2k+1}{2} \omega_l$$

où ω_l est la fréquence du balayage ligne. Comme ω_l est un multiple impair de ω_i , la condition (14) n'est qu'un cas particulier de (13).

Considérons maintenant les produits de réception π_1 et π_2 :

1° Cadre N° n .

On a vu, en (8) et (9) que l'on avait, à un coefficient constant près :

$$(15) \quad \begin{cases} \pi_1 = F_1 + M(F_1) \\ \pi_2 = F_2 - M(F_2) \end{cases}$$

2° Cadre N° $n + 2$.

Par suite du passage de (11) à (12), on a, pendant ce cadre :

$$(16) \quad \begin{cases} \pi_1 = F_1 - M(F_1) \\ \pi_2 = F_2 + M(F_2) \end{cases}$$

Désignons par $T_i = \frac{2\pi}{\omega_l}$ la période d'image conventionnelle. En additionnant (15) et (16), il vient :

$$(17) \quad \begin{cases} \pi_1(t) + \pi_1(t + T_i) = 2F_1 \\ \pi_2(t) + \pi_2(t + T_i) = 2F_2 \end{cases}$$

Il en résulte que, compte tenu de la persistance des impressions rétinienne, la succession de 4 trames fournit une image contenant autant d'informations que l'image fournie, en 2 trames, par les transmissions continues assurant actuellement les radio-diffusions en noir et blanc.

D. — Notion de « multiplex séquentiel ».

Examinons les propriétés d'une transmission multiplex binaire normale à impulsions modulées en amplitude, dont la fréquence de sondage soit ω_s .

La théorie de ce type de multiplex montre que la bande occupée par la transmission doit être $(0, \omega_s)$ et que la bande de chacune des informations transmises $f_1(t)$ et $f_2(t)$ ne doit pas excéder $(0, \omega_s/2)$.

L'étude spectrale faite dans le § B demeure valable, mais la figure 1 est à remplacer par la figure 2, les spectres normaux et transposés ne se recouvrant plus maintenant dans Σ .

À la réception, on désire maintenant récupérer des informations qui n'excèdent pas la bande $(0, \omega_s/2)$.

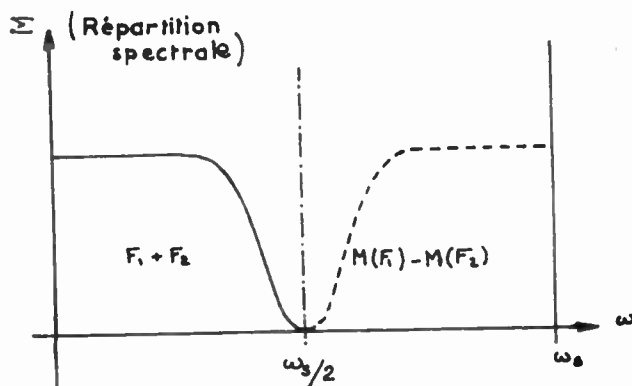


FIG. 2.

On remplacera donc la définition (3) des produits π_1 et π_2 par :

$$(18) \quad \pi_1 = [g_1 \Sigma]_{\omega_s/2} \quad \text{et} \quad \pi_2 = [g_2 \Sigma]_{\omega_s/2}$$

Si l'on se reporte alors à (16), qui donne π_1 et π_2 dans la bande $(0, \omega_s)$, on voit que l'on aura d'après (18) :

$$(19) \quad \begin{cases} \pi_1 = F_1 \\ \pi_2 = F_2 \end{cases}$$

car ici, les spectres F_1 et F_2 sont entièrement contenus dans la bande $(0, \omega_s/2)$ tandis que $M(F_1)$ et $M(F_2)$ sont en dehors de cette bande.

En conclusion, s'il est vrai qu'une transmission multiplex binaire normale présente un caractère séquentiel, ce caractère n'intéresse que la transmission elle-même, la reproduction étant continue.

Pour l'application de multiplex binaire définie au § A, la bande de chacune des informations transmises $f_1(t)$ et $f_2(t)$ excède l'intervalle $(0, \omega_s/2)$ et il en résulte un caractère séquentiel pour la reproduction elle-même.

Pour distinguer commodément cette application de la technique des multiplex à impulsions modulées en amplitude, des multiplex binaires classiques à reproduction continue, nous l'appellerons « multiplex binaire séquentiel ».

II. PROPRIÉTÉS GÉNÉRALES DES SIGNAUX ET DISPOSITIONS PARTICULIÈRES.

A. — Génération des signaux.

La figure 3 donne une vue schématique de l'installation permettant d'obtenir le signal Σ à partir des deux informations à transmettre f_1 et f_2 , que nous supposons toujours être des signaux video indépendants. Pour la commodité de l'exposé, nous

$$(2) \quad \begin{cases} g_1 = 1 + 2 \cos \omega_s t + 2 \cos 2 \omega_s t \\ g_2 = 1 - 2 \cos \omega_s t + 2 \cos 2 \omega_s t \end{cases}$$

On aura intérêt à procéder de la sorte dans les installations de haute qualité, où l'on veut obtenir le maximum de définition pour un encombrement donné de l'éther. Toutefois, dans la suite de cet exposé, nous nous bornerons à l'utilisation des fonctions de sondage (1) et à celle des filtres de la figure 4.

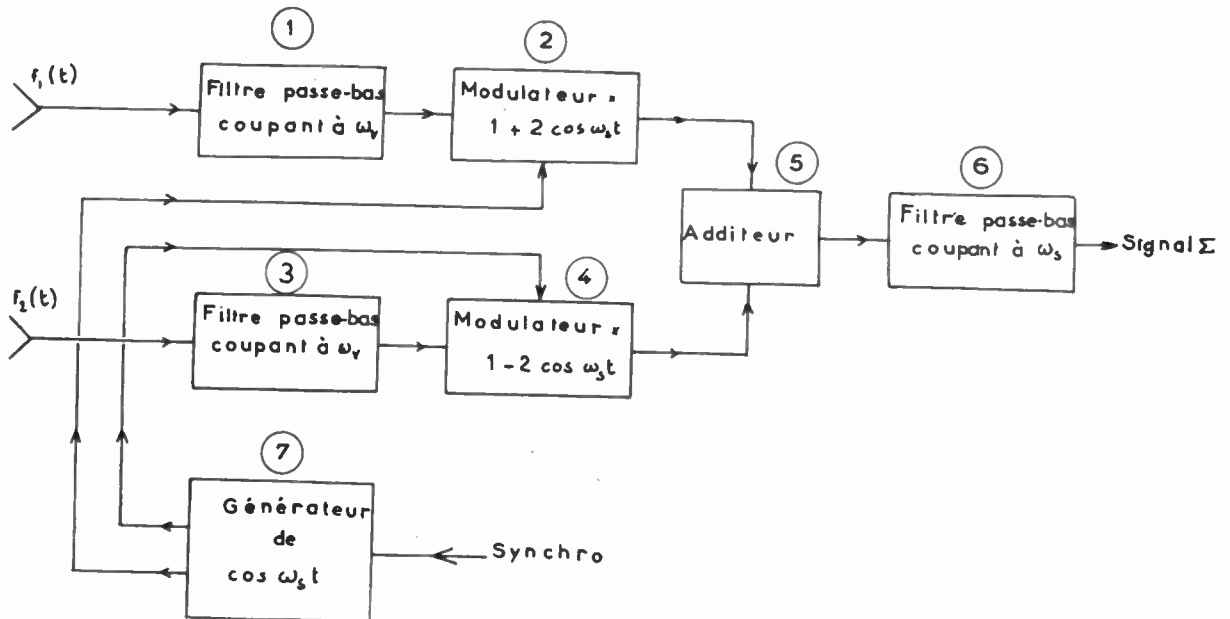


FIG. 3.

supposerons toujours que ces signaux sont positifs, chacun étant variable entre 0 et 1. Le niveau 0 correspondra arbitrairement au noir, et le niveau 1 au blanc.

Pour utiliser au mieux la bande allouée au canal de transmission, on choisira maintenant pour les filtres réels (1), (3) et (6) intervenant dans l'installation, filtres ayant nécessairement une coupure progressive, une disposition des courbes d'atténuation conforme à la figure 4.

La fréquence de coupure du filtre limitant Σ est la fréquence de sondage ω_s . Les filtres limitant f_1 et f_2 doivent présenter un gain pratiquement nul pour la fréquence pour laquelle l'atténuation du filtre de Σ apparaît. Du moins doit-il en être ainsi si les fonctions de sondage sont :

$$(1) \quad \begin{cases} g_1 = 1 + 2 \cos \omega_s t \\ g_2 = 1 - 2 \cos \omega_s t \end{cases}$$

On peut laisser transmettre ces filtres jusqu'à la fréquence ω_s ou même un peu au-dessus, si l'on fait intervenir dans les fonctions de sondage le second harmonique :

Sur la figure 3 nous avons représenté en (5) un système adducteur permettant la formation de Σ suivant l'équation :

$$(3) \quad \Sigma = [f_1 g_1 + f_2 g_2]_{\cos}$$

Cependant, on pourrait envisager aussi bien au même endroit un système soustracteur donnant :

$$(4) \quad \Sigma' = [f_1 g_1 - f_2 g_2]_{\cos}$$

En pratique, les propriétés essentielles de Σ et Σ' sont les mêmes. En particulier, les méthodes de réception déjà vues conduisent toujours à la séparation des signaux. Toutefois, au lieu de conduire, comme pour Σ à f_1 et f_2 , elles donnent pour Σ' , f_1 et $-f_2$, puisqu'aussi bien on peut dire que Σ' et Σ ne se différencient que par le changement de f_2 en $-f_2$.

Cependant, dans le cas particulier qui nous intéresse ici, pour lequel f_1 et f_2 sont toujours positives, les formes des signaux Σ et Σ' sont très différentes, et il en résulte que Σ' présente des propriétés spéciales dont la suite de l'exposé décrira quelques modes d'utilisation.

Sur la figure 3, nous avons représenté une entrée de synchronisation à l'oscillateur qui fournit les fréquences pilotes des modulateurs. Cette synchronisation a pour but d'établir la relation entre ω_s et la fréquence d'image ω_i , dont nous avons rappelé la nécessité au Chap. 1 § C.

B. — Propriétés des produits composants.

Nous entendons par « produits composants » de Σ chacune des quantités $[f_1 g_1]_{\omega_s}$ et $[f_2 g_2]_{\omega_s}$. Nous allons donner une description de l'un d'eux, $[f_1 g_1]_{\omega_s}$ par exemple.

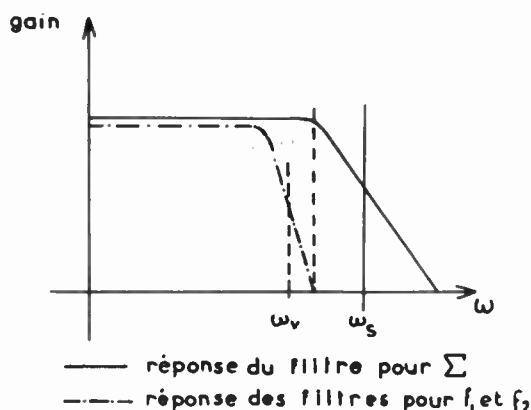


FIG. 4.

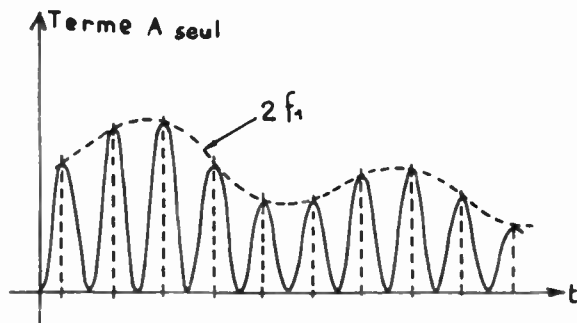


FIG. 5.

Pour ce faire, nous admettrons pour f_1 une variation quelconque, représentable par l'intégrale de Fourier :

$$(5) \quad f_1(t) = \int_0^{\omega_v} [x(\omega) \cos \omega t + y(\omega) \sin \omega t] d\omega$$

Il vient :

$$(6) \quad [f_1 g_1]_{\omega_s} = f_1(t) [1 + \cos \omega_s t] + \dots \rightarrow + \sin \omega_s t \int_0^{\omega_v} [x \sin \omega t - y \cos \omega t] d\omega$$

Le produit étudié apparaît ainsi comme la somme de deux termes :

$$(7) \quad A = f_1 [1 + \cos \omega_s t]$$

$$(8) \quad B = \sin \omega_s t \int_0^{\omega_v} [x \sin \omega t - y \cos \omega t] d\omega$$

En se livrant à des calculs numériques, on trouve que, si le premier terme peut osciller entre 0 et 2 fois le maximum de f_1 , le second n'atteint au maximum usuellement que 30 % environ du premier. De plus, comme c'est un terme en $\sin \omega_s t$, s'il peut perturber de manière notable la pseudo-phase de A , il altérera dans une proportion moindre son amplitude.

Il en résulte que pour une description « grosso-modo » de $[f_1 g_1]_{\omega_s}$ on pourra se contenter d'abord du terme A représenté figure 5. Le terme A donne l'image idéale d'un signal à « structure de points ». Il se présente comme une pseudosinusoïde de fréquence ω_s , ayant pour enveloppe inférieure l'axe des abscisses, et pour enveloppe supérieure $2 f_1$.

Pour les valeurs du temps :

$$(9) \quad t = \frac{2 k \pi}{\omega_s}$$

A est égal à $2 f_1$, et tangent à la courbe $2 f_1$.

Pour :

$$(10) \quad t = \frac{(2 k + 1) \pi}{\omega_s}$$

A est égal à 0 et tangent à l'axe des abscisses (minima). Pour les valeurs du temps (9) et (10), B est nul, de telle sorte que, pour les valeurs (9), $[f_1 g_1]_{\omega_s}$ demeure égal à $2 f_1$, et pour les valeurs (10), à 0. Ce sont là des propriétés fondamentales très générales pour les produits de sondage composant les signaux de multiplex à impulsions modulées en amplitude, et qui, pour un multiplex binaire, s'énoncent ainsi :

a) Pour les valeurs du temps correspondant aux impulsions de sondage du produit étudié, ce dernier est égal (ou proportionnel) à la fonction sondée.

b) Pour les valeurs du temps exactement entrelacées avec les précédentes, le produit est nul.

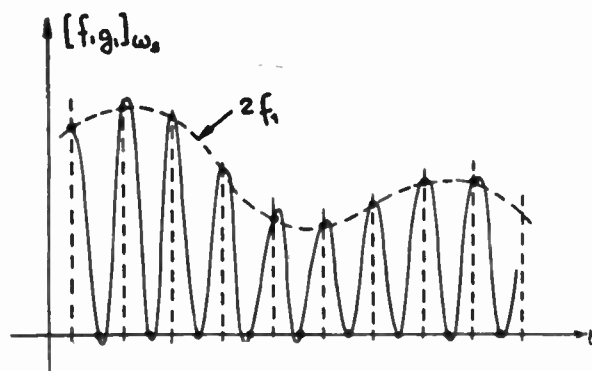


FIG. 6.

Mais pour les valeurs du temps (9) et (10) la dérivée de B par rapport au temps n'est pas en général nulle,

aussi les propriétés de contact avec l'axe des abscisses ou la courbe $2f_1$, établies pour A , ne sont-elles plus valables pour $[f_1 g_1]_{\omega_s}$.

La figure 6 donne l'allure de ce produit. On peut dire que l'allure générale reste celle de A . Mais les zéros ne sont plus exactement des minima, tandis qu'aux points sur la courbe $2f_1$, il n'y a plus contact, et ces points ne sont pas davantage des maxima.

A titre d'exemple, nous indiquerons que les calculs numériques montrent que dans le cas où la variation de f_1 est purement sinusoïdale, les minima, au-dessous de l'axe des abscisses, peuvent atteindre 12,5 % de l'amplitude maximum du signal au-dessus de cet axe. Et dans le cas où f_1 est un « signal unité », les minima peuvent atteindre 18 % de l'amplitude maximum du signal au-dessus de l'axe. Des différences relatives du même ordre de grandeur peuvent être observées entre la valeur des maxima et la valeur simultanée de $2f_1$.

Dans la suite, nous désignerons par « points de sondage », pour un produit tel que $[f_1 g_1]_{\omega_s}$, les points obtenus pour les valeurs (9) du temps, ou valeurs du temps pour les impulsions de sondage fournissant ce produit.

Nous appellerons « zéros » les points correspondant aux valeurs du temps exactement entrelacées avec les précédentes.

Nous avons étudié jusqu'à présent le produit $f_1 g_1$ seul. Il est clair que des conclusions identiques résulteraient de l'étude de $f_2 g_2$. Il convient, cependant, d'ajouter une propriété fondamentale, découlant de « l'entrelacement » des impulsions de sondage pour f_1 et f_2 , à savoir que les « points de sondage » de $[f_1 g_1]_{\omega_s}$ sont obtenus pour des valeurs de t qui donnent les « points de sondage » de $[f_2 g_2]_{\omega_s}$.

Signalons enfin que si, f_1 , par exemple, se réduit à une constante, le produit $[f_1 g_1]_{\omega_s}$ prend une forme très simple, qui est une sinusoïde tangente à l'axe des temps. Dans ce cas les « zéros » sont exactement des minima, et les « points de sondage » des maxima.

C. — Propriétés des signaux Σ et Σ' .

Lorsqu'on ajoute ou retranche, pour former Σ ou Σ' , les produits composants $[f_1 g_1]_{\omega_s}$ et $[f_2 g_2]_{\omega_s}$, les points de sondage de l'un concordent, quant au temps, avec les zéros de l'autre, et vice-versa. Il en résulte que pour la suite des temps (9), on aura :

$$(11) \quad \Sigma = \Sigma' = 2f_1$$

et pour la suite (10) :

$$(12) \quad \Sigma = -\Sigma' = 2f_2$$

Ce sont là les propriétés mathématiques fondamentales de ces signaux. Si l'on examine maintenant la répartition physique de l'énergie, on constate que, alors que l'énergie se trouve toujours du côté positif de l'axe des temps, pour Σ , qui est toujours

positif, elle peut se trouver soit au-dessus, soit en-dessous de cet axe, pour Σ' , qui, en général, change de signe au rythme de ω_s .

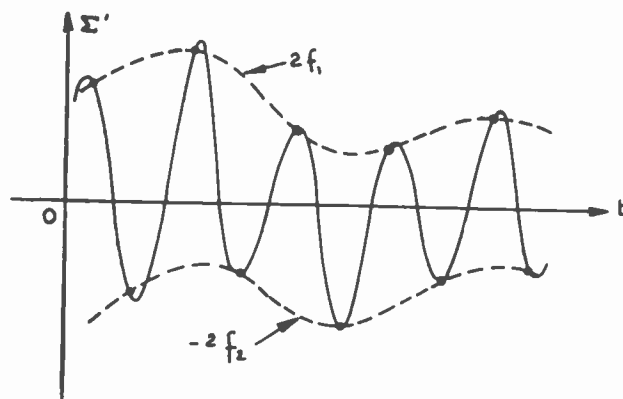


FIG. 7.

En effet, contrairement à Σ , dont l'allure générale peut être quelconque, Σ' conserve une structure pseudo-périodique très accentuée. Ce signal, en effet, présente au voisinage immédiat des points de sondage de f_1 , des crêtes positives pour lesquelles sa valeur est voisine de $2f_1$, tandis qu'au voisinage des points de sondage de f_2 , Σ' présente des crêtes négatives dont la valeur est voisine de $-2f_2$. Il en résulte que Σ' oscille, en première approximation, entre $2f_1$ et $-2f_2$, avec une période égale à $2\pi/\omega_s$.

La figure 7 donne l'allure générale de Σ' . Il est clair que l'on peut considérer, toujours en première approximation que la partie principale de l'énergie correspondant à la partie positive de Σ' est fonction de f_1 seulement, tandis que la partie principale de l'énergie correspondant à la partie négative de Σ' est fonction de f_2 seulement.

Certains chercheurs ont alors eu l'idée de proposer un procédé très simple pour la séparation de f_1 et f_2 à la station réceptrice, qui consiste à remplacer les produits π_1 et π_2 par la partie positive de Σ' , d'une part, et sa partie négative, d'autre part, ces deux parties étant obtenues à partir de Σ' au moyen de circuits à diodes convenablement dimensionnés.

Toutefois, il faut bien noter qu'un tel procédé de réception n'est qu'approché. Nous l'examinerons plus en détail au chapitre 4, et nous montrerons que les approximations ne sont nullement négligeables. Cependant, nous montrerons également qu'il est possible d'y porter remède dans une mesure suffisante pour permettre d'envisager des applications pratiques.

D. — Sur le caractère séquentiel des produits de réception.

Au chapitre 1, § B, nous avons laissé entendre que le caractère séquentiel des produits π_1 et π_2 était dû à la présence des spectres $M(F_1)$ et $M(F_2)$ dans ces produits. Or, en nous reportant au chapitre I, § B, nous voyons que les produits de réception π_1 et π_2 et les produits composants $[f_1 g_1]_{\omega_s}$ et $[f_2 g_2]_{\omega_s}$,

ont la même composition spectrale. On peut écrire, en négligeant des coefficients constants, par exemple :

$$(13) \quad \pi_1 = [f_1 g_1]_{\omega_s} = F_1 + M(F_1)$$

Or, reportons-nous maintenant au § B du présent chapitre où nous donnons une description graphique de $[f_1 g_1]_{\omega_s}$. Reprenons l'équation (6) :

$$(14) \quad [f_1 g_1]_{\omega_s} = f_1(t) + \frac{1}{2} f_1(t) \cos \omega_s t + \sin \omega_s t \int_0^{\omega_s} [x \sin \omega t - y \cos \omega t] d\omega$$

Le premier terme au second membre de (14) représente le spectre F_1 au dernier membre de (13). Les termes entre accolades de (14) représentent le spectre $M(F_1)$ du dernier membre de (13). Or, notre description a montré que ces termes en $\cos \omega_s t$ et $\sin \omega_s t$ sont précisément ceux qui amènent une « structure de points » sur $[f_1 g_1]_{\omega_s}$.

C'est cette « structure de points » qui est la marque du caractère séquentiel de la transmission. Le spectre $M(F_1)$ est donc bien seul responsable du caractère séquentiel des produits de réception, qui, en définitive, pour une reproduction d'images, sont les signaux appliqués aux kinoscopes.

III. — APPLICATION AU DOUBLEMENT DE LA CAPACITÉ D'UN FAISCEAU HERTZIEN.

A. — Introduction.

Nous avons montré que, si ω_s est la bande normalement allouée aux signaux video pour la télévision en noir et blanc au moyen d'un certain standard de balayage, il est possible, en utilisant le signal Σ , ou le signal Σ' , d'assurer la transmission de deux programmes indépendants, dans la même bande ω_s , la bande allouée à la video de chaque programme pouvant être très peu inférieure à ω_s .

On voit donc que si l'on attaque une liaison par faisceaux hertziens, non pas au moyen d'un signal video classique, mais au moyen de Σ ou Σ' , avec la même bande, il suffira d'assurer, à la station terminale, la séparation de f_1 et f_2 par les moyens déjà décrits, pour disposer ainsi de deux signaux video susceptibles de moduler deux chaînes d'émetteurs de radiodiffusion, distribuant donc deux programmes indépendants.

B. — Restitution des fonctions de sondage à la réception.

La formation des produits de réception π_1 et π_2 nécessite que l'on possède les fonctions de sondage g_1 et g_2 , et donc une onde en $\cos \omega_s t$, avec une phase fixe. Deux procédés nettement distincts peuvent être employés à cette fin :

1° On peut tirer l'onde $\cos \omega_s t$ des signaux transmis Σ ou Σ' .

2° On peut envisager un encombrement un peu plus grand de la transmission (10 à 20 %), et utiliser la portion de bande supplémentaire pour la transmission d'un signal pilote.

C. — Extraction de la sous-porteuse du signal transmis.

Nous appelons dans la suite ω_s la fréquence de la sous-porteuse du système. C'est, en effet, une sous-porteuse pour les spectres $M(F_1)$ et $M(F_2)$ transmis.

Pour faciliter l'exposé, nous supposons pour commencer que c'est le signal Σ' qui est utilisé pour la transmission. Considérons alors la quantité sans limitation de bande :

$$(1) \quad S' = f_1 [1 + 2 \cos \omega_s t] - f_2 [1 - 2 \cos \omega_s t] \\ = f_1 - f_2 + 2 [f_1 + f_2] \cos \omega_s t$$

Sur la figure 8, nous avons donné une décomposition spectrale de S' . En trait plein est représenté le spectre de $f_1 - f_2$. Nous le supposons pour l'instant

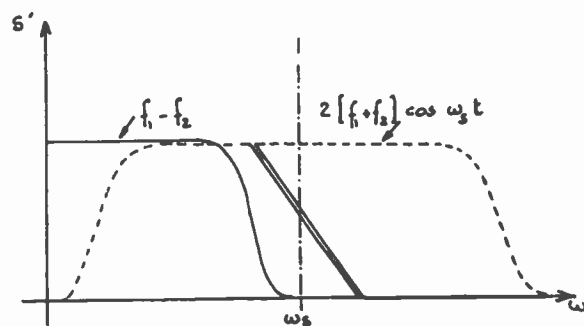


FIG. 8.

couvrir une bande un peu inférieure à ω_s , g_1 et g_2 étant réduits au premier harmonique.

En pointillé, nous avons donné la représentation spectrale complète de $2 [f_1 + f_2] \cos \omega_s t$. Ce spectre est rigoureusement symétrique par rapport à la fréquence ω_s .

Nous avons représenté, par un trait double continu, la courbe de réponse du filtre faisant passer de S' à Σ' dans la zone d'atténuation progressive.

On s'efforce évidemment, en général, de réduire au minimum cette zone d'atténuation progressive, mais, en pratique, elle demeure relativement large, et l'on peut considérer, pour se fixer les idées, qu'elle est de l'ordre de $0,2 \omega_s$ au moins.

Considérons alors une étroite bande de fréquence, de largeur au plus égale à $\omega_s/100$, et symétrique par rapport à ω_s . Le spectre tiré de S' , et contenu à l'intérieur de cette bande, est parfaitement symétrique. Le spectre tiré de Σ' , et contenu dans cette bande, n'est pas rigoureusement symétrique. Il l'est du point de vue de la fréquence des raies qui le constituent, mais pas du point de vue de leur amplitude, à cause de la variation de module du gain du filtre limitant la transmission. Cependant, pour les valeurs limites

données ci-dessus, la variation de module dans la bande sera de 10 % seulement.

Si donc on prélève, au moyen d'un filtre passe-bande lui-même bien symétrique, la partie de Σ' contenue dans la bande étroite en question, $\Delta\omega$, il sera possible d'obtenir ainsi un spectre bien symétrique si l'on fait intervenir un réseau correcteur de module convenable possédant lui-même une caractéristique de phase bien symétrique. Tout se passera donc, en définitive, au niveau près, comme si l'on avait prélevé la partie du spectre de $2[f_1 + f_2] \cos \omega_s t$ contenue dans la bande $\Delta\omega$ centrée sur ω_s .

Or, on a :

$$(2) \quad P = |f_1 + f_2| \cos \omega_s t_{\Delta\omega} = |f_1 + f_2|_{\Delta\omega/2} \cos \omega_s t$$

Le signal P se présente donc comme une onde modulée en amplitude par les composantes de $(f_1 + f_2)$ comprises dans la bande $[0, \Delta\omega/2]$. Le spectre de l'onde modulée est complet, de telle sorte que sa phase est exactement celle de $\cos \omega_s t$. Cette propriété permet d'utiliser le signal P comme signal de pilotage pour la fabrication, à la station réceptrice, de l'onde $\cos \omega_s t$.

Plusieurs procédés connus permettent le passage du signal P à une onde pure de même fréquence. En particulier, on peut faire subir au signal P un certain nombre d'écrêtages successifs bien symétriques, ce qui fournira un signal rectangulaire dont on extraira l'harmonique 1. On peut également employer le signal P pour synchroniser un oscillateur déjà assez bien calé sur la fréquence ω_s .

Toutefois, ces différents procédés exigent que l'amplitude du signal P demeure notable. Or, on voit, d'après (2), que P peut s'annuler, si f_1 et f_2 s'annulent simultanément, ce qui peut arriver.

Pour éviter ce phénomène, il convient de transmettre non pas f_1 et f_2 , variables de 0 à 1, mais $\alpha/2 + f_1$ et $\alpha/2 + f_2$, $\alpha/2$ étant une constante égale, par exemple, à 0,25.

On aura donc :

$$(3) \quad \Sigma' = |\alpha/2 + f_1| g_1 - |\alpha/2 + f_2| i_{\omega_s}$$

$$(4) \quad P = |f_1 + f_2 + \alpha|_{\Delta\omega/2} \cos \omega_s t$$

L'amplitude de P variera entre α et $2 + \alpha$, soit, pour les valeurs numériques choisies, de 0,5 à 2,5, ce qui conduit à un rapport 5. En pratique, on choisira la valeur de α minimum compatible avec un bon fonctionnement. En effet, la transmission de la constante α correspond, pour une puissance donnée de l'émetteur, à une diminution de la puissance disponible pour les signaux utiles f_1 et f_2 , et donc à une augmentation relative du bruit de fond.

Nous avons supposé, à l'occasion de la figure 8, que les fonctions g_1 et g_2 ne contenaient pas d'harmoniques supérieures à 1. On peut, comme nous l'avons déjà dit dans les précédents chapitres, introduire dans g_1 et g_2 les termes d'harmonique 2. On peut

alors étendre jusqu'à ω_s la bande offerte aux signaux vidéo f_1 et f_2 , ce qui évite de perdre une partie de la définition.

Le spectre recueilli dans la bande $\Delta\omega$ centrée sur ω_s demeure symétrique. Cette bande se trouve, en effet, envahie par des composantes de $f_1 - f_2$, mais également par des composantes, exactement symétriques des précédentes, du produit :

$$(5) \quad 2[f_1 - f_2] \cos 2\omega_s t$$

dû aux termes d'harmonique 2 des fonctions de sondage.

Supposons maintenant que la transmission soit assurée non par Σ' mais par Σ . Les raisonnements qui précèdent dans ce paragraphe demeurent valables, mais (2) doit être remplacée par :

$$(6) \quad P = |f_1 - f_2|_{\Delta\omega/2} \cos \omega_s t$$

On voit que le coefficient de $\cos \omega_s t$ peut varier entre -1 et $+1$, alors qu'il est indispensable, pour le maintien du pilotage, que ce coefficient demeure, par exemple, positif, et de plus, ne devienne jamais trop petit.

On voit qu'il sera nécessaire, pour obtenir ce résultat, de transmettre un signal Σ présentant la composition nouvelle :

$$(7) \quad \Sigma = |f_1 g_1 + f_2 g_2|_{\cos} + |1 + \alpha|_{\cos \omega_s t}$$

où α est la même constante qui intervient dans (3) et (4). L'amplitude de P variera alors, comme précédemment, entre α et $2 + \alpha$.

Il est à noter que l'emploi de Σ' selon (3), ou de Σ selon (7), conduit au même rendement de la puissance de l'émetteur. Si, en effet, nous choisissons des coefficients numériques tels que l'amplitude du signal Σ' de la théorie générale soit $f_1 + f_2$, dans les mêmes conditions celle de Σ est seulement f_1 ou f_2 . L'amplitude maximum est donc 2 pour Σ' et 1 pour Σ . Mais, avec les mêmes normes, d'après (3), l'amplitude de Σ' sera $\alpha + f_1 + f_2$ ou au maximum, $2 + \alpha$, et d'après (5) celle de Σ sera $1 + \alpha + f_1$ ou $1 + \alpha + f_2$, soit, au maximum, $2 + \alpha$.

Dans le cas de l'emploi du signal Σ , il convient d'éviter les perturbations que pourrait causer, à la réception, le terme additionnel $|1 + \alpha| \cos \omega_s t$. Aussi, à la station terminale, après reconstitution de l'onde $\cos \omega_s t$, retranchera-t-on du signal destiné aux modulateurs qui fabriquent les produits de réception π_1 et π_2 , le terme $|1 + \alpha| \cos \omega_s t$ fabriqué à partir de l'onde sous-porteuse reconstituée.

D. — Transmission supplémentaire d'un signal pilote.

Considérons une liaison par faisceau hertzien établie, par exemple, pour la transmission de signaux ordinaires de télévision. Cette liaison présente une

certaine bande passante B qui n'est pas nécessairement la bande utile B_0 . La transmission de signaux ordinaires de télévision nécessite, en effet, que la phase de la ligne ne soit pas trop distordue, et, en général, des distorsions de phase considérables apparaissent, pour les réseaux usuels, dans la partie de leur bande passante correspondant aux fréquences les plus élevées.

On peut donc considérer que dans beaucoup de liaisons par faisceau hertzien, une certaine bande reste disponible pour assurer la transmission de signaux peu sensibles aux distorsions de gain, tant d'ailleurs en module qu'en phase.

Dans le cas d'un faisceau hertzien où B_0 se confond avec B , on peut toujours envisager un petit élargissement de la bande, si aucune restriction n'est imposée concernant la régularité du gain à l'intérieur de l'extension ainsi opérée. On peut donc considérer que, dans de nombreux cas de liaisons par faisceaux hertiens, il est possible de disposer d'un petit excédent de bande, en plus de celle allouée à Σ ou Σ' , permettant la transmission d'un signal sinusoïdal.

On utilisera cette possibilité pour transmettre un signal pilote :

(8) $\pi = \cos \omega_p t$

avec :

(9) $\omega_p = \frac{p}{q} \omega_s$, p et q entiers, $p > q$

Pour fixer les idées, on considérera que, pour l'utilisation de filtres passe-bas usuels pour limiter la bande des signaux Σ ou Σ' , on peut prendre :

(10) $\omega_p \geq 1,2 \omega_s$

si l'on veut éviter toute perturbation de ω_p par le spectre du signal principal. On prendra alors par exemple :

(11) $p = 11 \qquad q = 9 \qquad \omega_p = \frac{11}{9} \omega_s$

à la station terminale, ω_p sera séparée au moyen d'un filtre à bande étroite, démultipliée par p , et ω_p/p sera multipliée par q . Ces diverses opérations peuvent être effectuées avec une sécurité considérable.

A l'émission, on peut envisager un oscillateur de base fournissant la fréquence ω_p , et les mêmes opérations pour obtenir ω_s .

La bande à l'intérieur de laquelle s'effectue la transmission de ω_p peut être très étroite, allant par exemple de $\frac{\omega_s}{100}$ à $\frac{\omega_s}{1\,000}$. Dans ces conditions, par rap-

port à l'amplitude maximum du signal principal Σ ou Σ' , l'amplitude du signal pilote peut être réduite dans un rapport de 10 à 30, si l'on désire conserver la même qualité de transmission.

Il en résulte que l'amplitude du signal pilote devient négligeable devant celle du signal principal.

L'avantage revient alors de manière nette à l'utilisation du signal Σ , qui conduit à un rapport signal/bruit deux fois plus grand que celui obtenu avec Σ' .

Ce qui précède suppose toutefois que la transmission principale n'est pas entachée de distorsion non-linéaire.

E. — Exigences sur les caractéristiques de transmission.

Les diverses opérations permettant, à partir de Σ ou Σ' , de récupérer indépendamment f_1 et f_2 à la station réceptrice, doivent être effectuées avec un certain degré de précision, si l'on veut éviter que des termes perturbateurs gênants n'apparaissent sur chaque canal.

Pour cela, en admettant que l'appareillage de réception soit correctement établi, il faut que les propriétés du signal Σ (ou Σ') à la sortie de la ligne de transmission soient assez voisines des propriétés de ce signal à l'entrée de la ligne.

Désignons par :

(12) $g = |g| e^{j\varphi}$

le gain de la ligne de transmission. Si $|g|$ était une constante en fonction de la fréquence, et si φ était fonction linéaire de ω , il n'y aurait aucune altération, le long de la ligne, des propriétés du signal transmis.

En pratique, le module du gain ne demeure pas rigoureusement constant dans la bande passante, et la phase n'est pas exactement fonction linéaire de la fréquence. Il en résulte toujours, à la réception, des phénomènes de diaphonie qui ne sont pas nécessairement gênants. Nous nous proposons, pour terminer ce chapitre, de calculer en pratique l'ordre de grandeur des termes de diaphonie que l'on peut rencontrer en présence de petites irrégularités des caractéristiques de transmission.

F. — Influence de la distorsion de phase.

Sur la figure 9, nous avons représenté une courbe arbitraire, donnant en valeurs positives le retard

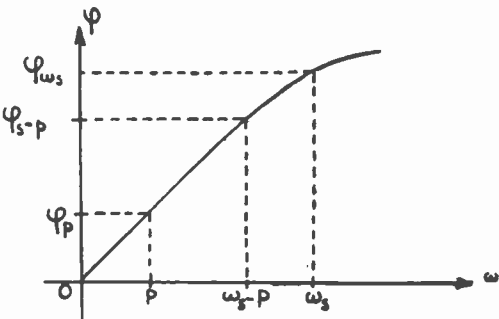


FIG. 9.

de phase φ pour la ligne de transmission. Généralement une telle courbe présente à partir de la fréquence zéro une partie linéaire assez longue, puis une courbure apparaît rapidement qui croît de plus en plus vite.

Nous supposons que les informations à transmettre sont :

$$(13) \quad f_1 = \cos pt \quad f_2 = 0$$

Il vient, en utilisant par exemple le signal Σ :

$$(14) \quad \Sigma = \cos pt + \cos (\omega_s - p)t$$

A l'arrivée au récepteur, le signal sera devenu :

$$(15) \quad \Sigma = \cos [pt - \varphi_p] + \cos [(\omega_s - p)t - \varphi_{\omega_s - p}]$$

Pour la séparation de f_1 et f_2 nous utiliserons :

$$(16) \quad \begin{cases} g_1 = 1 + 2 \cos (\omega_s - \varphi_s) \\ g_2 = 1 - 2 \cos (\omega_s - \varphi_s) \end{cases}$$

Calculons les produits de réception étendus à deux images consécutives :

$$(17) \quad \begin{cases} x_1 = \pi_1(t) + \pi_1(t + T_1) = 2 \cos (pt - \varphi_p) \\ \quad \quad \quad + 2 \cos (pt - \varphi_{\omega_s} + \varphi_{\omega_s - p}) \\ x_2 = \pi_2(t) + \pi_2(t + T_1) = 2 \cos (pt - \varphi_p) \\ \quad \quad \quad - 2 \cos (pt - \varphi_{\omega_s} + \varphi_{\omega_s - p}) \end{cases}$$

Si le déphasage était linéaire, par exemple de valeur :

$$(18) \quad \varphi = \varphi_p \frac{\omega}{p}$$

on aurait une séparation correcte à la réception, avec :

$$(19) \quad \begin{cases} x_1 = 4 \cos (pt - \varphi_p) \\ x_2 = 0 \end{cases}$$

On voit que la distorsion de phase conduit, d'une part, à une altération du signal reçu sur le canal normal x_1 , et, d'autre part, à des termes perturbateurs sur le second canal x_2 .

On voit également que tout se passerait bien si l'on avait :

$$(20) \quad \varphi_{\omega_s} = \varphi_p + \varphi_{\omega_s - p}$$

Traduite en langage géométrique la condition (20) implique que la courbe de phase de la figure 9 présente un centre de symétrie pour $\omega = \omega_s/2$.

Il n'est donc pas nécessaire, pour que la ligne de transmission conserve la possibilité de séparation des signaux sans diaphonie, que sa phase soit linéaire. Il suffit que sa variation présente un centre de symétrie pour $\omega = \omega_s/2$.

Nous allons maintenant évaluer de manière plus précise les phénomènes de diaphonie. Nous poserons :

$$(21) \quad \varphi_{\omega_s} - \varphi_{\omega_s - p} = \varphi_p + \varepsilon$$

d'où :

$$(22) \quad \begin{cases} x_1 = 2[1 + \cos \varepsilon] \cos (pt - \varphi_p) + 2 \sin \varepsilon \sin (pt - \varphi_p) \\ x_2 = 2[1 - \cos \varepsilon] \cos (pt - \varphi_p) - 2 \sin \varepsilon \sin (pt - \varphi_p) \end{cases}$$

On peut considérer, de manière arbitraire, que φ_p correspond, par exemple, à un retard normal, tandis que ε représente la distorsion de phase. Il est inutile de considérer la diaphonie autrement que lorsque ε est petit. Il vient alors, en ne conservant que les termes principaux :

$$(23) \quad \begin{cases} x_1 = 4 \cos (pt - \varphi_p) \\ x_2 = -2 \varepsilon \sin (pt - \varphi_p) \end{cases}$$

Nous obtenons ainsi de manière simple le terme de diaphonie provoqué par une composante sinusoïdale. Un tel cas peut se présenter en pratique lors de la transmission d'une mire à traits parallèles horizontaux ou verticaux. Les traits étant alternativement noirs et blancs, l'amplitude crête-à-crête du signal x sera 8. Les 2 canaux étant traités symétriquement, le passage du noir au blanc pour x_2 correspond également à une amplitude de 8. Or, sur ce canal nous trouvons un terme perturbateur d'amplitude crête-à-crête 4ε . Le coefficient de diaphonie est donc :

$$(24) \quad D = \frac{\varepsilon}{2}$$

Nous allons montrer la signification géométrique de ε . Sur la figure 10, nous avons repris la courbe

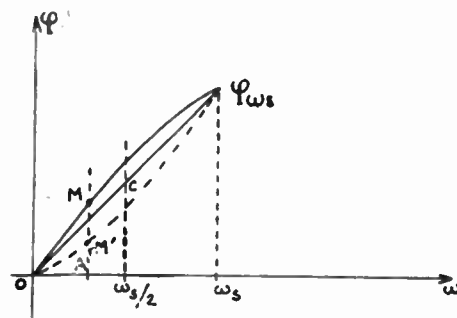


FIG. 10.

de phase de la figure 9, et nous avons porté en pointillé la courbe symétrique par rapport au milieu C de la corde joignant les points obtenus pour $\omega = 0$ et $\omega = \omega_s$.

On verra facilement que si la première courbe représente φ_p , la seconde représente : $\varphi_{\omega_s} - \varphi_{\omega_s - p}$

Soit donc les points M et M' pour la fréquence p.

On a :

$$(25) \quad \begin{cases} \overline{AM} = \varphi_p \\ \overline{AM'} = \varphi_{\omega_s} - \varphi_{\omega_s - p} \end{cases}$$

D'après (21), on aura :

(26) $\epsilon = \overline{AM'} - \overline{AM} = \overline{MM'}$

En conclusion, et en laissant de côté le signe, on peut dire que ϵ , erreur de phase qui intervient directement dans l'expression du coefficient de diaphonie, est, pour une fréquence à transmettre p , la différence des ordonnées de la courbe de phase, d'une part, et de sa symétrique par rapport à C , d'autre part.

G. — Influence des irrégularités de réponse en module.

Nous poserons arbitrairement que le gain est égal à l'unité aux très basses fréquences. Nous admettrons que la phase est linéaire, et nous représenterons le module du gain par :

(27) $|g| = 1 - \gamma(\omega)$

Considérons de nouveau le signal Σ d'après (11), et soit ce signal à la sortie de la ligne :

(28) $\Sigma = [1 - \gamma(p)] \cos pt + [1 - \gamma(\omega_s - p)] \cos(\omega_s - p)t$

Il vient alors :

(29) $\begin{cases} x_1 = 2 [2 - \gamma(p) - \gamma(\omega_s - p)] \cos pt \\ x_2 = 2 [\gamma(\omega_s - p) - \gamma(p)] \cos pt \end{cases}$

Le terme de diaphonie est ici de forme très simple puisque c'est une fraction du signal perturbateur. On introduira avantageusement une quantité γ jouant un rôle analogue à ϵ :

(30) $\gamma = |1 - \gamma(p)| - |1 - \gamma(\omega_s - p)|$
 $ = \gamma(\omega_s - p) - \gamma(p)$

On peut alors écrire, en négligeant l'altération de x_1 :

(31) $\begin{cases} x_1 = 4 \cos pt \\ x_2 = 2 \gamma \cos pt \end{cases}$

ce qui fait apparaître un coefficient de diaphonie :

(32) $D = \frac{\gamma}{2}$

La représentation graphique de γ est donnée sur la figure 11 où l'on a tracé en trait plein la courbe de module du gain, et en pointillé la courbe symétrique par rapport à l'axe vertical $\omega = \omega_s/2$. On a :

(33) $\gamma = \overline{M'M}$

On peut rechercher la valeur de γ conduisant à

un taux de diaphonie de 3 %. De (37) vient immédiatement :

(34) $\gamma = 2 D = 0,06$

Les courbes de module étant très généralement plates dans la première moitié de la bande, et la courbure apparaissant dans la deuxième moitié,

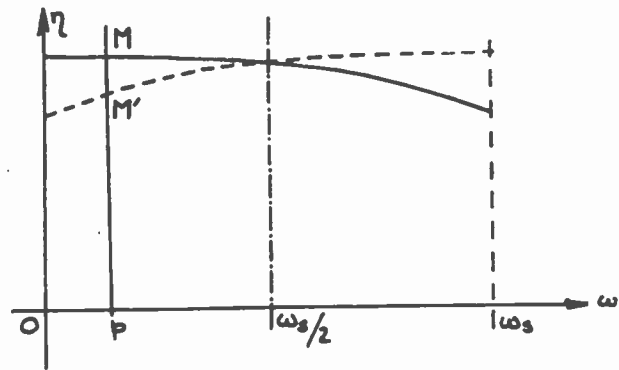


Fig. 11.

γ représente pour des fréquences assez basses ou assez élevées, la chute relative du gain de la ligne de transmission.

IV. — RÉCEPTION PAR SÉPARATION DES PARTIES POSITIVE ET NÉGATIVE.

A. — Introduction.

Nous avons indiqué au Chapitre II que le signal Σ' , formé de « points » alternativement positifs et négatifs, pouvait être utilisé d'une manière très simple à la réception. Un ensemble de deux circuits à diode permet d'envoyer sur un canal la partie positive du signal seulement, et sur un autre canal la partie négative. On peut alors considérer en première approximation que la partie positive du signal représente essentiellement f_1 , avec une certaine « structure de points », tandis que la partie négative représente dans des conditions analogues $-f_1$.

En pratique, ces propositions sont seulement approchées et chacun des signaux est obtenu avec une diaphonie qui peut être importante. Cette diaphonie existe, alors même, que la ligne de transmission assure le transport correct de Σ' , et nous l'appellerons pour cela « diaphonie intrinsèque » à la méthode de réception.

Si, d'autre part, la ligne de transmission n'est pas parfaite, une diaphonie supplémentaire peut apparaître, que nous appellerons « diaphonie de transmission ». Cette diaphonie de transmission joue ici un rôle analogue à celui de la diaphonie étudiée au Chapitre III pour une réception par la méthode classique des multiplex à impulsions modulées en amplitude.

Nous étudierons ici la « diaphonie intrinsèque » qui peut se décomposer en deux types :

- a) diaphonie intrinsèque en régime permanent,
- b) diaphonie intrinsèque en régime transitoire.

B. — Diaphonie intrinsèque en régime permanent.

En régime permanent, $f_1(t)$ et $f_2(t)$ sont des constantes, que nous désignerons simplement par f_1 et f_2 . On a alors :

$$(1) \quad \Sigma' = f_1 - f_2 + [f_1 + f_2] \cos \theta$$

Le signal Σ' est représenté sur la figure 12. Nous avons hachuré les parties positives des aires comprises entre le signal et l'axe. Lorsqu'on applique un tel signal à un kineoscope, et si l'on fait abstraction du Γ du tube, il est clair que l'on obtient une brillance proportionnelle à la valeur moyenne de l'aire hachurée (l'axe, ou base inférieure du signal, étant en coïncidence avec le cut-off du kineoscope). Nous allons calculer cette aire moyenne y_1 , qui sera considérée comme l'information obtenue par utilisation de la partie positive de Σ' .

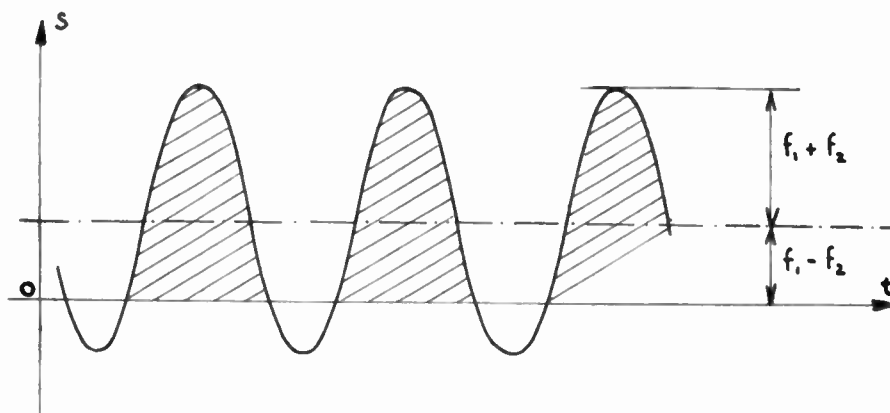


FIG. 12.

En utilisant les notations de la figure 13, on trouve :

$$(2) \quad y_1 = \frac{1}{2\pi} \int_{-\theta_0}^{\theta_0} \Sigma' d\theta$$

avec :

$$(3) \quad \theta_0 = \text{Arc cos} \frac{f_2 - f_1}{f_2 + f_1}, \quad 0 < \theta_0 < \pi$$

Il vient finalement :

$$(4) \quad y = \frac{\theta_0 + \sin \theta_0}{\pi} f_1 - \frac{\theta_0 - \sin \theta_0}{\pi} f_2$$

(4) est exprimé à la fois en fonction de θ_0 et f_1 et f_2 . En fait, d'après (3), θ_0 peut être exprimée en fonction de f_2/f_1 , par exemple, de telle sorte que y peut être exprimée en fonction de f_1 et f_2 seulement. Toutefois, nous conserverons l'écriture (4) plus simple, et aussi commode au moins pour les calculs numériques.

On voit que, lorsque $f_2 = 0$ ($\theta_0 = \pi$) le signal obtenu sur le canal f_1 est correct puisque nécessairement sans diaphonie. On a alors, d'après (4) :

$$(5) \quad y_1 = f_1$$

Pour mettre en évidence le terme de diaphonie, on écrira :

$$(6) \quad y = f_1 \left[1 + \frac{\theta_0 - \pi + \sin \theta_0}{\pi} - \frac{\theta_0 - \sin \theta_0}{\pi} \frac{f_2}{f_1} \right]$$

D'où le coefficient de diaphonie relative de f_2 sur f_1 :

$$(7) \quad D_1 = \frac{\theta_0 - \pi + \sin \theta_0}{\pi} - \frac{\theta_0 - \sin \theta_0}{\pi} \frac{f_2}{f_1}$$

La figure 14 donne une représentation de ce coefficient, qui est fonction de f_2/f_1 seulement. Notons que pour la valeur 1 de ce rapport D_1 atteint 37 %. La figure 15 donne une quantité plus suggestive, qui est le terme de diaphonie lui-même, quand

on se place dans les plus mauvaises conditions c'est-à-dire quand f_2 possède sa valeur maximum qui est 1. Le terme de diaphonie perturbant f_1 est alors donné en fonction de f_1 variant de 0 à 1.

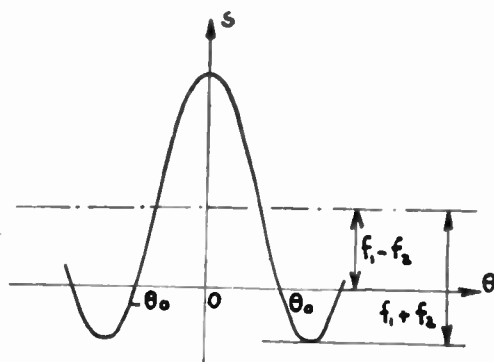


FIG. 13.

On voit que ce terme est maximum pour $f_1 = 1$, et décroît avec f_1 pour tendre vers 0 en même temps. C'est là une propriété intéressante de la diaphonie étudiée : il n'y a pas perturbation des plages noires, c'est-à-dire des parties de l'image où l'œil est le moins tolérant aux perturbations lumineuses.

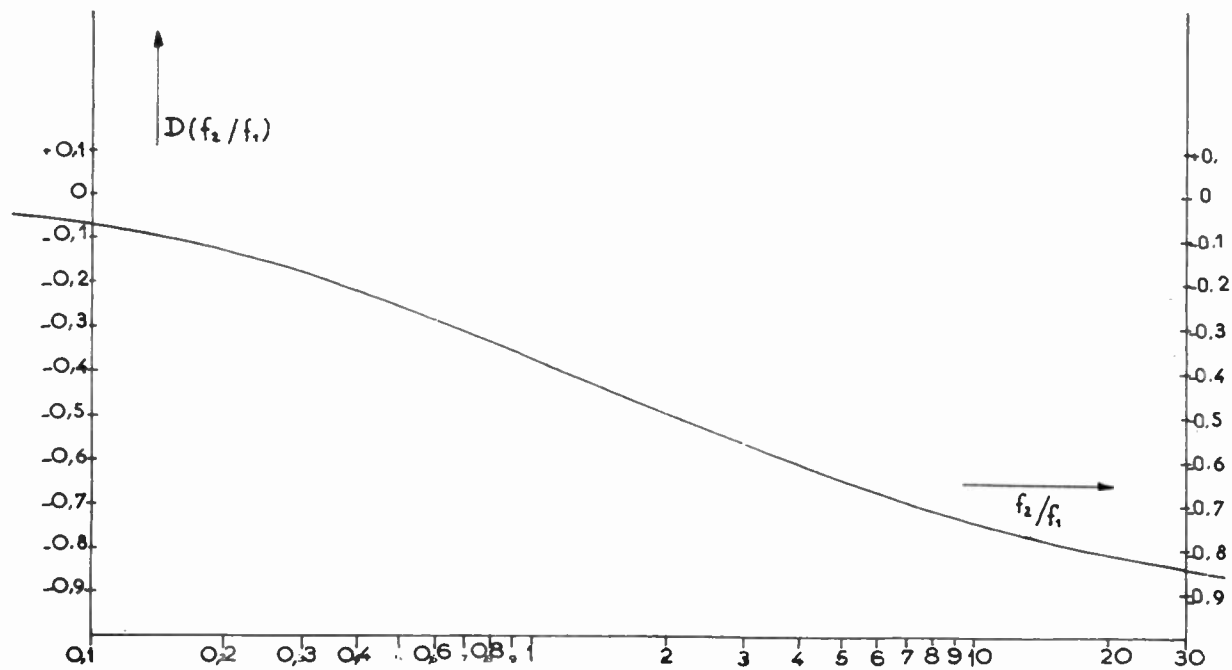


FIG. 14.

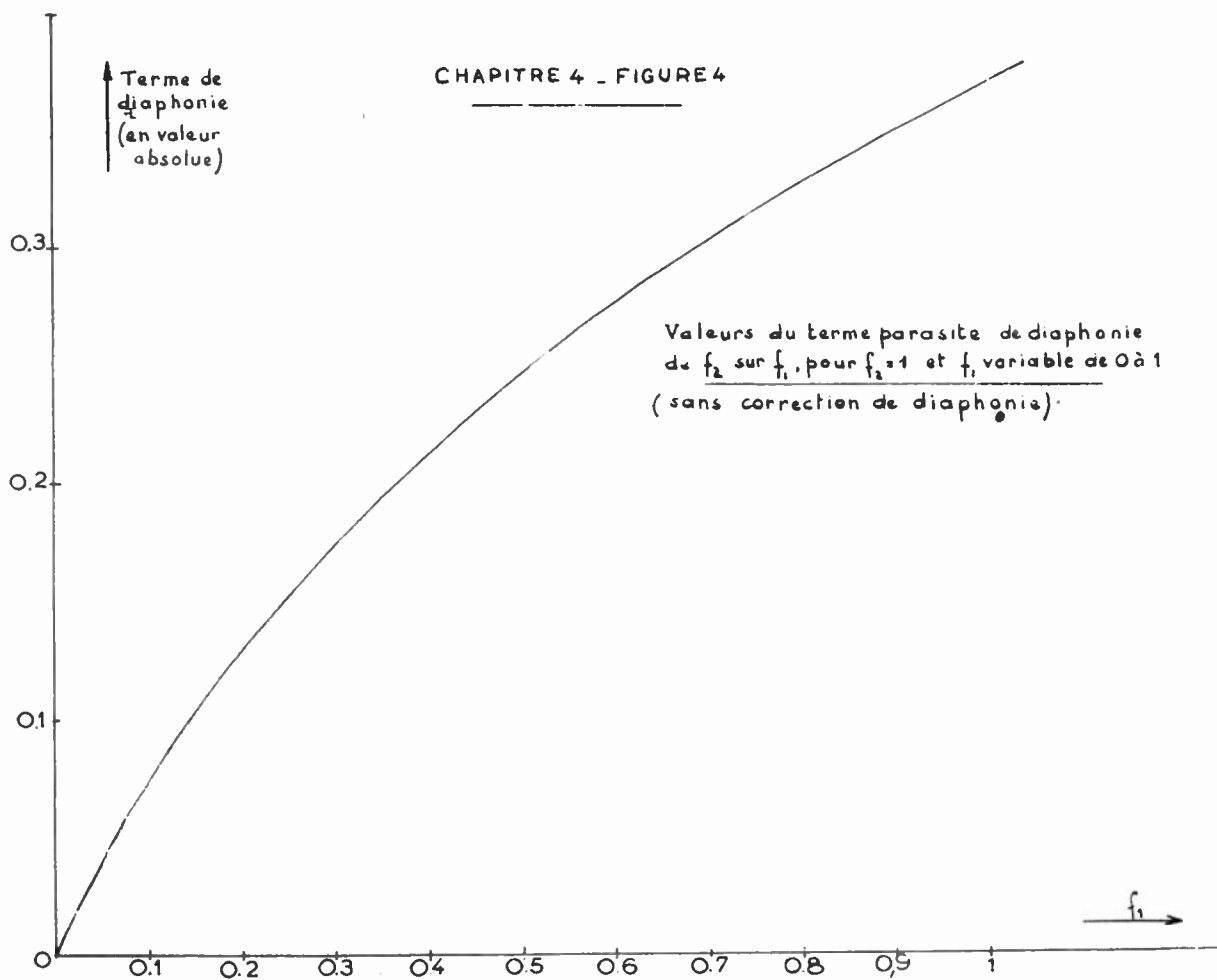


FIG. 15.

Les 2 voies jouant des rôles symétriques, les mêmes propriétés peuvent être établies pour la diaphonie créée par le signal f_1 sur le canal f_2 .

Il existe donc une unique fonction $D(x)$ qui permet d'exprimer chacun des 2 coefficients de diaphonie. Pour avoir le coefficient de diaphonie de la voie 2 sur la voie 1, D_1 , on remplacera x par f_2/f_1 , et l'on écrira :

$$(7a) \quad D_1 = D(f_2/f_1)$$

Pour avoir le coefficient de diaphonie de la voie 1 sur la voie 2, D_2 , on remplacera x par f_1/f_2 , et l'on écrira :

$$(7b) \quad D_2 = D(f_1/f_2)$$

C. — Diaphonie intrinsèque en régime transitoire.

En régime permanent, les produits composants, $[f_1 g_1] \omega_s$ et $[f_2 g_2] \omega_s$ sont de pures sinusoïdes. Nous avons indiqué au Chapitre II qu'en régime transitoire, les produits avaient approximativement 2 enveloppes, l'une est l'information correspondante f_1 ou f_2 , l'autre est l'axe des temps. Cependant, cette

Il apparaît ainsi que le signal f_1 , en régime transitoire, contribue à augmenter l'énergie de l'information y_2 recueillie en prélevant la partie négative de Σ' . Il s'agit là d'un phénomène de diaphonie supplémentaire, moins important en général que celui observé au § B, et de sens inverse. En effet, les termes de diaphonie observés au § B sont négatifs, et, par conséquent, y_2 par exemple est toujours plus petit que f_2 .

Nous avons dit que la diaphonie intrinsèque en transitoire était moins importante en général que celle observée en régime permanent. Il faut noter tout d'abord que la diaphonie en régime permanent existe, avec sensiblement la même valeur, en régime variable, et que la diaphonie en transitoire, telle que nous venons de la définir, s'y ajoute algébriquement.

D'autre part, la diaphonie en transitoire, de f_1 sur f_2 , par exemple, subsiste pour $f_2 = 0$, alors que dans ce cas, la diaphonie définie en régime permanent s'annule.

Nous ne dirons rien de plus sur la diaphonie particulière aux régimes transitoires, son évaluation étant compliquée, et fonction de la forme du signal.

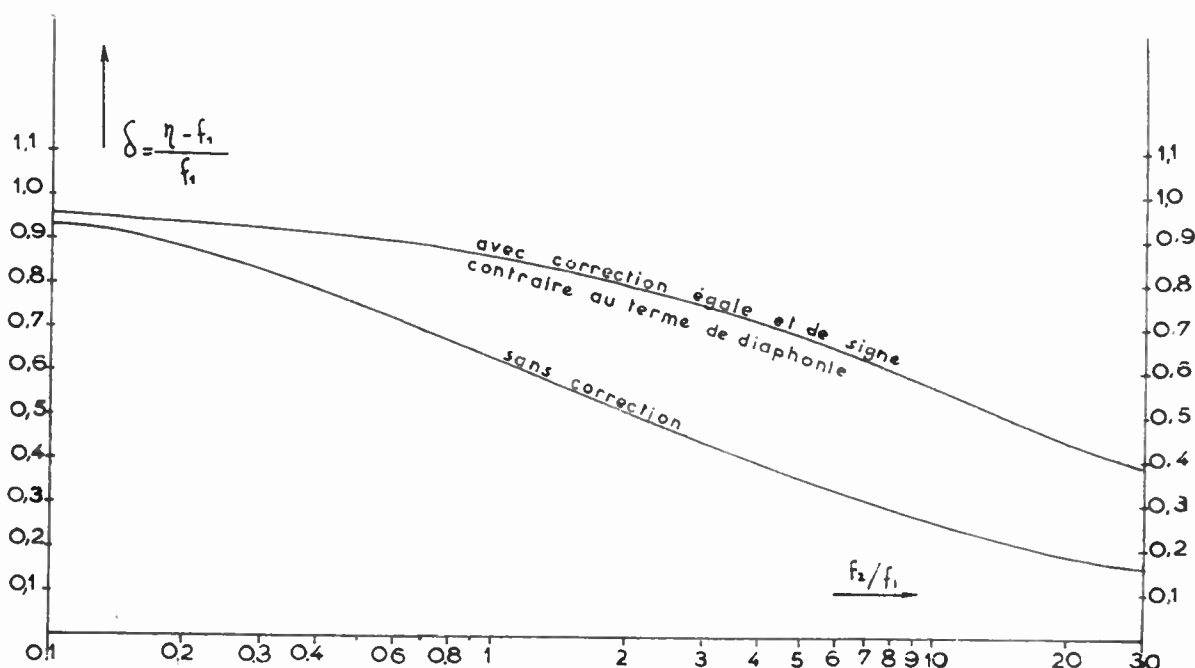


FIG. 16.

description n'est qu'approchée, et, en pratique, il y a toujours des crêtes de ces produits en-dessous de l'axe des temps.

Quand on fait l'opération :

$$(8) \quad \Sigma' = [f_1 g_1] \omega_s - [f_2 g_2] \omega_s$$

les crêtes négatives de $[f_1 g_1] \omega_s$, par exemple, sont à peu près en phase avec les crêtes négatives de $-[f_2 g_2] \omega_s$ et contribuent à accroître l'énergie disponible en-dessous de l'axe des temps pour Σ' .

D. — Correction des diaphonies intrinsèques.

Nous allons exposer ici un mode général de correction des diaphonies, plus particulièrement commode pour l'atténuation des diaphonies intrinsèques, mais également applicable aux diaphonies de transmission, dans la mesure où celles-ci sont bien définies.

Nous n'envisageons pas, dans cette étude rapide, de calculer les diverses diaphonies de transmission, le travail étant extrêmement long même si l'on se

borne à des irrégularités simples de la ligne de transmission.

Dans ces conditions, nous avons calculé numériquement un mode de correction de diaphonie dans le cas facile à traiter de la diaphonie intrinsèque en régime permanent.

Examinons donc la méthode proposée pour corriger la diaphonie de f_1 sur f_2 . Nous savons que, en l'absence de correction, l'émission étant faite sur le canal N° 1 :

(9) $y_1 = f_1 + f_1 D [f_2/f_1]$

De même, en émettant f_2 sur le canal N° 2, on reçoit sur ce canal :

(10) $y_2 = f_2 + f_2 D [f_1/f_2]$

On peut songer à atténuer les phénomènes de d'aphonie, en remplaçant f_1 et f_2 , à l'émission, par :

(11) $\begin{cases} \varphi_1 = f_1 - f_1 D [f_2/f_1] \\ \varphi_2 = f_2 - f_2 D [f_1/f_2] \end{cases}$

La figure 16 donne la représentation graphique de δ_1 . Sur la même figure on a reporté D comme terme de comparaison. On voit ainsi qu'une atténuation très notable a été obtenue. Toutefois, la valeur de 14 % observée pour δ_1 paraît encore élevée.

On peut alors envisager d'accentuer la correction faite préventivement à l'émission, en remplaçant les valeurs (11) par ;

(14) $\begin{cases} \varphi_1 = f_1 - \mu f_1 D [f_2/f_1] \\ \varphi_2 = f_2 - \mu f_2 D [f_1/f_2] \end{cases}$

avec :

(15) $\mu > 1$

Il y a lieu alors de rechercher, par tâtonnement, la valeur de μ qui conduit au phénomène de diaphonie le moins important. Nous n'avons pas effectué ce travail d'une manière très précise, et nous nous sommes bornés au calcul numérique de δ_1 pour la valeur :

(16) $\mu = 1,5$

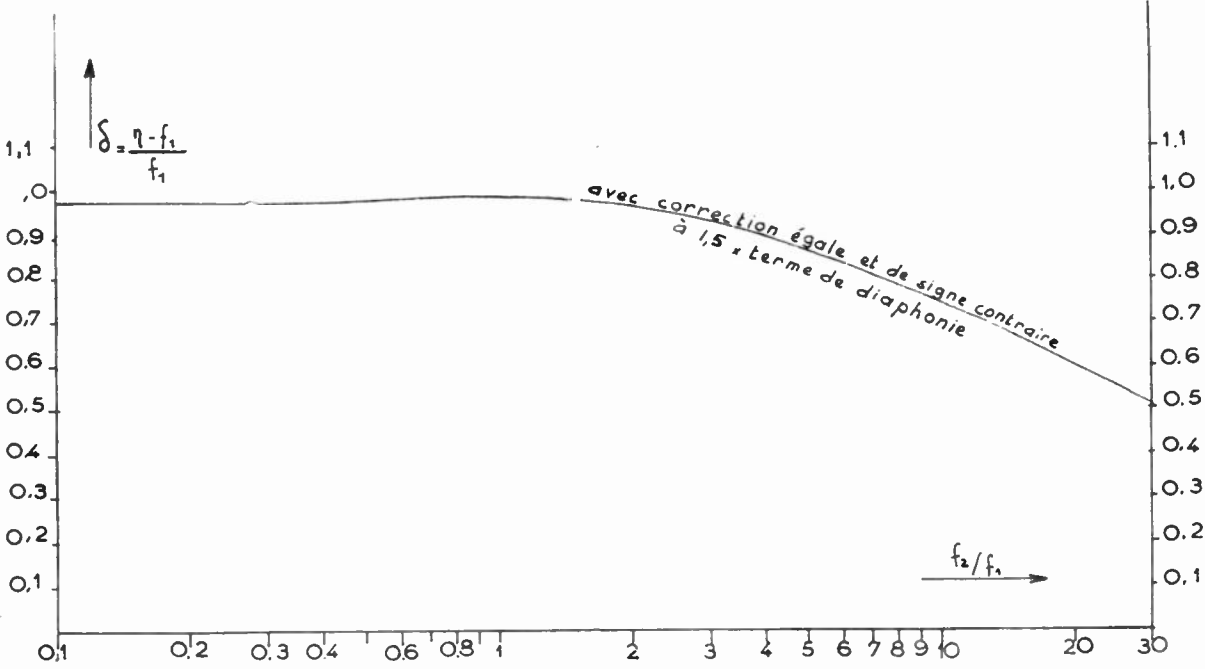


FIG. 17.

Le signal récupéré à la réception sur le canal 1, par exemple, sera :

(12) $\eta_1 = \left\{ f_1 - f_1 D [f_2/f_1] \right\} \left\{ 1 + D \left[\frac{f_2 - f_2 D (f_1/f_2)}{f_1 - f_1 D (f_2/f_1)} \right] \right\}$

On peut calculer numériquement η_1 en fonction de f_1 et f_2/f_1 , et en utilisant d'ailleurs pour D les valeurs numériques tirées de la figure 14. On peut ainsi obtenir le nouveau coefficient de diaphonie :

(13) $\delta_1 = \frac{\eta_1 - f_1}{f_1}$

Le résultat est représenté figure 17. On voit que δ_1 est maintenant très réduit, et la figure 18 donne, dans les mêmes conditions, le terme de diaphonie lui-même dans le cas où f_2 est transmis avec sa valeur maximum qui est l'unité, et pour f_1 variant de 0 à 1. On voit que ce terme ne dépasse pas 0,029. Dans ces conditions, les phénomènes de diaphonie peuvent, tout au moins pour certaines applications, et en particulier en télévision, être tenus pour négligeables.

On peut d'ailleurs constater sur la figure 17, que δ_1 demeure toujours négatif. La valeur $\mu = 1,5$ n'est pas exactement la valeur optimum, qui est un peu supérieure.

E. — Mode de réalisation de la correction.

Pour opérer la correction, il suffit de disposer des quantités $D [f_2/f_1]$ et $D [f_1/f_2]$.

Pour cela, à l'émission, on fabrique une première fois le signal Σ' , à partir de f_1 et f_2 , et ce signal est envoyé sur des circuits à diode qui envoient sur un premier canal la partie négative, c'est-à-dire respectivement y_1 et y_2 .

Par un procédé que nous décrirons plus loin, on se débarrasse de la « structure de points » qui affecte

nue est alors utilisée pour moduler le système de transmission.

Pour éviter la « structure de points » dans y_1 et y_2 , on fabrique le premier signal Σ'_1 , non pas au moyen de la fréquence de sondage ω_s , mais au moyen de la fréquence $2\omega_s$, par exemple. Dans ces conditions, les bandes de f_1 et f_2 étant limitées à ω_s , on réalise un signal de multiplex binaire non équentiel. Le mode de réception employé ici, qui n'est pas le mode normal d'un multiplex à impulsions modulées en amplitude, laisse subsister dans y_1

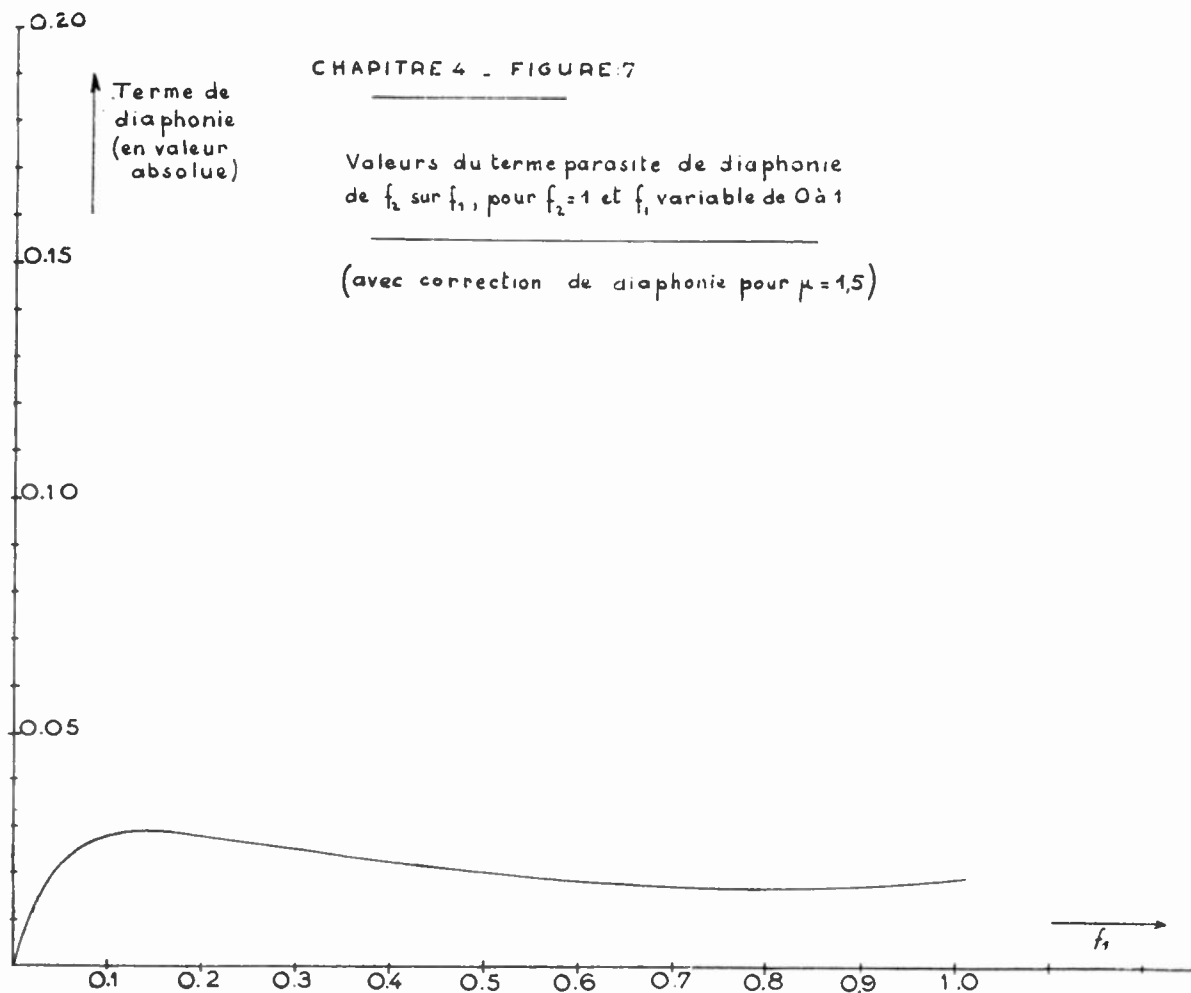


Fig. 18.

chacun de ces 2 signaux, et il ne reste plus qu'à effectuer les opérations

$$(17) \quad y_1 - f_1 \quad \text{et} \quad y_2 - f_2$$

pour obtenir les termes de diaphonie.

On fabrique alors :

$$(18) \quad f_1 - f_1 \mu D [f_2/f_1] \quad \text{et} \quad f_2 - f_2 \mu D [f_1/f_2]$$

qui sont maintenant considérées comme les deux informations à transmettre, et l'on élabore Σ'_2 à partir de ces signaux. La fonction Σ'_2 ainsi obtenue

est y_2 une structure de points, mais à la fréquence $2\omega_s$. Le passage de chacun de ces signaux à travers un filtre passe-bas coupant à ω_s élimine cette structure. On peut montrer, d'autre part, que, moyennant un choix convenable de la courbe d'atténuation du filtre limitant la bande de Σ'_2 à $2\omega_s$, les phénomènes de diaphonie intrinsèques sont les mêmes que si le signal Σ'_2 était obtenu par un sondage à la fréquence ω_s .

F. — Extensions de la correction.

L'efficacité de la correction est évidemment

ndépendante de la cause du phénomène de diaphonie considéré. Si donc une partie de la chaîne de transmission provoque de la diaphonie, on a intérêt à insérer, entre le générateur de Σ_1' et l'ensemble de circuits détecteurs, un dispositif simulant la partie de la chaîne de transmission qui provoque cette diaphonie.

Il conviendra, par exemple, d'examiner dans quelle mesure il y a lieu de tenir compte du Γ des kinescopes récepteurs dans l'établissement du système de correction. De même, le mode classique de réception de la télévision sur bande latérale unique peut amener d'autres termes de diaphonie.

Les dimensions du présent exposé étant limitées, nous ne pouvons envisager de développer davantage ces questions, qui font intervenir des calculs

généralement très longs. Dans certains cas, et en particulier pour l'évaluation de l'influence de la réception sur bande latérale unique, en régime transitoire, les calculs peuvent être si longs que l'expérimentation apparaisse comme un moyen moins coûteux d'obtenir les résultats cherchés.

Nous bornerons donc ici cet exposé en concluant par l'intérêt de recherches expérimentales. Ces recherches sont commencées, tant dans le domaine de la télévision en couleurs (utilisation de la transmission actuelle pour le transport de l'information verte et de l'information rouge) que dans le domaine de la transmission de 2 programmes indépendants en noir et blanc. Des résultats partiels sont, d'ores et déjà, obtenus, qui seront publiés ultérieurement à l'occasion de rapports d'ensemble sur l'aspect expérimental de la question.

PRINCIPES DU SYSTÈME DE TÉLÉVISION EN COULEUR SIMULTANÉE N.T.S.C.

PAR

M. Charles HIRSCH

INTRODUCTION :

L'objet de cet article est d'exposer les principes du système de télévision en couleur simultanée qui a été développé par le « *National Television System Committee* » (NTSC), et qui est maintenant en exploitation aux Etats-Unis. Ce système est « compatible ». Cela signifie, que les images de télévision en couleur peuvent également être reçues en noir et blanc, par les récepteurs existants, sans qu'il y ait lieu de les modifier. Inversement, les récepteurs de télévision en couleur peuvent aussi recevoir en noir et blanc les images transmises par les systèmes de télévision monochrome.

Brièvement, le principe du système consiste à ajouter la couleur aux images en noir et blanc, de la même manière, qu'un photographe colore des clichés à l'aide de crayons de couleur ou de peinture.

L'exposé qui suit se rapporte au standard utilisé aux Etats-Unis, dont les principales caractéristiques sont les suivantes :

- a) 525 lignes entrelacées une fois sur deux ;
- b) 60 trames, donc 30 images par seconde ;
- c) modulation négative ;
- d) largeur de bande du canal de transmission, égale à 6 Mc/s ;
- e) le son est transmis en modulation de fréquence à l'aide d'une fréquence porteuse supérieure de 4,5 Mc/s à la fréquence porteuse image.

CONSIDÉRATIONS GÉNÉRALES.

La reproduction de la couleur exige la définition des trois quantités indépendantes, que sont les intensités respectives des trois couleurs primaires, rouge, vert et bleu. Pour un système de télévision en couleur, il pourrait en résulter l'emploi de trois dispositifs complets de reproduction d'image, relatifs à chaque couleur primaire. Néanmoins, il est un fait, que si l'œil peut distinguer de petits changements de brillance (à la fois dans l'espace et dans le temps), il est, par contre, beaucoup moins sensible aux chan-

gements de couleur. C'est pourquoi, il n'est pas nécessaire que les images en couleur contiennent trois fois plus d'informations que les images monochromes.

En fait on peut obtenir une image en couleur satisfaisante avec très peu plus d'information (10 à 80 %), que pour la même image monochrome^{1,2,3}. Pour parvenir à ce résultat, la couleur est exprimée en une autre série de trois quantités : brillance, longueur d'onde dominante et pureté (qui correspondent à clarté, nuance et saturation) ; et l'image en couleur est transmise comme une image monochrome complètement définie, à laquelle on ajoute l'information minimum requise par la couleur.

Dans le système de télévision en couleur, développé par le « N T S C », l'information est transmise au moyen de deux signaux simultanés. L'un des deux est appelé signal de brillance (luminance), et porte toute l'information brillance (clarté). Il est transmis actuellement par le standard FCC pour la télévision en noir et blanc, et peut être reçu par un récepteur monochrome sans modification.

L'autre signal est appelé « Sous-porteuse couleur » (chrominance Subcarrier). Ce signal porte l'information couleur, qui peut être électriquement ajoutée au signal de brillance, et qui, appliqué à un tube récepteur tricolore, permettra la reproduction de l'image en couleur.

Ces signaux sont constitués de telle façon, que chaque élément d'image est reproduit instantanément dans sa propre couleur, et non séquentiellement. Ainsi, le pourpre résulte d'une combinaison simultanée du rouge et du bleu, et non de la superposition d'une trame rouge et d'une trame bleue comme dans le système de télévision séquentielle.

PARTAGE DE LA BANDE DE FRÉQUENCE ENTRE LES SIGNAUX DE BRILLANCE ET DE COULEUR.

Les signaux de brillance et de couleur, occupent la même bande de fréquence, que celle qui est normalement nécessaire à la transmission des images monochromes. Cela est possible, parce que le spectre des images de télévision se compose essentiellement de bandes étroites de fréquences, dont l'énergie

est en grande partie concentrée, près des harmoniques de la fréquence de balayage ligne (harmoniques pairs de la fréquence moitié de balayage ligne)^{2, 5, 6}.

Ce spectre résulte du découpage des images de télévision en lignes successives ; chaque image contient de ce fait, une très grande quantité de

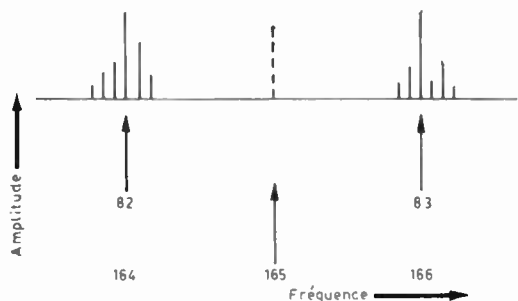
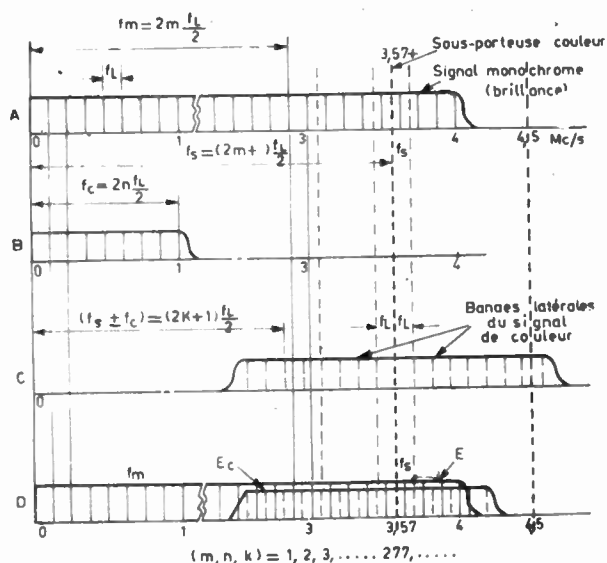


FIG. 1. — Concentration de l'énergie dans un signal de télévision. Les nombres du haut sont harmoniques de la fréquence ligne et les nombres du bas, harmoniques de la demi-fréquence ligne.

redondances ; et son spectre peut, par conséquent, être exprimé approximativement par une série de Fourier.

Le spectre du signal de couleur comporte également de semblables concentrations d'énergie, et il est possible, de les intercaler avec celles du spectre monochrome. Pour cela il suffit de les localiser aux



par l'œil était intégrale d'une image à la suivante. Pratiquement, ce n'est pas tout à fait le cas, et l'œil n'efface que la partie du signal de couleur, qu'il a retenu d'une image à l'autre. Pour cette raison, une partie seulement de l'amplitude de l'information « couleur » est éliminée.

Aussi, si le signal de couleur est transmis à niveau trop élevé, les crêtes de la sous-porteuse et des bandes latérales apparaîtront sous formes de « points » sur un récepteur monochrome.

C'est pourquoi, on a choisi une fréquence élevée (3,579545 Mc/s) comme sous-porteuse couleur. Comme cette fréquence est déjà considérablement atténuée dans les récepteurs monochromes existants on peut la transmettre sans inconvénient à niveau élevé. De cette façon, on rend pratiquement invisibles les « points » sur les récepteurs monochromes,

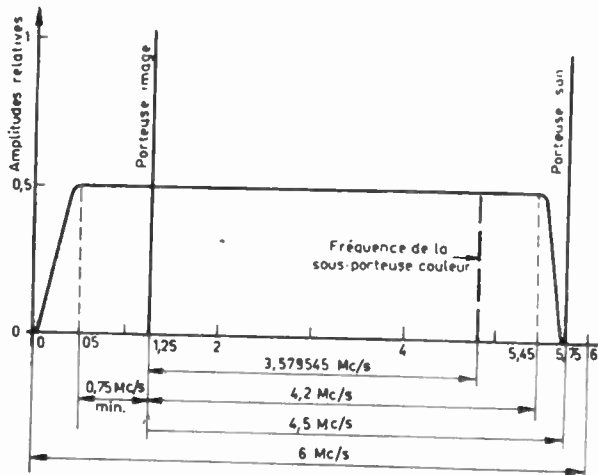


FIG. 4. — Caractéristique d'amplitude idéale de transmission d'image (Amplitude de la portuse-image égale à 1).

tout en assurant une valeur suffisante au rapport signal/bruit dans les récepteurs de télévision en couleur. La fréquence portuse choisie pour transmettre la couleur est égale au 455^e harmonique de la demi-fréquence ligne (voir figure 2).

La figure 4 montre la position des portuses images, et son, et de la sous-porteuse couleur, à l'intérieur des 6 Mc/s du canal de transmission.

SIGNAL DE BRILLANCE (luminance).

La tension du signal de brillance, E_Y , peut être obtenue directement, à l'aide d'une caméra, dont la sortie est proportionnelle à la brillance. Plus généralement, ce signal est obtenu, par combinaison des tensions (E_R , E_V , et E_B relatives aux primaires rouge, vert et bleu) que délivre une caméra tricolore. Dans ce dernier cas, les trois composantes sont combinées proportionnellement à leur contribution à la brillance totale :

$$E_Y = 0,59 E_V + 0,30 E_R + 0,11 E_B \quad (1)$$

Cette expression indique, que les primaires (1) vert, rouge et bleu, contribuent respectivement,

(1) Les primaires standard sont : rouge : $x = 0,670$; $y = 0,330$; vert : $x = 0,210$; $y = 0,710$; bleu : $x = 0,140$; $y = 0,080$ (voir figure 5).

pour 59, 30 et 11 pour cent, à la brillance de la lumière blanche (définie par les coordonnées chromatiques $x = 0,310$; $y = 0,316$, de la source Standard C).

Remarquons que la somme des facteurs numériques de l'équation (1) est égale à l'unité.

Le système est d'autre part proportionné de façon que le blanc soit obtenu quand $E_R = E_V = E_B$. En substituant dans l'équation (1), on obtient donc pour la lumière blanche :

$$E_Y = E_R = E_V = E_B \quad (1a)$$

Quand il n'y a pas de couleur, il est évidemment souhaitable que l'information « couleur » disparaisse.

Aussi la transmet-on à l'aide des deux composantes ($E_R - E_Y$) et ($E_B - E_Y$) qui sont appelées signaux « différence de couleur » (color-difference signals)¹⁰.

De l'équation (1) on tire dans le cas de la lumière blanche (pas de couleur) :

$$(E_R - E_Y) = 0 = (E_B - E_Y).$$

SIGNAUX « différence de couleur ».

C'est parce que l'œil est insensible aux détails fins, en ce qui concerne la couleur, que ces signaux ont été limités en bande de fréquence.

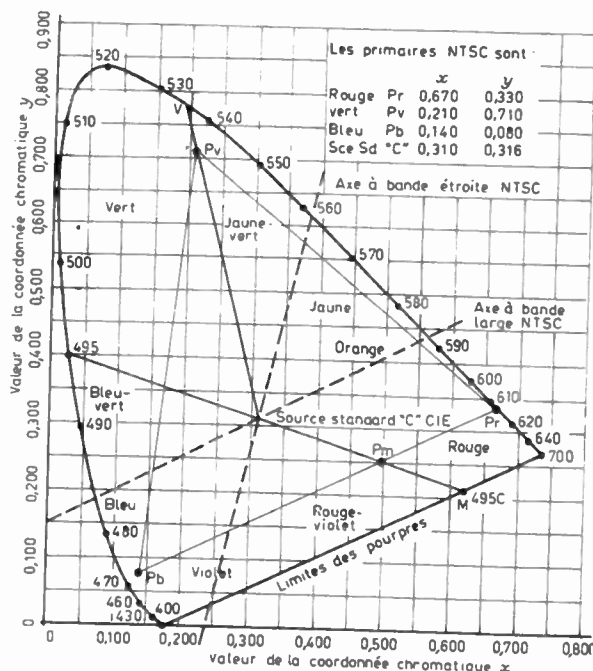


FIG. 5. — Diagramme chromatique CIE

On a trouvé que l'œil était moins sensible aux détails de couleur, dans la région moyenne reliant le jaune-vert au pourpre, axe chromatique (appelé axe Q), que pour celles situées sur l'axe reliant l'orange au cyan (bleu-vert) (appelé axe I) (voir figure 5). Les signaux « différence de couleur »

correspondant à ces axes sont respectivement appelés E_Q et E_I . Ils sont liés aux expressions $E_R - E_Y$ et $E_B - E_Y$ par les relations suivantes :

$$E_Q = 0,48 (E_R - E_Y) + 0,41 (E_B - E_Y) \quad (2a)$$

$$E_I = 0,74 (E_R - E_Y) - 0,27 (E_B - E_Y) \quad (2b)$$

Par conséquent E_Q peut être (et est) transmis avec une bande de fréquence d'environ 0,5 Mc/s, tandis que pour E_I on utilise une bande d'environ 1,3 Mc/s ou plus.

Cette façon de procéder, permet d'éliminer les défauts que l'on constaterait aux frontières séparant deux couleurs, comme on le montrera plus loin.

Puisque les dispositifs existants, ne peuvent fonctionner qu'en fonction des primaires rouge, vert et bleu, le récepteur doit reconvertir E_Q et E_I en $E_R - E_Y$, $E_B - E_Y$, et $E_V - E_Y$.

Bien qu'il n'apparaisse pas explicitement, le vert, quand il existe, est également transmis par ces signaux. Il est, selon l'équation (1), le principal composant de E_Y . Au récepteur, qui reçoit $(E_R - E_Y)$ et $(E_B - E_Y)$, on peut obtenir $(E_V - E_Y)$ par mélange de $-0,51 (E_R - E_Y)$ et $-0,19 (E_B - E_Y)$ de la façon exposée ci-dessous.

En remplaçant E_Y par sa valeur tirée de l'équation (1), nous obtenons :

$$\begin{aligned} & -0,51 (E_R - E_Y) - 0,19 (E_B - E_Y) \\ = & -0,51 (-0,59 E_V + 0,70 E_R - 0,11 E_B) \\ & -0,19 (-0,59 E_V - 0,30 E_R + 0,89 E_B) \\ = & 0,41 E_V - 0,30 E_R - 0,11 E_B \\ = & E_V - (0,59 E_V + 0,30 E_R + 0,11 E_B) \\ = & E_V - E_Y \end{aligned} \quad (3)$$

Dans le récepteur on ajoute aussi le signal de brillance à chaque signal « différence de couleur » de la façon suivante :

$$(E_R - E_Y) + E_Y = E_R \quad (4a)$$

$$(E_B - E_Y) + E_Y = E_B \quad (4b)$$

$$(E_V - E_Y) + E_Y = E_V \quad (4c)$$

Pour appliquer finalement les tensions E_R , E_B et E_V , respectivement entre les grilles de contrôle et les cathodes des trois canons à électron du tube tricolore, il suffit d'appliquer E_Y à une électrode et le signal « différence de couleur » à l'autre.

BRILLANCE CONSTANTE DU SYSTÈME ¹⁰.

Toutes les variations de brillance sont apportées par E_Y . Comme le montre le raisonnement suivant, les trois signaux « différence de couleur », agissant ensemble, ne contribuent pas aux variations de brillance.

Supposons que la variation de brillance, produite par chaque canon, soit directement proportionnelle à la tension appliquée. Dans ce cas les variations produites par les signaux « différence de couleur » sont :

$$a) \text{ pour le canon vert } Y_V = K \times 0,59 (E_V - E_Y)$$

$$b) \text{ pour le canon rouge } Y_R = K \times 0,30 (E_R - E_Y)$$

$$c) \text{ pour le canon bleu } Y_B = K \times 0,11 (E_B - E_Y),$$

où K est une constante de transfert qui a pour dimensions, des « foot-lamberts » par volt ; et les coefficients 0,59, 0,30, et 0,11 sont les mêmes que dans l'équation (1).

La variation totale de brillance, Y_{CD} , produite par les signaux « différence de couleur » est donc :

$$Y_{CD} = Y_V + Y_R + Y_B = K (0,59 E_V + 0,30 E_R + 0,11 E_B - E_Y)$$

Mais comme d'après (1) :

$$0,59 E_V + 0,30 E_R + 0,11 E_B = E_Y$$

On obtient finalement :

$$Y_{CD} = K (E_Y - E_Y) = 0$$

Cette discussion montre bien que dans le cas d'un système linéaire, les signaux « différence de couleur », lorsqu'ils agissent ensemble, ne contribuent pas aux variations de lumière. Pour cette raison, ce système est considéré comme étant à brillance constante.

CORRECTION GAMMA.

La discussion précédente est très simplifiée, parce que, la diffusion de la lumière (L) de l'image reproduite sur l'écran du tube n'est pas directement proportionnelle à la tension d'entrée (E), mais varie approximativement comme une puissance (γ) de cette entrée :

$$L = K E^\gamma \quad (5)$$

Les tensions appliquées au tube devront, en conséquence, être pré-correctées au moyen d'un procédé appelé « Correction Gamma ».

Un moyen d'effectuer cette correction consiste à transmettre les signaux suivants :

1) Un signal monochrome constitué par les tensions primaires corrigées tel que :

$$E'_Y = 0,59 E_V^{1/\gamma} + 0,30 E_R^{1/\gamma} + 0,11 E_B^{1/\gamma} \quad (6)$$

2) Les deux signaux « différence de couleur » :

$$(E_R^{1/\gamma} - E_Y^{1/\gamma}) \text{ et } (E_B^{1/\gamma} - E_Y^{1/\gamma}).$$

Note. On utilise aussi bien les symboles prime, tel que E' , que les symboles $E^{1/\gamma}$ pour désigner la correction gamma.

SIGNAL DE COULEUR COMPLET.

La figure 6, montre le schéma théorique d'un dispositif complet de production de signaux de télévision en couleur, dans lequel on utilise des corrections gamma individuelles sur chaque tension primaire. Ces tensions sont ensuite combinées pour former le signal monochrome et les signaux « différence de couleur ».

Ou encore ce qui est essentiellement la même chose, en constituant cette onde de deux composantes en quadrature, dont les amplitudes sont modulées respectivement par chaque information.

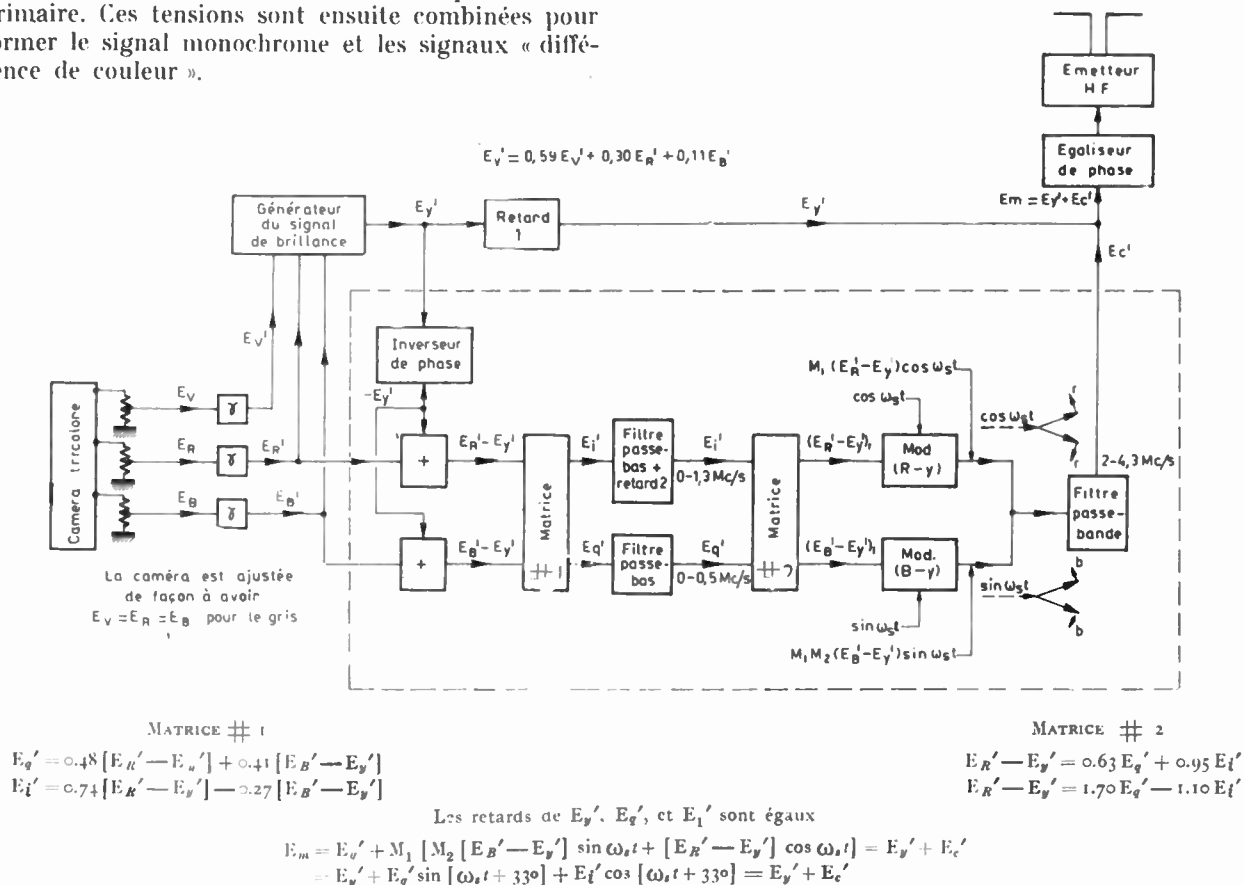


FIG. 6. — Schéma théorique d'un émetteur

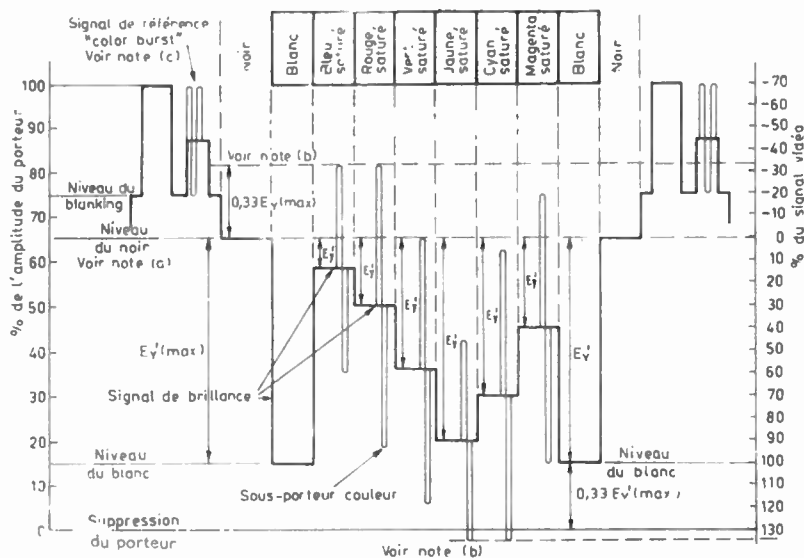


Fig. 7. — Signal de couleur N.T.S.C. complet pour une ligne dans le cas de barres de couleurs saturées de brillance maximum.

Il est possible de transmettre deux informations différentes avec une onde sinusoïdale. Par exemple, en modulant son amplitude avec l'une, et sa phase avec l'autre.

Chaque modulation, peut ensuite être reconstituée, à l'aide d'un montage hétérodyne, auquel on applique, d'une part l'onde modulée, et d'autre part une onde sinusoïdale de même fréquence en phase

NOTES

- (a) Le niveau du noir est arbitrairement choisi.
- (b) Le dépassement du niveau de blanking et du niveau de suppression du porteur, par les crêtes de modulation de la sous-porteuse couleur, dépend du choix des niveaux correspondant au noir et au blanc.
- (c) La fréquence de la sous-porteuse couleur, n'est pas représentée à l'échelle des temps.

avec le porteur correspondant à la modulation désirée. Ce procédé est quelquefois appelé « détection synchrone », et ne doit pas être confondu avec d'autres formes de détection, qui restituent l'enveloppe de modulation.

La référence de phase de l'onde modulée par les signaux de couleur est transmise à l'aide de quelques périodes de cette onde non modulée, durant le blanking qui suit les impulsions de synchronisation de ligne (voir figure 7)^{8, 15}.

Ce signal est appelé « color burst ». Sa fréquence est donc celle de la sous-porteuse couleur (3.579.545 c/s), et il est à 180° du porteur modulé par la composante « différence de couleur » ($E_B^{1/2} - E_Y$).

La figure 6 montre que $(E_R' - E_Y')$ et $(E_B' - E_Y')$ sont combinées dans la matrice $\#^1$, pour former E_Q' et E_I' , dont les bandes de fréquence, sont respectivement limitées à 0,5 et 1,3 Mc/s. E_Q' et E_I' sont ensuite mélangées de nouveau dans la matrice $\#^2$ pour reformer $(E_R' - E_Y')$ et $(E_B' - E_Y')$.

$(E_B^{1/2} - E_Y')$ module $\sin \omega t$, tandis que $(E_R^{1/2} - E_Y')$ module $\cos \omega t$. Des modulateurs équilibrés sont utilisés, de sorte que la sous-porteuse est supprimée dans le cas de la lumière blanche. Les sorties des deux modulateurs, sont d'abord combinées entre elles, pour former un seul signal « porteur de couleur », qui est ensuite lui-même combiné avec le signal monochrome, pour constituer le signal complet de l'image en couleur E_m dont l'équation est :

$$E_m = E_Y' + M_1 \{ M_2 (E_B^{1/2} - E_Y') \sin \omega t + (E_R^{1/2} - E_Y') \cos \omega t \}$$

(7)

Dans cette formule, la référence de phase est $\sin \omega t$ (elle est à 180° du signal de synchronisation de couleur). M_1 détermine l'amplitude de la sous-porteuse couleur par rapport à celle du signal monochrome E_Y' ; et M_2 fixe les proportions relatives des deux composantes « différence de couleur ».

COMPOSITION DU SIGNAL DE COULEUR.

Le sous-comité 13 du N T S C., a fixé la composition du signal de couleur complet, en accordance avec le critérium suivant :

« La composition du signal de couleur doit être telle, qu'à partir des signaux « différence de couleur », rouge et bleu, les primaires définies par le sous-comité 7, puissent être reconstituées directement, par démodulation de deux porteurs en quadrature. De plus les amplitudes relatives des deux signaux « différence de couleur » seront telles, que des surcharges de l'ordre de 1/3 correspondront aux maxima d'intensité des couleurs primaires, définies par le sous-comité 7, ou à leurs complémentaires ».

Ce résultat est obtenu en faisant $M_1 = 0,88$ et $M_2 = 0,56$ dans l'équation (6), qui devient maintenant :

$$E_m = E_Y' + 0,88 \{ 0,56 (E_B^{1/2} - E_Y') \sin \omega t + (E_R^{1/2} - E_Y') \cos \omega t \}$$

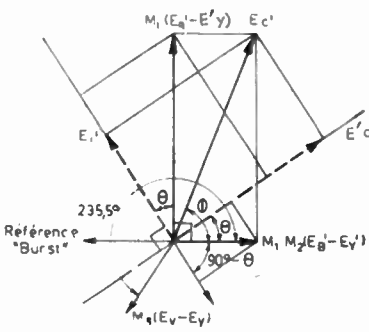
(7a)

Ce signal est montré figure 7, en ce qui concerne une ligne lors de la transmission de barres verticales noires et blanches, et de barres de couleurs saturées bleu, rouge, vert, jaune, cyan et magenta. Les intensités sont mesurées à partir du noir. Sur la figure le blanc correspond à 15 % et le noir à 65 % de l'intensité maximum du porteur ; c'est-à-dire une proportion de 50 % pour le signal video complet. Cette proportion n'a été prise que pour illustrer convenablement cet article. Si l'on fait l'amplitude du blanc égale à 1, les barres de couleur sont obtenues comme l'indique le tableau I.

TABEAU I.

Couleur des barres	noir	blanc	Couleur saturée					
			bleu	rouge	vert	jaune	cyan (bleu-vert)	magenta (pourpre)
composante								
$E_B^{1/2}$	0	1	1	0	0	0	1	1
$E_R^{1/2}$	0	1	0	1	0	1	0	1
$E_V^{1/2}$	0	1	0	0	1	1	1	0
E_Y'	0	1	0,11	0,30	0,59	0,89	0,70	0,41

Sur la figure 7, les valeurs de E_Y' forment une ligne brisée, et quelques périodes de la sous-porteuse couleur modulée sont superposées à E_Y' qui sert de base.



$$E'_C = M_1 \sqrt{M_2^2 [E'_B - E'_Y]^2 + [E'_R - E'_Y]^2} \angle \Phi = \sqrt{E'^2_q + E'^2_i} \angle \Phi$$

$$\Phi = 33^\circ \text{ ou } : \Phi = \tan^{-1} \frac{E'_R - E'_Y}{M_2 [E'_B - E'_Y]} = 33^\circ + \tan^{-1} \frac{E'_1}{E'_2}$$

$$E'_q = M_1 M_2 [E'_B - E'_Y] \cos \theta + M_1 [E'_R - E'_Y] \cos (90^\circ - \theta)$$

$$E'_1 = M_1 M_2 [E'_R - E'_Y] \cos (90^\circ - \theta) + M_1 [E'_B - E'_Y] \cos \theta$$

FIG. 8. -- Diagramme vectoriel du signal de couleur.

Le signal de couleur peut être calculé, en substituant les valeurs données par le tableau I, dans l'équation 7a.

La crête du vert saturé est exactement tangente

Le signal porteur de couleur complet, est appliqué à chaque démodulateur. Un porteur local synchronisé est produit au récepteur. Une composante ayant pour phase $\cos(\omega t + 33^\circ)$ est appliquée au démodulateur du signal, « différence de couleur » E_I . Une autre composante, de phase $\sin(\omega t + 33^\circ)$ est appliquée au démodulateur du signal « différence de couleur » E_Q .

DÉMODULATEURS SYNCHRONES.

Le fonctionnement du démodulateur synchrone est illustré par la figure 10. Un signal d'entrée $E(t) \sin \omega t$ est appliqué à l'une des grilles d'une lampe multi-grille, par exemple du type 6 AS 6.

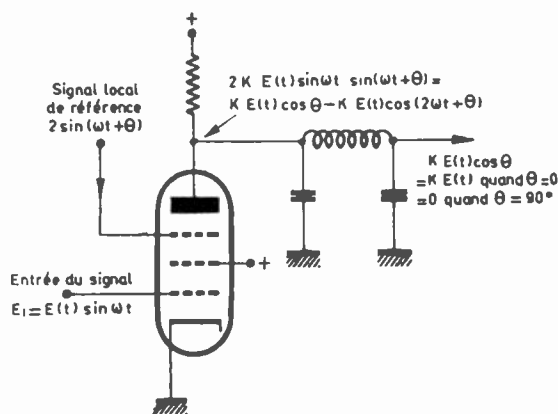


FIG. 10. — Principe de la démodulation synchrone.

Un signal de référence de phase, de même fréquence, mais dont la phase peut différer de $2 \sin(\omega t + \theta)$ est appliqué à une autre grille. Le courant plaque est proportionnel au produit des deux tensions et est égal à :

$$i_p = 2 K E(t) \sin \omega t \sin(\omega t + \theta)$$

Le produit des deux sinus donne ici :

$$i_p = 2 K E(t) [1/2 \cos \theta - 1/2 \cos(2\omega t + \theta)].$$

Le terme à fréquence double, $KE(t) \cos(2\omega t + \theta)$, est éliminé à l'aide du filtre passe-bas, inséré dans la plaque du modulateur, de sorte que l'on recueille uniquement :

$$i_p = K E(t) \cos \theta$$

En d'autres termes, la sortie du démodulateur est proportionnelle à l'amplitude de modulation, $E(t)$, du signal d'entrée, multipliée par le cosinus du déphasage θ , entre le porteur de couleur et le signal de référence. Le courant plaque est par conséquent proportionnel au déphasage de la composante du signal porteur de couleur, avec la direction de référence.

En se référant à l'équation 7b, on voit que les deux composantes du signal de couleur E'_Q et E'_I , modulent deux composantes en quadrature de la sous-porteuse couleur.

Les signaux « différence de couleur », E'_Q et E'_I , peuvent par conséquent être séparés, par hétérodyne du signal de couleur, dans un cas avec un signal ayant la phase de la composante de la sous-porteuse couleur, que module E'_Q (c'est-à-dire $\sin[\omega t + 33^\circ]$); et dans l'autre cas avec un signal ayant la phase de la composante de la sous-porteuse couleur que module E'_I (c'est-à-dire $\cos[\omega t + 33^\circ]$). C'est ce que montre la figure 9.

Si, néanmoins, le signal hétérodyne a un déphasage différent de 0 ou 90° par rapport aux deux composantes « différence de couleur », la sortie sera un mélange des deux.

En se référant de nouveau à la figure 9, on voit que les bandes de fréquence occupées par les signaux « différence de couleur » E'_Q et E'_I , à la sortie de leurs démodulateurs, sont limitées respectivement à 0,5 et 1,3 Mc/s.

E'_Q et E'_I sont ensuite mélangés dans les matrices $R - Y$ et $B - Y$, pour reconstituer $E'_R - E'_Y$ et $E'_B - E'_Y$.

Ces dernières étant elles-mêmes mélangées, pour reproduire $E'_V - E'_Y$.

Les trois signaux « différence de couleur » peuvent ensuite être appliqués aux cathodes respectives du tube tricolore de réception d'image.

Puisque le fonctionnement du tube, dépend de la tension grille-cathode, la sortie de la lumière pour le canon vert est :

$$L = K (E_V^{1/2} - E'_Y + E'_Y) = K E_V$$

Il en est de même pour les deux autres canons.

Remarquons que tous les éléments d'information, nécessaires à la définition d'un élément d'image, c'est-à-dire la brillance (E'_Y) et la couleur ($E_R^{1/2} - E'_Y$; $E_V^{1/2} - E'_Y$; et $E_B^{1/2} - E'_Y$), sont présents simultanément.

MÉLANGE ENTRE LES COULEURS (DIAPHOTIE), DU AU FONCTIONNEMENT A SIMPLE BANDE LATÉRALE DE MODULATION.

Nous avons vu qu'il était souhaitable d'utiliser une fréquence élevée (3.579.545 c/s) pour la sous-porteuse couleur, de façon à réduire son effet sur les récepteurs monochromes. Mais ceci impose une limitation de la bande latérale supérieure de modulation utilisée pour la composante du signal de couleur, puisque la bande totale video est limitée en pratique à moins de 4,5 Mc/s.

Par contre, la bande latérale inférieure, peut être beaucoup plus étendue (1 à 2 Mc/s), au-dessous de la sous-porteuse couleur.

De l'inégalité de ces bandes latérales, il résulte une diaphonie d'une composante sur l'autre du signal de couleur.

De la modulation à simple bande latérale, résulte le partage de la puissance, en deux parties égales,

entre les modulations d'amplitude et de phase ; ou si l'on préfère, il en résulte l'existence de deux bandes latérales égales, dont l'une est en phase, et l'autre en quadrature avec le porteur.

Ceci est montré figure 11.

Les figures 11 (a 1) et 11 (a 2) montrent, à l'aide de diagrammes de fréquence et de temps, la relation entre une sous-porteuse de couleur $E_o = \cos (\omega t + 33^\circ)$, qui peut être supprimée, et une bande latérale inférieure $E_L \cos [(\omega - \omega_L) t + 33^\circ]$.

La phase de E_o est arbitrairement prise comme référence de phase, qui, pour la simplicité de la discussion, se trouve être également la phase du canal « différence de couleur » I.

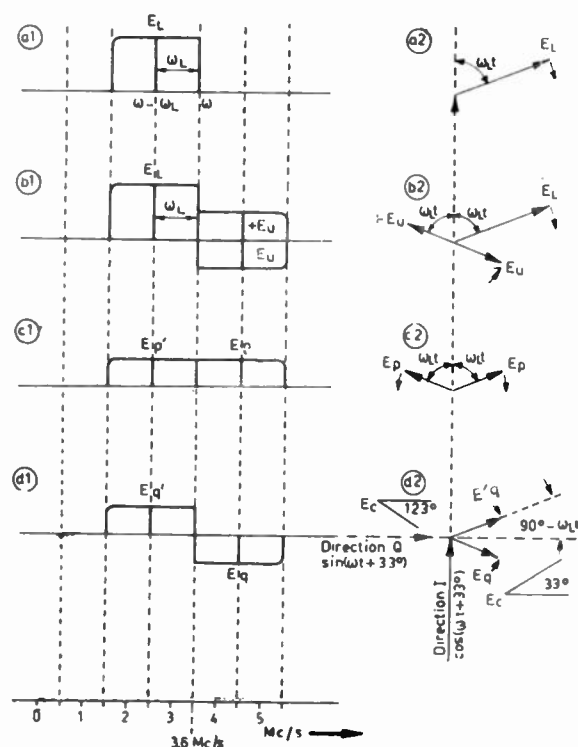


FIG. 11. — Analyse de la transmission à simple bande latérale.

Les figures 11 (b 1) et 11 (b 2), représentent le même signal, auquel on a ajouté, symétriquement autour de la référence de phase, la bande latérale supérieure :

$$+ E_U = \frac{E_L}{2} \cos [(\omega + \omega_L) t + 33^\circ],$$

dont l'amplitude est la moitié de celle de la bande inférieure.

On a également ajouté un autre signal :

$$- E_U = - \frac{E_L}{2} \cos [(\omega + \omega_L) t + 33^\circ],$$

égal, mais de phase opposée à $+ E_U$, de façon à rendre identique le signal de la figure 11 (b) à celui de la figure 11 (a).

Les bandes latérales, montrées dans la figure 11 (b), peuvent maintenant être décomposées en deux séries de bandes latérales égales. Une de ces séries (E_p et E_p') est montrée dans les figures 11 (c 1) et 11 (c 2). Elle est symétriquement disposée, autour de la phase de référence ($\cos \omega t + 33^\circ$). Elle représente une modulation d'amplitude pure de $\cos (\omega t + 33^\circ)$ (c'est le canal I de la sous-porteuse). La seconde série (E_q et E_q') est symétriquement disposée autour d'un second porteur, qui est en quadrature avec le porteur de référence. E_q et E_q' , représentent par conséquent une modulation d'amplitude de $\sin (\omega t + 33^\circ)$ (ce qui correspond au canal Q de la sous-porteuse).

Il est clair, que la somme des signaux des figures 11 (d) et 11 (c) est égale au signal de la figure 11 (b), et par conséquent à celui de la figure 11 (a).

Il est important de remarquer que l'enveloppe de la composante produite par E_q et E_q' , est en quadrature avec l'enveloppe de la composante produite par E_p et E_p' , et que les deux séries de bandes latérales en quadrature sont égales.

Ainsi, un signal du canal I, modulant en amplitude un porteur, et produisant deux bandes latérales de modulation, aura, en perdant l'une des bandes latérales, une partie de son énergie transférée à une composante, en quadrature avec le porteur, qui apparaîtra donc comme un signal parasite dans le canal Q. Ce signal parasite aura une forme différente de celui produit dans le canal I, parce que chaque fréquence le composant est déphasée de 90° par rapport à la fréquence correspondante du canal I.

ÉLIMINATION DE LA DIAPHOTIE ENTRE LES CANAUX DE COULEUR.

Cette diaphonie est éliminée de la façon suivante. En premier lieu, on transmet E_Q , à l'aide de bandes latérales symétriques. Ceci est possible, puisqu'il est limité à $0,5 \text{ Mc/s}$. Par conséquent, E_Q ne donnera pas de diaphonie sur E_I . D'autre part, E_I est transmis à l'aide de bandes latérales symétriques, pour les fréquences inférieures à $0,5 \text{ Mc/s}$ environ, et par simple bande latérale pour les fréquences

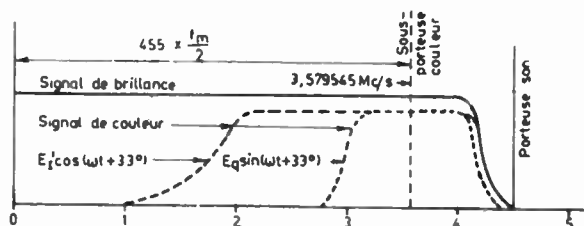


FIG. 12. — Bande de fréquence du signal de couleur.

supérieures à $0,5 \text{ Mc/s}$ (figure 12). Pour éviter la diaphonie de E_I sur E_Q , il suffit de limiter la bande de fréquence du canal E_Q à $0,5 \text{ Mc/s}$, au moyen d'un filtre passe-bas, placé à la sortie du démodulateur E_Q du récepteur.

FLUCTUATIONS DE BRILLANCE ET DE COULEUR.

L'étude des papillottements, dans le cas de lumière ayant différentes couleurs, montre que l'œil est moins sensible aux fluctuations chromatiques qu'à celles de brillance.

D'où l'idée de transformer les fluctuations de brillance dues au bruit et autres perturbations, si possible en fluctuations de couleur, puisqu'il en résultera une diminution de leur visibilité.¹⁰

Pour déterminer les valeurs relatives de la gêne apportée, dans les deux cas, le montage montré

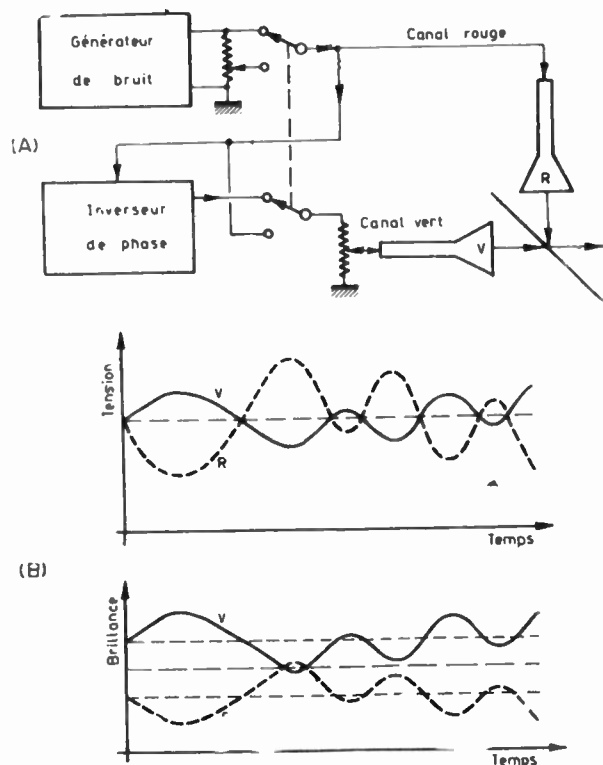


FIG. 13. — (A) Montage permettant d'évaluer les gênes relatives apportées par des fluctuations de brillance et de couleur.

(B) Fluctuations de tensions et fluctuations de brillance correspondantes dans un système de transmission à brillance constante.

figure 13 (A) a été fait. Les potentiomètres de contrôle des tubes vert et rouge sont ajustés de façon à obtenir une trame brillant d'un jaune uni. La couleur jaune est ajustée, de telle façon que par addition de bleu, on obtienne un blanc raisonnable. Puis du bruit de polarité opposée est appliqué aux deux tubes. Par ajustage de l'amplitude du bruit, appliqué au canal vert, un observateur normal trouve un point critique, pour lequel la gêne apportée par le bruit est réduite. Ce point correspond aux conditions de brillance constante, comme le montre la figure 13 (B).

On applique ensuite du bruit de même polarité, mais d'amplitude réduite, aux deux tubes, en inversant l'interrupteur double. Et l'observateur ajuste le niveau, de façon à retrouver la même gêne que celle constatée dans les conditions de brillance constante.

Les résultats obtenus par un petit groupe d'observateurs indiquent que l'on peut tolérer environ 8 dB de plus de bruit, pour éprouver la même gêne, quand ce bruit ne produit que les fluctuations chromatiques.

TRANSMISSION A BRILLANCE CONSTANTE.

Le canal monochrome d'un récepteur de couleur n'est pas plus sujet aux bruits et aux interférences qu'un récepteur monochrome de même définition.

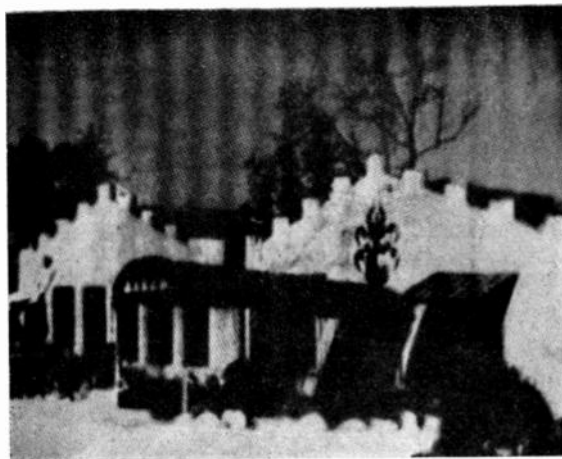


FIG. 14. — Effet visible d'interférence (C.W. interférence) dans le cas d'une transmission à amplitude constante.

La visibilité des barres est environ réduite de 8 db dans le cas d'une transmission à brillance constante.

Mais ces perturbations peuvent affecter le canal de couleur d'une façon appréciable, si des précautions spéciales ne sont pas prises¹¹. Les produits de démodulation, dus aux interférences, bruits, et composantes du signal de brillance, voisines de la sous-porteuse couleur, donnent naissance à des fréquences plus basses, qui deviennent plus visibles.

La figure 14 montre l'effet d'onde continue (C.W. interference) dû à l'interférence d'une fréquence inférieure de 500 kc/s à la sous-porteuse couleur, dans un récepteur correspondant à une version antérieure du signal de couleur. Dans cette version, une partie appréciable de la brillance était apportée par le canal de couleur. Ce récepteur utilise trois démodulateurs de gain égal, alimentés par trois tensions égales, décalées de 120°. Le battement à 500 kc/s est clairement visible sur la figure 14. Ce battement fait varier la brillance et la couleur.

Puisque les sorties des trois démodulateurs sont égales en amplitude, et déphasées de 120°, l'intensité lumineuse totale émise par le tube devrait s'annuler. Ce serait le cas, si l'œil était également sensible aux trois couleurs. Mais nous avons vu qu'il était en réalité plus sensible au rouge qu'au bleu, et encore plus au vert qu'au rouge. Il en résulte que des variations d'égale intensité, dans le vert, le rouge et le bleu, provoquent des sensations inégales, pour l'œil qui les combine, et par conséquent il n'y a pas annulation.

Cela est éclairé par la figure 15, qui montre comment les trois canons du tube, sont soumis au signal perturbateur, en fonction du temps. La partie inférieure de la figure montre comment cette fluctuation en fonction du temps est traduite en une fluctuation dans l'espace par le faisceau électronique.

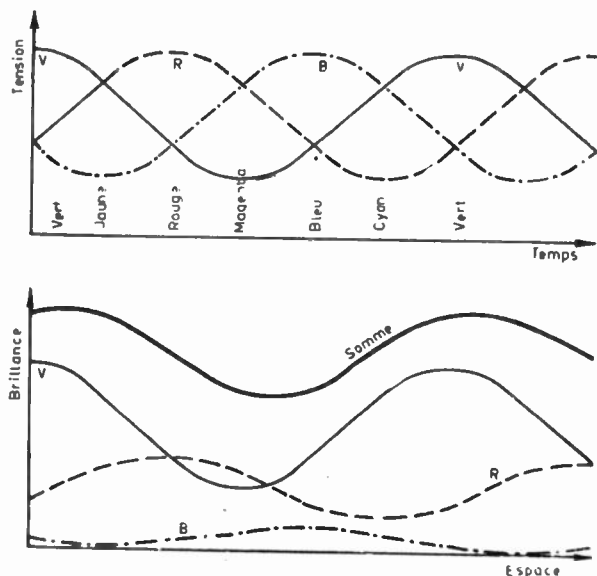


FIG. 15. — Courbes de tension et de brillance dans le cas de l'interférence montrée fig. 14. Remarquons que les variations de tension s'éliminent contrairement aux variations de brillance.

On montre aussi les sensations inégales de brillance, produites par le vert, le rouge et le bleu, quand ils sont excités de la même façon. Enfin, la figure 15 montre également la fluctuation totale de brillance résultante.

La visibilité de l'interférence peut être très réduite, en proportionnant le signal, de telle façon que la combinaison des sorties des démodulateurs n'engendre pas de fluctuations de brillance. Lorsqu'il en est ainsi, le batttement affecte seulement la chromaticité, et est par suite beaucoup moins visible. C'est pour cette raison que le signal de couleur complet est transmis de manière que le signal monochrome porte toute l'information brillance, tandis que le sous-porteur de couleur porte seulement les variations de couleur.

AVANTAGES DES SPÉCIFICATIONS N T S C.

Un système utilisant ces spécifications est capable de produire des images de télévision en couleur, avec une définition aussi grande que celle qui peut être obtenue, à ce jour, avec les systèmes de télévision monochrome. Il peut, de plus, transmettre toute l'information couleur, que l'œil est capable d'apprécier, cette transmission s'effectuant à l'intérieur de la bande de fréquences allouée actuellement à la télévision en noir et blanc.

Les canaux monochromes existants, dont la bande de fréquence est de 6 Mc/s, sont par conséquent propres à transmettre la couleur. Les émetteurs peuvent être très facilement convertis, pour la transmission de la couleur.

Cette transmission peut être reçue en image monochrome par les récepteurs noir et blanc existants, sans diminution de la qualité. Inversement les transmissions monochromes peuvent être reçues par les récepteurs de couleur. Ces spécifications désignent un système, qui tient compte des propriétés de l'œil du spectateur, qui en constitue actuellement l'équipement terminal.

BIBLIOGRAPHIE

1. A. V. BEDFORD, Mixed Highs in Color Television, *Proc. IRE*, sept 1950.
2. A. V. LOUGHREN and C. J. HIRSCH, An Analysis of Color Television Systems, *Electronics*, Feb. 1951.
3. Report of Panel 11, NTSC Pub. 11-114.
4. Report of Panel Actions, Panel 13, NTSC Pub. 13-162, oct. 23, 1951. Voir aussi Specifications for Color TV Field Tests, *Electronics*, Jan. 1952.
5. Pierre MERTZ, Television. — The Scanning Process. *Proc. IRE*, oct. 1941.
6. U. S. Patent 1, 769, 920, July 8, 1930, Frank Gray.
7. H. B. LAW, A Three-Gun Shadow Mask Color Kinescope, *Proc. IRE*, oct. 1951. Voir aussi *RCA Review*, Sept. 1951, Part II, pp. 443-648.
8. Report of Panel 14, NTSC Pub. 14-163.
9. R. D. KELL and C. C. FREDENDALL, Standardization of Transient Response, *RCA Review*, March 1949.
10. B. D. LOUGHLIN, Recent Improvement in Band-Shared Simultaneous Color Television, *Proc. IRE*, Oct. 1951.
11. Improvements in Dot-Sequential Color Television, *Electronics*, p. 154, Aug. 1950.
12. NYQUIST and PFLEGER, Effect of the Quadrature Component in Single Sideband Transmission, *Bell System Technical Journal*, Jan. 1940, p. 63.
13. W. F. BAILEY and C. J. HIRSCH, Quadrature Crosstalk in NTSC Color Television, *Proc. IRE*, Jan. 1954.
14. C. J. HIRSCH, W. F. BAILEY, B. D. LOUGHLIN, Principles of NTSC Compatible Color Television, *Electronics*, Feb. 1952.
15. D. RICHMAN, Color-Carrier Reference Phase Synchronization Accuracy in NTSC Color Television, *Proc. IRE*, Jan. 1954.

••

1. FCC Public Notes 53-1663 Mimeo 98948 (FCC Technical Standards Amended to Incorporate Color).
2. C. J. HIRSCH. — A Review of Recent Work in Color Television. Cet article se trouve dans un livre intitulé « *Advances in Electronics* » Vol. 5, 1953, Academic Press Inc., N. Y., N. Y. (Cet article comprend une importante bibliographie).
3. *Proc. IRE*. — Jan. 1954. Issue devoted entirely to color television.
4. *Proc. IRE*. — Oct. 1951. Issue devoted entirely to color television.
5. Report of Panel 12, Color System Analysis, NTSC P 12-369.
6. Report of Panel 13, Color Video Standards, NTSC P 13-376.
7. Some Supplementary References Cited in the National Television System Instruments Report, NTSC-P 13-422.
8. Report of Panel 14, Color Synchronizing Standard, NTSC-P 14-360.
9. Report of Panel 19, Definitions, NTSC-P 19-395.

ERRATA

Page 25, fig. 2, *lire* : $f_s = (2m + 1) \frac{f_L}{2}$

Page 29, fig. 8, 2^e ligne, *lire* : $= 33^\circ + \lg^{-1} \frac{E'_i}{E'_q}$

4^e ligne, *lire* : $E'_i =$

Page 31, 2^e col., 30^e ligne, *lire* : E'_v et non pas E^1_v

TÉLÉVISION EN COULEURS AUX U.S.A.

PAR

C. G. MAYER

Représentant technique de RCA en Europe

L'homme a été tenté de transmettre des images en couleurs dès le début de la télévision et depuis bien longtemps, on avait une idée grossière de la manière dont il fallait s'y prendre mais les moyens d'action étaient trop faibles et trop dispendieux pour pouvoir s'appliquer à la radiodiffusion telle que nous la concevons dans notre société actuelle. L'humanité progresse, du moins au point de vue technique, en apprenant à comprendre et à utiliser les secrets de la nature. De même que le raisonnement déductif suit l'expérience et que le plein développement scientifique est un résultat des deux, de même les images en couleurs suivent les images en noir et blanc dans le développement normal de la télévision.

Le problème fondamental consistait à trouver un procédé grâce auquel on pouvait ajouter la couleur dans les réseaux d'émission monochrome existants — émission qui atteint aujourd'hui presque 30 millions de récepteurs dans les foyers américains. Dans ces circonstances, il était évidemment très désirable que le passage à la couleur se fit progressivement et sans interrompre les bonnes relations entre organismes de radiodiffusion et public. Cependant le défi lancé à notre intelligence était d'imaginer une spécification d'émission permettant de recevoir les programmes en couleurs sur un récepteur pour la couleur et les émissions en monochrome sur les récepteurs actuels en noir et blanc, sans supplément de prix pour les usagers et sans changement ni supplément dans les installations. Le récepteur en couleurs devra également capter en noir et blanc les émissions standard en noir et blanc. On appelle un tel système de télévision, un système « entièrement compatible » et c'est un tel système qu'on a finalement adopté. En agissant ainsi, cependant, il restait à résoudre un autre problème majeur.

La F. C. C. qui a la responsabilité, dans l'intérêt du public, de réglementer le domaine des ondes hertziennes aux Etats-Unis, a établi les normes de télévision en noir et blanc pour un canal de 6 Mc/s dont l'aménagement est indiqué figure 1. Ces normes ont permis d'établir un service satisfaisant, et, par suite, la télévision s'est rapidement développée.

Il devint bientôt évident qu'afin de satisfaire les demandes de création de nouvelles stations, beau-

coup de canaux supplémentaires seraient nécessaires. Les bandes de fréquences ont été assignées aux différents services par un accord international, et à Atlantic City en 1947, les nations du monde entier se mirent d'accord pour que certaines bandes, dites les bandes I, III, IV et V soient réservées à la télévision. Aux Etats-Unis, les 12 canaux disponibles, dans les bandes I et III sont déjà complètement utilisés et la seule possibilité de créer les canaux

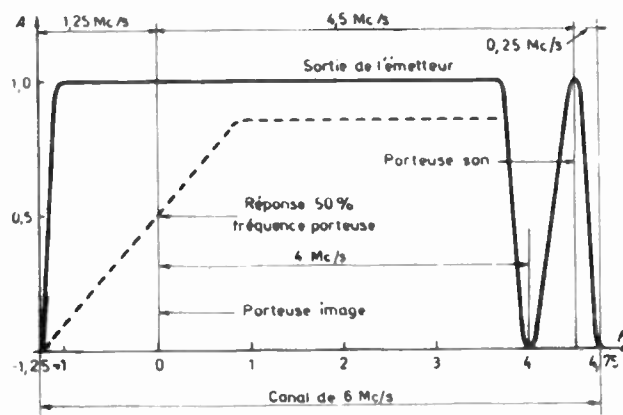


Fig. 1. — Normes de la télévision en couleurs aux Etats-Unis.

(525 lignes, 30 images, balayage entrelacé de 60 frames, polarisation horizontale, transmission négative; modulation d'amplitude pour la video, modulation de fréquence pour le son.

désirés étaient d'extrapoler le plan de fréquences jusque dans les bandes décimétriques peu explorées au-dessus de 400 Mc/s, c'est-à-dire dans les bandes IV et V. Ceci fut fait et donna en tout 70 nouveaux canaux, avec les mêmes normes de télévision que pour les bandes I et III, de telle sorte qu'il y a maintenant 258 stations fonctionnant dans ces bandes et 130 stations supplémentaires fonctionnant dans les bandes IV et V. 220 nouvelles autorisations dans ces 2 bandes ont été accordées et 250 autres sont en instance. Beaucoup de ces stations peuvent être interconnectées pour fonctionner en réseau et la majorité des spectateurs ont le choix entre plusieurs programmes.

L'importance de la protection du spectre radio ne peut être souligné trop fortement. Déjà il y avait une congestion importante, et devant l'accroissement incessant des services-radio dont chacun

revendique ses droits, on n'entrevoit guère de place vacante dans le spectre radio. Les caractéristiques de transmission des ondes hertziennes diffèrent considérablement suivant la fréquence utilisée et nous ne pouvons certainement pas nous permettre d'aggraver la congestion par une occupation prodigue du spectre radio.

Dans ces conditions, la R. C. C. considéra comme impensable d'attribuer des bandes supplémentaires pour la couleur. Ainsi pour avoir la télévision en couleurs, il fallait trouver un procédé pour comprimer la plus grande partie de l'information nécessaire et la faire rentrer dans la même largeur de canal que celle de la télévision en noir en blanc.

Il serait souhaitable d'examiner la solution du problème en 2 parties. D'abord une révision brève des caractéristiques de l'œil en ce qui concerne la couleur, et ensuite les techniques électriques permettant d'ajouter l'information supplémentaire nécessaire pour donner des images de télévision en couleurs satisfaisantes et agréables.

L'imprimerie en couleurs, la peinture et la photographie en couleurs dépendent du principe fondamental suivant lequel le mélange de trois couleurs convenables (habituellement le vert, le rouge et le bleu) permet d'obtenir presque toutes les sensations colorées possibles.

Ces techniques bien connues utilisent des pigments ou colorants tandis que dans la télévision on mélange des lumières, celles-ci doivent être converties par une caméra de télévision en signaux électriques pour la transmission H. F. et être reconverties en lumière par le tube reproducteur d'image ou « kinescope » comme nous l'appelons, du récepteur afin de les rendre à nouveau visibles. Comme l'œil est plus sensible au vert, moins sensible au rouge et encore moins au bleu, on obtient une brillance correcte en mélangeant ensemble les signaux de couleurs dans des proportions déterminées avec soin.

Vision de la couleur.

La vision est un phénomène complexe impliquant les yeux et le cerveau dans un processus qui est loin d'être parfaitement connu. Nous distinguons dans la sensation de couleur trois attributs principaux : la brillance, la teinte et la saturation. La brillance est une mesure de l'intensité lumineuse ou « luminance » et est la seule qualité commune aux objets à la fois colorés et non colorés. La teinte est caractérisée par la longueur d'onde dominante de la couleur qui détermine si elle est rouge, verte ou bleue. La saturation donne la vivacité des couleurs d'une même teinte et peut être reliée à la « pureté » ou absence de dilution avec le blanc.

La théorie de la colorimétrie a été construite d'après les principes du mélange des couleurs, et on a adopté le diagramme de chromaticité de la figure 2, par un accord international, ce qui permet de représenter les couleurs comme des points de ce diagramme en fonction du mélange des primaires standards. La position d'un point détermine complètement la teinte et la saturation mais ne donne aucune indication sur la brillance. Le spectre visi-

ble est représenté par une courbe fermée en forme de fer à cheval, toutes les couleurs possibles à pleine saturation se trouvant sur le périmètre, c'est-à-dire qu'en faisant le tour de la périphérie la teinte varie mais la saturation reste constante. Si nous nous déplaçons sur une ligne quasi-radiale à partir d'une teinte située sur la périphérie vers le blanc, la teinte reste constante alors que la saturation diminue.

Par un choix convenable des primaires on peut donner une représentation trichrome tout à fait satisfaisante pour les buts envisagés. Cependant, on a trouvé que, alors que le mélange trichrome est nécessaire pour les surfaces importantes, on peut obtenir une reproduction convenable par le mélange

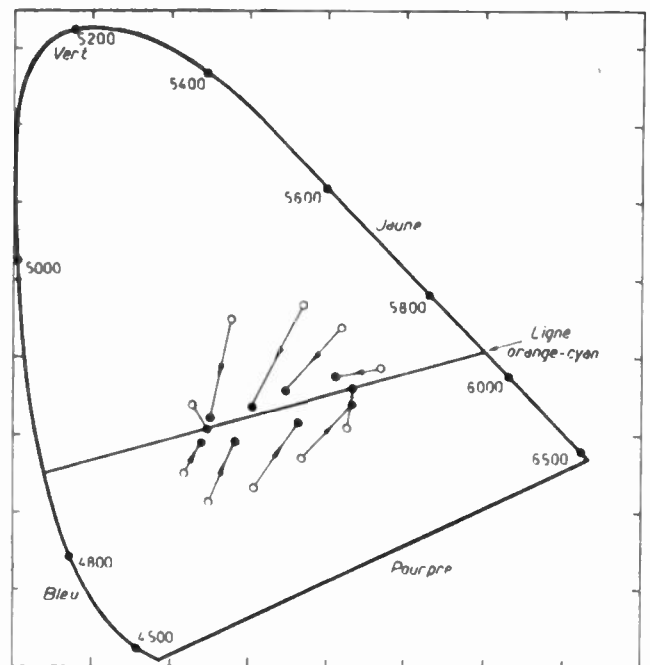


FIG. 2. — Triangle de couleur C.I.E. et équilibre des couleurs des échantillons. (Les nombres indiquent approximativement les longueurs d'onde en Angstrom).

de deux couleurs seulement pour une surface suffisamment petite, ce fait a été mis en relief par un grand nombre de données sur la vision découvertes seulement ces dernières années et dont l'importance n'est peut-être pas encore largement estimée.

En 1945, WILLMER et WRIGHT ont dit avoir trouvé que toutes les couleurs dans une bande suffisamment petite pouvait être reproduite en mélangeant seulement deux et non pas trois sources lumineuses primaires. Deux physiciens canadiens MIDDLETON et HOLMES effectuèrent l'expérience suivante : ils découpèrent des petits morceaux de feuilles colorées et demandèrent à des observateurs de comparer ces morceaux à différentes feuilles de grand format. On trouva que le meilleur équilibre était obtenu avec des morceaux de couleur différente de celle dont les pièces étaient originaires. On donne les résultats figure 2, les cercles représentent les chromaticités des morceaux d'origine et les points noirs les chromaticités des morceaux qui donnent le meilleur équilibre de couleur. Il est à remarquer

que pour les petits morceaux, le diagramme de chromaticité tend à se réduire à une seule ligne, et aussi que les deux primaires utiles pour cette ligne sont un rouge-orangé et un bleu-verdâtre (appelé « cyan »). Un autre chercheur anglais, le Dr HARTRIDGE a aussi fait un grand nombre de recherches variées sur l'acuité à la couleur pour de petits objets, lesquelles corroborent ces résultats. Une nouvelle confirmation est trouvée dans l'étude de la photographie à 2 couleurs. CLYNE note que, bien que d'autres paires de primaires soient capables de donner le noir et le gris, seule la combinaison orange-cyan peut représenter convenablement les teintes du ciel, de la chair et du feuillage si familières dans la nature. Ajoutons à ces données, les mesures d'acuité visuelle à la couleur et au contraste faites dans les laboratoires R. C. A., et il apparaîtra expérimentalement évident que pour les surfaces d'une image comportant des détails fins la représentation bichrome est largement suffisante. De plus, ces expériences mettent en évidence que l'œil semble avoir sa plus grande acuité lorsque les couleurs à différencier se trouvent au voisinage de l'axe orange-cyan. Il est intéressant de noter en passant que les gens avec une vision normale voient les petits objets colorés à peu près de la même manière que les daltoniens qui confondent le vert et le pourpre, voient tous les objets.

Maintenant qu'arrive-t-il lorsque les détails deviennent très fins ? Les essais entrepris depuis plusieurs années sur l'acuité visuelle à la couleur par A. V. BEDFORD établissent que l'œil ne voit pas du tout la couleur dans les détails très fins d'une image ; seule la perception de la brillance reste. En fait, l'acuité de l'œil pour le rouge est légèrement inférieure à celle pour le vert, et elle est encore moindre pour le bleu. L'œil est aussi plus sensible aux changements de brillance qu'il ne l'est aux variations de teinte et de saturation et à mesure que les détails de l'image deviennent plus fins, l'axe orange-cyan s'approche du point qui représente le blanc sur le diagramme de chromaticité. Il est aussi nécessaire de se rappeler que la théorie de l'information nous dit que la transmission des détails fins exige une plus grande largeur de bande qu'il ne faut pour transmettre un détail grossier (ou large surface). En général, un canal de transmission peut oublier le dernier signal transmis en un demi-cycle de la plus haute fréquence passante. Ainsi un canal de 4 Mc/s peut transmettre par seconde un maximum de 8 millions de signaux d'amplitudes entièrement indépendantes, bien qu'en pratique la vitesse à laquelle l'information peut être transmise sur un canal donné, soit déterminée par le rapport signal/bruit choisi.

A une certaine époque, on pensait que la transmission d'une image par 3 couleurs exigeait de transmettre 3 fois la somme d'information requise pour l'image monochrome donc nécessiterait une bande de fréquence trois fois plus large, mais il a été démontré depuis quelque temps que cela était inutile. En fait, on a démontré expérimentalement que la largeur de bande pour l'information de couleur seule peut être limitée à approximativement

1,5 Mc/s sans que l'œil perçoive aucune perte de couleur dans l'image transmise.

On sait par l'expérience de la télévision en noir et blanc que le pouvoir séparateur de l'œil pour les détails de brillance à la distance normale de vision est bien satisfait avec des images obtenues par le processus normal de balayage à 525 lignes. Ceci implique une largeur de bande vidéo de 4 Mc/s environ qui peut donc être regardée comme convenable pour l'information de luminance requise pour obtenir tous les détails de l'image.

En résumé, si on combine ces caractéristiques de vision avec la connaissance des transmissions électriques le signal de télévision en couleurs doit, en bref, répondre aux conditions suivantes :

— Pour des surfaces d'image relativement grandes, information par 3 couleurs (vert, rouge et bleu) exigeant une bande étroite de fréquence (0,5 Mc/s environ).

— Pour des surfaces relativement petites, information avec 2 couleurs (orange et bleu-vert) exigeant une bande de fréquence modérément large (1,5 Mc/s environ).

— Information de luminance (brillance) pour tout ce qui concerne les détails exigeant la même largeur de bande que celle des images en noir et blanc (environ 4 Mc/s).

On remarquera également que puisque on ne gagne rien en transmettant une information d'image que l'œil ne peut pas apprécier, tout système qui cherche à restituer tous les détails de tous les objets, quelles que soient leurs dimensions et leur couleur, doit être repoussé comme causant un gaspillage de spectre.

Techniques électriques.

Nous examinerons ici les techniques électriques nécessaires pour transmettre et recevoir trois signaux séparés satisfaisant aux exigences de la compatibilité et de l'encombrement prescrit.

L'envoi simultané de 2 ou plusieurs messages indépendants sur un système de transmission unique avec la possibilité de séparer ces messages à volonté, est un vieux problème de télécommunications. Cette technique est connue sous le nom de multiplexage et son emploi est très répandu. Il y a plusieurs méthodes qui toutes exigent une largeur de bande plus large que celle nécessaire pour un simple message. La méthode favorite pour multiplexer les 2 composantes de chrominance est celle de la division dans le temps, ce qui nécessite de partager l'utilisation du canal à intervalle de temps au moyen d'une commutation synchrone de l'émission et de la réception aux extrémités du circuit.

Dans l'application de cette méthode à la télévision en couleurs, on obtient un procédé de commutation intéressant en produisant deux porteuses de même fréquence et déphasées de 90°, chacune d'elles pouvant être modulée en amplitude indépendamment et simultanément pour produire une porteuse résultante unique à la même fréquence.

Le signal en phase ou signal I est destiné à porter l'information de saturation de couleur et le signal en quadrature ou signal Q donne l'information de la teinte. La méthode particulière de modulation employée a pour résultat une transmission à double bande avec suppression de la porteuse. Celle-ci constitue la sous-porteuse de chrominance qui par des moyens appropriés est directement ajoutée à la porteuse principale image.

Puisque la sous-porteuse est supprimée à l'émetteur dans le but de réduire la diaphonie, elle doit être re-introduite dans le récepteur afin de permettre de retrouver les signaux I et Q par détection synchrone. Cela entraîne l'existence dans le récepteur d'un oscillateur local de sous-porteuse qui travaille sur la même fréquence et la même phase que l'oscillateur correspondant à l'émission. Le résultat est obtenu en transmettant un échantillon de 9 cycles environ de la sous-porteuse de l'émetteur qui sert alors de signal de référence dans le récepteur pour verrouiller la sous-porteuse du récepteur avec celle de l'émetteur. On évite le brouillage avec les autres éléments constituant des signaux de télévision en introduisant l'échantillon dit « salve de couleur » dans la période de suppression qui suit l'impulsion normale de synchronisation horizontale.

Il est maintenant nécessaire de combiner les signaux de chrominance et de luminance pour obtenir le signal de couleur complet. La combinaison de ces signaux, par la méthode habituelle de juxtaposition, exigerait une bande passante minimum de 6 Mc/s ce qui est encore trop, vu la bande vidéo à notre disposition. On a donc recours de nouveau au multiplexage pour compresser l'information dans

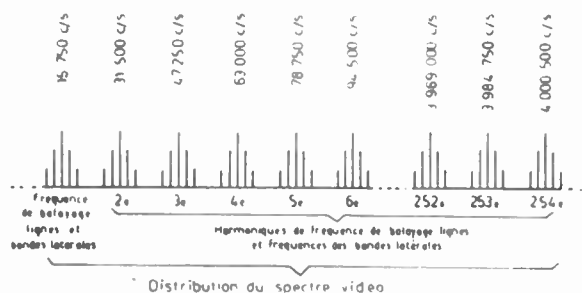


FIG. 3. — Distribution des fréquences vidéo.

4 Mc/s mais en agissant de la sorte, on se trouve en face de nouveaux problèmes de brouillage ou de diaphonie entre les signaux, ce qui rend plus difficile leur séparation correcte. Cette diaphonie a été un des problèmes majeurs de transmission dans l'évolution des systèmes de télévision en couleurs compatibles. Elle peut donner lieu à des défauts se traduisant par des moirures mouvantes sur l'image reçue.

Des perfectionnements récents apportés au système R. C. A. sont, cependant, arrivés à éliminer la plus grande partie de ce défaut.

Il y a de nombreuses années MERTZ et GREY étudièrent le problème de la transmission de 2 images de télévision à l'intérieur d'un canal donné. Ils découvrirent que la distribution d'énergie dans les

composantes du signal due au processus d'analyse de la télévision a pour effet (dans beaucoup de cas) de grouper la plus grande partie de l'énergie autour des fréquences multiples de la fréquence de balayage ligne, comme le montre la figure 3. Par contre aux fréquences voisines des multiples impairs de la moitié de la fréquence-ligne, le spectre est à peu près innoccupé. Une porteuse additionnelle, transportant une information d'image indépendante, peut remplir ces espaces de fréquence, pour peu que cette nouvelle porteuse soit décalée de la porteuse-image principale d'un multiple impair de la demi-fréquence de ligne. Ce processus est communément appelé « entrelacement de fréquence » et a pour résultat que la porteuse et ses bandes latérales se trouvent intercalées entre les harmoniques des fréquences-ligne et image. En pratique, on évite un émetteur supplémentaire en utilisant une sous-porteuse modulée par la nouvelle information à transmettre.

Une analogie commode pour illustrer la manière dont les lacunes du spectre sont remplies est de considérer le canal vidéo comme un peigne dont les dents représentent le signal de luminance. Le signal chrominance peut être alors considéré comme les dents d'un peigne semblable mais plus court, placé de telle sorte que les 2 peignes soient intercalés ; la position exacte du 2^e peigne étant déterminée par le choix de la fréquence de la sous-porteuse.

Un autre moyen de représentation est de considérer les phénomènes du point de vue du balayage entrelacé horizontal. C'est l'une des principales méthodes pour économiser de la bande et qui met en jeu l'entrelacement de 2 ensembles d'image le long de chaque ligne de balayage. Si, dans chaque trame, on ne transmet qu'un point sur deux de chaque ligne et qu'une ligne sur deux, il faut 4 trames de cette sorte pour obtenir une image complète. L'entrelacement par points réduit la dimension de l'élément d'image et réduit également le papillotement. La séquence d'entrelacement peut être aussi combinée de telle sorte que les balayages successifs ajoutent leurs effets ou bien on peut adopter un entrelacement de type équilibré. En principe, la compensation se produit lorsqu'on a pris la précaution d'inverser la polarité du signal entre deux balayages successifs de chaque élément de l'image. Les trames d'interférence correspondant à la diaphonie entre les signaux de chrominance et de luminance tendent alors à se compenser grâce à la persistance de la vision. L'entrelacement par points est combiné de telle sorte que les signaux d'élément d'image se répètent à une fréquence qui est un multiple impair de la moitié de la fréquence de ligne, ce qui se produit lorsque la période d'image comporte un nombre entier de périodes plus une demi-période de la sous-porteuse. Ce sont exactement les conditions pour obtenir l'entrelacement de fréquence, ainsi il est évident que l'entrelacement de fréquence donne toujours lieu à l'entrelacement de points et vice-versa, ce qui signifie que ces 2 notions sont équivalentes en pratique. Les principales considérations du système ont été traitées dans les grandes lignes et bien qu'un grand nombre de détails n'ait pas été examiné, il en a été dit assez

pour que nous puissions donner les spécifications générales du mode de transmission qui permet l'exploitation de la télévision en couleurs compatible, en utilisant le mieux possible le canal de 6 Mc/s de largeur de bande.

— (1) Signal total video de 4 Mc/s, représentant les variations de luminance du sujet, transmis par modulation d'amplitude de la porteuse-image :

— (2) Addition d'une sous-porteuse à l'intérieur de la bande video, à une fréquence choisie en fonction de l'« entrelacement de fréquence » :

— (3) Modulation simultanée avec suppression de cette sous-porteuse par les 2 composantes de chrominance.

a) Signal I (exigeant une largeur de bande de 1,5 Mc/s) représentant la saturation qui module la sous-porteuse en amplitude.

b) Signal Q (limité à environ 0,5 Mc/s de largeur de bande) représentant la teinte qui module la sous-porteuse en phase.

— (4) Addition du signal de référence de phase consistant en une « salve » — échantillonnant la fréquence de la sous-porteuse, incorporée pour sa transmission dans le signal classique de synchronisation.

— (5) Canal son en modulation de fréquence comme pour la télévision en noir et blanc.

De nombreux détails devraient être considérés avant d'établir une spécification mais faute de place nous n'en traiterons que quelques-uns ici.

Le choix de la fréquence de la sous-porteuse est un compromis entre plusieurs facteurs. Il est souhaitable d'avoir une fréquence aussi élevée que possible de telle sorte que les moirures de points soient fines et que les effets indésirables dus à une compensation imparfaite soient réduits. D'un autre côté, une fréquence relativement plus basse augmente la bande passante disponible pour transmettre le signal de chrominance. On a trouvé qu'il est suffisant dans ce but de prévoir une bande s'étendant à une fréquence à environ 0,5 Mc/s au-dessus de la fréquence de la sous-porteuse. Ainsi disposant d'une bande video de 4 Mc/s, la fréquence de la sous-porteuse doit se trouver vers 3,5 Mc/s. La fréquence exacte est déterminée par la condition d'« entrelacement de fréquence » et la considération de brouillage possible avec la porteuse-son dans un récepteur.

Des essais ont montré que tout battement entre porteuse-son et sous-porteuse de chrominance (dû à un affaiblissement insuffisant de la porteuse-son dans quelques récepteurs existants en noir et blanc) est beaucoup moins désagréable sur l'image si la fréquence de battement est un multiple impair de la demi-fréquence de ligne. Ceci peut être aussi obtenu en séparant les porteuses son et vision d'une quantité égale à un multiple de la fréquence de ligne, c'est-à-dire un multiple pair de la fréquence de demi-ligne. Il n'est guère possible de modifier les fréquences actuellement fixées pour le son à cause du grand nombre de récepteurs existants et

par suite pour obtenir le décalage désiré, les fréquences de ligne et de trame ont été très légèrement modifiées d'environ 0,1 %, ce qui reste dans les limites de tolérance du standard actuel noir et blanc. Un fonctionnement non synchrone de l'émetteur avec le secteur est devenu indispensable, ce qui d'ailleurs s'était déjà avéré nécessaire pour la télévision noir et blanc à cause des transmissions à grande distance et des équipements de reportages.

Afin de satisfaire à ces conditions, la fréquence de la sous-porteuse de chrominance se situe à 3,579545 Mc/s (avec une fréquence de balayage ligne de 15734,264 c/s et une fréquence de trame de 59,94 c/s). La figure 4 donne le spectre du signal de couleurs transmis.

Le choix de cette fréquence pour la sous-porteuse entraîne l'inégalité des bandes latérales puisqu'il réduit la bande supérieure alors que la

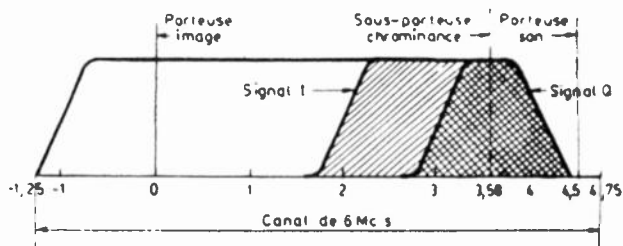


FIG. 4. — Spectre du signal de couleur transmis.

bande latérale inférieure peut s'étendre à environ 1,5 Mc/s au-dessous de la fréquence de la sous-porteuse. Comme indiqué plus haut, l'effet d'entrelacement est de faire une compensation de la diaphonie, mais un tel déséquilibre dans les bandes latérales a été une cause de diaphonie supplémentaire et il a fallu imaginer les moyens pour l'éviter.

La largeur de bande exigée pour les larges surfaces est inférieure à 1 Mc/s et par suite n'entraîne aucun problème de diaphonie. Par contre, pour les détails fins, un déséquilibre apparaît dans les bandes latérales et la diaphonie a tendance à se développer entre les composantes du signal de chrominance. Cela est évité en limitant la largeur de bande de Q à 0,5 Mc/s, de telle sorte que pour les détails fins où il y aurait diaphonie entre les 2 signaux, on transmet seulement le signal I, ce qui évidemment le libère de toute diaphonie.

Il a été déjà souligné qu'on se montrerait prodigue du spectre si on transmettait plus d'information que l'œil ne peut en voir et il est également évident que le même élément d'information ne doit pas être transmis plus d'une fois.

C'est ce qui se produirait si chaque composante de chrominance d'une scène à transmettre contenait une information relative à la brillance en même temps qu'à la couleur. On doit se souvenir que pour satisfaire aux exigences de la compatibilité, la brillance de la scène est transmise par modulation d'amplitude de la porteuse de la manière classique. Ainsi pour éviter la redondance, il est nécessaire de soustraire le signal de brillance de chaque signal de couleur, vert, bleu et rouge.

Mais le fait de transmettre 4 éléments d'information quand il n'en faut que trois entraînerait encore une redondance. Le signal représentant « vert-moins-brillance » peut être obtenu dans le récepteur en soustrayant la somme des signaux rouge et bleu du signal de brillance. On a donc fait en sorte de ne transmettre que les signaux représentant le « rouge-moins-brillance » et le « bleu-moins-brillance » en même temps que le signal de brillance. A titre de curiosité le terme de « couleur-moins-brillance » ne correspond à rien de physique, puisque l'œil a une réponse seulement à la couleur lorsqu'elle est accompagnée de brillance. En pratique les signaux « couleurs-moins-brillance » sont obtenus en soustrayant la valeur électrique du signal de brillance de la valeur électrique des signaux de couleur en utilisant un organe réalisant des combinaisons linéaires. L'information de brillance est rétablie au récepteur de couleur et on reconstitue les signaux vert, rouge et bleu nécessaires pour faire fonctionner le tube cathodique trichrome.

En utilisant ces deux signaux de différence de couleur pour moduler la sous-porteuse, on obtient une pureté plus grande en mélangeant chacun d'eux avec l'autre de telle sorte que le « bleu-moins-brillance » contient un peu de rouge et que le « rouge-moins-brillance » contient un peu de bleu. Ceci n'affecte pas les larges plages de l'image, mais pour les détails moyens on produit une couleur polluée ou mélangée allant de l'orange au bleu-vert qui donne une distorsion moins perceptible que si on transmettait une seule couleur pure. La réduction de distorsion sur les contours ainsi obtenue est due au fait que l'œil est moins sensible à la couleur sur les contours quand on évite un contraste brutal entre deux surfaces adjacentes.

Tout ceci concerne la télévision et il est intéressant de faire un parallèle avec l'art de la peinture. Dans l'art classique on crée d'abord la forme et ensuite on ajoute la couleur. Après avoir étalé la couleur de fonds, l'artiste peint habituellement la forme en monochrome généralement en noir ou brun et blanc apportant un soin particulier aux contours du dessin et il ajoute ensuite les touches de couleur. Ceci peut être observé dans les peintures du 16^e siècle environ et au-delà qu'on peut comparer avec les œuvres de quelques-uns de nos artistes impressionnistes modernes qui préfèrent laisser la forme se deviner d'elle-même et peindre directement en couleur dès le départ. Dans le système de télévision en couleurs décrit, le signal monochrome fournit le détail ou la forme et le coloriage est ajouté séparément ce qui peut être considéré comme une technique de l'art classique.

Compatibilité.

On doit considérer deux conditions : la 1^{re} concerne la réception des émissions en couleurs sur un récepteur noir et blanc usuel. Il est facile de voir que quand il n'y a pas de couleur dans l'émission, le récepteur répond entièrement au signal de luminance et le montre comme habituel avec

une qualité égale (ou supérieure) à celle fournie par le standard actuel de transmission en noir et blanc. Quand la scène originale comporte des couleurs l'information de chrominance est transmise, mais une compensation se produit d'elle-même par l'entrelacement soustractif du système de couleurs et, sauf quelques défauts résiduels dus à des effets non linéaires dans le tube cathodique, le récepteur n'est pas affecté par le signal de couleur. En choisissant convenablement la fréquence de la sous-porteuse on obtient une trame d'une texture tellement fine qu'à la distance normale de vision pour laquelle la structure du balayage des lignes habituelles disparaît, les moirures disparaissent également. Le spectateur regardant l'image est alors en principe incapable de trouver une différence à la réception entre ces images et celles transmises par l'émetteur normal en noir et blanc.

La 2^e condition est la réception des émissions en noir et blanc sur un récepteur de couleur. Ceci ne présente aucun problème puisque seul le signal de luminance demeure. Il est cependant essentiel que le tube cathodique trichrome soit construit avec soin et soit capable d'être réglé pour donner une trame blanche exempte d'une contamination de couleur due à un mauvais recouvrement ou une convergence incorrecte des faisceaux électroniques. Quand il est utilisé correctement, un tel tube donne des images en noir et blanc d'une qualité équivalente à celle des récepteurs en noir et blanc usuels.

Spécification du signal complet.

La spécification du signal complet pour les émissions de télévision en couleurs a été proposée à la F.C.C. par la RCA en Juin 1953 et peu de temps après par le Comité National de Télévision, représentant les industries radioélectriques. La spécification proposée a été vérifiée pratiquement par un nombre considérable d'expériences et le développement de la couleur exige un travail d'équipe de la part de nombreux ingénieurs électroniques et des physiciens hautement spécialisés des Etats-Unis. Cela représente plusieurs milliers d'heures de travail de la part des techniciens et des frais d'études et de recherches qui se montent pour RCA seule, à plus de dix millions de livres.

Témoin des démonstrations et après avoir considéré tous les facteurs mis en jeu, la F.C.C. décida d'adopter les recommandations de l'industrie et en Décembre 1953 établit les normes techniques de la télévision en couleurs compatibles, actuellement en vigueur.

Aucun système n'est parfait au commencement. Nous reconnaissons que le prix initial est élevé et qu'il y a certains défauts, mais aucun n'est cependant suffisant pour ôter le plaisir de voir des images en couleurs. Toutes les imperfections doivent être considérées comme provenant du dispositif et non pas inhérentes aux spécifications du signal. Cela a été reconnu dans les remarques suivantes qui apparaissent dans la décision finale de la F.C.C. sur la télévision en couleurs;

« Les spécifications du signal de télévision en couleurs proposées donnent une image raisonnablement satisfaisante avec une bonne qualité d'ensemble. La qualité de l'image n'est pas endommagée d'une manière appréciable par des effets tels que : mauvais recouvrement, ondulation des lignes, structure trop visible des points, sautellement. L'image est suffisamment brillante pour fournir une gamme de contraste satisfaisante dans les conditions favorables de lumière ambiante et peut être vue chez soi sans papillotement prohibitif. Les images de télévision peuvent être transmises d'une manière satisfaisante sur les relais existants de ville à ville et des améliorations entre villes peuvent être raisonnablement prévues.

Le succès de la couleur dépend du volume des ventes de récepteurs et tous les efforts doivent être faits pour abaisser les prix au niveau des possibilités de l'acheteur moyen. Tout effort doit être entrepris également pour concevoir l'équipement

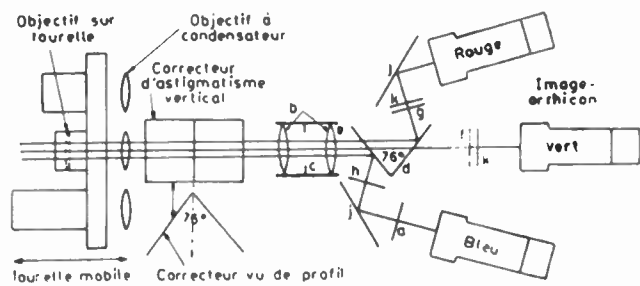


FIG. 5. — Schéma de la caméra de télévision RCA à 3 tubes.

d'une manière telle que l'on réduise le plus possible l'accroissement de vulnérabilité de l'exploitation au brouillage dans le cadre des spécifications proposées.

L'histoire a démontré que l'industrie américaine est capable de concevoir du matériel pratique et économique en vue d'une production de série. Nous avons l'assurance de l'industrie que les moyens considérables de conception et de production dont elle dispose seront concentrés sur ces problèmes qui demeurent.

Dispositifs terminaux.

Le standard F.C.C. pour la télévision en couleurs est le résultat d'une évolution s'étendant sur de nombreuses années d'études de procédés et de matériel. Nous décrirons seulement très brièvement ici les dispositifs utilisés aux extrémités du système, c'est-à-dire le tube de prises de vues de la caméra et le tube cathodique trichrome du récepteur.

Le type simple de caméra utilise trois tubes de prises de vues photosensibles connus sous le nom d'image-orthicon. Ce type de tube est, en fonctionnement normal, utilisé pour la télévision en noir et blanc et possède la haute sensibilité nécessaire pour la couleur. Le dispositif optique utilisé est montré sur la figure 5. On pourra voir qu'un sys-

tème diviseur d'image partage l'image réelle donnée par l'objectif, en trois composantes de couleurs, au moyen de miroirs dichroïques. Ils ont été spécialement préparés avec des plaques de verre qui sont sélectives à la couleur et qui permettent de transmettre une couleur et d'en réfléchir une autre. Des réglages précis sont prévus pour assurer des balayages identiques dans les trois tubes de prises de vues de telle sorte que la précision du recouvrement des balayages peut égaler celle du système optique. Une exploitation expérimentale poussée a montré que ce type de caméra donnait satisfaction. Ainsi chaque tube de prises de vues de la caméra voit l'image d'une certaine couleur et produit les signaux électriques correspondants. Ceux-ci sont transformés dans les équipements en un signal de luminance qui module l'émetteur de la manière habituelle et en signal de couleur qui module la sous-porteuse de chrominance. Certes, la caméra serait beaucoup plus simple si elle comportait un seul tube analyseur. On a recherché pendant un certain temps les méthodes pour y parvenir et des essais récents sur un tube de prises de vues trichrome développés dans les laboratoires RCA ont été très prometteurs. Bien qu'il reste beaucoup à faire, les principes de base de ce tube ont été établis, ainsi que sa possibilité de fabrication, de telle sorte que nous pouvons nous attendre à avoir une caméra avec un tube trichrome susceptible de remplacer la caméra actuelle à trois tubes. Les premiers récepteurs utilisaient trois tubes cathodiques avec filtres et miroirs optiques pour superposer les images. Une telle disposition était loin d'être idéale et des efforts spéciaux ont été faits pour développer un tube unique capable de donner une reproduction réelle de la couleur. Comme résultat, plusieurs types différents furent développés parmi lesquels le tube à masque trichrome à 3 faisceaux fut choisi pour la première fabrication.

Dans ce type de tube, l'écran lumineux de vision se compose d'un réseau régulier de très

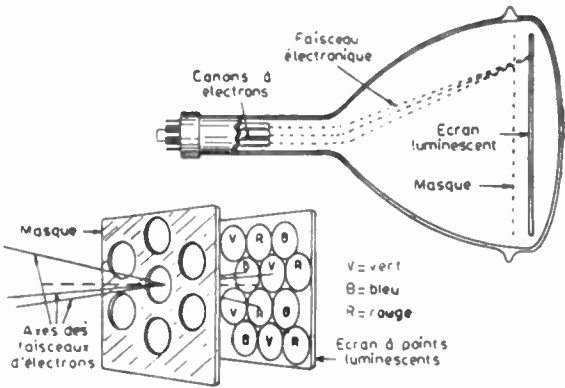


FIG. 6. — Schéma du tube cathodique trichrome RCA.

petites tâches phosphorescentes espacées légèrement et en groupes triangulaires disposées trois par trois déposées avec soin sur le verre comme l'indique la figure 6. Chaque point du trio est une tâche phosphorescente qui émet soit la lumière

rouge, soit la lumière verte, soit la lumière bleue lorsqu'elle est excitée par un faisceau électrique.

Trois canons électriques séparés fournissent les faisceaux qui provoquent leur luminescence et déterminent leur brillance. Un masque métallique contenant des ouvertures rondes en nombre égal au nombre de trios de points est interposé entre le canon électrique et la couche luminescente. La disposition géométrique des différentes parties est telle que les électrons issus du canon « rouge » atteignent seulement les taches luminescentes qui émettent de la lumière rouge et de la même manière les faisceaux des canons bleu et vert frappent seulement les substances correspondant à ces couleurs. Ainsi les signaux de couleur commandant les faisceaux électroniques produisent des images indépendantes de couleur primaire. Etant donné que les points sont très petits et très serrés, chaque trio se confond en un élément d'image unique d'où il résulte que l'œil voit l'image avec toutes les couleurs; un tube de ce type de 38 cm a presque 600 000 points, chacun ayant un diamètre de 0.35 mm approximativement et donne une image de 30 cm de large. La structure par points, à distance normale de vision est moins visible que la structure habituelle des lignes de balayage et le tube satisfait aux exigences du standard en blanc et noir.

Il y a eu une amélioration continue dans les performances du tube cathodique en couleurs portant sur une augmentation de lumière, un meilleur rendu des couleurs, un meilleur recouvrement et aussi dans les méthodes de fabrication, de telle sorte que le tube actuel donne d'excellents résultats. Il n'y a aucun doute toutefois que, vu le goût des Américains pour les récepteurs en noir et blanc à écran de grandes dimensions, il est indispensable de fabriquer des tubes plus grands pour la couleur.

On peut dire que des préparatifs sont déjà en route pour la fabrication d'un tube de 50 cm dans lequel de nouveaux perfectionnements seront apportés pour simplifier la fabrication et augmenter le rendement. Ce tube de 50 cm donne une image de près de 40 cm de large et a plus d'un million de points élémentaires.

Conclusion.

En conclusion, un bref compte rendu de l'état actuel de la télévision en couleurs en Amérique, a été donné. Les programmes en couleurs sont diffusés par la station NBC de NEW-YORK plusieurs fois par semaine, et les heures de programmes seront progressivement augmentées. On peut diffuser ces programmes sur tout le réseau. On fabrique des équipements pour la couleur et studios et émetteurs dans beaucoup d'endroits d'Amérique sont maintenant équipés pour diffuser la couleur. La RCA a offert de modifier gratuitement les émetteurs RCA existants, et sa filiale de radiodiffusion, la NBC diffuse des programmes en couleurs sans supplément de frais, de telle sorte que les annonceurs peuvent acquérir l'expérience de ce nouveau moyen.

Les récepteurs de couleurs et les tubes cathodiques trichromes sont fabriqués à la cadence de quelques milliers par mois et les récepteurs sont maintenant en vente dans le public au prix de 1 000 dollars (environ 350 livres). On estime que les demandes de récepteurs seront supérieures à la production cette année et l'année suivante et que vers 1957, l'industrie pourra vendre 3 millions de récepteurs en couleurs par an, ce chiffre montant jusqu'à 5 millions en 1958. Pour installer et dépanner ces récepteurs qui sont au moins 2 fois plus compliqués que les récepteurs en noir et blanc, un grand nombre de techniciens spécialistes sont en cours d'entraînement. Déjà plus de 30 000 d'entre eux ont suivi des cours spéciaux de « clinique » de la couleur organisés dans 65 villes importantes. Les circuits de relais hertziens transcontinentaux ont été égalisés en vue de la couleur et sont en usage régulier. De nombreux nouveaux circuits de transmission de cette sorte existeront à bref délai, de telle sorte qu'à la fin de 1954, plus de 100 stations dans le pays pourront être reliées en vue de faire un réseau d'émission d'images en couleurs, lesquelles pénétreront dans 75 % des foyers.

Après cela, nous comptons l'année prochaine, faire la démonstration d'un tube de 54 cm de forme rectangulaire dont le principe de fonctionnement est différent et qui donnera des images encore plus lumineuses. Le tube-couleur est réellement le cœur du système, car c'est son prix qui détermine essentiellement le prix du récepteur, c'est-à-dire le prix que le public doit payer pour voir des images en couleurs. Tout le reste doit être considéré comme d'importance secondaire. Beaucoup d'activité est donc déployée en beaucoup d'endroits pour faire baisser le prix du tube cathodique en couleurs et le succès plus ou moins grand de cet effort déterminera dans une large mesure le développement plus ou moins rapide de la radiodiffusion d'images en couleurs.

BIBLIOGRAPHIE

- [1] EPSTEIN, D.W. — Colorimetric Analysis of RCA Color Television System. *R.C.A. Rev.*, 14, 227 (1953).
- [2] WILLMER, E.N., WRIGHT, W.D. — Color Sensitivity of the Fovea Centralis. *Nature* 156, 119 (1945).
- [3] MIDDLETON, W.E.K., HOLMES, M.C. — The Apparent Colors of Surfaces of Small Subtense. A Preliminary Report. *J. Opt. Soc. Amer.*, 39, 582 (1949).
- [4] HARTRIDGE H. — The Visual Perception of Fine Detail. *Phil. Trans. Roy. Soc.* 232, 510 (1947).
- [5] BEDFORD, A.V. — Mixed Highs in Color Television. *Proc. Inst. Radio Engrs.* 38, 1003 (1950).
- [6] BROWN, G.H., LUCK, D.G.C. — Principles and Development of Color Television Systems. *R.C.A. Rev.*, 14, 144 (1953).
- [7] KELL R.D., SCHROEDER A.D. — Optimum Utilization of the Radio Frequency Channel for Color Television. *R.C.A. Rev.*, 14, 133 (1953).
- [8] Second Color Television Issue. *Proc. Inst. Radio Engrs.* 42, (1954).

PROCÉDÉ TECHNIQUE POUR L'AMÉLIORATION DES PERFORMANCES DES FAISCEAUX HERTZIENS EN TÉLÉPHONIE ⁽¹⁾

PAR

C. DUCOT

Laboratoires d'Electronique et de Physique Appliquées

Introduction.

On sait que les deux genres de modulations pouvant être employés dans les transmissions par faisceaux hertziens sont la modulation de fréquence et la modulation par impulsions. Dans son état actuel, la technique des impulsions n'est pas propre à la transmission d'information à grande vitesse qu'exigent à la fois la téléphonie multiplex, à partir d'un nombre moyen de voies, et la télévision, dans tous les cas. La modulation de fréquence reste donc le procédé normal de transmission sur faisceaux hertziens de ces types de signaux. Elle est seule notamment à pouvoir véhiculer un groupe de 60 voies téléphoniques ou plus comme un signal unique, ce qui est particulièrement intéressant lorsque les faisceaux hertziens doivent s'intégrer au réseau des grands circuits internationaux ou nationaux à courants porteurs.

Les limitations à la qualité des transmissions hertziennes proviennent de deux effets principaux de tous les systèmes : l'apparition d'un bruit venant se superposer au signal désiré indépendamment de celui-ci, et les déformations que ce signal peut subir au cours de la transmission. En téléphonie multiplex, l'effet de ces déformations est d'ailleurs, lorsqu'elles résultent de lois de transformation non-linéaires, d'ajouter au bruit précédent un autre bruit dit d'intermodulation qui dépend du signal transmis. Dans les systèmes à modulation de fréquence, le bruit indépendant du signal s'identifie pratiquement à un bruit de fluctuation, et le bruit d'intermodulation produit entre les voies une diaphonie dont la majeure partie est inintelligible, surtout lorsque le nombre de ces voies est assez élevé.

La présente communication a pour objet quelques considérations sur un procédé permettant d'améliorer dans certains cas les performances des faisceaux hertziens à modulation de fréquence pour la téléphonie multiplex.

Plus particulièrement, après quelques observations sur les recommandations du C.C.I.F. en matière de circuits téléphoniques, on étudiera les propriétés d'un système à deux modulations de fréquence échelonnées qui peut être employé avec avantage dans la transmission d'un signal téléphonique multiplex à nombre moyen de voies, par exemple de 60 à 240, et bien entendu aussi d'un nombre plus élevé réparti entre plusieurs ondes porteuses. Des résultats expérimentaux seront indiqués.

1^o Remarques préliminaires.

Les recommandations formulées par le C.C.I.F. à FLORENCE en octobre 1951 pour les faisceaux hertziens téléphoniques internationaux stipulent (Livre Jaune, tome III bis, p. 198) que « lorsque ceci peut être effectué, l'objectif doit être d'atteindre les qualités de transmission recommandées par le C.C.I.F. pour les circuits métalliques internationaux ».

La définition de ces qualités de transmission est rapportée à un « circuit fictif de référence » qui à la vérité ne convient pas parfaitement aux systèmes à faisceaux hertziens, car il laisse évidemment dans l'ombre les conditions particulières à ceux-ci, comme par exemple le régime des évanouissements. A ce propos le compte-rendu des réunions d'octobre 1953 à GENÈVE pose le problème :

Question nouvelle (3^e Commission d'études en coopération avec les 5^e et 4^e Commissions d'études) (catégorie A2) (très urgente) (correspond à la question n° 112 du C.C.I.R.).

Quelle est la variation (en fonction du temps) de l'écart entre signal et bruit que l'on peut admettre dans une communication téléphonique internationale ?

Remarque 1. - Le C.C.I.R. a besoin de divers renseignements pour établir les clauses essentielles d'une spécification pour les faisceaux hertziens ; dans des faisceaux hertziens de

(1) Communication faite à Milan, le 15 avril 1954, à l'occasion du Symposium d'Electronique et de Télévision.

construction moderne, on constate que l'équivalent ne varie pas sensiblement avec le temps, mais que les bruits peuvent parfois varier notablement (de l'ordre de plusieurs décibels).

.....
En attendant les résultats de l'étude de cette question, le C.C.I.F. communiquera au C.C.I.R. la spécification relative au bruit tolérable sur les systèmes modernes à courants porteurs sur lignes métalliques. Cette spécification comporte en particulier la clause que la limite fixée pour la puissance psophométrique peut être dépassée pendant 1 % du temps (Livre Jaune du C.C.I.F., tome III bis, page 121) ».

Le circuit fictif de référence sur paires coaxiales a une longueur de 2 500 kilomètres et comprend neuf couples de modulations de groupe secondaire. Il comporte en outre des couples de modulations de voie et de groupe primaire, mais ces couples n'intéressent pas le système à faisceaux hertziens.

La puissance psophométrique totale, incluant les bruits de fluctuation et d'intermodulation, introduite par le système à faisceaux hertziens sur ce circuit fictif, et ramenée en un point de niveau relatif zéro sur une voie, peut être assimilée à la puissance psophométrique due à la « ligne haute fréquence » coaxiale (avec ses répéteurs intermédiaires) ; elle ne doit donc pas dépasser pendant plus de 1 % du temps la valeur de 7 500 picowatts, si l'on veut se conformer à la recommandation initiale. Mais il est certain que ce pourcentage de temps a été établi pour tenir compte du caractère statistique du volume des conversations, et nullement pour prévoir des évanouissements éventuels dans la propagation hertzienne. Autrement dit, le calcul des bruits d'intermodulation devra pour nous se fonder sur ce pourcentage ; au contraire le calcul des bruits de fluctuation peut être affecté d'un coefficient encore indéterminé, qui d'ailleurs influencerait de la même façon sur la qualité des deux systèmes que nous voulons comparer, à savoir la simple et la double modulation de fréquence.

On peut admettre en gros que les puissances psophométriques sur une voie peuvent être calculées et mesurées comme des puissances moyennes sur une bande de 4 000 c/s, et qu'il suffit d'abaisser ces dernières d'environ 3 dB pour obtenir les puissances psophométriques avec une approximation suffisante.

D'autres conditions, par exemple de diaphonie intelligible, sont formulées, mais pour les caractéristiques des faisceaux hertziens à nombre de voies moyen ou élevé, c'est d'ordinaire la condition de bruit total qui demeure la plus déterminante.

Dans un système de transmission hertzienne à modulation de fréquence simple du type classique, il est bien connu que la grandeur caractéristique définissant le niveau du signal multiplex le long du faisceau est la déviation de fréquence correspondant à un signal dont la puissance est de 1 milliwatt en un point de niveau relatif zéro. On dit souvent que ce signal est au niveau 0 dBmo, signifiant par là qu'il se trouve au niveau de puissance 0 décibel par rapport à 1 milliwatt au point de niveau relatif zéro de la liaison.

La puissance du signal multiplex étant à chaque instant représentée sur le faisceau hertzien par le carré d'une déviation de fréquence, il est naturel d'exprimer celle-ci en valeur efficace. La déviation efficace de fréquence produite par un signal sinusoïdal est évidemment $\frac{1}{\sqrt{2}}$ fois la déviation de crête.

La grandeur caractéristique en question sera donc la déviation efficace d de fréquence sur la porteuse correspondant à un signal à 0 dBmo.

Il est évident que la valeur de d peut être déterminée librement de façon à rendre optimum la qualité de la transmission : moyennant l'adjonction ou le déplacement d'amplificateurs ou d'atténuateurs, on peut toujours, sans rien changer à l'essentiel du système à modulation de fréquence, faire varier d en maintenant constant l'équivalent de transmission.

Or, toutes choses égales d'ailleurs, l'augmentation de d entraîne, pour le faisceau hertzien et en un point de niveau relatif zéro, à la fois une diminution de la puissance de bruit de fluctuation et un accroissement de la puissance de bruit d'intermodulation. Un abaissement de d a des effets contraires. Il y a donc lieu de choisir d de façon à rendre minimum la somme de ces deux puissances de bruit.

Autrement dit il n'est pas logique de caractériser la qualité d'un système en donnant indépendamment le niveau de bruit de fluctuation et le niveau de bruit d'intermodulation en dBmo. Si ces deux chiffres étaient très différents, cela signifierait sans doute que d a été mal choisi.

On ne devra pas non plus comparer deux systèmes au point de vue du bruit de fluctuation seul ou du bruit d'intermodulation seul : c'est la puissance psophométrique de bruit total obtenue à la déviation optimum qui exprime adéquatement la qualité de chaque système.

Comme les intermodulations les plus importantes sont en général celles qui sont dues aux distorsions de non-linéarité du second ordre, sauf dans le cas assez exceptionnel où l'on peut situer le signal multiplex téléphonique à l'intérieur d'une octave en fréquence, on constate — c'est ce qu'ont fait STARR et WALKER dans leur magistrale communication « *Microwave Radio Links* » publiée dans les P.I.E.E. en septembre 1952 — que la puissance psophométrique totale sur une voie téléphonique, ramenée en un point de niveau relatif zéro, peut s'écrire en milliwatts sous la forme :

$$\frac{A}{d^2} + B d^2 \quad (1)$$

D'une part en effet le bruit de fluctuation physiquement présent à la sortie du dernier discriminateur et correspondant à la bande allouée à une certaine voie, en principe choisie la plus défavorisée, peut être représenté par un nombre A de (c/s)² efficaces de déviation de fréquence indépendant de la valeur de d qui sera adoptée. Dans la détermination de A interviennent la puissance émise, le gain

des aériens, la longueur et le nombre des tronçons, les conditions de propagation, le facteur de bruit des récepteurs et la fréquence moyenne de la bande allouée à la voie considérée.

D'autre part le bruit d'intermodulation du second ordre présent au même point et correspondant à la même bande peut être représenté par une grandeur Bd^4 , exprimée aussi en $(c/s)^2$ efficaces de déviation de fréquence, et dans laquelle le coefficient B est indépendant de d . Dans la détermination de B interviennent les diverses causes de distorsion, la puissance des signaux de conversation pour le nombre de voies fixé et en un point de niveau relatif zéro, et, comme pour A , la fréquence moyenne de la bande allouée à la voie considérée.

Pour avoir la puissance de bruit total en milliwatts au point de niveau relatif zéro, il suffit évidemment de diviser par d^4 la somme $A + Bd^4$, ce qui donne bien la quantité (1).

La déviation d optimum sera celle qui rend minimum cette quantité, c'est-à-dire :

$$d = \sqrt[4]{\frac{A}{B}}$$

La puissance de bruit total est alors :

$$2 \sqrt{AB}$$

et c'est cette grandeur qui caractérise la qualité intrinsèque de la liaison, étant d'autant plus petite que celle-ci est meilleure.

2° Le système à double modulation de fréquence.

Dans l'étude et la réalisation d'un système de faisceau hertzien à ondes décimétriques ou centimétriques et à simple modulation de fréquence, une des difficultés principales provient des conditions de linéarité à satisfaire pour que, compte tenu du bruit de fluctuation inévitable, la performance globale reste satisfaisante.

On peut évidemment réaliser la modulation dans des circuits à tubes classiques, triodes ou pentodes, qui offrent de nombreuses possibilités et une grande souplesse dans la constitution des circuits ; mais ces genres de circuits sont limités au domaine des fréquences inférieures, disons, à 150 Mc/s. On ne peut donc moduler, avec une bonne linéarité, il est vrai, qu'une porteuse auxiliaire qui devra être mélangée avec une onde à fréquence élevée pour donner la porteuse définitive. Le tube d'émission est alors nécessairement un tube amplificateur, et l'on sait que la puissance des tubes amplificateurs actuellement courants en hyperfréquences est assez limitée, de l'ordre du watt, ce qui limite la protection contre les bruits de fluctuation.

On peut au contraire préférer utiliser un tube d'émission oscillateur parce que ce genre de tubes atteint

en hyperfréquences, d'une façon générale et à niveau égal de la technologie, une puissance notablement supérieure, de l'ordre de 10 watts par exemple. Mais alors il faut pouvoir moduler directement le tube en fréquence avec une linéarité suffisante ; les caractéristiques de modulation des tubes oscillateurs de ce type permettent rarement de remplir cette condition. Comme la discrimination se fait toujours à des fréquences plus classiques, on pourrait bien réaliser un système à forte contre-réaction de fréquence, mais la mise au point d'une chaîne stable dans ce cas est souvent une opération délicate. De toute façon il faut répéter les démodulations et modulations à toutes les stations relais, ce qui est un inconvénient.

Un procédé permet de concilier l'emploi d'un oscillateur de puissance élevée et la modulation dans des circuits classiques dont la souplesse facilite le parachèvement d'une bonne linéarité. C'est le procédé de la double modulation de fréquence, proposé pour la première fois, à notre connaissance, par L.E. THOMPSON (cf. *P.I.R.E.*, 1946, 34, p. 936 : « *A Microwave Relay System* »), installé par la R.C.A. en 1945 sur une liaison expérimentale NEW-YORK-PHILADELPHIE (cf. *R.C.A. Review*, décembre 1946, « *Microwave Relay Communication System* », par GERLACH), puis employé par la WESTERN UNION TELEGRAPH CO sur une liaison NEW-YORK-WASHINGTON-PITTSBURGH (cf. *W.U.T.R.*, juillet 1952, « *A Radio Relay System employing a 4 000 Mc/s 3-cavity Klystron* », par J.J. LENEHAN). Ce procédé consiste à moduler en fréquence par le signal téléphonique multiplex une sous-porteuse, qui à son tour module en fréquence la porteuse engendrée par un oscillateur d'émission.

Incontestablement ce système présente une importante détérioration du quotient signal/bruit de fluctuation par rapport à une modulation simple qui utiliserait la même déviation de fréquence sur la porteuse, parce que le spectre de la sous-porteuse modulée ne peut être rapproché impunément de la fréquence zéro, ce qui rendrait impossible sa transmission correcte.

Cependant, comme il a été dit au paragraphe précédent, on ne peut comparer les systèmes du seul point de vue des bruits de fluctuation, et le système à double modulation présente notamment certains avantages :

1° La haute linéarité n'est imposée qu'à la première modulation et à la dernière démodulation qui s'opèrent à des fréquences relativement basses ; cet avantage est partagé par les systèmes à hétérodynes et amplificateurs, mais non par les systèmes à oscillateur d'émission directement modulé par le signal multiplex.

2° La déviation de fréquence sur la porteuse n'est plus limitée par des conditions de linéarité strictes, car la sous-porteuse traversera toujours un limiteur d'amplitude avant la seconde discrimination ; cela permet de pousser assez haut, par exemple à ± 5 Mc/s ou plus, cette déviation, ce qui vient en

compensation de la détérioration du quotient signal / bruit de fluctuation mentionnée précédemment.

3° Les distorsions non-linéaires de phase sur la porteuse ou la fréquence intermédiaire ne causent que des distorsions non-linéaires de phase considérablement plus faibles sur la sous-porteuse et n'agissent donc pratiquement pas sur l'intermodulation, ce qui est très favorable lors de la multiplication des étages à fréquence intermédiaire dans une liaison globale à un ou plusieurs tronçons, lorsqu'il y a désadaptation des guides d'antennes ou lorsqu'il y a propagation par trajets multiples : cet avantage n'est partagé par aucun des deux systèmes à simple modulation.

4° Il n'y a de seconde démodulation et de première remodulation qu'aux stations où le service l'exige ; les distorsions correspondantes ne sont donc multipliées que par le nombre de telles stations, et non par le nombre total de relais (environ 5 fois plus grand dans le circuit fictif) ; cet avantage concernant les distorsions est partagé naturellement par le système à hétérodynes et amplificateurs, mais seulement au moyen d'artifices onéreux et délicats (comme une très forte contre-réaction de fréquence) par le système à oscillateur d'émission directement modulé.

5° La surpuissance à l'émission, par rapport au système à hétérodynes et amplificateurs, donne, même à égalité de performance globale, une protection plus efficace contre les évanouissements, car il faudra un évanouissement plus profond pour que le quotient porteuse/bruit tombe au-dessous du seuil nécessaire au fonctionnement normal de la démodulation de fréquence : autrement dit certaines valeurs d'évanouissement qui pourraient causer de véritables interruptions sur le système sans surpuissance ne donneront qu'une détérioration de qualité sur le système avec surpuissance à l'émission.

6° Le système à double modulation offre en outre la possibilité d'introduire en simple modulation et en parallèle avec la sous-porteuse, normalement à fréquence plus basse qu'elle, une voie de service indépendante de la bande des signaux transmis avec une haute qualité : cette voie de service peut être notamment utilisée, totalement ou partiellement, pour la transmission, même permanente, de signaux d'essai et de contrôle facilitant la localisation des défauts lorsqu'il y a de nombreuses stations relais ne possédant pas de personnel fixe d'entretien.

Il faut dire également que le système à double modulation de fréquence se prêtait particulièrement au cas présent grâce à la possibilité d'utiliser un tube oscillateur d'émission exceptionnellement simple et robuste, le tube à réflexions multiples inventé par le Dr. F. COETERIER, des Laboratoires PHILIPS d'EINDHOVEN (v. p. ex. F. COETERIER, « *The Multi-reflection Tube — A new oscillator for very short waves* », Philips Technical Review, septembre 1946).

3° Comparaison quantitative entre les systèmes à double et à simple modulation de fréquence.

Quoi qu'il en soit, étant donné ce faisceau de caractéristiques du système à double modulation de fréquence, il semble maintenant indiqué de tenter une comparaison quantitative de principe entre sa performance et celle d'un système à simple modulation du type le plus moderne, c'est-à-dire à hétérodynes et amplificateurs.

Nous allons effectuer cette comparaison sur la base d'une application de chacun des deux systèmes au circuit fictif de référence de longueur 2 500 km comportant neuf couples de modulation et démodulation.

Nous supposons égaux tous les paramètres qui ne dépendent pas du choix de l'un ou de l'autre des deux systèmes, c'est-à-dire notamment le nombre et la longueur des tronçons, la longueur d'onde, le gain des aériens, les conditions de propagation et le facteur de bruit des récepteurs dans le spectre de la porteuse et de la fréquence intermédiaire.

Par contre nous supposons différentes les valeurs des paramètres liés au choix du système, c'est-à-dire notamment la puissance en watts du tube d'émission, P pour l'amplificateur du système classique, P' , ordinairement plus élevée, pour l'oscillateur du système à double modulation, et la déviation de fréquence en c/s efficaces μ sur la porteuse (en simple modulation) ou μ' sur la sous-porteuse (en double modulation) correspondant pour un signal modulateur sinusoïdal à un taux de 2^e harmonique fixé à l'avance, par exemple de 10^{-6} en puissance ou — 60 décibels.

Nous introduirons en outre dans le système à double modulation la déviation de fréquence de crête Δ c/s sur la porteuse, qui peut être choisie aussi grande que le permet une modulation saine du tube sans tenir compte de la stricte linéarité, et la fréquence moyenne F c/s de la sous-porteuse.

On peut montrer (les calculs sont donnés en appendice) que le coefficient d'amélioration de la performance par le système à double modulation peut alors s'écrire :

$$\gamma = \frac{1}{\sqrt{2}} \frac{\mu'}{\mu} \frac{\Delta}{F} \sqrt{\frac{P'}{P}}$$

Le coefficient γ représente le rapport entre les puissances de bruit total ramenées au point de niveau relatif zéro respectivement par le système à simple et par le système à double modulation de fréquence, la déviation efficace de fréquence pour 0 dBm étant supposée ajustée à sa valeur optimum tant sur la porteuse du système à simple modulation que sur la sous-porteuse du second système.

Dès que η est supérieur à 1, c'est le système à double modulation qui donne la meilleure qualité téléphonique. Si l'on trouve dans un cas particulier par cette formule une valeur η égale ou un peu inférieure à 1, on peut encore préférer le système à double modulation, du fait de sa protection beaucoup plus efficace contre un certain nombre de perturbations que nous avons négligées dans les calculs pour des raisons à la fois de simplicité et de concession à l'autre système, mais qui souvent ne seront pas négligeables dans le cas de la simple modulation, à savoir :

— les distorsions de phase des amplificateurs à fréquence intermédiaire ;

— les distorsions de phase dues à la désadaptation des guides d'antennes ; sur ce point, dans le système à double modulation, il suffit au bon fonctionnement que la condition de non-discontinuité de l'oscillation à l'émission en présence des effets d'entraînement soit satisfaite, condition évidemment bien moins sévère que les conditions de distorsion ;

— les distorsions dues au centrage imparfait de la fréquence intermédiaire dans la bande des amplificateurs ;

— les bruits de fluctuation dus à la limitation imparfaite (on voit tout de suite, par exemple sur la formule (3) de l'étude citée de STARR et WALKER, que les termes de bruit venant s'ajouter au bruit triangulaire sont d'un ordre de grandeur très inférieur en double modulation).

Il est raisonnable pour un système à 60 voies de prendre par exemple $\Delta = 5$ Mc/s et $F = 1.5$ Mc/s. Si le rapport $\frac{P'}{P}$ est alors d'environ 10, ce qui est dans l'ordre des réalisations actuelles, on voit que :

$$\eta = 7.5 \frac{\mu'}{\mu}$$

η sera supérieur à 1 dès que $\frac{\mu'}{\mu}$ dépassera 0,133.

Or une bonne valeur de μ correspondant à la technique actuelle des discriminateurs centrés vers 100 Mc/s est 0,35 Mc/s, ce qui revient à une déviation crête-à-crête de 1 Mc/s.

Si μ' est supérieure à 46,5 kc/s, c'est-à-dire si la déviation crête-à-crête sur la sous-porteuse correspondant à un taux de 2^e harmonique de — 60 dB est supérieure à 132 kc/s, ce qui peut être réalisé facilement pour des discriminateurs centrés autour de 1,5 Mc/s, on aura une valeur de η favorable. On peut aussi relever, par changement de fréquence à la réception, la fréquence moyenne de la sous-porteuse, par exemple jusqu'à 13 Mc/s, comme le faisait L. E. THOMPSON en 1945, mais cela devient inutile avec certains types de discriminateurs.

Il y a toujours intérêt, quand c'est possible, à augmenter la déviation Δ et à diminuer la fréquence sous-porteuse F .

4° Description d'un équipement expérimental à double modulation de fréquence.

Un équipement expérimental de faisceau hertzien téléphonique à double modulation de fréquence a été réalisé aux Laboratoires d'Electronique et de Physique appliquées.

Cet équipement se compose d'un poste émetteur et d'un poste récepteur, comprenant chacun une baie et une tête.

La figure 1 est une photographie de la baie et de la tête d'émission, la figure 2 une photographie de

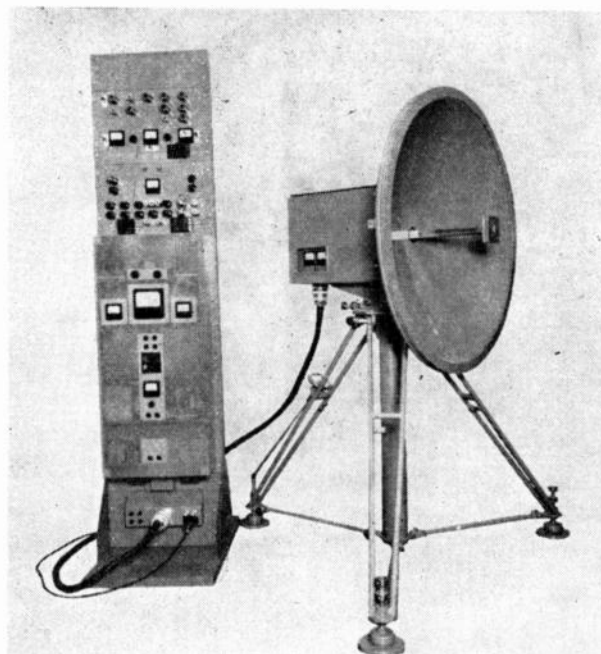


FIG. 1. — Emetteur expérimental pour faisceau hertzien à double modulation de fréquence.

la baie et de la tête de réception ; la figure 3 représente le schéma d'ensemble de la liaison expérimentale (les circuits purement séparateurs ne sont pas représentés).

Le poste émetteur comprend :

— un oscillateur à 9 Mc/s environ modulé en fréquence par le signal multiplex téléphonique à travers un tube à réactance ;

— un oscillateur fixe à 11 Mc/s ;

— un mélangeur fournissant la sous-porteuse à 2 Mc/s ;

— un filtre de sous-porteuse ;

— un stabilisateur de la fréquence sous-porteuse moyenne réagissant sur l'oscillateur modulé ;

— un amplificateur de sous-porteuse ;

— un émetteur à 3 500 Mc/s environ muni d'un tube C Z à réflexions multiples de puissance 10 watts ;

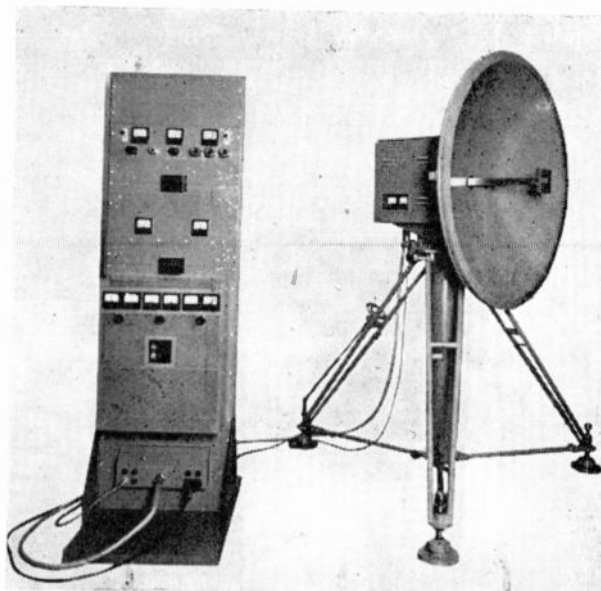


FIG. 2. — Récepteur expérimental pour faisceau hertzien à double modulation de fréquence.

— un stabilisateur de fréquence porteuse moyenne réagissant sur une électrode du tube C Z ; ce stabilisateur utilise deux cavités de référence dont les signaux de sortie sont commutés dans le temps sur un même cristal détecteur par un commutateur fondé sur l'absorption de pièces de ferrocube que l'on peut commander par un champ magnétique ;

— un aérien paraboloidal et son excitateur, du type à rayonnement arrière inventé par CUTLER.

Le poste récepteur comprend :

— un aérien et son collecteur, identiques à l'aérien et à l'excitateur d'émission ;

— un oscillateur local à tube triode et cavité extérieure ;

— un mélangeur produisant un signal à fréquence intermédiaire centrée autour de 58 Mc/s ;

— un amplificateur à fréquence intermédiaire muni d'un préamplificateur à faible bruit ; la bande globale est de $58 \text{ Mc/s} \pm 7 \text{ Mc/s}$ aux points situés à 3 dB du maximum ; l'amplificateur est du type à circuits décalés en fréquence ;

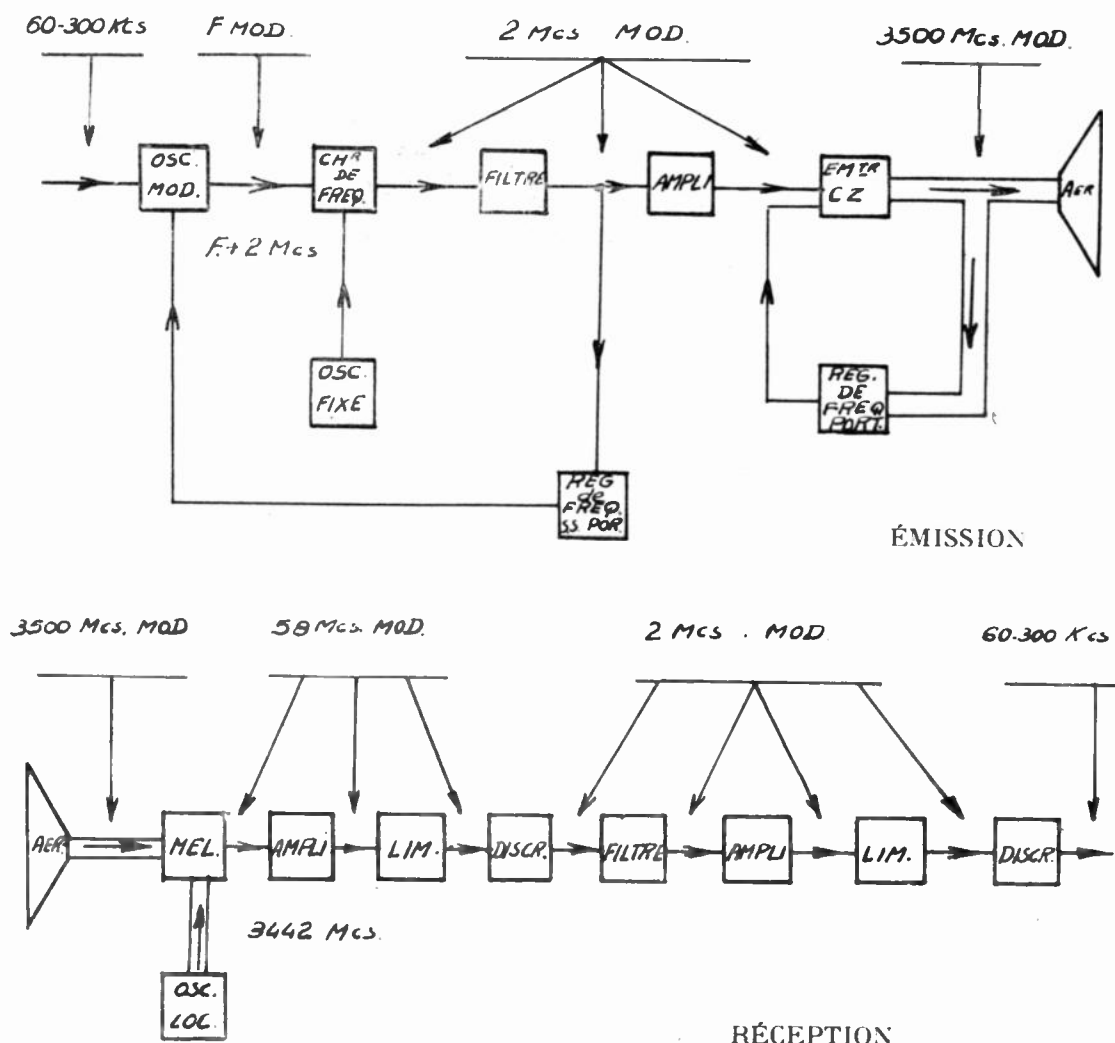


FIG. 3. — Schéma d'ensemble d'un équipement de faisceau hertzien téléphonique à double modulation de fréquence.

- un limiteur d'amplitude à 58 Mc/s ;
- un discriminateur restituant la sous-porteuse à 2 Mc/s ;
- un filtre de sous-porteuse ;
- un amplificateur de sous-porteuse ;
- un limiteur de sous-porteuse (terminé par un circuit sélectif pour atténuer les harmoniques du spectre de sous-porteuse) ;
- un discriminateur restituant le signal téléphonique multiplex initial.

5° Mesure des performances de l'équipement expérimental.

L'équipement a été essayé sur un tronçon dans les conditions suivantes :

- nombre de voies : 48 ;
- distance : 12,5 km ;
- aériens de diamètre 1 mètre ;
- puissance à l'émission réduite à 4 watts ;
- déviation de ± 5 Mc/s sur la porteuse ;
- affaiblissement systématique de 5 dB sur le signal à fréquence intermédiaire avant l'entrée du préamplificateur à faible bruit ;
- atténuation évaluée par recouplements à 10 dB environ sur le parcours du faisceau par suite d'un bâtiment faisant obstacle.

Autrement dit le bruit de fluctuation est le même que sur un tronçon de 50 km affecté d'un évanouissement de :

19 dB (les aériens de diamètre 1 m sont remplacés sur un tronçon de 50 km par des aériens semblables de diamètre 3 m).

- 12 dB (la distance étant quadruplée).
 - + 4 dB (l'affaiblissement à fréquence intermédiaire étant supprimé).
 - + 10 dB (l'obstacle étant supprimé).
- soit 21 dB d'évanouissement.

L'équipement était relié à un système à courants porteurs constitué comme suit : les voies sont transmises en bande latérale unique et situées dans des intervalles successifs de 4 kc/s. Le groupe I (12-60 kc/s) comprend 12 voies ; à l'intérieur de chacune les fréquences sont, quoique transposées bien entendu, rangées dans leur ordre naturel ; les voies sont numérotées de 1 (12-16 kc/s) à 12 (56-60 kc/s). Le groupe II (60-108 kc/s) comprend 12 voies où l'ordre des fréquences est inverse de l'ordre naturel ; elles sont numérotées de 1 (108-104 kc/s) à 12 (64-60 kc/s). Le groupe III (108-156 kc/s) comprend 12 voies inversées numérotées

de 1 (156-152 kc/s) à 12 (112-108 kc/s). De même le groupe IV (156-204 kc/s) comprend 12 voies inversées numérotées dans l'ordre des fréquences descendantes.

Le principe des mesures consistait à charger à l'émission 10 voies sur les 18 voies avec des signaux incohérents de puissances égales représentant les signaux de conversation et fournis par un générateur de bruit, puis à effectuer les mesures permettant de tracer la courbe donnant la puissance de perturbation totale à la réception dans l'une des 8 dernières voies, exprimée en dBmo en fonction de la puissance totale des signaux incohérents introduits à l'entrée exprimée elle aussi en dBmo.

En réalité le générateur de bruit était un amplificateur de groupe primaire transposé (bande 48 Kc/s), dont la démodulation sur quatre canaux en parallèle permettait d'alimenter en bruit simultanément les voies des groupes I, II, III et IV.

Théoriquement, il y a donc entre les voies d'un groupe l'incohérence mutuelle parfaite qui correspond probablement au service réel. En ce qui concerne les distorsions dues à la transmission H. F., l'incohérence est également assurée entre les voies, même homologues, de deux groupes déployés en sens inverses (groupe I et II) mais non de deux groupes de même sens (groupes II et III). Pratiquement la similitude des résultats dans ces deux cas laisse présumer qu'une incohérence de fait se trouve réintroduite, peut-être par le jeu de certains décalages en phase.

Avant de tracer la courbe du bruit total en fonction de la puissance de charge fictive, il convient de déterminer rapidement la valeur optimum de d , déviation efficace de fréquence sous-porteuse correspondant au niveau 0 dBmo, afin de centrer correctement la plage des niveaux parcourus.

Pour déterminer rapidement la valeur optimum de d , on introduit le signal global de charge fictive à l'entrée d'un atténuateur calibré A, dont la sortie est reliée au châssis de première modulation du faisceau hertzien. La sortie du châssis de seconde démodulation du faisceau hertzien est reliée, à travers un atténuateur calibré B, à l'entrée à bas niveau relatif du récepteur à courants porteurs.

En faisant varier d'un nombre de dB chaque fois égal et opposé les atténuations de A et B, on maintient constant l'équivalent (quel qu'il soit) de l'ensemble A + faisceau hertzien + B, tout en faisant varier la déviation produite par la charge fictive.

On constate alors que pour une certaine position des atténuateurs la perturbation totale sur une des voies non chargées est minimum ; le minimum est d'ailleurs peu aigu, et il faut souvent 2 ou 3 dB de variation sur la déviation pour que la perturbation varie sensiblement.

Ce minimum de perturbation totale correspond à une déviation efficace δ où devra s'établir le niveau + 1,2 dBmo, charge globale de conversation qui n'est pas dépassée plus de 1 % du temps dans un système à 48 voies non contrôlées en volume, selon le chiffre de HOLBROOK et DIXON (« Load Rating

Theory for Multi-Channel Amplifiers », B. S. T. J. octobre 1939) diminué de 5 dB conformément aux expériences récentes.

La déviation efficace d optimum correspondante est évidemment de 1,2 dB plus basse, soit 0,87 δ . On a trouvé comme moyenne des d optima (peu aigus) pour les voies I-1 et IV-1 la valeur 10 kc/s.

Les mesures de bruit total à charge fictive variable ont alors été effectuées. Les résultats sont indiqués pour la voie I-1 sur la figure 4, pour la voie II-1 sur la figure 5 et pour la voie IV-1 sur la figure 6. La mesure était faite en bande de voie uniforme, sans filtre psophométrique.

Le bruit total à la charge + 1,2 dBmo était à environ :

- 60 dBmo en voie I-1
- 59,5 dBmo en voie II-1
- 59 dBmo en voie IV-1

Le bruit à vide était à :

- 65 dBmo en voie I-1
- 64 dBmo en voie II-1
- 62 dBmo en voie IV-1

6° Conclusions des mesures précédentes.

A partir de l'une des courbes, en principe de la plus défavorable, c'est-à-dire de celle qui concerne

la voie IV-1, on peut tenter une extrapolation au « circuit fictif de référence » dont il a été question (50 tronçons de 50 km, 9 modulations et démodulations complètes).

En admettant la condition d'évanouissement mentionnée dans l'appendice (évanouissements de 20 dB sur 2 tronçons simultanément, propagation nominale sur les 48 autres tronçons) on doit prévoir un bruit de fluctuation 248 fois plus fort pour la liaison globale que pour le tronçon unique sans évanouissement, soit une augmentation de 24 dB, et donc de 3 dB par rapport aux conditions de l'expérience définies plus haut. On peut d'ailleurs penser que cette condition est sévère, car il semble par exemple que dans le système TD 2 les évanouissements supérieurs à 20 dB ne soient atteints pendant 1 % du temps que sur le moins favorisé des 107 tronçons.

Quant aux puissances d'intermodulation, elles doivent être multipliées par 9, c'est-à-dire relevées de 9,5 dB. Le résultat est la seconde courbe de la figure 6.

Le nouveau niveau + 1,2 dBmo optimum se trouverait à l'ancien niveau — 2 dBmo ; donc la déviation efficace d optimum devrait être diminuée de 3,2 dB, et la performance en bruit total à la (nouvelle) charge + 1,2 dBmo s'établirait (voir la figure) à — 55,7 dBmo + 3,2 dB = — 52,5 dBmo.

Ce chiffre doit être comparé à la valeur recommandée par le C.C.I.F. qui en mesure non psophométrique

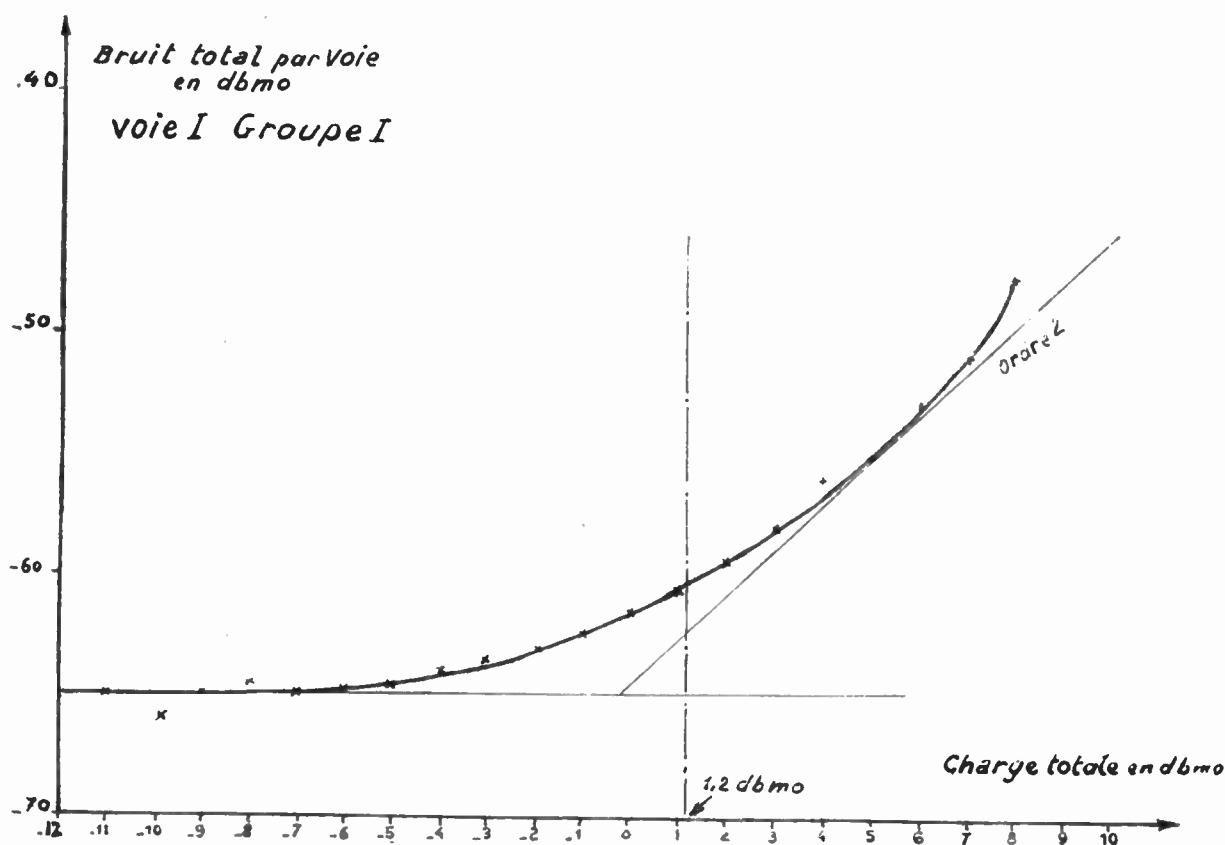


FIG. 4. — Puissance de bruit total en voie I du groupe I, en fonction de la charge : en groupes I, II, III, IV (12-204 kc/s).

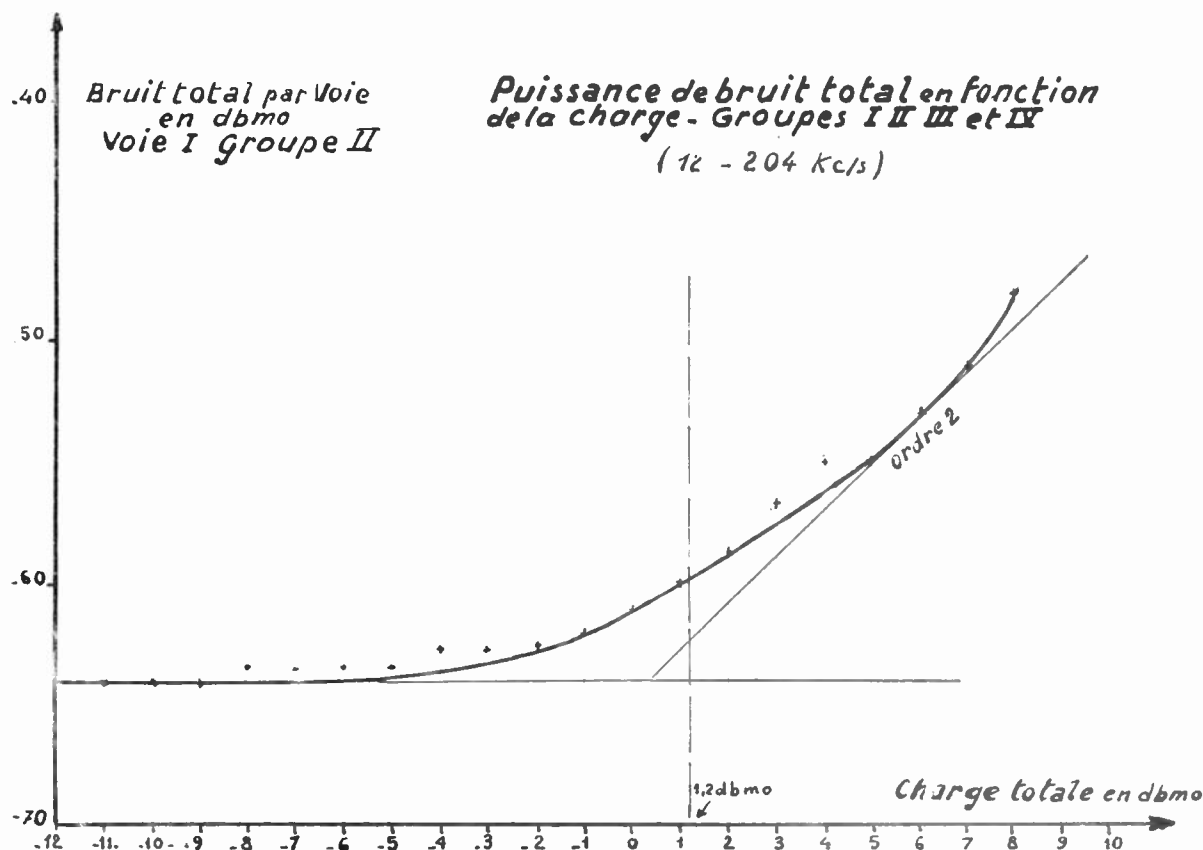


Fig. 5. — Puissance de bruit total en voie 1 du groupe II, en fonction de la charge en groupes I, II, III, IV (12-204 kc/s)

correspond à 15 000 picowatts au point de niveau relatif zéro, soit $-48,25$ dBmo. Ainsi, dans des conditions de propagation assez sévères, les recommandations du C.C.I.F. seraient satisfaites sur 48 voies avec une marge de 4 dB.

On peut raisonnablement présumer que des essais sur 60 voies donneraient encore un résultat conforme aux recommandations du C.C.I.F. En effet, le système est très loin de son point de surcharge ; la voie supérieure passerait de 204 à 300 kc/s, d'où une détérioration de 3,4 dB en bruit de fluctuation ; la charge selon HOLBROOK et DIXON s'élèverait de 0,3 dB et par suite, même dans l'hypothèse pessimiste où toutes les distorsions se comporteraient comme des distorsions de phase, le bruit d'intermodulation s'élèverait de $3,4 + 0,6 = 4$ dB, soit une détérioration moyenne de performance totale de 3,7 dB seulement.

Un dernier résultat intéressant des essais effectués sur la liaison expérimentale est la grande stabilité des performances en présence de dérives artificiellement provoquées sur certains paramètres de l'équipement hertzien. Ces essais ont été effectués lorsque cet équipement n'était encore associé qu'à un système à courants porteurs à 21 voies.

A titre indicatif il fallait faire varier la fréquence de l'oscillateur local de réception de 5 Mc/s au moins dans un sens ou dans l'autre pour obtenir une aug-

mentation de 3 dB. sur la puissance de bruit global dans la voie avantagée (la voie désavantagée étant moins sensible). La stabilité naturelle de l'oscillateur est de l'ordre du mégacycle/seconde.

Pour la même détérioration, il fallait faire dériver la fréquence moyenne sous-porteuse de 50 kc/s au moins dans un sens ou dans l'autre. Le circuit de stabilisation maintient cette fréquence dans des limites environ 5 fois plus étroites.

D'une façon générale on peut dire que le fonctionnement du système à double modulation de fréquence, malgré une complication apparente, se caractérise par une *stabilité de performance* et un « confort » remarquables.

On peut d'ailleurs ajouter que dans le système à hétérodynes et amplificateurs la complication fonctionnelle est en réalité aussi grande, car à l'échelonnement des deux modulations de fréquence correspond celui d'une modulation de fréquence et d'un changement de fréquence qui revient à une modulation d'amplitude.

Dans ces conditions une extension jusqu'à 240 voies sur une seule porteuse peut être envisagée moyennant l'emploi d'un tube d'émission de facteur de mérite environ 14 fois plus élevé, et dans le cas présent le facteur de mérite peut être représenté par le produit de la puissance P par le carré Δ^2 de la

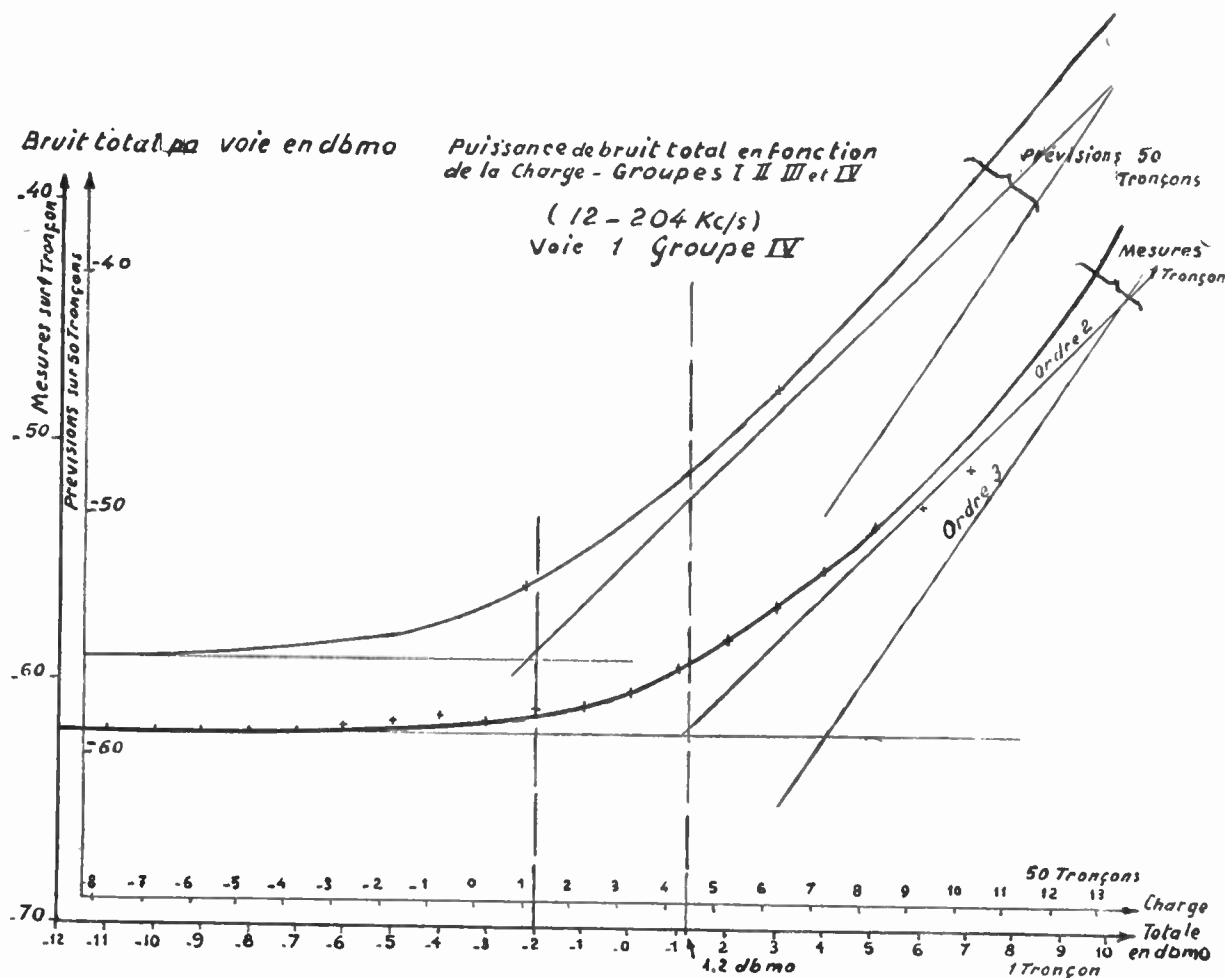


Fig. 6. — Puissance de bruit total en voie 1 du groupe IV, en fonction de la charge en groupes I, II, III, IV, (12-204 kc/s) et prévision pour tronçons.

déviations de crête utilisables, car aucune condition de distorsion n'est imposée pour cette modulation. Le facteur 14 provient de ce que pour transmettre 240 voies il faut en principe dans les bandes C.C.I.F. une bande 3,5 fois plus grande que pour en transmettre 60, et que si l'on admet une certaine loi d'homothétie à distorsion égale pour le second discriminateur (qui détermine la distorsion) le passage à une sous-porteuse de fréquence 3,5 fois plus élevée et modulée avec une déviation 3,5 fois plus élevée également donnerait, P et Δ restant inchangés, les changements suivants :

bruit de fluctuation : + 10,9 dB (F étant multiplié par 3,5),

— 10,9 dB (la déviation de sous-porteuse étant multipliée par 3,5),

+ 10,9 dB (la fréquence défavorisée du signal téléphonique multiplex étant multipliée par 3,5) ;

puissance de distorsion :

+ 7 dB (augmentation de la charge totale de 3,5 dB selon HOLBROOK et DIXON) ;

— 6 dB (répartition dans 4 fois plus de voies).

soit une détérioration de 10,9 dB en bruit de fluctuation et de 1 dB en bruit d'intermodulation, ensemble équivalent (par suite de l'interdépendance des deux bruits) à une détérioration de 11,5 dB du bruit de fluctuation seul, que l'on peut compenser par une multiplication par 14 du produit $P\Delta^2$.

Si l'on passait de la puissance 4 watts employée lors des mesures à une puissance de 20 watts par exemple, il suffirait que Δ soit multiplié par

$\sqrt{\frac{14}{5}} = 1,7$, c'est-à-dire que le tube soit modulable de $\pm 8,5$ Mc/s au lieu de ± 5 Mc/s.

L'auteur tient à remercier M. le Professeur G.A. BOUTRY, Directeur des Laboratoires d'Electronique et de Physique appliquées, d'avoir bien voulu autoriser la publication du présent article, ainsi que MM. G. ANDRIEUX, J. CAYZAC, R. ROY et leurs

collaborateurs, qui ont réalisé l'équipement expérimental ici décrit, pour leur participation importante à ce travail.

APPENDICE

A) Evaluation de la performance d'un système à simple modulation de fréquence.

Il nous faut supposer le système appliqué au circuit fictif de référence de longueur 2 500 km comportant neuf couples de modulation et démodulation.

On sait que pour les faisceaux hertziens la longueur moyenne des tronçons est limitée par la visibilité optique, condition nécessaire à la propagation correcte des ondes décimétriques ou centimétriques employées, à environ 50 kilomètres.

N'oublions d'ailleurs pas que le raccourcissement des tronçons et par suite l'augmentation de leur nombre sur une distance déterminée ont un effet favorable sur les bruits de fluctuation, puisque, toutes choses égales d'ailleurs, la puissance de bruit due à un tronçon est proportionnelle au carré de sa longueur tandis que le facteur par lequel elle se trouve multipliée sur la liaison totale n'est inversement proportionnel qu'à la longueur du tronçon moyen, au moins quand les divers tronçons sont à peu près égaux. Cependant cet effet favorable est très onéreux, et c'est pourquoi nous admettrons comme raisonnable la valeur de 50 km, qui conduit à 50 tronçons pour une liaison de 2 500 km.

Nous calculerons d'abord la puissance de bruit total correspondant à un seul tronçon de la liaison hertzienne et la mettrons sous la forme (1).

Plaçons-nous pour fixer les idées dans le cas d'un système destiné à transmettre 60 voies téléphoniques étagées de 4 en 4 kc/s entre 60 et 300 kc/s et commençons par le bruit de fluctuation.

Soient P la puissance rayonnée en watts, λ la longueur d'onde en centimètres, N le facteur de bruit du récepteur.

Supposons que les aériens sont du type paraboloïde et que par souci de stabilité on ne veuille pas rendre l'ouverture des faisceaux à 3 dB du maximum inférieure à 2°. Avec l'illumination habituelle le diamètre $\varnothing$ de l'ouverture du paraboloïde devra être tel que :

$$70^\circ \frac{\lambda}{\varnothing} = 2^\circ \quad \text{ou} \quad \varnothing = 35 \lambda$$

La surface de l'ouverture est alors :

$$\frac{\pi \varnothing^2}{4} = 306 \pi \lambda^2$$

Avec une efficacité ordinaire de 0,55 on a donc une surface effective de $528 \lambda^2$, et par suite la puissance reçue à 50 km est en watts :

$$p = P \frac{528 \lambda^2}{\lambda^2 (5 \cdot 10^4)^2} = 1,11 \cdot 10^{-8} P \lambda^2$$

D'autre part la puissance de bruit par c/s de bande dans le spectre de la porteuse à la réception est en watts, une fois majorée par le facteur de bruit :

$$b = 1,375 \cdot 10^{-23} \cdot 293 N = 4,03 \cdot 10^{-21} N$$

Le bruit par c/s de bande après discrimination, exprimé en (c/s)² efficaces de déviation sur la porteuse, et au voisinage de la fréquence f c/s dans le spectre du signal démodulé, peut s'écrire alors, les bruits dus à l'imperfection des limiteurs étant négligés ainsi que l'éventuelle préaccentuation (voir par exemple STARR et WALKER, loc. cit., p. 247) :

$$\beta = \frac{b}{p} f^2 = \frac{4,03 \cdot 10^{-21}}{1,11 \cdot 10^{-8}} \frac{N f^2}{P \lambda^2} = 3,63 \cdot 10^{-13} \frac{N f^2}{P \lambda^2}$$

Dans une bande de 4 kc/s autour de 300 kc/s (extrémité désavantagée du spectre) on aurait donc en (c/s)² efficaces de déviation un bruit égal à :

$$3,63 \cdot 10^{-13} \frac{N}{P \lambda^2} 9 \cdot 10^{10} \cdot 4 \cdot 10^3 = 130 \frac{N}{P \lambda^2}$$

Pour tenir compte de la mesure psophométrique on peut diviser ce chiffre par 2.

Par ailleurs si l'on veut calculer la performance non pas dans les conditions de propagation nominales ou en espace libre, mais dans la condition d'évanouissement qui n'est pas dépassée pendant plus d'un certain pourcentage du temps (encore à définir, comme nous l'avons dit précédemment), il faudra multiplier ce bruit par un coefficient α_1 correspondant : α_1 est évidemment égal à $10^{\theta/10}$ si θ est l'affaiblissement correspondant de la porteuse en décibels.

On exprimera alors la grandeur A_1 représentant en (c/s)² efficaces de déviation de fréquence le bruit de fluctuation sur la voie la plus élevée comme suit :

$$A_1 = 65 \frac{N \alpha_1}{P \lambda^2}$$

Si d est la déviation efficace de fréquence correspondant à 0 dBm₀, le bruit de fluctuation dans la voie désavantagée est, en milliwatts et ramené au point de niveau relatif zéro :

$$\frac{A_1}{d^2} = 65 \frac{N \alpha_1}{P \lambda^2 d^2}$$

La grandeur A_1 correspond à la grandeur A de l'expression (1), l'indice 1 signifiant seulement qu'un tronçon unique est considéré.

Calculons maintenant la puissance d'intermodulation du second ordre.

Nous admettrons que les intermodulations d'ordres supérieurs sont négligeables ; il suffit qu'elles le soient au niveau optimum auquel on sera amené à situer la déviation de fréquence, ce qui est le plus souvent réalisé.

Nous admettons en outre que les distorsions dues à la non-linéarité en phase des amplificateurs à fréquence intermédiaire ont été suffisamment réduites, par élargissement de la bande, par choix du circuit convenable ou par compensation, pour qu'on puisse les négliger — même, dans la suite, lorsqu'elles sont répétées sur 5 tronçons — devant les distorsions de non-linéarité du modulateur et du discriminateur. Cette hypothèse est peut-être un peu optimiste (sur un exemple, STARR et WALKER trouvent pour les nombres de tronçons supérieurs à 5 ou 6 une distorsion due aux répéteurs plus élevée que la distorsion de modulation et démodulation), mais comme notre propos est de souligner les avantages apportés dans certains cas par le système à double modulation de fréquence, qui n'est pas perturbé par les distorsions de phase des amplificateurs à fréquence intermédiaire, nous devons présenter le système plus classique sous son jour le plus favorable.

Pour évaluer la puissance d'intermodulation, il nous faut connaître d'abord en dBmo le niveau global des signaux de conversation sur 60 voies pour lequel il convient de calculer cette puissance, puis le taux de 2^e harmonique correspondant à une certaine déviation de fréquence ou, ce qui revient au même, la déviation efficace de fréquence μ pour laquelle la chaîne modulateur-démodulateur donne un taux de 2^e harmonique fixé à l'avance à une valeur arbitraire que nous choisirons égale à 10^{-6} en puissance ou — 60 dB ; enfin il nous faut connaître la loi reliant la puissance d'intermodulation du 2^e ordre dans une voie au taux de 2^e harmonique donné par un signal sinusoïdal.

La première indication (niveau de conversation sur 60 voies pour lequel il convient de chiffrer l'intermodulation) peut être tirée de la célèbre étude statistique publiée par HOLBROOK et DIXON dans le B.S. T.J. d'octobre 1939 sous le titre « *Load Rating Theory for Multi-Channel Amplifiers* ». Sans entrer dans le détail des multiples discussions que la question suscite encore, nous admettons avec la tendance actuelle que les résultats statistiques de cette étude restent valables aujourd'hui, mais que les puissances de conversation doivent toutes être, selon les récentes expériences, abaissées de 5 dB. par rapport aux chiffres qu'elle indique.

La mesure de la puissance de bruit devant être faite au psophomètre, appareil intégrateur, il nous suffira de considérer la puissance de conversation correspondant au volume qui n'est pas dépassé pendant plus de 1 % du temps, afin de nous conformer à la recommandation du C.C.I.F. pour les circuits métalliques.

Ce volume est donné par la fig. 6 de l'étude de HOLBROOK et DIXON, qui pour 60 voies, dont nous supposons encore les volumes individuels non contrôlés, indique environ + 4,3 dB Vol. Le volume de référence correspondant à 1,66 mW au point de niveau relatif zéro, soit à + 2,2 dBmo, le volume indiqué pour 60 voies correspond à + 6,5 dBmo, et puis-que nous admettons un abaissement de 5 dB par rapport à ces données, nous fixerons à + 1,5 dBmo le niveau global de conversation pour lequel nous

devons évaluer l'intermodulation, ce qui correspond à une déviation efficace de fréquence de 1,19 d .

La valeur de μ , déviation efficace de fréquence donnant un taux de 2^e harmonique de — 60 dB par rapport au fondamental, est une propriété du système.

Quant à la loi donnant la puissance d'intermodulation du second ordre dans la voie supérieure en fonction du taux de 2^e harmonique produit par un signal sinusoïdal, elle nous est fournie par exemple par l'étude de BROCKBANK et WASS intitulée « *Non-linear distortion in transmitter systems* » et publiée dans le J.I.E.E., 1945, 92, partie III, p. 45. Cette loi indique que la puissance d'intermodulation ne dépend que de la puissance totale des signaux de conversation et non de la distribution de cette puissance entre les diverses composantes présentes, pourvu que celles-ci soient en assez grand nombre.

La puissance d'intermodulation totale du second ordre, exprimée en cycles carrés efficaces de déviation sur la porteuse, sera d'après BROCKBANK et WASS :

$$4 \cdot 10^{-6} \frac{(1,19 d)^4}{\mu^2} = 8 \cdot 10^{-6} \frac{d^4}{\mu^2}$$

Dans la voie supérieure on aura sur la bande uniforme de 4 kc/s une puissance d'intermodulation égale à la précédente divisée par le nombre de voies, soit 60, et multipliée par le coefficient de distribution 0,5 indiqué par la figure 1 de l'étude de BROCKBANK et WASS, soit :

$$8 \cdot 10^{-6} \frac{d^4}{\mu^2} \cdot \frac{0,5}{60} = 6,66 \cdot 10^{-8} \frac{d^4}{\mu^2}$$

Pour tenir compte de la mesure psophométrique on peut diviser encore ce chiffre par 2, en sorte qu'en (c/s)² efficaces de déviation le bruit d'intermodulation sur la voie la plus élevée pourra être exprimé par :

$$B_1 d^4 = 3,33 \cdot 10^{-8} \frac{d^4}{\mu^2}$$

Ce bruit d'intermodulation, exprimé en milliwatts et ramené au point de niveau relatif zéro, sera donc :

$$B_1 d^2 = 3,33 \cdot 10^{-8} \frac{d^4}{\mu^2}$$

La puissance psophométrique totale de bruit de fluctuation et d'intermodulation produite dans la voie supérieure par le tronçon de faisceau hertzien est alors en milliwatts au point de niveau relatif zéro :

$$\frac{A_1}{d^2} + B_1 d^2 = 65 \frac{N \alpha_1}{P \lambda^2 d^2} + 3,33 \cdot 10^{-8} \frac{d^2}{\mu^2}$$

Il est correct d'avoir choisi la voie supérieure comme la plus désavantagée (alors que du point de vue

de l'intermodulation elle est au contraire une des plus avantageées), parce que, lorsqu'on abaisse la fréquence f de la voie choisie, on voit décroître le bruit de fluctuation plus vite (en f^2) que ne croît le bruit d'intermodulation (presque proportionnellement à 600 kc/s — f).

Pour obtenir la puissance psophométrique de bruit total sur le circuit fictif de référence, qui comporte 50 tronçons de 50 km, on devra d'abord multiplier par 50 le terme de bruit de fluctuation, en remplaçant toutefois le coefficient α_1 par un autre plus petit α , car s'il y a sur un tronçon un évanouissement de 10 log α_1 décibels pendant seulement un certain pourcentage du temps, il n'y aura pas pendant le même pourcentage du temps simultanément sur tous les tronçons des évanouissements tels que (les indices correspondant aux tronçons successifs) :

$$\alpha_1 + \alpha_2 + \dots + \alpha_{50} = 50 \alpha_1$$

La détermination de ce nouveau coefficient α est difficile car il dépend des lois statistiques de la propagation sur les divers tronçons. On admet quelquefois que sur 50 tronçons on peut considérer comme la condition la plus sévère dans laquelle les exigences doivent rester satisfaites la présence d'un évanouissement de 20 dB simultanément sur 2 tronçons. Dans ce cas, au lieu d'écrire pour le bruit de fluctuation du circuit fictif :

$$\frac{A}{d^2} = 3250 \frac{N \alpha}{P \lambda^2 d^2}$$

on peut écrire :

$$\frac{A}{d^2} = 248 \times 65 \frac{N}{P \lambda^2 d^2} = 16 \ 120 \frac{N}{P \lambda^2 d^2}$$

ce qui correspond à : $\alpha = 4,96$

Dans la suite nous conserverons pourtant la forme en α indéterminé, car cela ne change rien à la comparaison des systèmes.

Il faut ensuite multiplier par 9 la puissance d'intermodulation sur un tronçon pour obtenir celle qui apparaît dans le circuit fictif de référence ; encore admet-on là, comme nous l'avons vu, que la transmission peut se faire sans démodulation de fréquence dans les relais où celle-ci n'est pas exigée par le service et que les distorsions dues aux amplificateurs à fréquence intermédiaire ou aux trajets multiples de la propagation sont négligeables, multipliées par 50, devant des distorsions de modulation et de démodulation multipliées seulement par 9.

On aura donc sur le circuit fictif total une puissance d'intermodulation :

$$B d^2 = 9 \times 3,33 \cdot 10^{-8} \frac{d^2}{\mu^2} = 3 \cdot 10^{-7} \frac{d^2}{\mu^2}$$

et une puissance totale de perturbation :

$$\frac{A}{d^2} + B d^2 = 3,25 \cdot 10^3 \frac{N \alpha}{P \lambda^2 d^2} + 3 \cdot 10^{-7} \frac{d^2}{\mu^2}$$

La déviation efficace d optimum pour le niveau de signal 0 dBmo est :

$$d = \sqrt[4]{\frac{A}{B}} = \sqrt[4]{\frac{3,25 \cdot 10^3}{3 \cdot 10^{-7}} \frac{N \alpha \mu^2}{P \lambda^2}} = \sqrt[4]{1,08 \cdot 10^{10} \frac{N \alpha \mu^2}{P \lambda^2}}$$

et la puissance de perturbation en milliwatts au point de niveau relatif zéro est alors pour la voie désavantagée :

$$2 \sqrt{AB} = 2 \sqrt{3,25 \cdot 10^3 \frac{N \alpha}{P \lambda^2} 3 \cdot 10^{-7} \frac{1}{\mu^2}}$$

ou :

$$2 \sqrt{AB} = \frac{0,0624}{\lambda \mu} \sqrt{\frac{N \alpha}{P}}$$

B. — Evaluation de la performance d'un système à double modulation de fréquence et comparaison.

Commençons, comme pour le système précédent, par l'évaluation de la perturbation apportée par un tronçon.

Le début des calculs étant identique, on peut exprimer comme suit le bruit de fluctuation β' par c/s de bande après discrimination, mesuré en $(c/s)^2$ efficaces de déviation sur la porteuse, et au voisinage de la fréquence sous-porteuse F :

$$\beta' = 3,63 \cdot 10^{-13} \frac{N' F^2}{P' \lambda'^2}$$

Les lettres avec indice prime désignent les grandeurs du système à double modulation homologues des grandeurs désignées par les mêmes lettres sans indice dans le système à simple modulation. On peut dès maintenant considérer la longueur d'onde et le facteur de bruit de réception comme n'ayant aucune raison de différer dans les deux systèmes, et écrire :

$$\beta' = 3,63 \cdot 10^{-13} \frac{N F^2}{P \lambda^2}$$

Par une seconde application de la même formule, nous obtiendrons le bruit par c/s de bande après seconde discrimination, exprimé en $(c/s)^2$ efficaces de déviation sur la sous-porteuse, et au voisinage de la fréquence f c/s dans le spectre du signal démodulé :

$$\begin{aligned} \beta'' &= \frac{3,63 \cdot 10^{-13}}{\Delta^2/2} \frac{N F^2}{P' \lambda'^2} f^2 \\ &= 3,63 \cdot 10^{-13} \frac{N f^2}{P' \lambda^2} \frac{2 F^2}{\Delta^2} \end{aligned}$$

Δ étant la déviation de fréquence de crête sur la porteuse.

On voit que tous les résultats concernant les bruits de fluctuation vont être similaires à ceux du paragraphe précédent, les seules différences étant le remplacement de P par P' et la multiplication par le facteur $\frac{2 F^2}{\Delta^2}$

De même tous les résultats concernant les intermodulations du second ordre seront similaires à ceux du paragraphe précédent, à condition d'appeler d' la déviation efficace de fréquence sur la sous-porteuse correspondant à un signal à 0 dBmo. et μ' la déviation efficace de fréquence sur la sous-porteuse donnant un taux de 2^e harmonique de — 60 dB par rapport au fondamental.

On pourra donc écrire la puissance de bruit total due au circuit fictif global :

$$\frac{A'}{d'^2} + B' d'^2 = 6,5 \cdot 10^3 \frac{N \alpha F^2}{P' \lambda^2 \Delta^2 d'^2} + 3 \cdot 10^{-7} \frac{d'^2}{\mu'^2}$$

(α n'a pas de raison d'être différent pour le second système).

La déviation efficace d' optimum pour le niveau de signal 0 dBmo est :

$$d' = \sqrt[4]{2,16 \cdot 10^{10} \frac{N \alpha F^2 \mu'^2}{P' \lambda^2 \Delta^2}}$$

et la puissance psophométrique de bruit total en milliwatts au point de niveau relatif zéro pour la voie désavantagée sera :

$$2 \sqrt{A'B'} = 2 \sqrt{6,5 \cdot 10^3 \frac{N \alpha F^2}{P' \lambda^2 \Delta^2} 3 \cdot 10^{-7} \frac{1}{\mu'^2}}$$

ou :

$$2 \sqrt{A'B'} = \frac{0,088}{\lambda \mu'} \sqrt{\frac{N \alpha F}{P' \Delta}}$$

On voit que le rapport $\eta = \frac{2 \sqrt{AB}}{2 \sqrt{A'B'}}$ des puis-

sances de bruit total en simple et en double modulation, qui est le coefficient d'amélioration par le second système, s'écrit :

$$\eta = \frac{1}{\sqrt{2}} \frac{\mu'}{\mu} \frac{\Delta}{F} \sqrt{\frac{P'}{P}}$$

Si $\eta > 1$, le système à double modulation est le meilleur ; si $\eta < 1$ c'est le système à simple modulation.

Si la performance de l'un et de l'autre systèmes dépend bien évidemment de la constitution du circuit fictif de référence et du coefficient indéterminé α , le résultat de la comparaison, qui est la valeur de η , n'en dépend pas, au moins dans le cadre des hypothèses faites, notamment sur la petitesse de l'effet des distorsions de phase.

BIBLIOGRAPHIE

- BROCKBANK et WASS. — « Non-Linear Distortion in Transmitter Systems », *J.I.E.E.*, 1945, 92, partie III, p. 45.
- F. COETERIER. — « The Multireflection Tube. A New Oscillator for very short Waves », *Philips Technical Review*, septembre 1946.
- GERLACH. — « Microwave Relay Communication System », *R.C.A. Review*, décembre 1946.
- HOLBROOK et DIXON. — « Load Rating Theory for Multi-channel Amplifiers », *B.S.T.J.*, octobre 1939.
- J. J. LENEHAN. — « A Radio Relay System employing a 4 000 Mc/s 3-cavity Klystron », *W. U. T. R.*, juillet 1925.
- STARR et WALKER. — « Microwave Radio Links », *P.I.E.E.*, septembre 1952.
- L. E. THOMPSON. — « A Microwave Relay System », *P.I.R.E.*, 1946, 34, p. 936.

EFFET D'UN SIGNAL RADIOÉLECTRIQUE ISOLÉ SUR DES CIRCUITS RÉSONNANTS

PAR

Jean MARIQUE

Ingénieur A. I. Br. et Radio-ESE

*Secrétaire Général du Centre de Contrôle des Radiocommunications des Services Mobiles
(C. C. R. M.), Uccle-Bruxelles*

Dans un article récent [1], P. POINCELOT rappelle les difficultés qu'il y a à raccorder les notions de temps d'établissement d'un signal à celles de constante de temps et de largeur de bande d'un système. La lecture de cet article nous a incité à rassembler les résultats de calculs (dont certains remontent à plusieurs années) relatifs à l'effet d'un signal radio-électrique isolé sur des circuits résonnants, et à les présenter dans la note ci-dessous.

Nos calculs résultent de préoccupations un peu particulières qui concernent spécialement le fonctionnement des analyseurs de spectres radioélectriques et accessoirement l'étude des brouillages par des émissions manipulées. Ils n'ont donc pas un caractère aussi général que ce que l'on trouve dans des exposés classiques tels que ceux de POINCELOT lui-même [2] ou de BORG [3]. Mais par le fait même qu'ils concernent des cas spéciaux, ils complètent les résultats habituels, et nous croyons qu'ils méritent d'être mentionnés.

Le but spécial pour lequel nos calculs ont été entrepris nous a fait rechercher principalement les valeurs maxima de l'amplitude du courant dans le système : l'exposé ci-dessous se ressent évidemment de cette préoccupation.

I. — EFFET D'UN SIGNAL RADIOÉLECTRIQUE SUR UN CIRCUIT RÉSONNANT.

A. — Établissement du courant.

1. — Considérons un circuit résonnant $R L C$ série (figure 1) dont les caractéristiques sont :

pulsation propre : $\omega_1 = \sqrt{1/LC - (R/2L)^2}$

largeur de bande (à -3 db) : B c/s

amortissement : $\delta = R/2L = \pi B$

On admet que $\delta \ll \omega_1$.

Ce circuit étant au repos, appliquons-lui un signal radioélectrique

$$e(t) = E \exp(j\omega_0 t)$$

Appelons $\Delta = \omega_0 - \omega_1$ l'écart entre la pulsation ω_0 du signal et la pulsation propre du circuit ω_1 .

Posons enfin

$$z = \delta + j\Delta$$

On sait que l'on a très sensiblement :

$$2L i(t) = (E/z) (e^{j\omega_0 t} - e^{-\delta t + j\omega_1 t})$$

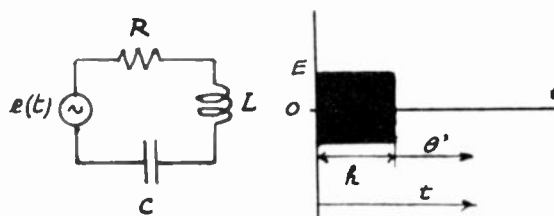


FIG. 1. — A gauche, circuit résonnant $R L C$. A droite, signal $e(t) = E \exp(j\omega_0 t)$ de durée h . Le temps θ' est compté à partir de la fin du signal et on a $t = h + \theta'$.

qu'on peut écrire :

$$2L i(t) = (E/z) e^{j\omega_0 t} (1 - e^{-\delta t - j\Delta t}) \quad (1)$$

L'amplitude instantanée $I(t)$ est donnée par le module de (1)

$$2LI(t) = (E/|z|) [1 + e^{-\delta t} (e^{-\delta t} - 2 \cos \Delta t)]^{1/2} \quad (2)$$

L'amplitude de régime I_r est donnée par :

$$2L I_r = E/|z|$$

2. — Dans le cas particulier où $\Delta = \omega_0 - \omega_1 = 0$ l'équation (2) donne :

$$2LI(t) = (E/\delta) (1 - e^{-\delta t}) \quad (4)$$

La valeur de régime n'est théoriquement jamais atteinte, et le temps nécessaire pour atteindre une amplitude donnée dépend uniquement de δ , donc de la largeur de bande du circuit. Le temps d'établissement du courant au sens du c. c. t. r. a une signification précise (1).

3. — Par contre quand $\Delta \neq 0$ il n'en est plus ainsi.

Pour toutes les valeurs de $\Delta t = (2k + 1)\pi$, l'équation (2) devient

$$2 LI(t) = (E/|z|) (1 + e^{-\delta t}) \quad (5)$$

et on vérifie que la plus grande amplitude I_m est obtenue pour $\Delta t_m = \pi$, c'est-à-dire pour $t_m = \pi/\Delta$.

Ce temps est indépendant de δ : il ne dépend que de Δ .

Ceci s'explique aisément si l'on se reporte à l'équation (1) dans laquelle le terme $\exp(-j\Delta t)$ révèle l'existence d'oscillations de battement entre les oscillations forcées et les oscillations libres (fig. 2).

Le temps nécessaire pour atteindre une fraction a de la valeur de régime (3) est donné par l'équation intrinsèque

$$1 + e^{-\delta t} (e^{-\delta t} - 2 \cos \Delta t) = a^2$$

qui montre que ce temps dépend non seulement de δ et de a , mais aussi de l'écart $\Delta = \omega_0 - \omega_1$. Le temps d'établissement du courant au sens du

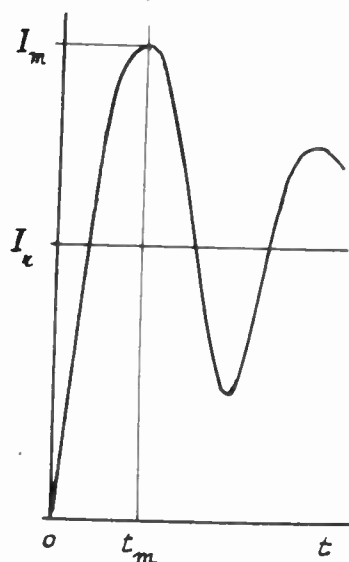


FIG. 2. — Variation de l'amplitude dans le circuit de la fig. 1 en fonction du temps pour $\Delta \neq 0$. Le premier maximum I_m est atteint au bout du temps $t_m = \pi/\Delta$.

(1) Le c.c.t.r. définit comme suit le temps d'établissement d'un signal : le temps pendant lequel le courant télégraphique passe du dixième aux neuf dixièmes (ou vice-versa) de la valeur qu'il atteint en régime établi (Avis N° 87 — Londres 1953).

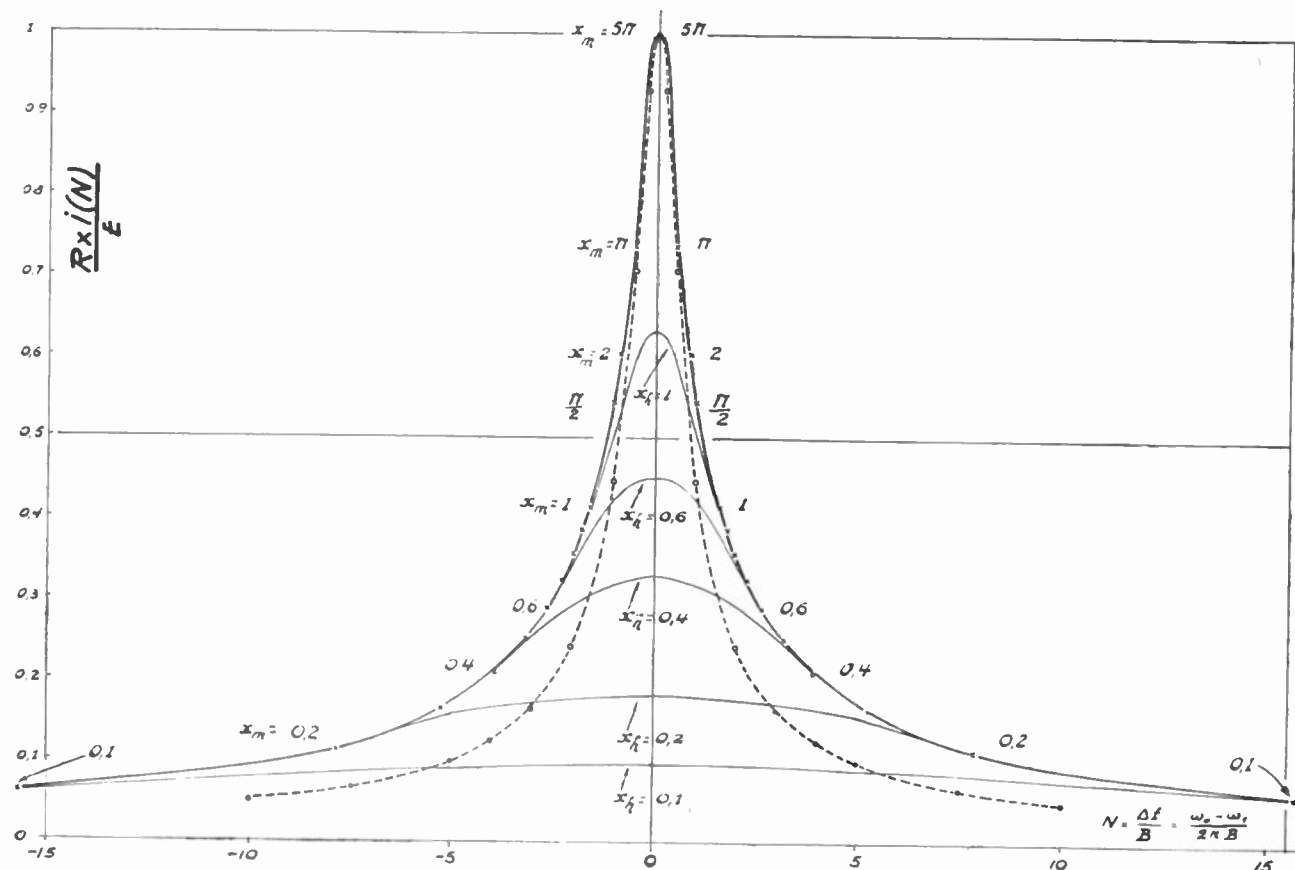


FIG. 3. — Variation de l'amplitude la plus élevée observée dans le circuit de la figure 1 en fonction de Δ en prenant $x_b = \delta b$ comme paramètre. On a $x_m = \delta t_m$. La courbe de sélectivité statique est représentée en trait interrompu.

c. c. r. ne dépend donc plus uniquement de la constante de temps du circuit.

B. — Cessation du courant.

4. — Supposons maintenant que le signal $E \exp(j\omega_0 t)$ cesse au temps $t = h$.

On a alors pour $t \geq h$

$$2Li(t) = (E/z) (-e^{-\delta t + j\omega_1 t} + e^{-\delta(t-h) + j[\omega_1(t-h) + \omega_0 h]}) \\ = (E/z) e^{j\omega_1 t} e^{-\delta t} (e^{\delta h + j\Delta h} - 1) \quad (5)$$

Il n'y a plus que des oscillations libres de pulsation ω_1 et la loi de décroissance dépend exclusivement de δ quel que soit l'écart Δ .

5. — Considérons ce qui se passe, lorsque la durée h du signal étant courte, on modifie l'écart Δ . Supposons que pour un certain écart Δ_m la durée h soit telle que le signal soit coupé au moment où l'amplitude atteint la valeur du premier maximum. On a alors, d'après le paragraphe 3

$$h = \pi / \Delta_m$$

Quand on fait varier l'écart Δ pour toutes les valeurs $\Delta \leq \Delta_m$ le premier maximum n'est jamais atteint, et c'est la durée h qui détermine entièrement le temps de croissance de l'amplitude.

Par contre, pour les écarts $\Delta \geq \Delta_m$, le premier maximum est toujours atteint. Ceci montre que, quand un signal est assez court, l'amplitude la plus grande qu'il provoque dans le circuit suit des lois différentes suivant que l'écart Δ est grand ou petit.

La figure 3 résume ces constatations. Pour établir les courbes, on a utilisé les paramètres réduits (sans dimension) suivants :

temps : $x = \delta t$

temps nécessaire pour atteindre le 1^{er} max. :
 $x_m = \delta t_m$

durée : $x_h = \delta h$

écart Δ : $N = \Delta / 2\pi B = (f_0 - f_1) / B$

La courbe en trait interrompu représente la courbe de sélectivité du circuit et correspond à l'état de régime. Les courbes en trait fin donnent, pour différentes valeurs de la durée réduite x_h du signal, la valeur de l'amplitude maximum du courant en fonction de l'écart réduit $N = \Delta / 2\pi B$. Enfin, la courbe en trait gras est l'enveloppe de toutes les courbes. On a indiqué, le long de cette enveloppe, les valeurs x_m du temps réduit nécessaire pour atteindre le premier maximum.

Comme l'enveloppe en trait gras est toute entière au-dessus de la courbe de sélectivité (trait interrompu), on en conclut que la sélectivité apparente du circuit en régime transitoire est moins bonne que la sélectivité en régime statique, ce qui est bien connu.

On remarque en passant qu'aucune des courbes x_h n'a l'allure en $\sin \alpha / \alpha$ qui serait celle du spectre d'un signal isolé.

II. — EFFET D'UN SIGNAL RADIOÉLECTRIQUE SUR UN SYSTÈME DE TROIS CIRCUITS RÉSONNANTS IDENTIQUES.

A. — Etablissement du courant.

6. — Pour différentes raisons, nous nous sommes intéressé depuis longtemps à un système sélectif comportant trois circuits résonnants identiques comme celui représenté à la figure 4. Chacun des trois circuits a les caractéristiques mentionnées au paragraphe 1.

Nous avons montré ([4] annexe 1) que, moyennant certaines simplifications permises en haute fréquence, et en particulier, quand les signaux ont la forme $E \exp(j\omega_0 t)$, on peut prendre pour l'admit-

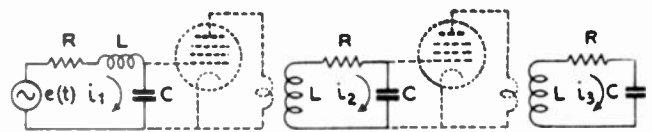


FIG. 4. — Schéma du système à trois circuits résonnants identiques.

tance impulsionnelle de ce système la valeur approchée

$$(1/16 L_3) t^2 \exp(-\delta t + j\omega_1 t)$$

pour l'établissement de laquelle on a posé égal à l'unité le facteur constant qui tient compte de l'effet des tubes et des couplages (1).

Quand on applique à un tel système supposé au repos un signal $e(t) = E \exp(j\omega_0 t)$ au temps $t = 0$, le courant dans le 3^e circuit est donné par

$$16 L^3 i(t) = E \int_0^t e^{j\omega_0 T} (t - T)^2 e^{(-\delta + j\omega_1)(t-T)} dT \quad (7)$$

et on peut écrire

$$16 L^3 i(t) = (E/z) [2e^{j\Delta t}/z^2 - e^{-\delta t} (t^2 + 2t/z + 2/z^2)] \quad (8)$$

On voit encore ici qu'il y a des oscillations de battement (terme $\exp(j\Delta t)$).

Lorsque l'état de régime est atteint, il vient

$$16 L^3 i(t) = 2 E e^{j\omega_1 t} / z^2 \quad (9)$$

expression qui ne contient que des oscillations forcées.

7. — Dans le cas particulier où $\omega_0 = \omega_1$ l'amplitude I du courant est donnée

$$16 L^3 I(t) = (E/\delta) [2/\delta^2 - e^{-\delta t} (t^2 + 2t/\delta + 2/\delta^2)]$$

et en remarquant que $\delta = R/2L$ il vient

$$R^3 I(t) = E [1 - e^{-t\delta} (\delta^2 t^2 / 2 + \delta t + 1)]$$

(1) REMARQUE IMPORTANTE. — Ce facteur, posé ici égal à l'unité pour simplifier les écritures, a en réalité les dimensions du carré d'une résistance pour le système considéré. Par conséquent, malgré les apparences, les deux membres des équations sont bien homogènes.

La courbe représentant I en fonction de t est tangente à l'axe des t à l'origine et croît vers la valeur de régime

$$R^3 I_r = E$$

qu'elle n'atteint théoriquement jamais.

Le temps nécessaire pour atteindre une amplitude donnée dépend uniquement de δ , comme dans le cas d'un seul circuit, et le temps d'établissement au sens du c. c. i. r. a une signification précise.

8. — Quand l'écart $\Delta = \omega_0 - \omega_1$ est suffisamment grand, les termes en $1/z$ et $1/z^2$ dans (8) peuvent être négligés devant le terme $t^2 \exp(-\delta t)$, tout au moins au voisinage de son maximum : les oscillations de battement, en particulier, y ont un effet négligeable, et on peut écrire pour l'amplitude $I(t)$

$$16 L^3 I(t) \approx (E/|z|) t^2 e^{-\delta t}$$

expression qu'on peut transformer en

$$R^3 I(t) = \frac{1}{2} E (\delta/|z|) \delta^2 t^2 e^{-\delta t} \quad (10)$$

Le maximum I_M de I a lieu pour $\delta t = 2$ et est donné par

$$R^3 I_M = 0,27 E (\delta/|z|) = 0,27 E / (1 + \Delta^2/\delta^2)^{1/2} \quad (11)$$

L'amplitude de régime tirée de (9) est donnée par

$$R^3 I_r = E (\delta/|z|)^3 = E / (1 + \Delta^2/\delta^2)^{3/2} \quad (12)$$

En rapprochant cette valeur de l'amplitude maximum tirée de (11) on a

$$I_M/I_r = 0,27 (1 + \Delta^2/\delta^2) \quad (13)$$

On voit que le rapport I_M/I_r pour le système considéré, est d'autant plus grand que Δ est plus grand. L'amplitude maximum peut donc être beaucoup plus grande que l'amplitude de régime quand l'écart Δ est grand, ce qui représente une diminution considérable de la sélectivité.

B. — Cessation du courant.

9. — Considérons maintenant que le signal $E \exp(j\omega t)$ cesse au temps $t = h$.

D'après ce qui précède, quand $\Delta = 0$ et si $\delta h \geq 2$, l'amplitude passe par un maximum au cours de la durée h du signal. Par contre si $\delta h \leq 2$ le signal est coupé avant que le courant ait atteint sa valeur maximum. Voyons donc ce qui se passe après la fin du signal.

Nous avons montré ([5] équation (27)) que l'amplitude du courant dans le 3^e circuit à partir de la fin du signal (1) est donnée par l'équation

$$\frac{16 L^3 I'}{h^3} = \left| \frac{E e^{-\delta \theta'}}{zh} \right\{ \mathcal{C} + 2\mathfrak{B}/zh + 2\mathfrak{A}/z^2 h^2 + \mathfrak{A} (\theta'/h)^2 + 2\mathfrak{B} (\theta'/h) + 2\mathfrak{A} (\theta'/h)/zh \} \quad (14)$$

où θ' représente le temps compté à partir du moment où le signal cesse, c'est-à-dire que $t = h + \theta'$ (fig. 1).

Dans le cas d'un signal unique de durée h , on a

$$\mathfrak{A} = 1 - e^{-\delta h}; \quad \mathfrak{B} = \mathcal{C} = -e^{-\delta h}$$

Les termes $\mathfrak{A}\mathfrak{B}\mathcal{C}$ sont réels quand Δh est un multiple entier de π , et on a

$$e^{-\delta h} = e^{-\delta h} \text{ quand } \Delta h = 2\pi m$$

$$e^{-\delta h} = -e^{-\delta h} \text{ quand } \Delta h = (2k+1)\pi$$

Considérons, comme nous l'avons fait au paragraphe 8, des écarts $\Delta = \omega_0 - \omega_1$ suffisamment grands pour que les termes en $1/z$ et $1/z^2$ soient négligeables tout au moins au voisinage du maximum des termes restants. Une valeur approchée de l'amplitude pour Δh multiple entier de π est alors donnée par

$$16 L^3 I'/h^3 = (E/|zh|) e^{-\delta \theta'} \{ \mathcal{C} + 2\mathfrak{B} (\theta'/h) + \mathfrak{A} (\theta'/h)^2 \} \quad (15)$$

Pour simplifier les écritures, nous posons $x = \delta h$ (pas nécessairement entier) et nous écrivons z au lieu de $|z|$, aucune confusion n'étant possible puisqu'il s'agit de modules.

10. — 1^{er} CAS : $\Delta h = 2\pi m$ (Δ grand)

On vérifie que le terme entre accolades dans (15) s'annule pour

$$\theta'/h = 1/(e^{1/2} - 1) \quad (16)$$

L'amplitude maximum I'_M après le signal est atteinte pour

$$\theta'_M/h = \frac{1 + (x-1)e^{-x} \pm [(1-e^{-x})^2 + x^2 e^{-x}]^{1/2}}{x(1-e^{-x})} \quad (17)$$

Comme θ' doit nécessairement être positif, seules les valeurs de $x = \delta h$ qui donnent des valeurs positives de θ'_M/h doivent être retenues.

Quand $x = \delta h > 2$, il n'y a qu'un maximum après la fin du signal, tandis que si $x = \delta h < 2$, il y en a deux.

Quand le signal est long (δh grand), on vérifie que

$$I'_M < I_M$$

$$\text{et} \quad \theta'_M > 2/\delta$$

mais au fur et à mesure que δh croît

$$I'_M \rightarrow I_M$$

$$\theta'_M \rightarrow 2/\delta$$

(1) Nous utilisons le signe (*) pour indiquer qu'il s'agit de valeurs rencontrées après la fin du signal.

Les valeurs limites sont déjà pratiquement atteintes pour $\delta h \geq 6$ (voir fig. 5).

Les variations de l'amplitude dans le cas où $\Delta h = 2\pi m$ (Δ grand) sont représentées à la figure 6 (c) et (d) pour $\delta h = 2$ et $\delta h = 6$. On rappelle que

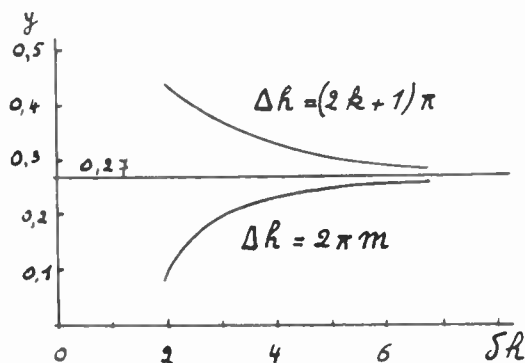


FIG. 5. — Valeur de l'amplitude maximum I'_M du courant dans le 3^e circuit de la figure 4 observée après la fin du signal, en fonction de δh , quand Δ est grand et Δh multiple entier de π . On a $R^3 I'_M = y E \delta / |z|$. L'ordonnée $y = 0,27$ correspond au maximum observé au temps $t = 2/\delta$.

ces amplitudes sont obtenues à partir des équations (10) et (15) et ne tiennent donc pas compte des termes en $1/z$ et $1/z^2$.

11. — 2^e CAS : $\Delta h = (2k + 1)\pi$ (Δ grand).

Le terme entre accolades de (15) ne s'annule plus pour aucune valeur finie de θ' , mais la dérivée de l'amplitude en fonction de x s'annule pour

$$\theta'_M/h = \frac{1 - (x-1)e^{-x} \pm [(1+e^{-x})^2 - x^2 e^{-x}]^{1/2}}{x(1+e^{-x})} \quad (18)$$

Quand $x = \delta h > 2$ la plus petite des valeurs de θ'_M correspond à un minimum d'amplitude, la plus grande à un maximum. Quand $x = \delta h < 2$ il n'y a plus qu'un maximum correspondant au signe + dans (18).

Quand le signal est long (δh grand) on vérifie que

$$I'_M > I_M$$

$$\text{et} \quad \theta'_M < 2/\delta$$

mais au fur et à mesure que δh croît

$$I'_M \rightarrow I_M$$

$$\theta'_M \rightarrow 2/\delta$$

Les valeurs limites sont pratiquement atteintes quand $\delta h \geq 6$ (voir fig. 5).

Les variations de l'amplitude dans le cas où $\Delta h = (2k + 1)\pi$ (Δ grand) sont représentées à la figure 6 (a) et (b) pour $\delta h = 2$ et $\delta h = 6$.

12. — On voit que quand δh est suffisamment grand, les variations de l'amplitude de courant après la fin du signal sont pratiquement identiques, que Δh soit un multiple entier pair ou impair de π .

On peut conclure des paragraphes 8, 10 et 11

que quand un signal assez long est appliqué à l'entrée d'un système tel que celui de la figure 4 et quand l'écart Δ est assez grand, un maximum d'amplitude I_M a lieu $2/\delta$ seconde après le début du signal et un second maximum I'_M égal au premier a lieu $2/\delta$ seconde après la fin du signal. Comme d'après (13), l'amplitude de régime pendant le signal peut être beaucoup plus petite que I_M , il est évident que le temps d'établissement au sens du c. c. i. r. n'a pas d'intérêt pratique.

Donnons un exemple numérique. Supposons que le système comporte trois circuits de largeur de bande $B = 200$ c/s ce qui donne $\delta = 200\pi$. On note en passant que la largeur de bande résultante du système est d'environ 100 c/s. Supposons que le système soit accordé à 1 kc/s de la fréquence du signal, ce qui donne $\Delta = 2000\pi$. On a alors $\Delta/\delta = 10$ et $\delta/z = 1/10$.

Prenons comme durée h du signal, la durée d'un point télégraphique transmis à la vitesse de 25 bauds. On a $h = 1/25$ et $\delta h = 200\pi/25 = 25$.

Dans ces conditions, l'amplitude de régime est donnée par

$$R^3 I_r = E (\delta/z)^2 = 0,001 E$$

tandis que l'amplitude des deux maxima est donnée par

$$R^3 I_M = R^3 I'_M = 0,27 E (\delta/z) = 0,027 E$$

L'amplitude de régime est donc 27 fois plus petite que l'amplitude maximum. Cet exemple fait clairement apparaître l'importante réduction de sélectivité qu'entraîne le régime transitoire.

13. — Comparons l'amplitude maximum I_M à l'amplitude maximum obtenue en réponse à une impulsion de tension $E_0 \delta(t)$. Pour la calculer, il faut prendre comme admittance impulsionnelle la valeur approchée :

$$(1/16 L^3) t^2 (e^{-\delta t + j\omega_1 t} + e^{-\delta t - j\omega_1 t})$$

qui diffère de celle employée au paragraphe 6 par l'adjonction d'une seconde exponentielle qui, quand les signaux ont la forme $E \exp. j\omega_0 t$ donne lieu à une réponse négligeable ([4]-annexe 1), mais qui ne peut pas être supprimée cette fois.

On obtient pour le courant $i''(t)$ dû à l'impulsion $E_0 \delta(t)$

$$16 L^3 i''(t) = 2 E_0 t^2 e^{-\delta t} \cos \omega_1 t$$

L'amplitude $I''(t)$ du courant est alors donnée par

$$R^3 I''(t) = \delta^2 t^2 e^{-\delta t} (\delta E_0) \quad (19)$$

L'amplitude maximum I''_M est obtenue pour $\delta t = 2$ et vaut

$$R^3 I''_M = 0,54 \delta E_0 \quad (20)$$

On rappelle que le second membre doit être multiplié par un facteur posé égal à l'unité (voir paragraphe 6), mais dont les dimensions sont celles du carré d'une résistance. Les dimensions du produit δE_0 sont

celles d'une tension, les dimensions de E_0 étant celles d'une tension multipliée par un temps.

En rapprochant (19) de (10) et (20) de (11) on voit que si $E_0 = E/2z$, les expressions sont identiques. On peut donc dire que, pour le système considéré, quand le signal est assez long (δh grand) et pour de grandes valeurs de l'écart Δ , les variations d'amplitude du courant sont les mêmes que

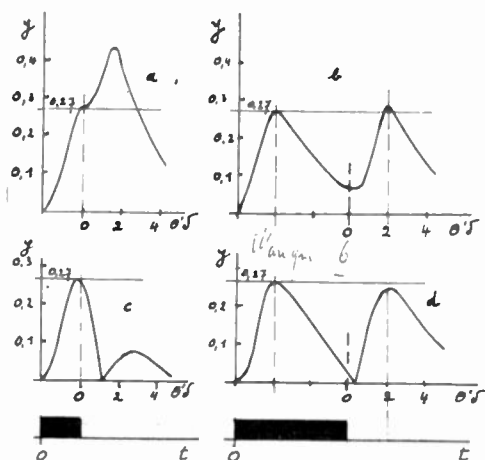


FIG. 6. — Variation de l'amplitude du courant dans le 3^e circuit en fonction du temps pour des signaux longs ($\delta b \geq 2$), quand Δ est grand, avec les conditions suivantes :

- (a) $\delta b = 2$ et $\Delta b = (2k + 1)\pi$; (b) $\delta b = 6$ et $\Delta b = (2k + 1)\pi$;
(c) $\delta b = 2$ et $\Delta b = 2\pi m$; (d) $\delta b = 6$ et $\Delta b = 2\pi m$.
On a $R^3 I = y E \delta / |z|$.

si le début du signal et la fin du signal étaient remplacés par des impulsions de tension ($E/2z$) $\delta(t)$ ayant lieu aux temps $t = 0$ et $t = h$.

14. — CAS D'UN SIGNAL COURT ($\delta h < 2$).

Dans ce cas, le signal est supprimé avant que l'amplitude ait atteint la valeur I_m , de sorte que le maximum est toujours observé après la fin du signal.

Les conditions (16) (17) (18) sont toujours d'application.

Pour de grandes valeurs de Δ , si $\Delta h = (2k + 1)\pi$ l'expression (15) ne donne qu'un maximum, tandis que pour $\Delta h = 2\pi m$ il y en a deux, dont le plus rapproché de la fin du signal est le plus élevé.

L'amplitude du courant varie comme l'indiquent les courbes de la figure 7 calculées pour $\delta h = 1$ et $\delta h = \pi/10$.

La figure 8 représente en fonction du désaccord réduit $\Delta h = (\omega_0 - \omega_1)h$ les amplitudes maxima correspondantes qui accompagnent un signal ($E \exp(j\omega_0 t)$) de durée telle que $\delta h = 1$ ou $\delta h = \pi/10$. Les calculs ont été faits à partir de l'équation complète (14).

On n'a représenté que les courbes relatives à $\Delta > 0$. Celles relatives à $\Delta < 0$ sont évidemment symétriques.

On a également représenté en trait interrompu le contour du spectre d'un tel signal (courbe en $\sin \alpha / \alpha$) et on voit que plus le signal est court, plus le lieu des maxima d'amplitude ressemble au contour du spectre. La figure 8 est à rapprocher de la figure 3 établie pour un seul circuit : en particulier la courbe

$x_h = 1$ de la figure 3 correspond à la courbe $\delta h = 1$ de la figure 8. Pour ces deux courbes, les abscisses

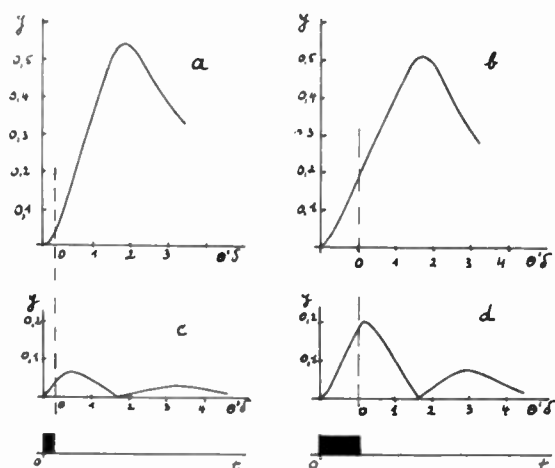


FIG. 7. — Variation de l'amplitude du courant de 3^e circuit en fonction du temps pour des signaux courts ($\delta b < 2$), quand Δ est grand, avec les conditions suivantes :
(a) $\delta b = \pi/10$ et $\Delta b = (2k + 1)\pi$; (b) $\delta b = 1$ et $\Delta b = (2k + 1)\pi$;
(c) $\delta b = \pi/10$ et $\Delta b = 2\pi m$; (d) $\delta b = 1$ et $\Delta b = 2\pi m$. On a $R^3 I = y E \delta / |z|$.

sont liées entre elles par la relation $N = 1/2 \Delta h$, c'est-à-dire que $N = 15,7$ correspond à $\Delta h = 10\pi$.

15. — CONCLUSIONS.

Le temps d'établissement (au sens du C. C. I. R.) du courant produit par un signal $E \exp(j\omega_0 t)$ dans un système composé de circuits résonnants ne se

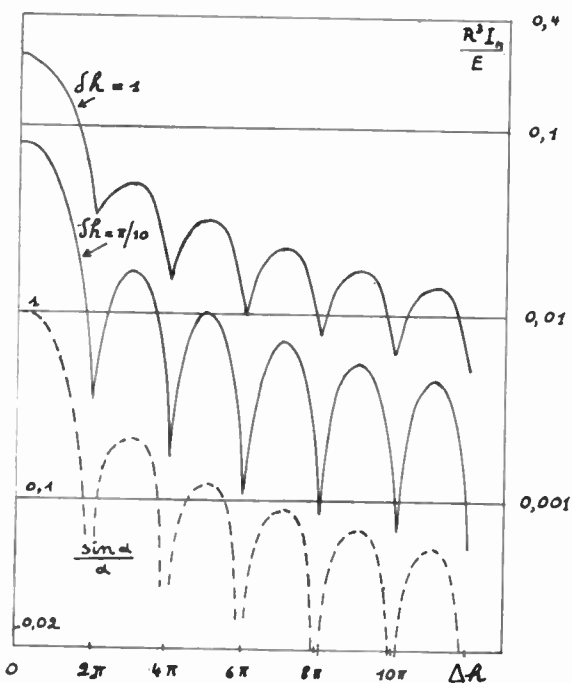


FIG. 8. — Signaux courts. Amplitudes maxima observées dans le 3^e circuit en fonction de l'écart réduit Δh , pour $\delta b = 1$ et $\delta b = \pi/10$ (échelle de droite). La courbe en trait interrompu (échelle de gauche) représente la forme du spectre continu d'un signal isolé (courbe en $\sin \alpha / \alpha$).

raccorde facilement à la constante de temps des circuits que dans le cas particulier où le signal a la même fréquence que la fréquence de résonance des circuits.

Par contre, dès que la fréquence du signal et celle du système diffèrent, ce raccord n'est plus possible sous une forme simple même quand il s'agit d'un simple circuit résonnant. En particulier, dans un système amplificateur tel que celui de la figure 4, le début et la fin d'un signal long tel que $\delta h \geq 6$ sont accompagnés de phénomènes de caractère impulsionnel dont les amplitudes sont d'autant plus grandes par rapport à l'amplitude de régime que l'écart entre les fréquences est grand, ce qui rend impossible toute relation simple entre le temps d'établissement au sens du C. C. I. R. et les caractéristiques des circuits.

Dans le cas de signaux courts tels que $\delta h \leq 2$, le lieu des maxima d'amplitude du courant de sortie du système de la figure 4 en fonction de l'écart entre les fréquences prend d'autant plus la forme en $\sin \alpha / \alpha$ du spectre du signal original, que δh est plus petit.

Notes complémentaires.

1. — Depuis la rédaction de l'étude ci-dessus, nous avons eu connaissance d'un intéressant article de D. G. TUCKER [6] qui couvre partiellement l'objet de la présente note.

TUCKER étudie particulièrement l'établissement du courant dans une suite de circuits résonnants disposés en cascade, mais il ne semble pas s'être préoccupé des phénomènes postérieurs à la fin du signal sauf dans le cas particulier où $\Delta = 0$ (d'après nos notations). Cependant, son article est illustré de deux beaux oscillogrammes expérimentaux qui montrent nettement l'existence du second maximum I'_M après la fin du signal quand $\Delta \neq 0$ (Les données de l'expérience sont, ainsi qu'on peut les reconstituer : $h = 50/1\ 000$; $\Delta h = \pm 30\pi$; $\delta = 140$; $\delta h = 7$ ce qui correspond approximativement à notre limite $\delta h \geq 6$).

On se convaincra facilement que notre étude, dans laquelle nous examinons les phénomènes postérieurs à la fin du signal, ainsi que les influences respectives de la durée du signal et de l'écart Δ , ne fait aucunement double emploi avec celle de TUCKER.

2. — Nous voudrions ajouter les considérations suivantes au sujet des signaux « longs » ($\delta h > 6$). Nous avons montré dans un article antérieur (voir [4] paragraphe 21) que, quand un système de circuits résonnants connectés en cascade est destiné à la réception de signaux radiotélégraphiques, la constante de temps est nécessairement telle que chaque signal peut être considéré comme agissant seul. Dans ces conditions, chaque signal d'un message est accompagné des deux maxima I_M et I'_M . Les expériences auxquelles nous avons procédé sur les étages HF de récepteurs industriels comportant 3 circuits accordés (deux tubes d'amplification HF avant le tube mélangeur) confirment pleinement l'existence de ces deux « impulsions » qui marquent respectivement le début et la fin de chaque signal ($\Delta \neq 0$).

On ne saurait donc assez insister sur ces phénomènes qui ont une grande importance pour l'étude des brouillages manipulés.

BIBLIOGRAPHIE

1. — P. POINCELOT. — Sur les relations entre la largeur de bande d'un circuit et sa constante de temps. *Annales des Télécommunications*. — Tome 9. — N° 6 (juin 1954), pp. 191-192.
2. — P. POINCELOT. — Les régimes transitoires dans les réseaux électriques. (livre) Gauthier-Villars, Paris 1953.
3. — H. BORG. — Théorie des circuits impulsionnels (livre). — *Éditions de la Revue d'Optique*. — Paris 1953.
4. — J. MARIQUE. — Actions de signaux télégraphiques périodiques sur des circuits résonnants RLC en cascade. *Revue HF*. Vol. II, N° 9, pp. 233-244 (Bruxelles 1954).
5. — J. MARIQUE. — Réponse des analyseurs de spectres radioélectriques à des signaux Morse non périodiques (*Annales des Télécommunications*). Tome 9, N° 7-8 (juillet-août 1954) pp. 215-223 et N° 9 (septembre 1954) pp. 247-255.
6. D. G. TUCKER. — Transient response of tuned circuit cascades. Vol. 23, N° 276, sept. 1946, pp. 250-258.

LE PLEURAGE DANS LES SYSTÈMES DE REPRODUCTION SONORE

PAR

Ing. Cesarina Bordone SACERDOTE
(Istituto Elettrotecnico Nazionale-Torino)

ET

Ing. Mario CACIOTTI
(Radiotelevisione Italiana)

ET

Prof. Gino SACERDOTE
(Istituto Elettrotecnico Nazionale-Torino)

Dans tous les systèmes d'enregistrement il y a un support sonore en mouvement relatif par rapport à un organe de réception qui peut être le pick-up, la tête magnétique, la cellule photoélectrique, etc...

Portons notre attention sur le cas de l'enregistrement magnétique : supposons que sur la bande soit enregistré un signal de fréquence parfaitement constante : on devrait constater que les différents maxima d'enregistrement sont strictement équidistants.

Supposons encore que nous utilisions la bande en question en la reproduisant dans un magnétophone : on observera alors, en général, que la fréquence n'est plus parfaitement constante, mais qu'elle subit des variations qui sont essentiellement de deux genres : une variation à caractère plutôt irrégulier disparaissant lentement, à laquelle vient se superposer une variation périodique. On peut dire que, en général, on distingue trois genres de variations de fréquence : une variation apériodique, c'est-à-dire une dérive, et des variations remarquablement périodiques ; ces dernières, quant à l'impression produite à l'oreille, se distinguent encore en deux catégories : si la fréquence de ce phénomène périodique est à peu près d'un Hertz on a le pleurage ; si au contraire cette fréquence dépasse les 10 Hz on a le phénomène du scintillement. Pour donner un exemple concret de ces variations dans la reproduction par disques, on peut avoir une variation régulièrement croissante ou décroissante de la fréquence, depuis le commencement jusqu'à la fin du disque, due à la variation du rayon de courbure du sillon, qui produit un couple variable à mesure que le lecteur avance vers le centre : ce couple exerce une influence sur la vitesse du moteur si celui-ci n'est pas assez puissant.

Si le disque n'est pas parfaitement centré, on peut avoir un pleurage dont la période est égale à la durée d'un tour.

Si enfin le moteur a de légers mouvements pendulaires, ces mouvements se reflètent à travers le système mécanique de démultiplication, et on a le scintillement.

Il faut cependant remarquer que ces distinctions sont surtout utiles dans un but de classification, et il peut se présenter le cas, particulièrement pour l'enregistrement magnétique, où l'on ne peut définir avec précision s'il s'agit d'une variation lente périodique, d'un pleurage, ou d'un scintillement, et dans certains cas des trois phénomènes en même temps.

Les paramètres qui définissent le pleurage sont : la fréquence et l'amplitude de modulation. La fréquence de modulation, comme nous venons de le dire, peut être un élément utile pour la définition de plusieurs genres de défauts. L'amplitude de modulation Δf est exprimée en général en pourcentage $\Delta f/f$. Δf est la variation maximum, c'est-à-dire $f_{\max} - f_{\min}$ (peak to peak). On doit remarquer que dans une bande enregistrée et reproduite par un système qui présente ces inconvénients, toutes les fréquences varient du même pourcentage, et on peut parler d'un pourcentage de variation lorsque la fréquence de modulation est beaucoup plus basse que la fréquence émise.

Les problèmes qui se posent dans l'étude de cette question sont essentiellement la mesure de la fréquence et de l'amplitude pour-cent de modulation du pleurage, l'examen des différentes causes qui le produisent, la détermination des limites de tolérance, qui dépendent des propriétés psycho-physiologiques de l'oreille.

On analyse le pleurage et le scintillement comme les ondes modulées en fréquence : pour le pleurage il se présente le cas classique d'une porteuse à fréquence élevée par rapport à la fréquence de modulation ; pour le scintillement au contraire, la fréquence porteuse et la fréquence de modulation peuvent avoir la même valeur. Dans les deux cas la variation pour-cent de fréquence est relativement petite (tout au plus quelques unités pour-cent).

L'expression analytique d'une grandeur U modulée en fréquence est :

$$U = A \sin(2\pi ft + m \sin 2\pi Ft)$$

où le facteur de modulation m est donné, en fonction

de la variation de fréquence, et des fréquences F (de modulation) et f (porteuse), par la relation :

$$m = \frac{1}{2} \frac{\Delta f}{f} \frac{f}{F}$$

En cas de *pleurage* ayant une amplitude de 1 % et une fréquence F d'environ 3 Hz, on a pour une porteuse de 1 000 Hz un facteur de modulation m d'environ 1,66.

Le spectre de fréquence montre que l'intensité de la porteuse f est affaiblie, tandis qu'à ses côtés apparaissent quelques composantes dont l'amplitude est du même ordre que celle de la porteuse, et qui sont séparées par un intervalle égal à la fréquence F de modulation. En cas de *scintillement* ayant encore une amplitude de 1 % mais une fréquence de 2 000 Hz, (relevée par WERNER comme fréquence propre des bandes magnétiques) le facteur de modulation m prend, pour une porteuse de 1 000 Hz, la valeur de 0,002 5. Dans le spectre de fréquence on remarque alors que la porteuse a une intensité beaucoup plus grande que celle des composantes, lesquelles sont distantes d'environ 2 000 Hz, et dont seules les deux premières ont une intensité non négligeable.

MÉTHODES DE MESURE

Pour mesurer le pleurage on enregistre sur le support sonore une fréquence pure d'amplitude constante (en général une fréquence élevée) et on mesure les petites différences à l'aide de fréquence-mètres électroniques spéciaux.

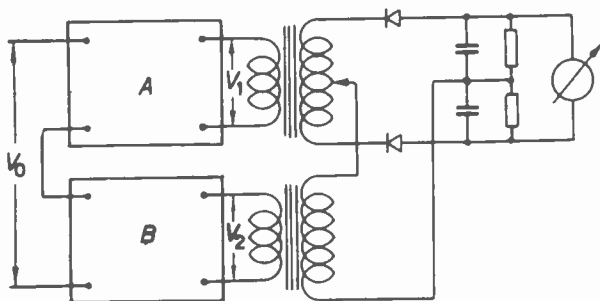


FIG. 1. — Circuit schématique d'un appareil pour la mesure du pleurage.

On peut effectuer cette mesure soit par une méthode de battements, soit par des circuits très semblables aux discriminateurs des appareils radio-récepteurs à modulation de fréquence.

Un autre type de « wowmeter » est basé sur le principe suivant : on rapporte la tension à une valeur constante et on l'envoie à deux circuits différents, qui présentent des caractéristiques spéciales d'impédance, c'est-à-dire dont l'amplitude et la phase varient en fonction de la fréquence. Dans la figure 1 A et B représentent les quadripôles types au moyen desquels on obtient ces variations. On envoie les sorties des deux quadripôles sur un pont

déphaseur, dont la tension de sortie est continue et proportionnelle à $V_1 V_2 \cos \varphi$ où V_1 et V_2 représentent les deux tensions d'alimentation du pont, déphasées entre elles de φ .

La tension continue de sortie peut indiquer directement les variations de fréquence avec une sensibilité remarquable ; on peut par exemple, lire ou enregistrer un pleurage de 1 %.

La tension continue de sortie est fonction des caractéristiques des deux quadripôles A et B, exprimées respectivement par les matrices mixtes :

$$\begin{vmatrix} A \\ B \end{vmatrix} = \begin{vmatrix} \alpha & Z \\ Y & \beta \end{vmatrix} \quad \begin{vmatrix} B \\ A \end{vmatrix} = \begin{vmatrix} \alpha' & Z' \\ Y' & \beta' \end{vmatrix}$$

Avec les notions de la figure 1 on a :

$$V_1 = \frac{Y'}{\alpha Y' + \alpha' Y} V_0 \quad V_2 = \frac{Y}{\alpha Y' + \beta' Y} V_0$$

où l'on suppose que les deux quadripôles travaillent à circuit ouvert.

Par conséquent la tension continue à la sortie du pont-déphaseur est proportionnelle à :

$$V_1 V_2 \cos \varphi = V_0^2 \frac{Y \times Y'}{|\alpha Y' + \alpha' Y|^2}$$

Si les deux quadripôles sont des simples admittances Y et Y' , alors $\alpha = \alpha' = 1$ et la tension continue de sortie est proportionnelle à :

$$V_0^2 = \frac{Y \times Y'}{|Y + Y'|^2}$$

Sur ce principe on a construit des appareils où les deux quadripôles sont respectivement : un circuit antirésonant sur la fréquence qui doit être mesurée, et un circuit de déphasage à résistance et capacité, en général composé de deux mailles.

Une particularité qui peut distinguer les différents circuits entre eux est la possibilité éventuelle d'être adaptés à la mesure sur différentes fréquences. Certains types d'appareils sous forme industrielle peuvent mesurer le pleurage seulement pour une fréquence donnée d'enregistrement. Au contraire, avec des quadripôles constitués de deux circuits antirésonants, on peut aisément régler l'appareil pour effectuer la mesure sur la fréquence voulue.

Tandis que pour l'exercice une mesure globale du pleurage suffit pour le contrôle des caractéristiques des enregistreurs, pour la recherche des causes du pleurage il est très important de connaître la fréquence de modulation et l'existence de ses variations éventuelles.

Il est donc utile de se servir d'un enregistreur, puisque la simple mesure des variations pour-cent de fréquence ne suffit point. On a effectué, avec les appareils du type mentionné, de nombreux enregistrements moyennant un enregistreur rapide à échelle linéaire.

Voici quelques exemples de ces enregistrements (fig. 2) :

a) Enregistrement du pleurage sur un magnétophone professionnel dans les conditions normales.

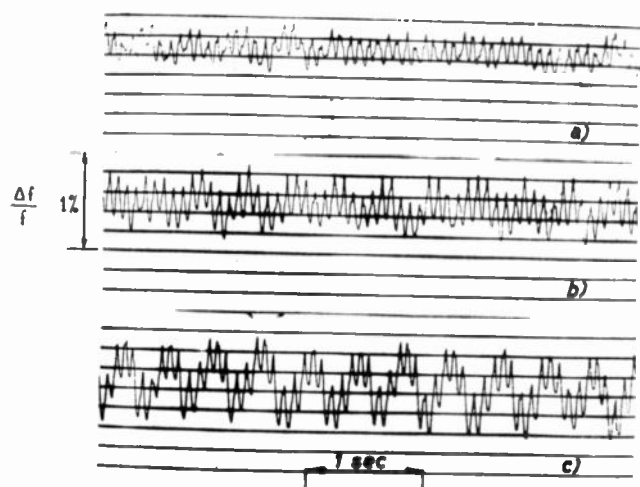


FIG. 2. — Enregistrement du pleurage d'un magnétophone professionnel : a) en conditions de fonctionnement normal ; b) à volant bloqué ; c) sans volant.

b) Pleurage de la machine à volant immobile ;

c) Enregistrement du même genre exécuté, le volant de l'appareil examiné étant supprimé.

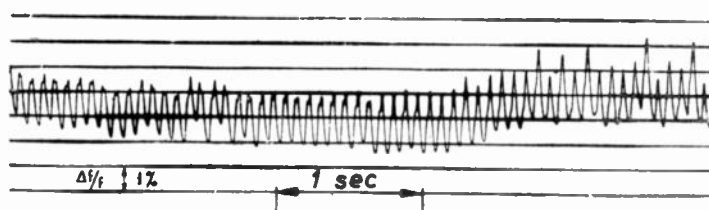


FIG. 3. — Enregistrement du pleurage d'un magnétophone dont le pivot présente une certaine excentricité.

La figure 3 représente le pleurage (enregistré) sur un appareil dont le pivot présente une excentricité notable.

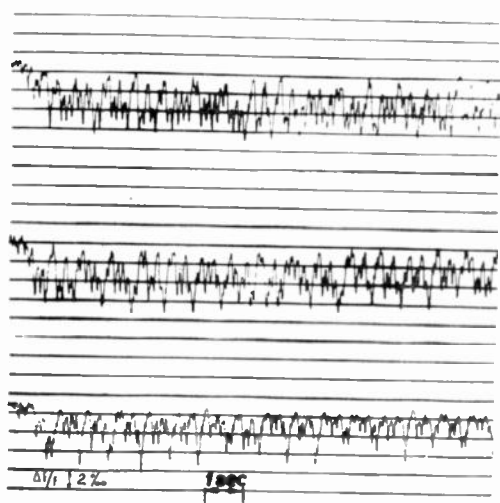


FIG. 4. — Enregistrements successifs du pleurage de la même bande fermée.

Pour examiner l'existence éventuelle d'oscillations se produisant au hasard, on a répété plusieurs fois des enregistrements sur bande continue : la figure 4 montre trois de ces enregistrements.

Pour un examen plus détaillé il peut être nécessaire de se servir d'enregistrements photographiques à l'oscillographe cathodique (fig. 5).

On peut avoir une analyse détaillée des petites variations de fréquence par un système dérivé de celui réalisé par GRÜTZMACHER pour l'enregistrement des courbes mélodiques de la voix. Le principe sur lequel l'appareil se fonde est le suivant (fig. 6) :

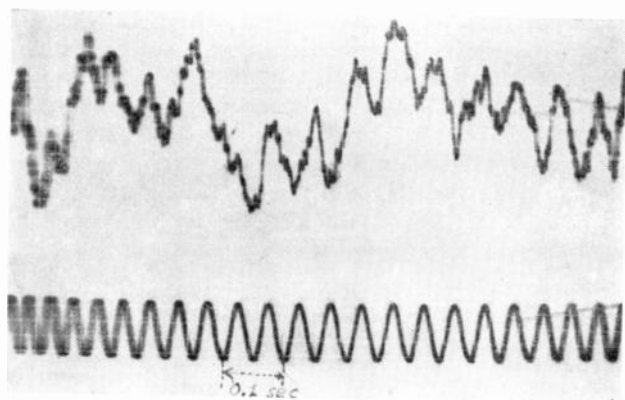


FIG. 5. — Enregistrement oscillographique du pleurage.

la fréquence à examiner passe dans un limiteur (en cas d'enregistrement de fréquences pures, celui-ci est assez simple, puisqu'il s'agit de limiter un signal d'amplitude à peu près constante) d'où elle sort sous forme d'onde carrée.

Le stade suivant est un intégrateur à la sortie duquel on obtient une onde à dent de scie : l'inter-

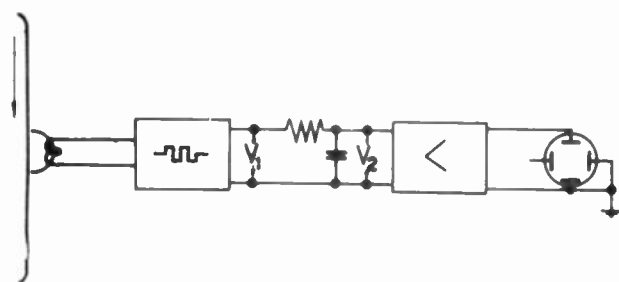


FIG. 6. — Circuit pour la détermination de la valeur instantanée du pleurage.

valle entre deux maxima successifs correspond à la période du signal examiné, l'amplitude de chaque maximum est inversement proportionnelle à la fréquence enregistrée. L'allure des maxima successifs suit, avec une remarquable sensibilité, l'allure de la fréquence à chaque période ; on photographie à l'oscillographe cathodique un nombre même grand, de périodes (fig. 7).

Ce dispositif a l'avantage de permettre la mesure sur une fréquence quelconque.

LES PROPRIÉTÉS MÉCANIQUES DES RUBANS

La reproduction magnétique peut être altérée à cause de l'élasticité du ruban, qui peut provoquer des vibrations pendant son déroulement. Ces vibrations donnent lieu à un scintillement dont la fréquence

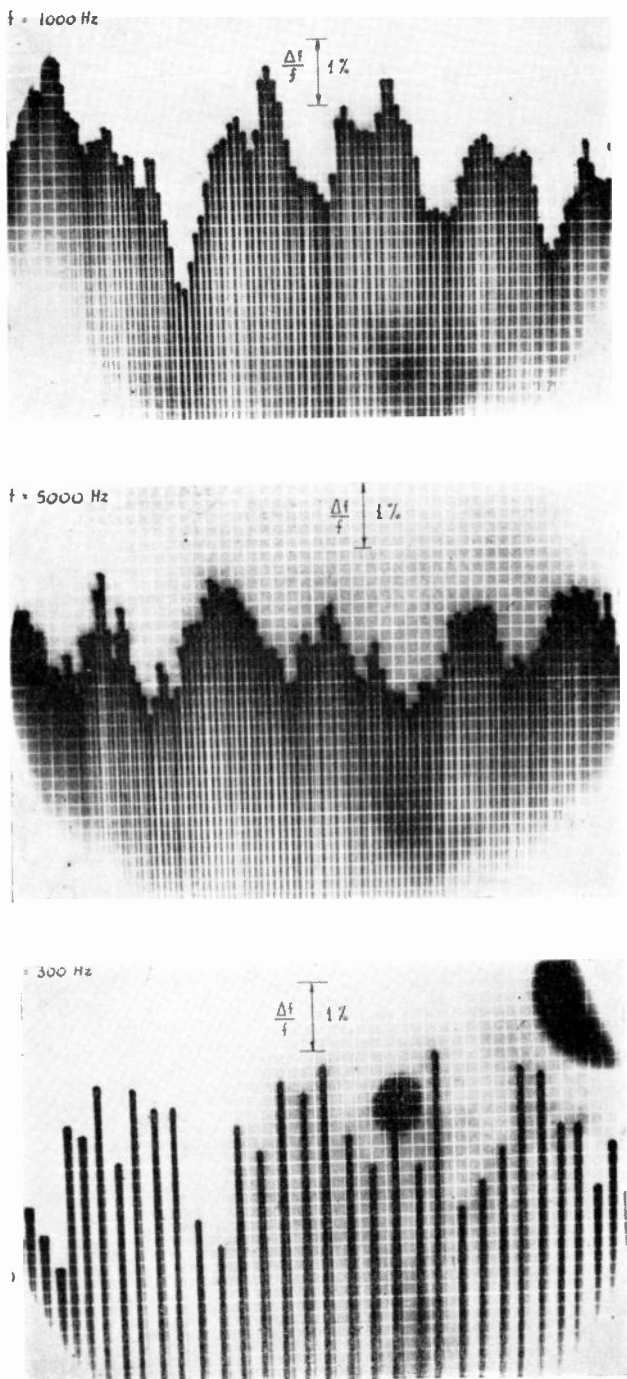


FIG. 7. — Enregistrement de la valeur instantanée du pleurage pour différentes fréquences du signal examiné.

n'est pas en relation avec le nombre de tours des organes mécaniques.

Pour examiner le comportement du ruban dans ces conditions il faut avant tout considérer ses

propriétés mécaniques statiques, c'est-à-dire l'allongement élastique du ruban en fonction de la force appliquée. Cette caractéristique a été déterminée par plusieurs auteurs et pour différentes qualités de ruban. Elle présente un premier trait assez linéaire et pour une certaine charge, à la limite d'élasticité, se manifeste un allongement à caractère plastique jusqu'à ce que l'on atteigne la charge de rupture.

Les études effectuées dans les Laboratoires Kodak ont montré que la limite d'allongement élastique des rubans de chlorure de polyvinyle est atteinte pour une charge de 1,8 kg et un allongement de 3 %, pour un ruban de triacétate de cellulose, elle est atteinte pour une charge de 2,8 kg et un allongement de 3 %. Ces valeurs changent sensiblement selon la température et l'humidité.

Le paramètre qui définit le comportement élastique du ruban est le module d'élasticité E , donné par l'expression suivante en newton/m² :

$$E = \frac{l}{\Delta l} \frac{F}{S} \quad \text{N/m}^2$$

où l est la longueur exprimée en mètres, Δl l'allongement en mètres, F la force appliquée en newtons,

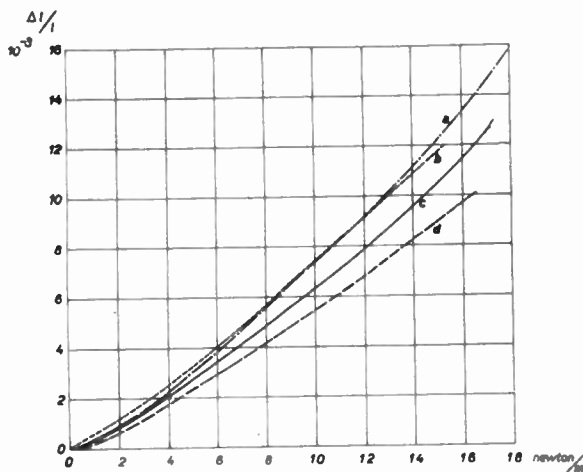


FIG. 8. — Courbes d'allongement statique de différents rubans magnétiques.

S la surface de la section du ruban en mètres carrés. Pour tous les rubans d'une largeur de 6,2-6,5 mm environ, cette surface varie de 0,30 à 0,35 10⁻⁶ m² ; suivant les données expérimentales relevées par WERNER à Berne, E varie de 3.10⁹ à 5.10⁹ newton/m². En conditions normales le ruban travaille avec une tension assez faible qui peut parfois atteindre des maxima d'environ 10 newtons, à laquelle correspondent des allongements d'environ 1 et 1,5 %.

Nous avons relevé la caractéristique élastique sur différents rubans à ces faibles charges (fig. 8) : cette mesure montre que le comportement élastique des rubans n'est pas parfaitement linéaire ; par conséquent le module d'élasticité dépend de la force appliquée, c'est-à-dire qu'il tend à diminuer lorsque la force appliquée augmente.

Après les mesures statiques, considérons les mesures dynamiques. Plusieurs de ces mesures ont été effectuées par WERNER qui a calculé qu'un morceau de ruban d'un mètre de long, fixé aux deux extrémités, vibre avec une fréquence propre variable de 675 à 830 Hz. Afin d'examiner le comportement

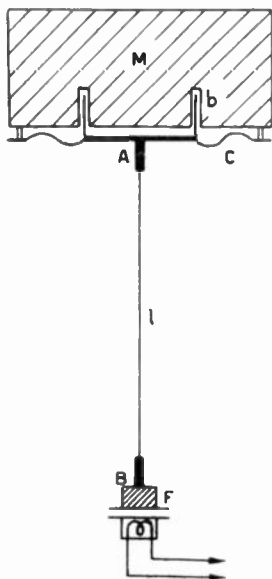


FIG. 9. — Appareillage pour la mesure des caractéristiques dynamiques de rubans magnétiques.

dynamique du ruban, on a réalisé le dispositif suivant (fig. 9) : un système électrodynamique comprenant un aimant permanent M et une bobine mobile b , tenue par un centreur élastique c , fait vibrer le point A , auquel est suspendu un morceau AB de ruban magnétique de longueur l . Au point B est appliqué un poids de F newtons. Le mouvement du point B est enregistré par un système mobile, qui indique la vitesse de déplacement. Si on fait passer dans la bobine un courant de fréquence f , le point A se met à vibrer et il communique au point B le mouvement par l'intermédiaire du ruban.

Si le système est linéaire on peut établir la relation suivante entre la fréquence de résonance F du système et les paramètres élastiques de la bande :

$$E = \frac{(2\pi f)^2 l F}{g S} = 3,63 \left(\frac{f^2 l F}{S} \right) \text{ N/m}^2$$

(on a négligé les effets du poids propre de la bande).

En traçant la courbe de vitesse de vibration en fonction de la fréquence, on peut obtenir soit le module d'élasticité dynamique soit le facteur de mérite mécanique de la bande, défini par $1/2\pi/ESR$ (R est la résistance de friction interne, visqueuse). Ce facteur qui indique les pertes mécaniques de la bande, peut rendre compte dans certains cas du comportement mécanique de la bande même.

Au cours des différentes déterminations, on a de suite constaté que par suite de la non linéarité des propriétés élastiques de la bande, son comportement dynamique présente une allure différente de celle

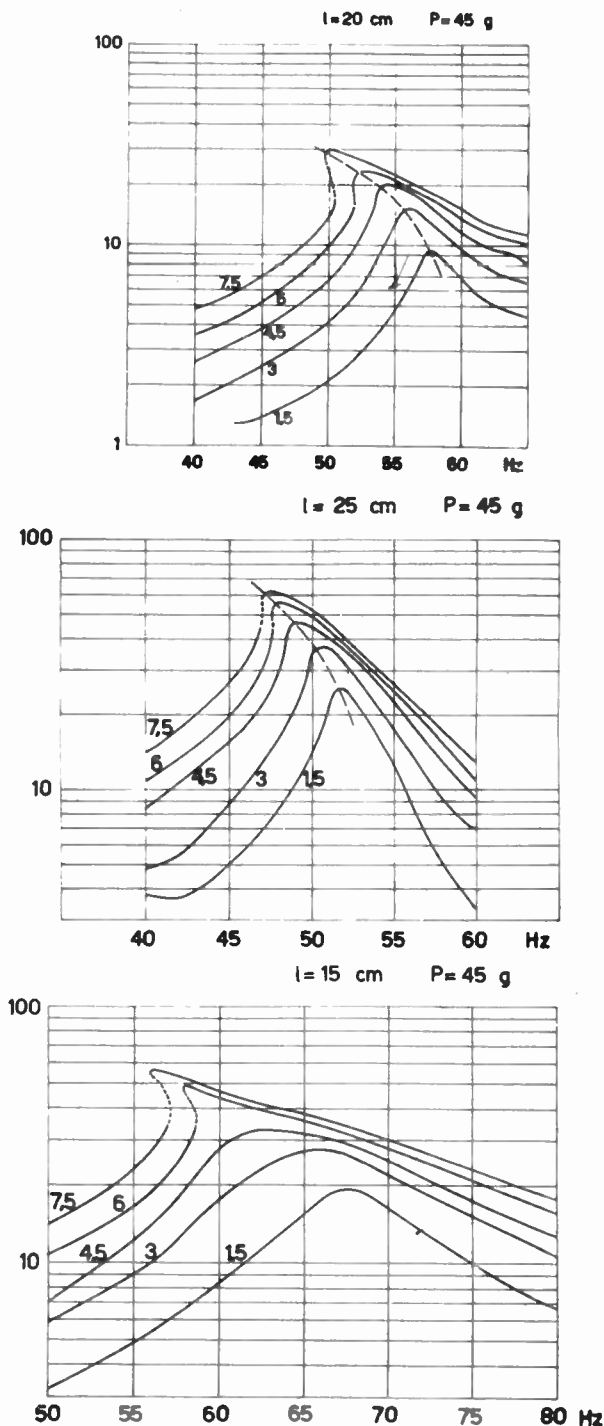


FIG. 10. — Caractéristiques dynamiques non linéaires de rubans magnétiques pour différentes valeurs de l'excitation et de la longueur de bande.

prévue, car la fréquence de résonance, dans les mêmes conditions de longueur et de force appliquée, est fonction de l'amplitude du mouvement du point traînant A (c'est-à-dire, de la tension appliquée à la bobine mobile). Si l'on augmente encore l'amplitude du mouvement du point A , on obtient des courbes caractéristiques qui présentent deux discontinuités (une pour les fréquences croissantes et une pour les fréquences décroissantes); celles-ci caractérisent les systèmes non linéaires.

La figure 10 montre quelques caractéristiques dynamiques.

Il faut tenir compte que dans le magnétophone la bande est soumise à des sollicitations mécaniques pendant de brefs instants seulement, correspondant au passage entre les deux tendeurs, mais les forces auxquelles elle est soumise peuvent être notables surtout à cause des grandes vitesses de l'enroulement de retour.

PERCEPTION DU PLEURAGE

Un problème fondamental dans l'étude du pleurage est celui d'en fixer les limites de perception.

Nous rappelons ici les résultats obtenus par SHOWER et BIDDULPH et par WERNER et WEDELL : la limite de la discrimination de fréquence est à peu près de 2,5 Hz, indépendamment de la fréquence, jusqu'à 2 000 Hz environ ; elle augmente ensuite en augmentant la fréquence et atteint 40 Hz pour des fréquences de 1 000 Hz. Au contraire la limite relative $\Delta f/f$ est à peu près constamment de 3 %, de 1 000 à 10 000 Hz ; elle augmente à mesure que la fréquence diminue jusqu'à la valeur de 4 % vers 80 Hz. La limite différentielle seuil de fréquence est fonction aussi du niveau de la sensation du son, c'est-à-dire qu'elle augmente à mesure que le niveau de sensation diminue. La limite différentielle dépend aussi de la fréquence avec laquelle les variations périodiques de hauteur se répètent de la fréquence plus haute à la plus basse ; précisément jusqu'à la fréquence de modulation de 3 Hz la limite est constante, aux fréquences plus grandes la limite croît.

Il s'agit maintenant de continuer ces recherches pour connaître l'effet du pleurage dans le cas de reproduction d'un programme musical ou de la voix parlée.

Pour effectuer ces essais on a créé un système qui produit artificiellement un pleurage dont on peut régler l'amplitude et la fréquence.

Comme il s'agit ici d'une évaluation tout à fait subjective, il faut se rapporter à l'opinion de nombreux observateurs et établir quelques degrés d'évaluation, comme on fait pour d'autres déterminations de ce genre.

Les degrés d'évaluation ont été fixés dans la manière suivante : pleurage à peine perceptible, nettement perceptible, gênant.

Il va sans dire que ces évaluations présentent une grande incertitude ; elles peuvent néanmoins donner quelques indications pour déterminer les effets du pleurage dans une reproduction sonore.

On peut tout d'abord observer que l'opinion des auditeurs est très différente selon qu'il s'agit de l'enregistrement de la parole ou de la musique et selon le caractère de la musique. Dans le cas de la voix une modulation de fréquence de quelques unités pour-cent n'est pas perceptible, car on remarque uniquement un changement de la qualité et de la couleur de la voix, qui dans la plupart des cas et dans certaines limites est même agréable. La voix paraît plus vibrante, et si on ne connaît pas la

voix originale on n'est pas en état de juger s'il s'agit d'une altération.

Pour la musique c'est différent, surtout s'il s'agit du piano, du violon, avec musique peu rythmée et avec des notes plutôt longues. Dans ce cas une oreille douée d'une bonne sensibilité musicale peut percevoir des modulations de fréquence à rythme lent de la valeur de 1 Hz, de $\Delta f/f$ à 6 % environ.

On se sert pour les dites déterminations de certains dispositifs qui permettent d'introduire dans la reproduction un pleurage à caractéristiques préétablies.

Une première méthode réalisée expérimentalement est basée sur l'effet Doppler : on enregistre la voix ou l'exécution musicale par le moyen d'un microphone mobile avec un mouvement alternatif d'amplitude et de fréquence connues. Le microphone est situé sur un petit guide mobile animé d'un mouvement alternatif par un dispositif à bielle et manivelle, qui permet d'altérer l'enregistrement selon les caractéristiques voulues.

On peut aussi se servir d'un enregistrement transmis par un haut-parleur animé d'un mouvement alternatif d'amplitude A et de fréquence F_0 , du même genre de celui que nous avons décrit tout à l'heure.

Suivant la théorie de l'effet Doppler si le son de fréquence f émis par une source fixe arrive à un observateur mobile avec vitesse v , celui-ci reçoit en effet un signal de fréquence f_1 , liée à la fréquence f de la source par la relation suivante :

$$f_1 = f \left(1 + \frac{v}{c} \right)$$

où c est la vitesse du son.

On a donc :

$$\frac{f_1 - f}{f} = v/c$$

Si v est la vitesse d'un mouvement alternatif de fréquence F_0 et d'amplitude A on a :

$$v = 2\pi F_0 A \sin 2\pi F_0 t$$

et par conséquent l'écart maximum de fréquence (peak to peak) est donné par :

$$\frac{\Delta f}{f} = \frac{4\pi}{c} A F_0 = 0,37 A F_0 \text{ ‰}$$

où A est exprimé en cm et c en cm/s.

D'après l'examen de cette relation on déduit que la méthode convient particulièrement pour produire de petits pleurages, mais tout de même facilement dosables dans les limites assez larges de fréquence.

En se servant du premier de ces deux systèmes, en enregistrant avec un microphone omnidirectionnel (si la distance de celui-ci de la source est grande par rapport à l'amplitude A du mouvement

alternatif) on peut aisément remplacer le mouvement à bielle alternatif par un mouvement circulaire du microphone, qui peut être plus facilement réalisé. Pour éviter les éventuels effets de modulation d'amplitude il suffit d'effectuer l'enregistrement et l'observation dans une chambre qui ne soit pas trop absorbante.

Il est possible d'obtenir une reproduction troublée par le pleurage, en agissant directement sur le système de reproduction ; on mentionne par exemple la méthode suivie par CORBINO en 1935, qui a réussi à obtenir une certaine modulation de fréquence en déplaçant d'une certaine valeur connue le centre

facile de constater que par cette méthode on peut aisément obtenir de très fortes valeurs de pleurage.

Les tableaux de la figure 19 présentent les résultats de ces déterminations effectuées par une trentaine de personnes, pendant la reproduction d'un passage lent de l'exécution au piano de la « Rêverie » de Schumann. La série de mesure a été effectuée avec $A = 0,5$ mm, n variable.

LE PLEURAGE DANS LE RÉENREGISTREMENT

Il y a une question très intéressante dans l'étude du pleurage : rechercher son comportement dans

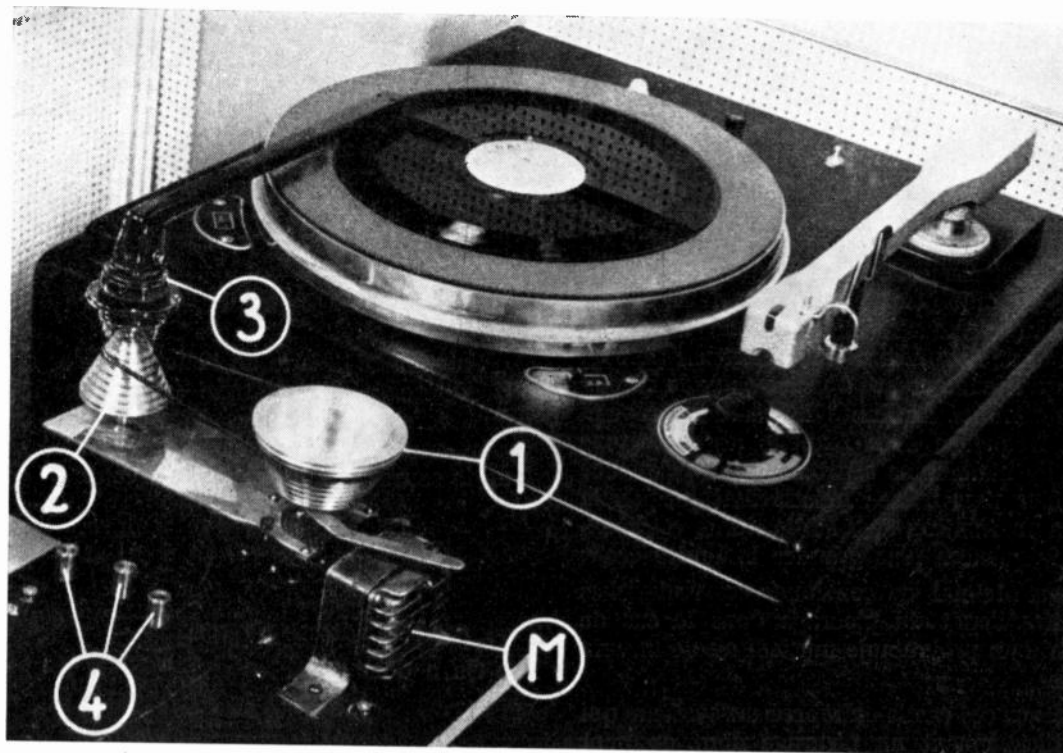


FIG. 11. — Dispositif pour créer un pleurage de caractéristiques voulues dans la reproduction par tourne-disques.

du disque par rapport à celui du plateau tournant ; on obtient ainsi un pleurage d'amplitude qui peut être réglé avec l'excentricité, et de fréquence constante liée au nombre de tours du disque.

Avec un mouvement alternatif du pick up on obtient une reproduction troublée par un pleurage dont les caractéristiques dépendent de celles du mouvement du pick up ; on obtient ceci avec une série de cames ayant une excentricité variable, actionnée par un moteur à vitesse réglable (fig. 11).

Si le déplacement maximum total du centre de rotation du bras du pick up est représenté par A , et par n le nombre de tours par minute, par N le nombre de tours du plateau tournant, par R le rayon du sillon en question, on a :

$$\frac{\Delta f}{f} = \frac{A}{R} \frac{n}{N}$$

Pour un plateau tournant à 78,4 tours/m sur le rayon moyen de 110 mm on a $0,116 A n \text{ ‰}$. Il est

des enregistrements successifs effectués avec des appareils qui présentent à peu près les mêmes propriétés.

Le problème a d'abord un intérêt pratique, puisqu'il arrive souvent de réenregistrer plusieurs fois un même enregistrement ; il est donc important de pouvoir prévoir à chaque répétition quelle est le taux du pleurage. La question est aussi intéressante au point de vue des mesures, car on sait que le réenregistrement répété peut nous donner un élément pour évaluer un appareil.

Si l'on veut considérer le problème du point de vue analytique on peut avant tout rappeler que la recherche de la composition spectrale d'une onde modulée en fréquence modulée à son tour en fréquence, présente des difficultés telles qu'elles en viennent à ôter toute valeur pratique aux résultats qu'on peut obtenir.

Nous pouvons tout de même limiter notre examen au cas d'un réenregistrement sur des appareils ayant

les mêmes propriétés, c'est-à-dire la même fréquence et la même profondeur de modulation.

Suivant un procédé pas tout à fait exact (cependant utile pour justifier les résultats que l'on obtient) après le premier enregistrement la fréquence f_1 est donnée par l'expression suivante :

$$f_1 = f (1 + h \sin 2\pi Ft)$$

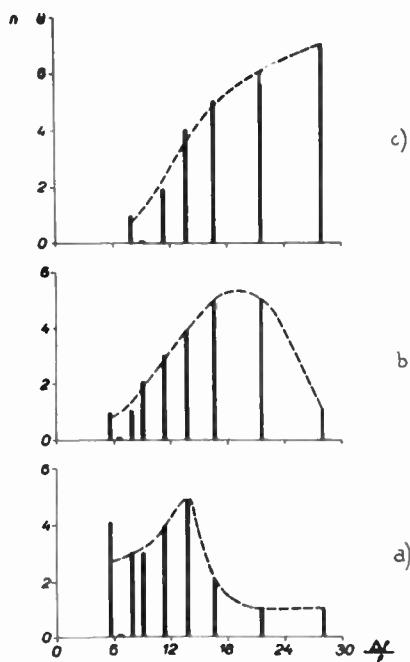


FIG. 12. — Statistique de la perception du pleurage dans la musique : a) seuil de perception ; b) perception nette ; c) gêne.

Après le deuxième enregistrement on peut admettre que l'on a :

$$f_2 = f_1 (1 + h \sin (2\pi Ft + \varphi)) ; \\ f (1 + h \sin 2\pi Ft) (1 + h \sin (2\pi Ft + \varphi)) ;$$

où φ est un angle classique de déphasage qui peut prendre une valeur quelconque.

Dans ce cas la nouvelle fréquence est donnée par :

$$f_2 = f (1 + h) (1 + \cos \varphi) \sin 2\pi Ft + \sin \varphi \cos 2\pi Ft - \\ = f \left[1 + 2h \cos \frac{\varphi}{2} \sin (2\pi Ft + \varphi_1) + \dots \right]$$

où le deuxième terme représente une fréquence modulée avec une profondeur égale à $2h \cos \frac{\varphi}{2}$; on néglige les termes en h^2 , puisque l'on suppose h assez petit.

Le déphasage φ peut être un angle quelconque ; comme les fréquences enregistrées qui en résultent sont légèrement différentes dans les deux cas, on peut supposer les valeurs de φ continuellement variables et distribuées au hasard ; la valeur moyenne de la profondeur de modulation est alors :

$$\frac{1}{\pi} \int_0^\pi 2h \cos \frac{\varphi}{2} d\varphi = \frac{4h}{\pi} = 1,27h$$

Par conséquent le deuxième enregistrement présente une variation pour-cent de fréquence 1,27 fois plus grande que le premier.

On ne peut pas répéter à l'infini le raisonnement que l'on vient de faire ; on peut pourtant retenir comme valable que pour les premiers réenregistrements il y a pour les pourcentages de variation de fréquence la relation

$$\left(\frac{\Delta f}{f} \right)_{n+1} = 1,27^n \left(\frac{\Delta f}{f} \right)_0$$

Si h_n est la profondeur de modulation à une répétition n , on obtient par le même raisonnement :

$$h_{n+1}^2 = h_n^2 + h^2 + 2h_n h \cos \varphi_n$$

Cette relation peut nous donner une indication sur les écarts possibles, en fonction de φ_n .

C'est-à-dire que la profondeur de modulation h_{n+1} à la $(n+1)^e$ répétition est représentée vectoriellement par le troisième côté d'un triangle dont les deux premiers côtés sont h et h_n , formant l'angle φ_n entre eux (fig. 20) ; ceci représente le déphasage

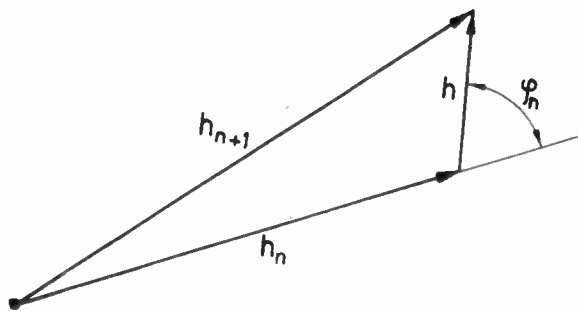


FIG. 13. — Augmentation de la profondeur de modulation de la n à la $(n+1)^e$ répétition de l'enregistrement.

de la fréquence de modulation dans les répétitions successives.

Si l'on suit le phénomène du premier réenregistrement, on remarque (fig. 21) que le vecteur h_n compose une ligne polygonale formée par $(n+1)$ vecteurs de module h , dont chacun est déphasé d'un angle φ_i par rapport au précédent.

Le problème de la limite de la profondeur de modulation après un nombre n de répétitions est analogue à celui étudié par RAYLEIGH de la composition de plusieurs vecteurs égaux, déphasés entre eux par des angles distribués au hasard avec une égale probabilité.

RAYLEIGH a calculé que, en augmentant le nombre n des vecteurs de module h , la valeur la plus probable du module de la résultante tend vers $h\sqrt{n}$.

La probabilité pour que le carré du module de la résultante soit compris entre r et $r+dr$ est donnée, toujours selon RAYLEIGH, par :

$$\frac{2}{n h^2} e^{-22/n h^2} dr$$

On a effectué quelques mesures de la variation en pour-cent de la fréquence dans les réenregistrements répétés avec deux magnétophones qui présentaient les mêmes caractéristiques de pleurage.

Les résultats de ces mesures, confirmés par les données d'autres couples de magnétophones, sont représentés dans la figure 22.

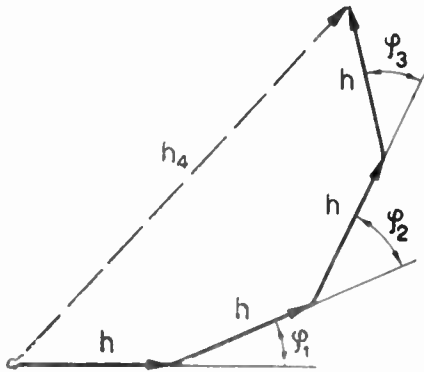


FIG. 14. — Représentation de la profondeur de modulation pour $n + 1$ réenregistrements successifs.

On remarque que selon la théorie que l'on vient d'exposer, la variation en pour-cent de fréquence des premiers réenregistrements suit une loi de progression géométrique, représentée dans la même figure par la ligne droite.

Dans la figure 22a est représenté le comportement du $\frac{\Delta f}{f}$ pour un grand nombre de réenregistrements,

d'après la loi que RAYLEIGH a trouvée pour des phénomènes analogues. On remarque que les deux diagrammes peuvent très bien être comparés.

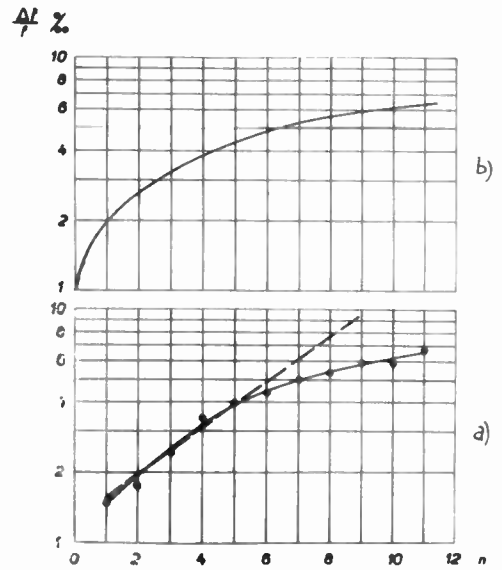


FIG. 15. — Pleurage de réenregistrements répétés : a) mesures expérimentales ; b) données suivant la théorie de Rayleigh.

Dans la figure 23 on représente l'enregistrement du pleurage effectuée par un traceur à mouvement asservi à échelle linéaire, pour le deuxième, le quatrième et le septième réenregistrement.

MESURES DES PROPRIÉTÉS ÉLECTRIQUES DES SEMI-CONDUCTEURS

PAR

B. PISTOULET

Docteur ès-Sciences, Ingénieur à la Compagnie Française Thomson-Houston.

I. — INTRODUCTION.

Les caractéristiques électriques d'un semi-conducteur sont conditionnées par un certain nombre de facteurs que l'on peut classer en deux groupes : dans le premier figurent des grandeurs physiques caractéristiques de l'espèce physique considérée (germanium, silicium, etc...) et sensiblement indépendantes, tout au moins dans une certaine mesure, de la pureté du matériau ; c'est par exemple le cas de la mobilité des porteurs et de la constante de diffusion. Dans le second groupe on trouve un certain nombre de grandeurs qui dépendent directement du degré de contamination du matériau, de la nature des impuretés, et de la perfection du réseau cristallin ; c'est le cas de la résistivité, de la constante de Hall, et de la durée de vie des porteurs. Ce sont ces facteurs, différents suivant les échantillons, qu'il est important de connaître, d'une part aux différents stades de la préparation du matériau, pour laquelle on se reportera à l'exposé détaillé de J. M. MERCIER ⁽¹⁾, d'autre part au stade final de l'utilisation pour savoir si le matériau obtenu a bien les propriétés requises pour la construction du dispositif à semi-conducteur envisagé (diode, redresseur, transistor, etc...). Il sera seulement question ici des mesures de résistivité, de constante de Hall et de durée de vie. Les chiffres donnés ont uniquement trait au germanium, mais il est évident que les méthodes décrites s'appliquent quel que soit le semi-conducteur considéré.

II. — RAPPEL DE QUELQUES RELATIONS FONDAMENTALES.

Nous rappellerons d'abord rapidement quelques définitions et quelques relations fondamentales ⁽²⁾.

La mobilité μ d'un porteur de charge dans un cristal est définie par la relation :

$$\mu = \frac{v}{E} \quad (1)$$

v étant la vitesse moyenne prise dans le cristal, par le porteur, sous l'influence d'un champ électrique E . Pour des concentrations d'impuretés inférieures à 10^{18} atomes par cm^3 correspondant pour le Germanium à des résistivités de l'ordre de quelques ohms-cm, cette mobilité est principalement conditionnée par la diffusion des porteurs par les vibrations thermiques du réseau cristallin.

Parmi les travaux récents sur cette question, citons ceux de M. B. PRINCE ⁽³⁾ qui indique les valeurs maxima suivantes pour des mobilités des électrons et des trous dans le Germanium très pur à 300°K :

$$\mu_n = 3\,900 \pm 100 \text{ cm}^2/\text{volt. sec}$$

$$\mu_p = 1\,900 \pm 50 \text{ cm}^2/\text{volt. sec}$$

les indices n et p se rapportant respectivement ainsi que dans ce qui suit aux électrons et aux trous.

La conductibilité σ est définie par la relation :

$$\sigma = \frac{J}{E} \quad (2)$$

où J est la densité de courant, E le champ électrique. S'il y a n électrons et p trous libres par unité de volume, de charge $\pm e$, on voit aisément que σ a pour valeur :

$$\sigma = e(n\mu_n + p\mu_p) \quad (3)$$

Un effet particulièrement intéressant dans l'étude des semi-conducteurs est l'effet Hall. Soit (fig. 1) un échantillon parallélépipédique soumis à l'action d'un champ électrique uniforme E dirigé suivant

⁽¹⁾ J. M. MERCIER, La Technologie du Germanium, Onde Electrique, N° 328, juillet 1954, p. 559-572.

⁽²⁾ Voir par exemple : SHOCKLEY, Electrons and Holes in Semiconductors, D. Van Nostrand Company Inc. New York, N. Y. 1950.

⁽³⁾ M. B. PRINCE, « Drift mobilities in semiconductors, I. Germanium », *Physical Review*, 92, 3, nov. 1953, P. 681-687.

l'axe x , et d'un champ magnétique uniforme H dirigé suivant l'axe z . Les charges mises en mouvement par E sont soumises lorsqu'on applique H à une force $F_1 = \pm e v \wedge H$. Sous l'action de cette force elles viennent s'accumuler sur les faces de l'échantillon normales à l'axe y . Cette accumulation produit à l'intérieur du spécimen un champ électrique E_y exerçant sur les charges en mouvement une force de sens opposé à la force de Laplace. L'équilibre est atteint lorsque la densité superficielle de charges est telle que ces deux forces s'équilibrent, condition nécessaire pour que le courant I circule suivant Ox . Le champ électrique résultant et le courant font alors entre eux un angle θ . Lorsqu'on est en présence d'un seul type de porteurs, trous ou électrons, on a :

$$\theta \simeq \tan \theta = \frac{E_y}{E} = \pm \mu H \quad (4)$$

le signe plus s'appliquant pour les trous, le signe moins pour les électrons. La mesure de θ permet donc de déterminer la nature des porteurs et leur mobilité. En pratique, on mesure la différence de potentiel V

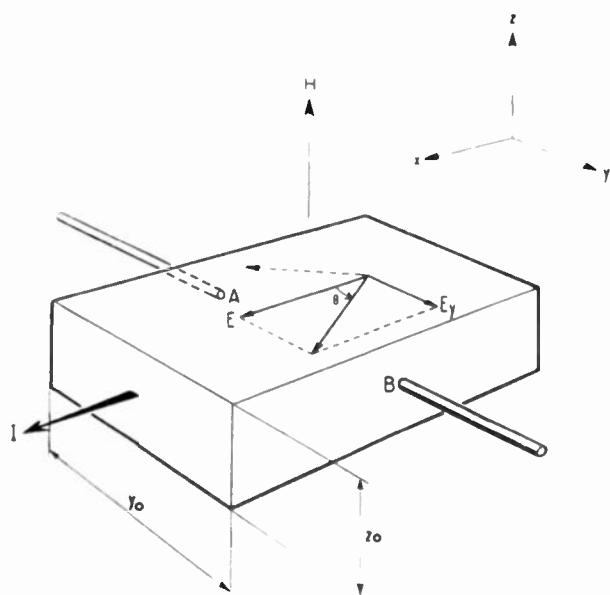


FIG. 1

apparaissant entre les points A et B et on détermine la valeur de la constante de Hall R , égale par définition au rapport de E_y par le produit de l'intensité du champ magnétique H et de la densité de courant J :

$$R = \frac{E_y}{HJ} = \pm \frac{\mu}{\sigma} = \pm \frac{1}{ne} \quad (5)$$

Un traitement plus rigoureux de la question (*) conduit à la relation :

(*) Voir par exemple F. SEITZ, *The Modern Theory of Solids*, Mc Graw-Hill, New York, 1940.

$$R = \frac{3\pi}{8} \cdot \frac{1}{ne} \quad (5 \text{ bis})$$

Ce coefficient fournit donc une mesure directe du nombre de porteurs présents par unité de volume. R est relié à la tension V , à l'intensité du courant I , à H , et à la dimension z_0 de l'échantillon par la relation :

$$R = \frac{V z_0}{H I} \quad (6)$$

Si on exprime V en volts, I en ampères, H en gauss, z_0 en cm, la valeur de R en $\text{cm}^2/\text{coulomb}$ est donnée par :

$$R = 10^8 \frac{V z_0}{H I} \quad (7)$$

Lorsque les deux types de porteurs, trous et électrons sont présents, on montre que :

$$R = - \frac{3\pi}{8} \cdot \frac{n b^2 - p}{(n b + p)^2 e} \quad (8)$$

$$\text{avec } b = \frac{\mu_n}{\mu_p} \quad (9)$$

En réalité, on se trouve toujours en présence de porteurs des deux types. Lorsqu'il n'existe dans le cristal ni défauts cristallins, ni atomes d'impuretés, les seuls porteurs libres sont constitués par un certain nombre d'électrons dans la bande de conduction et un nombre égal de trous dans la bande de valence. Cette concentration de porteurs appelée concentration intrinsèque n_i n'est fonction que de la température. Lorsque le semi-conducteur contient une certaine proportion d'atomes donneurs ou accepteurs les concentrations d'électrons et de trous sont reliées par la relation :

$$np = n_i^2 \quad (10)$$

Les relations (3), (10), et (8), permettent de calculer σ en fonction de R ou inversement, si l'on connaît μ_n , μ_p et n_i . Cette question est développée dans une étude de P. G. HERKART et J. KURSHAN (*) ; les valeurs de n_i , μ_n et μ_p n'étant pas connues de manière très précise, les relations numériques entre R et σ qu'ils en déduisent ne sont sans doute pas absolument valables ; d'autre part, il existe un certain désaccord (6) entre les valeurs des mobilités intervenant dans les relations (1) et (4). Pour fixer les idées, disons simplement qu'il est courant de mesurer sur du Germanium, suivant son degré de

(*) P. G. HERKART and J. KURSHAN, Theoretical resistivity and Hall coefficient of impure Germanium near room temperature, *RCA Rev.* 1953, 14, n° 3, sept. p. 427-440.

(6) J. R. HAYNES and W. SHOCKLEY, The mobility and life of injected holes and electrons in germanium, *Phys. Rev.* 81, 5, Mars 1951, p. 835-843.

pureté, des valeurs de la constante de Hall comprises entre 1 000 et 80 000 cm³/coulomb).

Une autre caractéristique importante est la durée de vie des porteurs dans le germanium : lorsque l'on injecte, par exemple, au moyen d'une jonction ou d'un contact ponctuel, des porteurs minoritaires dans un semi-conducteur on crée une situation différente de la situation d'équilibre thermodynamique ; la vitesse avec laquelle cet équilibre se rétablit est caractérisée par la durée de vie τ des porteurs qui est définie par la relation :

$$p(t) = p_0 e^{-\frac{t}{\tau}} \quad (11)$$

où p_0 est le nombre de porteurs excédentaires présents au temps zéro, et $p(t)$ le nombre de porteurs subsistant au temps t . La durée de vie peut être considérablement réduite par des imperfections cristallines et par certaines impuretés qui se comportent alternativement comme des centres de recombinaison ou de génération des paires électron-trou.

En pratique on a intérêt à ce que la durée de vie soit la plus grande possible, aussi bien pour réduire le courant inverse des jonctions que pour améliorer le gain des transistors, par exemple. Les valeurs les plus élevées citées à l'heure actuelle dans la littérature sont de l'ordre de la milliseconde.

III. — DISPOSITIFS DE MESURE.

a) Résistivité.

Le moyen le plus simple de mesurer la résistivité d'un barreau de germanium est de mesurer sa résistance entre deux sections suffisamment rapprochées pour que l'on puisse considérer sa résistivité comme

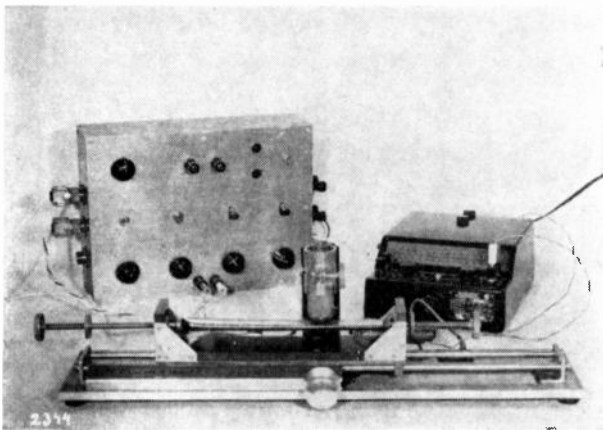


Fig. 2

constante entre les deux sections. Il suffit, pour cela, de mesurer par une méthode potentiométrique la différence de potentiel apparaissant entre deux points d'une génératrice lorsque le barreau est parcouru par un courant donné, la répartition des lignes de courant étant uniforme. La fig. 2 représente le dispositif utilisé à cet effet.

Cette méthode présente l'inconvénient de n'être praticable que sur des spécimens de section constante, et suffisamment longs pour que la mesure se fasse dans une région éloignée des connexions, au voisinage desquelles se produisent des perturbations notables dans la répartition des filets de courant. Ces conditions se trouvent rarement remplies dans les très nombreux cas, rencontrés en pratique, où il s'agit de mesurer la résistivité de spécimens destinés à différents usages, et pouvant avoir des formes et des dimensions variées ; très souvent par exemple, il s'agit de disques d'épaisseurs diverses, découpés perpendiculairement à l'axe d'un monocristal. Il est donc d'un grand intérêt de disposer d'une méthode de mesure non destructive, pouvant être directement mise en œuvre sur les spécimens mêmes qui doivent servir à la réalisation des dispositifs à semi-conducteur envisagés. Aussi avons-nous porté nos efforts à la Compagnie Française Thomson-Houston sur la mise au point d'une méthode satisfaisant à ces conditions, la méthode des pointes. La théorie complète, ainsi que les calculs relatifs aux différents cas possibles, font l'objet d'une étude détaillée de J. LAPLUME (7), à laquelle nous renvoyons le lecteur. Nous nous bornerons simplement à en rappeler le principe, et à décrire la réalisation de l'appareillage. Supposons que l'on applique sur la surface d'un échantillon deux pointes reliées à un générateur, de manière à produire le passage d'un certain courant ; si l'on sait déterminer la répartition du potentiel dans l'échantillon, la mesure de la d.d.p. apparaissant entre deux autres pointes, appuyées en deux points déterminés de sa surface, permet d'en calculer la résistivité.

Nous avons choisi deux dispositions de pointes : dans la première les quatre pointes sont alignées et équidistantes, le courant étant injecté par les deux pointes extrêmes, et la tension mesurée entre les deux pointes centrales, ou inversement ; dans l'autre, les quatre pointes sont disposées aux sommets d'un carré, les deux paires de pointes étant disposées suivant deux côtés opposés du carré.

Dans les deux cas, les pointes sont supportées par une tête isolante comportant quatre logements cylindriques dans lesquels les pointes coulisent avec un jeu minimum. Chaque pointe est maintenue en saillie hors de la tête par un ressort qui sert également de connexion électrique. Lorsqu'on applique la tête sur un échantillon, la pression de contact de chaque pointe est fournie par le ressort correspondant. Les pointes sont acérées, de sorte que les contacts pointe semi-conducteur peuvent être assimilés à des contacts ponctuels (le diamètre de la plage de contact est en effet très petit devant l'écartement des pointes qui est de l'ordre de quelques millimètres). L'échantillon reposant sur un support isolant, on vient appuyer, sur sa face supérieure, la tête porte-pointes au moyen d'un dispositif à crémaillère. Le générateur est à courant constant, le courant injecté étant le courant anodique d'une

(7) J. LAPLUME. — Bases théoriques de la mesure de la résistivité et de la constante de Hall par la méthode des pointes. *Onde Electrique* (à paraître).

pentode de résistance interne très élevée devant les résistances de contact. La mesure de la différence de potentiel apparaissant entre les deux autres pointes se fait par une méthode d'opposition. La figure 3 représente l'appareil utilisé.

Ce dispositif rend de très grands services, en per-

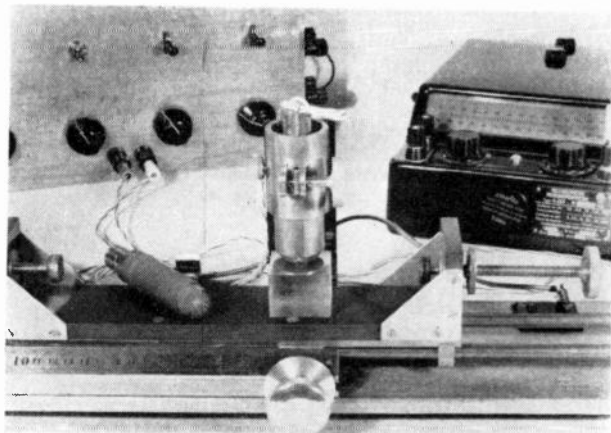


FIG. 3

mettant la mesure rapide et précise d'échantillons très divers, à condition qu'ils soient monocristallins ; s'il n'en est pas ainsi, la présence des joints intercrystallins peut modifier profondément le trajet des filets de courant, et fausser considérablement les mesures. Si les mesures effectuées en différents points d'un échantillon conduisent à des résultats dont la dispersion est notable, cela permet de conclure que l'échantillon présente des défauts d'homogénéité qu'il est très important de déceler, car ces défauts rendent presque toujours le matériau inapte à être utilisé dans la construction de dispositifs à semi-conducteurs. Les échantillons à mesurer de manière courante peuvent être classés en deux catégories : d'une part, des lingots de dimensions grandes devant l'écartement des pointes ; dans ce cas leur forme intervient fort peu dans la répartition des lignes de courant et on peut ramener ce cas à celui d'un solide de très grandes dimensions, limité par une surface plane ; d'autre part, des disques à faces parallèles grossièrement circulaires et d'épaisseurs diverses, à l'intérieur desquels la répartition des filets de courant est sensiblement identique à celle existant dans un bloc parallélépipédique de même épaisseur, et ayant pour base un carré ou un rectangle circonscrit au disque, ce qui permet d'appliquer les calculs (*) correspondant à ce cas.

b) Constante de Hall.

Généralement on mesure la constante de Hall R , conformément à la méthode décrite au paragraphe II, en utilisant des échantillons parallélépipédiques. Des raisons, analogues à celles que nous avons exposées, au sujet de la mesure de la résistivité, nous ont fait rechercher une méthode de mesure non destructive pouvant être mise en œuvre sur des spécimens de formes et de dimensions variées. La solution adoptée

présente de grandes analogies avec celle qui a été décrite ci-dessus pour la mesure de la résistivité, puisque l'on utilise encore quatre pointes, deux servant à l'injection du courant, et les deux autres à la mesure de la tension de Hall (*). Nous renvoyons à l'article de J. LAPLUME (7) pour l'exposé théorique de la méthode. La figure 4 représente la disposition des pointes sur l'échantillon. Le courant est injecté par les pointes 1 et 2 et la tension de Hall recueillie entre les pointes 3 et 4, le champ magné-

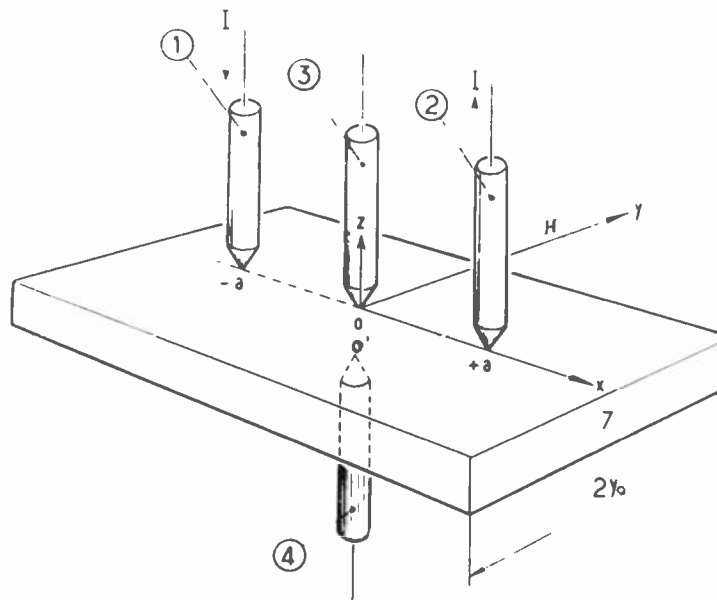


FIG. 4

tique uniforme H étant appliqué normalement au plan défini par les axes des quatre pointes. Sous l'action de H , un certain nombre de charges s'accroissent sur les faces de l'échantillon, et l'équilibre

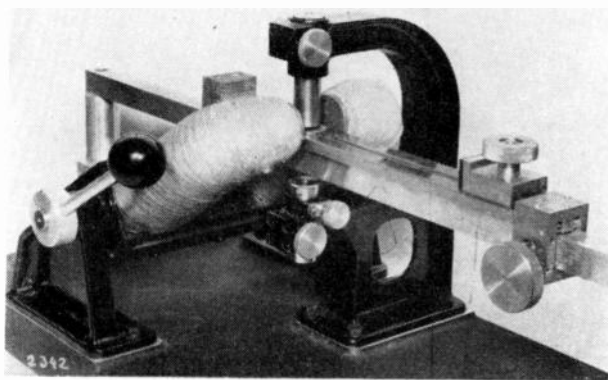


FIG. 5

est atteint lorsque le champ électrique créé par cette répartition de charges exerce sur les charges en mouvement une force égale et opposée à la force de Laplace. La tension de Hall est égal à la circulation de ce champ entre les points 0 et 0'.

(*) Brevet Français N° 647.145 du 27-4-53.

La figure 5 est une vue de l'appareil. La tête supérieure est semblable à celle de l'appareil de mesure de la résistivité décrite précédemment, la seule différence provenant de ce qu'elle porte trois pointes alignées au lieu de quatre. La tête inférieure porte une seule pointe. Deux crémaillères permettent de déplacer verticalement les deux têtes et de les appliquer sur les faces supérieure et inférieure de l'échantillon à mesurer. Un soin particulier a été apporté à la réalisation de l'ensemble, afin que la pointe centrale de la tête supérieure et la pointe de la tête inférieure soient rigoureusement coaxiales. Le chariot sur lequel est fixé l'échantillon se déplace horizontalement par rapport à l'ensemble de l'appareil, afin de permettre d'effectuer des mesures en différentes régions d'un échantillon.

Des mesures, effectuées sur un certain nombre d'échantillons de Germanium, nous ont montré que les résultats obtenus par la méthode des pointes, et par les méthodes courantes, différaient au maximum de quelques pour cent, c'est-à-dire que les écarts observés étaient de l'ordre des erreurs d'expérience.

Sur la fig. 6 on a porté, à titre d'exemple, les valeurs de la constante de Hall R , et de la résistivité ρ , d'un monocristal, mesurées à différentes abscisses. Ce specimen étant assez fortement contaminé, R et ρ sont sensiblement proportionnels.

c) Durée de vie.

Différentes méthodes sont employées pour mesurer la durée de vie des porteurs minoritaires injectés dans les semi-conducteurs. Dans la méthode de Morton et Haynes, un pinceau lumineux, tombant sur la surface convenablement traitée d'un cristal de germanium, provoque la libération de paires

électron-trou ; la mesure de la concentration des porteurs minoritaires en fonction de la distance au point d'impact permet de déterminer la valeur de la durée de vie de ces porteurs, la constante de diffusion étant connue. On trouvera une étude détaillée de la question faite par L. B. VALDES ⁽⁹⁾, dans le cas où la zone d'impact du faisceau lumineux est une raie rectiligne ; le pinceau lumineux est haché par un disque opaque tournant percé de trous ; une pointe collectrice, appuyée en un point de la surface du germanium, capte un courant proportionnel à la densité de porteurs en ce point ; on peut alors tracer sur un graphique la courbe donnant le nombre de porteurs présents au voisinage du collecteur en fonction de la distance r séparant le collecteur de la raie lumineuse. Cette relation s'écrit, tout au moins pour les valeurs de r pas trop voisines de zéro :

$$\frac{p}{p_0} = j H_0^{(1)} \left(j \frac{r}{\sqrt{D\tau}} \right) \tag{12}$$

p étant le nombre de porteurs excédentaires subsistant à la distance r de la raie, p_0 le nombre de porteurs excédentaires injectés, $H_0^{(1)}$ la fonction de Hankel d'ordre zéro, D la constante de diffusion, τ la durée de vie. D étant connu (pour le germanium $D_n \simeq 93 \text{ cm}^2/\text{sec}$, $D_p \simeq 13 \text{ cm}^2/\text{sec}$), on en déduit alors la valeur de τ pour l'échantillon considéré.

Un autre procédé de mesure de la durée de vie a été décrit par D. NAVON, R. BRAY, et H. Y. FAN ⁽¹⁰⁾. Il est basé sur le fait que l'injection de porteurs excédentaires dans un specimen abaisse sa résistivité. Le nombre de porteurs excédentaires présents est déterminé par la variation de conductance produite. Le nombre de ces porteurs subsistant après une injection décroissant exponentiellement avec le temps, on peut donc déduire τ du relevé de la décroissance de la conductance du specimen en fonction du temps. Cette méthode nécessite un appareillage spécial (générateur d'impulsions à fort niveau, synchroscope...), et paraît plus délicate à mettre en œuvre que la précédente.

Enfin, comme l'ont montré P. AIGRAIN et H. BULLIARD ⁽¹¹⁾ dans leurs travaux sur l'effet photomagnétoélectrique, cet effet fournit un autre moyen de mesure de la durée de vie des porteurs excédentaires.

Nous avons choisi une méthode analogue à celle de L. B. VALDES sauf en ce qui concerne le mode d'illumination. Il est nécessaire de séparer le courant collecteur dû aux porteurs minoritaires injectés, du courant de repos. Ceci est obtenu dans le dispositif de L. B. VALDES en utilisant un pinceau lumineux haché à une fréquence acoustique ; le courant alternatif apparaissant au collecteur est alors amplifié

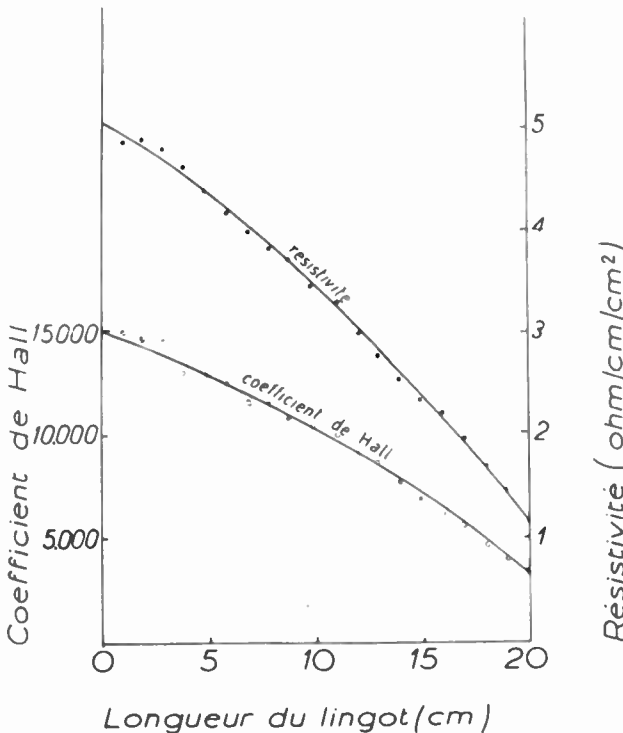


FIG. 6

⁽⁹⁾ L. B. VALDES, *P.I.R.E.* 40, Nov. 1952, N° 11, P. 1420-1423.
⁽¹⁰⁾ D. NAVON, R. BRAY, and H. Y. FAN, *P.I.R.E.* 47, Nov. 1952, 11, P. 1342-1347.
⁽¹¹⁾ P. AIGRAIN et H. BULLIARD, Sur la théorie de l'effet photomagnétoélectrique, *C. R. Acad. Sc.*, 236, 6, 9 Févr. 1953, P. 595-596 et P. AIGRAIN et H. BULLIARD, Résultats expérimentaux de l'effet photomagnétoélectrique, *C. R. Acad. Sc.*, 236, 7, 16 Février 1953, p. 672-674.

et détecté. Dans l'appareil représenté sur la figure 7, nous utilisons un flux lumineux continu ; le collecteur est polarisé et l'appareil de mesure est constitué par un galvanomètre sensible, intercalé dans le circuit collecteur, et soustrait à l'influence du courant de repos au moyen d'un dispositif de compensation. L'équilibre étant réalisé en l'absence d'illumination, on projette sur la surface de l'échantillon l'image d'un filament rectiligne incandescent ; la déviation du galvanomètre qui en résulte est proportionnelle à la concentration de porteurs excédentaires présents au voisinage de la pointe collectrice.

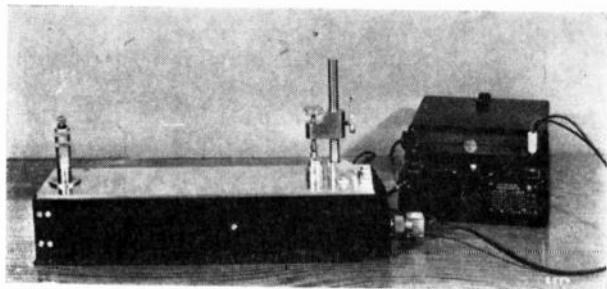


FIG. 7.

Dans un second dispositif, la pointe collectrice n'est pas polarisée et on connecte directement un galvanomètre entre la pointe et le germanium. Avec un

appareil sensible, le courant photoélectrique recueilli par le collecteur est suffisamment faible pour ne pas modifier la hauteur de la barrière de potentiel. Dans ce cas comme dans le précédent l'emploi d'un flux lumineux continu permet de s'affranchir du dispositif mécanique destiné à hacher la lumière, de l'amplificateur, et d'éliminer la gêne causée par le bruit de fond.

Le déplacement de la raie lumineuse est produit par déplacement du filament incandescent au moyen d'une vis micrométrique.

CONCLUSIONS.

Dans la technologie des semi-conducteurs, les mesures revêtent une très grande importance. Elles doivent permettre de contrôler les propriétés et les qualités du matériau à chacun des stades de sa préparation. Au stade final elles permettent de sélectionner les échantillons suivant les applications auxquelles ils sont destinés, les caractéristiques exigées du semi-conducteur variant considérablement suivant ces applications. Il en résulte qu'au cours de l'élaboration du matériau on est conduit à effectuer un grand nombre de mesures et de tests ; il est donc très souhaitable de disposer de méthodes pratiques, rapides, et non destructives, et c'est là que réside le très gros avantage de la méthode des pointes pour les mesures de résistivité et de constante de Hall.

VIE DE LA SOCIÉTÉ

Période du 16 novembre au 31 décembre 1954

RÉUNION DU BUREAU

Le Bureau s'est réuni le mercredi 15 décembre 1954 sous la présidence de M. G. RABUTEAU, *Président de la Société des Radioélectriciens*.

Etaient présents : MM. ANGOT, AUBERT, P. BESSON, BOUTHIL-
LON, BUREAU, CABESSA, DAVID, FROMY, de MARE, MATRAS,
MOULON, PARODI, PICAULT, PONTE, R. RIGAL.

Assistait à la séance : M. TESTEMALE.

Etaient excusés : MM. LESCHI, LIBOIS.

Au cours de cette séance les principaux points suivants ont été examinés :

- 1^o Attribution de la médaille René MESNY ;
- 2^o Election partielle du Conseil pour 1955 et préparation de l'Assemblée Générale ;
- 3^o Remplacement de M. L. CAHEN, Président du Comité de Rédaction, démissionnaire ;
- 4^o Publication dans l'Onde Electrique de la revue des activités de l'année. Cette question sera étudiée par le Comité de Rédaction ;
- 5^o Inauguration du Groupe de l'Est. — M. G. RABUTEAU a inauguré le 4 décembre 1954 le Groupe de l'Est créé à Nancy, par M. A. GOUDET.

Au cours de cette inauguration les membres du Groupe ont assisté aux communications de MM. LORACH et FAYARD et à la présentation de matériel de Télévision avec tube de prise de vue de très petites dimensions.

RÉUNION DU CONSEIL

Le Conseil s'est réuni le mercredi 29 décembre 1954 sous la présidence de M. G. RABUTEAU, *Président de la Société des Radioélectriciens*.

Etaient présents : MM. ANGOT, AUBERT, P. BESSON, BOUTHIL-
LON, BUREAU, CAHEN, CARBENAY, DANZIN, DAUPHIN, FLAMBARD,
FROMY, ICOLE, LIZON, LOCHARD, de MARE, MARIQUE,
MATRAS, PONTE, POTIER, PICAULT, PARODI, SOLLIMA, STEINBERG,
THIEN CHI, de VALROGER.

Assistait à la Séance : M. LIBOIS.

Excusés : MM. CABESSA, CHAVASSE, CHÉDEVILLE, P. DAVID,
FREYMAN, LESCHI, F. RAYMOND, R. RIGAL.

Au cours de la séance les principaux points suivants ont été examinés :

- 1^o Compte rendu de l'Inauguration du Groupe de l'Est ;
- 2^o Attribution de la médaille René MESNY.

Le Conseil décide à l'unanimité d'attribuer à M. P. DAVID la médaille René MESNY ; la remise de cette médaille aura lieu au cours de l'Assemblée Générale du 22 janvier 1955 ;

- 3^o Election partielle du Conseil pour 1955 et préparation de l'Assemblée Générale (rapport moral et rapport du Trésorier) ;
- 4^o Questions diverses.

Mr. MARIQUE signale « les Journées Internationales de Calcul Analogique » qui seront organisées par la Société Belge des Ingénieurs des Télécommunications et d'Electronique (S.I.T.E.L.) du 27 septembre au 1^{er} octobre 1955 à Bruxelles.

RÉUNIONS EN SORBONNE

Réunion du 27 octobre 1954.

Au cours de cette séance, présidée par M. G. Rabuteau, Président de la Société des Radioélectriciens, M. J. MOULON, Ingénieur des Télécommunications au C.N.E.T. fit un exposé sur « *Les Transistrons dans les Télécommunications* ».

Après avoir rappelé les propriétés élémentaires des transistrons à pointe et de jonction le conférencier étudie plus particulièrement les caractéristiques des transistrons de jonction : il met en évidence les principaux paramètres intervenant dans leur fonctionnement et notamment : la relation $I = bi$, liant le courant de sortie I au courant d'entrée i , l'influence du courant I_0 pour une tension nulle et sa variation en fonction de la température, M. MOULON décrit ensuite le schéma électrique équivalent du transistor et signale l'influence des capacités et du temps de transit des électrons. Les différents types de transistrons sont ensuite passés en revue : transistor au germanium, au silicium, transistor à barrière de surface, transistor $p m p$ etc.

Les transistrons peuvent remplacer les tubes électroniques dans de nombreux montages, tels que amplificateurs *BF* et *HF*, postes de radio. Leur emploi se caractérise par un taux de contre-réaction élevé.

Les transistrons se prêtent également à des applications qui leur sont propres et peuvent plus particulièrement être utilisés comme des éléments actifs dans des réseaux passifs (filtres, amplificateurs à impédance négative).

La communication de M. MOULON donne lieu à divers échanges de vue et remarques notamment, en ce qui concerne la loi de variation du bruit de fond en fonction de la fréquence et l'homogénéité des fabrications.

Réunion du samedi 11 décembre 1954.

Cette séance présidée par M. PICAULT, ancien Président de la Société des Radioélectriciens était consacrée à un exposé de MM. NGUYEN THIEN CHI et J. VERGNOLLE sur « *Les condensateurs électrolytiques au tantale* ».

Les condensateurs au tantale existent sous deux formes : type bobiné et type à anode frittée. Le type à anode frittée est composé d'une anode massive en tantale frittée. Cette anode est poreuse et offre ainsi une surface réelle très importante. La cathode est constituée par le boîtier du condensateur lequel est rempli par l'électrolyte constitué par de l'acide sulfurique ou du chlorure de lithium. Le conférencier étudie ensuite la polarisation anodique du tantale et donne des indications sur les performances dynamiques, la fabrication et les mesures des caractéristiques. Il cite quelques exemples de réalisation dans les Laboratoires CSF et insiste sur l'intérêt du condensateur marmite.

On arrive à réaliser des condensateurs capables d'une charge de 1 à 3 microcoulombs par mm^2 .

Les condensateurs électrolytiques au tantale sont avantageusement employés ou utilisés chaque fois qu'il faut prévoir des condensateurs de très petites dimensions devant couvrir une large plage de température. Ces condensateurs sont plus particulièrement intéressants dans les circuits électroniques équipés de transistrons et fonctionnant sous des tensions relativement basses.

Réunion du samedi 18 décembre 1954.

Cette réunion constituait la première d'une série de séances communes avec le Comité National Français de Radioélectricité

Scientifique consacrées aux comptes rendus de la XI^e Assemblée Générale de l'Union Radio Scientifique Internationale (U.R.S.I.) tenue en août 1954 à La Haye.

Le programme de cette réunion était le suivant :

Introduction Générale, par M. G. LEHMANN, Président du Comité National Français de Radioélectricité Scientifique.

Commission VII « Electronique » par M. G. LEHMANN.

Commission I « Mesures », par M. P. ABADIE, Ingénieur en chef des Télécommunications.

Commission VI « Théorie de l'Information, Circuits et Antennes », par M. A. ANGOT, Ingénieur Militaire en Chef des Télécommunications d'Armement.

et par M. G. FOLDES, Ingénieur au L.N.R.

Les autres séances auront lieu aux dates suivantes :

Samedi 22 janvier 1955 à 17 h.

Commission V « Radio-Astronomie », par M. M. LAFFINEUR, Président de la Commission V de l'U.R.S.I.

Samedi 12 février 1955 à 17 h.

Commission IV « Atmosphériques », par M. R. RIVAUT, Ingénieur au L.N.R.

Commission II « Troposphère », par M. J. VOGÉ, Ingénieur des Télécommunications.

Commission III « Ionosphère », par M. D. LEPECHINSKY, Ingénieur en Chef au Bureau Ionosphérique Français.

Résumés des communications de la séance du 18 décembre 1954 :

Communication de M. G. Lehmann, Président du Comité National Français de Radioélectricité Scientifique.

Le conférencier rappelle d'abord que le R. P. LEJAY a été réélu Président de l'U.R.S.I., M. M. LAFFINEUR, Président de la « Commission Internationale de Radioastronomie » et M. B. DECAUX, Président de la « Commission Mesures ».

L'U.R.S.I. fondée en 1920 a pour mission de faciliter et de diffuser les études relatives aux propriétés scientifiques fondamentales de radioélectricité nécessitant une coopération internationale.

M. G. LEHMANN donne ensuite des indications sur l'organisation de l'U.R.S.I., les buts et les méthodes de travail des Assemblées Générales réunies jusqu'à présent tous les deux ans et qui ne seront plus convoquées à partir de 1954 que tous les trois ans. La prochaine Assemblée Générale aura lieu en 1957 aux U.S.A.

Dans la deuxième partie de son exposé, M. G. LEHMANN traite des travaux de la Commission VII « Electronique ». Il résume rapidement les différents sujets traités dans les domaines de l'Electronique des solides, de l'Electronique des gaz, de l'émission électronique et de la production des ondes centimétriques et millimétriques.

Dans ce dernier domaine un exposé de M. GUÉNARD sur la situation des tubes générateurs d'hyperfréquence (1 000 à 40 000 Mc/s) montra que la France était à la pointe du progrès et que certaines de ses études étaient plus près du stade de la réalisation que des études similaires faites aux Etats-Unis.

M. ORTUSI fit un parallèle entre les cathodes à oxyde et les semi-conducteurs mixtes.

Il naît de nouveaux types de cathodes à réserve de matière émissive qui vient affluer à la surface en lieu et temps utile.

Communication de M. P. Abadie :

M. P. ABADIE résume les travaux de la Commission I « Mesures » notamment sur les points suivants :

Mesures de Fréquences :

Les progrès accomplis permettent des mesures à 10^{-9} près dans la gamme de 100 kc/s, et à 10^{-4} près dans la gamme des ondes millimétriques. Un étalon de fréquence atomique permet d'obtenir une précision de 5×10^{-10} . On espère atteindre en utilisant le cæsium 10^{-11} et même 10^{-12} .

Vitesse des ondes électromagnétiques dans le vide.

La vitesse retenue à l'Assemblée Générale de l'U.R.S.I. de SYDNEY était de $299\,792 \pm 2$ km/s.

De nouvelles mesures effectuées ont confirmé le bien fondé de cette valeur qui continue à être recommandée.

Emission de fréquences étalons et de faisceaux horaires.

Mesure des puissances en ondes centimétriques.

Les précisions de mesures pour des puissances faibles est de l'ordre de 2 à 5 %.

Il est recommandé aux membres de l'Union de comparer leurs appareils étalons.

Mesures de champ, mesures de bruits et mesures d'impédances.

Communication de M. Angot sur les travaux de la Commission VI « Théorie de l'Information » — « Circuits et Antennes ».

M. ANGOT signale tout d'abord que les sujets relevant de la Commission VI étaient très variés et indépendants les uns des autres ce qui a conduit à créer des sous-commissions travaillant indépendamment.

Il énumère ensuite les sujets des différents exposés en mettant en évidence les travaux présentés par les représentants français. M. LOEB sur les servo mécanismes considérés comme filtres non linéaires, et l'étude du codage à deux symboles.

M. LAPOSTOLLE sur la propagation des ondes électromagnétiques le long d'une hélice.

M. ROBIN sur l'antenne biconique.

M. POINCELOT sur la répartition du courant le long d'une antenne cylindrique, sur l'inexistence de l'Onde de sol Sommerfeld, et sur la notion de fréquence instantanée (en collaboration avec M. ROBIN). Il signale enfin l'orientation des futures recherches.

Il résume les travaux de M. BLANC-LAPIERRE ainsi que ceux de MM. BOURRASSIN et COLOMBO sur les systèmes de transmissions de télévision à bandes latérales asymétriques.

Communication de M. G. Foldès : sur les travaux de la Commission mixte IV et VI.

M. FOLDES examine les travaux relatifs au bruit atmosphérique, cause de limitation des systèmes radioélectriques et à la recherche de paramètres susceptibles de caractériser son action.

Il développe plus particulièrement le rapport présenté par M. BLANC-LAPIERRE sur la théorie complète d'un bruit complexe.

Enfin il signale les différentes techniques de mesures avant détection et après détection pratiquées dans les divers pays.

ACTIVITÉ DES SECTIONS

Troisième Section : « Electroacoustique ».

Réunion du 13 décembre 1954. M. CHAVASSE, Président du Groupement des Acousticiens de Langue Française, GALT, avait invité les membres de la troisième Section dont il est Président à assister aux deux exposés suivants :

M. R. G. BUSNEL « **Caractères Physiques des signaux acoustiques provoquant des réactions sur les orthophères** ».

M. P. BUGARD : « **Action biologique des bruits intenses** ».

Cinquième Section « Hyperfréquences » et sixième Section « Electronique ».

Les cinquième et sixième Sections avaient organisé, en commun avec la huitième Section de la Société Française des Electriciens (S.F.E.) et la Société Française des Ingénieurs et Techniciens du Vide, une séance le 17 décembre 1954, présidée par M. SUEUR et au cours de laquelle M. P. GUÉNARD fit un exposé sur les

«Développements récents des tubes à ondes progressives».

Le conférencier décrit huit tubes en les illustrant avec de nombreuses projections. Cette communication donne lieu à diverses remarques notamment en ce qui concerne les tubes TPO 851 et TPO 920 et l'influence du procédé de focalisation sur le bruit de fond. Des précisions sont également demandées sur les taux d'ondes stationnaires et la bande passante.

Distinction Honorifique :

Nous avons le plaisir d'annoncer que notre Secrétaire Général, M. J. MATRAS, Ingénieur Général des Télécommunications a été nommé Chevalier dans l'ordre de la Légion d'Honneur.

NÉCROLOGIE**LE COLONEL HENRY BÉDOURA**

Officier et Ingénieur, aimé et estimé de tous, le Colonel Henry BÉDOURA, qui fut Vice-Président de la Société des Radioélectriciens de 1937 à 1939, a été arraché à l'affection des siens le 8 mai 1954, après une brève mais implacable maladie. Chef de la Section de Radioélectricité de l'Ecole Supérieure d'Electricité, puis Sous-Directeur de l'Ecole et Directeur des Etudes, il était peu d'ingénieurs de nos spécialités qui n'aient eu affaire à lui, et qui n'aient apprécié son accueil toujours si bienveillant et si cordial. Mais, sans doute, beaucoup d'ingénieurs n'ont-ils que soupçonné la brillante carrière militaire du Colonel BÉDOURA, et beaucoup de ses anciens compagnons d'arme ignorent-ils encore les services que leur camarade a rendus, depuis 1929, pour la formation des ingénieurs radioélectriciens. Au cours de cette double carrière, que l'on évoquera ici, le Colonel BÉDOURA a toujours su se placer au rang des meilleurs, et justifier la confiance que, sur le plan militaire comme sur le plan technique, ses chefs avaient placée en lui.

Né à Morlanne, dans les Basses-Pyrénées, en 1878, Henry BÉDOURA, après des études secondaires, avait été reçu à Saint-Cyr



Le Colonel Henry BÉDOURA

Ancien Vice-Président de la Société des Radioélectriciens

en 1899. Sorti dans l'Infanterie, il avait franchi les divers grades jusqu'à celui de Capitaine et, au début de la guerre de 1914, commandait la 1^{re} Compagnie du 119^e Régiment d'Infanterie, en garnison à Paris.

Il aimait rappeler son baptême du feu, aux premiers jours du mois d'août, lorsque, après avoir été magnifiquement accueilli

par les Belges, il avait dispersé une colonne allemande qui, tambours en tête, s'avancait sur le remblai du chemin de fer, près de la gare de Charleroi. Un mois à peine après, il recevait sa première blessure.

En 1915, il est à nouveau touché, et les attaques de septembre auxquelles il participe lui valent sa première citation, et le grade de Chevalier de la Légion d'Honneur.

Nommé Chef de Bataillon en 1916, une seconde fois cité, il est à Verdun, puis, en 1917, au Chemin des Dames. En 1918, il est encore légèrement blessé au cours de l'offensive qui devait nous conduire à la victoire, dans une opération qui, des abords de l'Aisne le mène jusqu'à Sissonne, où sa progression est arrêtée par l'Armistice. Deux citations venaient encore ajouter des palmes à sa Croix de Guerre. Il rentre à Paris, seul Officier survivant parmi ceux du 119^e R. I. qui avaient quitté la Capitale cinquante-deux mois auparavant.

En 1920, promu Officier de la Légion d'Honneur, il est chargé d'organiser les cours d'élèves-officiers du Groupe d'Armée Fayolle.

Dès 1923, le Général FERRIÉ avait pressenti les services que le Commandant BÉDOURA pouvait rendre à la Radiotélégraphie Militaire, et, sur son initiative, l'Ecole Supérieure d'Electricité l'avait accueilli dans sa Section de Radioélectricité, il en était sorti diplômé en 1925, avec les félicitations du Jury.

Mis ensuite à la disposition du Général Commandant le théâtre d'opérations du Maroc, il prit le commandement du 39^e Régiment de Tirailleurs Algériens, et participa à la campagne du Rif, aboutissant à la reddition d'Abd-el-Krim. Affecté en 1927 à la direction du Centre d'Information de Meknès, il occupa ensuite les fonctions de Chef d'Etat-Major du Général GIRAUD, puis il revint l'année suivante, en France comme Chef d'Etat-Major du Général FERRIÉ, sans vouloir d'ailleurs quitter son arme d'origine, l'Infanterie, dont il était si fier.

C'est alors qu'en 1929 commença sa collaboration avec l'Ecole Supérieure d'Electricité, où il remplit tout d'abord les fonctions de Chef du Service des Travaux Pratiques de la Section de Radioélectricité, ce qui le conduisit à étudier et à réaliser dans les bâtiments mis à la disposition de l'Ecole par la Radiotélégraphie militaire, aux Invalides, des séries de manipulations sur les hautes fréquences, les ondes radioélectriques, etc... En 1936, il fut nommé Directeur des Etudes de cette Section et eut à mettre au point les programmes de l'enseignement, à préparer les emplois du temps, à étudier les projets d'extension des laboratoires.

Ayant pris sa retraite de Colonel en 1937, quelques mois après avoir été promu Commandeur de la Légion d'Honneur, il devint Sous-Directeur de l'Ecole et Directeur des Etudes, ce qui le conduisit à s'installer dans les nouveaux bâtiments de Malakoff, et à s'occuper de l'enseignement donné à la Section Normale, devenue depuis la Division « Electricité ». Il participa efficacement à tous les échanges de vues poursuivis depuis 1937 pour maintenir les études au niveau nécessité par les développements de la technique, pour créer certains cours, pour accroître l'importance des travaux pratiques, etc...

Mobilisé en 1939 au Commandement Supérieur des Transmissions, il eut à s'occuper de problèmes d'organisation. Après les événements de 1940, il contribua à la mise en place de la Section de Radioélectricité à Lyon, car l'ennemi avait interdit tout enseignement des techniques radioélectriques dans la zone occupée. Administrateur de la Société « Les Laboratoires Radioélectriques » — dont il devait devenir plus tard Président, et dont le Directeur, Mario NIKIS, Ingénieur Radio E. S. E., devait être arrêté avec nombre de ses ingénieurs et mourir en déportation — il contribua à réaliser clandestinement du matériel de télécommunication et de radionavigation destiné à la Résistance, et, en particulier, à son ami le Colonel LABAT, Ingénieur Radio E. S. E., arrêté lui aussi et abattu par les Allemands au camp de Struthof.

Les élèves des promotions qui se succédèrent à Malakoff pendant l'occupation, ne peuvent oublier avec quelle compréhension il les aida et les conseilla au milieu des multiples tracasseries des services de la main-d'œuvre, et que c'est beaucoup grâce à lui qu'un grand nombre d'entre eux purent éviter un départ en Allemagne.

Dès sa sortie de l'E. S. E., en 1925, le Commandant BÉDOURA avait adhéré à la Société des Radioélectriciens et, en 1936, il fut appelé à faire partie de son Conseil. Il fut élu Vice-Président, de 1937 à 1939, et put ainsi faire bénéficier la Société de ses

connaissances en Radioélectricité, et de ses qualités d'organisateur et d'administrateur.

Ce n'est certainement pas sans regrets que le Colonel BÉDOURA fixa lui-même à la fin de l'année scolaire 1952-53 le terme de son activité dans l'enseignement, et la manifestation dont il fut alors l'objet de la part de tout le personnel de l'Ecole Supérieure d'Electricité, lui montra, s'il en était besoin, combien il avait su se créer d'amis.

Membres du Corps enseignant, Ingénieurs et Agents de l'Ecole, Elèves des deux Divisions, tous souhaitaient le voir profiter pendant de longues années d'une retraite si bien gagnée, et le retrouver d'ailleurs souvent à l'Ecole à l'occasion de diverses manifestations annuelles. C'est avec une stupeur attristée qu'ils apprirent sa disparition, en mai 1954.

A Madame BÉDOURA et à ses Enfants, la Société des Radioélectriciens renouvelle ses bien vives condoléances, et tient à redire combien le souvenir de son Ancien Vice-Président restera gravé dans la mémoire de tous ses Membres.

INFORMATIONS

Prix Général Ferrié

Le Comité National Ferrié rappelle qu'il a créé un *prix annuel* de 100 000 F (cent mille) destiné à récompenser un jeune Français ayant présenté une étude de nature à contribuer au progrès de la Radioélectricité.

Le Jury est composé de hautes personnalités civiles et militaires.

Les candidats doivent avoir accompli leur service dans l'Armée des Transmissions, ou faire partie des Services de Transmissions de la Défense Nationale, et présenter — avant le 1^{er} avril de chaque année — un travail effectué dans un délai maximum de dix ans après la date normale de libération du Service.

Pour tous renseignements, s'adresser au Comité National Ferrié, 23, rue de Lubeck, à Paris (16^e).

DOCUMENTATION. — La Société des Radioélectriciens a reçu gracieusement pour sa bibliothèque les ouvrages suivants :

— « L'Ingénieur du Son » en Radiodiffusion, Cinéma, Télévision par M. V. JEAN-LOUIS, Ingénieur E. E. M. I. Préface du Général LESCHI, Directeur des Services Techniques de la Radiodiffusion Télévision Française (Editions Chiron).

— « Electricité et Optique ». (La lumière et les théories électrodynamiques) par H. POINCARÉ. Edition revue et complétée par MM. Jules BLONDIN, agrégé de l'Université et Eugène NECULCEA, licencié ès-Sciences. (Gauthier-Villars, Editeur).

— « Tubes d'émission » (l'emploi de pentodes, de tétrodes et de triodes dans les montages d'émission) par J. P. HEYBOER et P. ZIJLSTRA. (Bibliothèque Technique Philips).

— « Les fonctions de Bessel et leurs applications en physique », par Mr. G. GOUDET. (Masson et Cie, Editeurs).

— « Memorial des Sciences Physiques ». (Gauthier-Villars, Editeur).

Fascicule LVII. — Conductibilité électrique des lames minces, par M. A. BLANC-LAPIERRE et M. M. PERROT, Professeurs à la Faculté des Sciences d'Alger.

Fascicule LVIII. — Propriétés magnétiques des lames métalliques minces, par M. A. COLOMBANI, Professeur à la Faculté des Sciences de Caen.

Fascicule LIX. — Les aspects modernes de la cryométrie, par M. Y. DOUCET, Maître de Conférences à la Faculté des Sciences de Dijon.

NOUVEAUX MEMBRES

MM.	Présentés par MM.
AUDEBERT Michel, élève à l'E. S. E. (division Radio-électricité et électronique)	GAUSSOT et Mme HUTER.
AZAM Jean, élève à l'E. S. E. (division Radio-électricité et électronique)	KEVORKIAN et G. AZAM.
BAUDET Pierre, élève à l'Ecole de Radioélectricité de Bordeaux	CAU et COMBE D'ALMA.
BAUDOIN Robert, Ingénieur E. S. E. à la Sté « Le Matériel Téléphonique »	CABESSA et LOEFFLER.
BEAUSSAC André, Officier d'Artillerie coloniale, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.
BEAUSSIER Jacques, attaché au C. N. R. S., Licencié ès-Sciences, Ingénieur E. S. E.	DAVID et GRIVET.
BÉDOURA Jacques, Ingénieur Militaire des Télécommunications	P. BESSON et DAUPHIN.
BERRADA Abderrazak, élève à l'E. S. E. (division Radioélectricité et Electronique)	Mme HUTER et DAUPHIN.
BOURDAROT Guy, Jean, Louis, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.
BOURGUIGNON Jean, Paul, élève à l'E. S. E. (division Radioélectricité et Electronique)	GAUSSOT et Mme HUTER.
BOVAGNE Henri, Pierre, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et Mme HUTER.
BROSSARD Pierre, Claude, élève à l'Ecole de Radio-électricité de Bordeaux	CAU et COMBE D'ALMA.
CABANNES René, Jean, Louis, élève à l'E. S. E. (division Radioélectricité et Electronique)	GAUSSOT et Mme HUTER.
CALDERO Jean-Claude, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.
CANAC René, Jean, Joseph, Lieutenant de Vaisseau, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.
CROITORU Zicov, Ingénieur E. S. E., Ingénieur au Laboratoire Central des Industries Electriques (L. C. I. E.)	GAUSSOT et Mme HUTER.
DELPHIN Pierre, Paul, élève à l'E. S. E. (division Radioélectricité et Electronique)	GAUSSOT et Mme HUTER.
DEMAN Pierre, Jehan, Ingénieur des P. T. T., Direction des Services Radioélectriques	LIROIS et PIACE.
DUCHET Jean, élève à l'E. S. E. (division Radio-électricité et Electronique)	Mme HUTER et DAUPHIN.
Mlle DUFAURE DE LA JARTE Simone, Ingénieur E. S. E. à la Sté « La Radiotechnique »	DAUPHIN et GAUSSOT.
FABÈRES Pierre, Ingénieur Radioélectricien E. R. B.	CAU et COMBE D'ALMA.
GEY Albert, Armand, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.
GHAZALA Nidhat, Youssef, Ingénieur de l'Ecole Polytechnique du Caire, Ingénieur Radio E. S. E., Ingénieur à la Compagnie J. Visseaux	HEINISCH et AROUETE.
GLAUDE Vital, Max, Marie, Licencié ès-Sciences, Ingénieur Radio E. S. E.	DAUPHIN et GAUSSOT.
GRIMALDI François, Ingénieur Radio E. S. E.	Mme HUTER et GAUSSOT.

MM.	Présentés par MM.	MM.	Présentés par MM.
GUÉRINEAU Jo, Ingénieur E. I. N., Directeur des Ets J. Guérineau et Fils	GOUDET et GUDEFIN	POËLLE Bernard, André, Ingénieur Civil E. N. S. T. à la Sté Française Radioélectrique	R. RIGAL et FRANÇOIS.
HARRISON Charles, Lieutenant-Colonel, adjoint de l'attaché Militaire à l'Ambassade des U. S. A. à Paris	AUBERT et BRAILLARD.	PRUVOT François, Xavier, élève à l'Ecole Centrale de T. S. F.	QUINET et CHRÉTIEN
HAUSEUX Bernard, élève à l'E. S. E. (division Radioélectricité et Electronique)	Mme HUTER et GAUSSOT.	RICKLIN Philippe, Michel, élève à l'Ecole de Radioélectricité de Bordeaux	CAU et COMBE D'ALMA.
HUG Gilbert, André, Agent Technique	CAZALAS et NEUVY.	ROUGERON Jacques, Agent technique Radio ...	PALAZO et POINCELOT.
JOLY Léon, Pierre, Ingénieur E. C. P. à la Sté « Le Matériel Téléphonique »	RABUTEAU et LIZON.	ROULIN Jean, élève à l'E. S. E. (division Radio-électricité et Electronique)	DAUPHIN et GAUSSOT.
KAHN Alain, Georges, Ingénieur E. P. C. I., Licencié ès-Sciences, Chef des Services au Laboratoire Industriel d'Electricité	GI CHARLES et MASSELIN.	ROUX Guy, Agent Technique à la Sté Les Lignes Télégraphiques et Téléphoniques	R. RIGAL et CHIRON.
LAFFAIRE Michel, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.	SAINT-ETIENNE Jean, Pierre, élève à l'E. S. E. (division Radioélectricité et Electronique)....	Mme HUTER et DAUPHIN.
LAISNE André, élève à l'E. S. E. (division Radio-électricité et Electronique)	DAUPHIN et GAUSSOT.	SANDER André, Emile, Ingénieur E. C. P., élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.
LASNERET, Ingénieur E. S. E. au C. R. T. Paris E. de F.	P. BESSON et VARRET.	SAUVAN Jean-Pierre, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.
LISIMAQUE Jean, Marcel, Pierre, Ingénieur E. P. C. I. à la Sté « Le Matériel Téléphonique » ..	LOEFFLER et CABESSA.	VERDIER Pierre, Raymond, Ingénieur de l'Ecole spéciale des Travaux Publics (section mécanique et électricité), Ingénieur au Laboratoire « Hyperfréquences » de la Sté L. E. P.	CAYZAC et ASTOR.
MARTIN Claude, Louis, Ingénieur Radio E. S. E.	DAUPHIN et GAUSSOT.		
MENACHE Henri, élève à l'Ecole de Radioélectricité de Bordeaux.....	CAU et COMBE D'ALMA.		
NOEL Jean, Capitaine d'Artillerie Coloniale, élève à l'E. S. E. (division Radioélectricité et Electronique)	DAUPHIN et GAUSSOT.		

OFFRES D'EMPLOIS

0.46. — Agents Techniques 2^e ou 3^e échelon radio-électriciens spécialement réceptions ondes moyennes et asservissements à distance pour appareils de navigation aérienne, et ayant au moins 3 ans de pratique dans l'Industrie. Ecrire à la Société qui transmettra.

Pour calculer et réaliser vos circuits de télévision, émission et réception...
Pour calculer et réaliser vos circuits de radar...

l'ouvrage d'HENRI ABERDAM

AMPLIFICATEURS A LARGE BANDE

est indispensable dans votre bureau d'études

- I. — AMPLIFICATEURS VIDEO-FRÉQUENCE. Calcul des corrections d'amplitude et de phase en fréquences basses, en fréquences élevées. Selfs-série, capacités, filtres passe-bas, compensations mixtes, quadripôles en pi et en T.
- II. — AMPLIFICATEURS A FRÉQUENCE INTERMÉDIAIRE. Circuits classiques, couplés, accordés, désaccordés, accords décalés. Circuits spéciaux. Construction, choix des tubes, des circuits, faible variation de gain. Contre-réaction, étude, calculs, doublets, triplets. Tubes à grille à la masse, temps de transit, applications,
- III. — COMPLÉMENTS. Conservation de la largeur de bande. Calculs sur les circuits couplés. Calculs exacts d'impédances avec accords décalés. C. R. sur circuit anti-résonnant. Etudes complètes d'amplificateurs par méthode algébrique et par méthode graphique.

Un volume de 212 pages 15 X 24 cm, broché : 2.770 fr. franco de port ; relié pleine toile : 3.100 fr. franco de port.

ÉDITIONS CHIRON, 40, Rue de Seine, PARIS-6^e - C. C. P. PARIS 53.35

SFME
PANTIN

TRANSFORMATEUR
d'ISOLEMENT
120.000 volts

VILLAUME (S.I.E.U.-I.G.F.)

pour balisage d'antenne

62, RUE DENIS-PAPIN ★ Tél. NORD. 47-62

F. GUERPILLON & C^{ie}
Société à responsabilité limitée au capital de 27 millions

APPAREILS DE MESURES ÉLECTRIQUES

BUREAUX
ET
ATELIERS
64, Av. A. Briand
MONTRouGE
(Seine)
Tél. ALEsia + 29.85
(3 lignes)

Pour la BELGIQUE :

S^{TE} BELGE GUERPILLON
11, Rue Bara, à BRUXELLES - MIDI — Tél. 21-06-01

APPAREILS
DE TABLEAUX DE CONTROLE ET DE LABORATOIRE
APPAREILS SPÉCIAUX TROPICALISÉS

Nouveau catalogue franco sur demande
PUBL. RAPP

du nouveau!

LE UGON 2
AUX PERFORMANCES AMÉLIORÉES

RELAIS SUBMINIATURES UGON
BREVETÉS S.G.D.G.

- SENSIBILITÉ 2 milliwatts
- POUVOIR DE COUPURE 24 v. 0,5A
- TROPICALISÉ (Soudures métal-verre)
- MONTAGE A VOLONTÉ sur support subminiature rond normal ou fils à souder

LE PROTOTYPE MÉCANIQUE - 16, bis, Rue Georges-Pitard
Paris-15^e — VAU 38-03
PUBL. RAPP

V. JEAN-LOUIS

L'INGÉNIEUR DU SON
Radiodiffusion — Cinéma — Télévision

I. — Acoustique Psychotechnique
II. — La prise de Son.
III. — Les fonctions.

Un ouvrage de 300 pages — illustré

Editions CHIRON, 40, rue de Seine, PARIS-6^e

POTENTIOMÈTRES

- GRAPHITÉS OU BOBINÉS
- ÉTANCHES OU STANDARDS
- A PISTE MOULÉE

Variohm

Rue Charles-Vapereau, RUFIL-MALMAISON (S.-&-O.) - Tél. MAL. 24-54
Publ. RAPP

L'ONDE ÉLECTRIQUE

35^e ANNÉE. N° 335

FÉVRIER 1955

PRIX : 250 FRANCS

REVUE DE LA SOCIÉTÉ DES RADIOÉLECTRICIENS
ÉDITIONS CHIRON, 40, RUE DE SEINE, PARIS - 6^e



LE RMD 34
PAR DE
NAVIGATION MARITIME
DES LABORATOIRES DERVEAUX

LIRE DANS CE NUMÉRO

Enregistrement magnétique et machines à calculer, F. RAYMOND — Asservissements avec chyratrons, P. BONNET — Coefficient de transfert d'ondes, F. RAYMOND — Résistivité et loi de Hall, LAPLUME — Ondes électromagnétiques polarisées, M. BOUX — Mesure de phases pour guidage, J. ZACKHEIM — Multiplieur à découpage temporel, M. LILAMAND — Correcteur automatique de fréquence, J. CAYZAC — Modulation d'espace, G. POTIER — Circuit déclencheur ferro-résonant, J. SAN-
TESMAKES

TOUS LES TUBES
CATHODIQUES
POUR LA TV ET
L'INDUSTRIE



*You Can be Sure... if it's
Westinghouse*



TOUS LES TUBES
SPÉCIAUX

RADARS • KLYSTRONS
MAGNÉTRONS • PENCILS
TUBES • TRANSISTORS
DIODES, TRIODES, V.H.F.
ET LES
INSTRUMENTS DE MESURE
DES GRANDES MARQUES

GRANDE RAPIDITÉ
DE LIVRAISON

TÉLÉPHONE

ANJOU 85-00



YOUNG ELECTRONIC

9 et 11, RUE ROQUÉPINE - PARIS-8^E - TÉLÉPHONE : ANJOU 85-00

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 2500 F
ETRANGER 2800 F

ÉDITIONS CHIRON

40, Rue de Seine — PARIS (6°)

C. C. P. PARIS 53-35

Prix du

numéro :

250 francs

Vol. XXXV

FÉVRIER 1955

Numéro 335

SOMMAIRE

	Pages
Considérations sur l'enregistrement magnétique dans le domaine des machines à calculer	F.H. RAYMOND.... 89
Étude du comportement dynamique des asservissements comportant des thyatrones	P. BONNET..... 97
Sur la réalisation d'un coefficient de transfert donné ou les montages d'amplificateurs à réaction.....	F.H. RAYMOND.... 105
Bases théoriques de la mesure de la résistivité et de la constante de Hall par la méthode des pointes.....	J. LAPLUME..... 113
Ondes électromagnétiques polarisées elliptiques et circulaires	M. BOUX..... 126
Utilisation de la mesure de phases dans un système de guidage et de localisation des mobiles.....	J. ZAKHEIM..... 137
Multiplieurs à découpage temporel	Madame LILAMAND. 142
Correcteur automatique de fréquence pour émetteur de télécommunications sur ondes centimétriques	J. CAYZAC..... 151
La modulation d'espacement d'impulsions.....	G. X. POTIER..... 159
Circuit déclencheur ferorésonnant parallèle à réaction.....	J. G. SANTESMASES M. R. VIDAL..... 165
Vie de la Société — Assemblée Générale du 22 janvier 1955	174

Sur la couverture :

Le Radar de Marine RMD 30 des Laboratoires Derveaux spécialement conçu pour les navires de faible et de moyen tonnages. — Laboratoires Derveaux 6, rue Jules-Simon — Boulogne-sur-Seine (Seine).

Les opinions émises dans les articles ou comptes rendus publiés dans L'Onde Electrique n'engagent que leurs auteurs

SOCIÉTÉ DES RADIOÉLECTRICIENS

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

2° 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° Service gratuit de la revue mensuelle *L'Onde Electrique* ;

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

3° Visites de diverses installations radioélectriques : stations d'émission et de réception, postes de navires et d'avions, laboratoires, radio-phares, expositions, studios, etc ;

4° Renseignements divers (joindre un timbre pour la réponse).

COTISATIONS

1° Membres titulaires, particuliers 1 600 F.

— sociétés ou collectivités 10 000 F. } au gré
25 000 F. } de la société
ou 50 000 F. } ou collectivité

2° Membres titulaires, âgés de moins de vingt-cinq ans, en cours d'études en France..... 800 F

Les membres de ces deux catégories, résidant à l'étranger, doivent verser en plus un supplément pour frais postaux de 400 F.

3° Membres à vie :

Les particuliers, membres titulaires peuvent racheter leur cotisation annuelle par un versement unique égal à dix fois le montant de cette cotisation soit 16 000 F.

4° Membres donateurs :

Seront inscrits en qualité de donateurs, les membres qui feraient don à la Société, en plus de leur cotisation, d'une somme d'au moins 5 000 F.

5° 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 15 000 F

Adresser la correspondance administrative et technique, et effectuer le versement des cotisations au Secrétariat de la Société des Radioélectriciens

SOCIÉTÉ DES RADIOÉLECTRICIENS

10, Avenue Pierre-Larousse, Malakoff (Seine)

Tél. ALÉSIA 04-16 — Compte de chèques postaux Paris 697-38

CHANGEMENTS D'ADRESSE : Joindre 20 francs à toute demande

SOCIÉTÉ DES RADIOÉLECTRICIENS

PRÉSIDENTS D'HONNEUR

† R. MESNY (1947) — † H. ABRAHAM (1947)

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 VALBREUZE, 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. BÉREAU, Agrégé de l'Université, Docteur es-Sciences.
 1938 P. FRANCK, Ingénieur général de l'Air.
 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, Secrétaire perpétuel de l'Aca-
 démie des Sciences.
 1949 M. PONTE, Directeur Général Adjoint de la Cie Gle de T.S.F.
 1950 P. BESSON, Ingénieur en Chef des Ponts et Chaussées.
 1951 Général LESCHI, Directeur des Services Techniques de la
 Radiodiffusion — Télévision Française.
 1952 J. de MARE, Ingénieur Conseil.
 1953 P. DAVID, Ingénieur en chef à la Marine.
 1954 G. RABUTEAU, Directeur Général de la Sté « Le Matériel
 Téléphonique ».

BUREAU DE LA SOCIÉTÉ

Président (1955)

M.H. PARODI, Membre de l'Institut, Professeur au Conservatoire National
des Arts et Métiers.

Président désigné pour 1956 :

M.R. RIGAL, Ingénieur Général des Télécommunications.

Vice-Présidents :

MM. E. FROMY, Directeur de la Division Radioélectricité du L.C.I.E.
 A. ANGOT, Ingénieur militaire en Chef, Directeur de la Section
 d'Etudes et de Fabrications des Télécommunications.
 C. BEURTHERET, Ingénieur en Chef à la C. F. T. H.

Secrétaire Général :

M. J. MATRAS, Ingénieur Général des Télécommunications.

Trésorier :

M. R. CABESSA, Ingénieur à la Société L.M.T.

Secrétaires :

MM. R. CHARLET, Ingénieur des Télécommunications.
 J.M. MOULON, Ingénieur des Télécommunications
 P. DEMAN, Ingénieur des Télécommunications.

SECTIONS D'ÉTUDES

N°	Dénomination	Présidents	Secrétaires
1	Etudes générales.	Colonel ANGOT	M. TOUTAN.
2	Matériel radioélectr.	M. LIZON	M. GAMET.
3	Electro acoustique.	M. CHAVASSE.	M. POINCELOT
4	Télévision.	M. MALLEIN	M. ANGEL
5	Hyperfréquences.	M. WARNECKE	M. GUÉNARD
6	Electronique.	M. CAZALAS.	M. PICQUENDAR
7	Documentation.	M. CAHEN.	Mme ANGEL.
8	Electronique appliq.	M. RAYMOND.	M. LARGUIER.

GROUPE DE GRENOBLE

Président. — M.-J. BENOIT, Professeur à la Faculté des Sciences
 de Grenoble, Directeur de la Section de Haute Fréquence à
 l'Institut Polytechnique de Grenoble.

Secrétaire. — M. J. MOUSSEIGT, Chef de Travaux à la Faculté des
 Sciences de Grenoble.

GROUPE D'ALGER

Président. — M. H. CORBERY, ingénieur en chef à l'Electricité
 et Gaz d'Algérie.

Secrétaire. — M. P. CACHON, Assistant à la Faculté des Sciences
 d'Alger.

GROUPE DE L'EST

Président. — G. GOUDET, Directeur de l'Ecole Nationale Supérieure
 d'Electricité et de Mécanique de Nancy.

Secrétaire. — E. GUDEFIN, Assistant à l'E. N. S. E. M.

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).

CONSEIL DE LA SOCIÉTÉ DES RADIOÉLECTRICIENS

- J. BOULIN, Ingénieur des Télécommunications à la Direction des
 Services, Radioélectriques.
 F. CARBENAY, Ingénieur en Chef au Laboratoire National de Radio-
 électricité.
 G. CHEDEVILLE, Ingénieur Général des Télécommunications.
 R. FREYMAN, Professeur à la Faculté des Sciences de Rennes.
 J. MARIQUE, Secrétaire Général du C.C.R.M. à Bruxelles.
 F.H. RAYMOND, Directeur de la Société d'Electronique et d'Auto-
 matisme.
 J.L. STEINBERG, Maître de Recherches au C.N.R.S.
 L. DE VALROGER, Directeur du Département Radar-Hyperfréquen-
 ces de la Cie Française Thomson-Houston.
 I. ICOLE, Ingénieur en chef des Télécommunications, Chef du Dé-
 partement Faisceaux-Hertziens, Direction des Lignes Souterraines
 à Grande Distance.
 J. LOCHARD, Lieutenant Colonel, Chef des Services Techniques du
 Groupe de Contrôle Radioélectrique.
 N'GUYEN THIEN CHI, Chef de Département à la Cie Gle de T.S.F.,
 Ingénieur-Conseil Cie Industrielle des Métaux électroniques.

- MM. G. POTIER, Ingénieur à la Société « Le Matériel Téléphonique ».
 P. RIVÈRE, Chef du Service « Multiplex » de la Sté Française Ra-
 dioélectrique.
 M. SOLLIMA, Directeur du Groupe Electronique de la Cie Française
 Thomson-Houston.
 H. TESTEMALE, Ingénieur des Télécommunications.
 A. VIOLET, Chef de Groupe à la Sté « Le Matériel Téléphonique »
 l'Ingénieur J. BLOESMMA, Ingénieur Radio E. S. E.
 A. CHARLES, Général du C. R., Conseiller technique à la Société
 Kodak-Pathé.
 A. DIDIER, Professeur au C. N. A. M.
 G. GOUDET, Directeur de l'Ecole Nationale Supérieure d'Elec-
 tricité et de Mécanique de Nancy.
 M. D. INDJOUJIAN, Ingénieur des Télécommunications, chargé
 du Département Télécommande du C. N. E. T.
 P. MANDEL, Ingénieur en Chef à la Société Nouvelle de l'outillage
 R.B.V. et de La Radio-Industrie.
 A. PAGES, Ingénieur à la Compagnie Générale de T. S. F.
 H. TANTER, Ingénieur au L. C. T.

RÉSUMÉS DES ARTICLES

UTILISATION DE LA MESURE DE PHASE DANS UN SYSTÈME DE GUIDAGE ET DE LOCALISATION DES MOBILES, par J. ZAKHEIM, *Ingénieur E.S.E. et I.E.T., Chef de Groupe de Recherches à l'ONERA*. *Onde Electrique* de février 1955 (pages 137 à 141).

On décrit un dispositif radio mis au point à l'O.N.E.R.A. et depuis un an en exploitation dans cet organisme. L'appareillage est destiné à la détermination de la direction dans laquelle se trouve un mobile.

Après avoir passé en revue les caractéristiques essentielles de cet appareillage, on décrit les différents résultats que ce matériel permet d'atteindre tels que : le guidage dans un plan vertical, la mesure des vitesses en vol horizontal, la restitution des trajectoires.

CORRECTEUR AUTOMATIQUE DE FRÉQUENCE POUR ÉMETTEUR DE TÉLÉCOMMUNICATIONS SUR ONDES CENTIMÉTRIQUES, par J. CAYZAC, *Ingénieur aux Laboratoires d'Electronique et de Physique Appliquées*. *Onde Electrique* de Février 1955 (pages 151 à 158).

L'article décrit un correcteur automatique de fréquence fonctionnant spécialement dans les gammes d'ondes centimétriques. Il permet de stabiliser notamment un oscillateur modulé en fréquence ou en amplitude dont le spectre occupe une large bande de fréquences. Une description est donnée d'un tel dispositif utilisé dans un émetteur de télécommunications. Le procédé de discrimination de fréquence fait appel à deux cavités résonantes accordées sur des fréquences symétriquement décalées par rapport à la fréquence désirée. Un seul et même cristal détecteur délivre, à partir de signaux sortant alternativement des deux résonateurs, un signal contenant les informations d'amplitude et de signe de la correction.

ONDES ÉLECTROMAGNÉTIQUES POLARISÉES ELLIPTIQUES ET CIRCULAIRES, par Maurice BOUX, *Docteur ès-Sciences*. *Onde Electrique* de février 1955 (pages 126 à 136).

L'utilité des ondes électromagnétiques polarisées non rectilignes s'introduit d'une part lorsqu'il s'agit d'assurer la permanence de certaines liaisons radioélectriques entre terre et avion par exemple ou entre avion et fusée bien que les antennes de réception et d'émission puissent occuper des positions respectives assez variées ; d'autre part en radar lorsqu'il s'agit par exemple de différencier l'écho d'un avion des échos dus à la pluie : si l'émission est en polarisation circulaire, l'écho de pluie est en polarisation circulaire de sens inverse et n'est en principe pas reçue par le radar, alors que l'écho de l'avion est elliptique quelconque et est bien reçue.

On passe en revue divers dispositifs où intervient la polarisation non rectiligne : guides d'ondes, lames de diélectriques ou ailettes métalliques, déphaseurs, gyrateurs, puis on décrit divers dispositifs d'antennes qui fonctionnent en polarisation elliptique ou circulaire ; réseaux de lames à 45°, antennes en hélice, et on indique leurs emplois principaux. On donne quelques résultats expérimentaux d'élimination des échos de pluie. On termine par un résumé succinct des études théoriques relatives à la polarisation elliptique.

ETUDE DU COMPORTEMENT DYNAMIQUE DES ASSERVISSEMENTS COMPORTANT DES THYRATRONS, par P. BONNET, *Ingénieur Militaire de 1^{re} classe des Fabrications d'Armements*. *Onde Electrique* de février 1955, (pages 97 à 104).

L'article comprend deux parties.

— Dans la première, l'auteur traite le cas d'un moteur à excitation séparée dont l'induit est alimenté par des thyratrons, lorsqu'on peut légitimement supposer que les différents paramètres varient linéairement autour d'une valeur moyenne.

Si la conduction est discontinue, on ne doit pas tenir compte de la self de l'induit et l'ensemble thyatron-moteur forme un bloc dont on doit calculer globalement la fonction de transfert.

Si la conduction est continue, tout se passe comme si le groupe des thyratrons était simplement une alimentation à tension variable, en série avec le moteur, et on doit prendre la self du moteur avec sa valeur réelle.

— Dans la deuxième partie, l'auteur traite le cas plus général où on ne peut plus linéariser les phénomènes. Il expose une méthode, basée sur l'assimilation classique des phénomènes périodiques à leur fondamental, qui permet, en principe, de trouver la fonction de transfert globale de l'ensemble thyratrons-moteur. Malheureusement, l'établissement des graphiques qui semblent absolument nécessaires pour caractériser certains éléments de l'asservissement, nécessite un travail très important.

LA MODULATION D'ESPACEMENT D'IMPULSIONS, par G. X. POTIER, *Ingénieur à la Société Le Matériel Téléphonique*. *Onde Electrique* de Février 1955 (pages 159 à 164).

Description d'un procédé de modulation permettant de transmettre une centaine de communications avec une bande passante du même ordre que celle utilisée pour les équipements à modulation de position à 24 voies.

Ce résultat est obtenu en utilisant de façon plus rationnelle le temps disponible pour la transmission de l'information et fait bénéficier les liaisons par impulsions de l'effet statistique utilisé dans les équipements à courants porteurs.

Le principe même de la modulation d'espacement donne naissance à un bruit d'origine diaphonique qui peut être éliminé à l'aide de dispositifs enregistreurs.

CONSIDÉRATIONS SUR L'ENREGISTREMENT MAGNÉTIQUE DANS LE DOMAINE DES MACHINES À CALCULER, par F.H. RAYMOND, *Société d'Electronique et d'Automatisme*. *Onde Electrique* de Février 1955, (pages 89 à 96).

L'exposé se propose de mettre en évidence les conditions physiques liées à l'emploi de l'enregistrement magnétique dans le domaine particulier des machines à calculer numériques.

Il met en évidence l'origine des relations qui existent entre le dimensionnement géométrique des têtes de lecture et d'écriture et la densité d'informations binaires pouvant être enregistrées.

SUMMARIES OF THE PAPERS

STUDY OF THE DYNAMIC BEHAVIOUR OF SERVO MECHANISMS USING THYRATRONS, by P. BONNET, *Ingénieur Militaire de Première Classe des Fabrications d'Armements*, Onde Electrique February 1955 (pages 137 to 141).

The article consists of two parts.

In the first the author deals with the case of a separately excited motor where the armature is supplied from thyratrons when it can be assumed the different parameters vary linearly about a mean value.

If conduction is discontinuous it is not necessary to consider the self inductance of the armature circuit, but the thyatron-motor system must be considered as a whole and the transfer constant calculated accordingly.

If the conduction is continuous, the behaviour of the group of thyratrons is that of a variable voltage source in series with the motor, and the actual value of the self inductance of the motor must be taken into account.

In the second part the author deals with the more general case where the linearity of the behaviour can no longer be assumed. A method is suggested based on relating periodic phenomena to their fundamental repetition rate in the standard manner, whereby in principle an overall transfer function is deduced of the thyatron-motor combination. Unfortunately the preparation of certain graphs essential for demonstrating the behaviour of some parts of the servo mechanism requires a formidable amount of work.

PULSE SPACE MODULATION, by G. X. POTIER, *Ingénieur à la Société Le Matériel Téléphonique*, Onde Electrique, February 1955 (pages 159 to 164).

The article describes a modulation process which permits the transmission of 100 communication channels in a bandwidth of the same order as that used for 24 channels pulse position modulation systems.

This result is obtained by using more rationally the available time for the transmission of information, and by making use in pulse systems of the statistical information used in carrier systems.

The principle of space modulation gives rise to crosstalk noise which can be eliminated by storage devices.

REQUIREMENTS OF MAGNETIC RECORDING ASSOCIATED WITH CALCULATING MACHINES, by F.H. RAYMOND, *Société d'Electronique et d'Automatisme*, Onde Electrique, February 1955 (pages 89 to 96).

The article describes the physical conditions associated with the use of magnetic recording for numerical calculating machine.

It shows the basis of the relations between the geometrical dimensions of the recording and reproducing heads and the density of the binary information which can be recorded.

USE OF PHASE MEASUREMENTS IN A SYSTEM FOR GUIDING AND LOCATING MOVING OBJECTS, by J. ZAKHEIM, *Ingénieur ESE and IET, Chef de Groupe de Recherches à l'ONERA*, Onde Electrique, February 1955 (pages 137 to 141).

A radio equipment developed by, and since last year, produced by O.N.E.R.A. is described. The apparatus is intended to determine the direction of a mobile object.

After reviewing the essential characteristics of this apparatus, the different results which it is possible to obtain are described; for example, guidance in a vertical plane, the measurement of the speed in horizontal flight, the restoration of trajectories.

AUTOMATIC FREQUENCY CONTROL FOR CENTIMETRIC WAVE TRANSMITTER, by J. CAYZAC, *Ingénieur aux Laboratoires d'Electronique et de Physique Appliquées*, Onde Electrique, February 1955 (pages 151 to 158).

The article describes an automatic frequency control system designed for working in the centimetric wave band. It is particularly suitable for stabilising an oscillator which is either frequency or amplitude modulated and a wide spectral frequency band is occupied. A description is given of a similar arrangement for a telecommunication transmitter. The process of frequency discrimination is referred to two resonant cavities tuned to frequencies symmetrically located with respect to the desired frequency. By virtue of the signals alternately emitted from the resonators, one crystal detector provides the amplitude and sign of the required correction.

ELLIPTIC AND CIRCULAR POLARISED ELECTROMAGNETIC WAVES, by Maurice BOUX, *Docteur ès-Sciences*, Onde Electrique, February 1955 (pages 126 to 136).

The value of non-rectilinear polarised electromagnetic waves is evident when it is necessary to ensure the reliability of certain radio communication channels; for example, from the ground to an aeroplane, or from an aeroplane to a rocket, in spite of the fact that the sending and receiving antennae may occupy rather variable positions relative to each other. These waves are also valuable in radar, for example, in differentiating between the echo from an aeroplane and those caused by rain. If the transmitter uses circular polarisation, the echoes from rain will also have circular polarisation in the opposite sense and will not affect the radar, whilst the echo from an aeroplane will exhibit some degree of elliptical polarisation and will be well received.

Various items taking advantage of non-rectilinear polarisation are reviewed, including wave guides, sheets of dielectrics, or vanes of metal, phase shifters, and gyrators. Following this, various antenna arrangements based on elliptical or circular polarisation, such as arrays of vanes at 45°, and helical antennae, are described and their chief applications indicated.

The results of experiments on the elimination of echoes from rain are given, and the paper concludes with a short resumé of the theory of elliptical polarisation.

RÉSUMÉS DES ARTICLES (Suite)

SUR LA RÉALISATION D'UN COEFFICIENT DE TRANSFERT DONNÉ (OU : LES MONTAGES D'AMPLIFICATEURS A RÉACTION). par F.H. RAYMOND, *Société d'Electronique et d'Automatisme*. Onde Electrique de février 1955 (pages 105 à 112).

On se propose d'indiquer dans cet article comment la généralisation naturelle des montages d'amplificateurs à réaction permet, avec l'utilisation des réseaux formés de résistances et de capacités, la réalisation de coefficients de transfert des plus généraux.

Il est signalé, en dehors du calcul analogique, la possibilité d'emploi de ces méthodes dans le filtrage des phénomènes très lents.

MULTIPLIEUR A DÉCOUPAGE TEMPOREL, par Madame LILAMAND, *Ingénieur à la Société d'Electronique et d'Automatisme*. Onde Electrique de février 1955 (pages 142 à 150).

Plusieurs types de multiplieurs analogiques ont déjà été envisagés. Le multiplieur à découpage temporel est particulièrement intéressant parce qu'il nécessite peu de matériel et parce qu'il permet d'obtenir plusieurs produits ayant un facteur commun.

Cet appareil est composé de deux éléments : un modulateur et un multiplieur. Le modulateur, qui peut commander plusieurs multiplieurs, fournit une tension en créneaux d'amplitude constante et de fréquence fonction de l'une des variables ; le multiplieur module l'amplitude de ces créneaux en fonction de l'autre variable et en prend la valeur moyenne.

La mise au point d'un commutateur électronique de précision qui est l'organe essentiel du multiplieur et la compensation des capacités parasites des lampes qui ne sont plus négligeables à la fréquence de découpage de 2 kHz ont permis d'obtenir une précision de 1/1 000.

MESURE DE LA RÉSISTIVITÉ ET DE LA CONSTANTE DE HALL PAR LA MÉTHODE DES POINTES, par J. LAPLUME, *Docteur ès-sciences, Ingénieur à la Cie Fse Thomson-Houston*. Onde Electrique de février 1955 (pages 113 à 125).

L'auteur étudie la répartition du champ électrique créé à l'intérieur d'un milieu conducteur homogène lors du passage d'un courant électrique injecté par deux pointes. Il en déduit les formules permettant de calculer la résistivité et la constante de Hall dans la méthode des pointes. Les calculs sont développés dans le cas d'un échantillon parallélépipédique. On montre que la forme exacte de l'échantillon influe peu sur les résultats, pourvu que certaines dimensions soient supérieures à l'écartement des pointes.

SUMMARIES OF THE PAPERS (Continued)

GATING MULTIPLIER, by Madame LILAMAND, *Ingénieur à la Société d'Electronique et d'Automatisme*. *Onde Electrique*, February 1955 (pages 142 to 150).

Several types of analog multiplier have already been envisaged.

The multiplier with a time cut-off is particularly interesting because it requires few components and because many products, with a common factor, can be obtained from it.

This apparatus consists of two elements, a modulator and a multiplier. The modulator, which can control a number of multipliers, provides a castellated wave voltage of constant amplitude and repetition frequency of one of the variables; the multiplier modulates the amplitude of the wave in accordance with a function of the other variable and a mean value is taken.

The development of a precise electronic commutator, which is the essential feature of the multiplier, and the compensation of parasitic capacities of valves which are no longer negligible at a gating frequency of 2 000 c/s enables a precision of 1 in 1 000 to be obtained.

THE REALISATION OF A GIVEN TRANSFER COEFFICIENT (OR FEEDBACK AMPLIFIER SYSTEMS), by F.H. RAYMOND, *Société d'Electrique et d'Automatisme*. *Onde Electrique*, February 1955 (pages 105 to 112).

This article sets out to show that transfer characteristics of a most general form can be realised by feedback amplifier systems using capacity — resistance networks.

The possibility of employing these methods for filtering low frequency phenomena is pointed out.

MEASUREMENT OF RESISTIVITY AND THE HALL CONSTANT BY POINT METHODS, by J. LAPLUME, *Docteur ès-Sciences, Ingénieur à la Compagnie Française Thomson-Houston*. *Onde Electrique*, February 1955 (pages 113 to 125).

The author studies the distribution of the electric field produced in the interior of a homogenous conductive medium during the passage of an electric current injected at two points. Formulae for deducting the resistivity and Hall constant are deduced by the point method. The calculation is developed for the case of a specimen of parallelepiped specimen. It is shown that the exact form of the specimen has little effect on the results provided certain dimensions are larger than the separation of the points.

CONSIDÉRATIONS SUR L'ENREGISTREMENT MAGNÉTIQUE DANS LE DOMAINE DES MACHINES A CALCULER*

PAR

F. H. RAYMOND

Directeur de la Société d'Electronique et d'Automatisme

1. — INTRODUCTION.

L'idée d'utiliser l'enregistrement magnétique pour constituer la mémoire d'une machine à calculer date de plusieurs années.

L'apparition de techniques diverses nouvelles permettant la réalisation des organes de mémoire a modifié les conditions d'emploi de l'enregistrement magnétique dans le domaine des machines à calculer, mais lui a donné une importance indiscutable à ce jour.

En effet, on considérait que la recherche de l'optimum pour un tambour magnétique devait coïncider avec le meilleur prix de revient compatible avec le plus petit temps d'accès (on appelle temps d'accès la durée d'attente pour recueillir une information sur une piste du tambour lorsqu'on définit son emplacement). Déjà des méthodes de programmation, associées à une répartition des adresses des emplacements de mémoires sur une piste, permettent de s'affranchir des conditions trop strictes relatives au temps d'accès des pistes d'un tambour magnétique, mais leur usage ne peut être optimum dans tous les programmes qu'aura à accomplir une calculatrice numérique universelle.

Cette méthode est utilisée sous une forme simple, en particulier dans la machine C.U.B.A. (1) dans laquelle le pas des adresses dans l'ordre naturel peut être programmé (ce fonctionnement est appelé « stroboscopie programmée »).

La possibilité d'utiliser une armoire à accès direct relativement importante (possibilité déjà offerte par les tubes cathodiques de mémoires type Williams, par exemple) conduit à rechercher l'accélération ou le non ralentissement d'une machine à calculer dans une organisation convenable du transfert des informations des pistes du tambour magnétique

à cette mémoire auxiliaire. C'est ainsi que dans le calculateur C. A. B. 2022, le temps d'accès du tambour joue un rôle secondaire, une mémoire à accès direct de 128 mots, formée de matrices de tores en ferrite, permet l'utilisation optimum des instructions de la machine et de son organe de calcul.

Cet exposé a pour objet de mettre en évidence les conditions caractéristiques de l'enregistrement magnétique dans le domaine des machines à calculer. Les informations à mémoriser sont de nature binaire, le mode d'enregistrement magnétique étant par essence un mode d'enregistrement série, une transformation du mode de représentation des informations binaires est donc nécessaire si on désire que la substitution d'une information à une autre s'effectue sans aucune incidence sur la sécurité de fonctionnement du dispositif; en particulier, l'effacement se réalise automatiquement lorsqu'on écrit une information différente de celle qui était enregistrée dans un même emplacement.

Nous justifierons plus tard le choix de la modulation de phase.

Les études conduites à la S.E.A., compte tenu des observations qui précèdent, ont eu comme objectif de rechercher une solution voisine de l'optimum d'un tambour magnétique, en mettant en jeu les paramètres suivants :

— instabilité mécanique du tambour (liée à son inertie, donc à la vitesse linéique de la couche d'enregistrement),

— le minimum de tubes électroniques nécessaires pour les opérations de commutation.

A dire vrai, nous considérons que la réalisation actuelle du tambour C.A.B. (1), dont quatre exemplaires sont en cours ou d'expérimentation ou de terminaison (2), constitue un pas vers une solution

(*) Communication présentée au Congrès sur les Procédés d'Enregistrement (Paris avril 1954).

(1) Calculatrice Universelle Binaire de l'Armement, en cours d'installation au Laboratoire Central de l'Armement au Fort de Montrouge.

(1) Calculatrice Arithmétique Binaire, universelle.

(2) Cette réalisation simultanée de quatre tambours devant permettre de mieux définir les caractéristiques mécaniques.

plus évoluée dans laquelle le problème de commutation sera résolu, à notre avis, de manière plus logique, la solution actuelle étant une première étape conditionnée d'ailleurs par les matériaux magnétiques disponibles en France au début de cette entreprise.

Bien entendu, si le temps d'accès à une piste joue maintenant un rôle secondaire, ceci n'est vrai que dans les limites qu'il n'est pas possible de fixer ici lorsqu'on sépare le tambour de l'organisation d'ensemble d'une calculatrice.

2. — ASPECT PHYSIQUE DE L'ENREGISTREMENT MAGNÉTIQUE.

Délaissant la modulation de position qui fut évoquée il y a quelques années, on dispose, *a priori*, de deux méthodes pour magnétiser une couche sensible recouvrant ou formant la surface extérieure d'un cylindre tournant autour de son axe.

Dans le premier procédé (fig. 1) le flux de fuite de la tête magnétise le milieu, les lignes de force sont parallèles à l'axe du tambour.

Dans le second procédé (fig. 2), au contraire, les lignes de force sont dirigées dans le sens de la vitesse relative du tambour par rapport à la tête.

Dans la première méthode, il est évident que l'épaisseur de la tête ne peut dépasser l'emplacement réservé à une unité d'information sur une piste. En conséquence, la section du flux sera petite et sa reluctance élevée, ce qui fait préférer la seconde méthode.

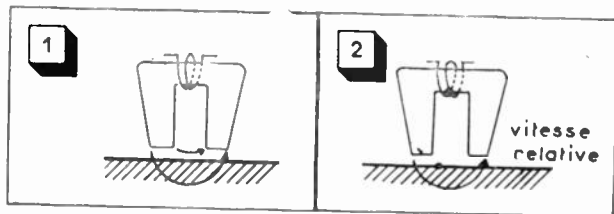


FIG. 1. — Coupe Diamétrale.

FIG. 2. — Coupe perpendiculaire à l'axe du tambour.

Considérons figure 3 l'extrémité d'une tête d'enregistrement au voisinage du milieu magnétique d'inscription.

On a tracé sur la figure la forme des lignes de force, dont la densité va rapidement en décroissant à partir de l'entrefer de la tête. Si on considère, par exemple, la surface extérieure de la couche magnétique, le module du champ suit la loi du type de celle indiquée figure 4.

Si on considère le déplacement de la tête devant le support, suivant une translation parallèle à ce dernier, de vitesse V_0 , du fait du cycle d'hystérésis et des cycles de recul caractérisant le milieu d'enregistrement, l'intensité d'aimantation en un point est une fonctionnelle qui peut s'écrire :

$$\vec{I}(x, t) = \vec{F} \left\{ f(\tau) \prod_{\tau=-\infty}^{\tau=t} (x - V_0 \tau) \right\}$$

ou $f(t)$ est la loi des ampères/tours dans la bobine de la tête en fonction du temps t .

Il est évident qu'une théorie convenable de l'enregistrement magnétique devrait être basée sur l'étude des propriétés mathématiques de cette fonctionnelle et de ses dérivées (dérivées au sens de la théorie des fonctionnelles).

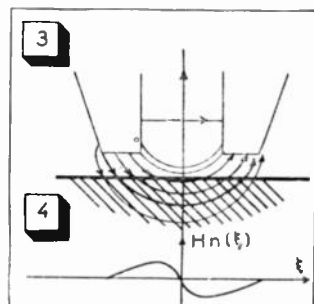


FIG. 3.

FIG. 4.

Cette remarque nous conduit à manifester une certaine méfiance à l'égard des procédés de calculs mathématiques (ceux utilisant en particulier un principe de réciprocité valable en théorie linéaire) pour déterminer le rapport des signaux entre la lecture et l'inscription. Ce n'est pas notre propos de discuter la validité des calculs habituellement effectués dans le cas de l'enregistrement sonore.

Dans le cas qui nous intéresse, recherchant le maximum d'hérédité des phénomènes — le mot hérédité étant bien, dans ce cas, synonyme de mémoire — nous abandonnons l'idée d'effectuer des calculs mathématiques, du moins pour l'instant et nous nous contenterons dans cette étude de considérations physiques.

On ouvrira ici une parenthèse visant l'enregistrement des informations, en général.

Nous pensons, toutes les fois que l'on sera conduit à augmenter la densité d'informations par unité de surfaces et à augmenter le spectre de fréquence de ces informations, qu'à côté des procédés classiques d'enregistrement, le procédé consistant à coder préalablement le signal (c'est-à-dire à le transformer en signaux discrets en numération binaire ou analogue), à l'enregistrer sous cette forme et à le décoder après lecture, devrait constituer une solution capable de supporter la concurrence avec des procédés plus coutumiers en ayant l'avantage d'avoir recours, pour chaque fonction, à une technique bien adaptée.

Nous imaginons donc que l'enregistrement magnétique tel qu'il se développe pour les machines à calculer sera appelé à contribuer à la résolution de bien d'autres problèmes.

2.1. — Mécanisme d'écriture.

Nous nous proposons tout d'abord d'insérer une succession récurrente d'impulsions ; on suppose que les moyens techniques sont mis en œuvre pour que la loi des ampères/tours $f(t)$ soit elle-même une succession d'impulsions : plus précisément la force

magnétomotrice dans la tête d'inscription est, T désignant la période de récurrence, égale à M durant $\frac{T}{2}$ puis prend brusquement la valeur $-M$ durant la demi-période $\frac{T}{2}$ suivante et ainsi de suite.

La distribution des lignes de force est donc invariable durant chaque demi-période, leur sens bascule d'une demi-période à la suivante.

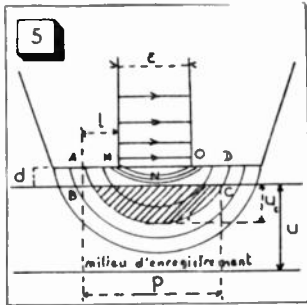


FIG. 5.

Considérons l'un de ces états, le milieu d'enregistrement étant pris à l'état désaimanté (pour la simplicité de l'exposé). Très grossièrement sont tracées sur la figure 5, ci-dessus, les lignes de force créées par la tête. Soit H_0 le champ de saturation. Supposons que soient réalisées les conditions d'alimentation de la tête de manière que le champ à l'intérieur d'une ligne de force $A B C D$ soit égal ou supérieur à H_0 . Ce domaine est borné du côté tête par une ligne de force $M N O$.

L'énergie à dépenser pour réaliser ce spectre est pour l'instant secondaire, elle croit évidemment avec d , distance tête-tambour, sans qu'une loi simple de dépendance par rapport à d soit établie, le flux de fuite considéré dépendant de la forme de la tête (laquelle étant relativement petite ne permet pas d'être réalisée suivant un dessin rigoureux, en particulier lorsqu'il est fait usage de ferrites).

Dans le domaine considéré à l'instant, hachuré sur la figure, le milieu est saturé, nous dirons aimanté. On doit obtenir ce résultat sans saturer les becs de la tête ce qui aurait pour effet d'augmenter d donc d'allonger les lignes de fuite. Néanmoins, ceci peut être compatible avec le résultat désiré lorsque, ce qui est le cas pratique, la tête se sature par des valeurs du champ supérieures à celle qui sature le milieu d'enregistrement.

Supprimons la force magnétomotrice dans la tête. Si le milieu est assimilable à un aimant parfait, le domaine $(B C N)$ reste aimanté avec l'intensité $\vec{I}_r$ (I_r : aimantation rémanente).

On remarque immédiatement que l'épaisseur du milieu d'enregistrement n'a pas besoin d'être supérieure à l'épaisseur du domaine saturé u_c .

Le champ créé par le milieu ainsi aimanté est celui que créent la densité volumique de charge magnétique égale à $\text{div } \vec{I}$ et la densité superficielle $-I_n$ (I_n composante normale de I). Comme $\vec{I}$ est constant en module (hypothèse précédente); $\text{div } \vec{I} = 0$ dans tout le domaine saturé. Sur la surface intérieure

$B C$ la divergence conduit, (un calcul simple le confirme) en fait, à définir une densité égale à $-I_n$, or BC coïncide avec une ligne de force, donc le champ créé par le milieu est équivalent à celui créé par la densité $\sigma = -I_n$ sur la face extérieure BNC du milieu. La loi de distribution de σ en fonction de l'abscise ξ sur BNC ($\xi = 0$ sur l'axe de symétrie) est donc déterminée par l'inclinaison des lignes de force au passage dans le milieu.

Il semblerait que l'épaisseur du milieu puisse prendre n'importe quelle valeur.

En réalité, si elle est inférieure à u_c , des charges superficielles apparaissent sur la face intérieure du milieu (vers l'axe du tambour) de signes opposés à celles qui sont apparues sur la face tournée vers la tête. Ceci aura pour effet de diminuer le champ créé par le milieu ainsi aimanté vers la tête. Si on prend $u \gg u_c$ l'aimantation, sans saturation, du domaine extérieur à celui considéré plus haut absorbera une certaine énergie qu'il faudra fournir à la tête, il n'est donc pas souhaitable d'avoir une épaisseur de milieu d'enregistrement supérieure, ou très supérieure à l'épaisseur critique u_c .

L'aspect physique approximatif des lignes de force, mais non foncièrement inexact, représenté fig. 5, conduit à la densité superficielle $\sigma(\xi)$ de la fig. 6 et la distance des maxima σ_m est voisine de p (avec $p = \epsilon + 2l$)

ϵ et l ayant les significations portées sur la figure 5. Déplaçons la tête de $dx = V_0 dt$ de la droite vers la gauche sur nos figures.

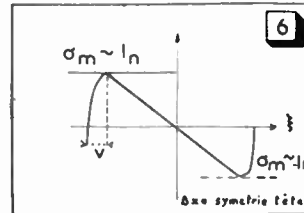


FIG. 6.

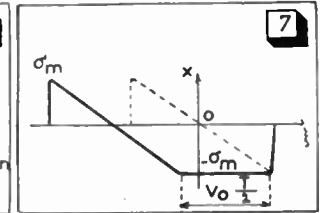


FIG. 7.

On voit que la limite de la zone saturée, se déplace de dx dans le même sens.

Pour ce qui intéresse le champ qui sera créé par l'aimantation du milieu, la densité superficielle est celle qu'on obtient en déplaçant vers la droite la courbe de la figure 6. A la fin d'une demi-période la densité est celle de la figure 7, les régions ayant été saturées le restent lorsqu'elles sont atteintes par des lignes de force situées à l'arrière (par rapport au sens du mouvement) dont le champ est inférieur au champ coercitif, or il en est ainsi le champ décroissant de manière monotone de part et d'autre de l'axe de symétrie de la tête (fig. 5.)

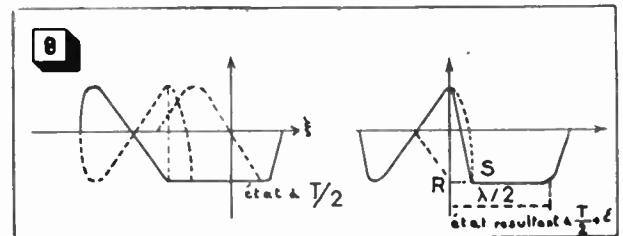


FIG. 8.

Les conditions d'alimentation de la tête étant supposées idéalement remplies à $\frac{T}{2}$ on inverse la

courbe $\sigma(\xi)$ d'où le résultat présenté figure 8.

L'état intermédiaire $R'S$ est assez difficile à définir, du moins pouvons-nous affirmer qu'il existe. La répétition du phénomène conduit au diagramme de la figure 9.

On réalise donc au pas $\lambda = V_0 T$, des pôles de densité $\pm \sigma_m$ et de longueur $\frac{\lambda}{2} - V$, V est difficile à apprécier : cette quantité dépend des cycles d'hystérésis du matériau recouvrant le tambour et de la forme des becs de la tête. Par ailleurs, si l'inversion de la force magnétomotrice dure un temps $\tau \ll \frac{T}{2}$ alors V est augmentée d'au moins τV_0 .

En réalité, l'effet du champ démagnétisant tend à accroître les fronts donc la longueur V est sous la dépendance de la matière du milieu de manière d'autant moins importante que le champ coercitif de ce milieu est grand.

D'après l'analyse qui précède, les dimensions de la tête n'interviennent pas directement sur le mode d'enregistrement. La densité d'information définie ici par $\delta = 1/\lambda$ est donc indépendante des dimensions et de la forme de la tête : celles-ci n'interviennent qu'en raison de la lecture si cette fonction est réalisée

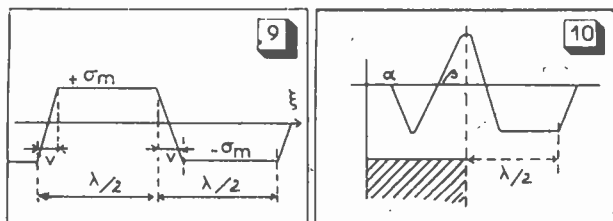


FIG. 9.

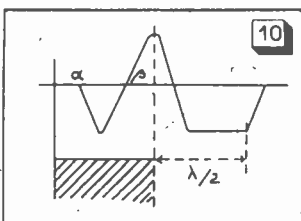


FIG. 10.

avec la même tête ou s'il est prévu une tête de lecture distincte de celle d'écriture, la densité σ_m croît avec $1/d$ puisque lorsque d décroît, l'inclinaison α des lignes de force au passage dans le milieu diminue de sorte que $\sigma_m = I_n = I_{\max} \sin \alpha$ croît avec α . Mais un second phénomène intervient, caractéristique du domaine d'application considéré ici. Si on considère en effet un certain nombre d'emplacements réservés à des informations binaires, il est nécessaire que le remplacement d'une information par une autre dans un seul de ces emplacements soit réalisé sans détruire ou perturber les informations contenues dans les cases voisines.

Pour examiner cette question, il n'est plus possible de supposer le milieu à l'état disons vierge. Puisque c'est la « queue » du spectre (vu en densité σ superficielle) de fuite qui définit l'aimantation, l'ensemble du spectre, plus précisément la longueur $\alpha \beta$ de celui-ci peut perturber la zone en aval (fig. 10). Il est donc nécessaire que $\alpha \beta$ soit aussi petit que possible. L'optimum semble être que $\alpha \beta$ soit égal à :

$$V + V_0 \tau \text{ or } \alpha \beta = p + V$$

d'où la condition :

$$p + V \sim V + V_0 \tau$$

$$\text{où } p \cong \varepsilon + 2l \cong V_0 \tau$$

On peut encore écrire :

$$\frac{p}{\lambda} = \frac{\tau}{T}$$

La quantité p sera désignée ci-après *entrefer équivalent*. Sa valeur dépendant de l , on notera que cette quantité dépend beaucoup de la forme de l'entrefer, si les dièdres de l'entrefer, tournés vers le milieu d'enregistrement, ne sont pas parfaitement réalisés, il en résulte une augmentation de l .

Des mesures effectuées à la S.E.A. sur des têtes en ferrites indiquent que $p = 1,2 \varepsilon$ environ.

Il est clair que, toutes choses égales par ailleurs, le rapport $\frac{p}{\varepsilon}$ croît lorsque d diminue, de sorte que le

dimensionnement optimum d'une tête en vue de réaliser une densité δ donnée est affaire d'expériences guidées par le calcul du spectre (ce dernier calcul est en cours à la S.E.A.).

On notera, en outre, que l'on peut admettre à l'extrême limite que $\frac{\tau}{T} = \frac{1}{4}$, de sorte que $p = \frac{\lambda}{4}$.

la répartition de σ devenant presque sinusoïdale, la loi temporelle des ampères-tours étant plus voisine d'une loi harmonique que d'une série de crêteaux rectangulaires.

Un second mode d'utilisation de l'enregistrement magnétique consiste à remarquer que nous n'avons jamais à modifier le contenu d'un seul emplacement élémentaire de mémoire, mais un groupe de tels emplacements contigus. De sorte que la perturbation examinée ci-dessus ne peut jouer que d'un groupe au suivant : confondant dans notre langage le contenant et le contenu, nous dirons que nous plaçons entre *deux mots* (binaires) consécutifs un emplacement de mémoire de 1 digit : dans ces conditions, nous pouvons prendre comme base de dimensionnement la relation : $p \leq \lambda$. On prend, en général, $p < \lambda$ afin que la variation de $\bar{d}$ lors de la rotation du tambour et en fonction de la dilatation des éléments du montage mécanique ne puissent modifier p dans un sens risquant d'inverser l'inégalité considérée.

2.2. - Mécanisme de lecture (Fig. 11).

Supposons une tête, à priori différente, placée devant le milieu aimanté et se déplaçant, par rapport à lui, à la vitesse V_0 .

Supposons que la répartition $\sigma(\xi)$ soit sinusoïdale : ceci est très voisin de la réalité si $p \sim \frac{\lambda}{4}$ et si $p < \lambda$

avec un espace entre mots et, dans ce cas, si la loi des ampères-tours est sinusoïdale. D'ailleurs, si utile, nous reviendrons sur cette hypothèse.

On peut calculer assez aisément le champ créé par $\sigma(\xi)$ dans l'entrefer d'épaisseur d . Sa composante normale à mm' (le milieu bordant cet entrefer,

côté tête étant supposé à perméabilité infinie) est :

$$Hy = \sigma_m e^{-\frac{2\pi d}{\lambda}} \cos \frac{2\pi \xi}{\lambda}$$

avec $\xi = V_0 t + \xi_0$ et lorsque $\sigma(\xi) = \sigma_m \cos \frac{2\pi \xi}{\lambda}$.

Les droites xx' , yy' sont des axes de symétrie. Le flux traversant un segment de longueur Δ est, par un calcul facile :

$$\Phi = \frac{\lambda}{2\pi} e^{-\frac{2\pi d}{\lambda}} \sin \frac{\pi \Delta}{\lambda} \cos 2\pi f t$$

puisque la fréquence f est donnée par $f = V_0 \lambda$ (et en prenant arbitrairement une phase $\xi_0 = 0$).

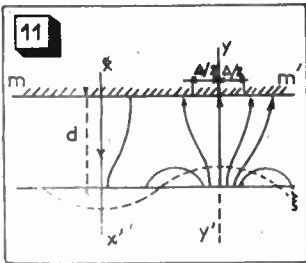


FIG. 11.

Bien qu'il serait nécessaire de perfectionner ces calculs pour tenir compte de l'entrefer de la tête, ϵ , on peut préciser (des expériences en cours ont pour objet de déterminer cette influence) que tout le temps que :

$\frac{\epsilon}{d}$ et $\frac{\epsilon}{\lambda}$ sont assez petits les résultats élémentaires ci-dessus donnent une vue qualitative suffisamment instructive pour être appliquée à la discussion de quelques résultats expérimentaux (fig. 12).

On pourra introduire, ainsi que cela est coutumier en électrotechnique, lorsque $\epsilon \neq 0$, un coefficient d'Hopkinson β et écrire par conséquent :

$$\Phi = \frac{\lambda}{2\pi} \cdot \beta e^{-\frac{2\pi d}{\lambda}} \sin \frac{\pi \Delta}{\lambda} \cos \pi f t$$

Le module de cette fonction harmonique est sera

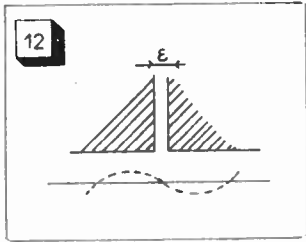


FIG. 12.

désigné, par la suite, Φ_M flux maximum capté par un pôle de la tête.

Nous allons examiner, dans le même esprit, le mécanisme de lecture en considérant trois cas typiques suivant l'importance relative de λ .

Cas λ très grand :

Le flux par pôle qu'il est possible de conduire à travers la tête est une fraction du flux par pôle mesuré sur la surface d'enregistrement (fig. 13).

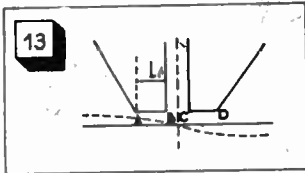


FIG. 13.

On ne commet pas une erreur en écrivant que c'est une fraction $\frac{l_1}{\lambda}$ de celui-ci, car l_1 est assez difficile à définir avec précision, la notion de grande longueur d'onde étant précisément en rapport avec la forme de la tête : on entre dans cette hypothèse dès que la longueur d'onde λ est supérieure à la partie $ABCD$ parallèle à la surface aimantée, de la tête.

Ce cas n'a pas été examiné car il conduit à des densités trop faibles pour être intéressantes.

Passons au deuxième cas typique.

Cas λ moyen :

Désignons par D l'axe des pôles de la tête (fig. 14).

Dans la position de la figure les pôles du milieu aimanté créent des flux égaux maxima et de signes

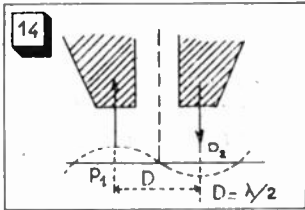


FIG. 14.

contraires, d'autant plus voisins du maximum théorique mentionné plus haut que d/ϵ est grand. On peut écrire que le flux utile est la différence des flux des deux pôles P_1 et P_2 , d'où la relation immédiate :

$$\Phi = k \left(\frac{d}{\epsilon} \right) \Phi_M \cdot e^{-\frac{2\pi d}{\lambda}} \sin \frac{\pi D}{\lambda} \cos \frac{2\pi V_0 t}{\lambda}$$

obtenue en comparant de Φ_M flux périodiques en t déphasés entre eux de $\frac{\pi D}{\lambda}$.

La figure 14 correspond au cas $D = \frac{\lambda}{2}$ à $t = 0$, et la figure 15 au cas $D = \lambda$ où $\Phi = 0$. Φ_M a, ici, la signification évidente suivante, d'après ce qui

précède (l'effet de Δ disparaissant, du moins pour simplifier les choses) :

$$\Phi_M = \beta \cdot \frac{\lambda}{2\pi} e^{-\frac{2\pi d}{\lambda}}.$$

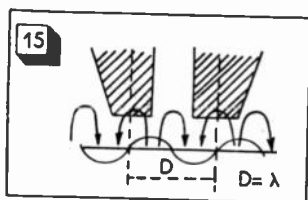


FIG. 15.

Cas λ petit (fig. 16 a et b).

Nous aurons affaire à ce cas si $\varepsilon > \lambda/2$.

La déformation des lignes de force qu'entraîne le trou de longueur ε modifie la manière dont la tête capte le flux, ou une partie de ce flux, produit par deux pôles consécutifs sur le milieu aimanté. L'entrefer ε (plus précisément un entrefer fictif légèrement supérieur à ε en raison de l'épanouissement des lignes de force aux alentours des dièdres O_1 et O_2) joue à l'égard du mécanisme de captation du flux un rôle analogue à D du cas précédent.

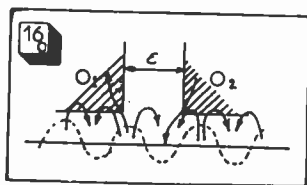


FIG. 16 a.

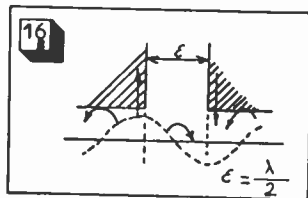


FIG. 16 b.

On dira que nous sommes placés dans le cas λ petit si l'effet de l'entrefer l'emporte sur les effets considérés jusqu'à présent.

Il est évident que la limite inférieure $\lambda = \varepsilon$ n'a qu'une valeur théorique ; de manière très grossière le flux maximum capté par un pôle de la tête est :

$$\Phi_M \frac{\frac{\lambda}{2} - \frac{\varepsilon}{2}}{\frac{\lambda}{2}} = \Phi_M \left(1 - \frac{\varepsilon}{\lambda}\right)$$

Φ_M ayant la même signification que précédemment.

Une loi plus exacte devrait s'exprimer comme suit :

$$\Phi = \Phi_M f \left(1 - \frac{\varepsilon}{\lambda}\right)$$

f étant une fonction décroissante avec l'argument $1 - \frac{\varepsilon}{\lambda}$.

Une étude plus précise, actuellement en cours, a pour objet de déterminer le rapport $\frac{\varepsilon}{d}$ au-dessus duquel l'effet analysé ici devient prépondérant.

La récapitulation des résultats précédents se fera en introduisant un certain nombre de spires sur le circuit magnétique reliant entre eux les deux bords de l'entrefer examiné jusqu'à présent, la force électromotrice produite étant, au signe près, égale à $\frac{d\Phi}{dt}$.

Nous ferons une première discussion en supposant que la fréquence des signaux est constante et égale à f . Le paramètre, lié à V_0 , sera la densité δ d'informations par piste.

L'amplitude de la force électromotrice est :

$$E = \beta \frac{\lambda}{2\pi} e^{-\frac{2\pi d}{\lambda}} G \left(\frac{\pi D}{\lambda}, \frac{\pi \Delta}{\lambda} \right) \cos 2\pi f$$

où $G = \sin \frac{\pi D}{\lambda}$ dans le cas λ moyen et $G = \sin \frac{\pi \Delta}{\lambda}$ si la longueur des becs de la tête n'est pas grande devant λ et ε .

En discutant en fonction de $\delta = \frac{1}{\lambda}$ et à $f = C^te$ on a donc, cas λ moyen :

$$E(\delta) = \frac{1}{\delta} e^{-2\pi d \delta} \cdot G(\pi D \delta) \sim \frac{1}{\delta} e^{-2\pi d \delta} \sin \pi D \delta$$

à des constantes multiplications près (sauf qu'en toute rigueur le coefficient d'Hopkinson β est fonction de δ , nous pensons d'après ce qu'on a vu, dans le cas λ moyen qu'on peut supposer β indépendant de δ).

La courbe représentative est donnée sur la planche I pour trois valeurs typiques du rapport $\frac{d}{D} \left(\frac{1}{10}, \frac{1}{2}, 1 \right)$.

Mode d'inscription des informations.

Nous avons étudié jusqu'à présent l'enregistrement et la lecture d'un signal périodique de fréquence f . L'enregistrement dans une longueur d'onde λ d'une information binaire signifie que l'onde considérée doit être modulée : une modulation d'amplitude, de fréquence (ou de phase) est donc applicable avec cette remarque que chaque période du porteur contenant une information, la fréquence la plus grande de celles-ci est $\frac{f}{2}$. La largeur de bande

du dispositif d'inscription et de lecture doit donc être de $\frac{f}{2}$. En modulation de phase par tout ou rien nous aurons donc la présence d'un signal à fréquence instantanée f ou $f/2$ (la fréquence étant définie comme la dérivée, à 2π près de la phase du signal).

Le dimensionnement optimum, ou le choix de la densité est donc déterminé en traçant les courbes à f et $f/2$ (représentées planches II, III et IV pour les valeurs $\frac{d}{D}$ prises comme exemple pour tracer les courbes $of = I \tau$ de la planche I) et en associant à la

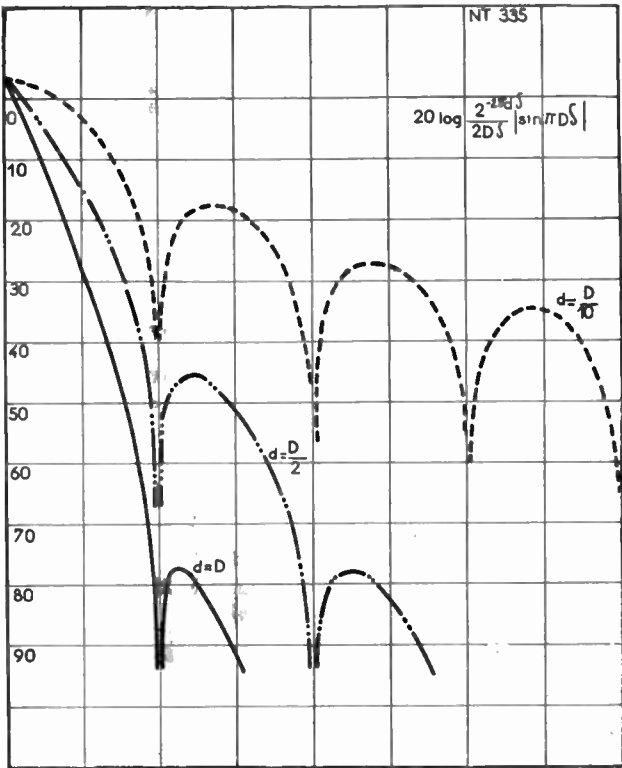


FIG. 17

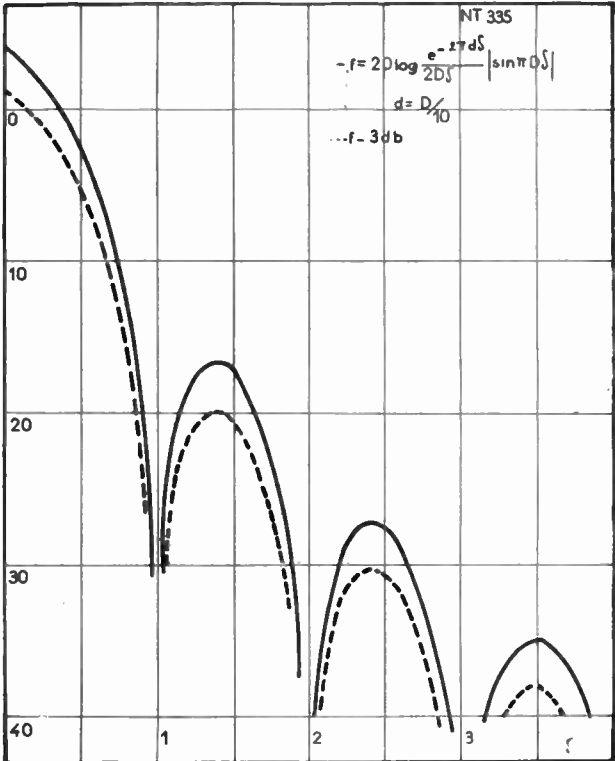


FIG. 18

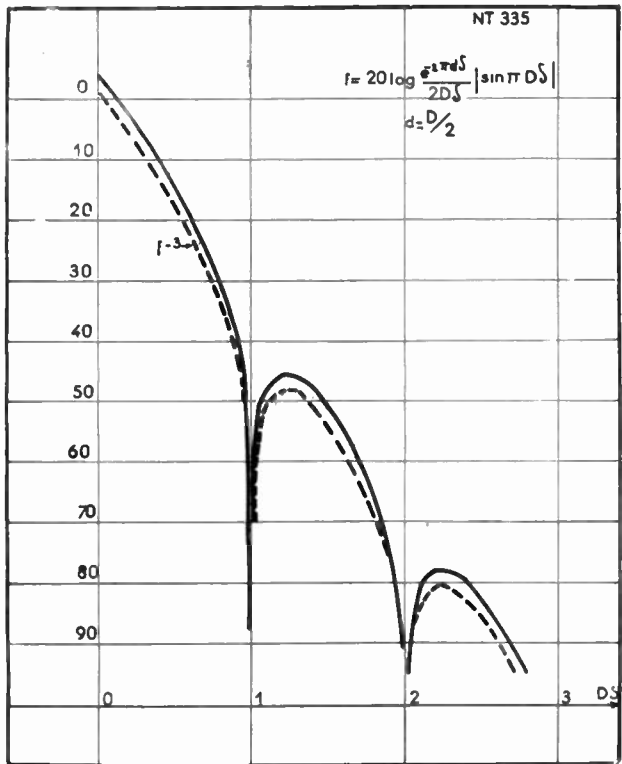


FIG. 19

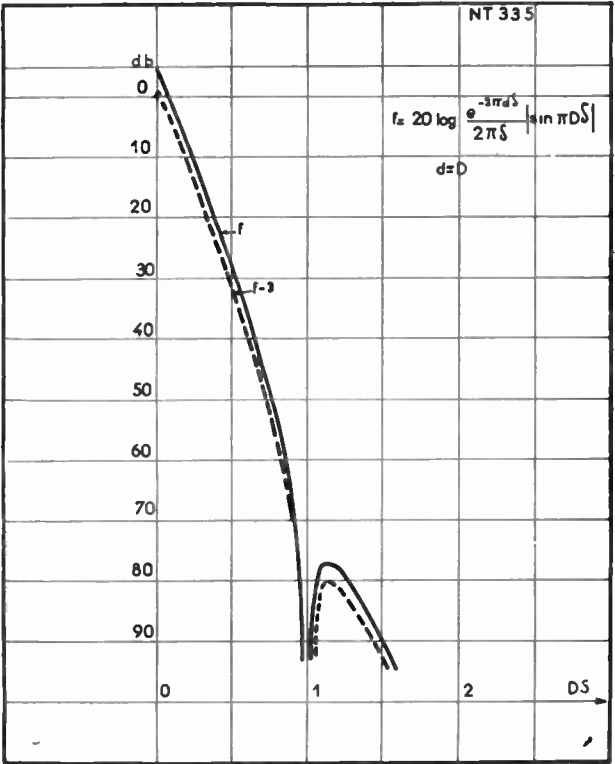


FIG. 20

fréquence f la densité δ , à $f/2$ la densité $\delta/2$. Le choix revient à avoir un signal ayant la même valeur aux deux extrémités de l'intervalle ($\delta/2$, δ) choisi en faisant intervenir, d'une part les courbes précédentes et, d'autre part, les caractéristiques du transformateur associé à la tête de lecture.

Cette discussion fera l'objet d'une autre publication.

Pour terminer cette communication, les deux photographies ci-dessous représentent, l'une (fig. 18)

le tambour isolé d'une capacité d'environ 200 000 digits (1) répartie en 64 pistes, l'autre le sous-ensemble d'une machine universelle type C.A.B. (2).

La commutation à l'écriture est réalisée par des matrices de tores, à cycle magnétique rectangulaire (partie haute de gauche).

(1) Capacité qui sera doublée.

(2) Calculateur Arithmétique Binaire.



FIG. 21

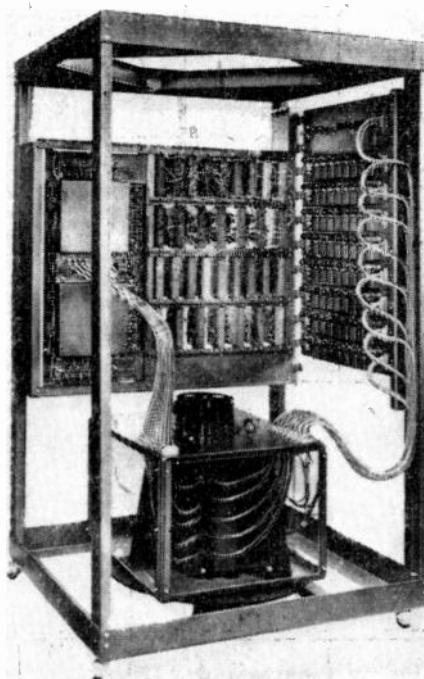


FIG. 22

ÉTUDE DU COMPORTEMENT DYNAMIQUE DES ASSERVISSEMENTS COMPORTANT DES THYRATRONS

PAR

P. BONNET

*Ingénieur Militaire de 1^{re} classe des Fabrications
d'Armements*

Les thyratrons constituent, dans les servomécanismes, des organes de puissance souples, légers et dénués pratiquement de constante de temps, tout au moins lorsqu'ils sont excités par une porteuse à fréquence suffisamment élevée. Cependant leur emploi s'accompagne souvent de l'impression qu'on ne sait pas très bien où l'on va, car leur étude par le calcul est délicate, et ce pour les raisons suivantes : ils ne sont pas linéaires, présentent plusieurs modes de fonctionnement, selon que la conduction est continue ou discontinue, et font apparaître dans certains cas une sorte de contre-réaction à discriminateur non-linéaire.

On est finalement souvent amené à procéder à une mise au point très empirique.

Le but de la présente communication n'est pas d'élucider complètement le problème du calcul des asservissements à thyratrons, qui est très compliqué, mais d'exposer quelques idées qui permettront peut-être de faire avancer la question.

Dans une première partie, nous étudierons le cas où la non-linéarité n'est pas trop accentuée, et où l'on peut légitimement considérer qu'autour d'une vitesse moyenne les accroissements des différentes variables suivent des lois linéaires. Ceci nous conduira à des schémas équivalents pour les deux modes de fonctionnement, continu et discontinu.

Dans la seconde partie, nous aborderons le problème plus général où on ne peut pas linéariser. C'est par exemple le cas d'un moteur à deux sens de marche alimenté par deux groupes de thyratrons débitant en sens inverses, si l'on ménage un seuil autour du point zéro pour éviter que les deux groupes débitent en même temps l'un dans l'autre.

On montrera qu'une étude complète de ce cas — qui peut d'ailleurs s'appliquer à d'autres systèmes de commande non-linéaires que les thyratrons — est soluble par une méthode graphique de principe relativement simple, mais que l'établissement des

graphiques nécessite un travail important, et particulier à chaque cas.

En conséquence, cette méthode ne semble pas à conseiller pour les servomécanismes de petite puissance, pour lesquels il semble plus rapide de procéder à la mise au point expérimentale. Mais il est possible qu'elle puisse rendre des services dans le cas d'asservissements importants.

I. — COMPORTEMENT DYNAMIQUE D'UN MOTEUR ALIMENTÉ PAR THYRATRONS, AUTOUR D'UNE VITESSE MOYENNE (1).

Nous partirons du schéma classique de l'étage de puissance représenté sur la figure 1. Le moteur *M*

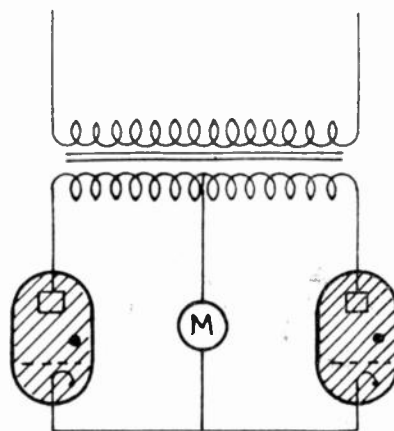


FIG. 1

est un moteur à excitation séparée dont l'induit est alimenté par les thyratrons.

(1) Voir KP Puchlowsky : Electronic Motor Control.

Industrial reference book.
John Wiley et Sons Inc. New York.

Les notations employées sont les suivantes :

E_e : f.e.m. efficace appliquée par chaque moitié du secondaire du transformateur.

e_c : f.c.e.m. du moteur.

Ω : vitesse du moteur.

$$K_e : \frac{e_c}{\Omega}$$

$$a = \frac{e_c}{\sqrt{2} E_e}$$

L, R : self et résistance d'induit.

i : courant instantané d'induit.

I : courant moyen d'induit.

ω : pulsation du courant.

$x = \omega t$.

x_a : angle d'allumage des thyratrons.

x_e : angle d'extinction des thyratrons.

$$\theta = A \operatorname{rctg} \frac{L \omega}{R}$$

$$E = I R$$

$$e = i R$$

$$\bar{E} = \frac{I R}{\sqrt{2} E_e}$$

On néglige la chute de tension dans les thyratrons.

L'équation donnant le courant instantané i est évidemment :

$$iR + e_c + L \frac{di}{dt} = \sqrt{2} E_e \sin \omega t$$

soit :

$$(1) \quad iR + e_c + L \omega \frac{di}{dx} = \sqrt{2} E_e \sin x$$

1) Régime de conduction discontinue.

En régime de conduction discontinue, les deux thyratrons sont éteints en même temps pendant un

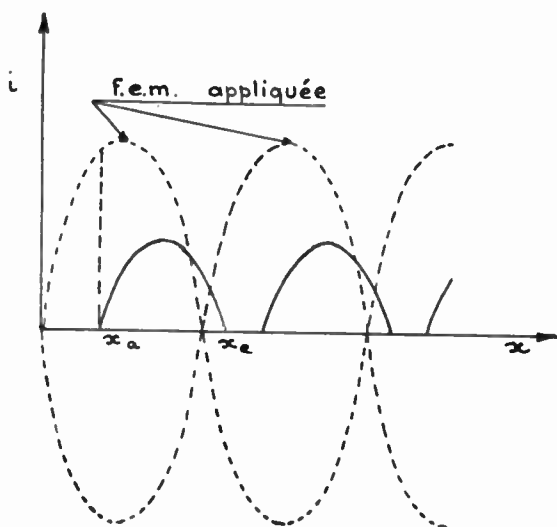


FIG. 2

certain laps de temps. Le courant a la forme indiquée sur la figure 2. Chaque thyatron conduit indé-

pendamment de l'autre et chaque impulsion de courant est indépendante des précédentes.

Si l'on intègre l'équation 1 en supposant que le thyatron commence à conduire pour $x = x_a$, on obtient l'équation du courant instantané i sous la forme :

$$(2) \quad Ri = \sqrt{2} E_e \cos \theta [\sin (x - \theta) - e^{-\frac{x-x_a}{tg \theta}} \sin (x_a - \theta)] - e_c \left(1 - e^{-\frac{x-x_a}{tg \theta}} \right)$$

L'angle d'extinction x_e est déterminé en faisant $i = 0$ pour $x = x_e$. En intégrant Ri entre les limites x_a et x_e , et compte tenu du fait que x_e est déterminé de la façon ci-dessus indiquée, on trouve :

$$\int_{x_a}^{x_e} Ri \, dx = \sqrt{2} E_e (\cos x_a - \cos x_e) - e_c (x_e - x_a)$$

La valeur moyenne du courant, en régime permanent, est évidemment obtenue en divisant par

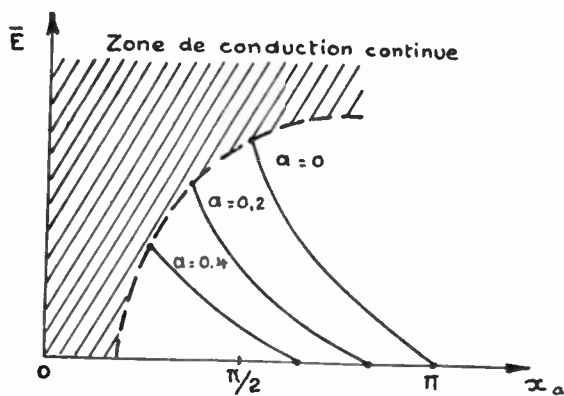


FIG. 3

π cette intégrale, puisque l'impulsion de courant se reproduit chaque fois que x_a augmente de π . On obtient l'expression connue :

$$(3) \quad RI = \frac{\sqrt{2} E_e}{\pi} (\cos x_a - \cos x_e) - e_c \frac{x_a - x_e}{\pi}$$

RI est donc en fonction de x_a et e_c (x_e n'est pas une variable indépendante. C'est une fonction de x_a et e_c). On peut tracer les courbes donnant RI en fonction de x_a , e_c étant le paramètre des courbes. En général, dans un souci de normalisation, pour avoir des courbes universelles, on trace plutôt les

courbes donnant $\bar{E} = \frac{RI}{\sqrt{2} E_e}$ en fonction de x_a . Le

paramètre est alors $a = \frac{e_c}{\sqrt{2} E_e}$ (fig. 3). Ces courbes sont limitées à une zone au delà de laquelle la conduction cesse d'être discontinue et devient continue. Seule la zone de conduction discontinue nous intéresse en ce moment.

Ces courbes sont valables pour un régime permanent, c'est-à-dire lorsque x_a , e_c , et RI sont constants.

Nous allons maintenant chercher ce qui se passe lorsque x_a est variable.

Lorsque x_a varie, la vitesse du moteur varie en général aussi puisque le couple électromagnétique, proportionnel à I , varie. Pour nous permettre de linéariser le problème, nous supposons que la variation de x_a est petite. Celles de e_c et de RI le seront aussi :

Finalement, les variations de RI seront fonctions linéaires de celles de x_a et e_c .

Nous poserons :

$$\begin{aligned}x_a &= x_{a0} + \Delta x_a \\e_c &= e_{c0} + \Delta e_c \\RI &= RI_0 + \Delta(RI)\end{aligned}$$

Δx_a , Δe_c et $\Delta(RI)$ sont des fonctions du temps. Il faut revenir un instant sur la définition de RI . En régime permanent, nous avons défini RI , chute de tension ohmique moyenne, comme la moyenne de Ri prise sur une variation égale à π de l'angle x , ce qui était légitime car les impulsions de courant se succédaient à une distance de π les unes des autres. Ici, nous ne devons plus prendre la moyenne sur une variation π de x , car, x_a étant variable l'intervalle qui sépare le début de deux impulsions successives n'est plus π mais $\pi + d(\Delta x_a)$, $d(\Delta x_a)$ étant la variation de l'angle d'allumage entre deux allumages consécutifs.

On écrira donc :

$$RI = \frac{\int_{x_a}^{x_e} Ri \, dx}{\pi + d(\Delta x_a)} \approx \frac{1 - \frac{d(\Delta x_a)}{\pi}}{\pi} \int_{x_a}^{x_e} Ri \, dx$$

La chute ohmique RI ainsi définie sera considérée comme la valeur instantanée de la chute ohmique dans l'induit vis-à-vis de la modulation de x_a .

Nous considérerons les variations Δx_a de x_a comme lentes vis-à-vis de la fréquence du courant d'alimentation, c'est-à-dire $d(\Delta x_a)$ petit devant π .

Nous pouvons récrire l'équation 3 pour RI_0 :

$$RI_0 = \frac{\sqrt{2} E_e}{\pi} (\cos x_{a0} - \cos x_{e0}) - e_{c0} \frac{x_{a0} - x_{e0}}{\pi}$$

Au temps t , x_a a pris l'accroissement Δx_a , e_c , l'accroissement Δe_c et RI l'accroissement $\Delta(RI)$. La moyenne de RI n'est plus calculée sur π , mais sur $\pi + d(\Delta x_a)$. On pourra écrire :

$$\begin{aligned}RI &= \frac{1 - \frac{d(\Delta x_a)}{\pi}}{\pi} \sqrt{2} E_e [\cos x_a - \cos x_e] \\&\quad - \frac{1 - \frac{d(\Delta x_a)}{\pi}}{\pi} e_c (x_e - x_a)\end{aligned}$$

avec $x_a = x_{a0} + \Delta x_a$

$$e_c = e_{c0} + \Delta e_c$$

$$RI = RI_0 + \Delta(RI)$$

En négligeant les termes du second ordre, il viendra finalement :

$$(4) \quad \Delta(RI) = \left[\frac{\partial(RI)}{\partial x_a} \right]_{x_{a0}, e_{c0}} \Delta x_a + \left[\frac{\partial(RI)}{\partial e_c} \right]_{x_{a0}, e_{c0}} \Delta e_c - RI_0 \frac{d(\Delta x_a)}{\pi}$$

La variation de RI sera donc fonction non seulement des variations de x_a et e_c , mais aussi de $d(\Delta x_a)$. Nous allons chercher la valeur de ce terme. C'est la variation de Δx_a entre deux allumages consécutifs ; $d(\Delta x_a) = \frac{\partial(\Delta x_a)}{\partial t} \tau$, τ étant le temps qui sépare deux allumages consécutifs. Comme les variations de Δx_a sont lentes, on peut prendre pour τ la valeur $\frac{\pi}{\omega}$. On aura donc $d(\Delta x_a) = \frac{\pi}{\omega} \frac{\partial(\Delta x_a)}{\partial t}$ soit si nous employons la notation symbolique : $d(\Delta x_a) = \frac{\pi}{\omega} p \Delta x_a$.

L'équation 4 représente donc finalement la relation cherchée reliant ΔRI à Δx_a et Δe_c .

Pour la mettre sous une forme plus classique, nous ferons les remarques suivantes :

1° Le terme $\frac{\partial(RI)}{\partial x_a}$ est la pente de la courbe de paramètre e_{c0} au point x_{a0} . Elle est négative, nous l'appellerons $(-\alpha)$.

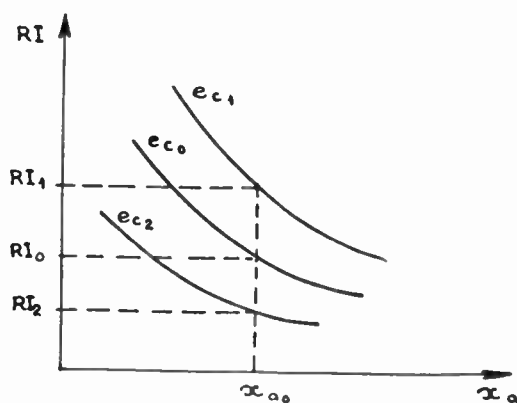


FIG. 4

2° Le terme $\frac{\partial(RI)}{\partial e_c}$ est facile à mettre en évidence sur les courbes donnant RI en fonction de x_a , paramétrées en e_c (fig. 4). C'est le rapport $\frac{RI_1 - RI_2}{e_{c1} - e_{c2}}$. Il est également négatif, car les e_c décroissent du bas vers le haut. Nous l'appellerons $(-\beta)$.

3° Le fait que RI décroisse en fonction de x_a invite à prendre comme grandeur de commande non pas Δx_a mais $(-\Delta x_a)$.

On écrira donc finalement :

$$(5) \quad \Delta(RI) = (\alpha + \frac{RI_0}{\omega} p) (-\Delta x_a) - \beta \Delta e_c$$

et on déduit immédiatement de cette relation le schéma équivalent de l'ensemble thyratrons-moteur (fig. 5).

(— Δr_a) est la grandeur de commande. Elle arrive, multipliée par le terme $\alpha + RI_0 p$, au discriminateur, où lui est soustraite la grandeur $\beta \Delta e_c = \beta K_e \Delta \Omega$. La différence, $\Delta(RI)$, est appliquée au moteur M . La variation $\Delta \Omega$ de la vitesse est liée à $\Delta(RI)$ par

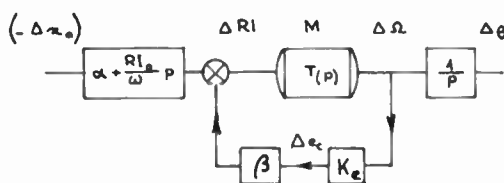


FIG. 5

la fonction de transfert $T(p)$ qui est en général facile à calculer. Par exemple, si l'inertie totale sur l'arbre est I , le frottement visqueux f , et le coefficient de couple $k_T = \frac{\Gamma}{I}$, on aura :

$$T(p) = \frac{\Delta \Omega}{\Delta(RI)} = \frac{K_T}{J_p + f}$$

Dans ce schéma, la self du moteur n'apparaît pas. Ce fait est dû à la conduction discontinue des thyratrons qui empêche à chaque instant le courant de « se souvenir » de ses états précédents.

Il est intéressant de constater, dans la chaîne directe, la présence du terme $\frac{RI_0 p}{\omega}$ qui favorise

la stabilité. Cependant, il faut remarquer que $\frac{RI_0}{\omega}$ est nécessairement assez petit en raison de la présence du terme ω au dénominateur. L'effet stabilisant ne se manifeste donc que pour les valeurs de p assez grandes, c'est-à-dire pour les asservissements à bande passante étendue, et dont la fréquence de coupure est une fraction importante de ω .

Si l'on utilise les courbes normalisées de la figure 3 donnant $\bar{E} = \frac{IR}{\sqrt{2} E_e}$ en fonction de x_a et de $a = \frac{\sqrt{2} E_e}{e_c}$ on relèvera directement les coefficients :

$$\frac{\partial \bar{E}}{\partial x_a} = -A \quad \frac{\partial \bar{E}}{\partial a} = -B$$

On en déduira α et β par les relations :

$$\alpha = A \sqrt{2} E_e \quad \beta = B$$

2) Régime de conduction continue.

Lorsque l'angle d'extinction x_e d'un thyatron, calculé à partir de la formule 2, où l'on fait $Ri = 0$, devient supérieur à $x_a + \pi$ il se produit d'un thyatron à l'autre un phénomène de commutation qui

fait que l'allumage du thyatron qui était éteint provoque l'extinction du thyatron qui conduisait. En effet, l'allumage d'un thyatron provoque une inversion de la force contre-électromotrice due à la self, qui s'opposait à l'extinction du thyatron conducteur. Celui-ci s'éteint donc aussitôt. Il n'y a donc plus arrêt de la conduction totale, mais commutation de la conduction d'un thyatron sur l'autre aux instants successifs déterminés par l'angle d'allumage x_a . Le courant ne s'annule plus. Il a la forme indiquée sur la figure 6. L'angle réel d'extinction d'un thyatron n'est donc pas l'angle calculé ; il vaut : $x_e = x_a + \pi$.

L'intégration de l'équation 1 doit tenir compte du fait que le courant initial n'est plus nul, mais a une valeur non nulle i_m . Dans ces conditions, on trouve comme expression de la chute ohmique instantanée :

$$(6) \quad iR = \sqrt{2} E_e \cos \theta [\sin(x - \theta) - e^{-\frac{x-x_a}{\tau g \theta}} \sin(x_a - \theta)] - e_c \left(1 - e^{-\frac{x-x_a}{\tau g \theta}}\right) + Ri_m e^{-\frac{x-x_a}{\tau g \theta}}$$

Le régime permanent se calcule en écrivant que pour $x = x_a + \pi$, on a $iR = i_m R$. On trouve ainsi

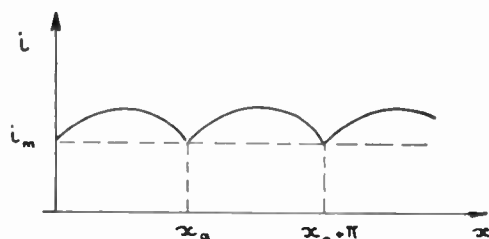


FIG. 6

la valeur de Ri_m en fonction de x_a . A partir de cette valeur de Ri_m , si l'on calcule, en régime permanent, la moyenne de iR sur une période, on trouve :

$$IR = \frac{1}{\pi} \int_{x_a}^{x_a + \pi} iR dx = \frac{2 \sqrt{2} E_e}{\pi} \cos x_a - e_c$$

C'est la même équation que celle d'un moteur dont l'induit est alimenté par une tension continue V , en posant :

$$V = \frac{2 \sqrt{2} E_e}{\pi} \cos x_a$$

Si l'on passe au régime variable, il semble, a priori, que, contrairement au cas de la conduction discontinue, il faudra tenir compte de la self, car une impulsion est liée à la précédente par la valeur du courant instantané au moment de la commutation. Mais, en raison de la superposition d'une composante continue et d'un courant pulsé, on peut se demander s'il ne sera pas nécessaire de multiplier cette self par un coefficient convenable pour tenir compte de ce fait.

Pour étudier l'action de la self, nous allons supposer que, le régime permanent étant atteint, on donne

un brusque accroissement à x_a , tout en maintenant la vitesse (c'est-à-dire e_c) constante. On crée ainsi un échelon de la grandeur de commande.

On va étudier la forme du courant qui en résulte (fig. 7).

Soient x_{a0} , x_{a1} les angles d'allumage initial et final.

Nous négligerons la période comprise entre x_{a0} et x_{a1} . À partir de x_{a1} , la chute de tension ohmique aux instants successifs de commutation s'exprime

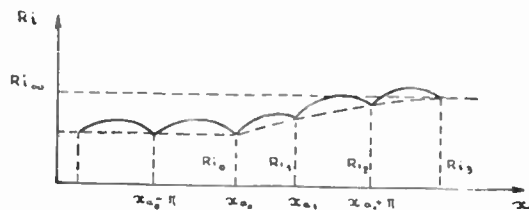


FIG. 7

par la somme d'un terme qui ne dépend que de x_{a1} , c'est-à-dire d'un terme constant que nous appellerons A , et d'un terme dépendant du courant de commutation précédent (équation 6). On peut écrire :

$$Ri_{m,n} = A + Ri_{m,n-1} e^{-\frac{\pi}{tg \theta}}$$

On tire de cette expression :

$$Ri_{m,n} = A \frac{1 - e^{-\frac{(n-1)\pi}{tg \theta}}}{1 - e^{-\frac{\pi}{tg \theta}}} + Ri_{m,1} e^{-\frac{(n-1)\pi}{tg \theta}}$$

Cette expression tend vers :

$$Ri_{m,\infty} = \frac{A}{1 - e^{-\frac{\pi}{tg \theta}}}$$

qui représente la nouvelle chute ohmique de commutation permanente. Si l'on forme l'expression $\frac{Ri_{m,\infty} - Ri_{m,n}}{Ri_{m,\infty} - Ri_{m,1}}$, on s'aperçoit qu'elle vaut $e^{-\frac{(n-1)\pi}{tg \theta}}$.

Comme on a $\pi = \omega T$ (T étant la demi-période du courant d'alimentation) $tg \theta = \frac{L\omega}{R}$, si l'on pose $\frac{L}{R} = \tau$ et $(n-1)T = t$, on a finalement :

$$\frac{Ri_{m,\infty} - Ri_{m,n}}{Ri_{m,\infty} - Ri_{m,1}} = e^{-\frac{t}{\tau}}$$

C'est-à-dire que le courant de commutation s'établit suivant la même loi qu'un courant continu passant dans l'induit.

Si maintenant nous passons au courant moyen (nous appellerons RI_n le courant moyen entre les instants n et $n+1$), on voit de même en intégrant l'expression 6 entre x_{a1} et $x_{a1} + \pi$ que RI_n sera composé d'un terme constant et d'un terme comprenant $Ri_{m,n}$.

D'où l'on tire immédiatement :

$$\frac{RI_{\infty} - RI_n}{RI_{\infty} - RI_1} = \frac{Ri_{m,\infty} - Ri_{m,n}}{Ri_{m,\infty} - Ri_{m,1}} = e^{-\frac{t}{\tau}}$$

Le courant moyen s'établit donc suivant la même loi que le courant de commutation, et il faut bien prendre, dans l'étude du comportement en régime variable du système, la self de l'induit avec sa valeur réelle. Comme la force électromotrice est appliquée en permanence, la chute de tension moyenne dans l'induit, soit :

$$RI + e_c \text{ moyen} + L \frac{dI}{dt} \text{ moyen}$$

est égale à la moyenne de la tension appliquée. Si x_a varie autour d'une valeur moyenne x_{a0} , on pourra écrire :

$$x_a = x_{a0} + \Delta x_a$$

$$e_c = e_{c0} + \Delta e_c$$

$$RI = RI_0 + \Delta(RI)$$

La conduction commence à $x_{a0} + \Delta x_a$ et finit à $x_{a0} + \pi + \Delta x_a + d(\Delta x_a)$, $d(\Delta x_a)$ étant la variation de Δx_a pendant une période de conduction

(on a vu que $d(\Delta x_a) = \frac{\pi}{\omega} p \Delta x_a$).

On écrira donc :

$$\begin{aligned} [RI_0 + \Delta(RI)] + (e_c + \Delta e_c) + L \frac{d(\Delta I)}{dt} \\ = \frac{E e \sqrt{2}}{\pi + d(\Delta x_a)} \int_{x_{a0} + \Delta x_a}^{x_{a0} + \pi + \Delta x_a + d(\Delta x_a)} \sin x \, dx \end{aligned}$$

En négligeant les termes du second ordre, on trouve, en employant la notation symbolique et en posant $V_0 = \frac{E e \sqrt{2}}{\pi} \cos x_{a0}$:

$$\Delta(RI) + \Delta e_c + pL \Delta I = \left(\frac{\partial V_0}{\partial x_{a0}} - \frac{V_0}{\omega} p \right) \Delta x_a$$

En remarquant encore que $\frac{\partial V_0}{\partial x_{a0}}$ est négatif et que $(-\Delta x_a)$ est la véritable grandeur de commande, on peut écrire, en posant $\frac{\partial V_0}{\partial x_{a0}} = -\alpha$:

$$\Delta(RI) + \Delta e_c + pL \Delta I = \left(\alpha + \frac{V_0}{\omega} p \right) (-\Delta x_a)$$

À part la présence du terme $\frac{V_0 p}{\omega}$, cette équation est la même que celle d'un moteur à courant continu alimenté par une tension variable ΔU , dont la valeur serait $(-\alpha \Delta x_a) = \Delta U$. Le fait d'alimenter le moteur par des thyratrons ne change donc rien à son fonctionnement habituel au point de vue

fréquentiel. Il a seulement comme résultat d'introduire une tension d'alimentation variable avec Δx_a .

Si l'on tient compte du terme $\frac{V_0}{\omega} p$, on s'aperçoit qu'en réalité la source de tension présente une légère avance de phase sur la grandeur de commande Δx_a . Cependant, cette avance de phase est toujours faible, car, comme nous l'avons remarqué au paragraphe précédent, la présence du facteur ω au dénominateur rend le coefficient de p petit. On peut d'ailleurs calculer de façon plus précise le terme $y = (\alpha + \frac{V_0}{\omega} p)$. On a :

$$\alpha = - \frac{\partial V_0}{\partial x_{a0}} = \frac{E e \sqrt{2}}{\pi} \sin x_{a0}$$

$$\text{Donc } y = \frac{E e \sqrt{2}}{\pi} \left(\sin x_{a0} + \frac{\cos x_{a0}}{\omega} p \right)$$

L'influence du terme en p ne devient donc sensible que pour les valeurs de p voisines de $\omega \operatorname{tg} x_{a0}$. Comme, en régime de conduction continue, on a toujours $x_{a0} < \frac{\pi}{2}$ et que d'autre part x_{a0} ne peut jamais être très petit car on doit avoir $x_{a0} > \operatorname{arc} \sin a$ (sans quoi la tension appliquée sur la plaque du thyatron au moment de l'impulsion de commande

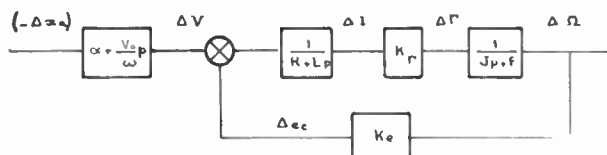


FIG. 8

qui est la différence entre la *f.e.m.* sinusoïdale et la *f.c.e.m.* du moteur, serait négative et le thyatron ne s'allumerait pas), $\omega \operatorname{tg} x_{a0}$ est de l'ordre de grandeur de ω et p est en général petit devant ce terme. On pourra donc souvent négliger le terme en p devant le terme constant.

Le schéma équivalent de l'ensemble thyratrons-moteur pour ce mode de conduction est représenté sur la figure 8.

C'est le schéma d'un moteur continu alimenté par une tension ΔV . $\Delta \Gamma$ représente la variation de couple, k_r le coefficient de couple, J et f l'inertie et le frottement visqueux de l'induit.

II. — ETUDE DU CAS GÉNÉRAL.

Dans cette deuxième partie, nous supposerons que les conditions sont telles qu'on ne peut pas linéariser.

Nous supposerons que nous avons affaire, par exemple, à un montage à deux sens de marche du type indiqué sur la figure 9. Si l'on veut éviter que les deux groupes de thyratrons débitent l'un dans l'autre au point zéro, on peut être amené à aménager un seuil qui rend la commande très non-linéaire.

Nous appellerons x la grandeur de commande à laquelle sont reliés les angles d'allumage des différents thyratrons. Selon que x est positif ou négatif, une paire de thyratrons ou l'autre conduit. La conduction est supposée discontinue, c'est-à-dire que la self du moteur n'intervient pas.

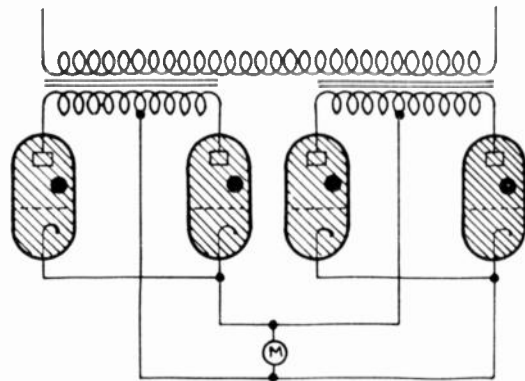


FIG. 9

Si l'on trace sur un graphique la valeur de la chute ohmique $IR = E$ en fonction de x et de e_c , ou plutôt, dans le même but de normalisation que précédemment, la fonction $\bar{E} = \frac{E}{\sqrt{2} E e}$ en fonction

de x et de $a = \frac{e_c}{\sqrt{2} E e}$ on obtiendra un réseau de courbes du type de celui représenté sur la figure 10. Il faut bien faire attention, ici, que a peut prendre des valeurs négatives, car la force contre électromotrice du moteur peut s'ajouter à, et non seulement se retrancher de, la tension appliquée, lorsqu'un groupe de thyratrons débite dans un sens tendant à freiner le moteur.

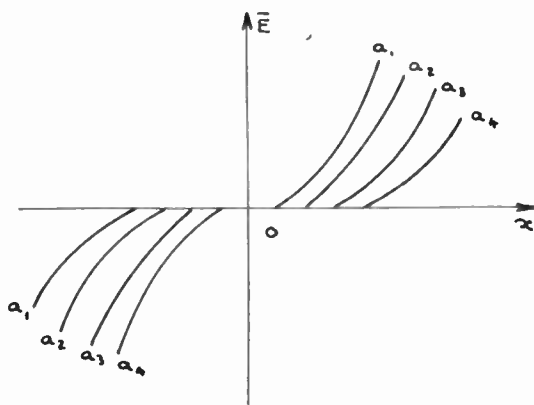


FIG. 10

On voit qu'autour du point 0 il est difficile de linéariser. Le schéma équivalent de l'ensemble thyratrons-moteur se présente sous la forme indiquée figure 11.

D est une sorte de « discriminateur non linéaire » qui doit rendre compte de l'effet combiné de la grandeur de commande x et de la grandeur en réaction a . Les propriétés de D sont connues par les courbes de la figure 10.

Le problème étant non linéaire, nous appliquerons la méthode consistant à assimiler toute variation périodique à son fondamental. On sait que cette façon d'opérer est justifiée par le fait que la fonction de transfert du moteur, $\mu(p)$, se comporte comme celle d'un filtre passe-bas qui affaiblit les harmoniques.

Nous cherchons à déterminer la fonction de transfert de la boucle fermée de la figure 11 pour pouvoir l'intégrer dans une boucle d'asservissement.

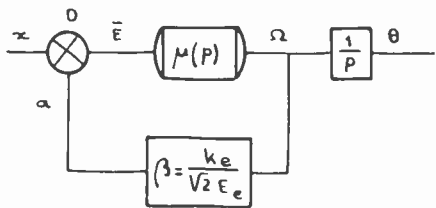


FIG. 11

C'est-à-dire que si l'on donne à x une variation sinusoïdale :

$$x = X \sin \omega t \quad (\text{ne pas confondre } \omega \text{ avec la pulsation de la tension appliquée du paragraphe précédent})$$

Ω variera suivant une loi périodique que l'on assimilera à son fondamental

$$\Omega = \Omega_0 \sin (\omega t + \varphi)$$

et on cherche à connaître les deux quantités $\frac{\Omega_0}{X}$ et φ en fonction de ω .

Il revient au même d'ailleurs d'étudier la variation de a en fonction de x puisque Ω et a sont reliés par un coefficient constant, et cela est plus facile puisqu'on atteint a directement sur les courbes de la figure 10.

On posera donc : $a = A \sin (\omega t + \varphi)$ et on étudiera les quantités $\frac{A}{X}$ et φ en fonction de ω . Il est évident que, le réseau des courbes de la figure 10 n'étant pas linéaire, ces quantités dépendront de l'amplitude X ; finalement, la fonction de transfert de la boucle fermée, dont le module est $\frac{A}{X}$ et l'argument φ , sera fonction de X et ω . On sait que la connaissance d'une telle fonction de transfert dans une boucle d'asservissement permet d'en étudier la stabilité.

Il est nécessaire, et c'est là la principale difficulté de la méthode que nous allons exposer, de connaître le comportement du discriminateur D au point de vue fréquentiel, c'est-à-dire de connaître le fondamental de $\bar{E}$ lorsqu'on l'excite par des valeurs de x et a sinusoïdales de phases et d'amplitudes quelconques. On aura, en assimilant $\bar{E}$ à son fondamental :

$$\bar{E} = \bar{E}_m \sin (\omega t + \psi)$$

lorsque $x = X \sin \omega t$ et $a = A \sin (\omega t + \varphi)$.

La connaissance du discriminateur D du point de vue fréquentiel nécessite donc un double jeu de réseaux de courbes donnant :

$$\begin{aligned} \bar{E}_m &= \bar{E}_m (X, A, \varphi) \\ \psi &= \psi (X, A, \varphi) \end{aligned}$$

($\bar{E}_m$ et ψ ne dépendront pas de ω puisque, la conduction étant supposée discontinue, les thyratrons ne présentent pas de constante de temps).

On aura intérêt, en pratique, à tracer les courbes représentant $\frac{A}{\bar{E}_m}$ et ψ en fonction de φ , A étant un paramètre, pour différentes valeurs de X . On peut les calculer à partir des courbes de la figure 10.

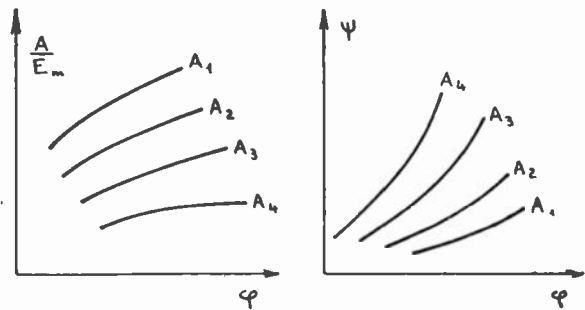


FIG. 12

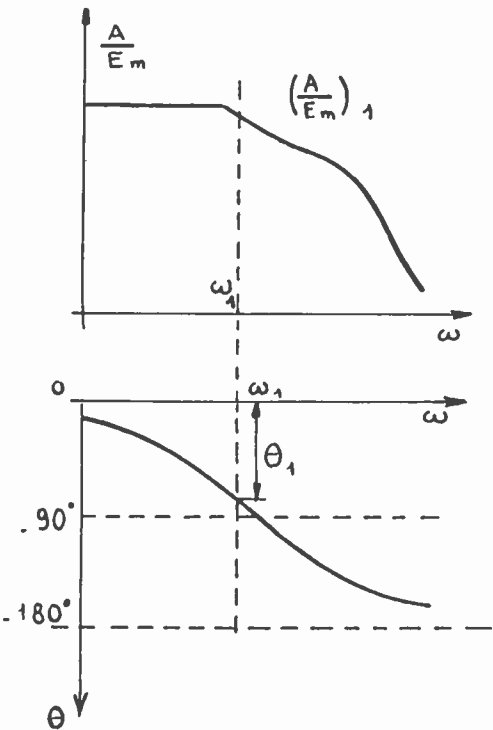


FIG. 13

Pour une valeur de X égale à X_1 , on aura donc deux réseaux de courbes du type indiqué sur la figure 12.

En plus de ces courbes qui caractérisent le « discriminateur non linéaire » D , nous connaissons les courbes de gain et de phase de la fonction de trans-

fert $\beta\mu(p)$, c'est-à-dire le rapport $\frac{A}{E_m}$ et la phase θ de a par rapport à E , soit $\theta = \varphi - \psi$ en fonction de ω (fig. 13).

Pour déterminer les quantités qui nous intéressent, c'est-à-dire le rapport $\frac{A}{X}$ et l'angle φ , en fonction de ω et X , nous opérerons de la façon suivante. Parmi les graphiques caractéristiques du discriminateur D , nous choisissons ceux qui correspondent à la valeur $X = X_1$ (fig. 14).

Sur les graphiques donnant $\frac{A}{E_m}$ et θ en fonction de ω (fig. 13), nous relevons les valeurs $\left(\frac{A}{E_m}\right)_1$ et θ_1 qui correspondent à une valeur ω_1 de ω .

Sur les graphiques de la figure 14 nous traçons respectivement les droites $\left(\frac{A}{E_m}\right) = \left(\frac{A}{E_m}\right)_1$ et $\psi = \varphi - \theta_1$. Par interpolation entre les valeurs du paramètre A , on trouve facilement sur chaque

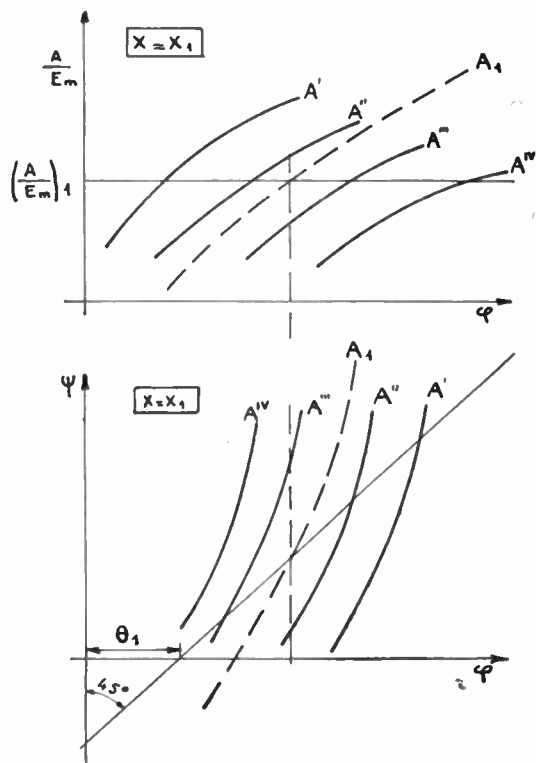


Fig. 14

graphique le point qui correspond au même A et au même φ ; soient A_1 et φ_1 , ces valeurs.

On a donc ainsi les valeurs A_1 et φ_1 , qui correspondent à X_1 et ω_1 , d'où $\frac{A_1}{X_1}$ et φ_1 correspondant à X_1 et ω_1 . En répétant cette opération pour plusieurs valeurs de X_1 , puis de ω_1 , on pourra finalement tracer les réseaux de courbes $\frac{A}{X}(\omega)$ et $\varphi(\omega)$, paramétrées en X , qui représentent la fonction de transfert cherchée (fig. 15).

Conclusion.

La méthode exposée ci-dessus permet de résoudre théoriquement le problème des asservissements à thyatronns quand il est impossible de linéariser de façon satisfaisante. Elle est valable également dans tous les cas où il existe dans un asservissement un « discriminateur non linéaire », c'est-à-dire un organe à deux entrées, dont l'une est attaquée directement et l'autre en réaction, et qui fournit une grandeur fonction non linéaire des deux précédentes.

Cette méthode n'est pas très compliquée si l'on dispose des réseaux de courbes de la figure 12, car la détermination d'un point du réseau de courbes de la figure 15 est simple et rapide.

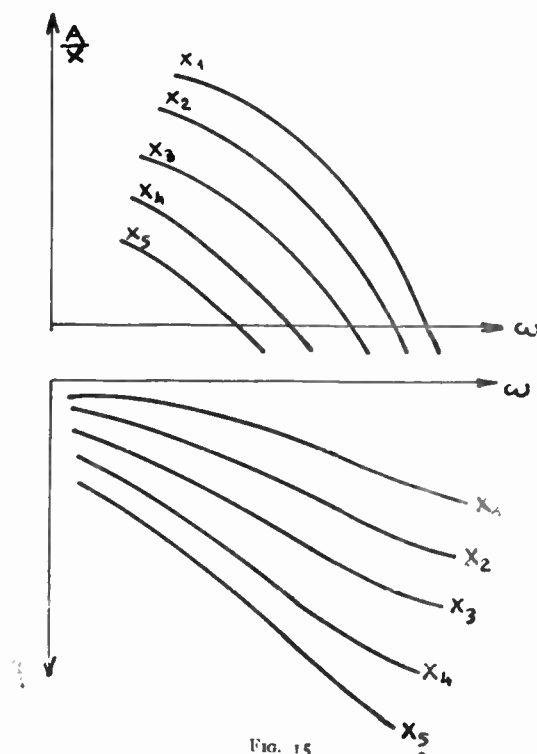


Fig. 15

Mais la détermination des courbes de la figure 12 (qui semblent cependant absolument nécessaires puisqu'elles correspondent à la connaissance complète du comportement du discriminateur en régime fréquentiel) nécessite un travail souvent hors de proportion avec le résultat à atteindre, et on préférera dans la plupart des cas réaliser le montage et relever expérimentalement les courbes de la figure 15 plutôt que d'effectuer ces calculs compliqués.

Il semble donc qu'on soit ici à la limite des possibilités de généralisation de l'analyse fréquentielle aux systèmes non-linéaires.

Il serait intéressant de savoir si une telle méthode, malgré les calculs importants qu'elle nécessite, peut rendre des services dans le cas d'asservissements importants très non-linéaires.

On peut cependant remarquer que l'établissement du réseau de courbes du discriminateur non linéaire peut se faire expérimentalement en étudiant sa réponse à deux excitations sinusoïdales d'amplitude et phase variables. Mais de toutes façons cette étude serait très longue et devrait être justifiée par une impossibilité d'opérer autrement.

SUR LA RÉALISATION D'UN COEFFICIENT DE TRANSFERT DONNÉ

ou les montages d'amplificateurs à réaction

PAR

F. H. RAYMOND

Directeur de la Société d'Electronique
et d'Automatisme

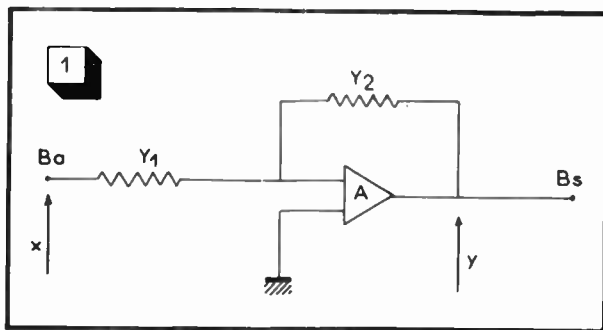
I. — INTRODUCTION.

La synthèse par un réseau électrique d'une fonction de transfert est un problème général dont les applications sont nombreuses. On notera que les procédés mathématiques ou pratiques utilisés sont adaptés à des domaines définis : filtres, réseaux correcteurs des amplificateurs à réaction ou servomécanismes par exemple.

Directement ou indirectement un problème peut être traduit par la donnée des singularités et zéros d'une fonction de transfert. Nous signalerons une tentative importante de WILLIAM H. KAUTZ (1) destinée à réaliser cette synthèse lorsque la donnée est une fonction temporelle (la réponse indicielle par exemple).

Dans le cas des asservissements ou des filtres à très basses fréquences, en particulier, l'objet de cette note est de proposer une méthode générale dont le principe a été donné à des fins différentes, dans un travail antérieur (2).

On a l'habitude dans le domaine du calcul analogique de considérer des schémas analogues à celui de la figure 1 où A est un amplificateur à grand gain. Les admittances isomorphes Y_1 et Y_2 sont placées entre les bornes d'accès B_a et de sortie B_s du système, respectivement, et la borne d'entrée de l'amplificateur. En admettant que sont infinis l'admittance d'entrée de l'amplificateur et son gain



interne, nous écrivons (relation entre transformées de Laplace de $x(t)$ et $y(t)$) :

$$x Y_1 + y Y_2 = 0$$

d'où le coefficient de transfert :

$$(1) \quad f(p) = \frac{y}{x} = - \frac{Y_1(p)}{Y_2(p)}$$

En rejetant toute admittance Y_1 qui aurait un zéro en p , impédance infinie pour les signaux de fréquence nulle, le tableau ci-dessous donne un certain nombre de cas typiques d'usage courant. (Voir page 106).

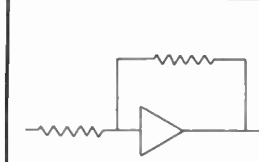
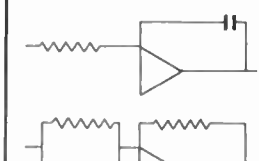
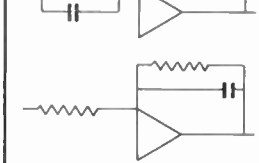
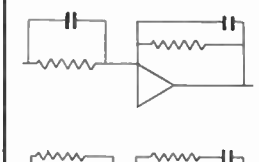
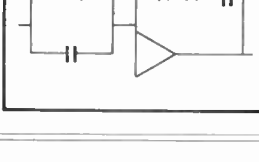
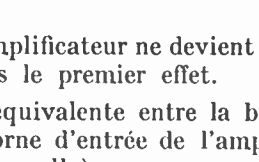
Dans une certaine mesure, le module des admittances Y_1 et Y_2 peut être réglé par un potentiomètre, ainsi qu'il est représenté figure 2 ; on a alors :

$$(2) \quad \frac{y}{x} = - \frac{a}{b} \frac{Y_1}{Y_2}$$

Ceci n'est valable que dans la mesure où la résistance du potentiomètre n'est pas trop grande et

(1) William H. Kautz « Network Synthesis for specified transient response » *Tech. Report N° 209* — 25 avril 1952 — Research Lab. of Electronics, M.I.T.

(2) F. H. RAYMOND. — C. R. du Colloque « Les machines à calculer et la pensée humaine » Paris 8-13 janvier 1951 — p. 185 « Conceptions générales d'opérateurs mathématiques électroniques » — Voir aussi supplémento à « *Ricerca Scientifica* » ann° 22° — 1952 — page 77 (Chapitre III).

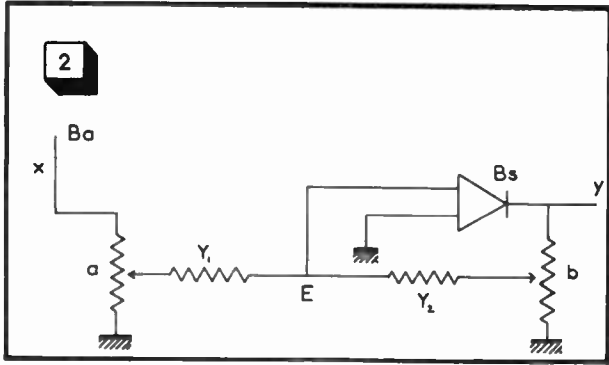
	entrée Y_1 (ou Z_1)	réaction Y_2 (ou Z_2)	coefficient de transfert (au signe près) $-f = \frac{Y_1}{Y_2} = Y_1 Z_2 = \frac{1}{Z_1 Y_2} = \frac{Z_2}{Z_1}$
	$Z_1 = R_1$	$Z_2 = R_2$	$\frac{R_2}{R_1}$
	$Z_1 = R_1$	$Y_2 = C_2 p$	$\frac{1}{R_1 C_2 p}$
	$Y_1 = \frac{1}{R_1} + C_1 p$	$Z_2 = R_2$	$\frac{R_2}{R_1} + R_2 C_1 p = \frac{R_2}{R_1} (1 + R_1 C_1 p)$
	$Z_1 = R_1$	$Y_2 = \frac{1}{R_2} + C_2 p$	$\frac{1}{\frac{R_1}{R_2} + R_1 C_2 p} = \frac{R_2}{R_1} \frac{1}{1 + R_2 C_2 p}$
	$Y_1 = \frac{1}{R_1} + C_1 p$	$Y_2 = \frac{1}{R_2} + C_2 p$	$\frac{R_2}{R_1} \frac{1 + R_1 C_1 p}{1 + R_2 C_2 p}$
	$Y_1 = \frac{1}{R_1} + C_1 p$	$Z_2 = R_2 + \frac{1}{C_2 p}$	$\left(\frac{1}{R_1} + C_1 p \right) \left(R_2 + \frac{1}{C_2 p} \right) = \frac{R_2}{R_1} \cdot \frac{(1 + R_1 C_1 p)(1 + R_2 C_2 p)}{R_2 C_2 p}$

où le bruit de l'amplificateur ne devient pas inadmissible. Considérons le premier effet.

L'admittance équivalente entre la borne *Ba* par exemple et la borne d'entrée de l'amplificateur est (théorème de Kennelly) :

$$Y_1(a) = \frac{\frac{1}{(1-a)\rho} Y}{\frac{1}{(1-a)\rho} + \frac{1}{a\rho} + Y}$$

où ρ est la résistance totale du potentiomètre, l'index définit une résistance $a\rho$ entre lui et



le point bas du potentiomètre et une résistance $(1-a)\rho$ entre lui et le point haut. On peut écrire :

$$Y_1(a) = aY \frac{1}{1 + (1-a)a\rho Y}$$

Si on pose : $\epsilon = \rho Y$

$$(3) \quad Y_1(a) = aY \frac{1}{1 + (1-a)a\epsilon}$$

Une relation analogue est valable pour $Y_2(b)$.

Cette relation met en évidence l'approximation admise dans (2), elle implique :

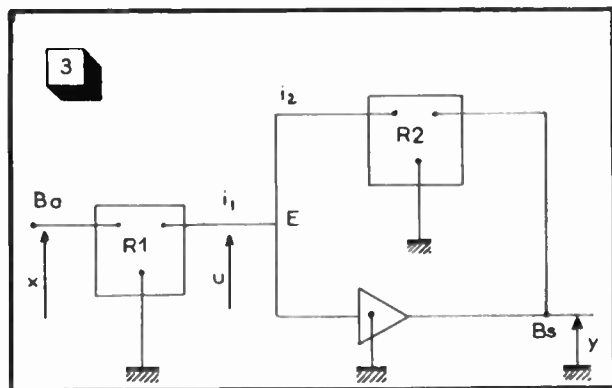
$$| (1-a)a\epsilon(p) | < \mu \% \text{ où } \mu \text{ est fixé.}$$

Nous pensons superflue toute démonstration dans laquelle on établirait qu'avec des montages des types précédents mettant en œuvre des admittances formées par des combinaisons quelconques de résistances et de capacités, le coefficient de transfert $f(p)$ est une fonction rationnelle dont les zéros et pôles sont réels et négatifs. La démonstration se ferait en usant des théorèmes connus

sur les propriétés et la synthèse des *impédances* ne comportant que des éléments R et C .

Le recours à des quadripôles ne modifie pas considérablement l'aspect du problème (référons-nous à la fig. 3).

Si on exprime pour chaque quadripôle supposé



inéditable et passif le courant y pénétrant par la borne E dont le potentiel est u , nous avons :

$$(8) \quad \begin{cases} i_1 = Y_{11} u + Y_{12} x \\ i_2 = Y_{11} u + Y'_{12} y \end{cases}$$

Le fonctionnement avec l'amplificateur supposé idéal (gain infini, admittance d'entrée infinie) impose les conditions $i_1 + i_2 = 0$ et $u = 0$, on a donc :

$$(4) \quad Y_{12} x + Y'_{12} y = 0$$

donc le coefficient de transfert réalisé à pour valeur :

$$(5) \quad f(p) = - \frac{Y_{12}}{Y'_{12}}$$

La synthèse de chacun des quadripôles revient donc à la synthèse de réseaux ayant entre deux bornes une admittance de transfert donnée. C'est un problème classique.

Dans un certain nombre d'applications où les fréquences sont basses ou très basses — servomécanismes en particulier — il est recommandé ou nécessaire impérativement d'avoir recours uniquement à des résistances et capacités pour réaliser la synthèse d'un quadripôle. Dans ce cas, on démontre ⁽¹⁾ que Y_{12} est une fonction rationnelle $Y_{12} = \frac{P(p)}{Q(p)}$ dont les pôles sont réels et négatifs et les zéros dans $R_p < 0$. Si donc on pose $Y'_{12} = \frac{P'}{Q'}$ on a :

$$f(p) = - \frac{P}{Q} \frac{Q'}{P'} \quad (p = \sigma + j\omega)$$

dont les zéros sont :

(1) Voir en particulier Aaron Fialkow et Erving Gerst « The transfert function of general two terminal-pair RC network », Quarterly of Applied Mathematics — Vol X — July 52 — N° 2 — pp. 113-127.

ceux de P dans $R_p < 0$

ceux de Q' sur le demi-axe $\omega = 0, \sigma < 0$

dont les pôles sont :

les zéros de Q : sur la demi-axe $\omega = 0, \sigma < 0$;

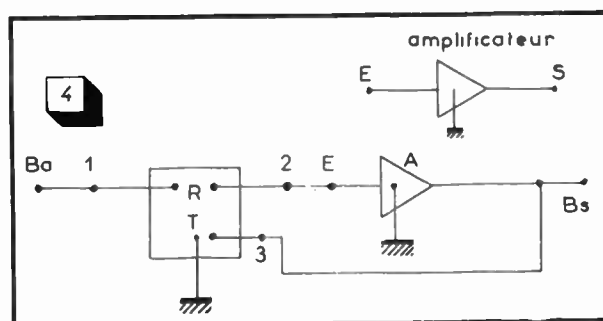
les zéros de P' dans σ ou $R_p < 0$

On peut donc réaliser une fonction de transfert théoriquement quelconque. Un ajustement des valeurs des éléments constituant les deux réseaux permet, en effet, de poser $Q = Q'$, dès lors, on atteint le degré le plus complet de généralité. Mais il est évident que cette solution n'a aucune valeur pratique sauf si on a recours à des éléments infiniment stables dans le temps.

II. — SOLUTION GÉNÉRALE.

Considérons un réseau passif, à liaisons holonomes, R , une borne T sert de référence aux potentiels, les bornes 1, 2 et 3 sont disponibles pour lier le réseau aux dispositifs choisis.

L'élément actif est un amplificateur A , dont l'entrée est connectée à la borne 2 et dont la sortie est reliée à 3 : ce faisant, mais sous la limitation que



l'amplificateur doit être considéré comme un *tripôle* dont une borne est un potentiel de T , le schéma considéré est le plus général qui permet d'introduire un élément actif dans un réseau R quelconque.

On notera que si l'amplificateur est à liaison par capacité ou transformateurs que le potentiel de référence pour A est sans influence sur les résultats. L'exposé qui est fait ici suppose que A est un amplificateur à liaisons directes (dit à courant continu).

Désignons par $g(p)$ le gain de A (gain interne en tension) de l'amplificateur. Soit y_e son admittance d'entrée et Z_s son impédance de sortie.

Soit V_g la tension d'entrée, la tension de sortie est, en désignant par i_s le courant débité par B_s :

$$g(p) V_g - Z_s i_s$$

L'équation générale du réseau peut être écrite (convention classique pour les i et les E) :

$$i_1 = y_{11} E_1 + y_{12} E_2 + y_{13} E_3$$

$$i_2 = y_{21} E_1 + y_{22} E_2 + y_{23} E_3$$

$$i_3 = y_{31} E_1 + y_{32} E_2 + y_{33} E_3$$

Dans l'hypothèse d'un amplificateur idéal, on a :

$$i_2 = 0, E_2 = 0$$

d'où :

$$(1) \quad f(p) = \frac{E_3}{E_1} = -\frac{y_{21}}{y_{23}}$$

Or, si on rappelle que y_{21} est le quotient du mineur Δ_{21} , du déterminant $\Delta(p)$ relatif aux sommets 1 et 2 du réseau, pareille remarque s'appliquant à y_{23} , on voit que $f(p)$ peut s'écrire :

$$(2) \quad f(p) = -\frac{\Delta_{21}}{\Delta_{23}}$$

qui est une fonction rationnelle dont les zéros et les pôles sont dans $Rp < 0$ et non localisés sur l'axe réel. En bref, ce qui se réalisait dans le montage de la figure 3 par une certaine identité des réseaux R_1 et R_2 (éventuellement par leur identité) est automatiquement obtenu ici, ce qui laisse prévoir, en outre, une économie d'éléments entrant dans le montage, toutes choses égales d'ailleurs.

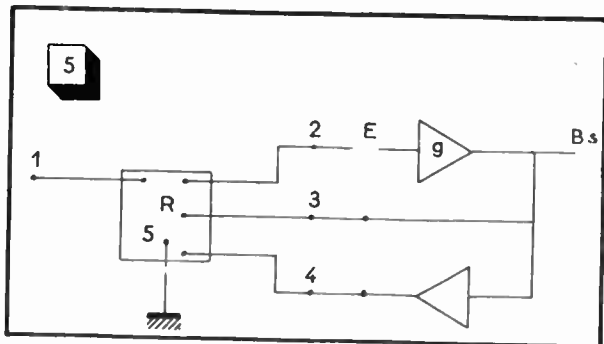
Le problème de synthèse est donc posé, pour le réseau R , il n'est pas classique et mériterait donc l'aide des spécialistes. Ainsi, les problèmes de réseaux correcteurs et de constitution de boucles secondaires se trouvent modifiés et d'ailleurs améliorés, par la considération du principe général proposé ici.

On peut noter une généralisation inspirée d'un travail antérieur — (2) page 105. Si on utilise un amplificateur de gain égal avec une excellente approximation à -1 (d'usage courant dans le calcul analogique, obtenu en prenant $R_1 = R_2$ dans la première ligne du tableau du début) le schéma de la figure 5

conduit à :

$$f(p) = \frac{y_{21}}{y_{24} - y_{23}}$$

On peut encore généraliser par l'introduction d'une



borne 5 à laquelle on applique $-E_1$, dès lors :

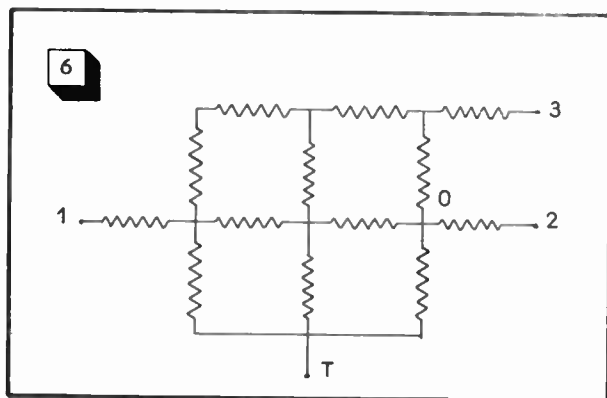
$$f(p) = \frac{y_{21} - y_{25}}{y_{24} - y_{23}}$$

Cette relation montre donc qu'avec un réseau R , C , il est possible de réaliser une fonction de transfert rationnelle en p à coefficients réels ou positifs et dont les singularités sont n'importe où dans le plan p .

Bien évidemment, ceci ne semble être utile que dans le domaine du calcul analogique. Nous en donnerons un exemple plus loin.

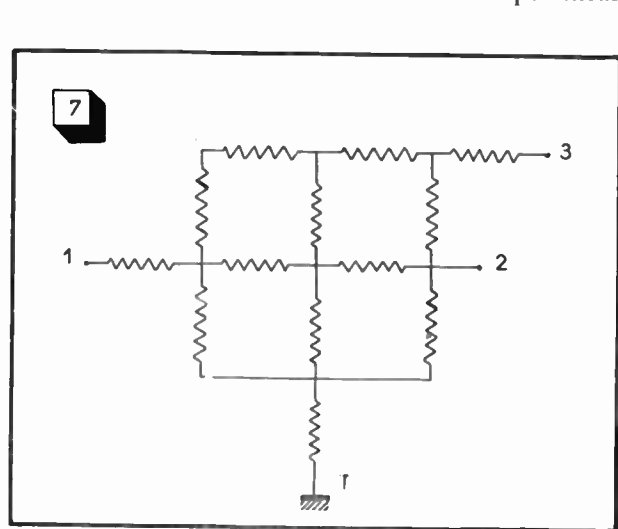
III. — UNE APPLICATION GÉNÉRALE.

La méthode de Foster a donné un rôle important dans la théorie des réseaux aux réseaux en échelle. La généralisation la plus naturelle consiste à examiner les propriétés d'une double échelle telle que représentée figure 6, utilisée dans un procédé général



de synthèse de quadripôles par le Professeur Guillemin (1).

Les conditions de fonctionnement $u = 0, i = 0$ montrent que l'admittance placée entre E et 0 ne joue aucun rôle. Il en est de même de celle qui joint le point 0 à la borne T . Elle peut donc être également supprimée, sauf si la borne T est réunie à la masse par une admittance non infinie. Nous étudierons, à titre d'exemple, le schéma de la figure 7.



tif du type de réseau, les résultats obtenus méritant d'être appliqués, nous remettons à plus tard une étude plus générale et plus satisfaisante.

(1) Advances in Electronics, t. III, pp. 261-303 « A Summary of Modern Methods of Methods of Network synthesis » par E. A. Guillemin.

On exprimera les courants pénétrant par les bornes 1, 2 et 3 en fonction de E_1 , E_2 et E_3 , en fonction de la matrice admittance.

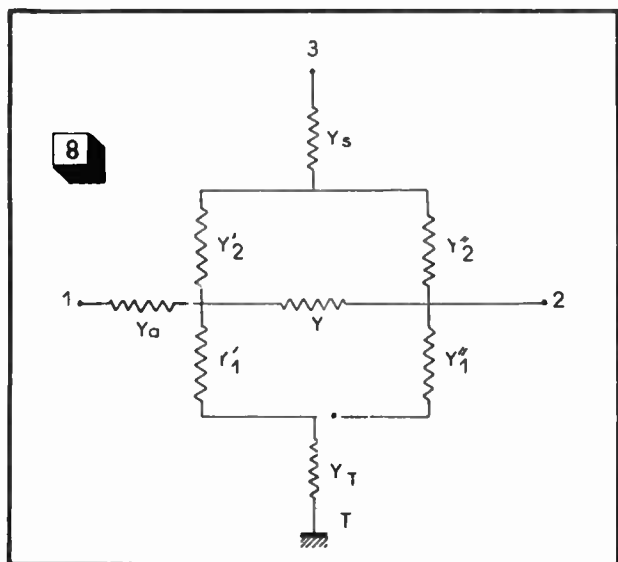
Les calculs à effectuer concernent y_{21} et y_{23} comme il a été vu plus haut. Or, $i_3 = y_{31} x + y_{32} u + y_{33} y$, donc pour les conditions de fonctionnement :

$$E_1 = E_3 = 0, E_2 = 1 : i_3 = y_{32}.$$

Pareille remarque s'applique à i_1 . Tel est le procédé de calcul le plus rapide. D'ailleurs, si on calcule le rapport $\frac{i_1}{i_3}$ mais pour $y = x = 0$ et E_2 quelconque, on a :

$$(1) \quad f(p) = \frac{i_1}{i_3}$$

Nous avons appliqué (1) ces considérations au cas d'un réseau n'ayant que 9 éléments qui jouit



d'une certaine symétrie favorable au calcul (fig. 8). Tous calculs faits, on a :

$$(2) \quad f(p) = \frac{Y_a}{Y_s} \frac{h(p)}{k(p)}$$

où l'on a posé successivement :

$$(3) \quad \begin{cases} h(p) \equiv P_T P_S Y + P_S Y'_1 Y''_1 + P_T Y'_2 Y''_2 \\ k(p) \equiv Y'_1 [Y''_2 (Y_T + Y''_1) + Y'_2 Y''_1] \\ \quad + P_T [Y''_2 (Y + Y_a + Y'_2) + Y Y'_2] \end{cases}$$

avec :

$$P_T = Y_T + Y'_1 + Y''_1, P_S = Y_S + Y'_2 + Y''_2$$

Les admittances ont les valeurs indiquées sur le schéma 6.

Faisons ici une remarque.

En exprimant $f(p)$ comme une fonction entière

des admittances, nous avons en vue des applications dans lesquelles chaque admittance serait une fonction linéaire de p (chaque admittance étant soit une résistance, soit une capacité, soit les deux en parallèle). Ce point de vue ne semble pas limitatif pour l'instant, tout progrès dans ce domaine devant être recherché par une analyse mathématique générale et non au prix de calculs élémentaires, mais laborieux lorsque le réseau est tant soit peu compliqué.

Observations sur les relations (2) et (3)

Des raisons technologiques imposent pratiquement que Y_a et Y soient des résistances. Y_a intervient comme facteur dans $f(p)$ et en produit avec $P_T Y$; dans les expressions de h et k ; P_T étant du premier degré en p , le produit $P_T Y$ peut donc être du second degré. Mais ce n'est pas le seul produit qui puisse être du second degré, il suffit d'examiner les relations (3) pour s'en rendre compte. Nous supposons donc dans la suite que $Y = a + bp$ avec a toujours différent de zéro.

Les termes de plus haut degré de h et k sont donc en p^3 . La fonction rationnelle $\frac{h}{k}$ est donc du 3^e degré.

Comme on peut prendre pour Y_s une capacité ou une résistance, finalement, on peut réaliser (entre autres) :

$$\frac{\sigma}{p} \frac{p^3 + ap^2 + \dots}{p^3 + \lambda_1 p^2 + \dots} \quad \text{ou} \quad \omega \frac{p^3 + ap^2 + \dots}{p^3 + \lambda_1 p^2 + \dots}$$

Considérons divers cas simples.

Cas $Y_s = \infty$. On a alors $P_S = Y_S = \infty$ donc :

$$f = Y_a \frac{P_T Y + Y'_1 Y''_1}{k(p)}$$

On observe que les admittances Y'_2 et Y''_2 peuvent être modifiées en module par un potentiomètre et en signe par un amplificateur à contre-réaction de gain moins un (voir fig. 9) :

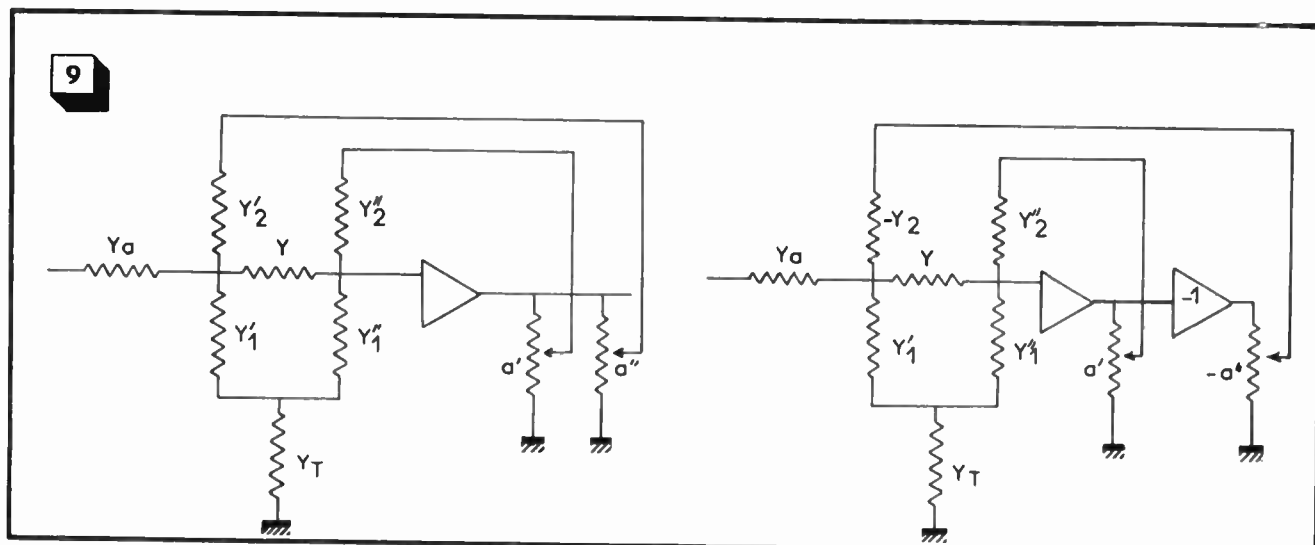
En examinant $k(p)$, on voit que cette fonction ne dépend pas de Y_s , la discussion aura ici une portée générale (en ce qui concerne $k(p)$ bien entendu). Si on désire que $k(p)$ soit un polynôme complet de degré 2 ou 3, diverses observations en résultent quant au choix des éléments du réseau.

Une discussion générale consisterait à remplacer les impédances par des formes linéaires en p . Nous n'examinerons ici que quelques cas typiques. On remarque tout d'abord que le numérateur :

$$P_T Y + Y'_1 Y''_1$$

ne dépend que du choix des admittances Y'_1 , Y''_1 , Y_T et Y . Ceci laisse penser que l'usage du circuit considéré ne conduira pas à des calculs difficiles. Le polynôme $k(p)$ du troisième degré dépend de Y_T et Y''_1 , mais principalement des autres admittances.

(1) N. T. 264 — 14 avril 1953 — de la Société d'Electronique & d'Automatisme (S. E. A.).



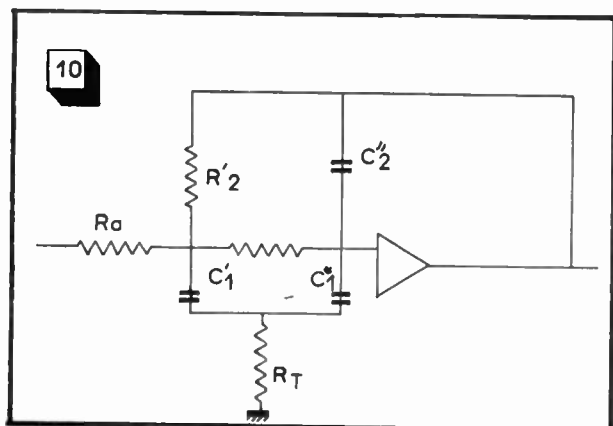
tances. Le choix des éléments permet donc de réaliser la fonction de transfert :

$$f = -Y_a \frac{p^2 + 2\omega_3 p + \omega_0^2}{p^3 + \lambda p^2 + \mu p + \tau}$$

On voit qu'il faut que :

$$\begin{cases} P_T = Y_T + Y'_1 + Y''_1 \text{ soit de la forme } \frac{1}{R_T} + (C'_1 + C''_1) p \\ Y = \frac{1}{R} \end{cases}$$

ainsi :



$$\begin{aligned} P_T Y + Y'_1 + Y''_1 &= \frac{1}{R R_T} + \frac{(C'_1 + C''_1)}{R} p \\ + C'_1 C''_1 p^2 &= \frac{1}{R R_T} [1 + R_T (C'_1 + C''_1) p \\ &\quad + C'_1 C''_1 R R_T p^2] \end{aligned}$$

Pour que le terme en p^3 existe, il faut alors que

$Y''_2 = C''_2 p$ mais si $Y'_2 = \frac{1}{R'_2}$ alors le terme constant est $\frac{1}{R_T} \frac{1}{R'_2} \frac{1}{R}$ (tiré du produit $P_T Y Y'_2$). Le schéma est donné figure 10.

Si $Y_T = \infty$, $Y_S = \infty$, alors on a :

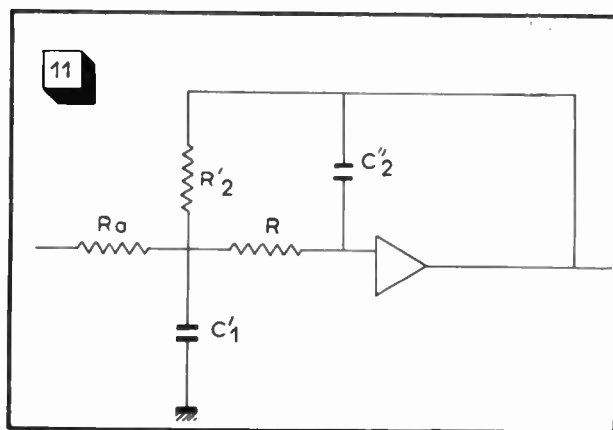
$$f = -Y_a \frac{P_T Y}{Y_T Y'_1 Y''_2 + P_T [Y''_2 (Y + Y_a + Y'_2) + Y Y'_2]}$$

et pour $P_T \rightarrow \infty$:

$$f = -Y Y_a \frac{1}{Y''_2 [Y + Y_a + Y'_2 + Y'_1] + Y Y'_2}$$

d'où avec les admittances de la figure 11 :

$$f = -\frac{1}{R R_a} \frac{1}{C''_2 p \left[\frac{1}{R} + \frac{1}{R_a} + \frac{1}{R'_2} + C'_1 p \right] + \frac{1}{R R'_2}}$$

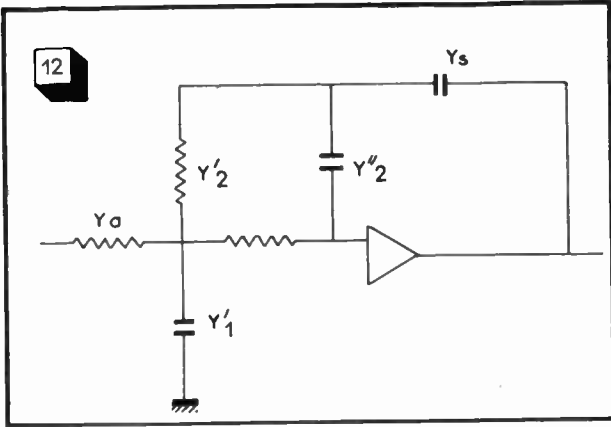


$$f = - \frac{R'_2}{R_a} \frac{1}{1 + RR'_2 \left(\frac{1}{R} + \frac{1}{R_a} + \frac{1}{R'_2} \right) C''_2 p + C''_2 C'_1 R R'_2 p^2}$$

Examinons le rôle de Y_s dans le cas simple $Y_T \rightarrow \infty$ (fig. 12).

On a successivement :

$$f = - \frac{Y_a}{Y_s} \frac{P_T [P_s Y + Y'_2 Y''_2]}{Y_T [Y''_2 (Y + Y_a + Y'_1 + Y'_2) + Y Y'_2]}$$



et à la limite $Y_T = P_T \rightarrow \infty$

$$f = - \frac{Y_a}{Y_s} \frac{P_s Y + Y'_2 Y''_2}{Y''_2 (Y + Y_a + Y'_1 + Y'_2) + Y Y'_2}$$

Compte tenu du calcul qui vient d'être effectué ci-dessus, on voit que si $Y_s = C_s p$, on a :

$$P_s = (C_s + C''_2) p + \frac{1}{R'_2}, \quad Y'_2 Y''_2 = \frac{C''_2}{R'_2} p$$

et :

$$f = - \frac{1}{R_a C_s p} \frac{\sigma + \sigma_2 p}{\text{polynôme second degré}}$$

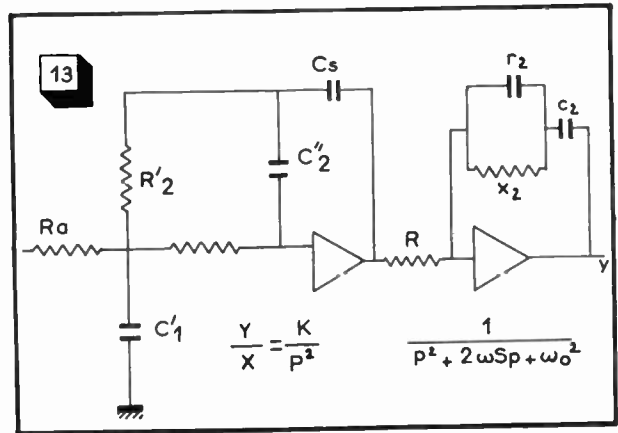
On notera que cette forme ne manque pas d'intérêt pratique dans l'étude des servomécanismes, on voit, en particulier que la réalisation de l'opération :

$$K = \frac{M}{p^2} \frac{1}{\text{polynôme second degré}}$$

peut s'exécuter ainsi avec deux amplificateurs ou deux étages d'amplifications comme indiqué sur la figure 13.

Certes, on peut prévoir obtenir ce résultat avec un seul amplificateur, mais il faudrait posséder une règle commode de construction du réseau en double échelle considérée au début.

Les applications que suggèrent les considérations



précédentes sont nombreuses, croyons-nous. Nous pensons que la est notre excuse de publier ce papier sans résoudre le problème général de synthèse — qui est le plus intéressant, mais aussi le plus difficile.

IV. — PROBLÈMES DE FILTRAGE.

Nous mentionnerons ici que l'un des schémas examinés plus haut a été expérimenté au Laboratoire de Calcul Electronique du Service Technique Aéronautique (1), en vue de déterminer, par mesure directe, la densité spectrale d'une grandeur aléatoire. Une telle mesure exige un filtre de bande étroite dont la fréquence centrale et la largeur de bande doivent être réglables facilement dans une grande gamme et, pour les applications usuelles des méthodes statistiques aux problèmes de l'industrie ou de la mécanique du vol, dans une gamme de fréquences très basses.

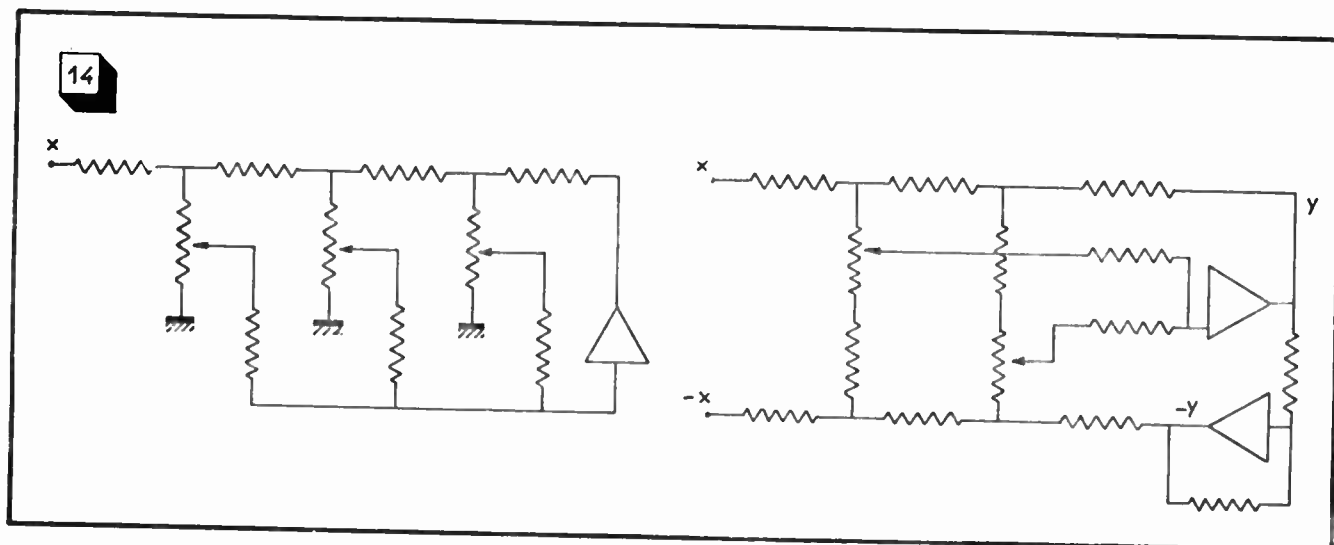
V. — UN SCHÉMA INSPIRÉ PAR LE CALCUL ANALOGIQUE.

L'application d'un réseau en échelle dans le cadre des schémas habituels des calculateurs analogiques à usage général (2) a été faite par M. BRODIN (3). Nous nous contenterons d'indiquer ici le schéma utilisé. La théorie se ferait aisément à partir des propriétés des réseaux en échelle, soit avec les déterminants, soit mieux à l'aide des polynômes électrosphériques. La structure du réseau est une échelle (dans le cas le plus simple, elle est formée de résistances en série et de capacités en parallèle, ou, dans le cas expérimenté par M. BRODIN, les capacités sont en série et les résistances en dérivation). Dans chaque branche parallèle, on prélève une fraction de la tension au nœud correspondant à l'aide d'un potentiomètre. Ces tensions sont appliquées à un réseau de sommation. Evidemment, la sommation pondérée des tensions aux nœuds de l'échelle peut être réalisée par un réseau de résistances inégales de valeurs convenables. L'une des bornes d'accès

(1) à la Société d'Electronique et d'Automatisme (S. E. A.).

(2) voir F. H. RAYMOND, loc. cit chapitre I.

(3) F. BRODIN « Simulation d'équations différentielles du second ordre » note SABA 54 du 24-4-54 (non publiée) du Laboratoire de Recherches Balistiques et Aérodynamiques (L.R.B.A.) de la Direction des Etudes et Fabrications d'Armement (Vernon-Eure).



de l'échelle reçoit une tension arbitraire, l'autre est connectée à la sortie de l'amplificateur.

L'introduction de termes négatifs est obtenu dans le cas BRODIN, en réalisant une échelle double dont la borne d'accès de l'une d'elles est alimentée par $-y$ (fig. 14).

VI. — NOTE FINALE.

Nous nous proposons de développer plus avant, cette note, empruntée à un document intérieur, de la

SEA (1). L'arrivée récente à notre bibliothèque des comptes-rendus de la « National Telemetering Conference » tenue à Chicago en mai 1953, nous autorise, pensons-nous, à abréger ce travail. Dans ces comptes-rendus, en effet, sous le titre « filters for telemetry » (pp. 159-176), L. L. RAUCH (*University of Michigan*) et C.E. HOWE (*Oberlin College*) proposent des schémas analogues aux nôtres, ils indiquent, en outre, que l'un d'eux avait proposé deux schémas de cette nature en 1949.

(1) Référence N.T. 324 du 25 mars 1954 faisant suite à la note N. T. 264 du 14 avril 1953.

BASES THÉORIQUES DE LA MESURE DE LA RÉSISTIVITÉ ET DE LA CONSTANTE DE HALL PAR LA MÉTHODE DES POINTES

L'AR

J. LAPLUME

Docteur ès-Sciences

Ingénieur à la Compagnie Française Thomson-Houston

La méthode dite « des pointes » consiste à injecter le courant dans une substance conductrice ou semi-conductrice et à recueillir la tension ohmique ou la tension de Hall au moyen de pointes rapprochées. Cette méthode diffère des méthodes classiques par le fait que la répartition des lignes de courant n'est pas uniforme. Ses avantages, ainsi que les détails de mise en œuvre, sont exposés dans un article de M. B. PISTOULET (1).

La résistivité ou la constante de Hall se calculent à partir du courant injecté et de la tension recueillie. Les formules qui permettent ce calcul dans le cas de la répartition uniforme sont bien connues. Lorsqu'on emploie la méthode des pointes, ces formules doivent être modifiées pour tenir compte de la différence de répartition des lignes de courant. L'établissement de ces nouvelles formules constitue l'objet du présent article.

La répartition des lignes de courant est essentiellement fonction de la disposition des pointes et de la forme de l'échantillon. Lorsque les dimensions de celui-ci sont grandes vis-à-vis de l'écartement des pointes, tout se passe pratiquement comme si le milieu conducteur ou semi-conducteur s'étendait indéfiniment dans les trois dimensions. Si tel n'est pas le cas, il faut apporter des corrections aux formules valables dans ce cas limite. Nous traiterons le cas d'un échantillon indéfini suivant une dimension (repérée par la coordonnée x), et dont la section perpendiculaire à l'axe des x est rectangulaire. Ces conditions sont moins restrictives qu'elles ne paraissent. Certains échantillons s'écartent de cette forme idéale, mais leur étendue suivant tel axe de coordonnées est grande devant l'espacement des pointes. La correction par rapport au cas limite est alors faible et on pourra introduire dans les formules une dimension transversale moyenne sans commettre d'erreur importante.

(1) B. PISTOULET. Mesure des propriétés électriques des semi-conducteurs.
Onde Electrique de janvier, p. 73.

Les corrections à apporter dans le cas où les dimensions de la plaquette suivant la coordonnée x ne sont pas très grandes sont établies en Appendice IV.

Les hypothèses admises tout au long des calculs sont les suivantes :

1^o La section de contact pointe-conducteur est très faible devant l'espacement des pointes. Le contact peut alors être considéré comme ponctuel.

2^o Le conducteur est supposé homogène dans toute la région située au voisinage des pointes, et à l'intérieur de laquelle passe la quasi-totalité du courant injecté.

On sait que la répartition du champ électrique et du vecteur courant à l'intérieur d'un milieu conducteur homogène et suivant la loi d'Ohm est laplacienne. La détermination de la répartition des lignes de courant se ramène donc à la répartition du champ électrostatique créé par deux points-sources. Le potentiel électrostatique créé par un point-source est de la forme

$$(1) \quad V = K/r,$$

où r est la distance par rapport au point-source. K est une constante liée à l'intensité injectée ou recueillie. Au voisinage d'un point, le potentiel créé par les autres points est négligeable. Le champ électrique est alors radial et d'intensité

$$E = - \frac{\partial V}{\partial r} = K/r^2.$$

Si ρ désigne la résistivité du milieu, le vecteur courant i pour valeur

$$(2) \quad |\vec{i}| = \frac{|\vec{E}|}{\rho} = K/\rho r^2.$$

On en déduit le courant total I injecté par la pointe au contact du conducteur en intégrant sur une

demi-sphère de rayon r centrée sur le point de contact (fig. 1).

$$(3) \quad I = |\vec{i}| \times 2\pi r^2 = 2\pi K/\rho.$$

On en déduit

$$(4) \quad K = \rho I/2\pi$$

et par comparaison avec (1)

$$(5) \quad V = \frac{\rho I}{2\pi} \frac{1}{r}.$$

Si le courant I est recueilli, la même formule est valable avec le signe négatif.

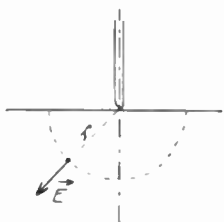


FIG. 1

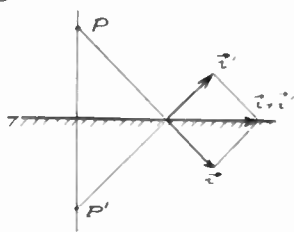


FIG. 2

Les conditions aux limites imposent que les lignes de courant soient tangentes à la surface du conducteur. Dans un solide limité par des surfaces planes, on satisfait automatiquement à ces conditions en appliquant le principe des images spéculaires. On associe à chaque point-source son symétrique par rapport à chacune des surfaces planes limites; ce point-image est supposé injecter ou recueillir le même courant (fig. 2). On constate aisément que la résultante des deux courants radiaux i et i' est située dans le plan de symétrie.

Le potentiel en un point quelconque s'obtient donc en sommant les contributions des points-sources réels et de leurs images, chacune des contributions étant calculée en appliquant la formule (5).

I. — MESURE DES RÉSISTIVITÉS.

A. — Méthodes des quatre points alignées.

Les quatre pointes $ABCD$ sont situées dans le plan de symétrie $y = 0$ de l'échantillon (fig. 3). Le courant I est injecté par A et recueilli par B . La tension ohmique est mesurée entre les points C et D . Les notations sont indiquées sur la figure.

Les coordonnées des points-sources A et B sont $(-l, 0, 0)$ et $(+l, 0, 0)$. Les coordonnées des images de A et B par rapport aux plans limites $z = 0$, $z = -z_0$, $y = y_0$, $y = -y_0$, sont données par la formule générale

$$x = \pm l, \quad y = 2my_0, \quad z = 2nz_0,$$

où m et n sont deux entiers prenant toutes les valeurs de $-\infty$ à $+\infty$. La valeur 0 correspond aux points-sources.

On en déduit l'expression du potentiel en tout point x, y, z situé à l'intérieur de l'échantillon

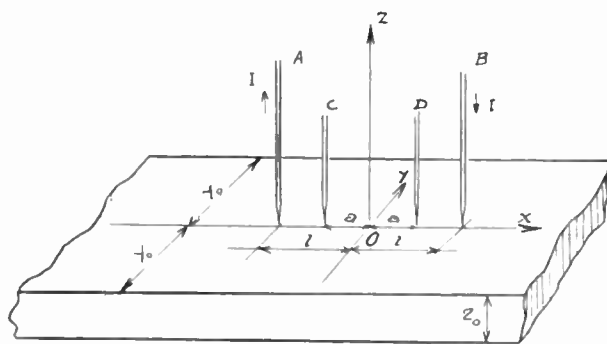


FIG. 3.

$$(6) \quad V(x, y, z) = \frac{\rho I}{2\pi} \sum_{m=-\infty}^{+\infty} \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{(l-x)^2 + (2my_0 - y)^2 + (2nz_0 - z)^2}} - \frac{1}{\sqrt{(l+x)^2 + (2my_0 - y)^2 + (2nz_0 - z)^2}} \right].$$

La tension mesurée entre C et D est

$$v = V(a, 0, 0) - V(-a, 0, 0) = 2V(a, 0, 0),$$

cette dernière égalité résultant de l'antisymétrie de la formule (6) par rapport à x . Finalement

$$(7) \quad v = \frac{\rho I}{\pi} \sum_{m=-\infty}^{+\infty} \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{(l-a)^2 + (2my_0)^2 + (2nz_0)^2}} - \frac{1}{\sqrt{(l+a)^2 + (2my_0)^2 + (2nz_0)^2}} \right].$$

Cette somme ne se ramène à aucune fonction tabulée. Il est donc nécessaire d'en effectuer le calcul numérique.

Remarquons tout d'abord que y_0 et z_0 jouent un rôle symétrique dans la formule (7). Nous pouvons donc supposer $z_0 \leq y_0$ et calculer une valeur approchée de la somme valable dans ce cas. En permutant z_0 et y_0 , nous obtiendrons une valeur approchée valable pour $y_0 \leq z_0$, et nous aurons ainsi couvert la totalité des cas.

Nous décomposons comme suit la somme (7) :

$$(8) \quad v = \frac{\rho I}{\pi} \left\{ \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{(l-a)^2 + (2nz_0)^2}} - \frac{1}{\sqrt{(l+a)^2 + (2nz_0)^2}} \right] + 2 \sum_{m=1}^{+\infty} \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{(l-a)^2 + (2my_0)^2 + (2nz_0)^2}} - \frac{1}{\sqrt{(l+a)^2 + (2my_0)^2 + (2nz_0)^2}} \right] \right\},$$

c'est-à-dire que nous isolons la contribution des termes $m = 0$. La parité de la fonction v par rapport à y , nous permet de limiter la sommation sur m aux valeurs positives de l'indice moyennant l'introduction d'un facteur 2.

Portons tout d'abord notre attention sur la somme double. Elle peut s'écrire, en extrayant $2z_0$ des radicaux :

$$(9) \quad \frac{1}{2z_0} \sum_{m=1}^{+\infty} \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{\frac{(l-a)^2 + (2my_0)^2}{(2z_0)^2} + n^2}} - \frac{1}{\sqrt{\frac{(l+a)^2 + (2my_0)^2}{(2z_0)^2} + n^2}} \right].$$

La sommation sur n peut être remplacée par une intégration avec une bonne approximation. Nous démontrons en effet, en appendice (formule A5), que :

$$(10) \quad \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{A^2 + n^2}} - \frac{1}{\sqrt{B^2 + n^2}} \right] \approx \text{Log} \frac{B^2}{A^2}$$

pourvu que A et B ne soient pas inférieurs à l'unité. Tel est bien le cas ici, puisque $z_0 \leq y_0$ et $m \geq 1$, et que, par conséquent :

$$A^2 = \frac{(l-a)^2 + (2my_0)^2}{(2z_0)^2} \geq 1$$

$$B^2 = \frac{(l+a)^2 + (2my_0)^2}{(2z_0)^2} \geq 1.$$

La sommation sur n fournit donc le résultat approché :

$$(11) \quad \text{Log} \frac{(l+a)^2 + (2my_0)^2}{(l-a)^2 + (2my_0)^2}.$$

Le résultat de la sommation sur m se déduit de la formule établie en appendice (formule A 17)

$$\sum_{m=1}^{+\infty} \text{Log} \left(1 + \frac{l^2}{m^2} \right) = \text{Log} \frac{\text{Sh} \pi l}{\pi l}.$$

On en déduit :

$$\sum_{m=1}^{+\infty} \text{Log} \frac{(l+a)^2 + (2my_0)^2}{(l-a)^2 + (2my_0)^2} = \sum_{m=1}^{+\infty} \left\{ \text{Log} \left[1 + \left(\frac{l+a}{2my_0} \right)^2 \right] \right.$$

$$\left. - \text{Log} \left[1 + \left(\frac{l-a}{2my_0} \right)^2 \right] \right\} = \text{Log} \frac{\text{Sh} \pi \frac{l+a}{2y_0}}{\pi \frac{l+a}{2y_0}}$$

$$- \text{Log} \frac{\text{Sh} \pi \frac{l-a}{2y_0}}{\pi \frac{l-a}{2y_0}} = \text{Log} \frac{l-a}{l+a} \frac{\text{Sh} \pi \frac{l+a}{2y_0}}{\text{Sh} \pi \frac{l-a}{2y_0}}$$

Finalement, la somme double a pour valeur approchée :

$$(12) \quad \frac{1}{2z_0} \text{Log} \frac{l-a}{l+a} \frac{\text{Sh} \pi \frac{l+a}{2y_0}}{\text{Sh} \pi \frac{l-a}{2y_0}}$$

Passons maintenant à la somme simple de l'équation (8).

Il faut ici distinguer deux cas.

1° *Echantillon mince*, $2z_0 \leq (l-a)$.

On peut utiliser l'approximation de l'intégrale (10) qui nous donne :

$$(13) \quad \frac{1}{2z_0} \text{Log} \left(\frac{l+a}{l-a} \right)^2 = \frac{1}{z_0} \text{Log} \frac{l+a}{l-a}.$$

En rassemblant les expressions (12) et (13), nous obtenons :

$$(14) \quad v = \frac{\rho l}{\pi z_0} \text{Log} \frac{\text{Sh} \pi \frac{l+a}{2y_0}}{\text{Sh} \pi \frac{l-a}{2y_0}}.$$

Cette formule se simplifie quelque peu dans le cas usuel où les quatre pointes sont équidistantes. Alors :

$$(15) \quad l = 3a, \quad l+a = 4a, \quad l-a = 2a.$$

$$\text{Sh} \pi \frac{l+a}{2y_0} = \text{Sh} \pi \frac{2a}{y_0} = 2 \text{Sh} \frac{\pi a}{y_0} \text{Ch} \frac{\pi a}{y_0},$$

$$\text{et} \quad \text{Sh} \pi \frac{l-a}{2y_0} = \text{Sh} \frac{\pi a}{y_0}.$$

Dans ce cas :

$$(16) \quad v = \frac{\rho l}{\pi z_0} \text{Log} 2 \text{Ch} \frac{\pi a}{y_0}.$$

2° *Echantillon épais*, $2z_0 > (l+a)$.

L'approximation de l'intégrale n'est plus valable. Il faut effectuer le calcul direct terme à terme.

Nous nous sommes borné au cas des pointes équidistantes, pour lequel l et a sont liés par les relations (15).

La somme prend alors la forme particulière :

$$(17) \quad \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{(2a)^2 + (2nz_0)^2}} - \frac{1}{\sqrt{(4a)^2 + (2nz_0)^2}} \right]$$

$$= \frac{1}{4a} + \frac{1}{z_0} \sum_{n=1}^{\infty} \left[\frac{1}{\sqrt{\left(\frac{a}{z_0}\right)^2 + n^2}} - \frac{1}{\sqrt{\left(\frac{2a}{z_0}\right)^2 + n^2}} \right]$$
$$= \frac{1}{4a} K(a/z_0)$$

avec : (18)

$$K(s) = 1 + 4s \sum_{n=1}^{\infty} \left[\frac{1}{\sqrt{s^2 + n^2}} - \frac{1}{\sqrt{(2s)^2 + n^2}} \right].$$

La fonction $K(s)$ est représentée sur la figure 4.

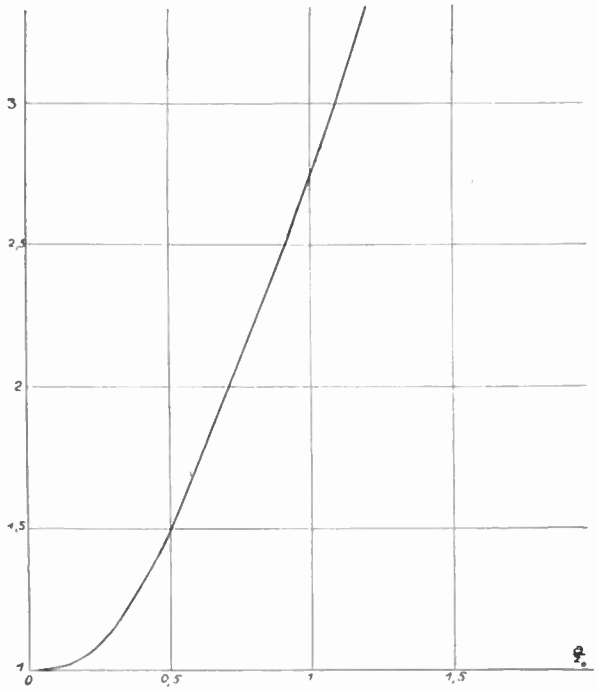


FIG. 4. — Facteur $K\left(\frac{a}{z_1}\right) \approx 2,773 \frac{a}{z_1}$ pour $\frac{a}{z_1} > 1$

La méthode de calcul est indiquée en appendice.
Finalement, l'expression de v est, dans ce cas,

$$(19) \quad v = \frac{\rho I}{\pi} \left[\frac{1}{4a} K(a/z_0) + \frac{1}{z_0} \text{Log Ch } \frac{\pi a}{y_0} \right].$$

Le terme correctif $\text{Log Ch } \frac{\pi a}{y_0}$ est fourni par la courbe de la figure 5.

B. — Méthode des quatre pointes en carré.

Les quatre pointes sont disposées aux quatre sommets d'un carré. Le courant I est injecté en B et recueilli en A . La tension ohmique est recueillie entre C et D (fig. 6).

La répartition $V(x, y, z)$ du potentiel est toujours fournie par la formule (6) où l est remplacé par a , conformément au changement de notation. La tension recueillie est évidemment :

$$(20) \quad v = V_D - V_C = V(a, 2a, o) - V(-a, 2a, o) = 2V(a, 2a, o).$$

Cette dernière relation résulte immédiatement de la symétrie du problème, qui entraîne l'antisymétrie de la fonction $V(x, y, z)$ par rapport à a .
Par conséquent :

$$(21) \quad v = \frac{\rho I}{\pi} \sum_{m=-\infty}^{+\infty} \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{(2my_0 - 2a)^2 + (2nz_0)^2}} - \frac{1}{\sqrt{(2a)^2 + (2my_0 - 2a)^2 + (2nz_0)^2}} \right].$$

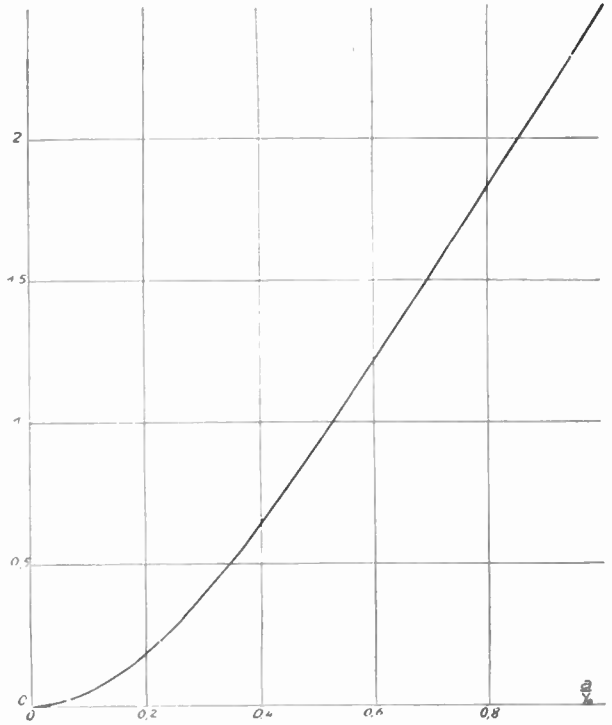


FIG. 5. — Facteur $\text{Log Ch } \frac{\pi a}{y_0} \approx 3,14 \frac{a}{y_0} - 0,693$ pour $\frac{a}{y_0} > 1$

Isolons le terme $m = 0$ et sortons le facteur 2 du radical :

$$(22) \quad v = \frac{\rho I}{2\pi} \left\{ \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{a^2 + (nz_0)^2}} - \frac{1}{\sqrt{2a^2 + (nz_0)^2}} \right] \right\}$$

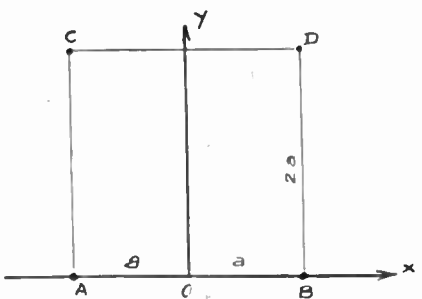


FIG. 6.

$$\begin{aligned}
& + \sum_{m=-1}^{+\infty} \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\sqrt{(my_0 - a)^2 + (nz_0)^2}} \right. \\
& \quad - \frac{1}{\sqrt{a^2 + (my_0 - a)^2 + (nz_0)^2}} \\
& \quad + \frac{1}{\sqrt{(my_0 + a)^2 + (nz_0)^2}} \\
& \quad \left. - \frac{1}{\sqrt{a^2 + (my_0 + a)^2 + (nz_0)^2}} \right] \Bigg\}.
\end{aligned}$$

Ces deux derniers termes correspondent à la sommation sur les valeurs négatives de l'indice m .

Occupons-nous d'abord de la somme double. Elle est du type (9). La sommation sur n peut être remplacée par une intégrale dans le cas où z_0 est $\leq y_0$, et fournit le résultat suivant, par application de la formule (10) :

$$(23) \quad \frac{1}{z_0} \text{Log} \frac{a^2 + (my_0 - a)^2}{(my_0 - a)^2} \frac{a^2 + (my_0 + a)^2}{(my_0 + a)^2}.$$

Le calcul de la somme sur m est effectué en appendice et conduit au résultat (formule A 19) :

$$(24) \quad \frac{1}{z_0} \text{Log} \frac{Ch \frac{2\pi a}{y_0} - \cos \frac{2\pi a}{y_0}}{4 \sin^2 \frac{\pi a}{y_0}}.$$

Pour le calcul de la somme simple, il faut distinguer deux cas :

1° *Echantillon mince*, $z_0 \leq a$.

L'application de la formule (10) conduit au résultat :

$$\frac{1}{z_0} \text{Log} 2.$$

D'où finalement pour l'expression de v :

$$(25) \quad v = \frac{\rho I}{2\pi z_0} \text{Log} \frac{Ch \frac{2\pi a}{y_0} - \cos \frac{2\pi a}{y_0}}{2 \sin^2 \frac{\pi a}{y_0}}.$$

2° *Echantillon épais*, $z_0 > a$.

Il faut effectuer le calcul direct de la somme. Celle-ci se met sous la forme équivalente :

$$\begin{aligned}
(26) \quad & \frac{1}{a} \left(1 - \frac{1}{\sqrt{2}} \right) + \frac{2}{z_0} \sum_{n=1}^{\infty} \left[\frac{1}{\sqrt{(a/z_0)^2 + n^2}} \right. \\
& \quad \left. - \frac{1}{\sqrt{(\sqrt{2} a/z_0)^2 + n^2}} \right] \\
& = \frac{1}{a} \frac{2 - \sqrt{2}}{2} \left[1 + \frac{4}{2 - \sqrt{2}} \frac{a}{z_0} \sum_{n=1}^{\infty} \left[\frac{1}{\sqrt{(a/z_0)^2 + n^2}} \right. \right.
\end{aligned}$$

$$\left. - \frac{1}{\sqrt{(\sqrt{2} a/z_0)^2 + n^2}} \right] \Bigg] = \frac{1}{a} \frac{2 - \sqrt{2}}{2} G(a/z_0)$$

avec :

$$(27) \quad G(s) = 1 + \frac{4s}{2 - \sqrt{2}} \sum_{n=1}^{+\infty} \left[\frac{1}{\sqrt{s^2 + n^2}} - \frac{1}{\sqrt{(\sqrt{2}s)^2 + n^2}} \right]$$

Le calcul s'effectue de la même manière que celui de K et permet de construire la courbe de la fig. 7.

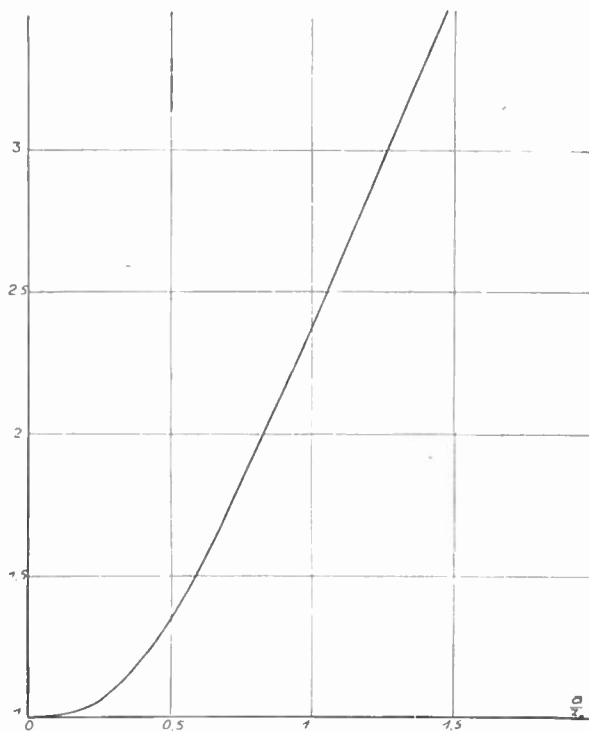


FIG. 7. — Facteur $G\left(\frac{a}{z_0}\right) \approx 2,37 \frac{a}{z_0}$ pour $\frac{a}{z_0} > 1$.

Finalement :

$$\begin{aligned}
(28) \quad v = & \frac{\rho I}{2\pi} \left[\frac{2 - \sqrt{2}}{2a} G\left(\frac{a}{z_0}\right) \right. \\
& \left. + \frac{1}{z_0} \text{Log} \frac{Ch \frac{2\pi a}{y_0} - \cos \frac{2\pi a}{y_0}}{4 \sin^2 \frac{\pi a}{y_0}} \right].
\end{aligned}$$

Le terme correctif $\text{Log} \frac{Ch \frac{2\pi a}{y_0} - \cos \frac{2\pi a}{y_0}}{4 \sin^2 \frac{\pi a}{y_0}}$ est

fourni par la courbe de la fig. 8.

Nous n'étudierons pas le cas où z_0 est $\geq y_0$. L'intérêt pratique en est faible, car on peut toujours se placer dans le cas $z_0 \leq y_0$, en orientant convenablement l'échantillon.

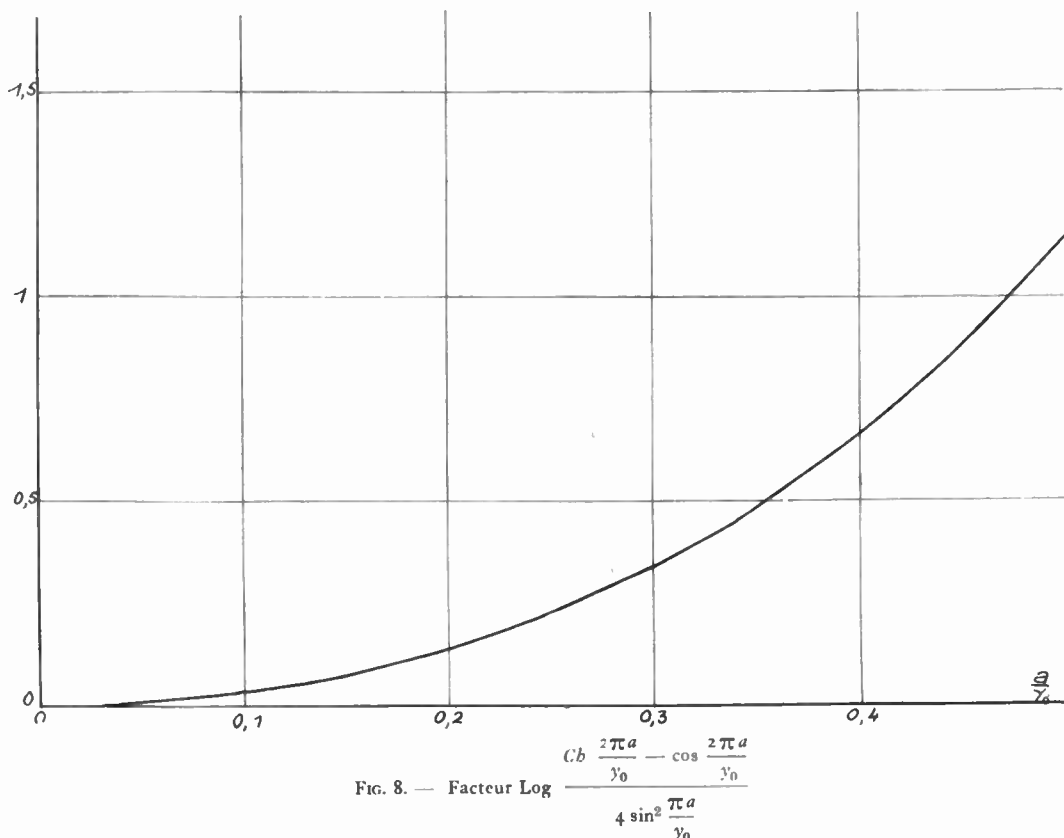


FIG. 8. — Facteur Log

II. — MESURE DE L'EFFET HALL.

Le courant est injecté et recueilli par les pointes *B* et *A* respectivement d'abscisses $+a$ et $-a$. Le champ magnétique uniforme H est appliqué suivant la direction oy . La tension de Hall est recueillie entre deux pointes situées l'une à l'origine o , l'autre au point de cote $-z_0$ sur l'axe des z .

Le potentiel électrique est toujours donné par la formule (6) dans laquelle on fait $l = a$. On en déduit la composante E_x du champ électrique par dérivation :

$$(29) \quad E_x = -\frac{\partial V}{\partial x} = \frac{\rho I}{2\pi} \sum_{m=-\infty}^{+\infty} \sum_{n=-\infty}^{+\infty} \left\{ \frac{a-x}{[(a-x)^2 + (2my_0-y)^2 + (2nz_0-z)^2]^{3/2}} + \frac{a+x}{[(a+x)^2 + (2my_0-y)^2 + (2nz_0-z)^2]^{3/2}} \right\}$$

En particulier, sur l'axe Oz où $x = y = 0$

$$(30) \quad E_x = \frac{\rho I a}{\pi} \sum_{m=-\infty}^{+\infty} \sum_{n=-\infty}^{+\infty} \frac{1}{[a^2 + (2my_0)^2 + (2nz_0-z)^2]^{3/2}}$$

Un porteur de charge de mobilité μ acquiert sous l'action de E_x une composante de vitesse μE_x . Le champ magnétique H exerce sur le porteur une force de Laplace $e\mu E_x H$ dirigée suivant Oz . Les charges déviées par le champ magnétique s'accumulent sur les faces $z = 0$ et $z = -z_0$ et un état d'équilibre

s'établit lorsque le champ E_H créé par ces charges égale la force de Laplace. Par conséquent :

$$(31) \quad E_H = \mu E_x H.$$

La tension de Hall V_H s'obtient en calculant la

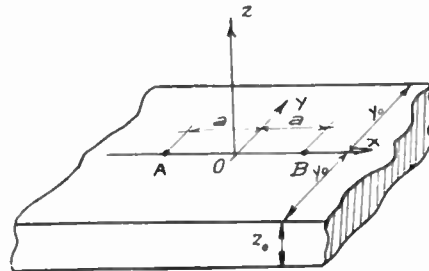


FIG. 9.

circulation de ce champ entre les deux plans $z = 0$ et $z = -z_0$.

Sur l'axe Oz :

$$(32) \quad V_H = \int_{-z_0}^0 E_H dz = \mu H \int_{-z_0}^0 E_x(0, 0, z) dz.$$

D'après (30) :

$$(33) \quad V_H = \frac{\rho I \mu H a}{\pi} \int_{-z_0}^0 \sum_{m=-\infty}^{+\infty} \sum_{n=-\infty}^{+\infty} \left\{ \frac{1}{[a^2 + (2my_0)^2 + (2nz_0+z)^2]^{3/2}} \right\} dz.$$

Nous avons, dans la somme, remplacé $2 n z_0$ par $- 2 n z_0$ sans changer le résultat, puisque la sommation porte sur toutes les valeurs entières, positives comme négatives, de n .

Effectuons d'abord l'intégration sur z .

(34)

$$\int_{-z_0}^0 \frac{dz}{[a^2 + (2my_0)^2 + (2nz_0 + z)^2]^{3/2}}$$
$$= \frac{1}{a^2 + (2my_0)^2} \int_{-z_0}^0 \frac{d\left(\frac{2nz_0 + z}{\sqrt{a^2 + (2my_0)^2}}\right)}{\left[1 + \left(\frac{2nz_0 + z}{\sqrt{a^2 + (2my_0)^2}}\right)^2\right]^{3/2}}$$
$$= \frac{1}{a^2 + (2my_0)^2} \left[\frac{2nz_0 + z}{\sqrt{a^2 + (2my_0)^2 + (2nz_0 + z)^2}} \right]_{-z_0}^0$$
$$= \frac{1}{a^2 + (2my_0)^2} \left[\frac{2nz_0}{\sqrt{a^2 + (2my_0)^2 + (2nz_0)^2}} - \frac{(2n - 1)z_0}{\sqrt{a^2 + (2my_0)^2 + ((2n - 1)z_0)^2}} \right].$$

Cette expression est de la forme

(35)

$$\frac{1}{a^2 + (2my_0)^2} (A_n - A_{n-1})$$

avec

$$A_n = \frac{2nz_0}{\sqrt{a^2 + (2my_0)^2 + (2nz_0)^2}}.$$

Si nous effectuons la sommation sur les valeurs entières de n à partir de N jusqu'à $N + p$, nous obtenons

$$A_N - A_{N-1} + A_{N+1} - A_N + A_{N+2} - A_{N+1} + \dots + A_{N+p} - A_{N+p-1} = - A_{N-1}.$$

On en déduit que

(36)

$$\sum_{n=-\infty}^{+\infty} (A_n - A_{n-1}) = - \lim_{n \rightarrow -\infty} A_n = 1.$$

Le résultat de l'intégration sur z et de la sommation sur n se réduit donc finalement à $\frac{1}{a^2 + (2my_0)^2}$.

Donc (37)

$$V_H = \frac{\rho I \mu H a}{\pi} \sum_{m=-\infty}^{+\infty} \frac{1}{a^2 + (2my_0)^2}.$$

Nous démontrons en appendice (formule A 11) que

(38)

$$\sum_{m=-\infty}^{+\infty} \frac{1}{\lambda^2 + m^2} = \frac{\pi}{\lambda} \operatorname{Coth} \pi \lambda.$$

Donc (39)

$$\sum_{m=-\infty}^{+\infty} \frac{1}{a^2 + (2my_0)^2} = \frac{1}{4y_0^2} \sum_{m=-\infty}^{+\infty} \frac{1}{\left(\frac{a}{2y_0}\right)^2 + m^2}$$
$$= \frac{\pi}{2ay_0} \operatorname{Coth} \frac{\pi a}{2y_0}$$

Finalement :

(39)

$$V_H = \frac{\rho I \mu H}{2y_0} \operatorname{Coth} \frac{\pi a}{2y_0},$$

ou encore, en introduisant la constante de Hall :

(40)

$$R = \mu \rho :$$

(41)

$$V_H = \frac{RHI}{\pi a} \frac{\pi a}{2y_0} \operatorname{Coth} \frac{\pi a}{2y_0}.$$

Le facteur correctif $\frac{\pi a}{2y_0} \operatorname{Coth} \frac{\pi a}{2y_0}$ est représenté sur la fig. 10.

Ces calculs ont été effectués avec la collaboration de M. B. PISTOULET.

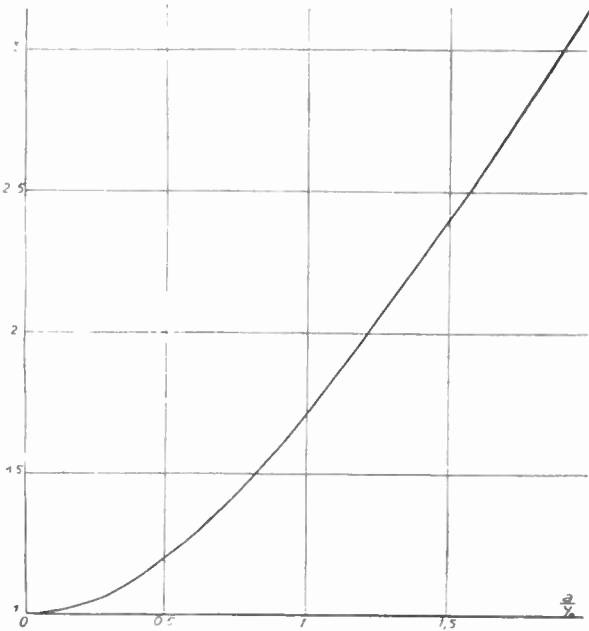


FIG. 10. — Facteur $\frac{\pi a}{2y_0} \operatorname{Coth} \frac{\pi a}{2y_0} \approx 1,57 \frac{a}{y_0}$ pour $\frac{a}{y_0} > 2$

Depuis la rédaction de cet article, nous avons pris connaissance d'une étude de L. B. VALDES publiée dans les *P.I.R.E.* de février 1954, p. 420-427. On trouvera dans cette publication des courbes permettant le calcul de la résistivité par la méthode des quatre pointes alignées. Toutefois, les cas envisagés par M. VALDES sont ou différents ou moins généraux que ceux exposés dans le cours de cet article.

APPENDICE I

Calcul des sommes de la forme

$$(A1) \quad S(A, B) = \sum_{-\infty}^{+\infty} \left[\frac{1}{\sqrt{A^2 + n^2}} - \frac{1}{\sqrt{B^2 + n^2}} \right],$$

où n prend toute valeur entière, zéro compris.

Si nous essayons d'évaluer par la méthode des trapèzes l'intégrale :

$$(A2) \quad J(A, B) = \int_{-\infty}^{+\infty} \left[\frac{1}{\sqrt{A^2 + n^2}} - \frac{1}{\sqrt{B^2 + n^2}} \right] dn = \int_{-\infty}^{+\infty} Y(n) dn$$

en découpant l'intervalle d'intégration en sous-intervalles par les points d'abscisse entière, nous sommes conduit précisément à la somme $S(A, B)$. On sait que la méthode des trapèzes est d'autant plus précise que la dérivée seconde de la fonction à intégrer est plus faible. Or, on calcule aisément que

$$(A3) \quad Y''(n) = \frac{2n^2 - A^2}{(A^2 + n^2)^{5/2}} - \frac{2n^2 - B^2}{(B^2 + n^2)^{5/2}}.$$

Cette dérivée seconde tend vers 0 lorsque A et B augmentent indéfiniment. Par conséquent, pour les grandes valeurs de A et B , on peut admettre que

$$(A4) \quad S(A, B) \# J(A, B) = \left| \operatorname{Log} \frac{n + \sqrt{A^2 + n^2}}{n + \sqrt{B^2 + n^2}} \right|_{-\infty}^{+\infty}$$

Pour la limite supérieure, la fraction est équivalente à $2n/2n$ et le logarithme est nul. Pour la limite inférieure, nous pouvons écrire $n = -N$

$$\begin{aligned} \operatorname{Log} \frac{\sqrt{A^2 + N^2} - N}{\sqrt{B^2 + N^2} - N} &= \operatorname{Log} \frac{-1 + \sqrt{1 + \frac{A^2}{N^2}}}{-1 + \sqrt{1 + \frac{B^2}{N^2}}} \\ &\approx \operatorname{Log} \frac{\frac{1}{2} \frac{A^2}{N^2}}{\frac{1}{2} \frac{B^2}{N^2}} = \operatorname{Log} \frac{A^2}{B^2}. \end{aligned}$$

Finalement

$$(A5) \quad J(A, B) = \operatorname{Log} \frac{B^2}{A^2}.$$

Nous allons évaluer l'ordre de grandeur de l'erreur commise en assimilant S à J lorsque A et B ne sont pas très grands, disons de l'ordre de 1 seulement.

L'allure de la fonction $Y(n)$ est représentée sur la fig. 11. Elle est symétrique en n , et il nous suffit de considérer seulement les valeurs positives de n . La courbe $Y(n)$ présente un point d'inflexion pour une certaine valeur n_0 de la variable. Cette valeur est

celle qui annule $Y''(n)$, c'est-à-dire l'expression (A3).

Pour $0 \leq n \leq n_0$, la méthode des trapèzes fournit une valeur par défaut de l'intégrale. Pour $n > n_0$, elle fournit une valeur par excès. Supposons $B > A$.

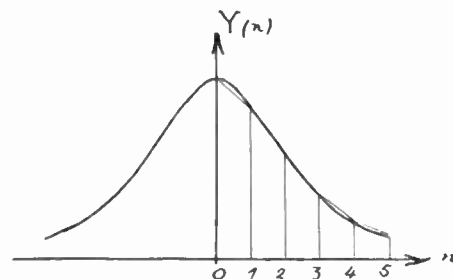


FIG. 11.

Pour $n = 0$, $Y'' = 1/B^3 - 1/A^3 < 0$. Pour $B = \sqrt{2}$,

Y'' se réduit pour $n = 1$ à $(2 - A^2)/(A^2 + 1)^{5/2} > 0$;

par conséquent, n_0 est plus petit que 1 et la méthode des trapèzes fournit une valeur par excès dès que n est égal ou supérieur à 1. Isolons donc le terme $n = 0$ dans la somme (A1)

$$(A6) \quad S(A, B) = \frac{1}{A} - \frac{1}{B} + 2 \sum_{n=1}^{\infty} Y(n).$$

La méthode des trapèzes appliquée au calcul de l'intégrale

$$I = \int_1^{\infty} Y(n) dn$$

conduit à la somme

$$\frac{1}{2} Y(1) + Y(2) + Y(3) + \dots = \sum_{n=1}^{\infty} Y(n) - \frac{1}{2} Y(1).$$

D'où

$$(A7) \quad 2 \sum_{n=1}^{\infty} Y(n) \# Y(1) + 2 \int_1^{\infty} Y(n) dn$$

et

$$S(A, B) \# \frac{1}{A} - \frac{1}{B} + \frac{1}{\sqrt{A^2 + 1}} - \frac{1}{\sqrt{B^2 + 1}} + 2I.$$

En se reportant à (A4), on constate aisément que

$$2I = 2 \operatorname{Log} \frac{\sqrt{B^2 + 1} + 1}{\sqrt{A^2 + 1} + 1}$$

D'où une nouvelle expression de $S(A, B)$:

$$(A8) \quad S(A, B) = \frac{1}{A} - \frac{1}{B} + \frac{1}{\sqrt{A^2 + 1}} - \frac{1}{\sqrt{B^2 + 1}} + 2 \operatorname{Log} \frac{\sqrt{B^2 + 1} + 1}{\sqrt{A^2 + 1} + 1}.$$

L'erreur commise est égale au double de la différence entre la valeur exacte de l'intégrale I et son approximation

$$\sum_{n=1}^{\infty} [Y(n) - \frac{1}{2} Y(1)].$$

D'après la formule d'Euler-MacLaurin, (voir par exemple J. HADAMARD, Cours d'Analyse, Tome I, § 105, Hermann Ed., Paris 1927) cette différence est de l'ordre de grandeur de $\frac{1}{12} [Y'(1) - Y'(\infty)]$.

On en déduit l'ordre de grandeur de l'approximation de la formule (A 8)

$$\begin{aligned} \frac{1}{6} Y'(1) &= \frac{1}{6} \left[\frac{-n}{(A^2 + n^2)^{3/2}} + \frac{n}{(B^2 + n^2)^{3/2}} \right]_{n=1} \\ &= \frac{1}{6} \left[-\frac{1}{(A^2 + 1)^{3/2}} + \frac{1}{(B^2 + 1)^{3/2}} \right]. \end{aligned}$$

D'où une approximation encore meilleure de $S(A, B)$:

$$\begin{aligned} (A9) \quad S(A, B) &= \frac{1}{A} - \frac{1}{B} + \frac{1}{\sqrt{A^2 + 1}} - \frac{1}{\sqrt{B^2 + 1}} \\ &+ 2 \log \frac{\sqrt{A^2 + 1} + 1}{\sqrt{B^2 + 1} + 1} + \frac{1}{6} \left[\frac{1}{(A^2 + 1)^{3/2}} - \frac{1}{(B^2 + 1)^{3/2}} \right]. \end{aligned}$$

Nous allons comparer les résultats des formules (A5), (A8) et (A9) pour des valeurs particulières pas trop élevées de A et B , par exemple $A = 1$ et $B = 2$.

$$(A5) \rightarrow S(1, 2) = 1,3862$$

$$(A8) \rightarrow 1,3468$$

$$(A9) \rightarrow 1,3908.$$

L'écart entre les résultats de (A5) et (A9) est seulement de 3,3 ‰ en valeur relative.

On en conclut qu'en dépit de sa simplicité la formule (A 5) fournit une excellente approximation de la somme $S(A, B)$, même pour des valeurs de A et B voisines de l'unité.

APPENDICE II

$$1. - \text{Calcul de } \sum_{m=1}^{\infty} \frac{1}{\lambda^2 + m^2}.$$

Considérons dans le plan complexe la fonction

$$(A10) \quad \varphi(z) = \frac{\pi}{\lg \pi z} \frac{1}{\lambda^2 + z^2}.$$

Les seules singularités de cette fonction sont les pôles réels $z = m$ (m entier positif, négatif ou nul), et les pôles imaginaires $z = \pm i\lambda$. Lorsqu'on augmente indéfiniment $|z|$ par valeurs discrètes en évitant les valeurs réelles entières, $\varphi(z)$ tend vers 0

comme $1/|z|^2$. L'intégrale de $\varphi(z)$ sur le cercle de l'infini est donc nulle. Il en résulte, d'après le théorème de Cauchy, que la somme des résidus de $\varphi(z)$ est nulle.

Le résidu pour $z = m$ est égal à $1/(\lambda^2 + m^2)$.

Le résidu au point $i\lambda$ a pour valeur

$$\frac{\pi}{\operatorname{tg} \pi i \lambda} \frac{1}{2 i \lambda} = -\frac{\pi}{2 \lambda \operatorname{Th} \pi \lambda}.$$

Le résidu au point $-i\lambda$ a même valeur, comme on le voit en changeant λ en $-\lambda$: Finalement

$$(A11) \quad \sum_{m=-\infty}^{+\infty} \frac{1}{\lambda^2 + m^2} = \frac{2\pi}{2\lambda \operatorname{Th} \pi \lambda} = \frac{\pi}{\lambda} \operatorname{Coth} \pi \lambda.$$

$$2. - \text{Calcul de } \sum_{m=1}^{\infty} \operatorname{Log} \left(1 + \frac{t^2}{m^2} \right).$$

$$\text{Posons (A12) } f(t^2) = \sum_{m=1}^{\infty} \operatorname{Log} \left(1 + \frac{t^2}{m^2} \right),$$

et dérivons par rapport à t^2 :

$$(A13) \quad \frac{d f(t^2)}{d(t^2)} = \sum_{m=1}^{\infty} \frac{1/m^2}{1 + t^2/m^2} = \sum_{m=1}^{\infty} \frac{1}{m^2 + t^2}.$$

La dérivation est justifiée par le fait que la série (A13) est uniformément convergente.

D'après (A11):

$$\frac{\pi}{t} \operatorname{Coth} \pi t = \sum_{m=-\infty}^{+\infty} \frac{1}{m^2 + t^2} = \frac{1}{t^2} + 2 \sum_{m=1}^{\infty} \frac{1}{m^2 + t^2}$$

$$\text{d'où (A14) } \sum_{m=1}^{\infty} \frac{1}{m^2 + t^2} = \frac{\pi}{2t} \operatorname{Coth} \pi t - \frac{1}{2t^2}.$$

En comparant (A13) et (A14), nous voyons que

$$(A15) \quad \frac{d f(t^2)}{d(t^2)} = \frac{\pi}{2t} \operatorname{Coth} \pi t - \frac{1}{2t^2}.$$

Par intégration

$$\begin{aligned} (A16) \quad f(t^2) &= \int \left(\frac{\pi}{2t} \operatorname{Coth} \pi t - \frac{1}{2t^2} \right) d(t^2) \\ &= \int \left(\pi \operatorname{Coth} \pi t - \frac{1}{t} \right) dt = \operatorname{Log} \frac{\operatorname{Sh} \pi t}{t} + C^c. \end{aligned}$$

On voit manifestement sur (A12) que f s'annule pour $t = 0$.

On en déduit que la constante vaut $-\text{Log } \pi$,
et par conséquent

$$(A17) \quad f(t^2) = \sum_1^{\infty} \text{Log} \left(1 + \frac{t^2}{m^2} \right) = \text{Log} \frac{\text{Sh } \pi t}{\pi t}.$$

3. — Calcul de

$$\sum_1^{\infty} \text{Log} \frac{a^2 + (my_0 - a)^2}{(my_0 - a)^2} \frac{a^2 + (my_0 + a)^2}{(my_0 + a)^2}.$$

On a identiquement

$$a^2 + (my_0 + a)^2 = y_0^2 \left[m + \frac{a}{y_0} (1 + i) \right] \left[m + \frac{a}{y_0} (1 - i) \right],$$

et

$$a^2 + (my_0 - a)^2 = y_0^2 \left[m - \frac{a}{y_0} (1 + i) \right] \left[m - \frac{a}{y_0} (1 - i) \right].$$

Faisons le produit membre à membre :

$$[a^2 + (my_0 + a)^2] [a^2 + (my_0 - a)^2]$$

$$= y_0^4 \left[m^2 - \frac{a^2}{y_0^2} (1 + i)^2 \right] \left[m^2 - \frac{a^2}{y_0^2} (1 - i)^2 \right]$$

et

$$\frac{[a^2 + (my_0 + a)^2] [a^2 + (my_0 - a)^2]}{(my_0 + a)^2 (my_0 - a)^2}$$

$$= \frac{\left[m^2 - \frac{a^2}{y_0^2} (1 + i)^2 \right] \left[m^2 - \frac{a^2}{y_0^2} (1 - i)^2 \right]}{\left(m^2 - \frac{a^2}{y_0^2} \right)^2}$$

$$= \frac{\left[1 - \frac{1}{m^2} \frac{a^2}{y_0^2} (1 + i)^2 \right] \left[1 - \frac{1}{m^2} \frac{a^2}{y_0^2} (1 - i)^2 \right]}{\left(1 - \frac{a^2}{m^2 y_0^2} \right)^2}.$$

Par application de la formule (A17), nous obtenons pour la somme des logarithmes de ces fractions, la formule :

(A18)

$$\text{Log} \frac{\text{Sh} \left[\pi i \frac{a}{y_0} (1 + i) \right] \cdot \text{Sh} \left[\pi i \frac{a}{y_0} (1 - i) \right]}{\left| \frac{\pi i a}{y_0} (1 + i) \right| \left| \frac{\pi i a}{y_0} (1 - i) \right|} = 2 \text{Log} \frac{\text{Sh} \frac{\pi i a}{y_0}}{\frac{\pi i a}{y_0}}.$$

$$\begin{aligned} \text{Or : } & \text{Sh} \left[\pi i \frac{a}{y_0} (1 + i) \right] \cdot \text{Sh} \left[\pi i \frac{a}{y_0} (1 - i) \right] \\ &= \frac{1}{2} \text{Ch} \left[\pi i \frac{a}{y_0} (1 + i) + \pi i \frac{a}{y_0} (1 - i) \right] \\ &= \frac{1}{2} \text{Ch} \left[\pi i \frac{a}{y_0} (1 + i) - \pi i \frac{a}{y_0} (1 - i) \right] \end{aligned}$$

$$\begin{aligned} &= \frac{1}{2} \text{Ch} \frac{2 \pi i a}{y_0} - \frac{1}{2} \text{Ch} \left(- \frac{2 \pi a}{y_0} \right) \\ &= \frac{1}{2} \left[\cos \frac{2 \pi a}{y_0} - \text{Ch} \frac{2 \pi a}{y_0} \right]. \end{aligned}$$

Reportons dans (A18), nous obtenons

$$\begin{aligned} (A19) \quad & \text{Log} \frac{\frac{1}{2} \left[\cos \frac{2 \pi a}{y_0} - \text{Ch} \frac{2 \pi a}{y_0} \right]}{-2 \pi^2 \frac{a^2}{y_0^2}} - 2 \text{Log} \frac{\sin \frac{\pi a}{y_0}}{\frac{\pi a}{y_0}} \\ &= \text{Log} \frac{\text{Ch} \frac{2 \pi a}{y_0} - \cos \frac{2 \pi a}{y_0}}{4 \sin^2 \frac{\pi a}{y_0}}. \end{aligned}$$

APPENDICE III

Calcul des fonctions K et G pour les valeurs de l'argument inférieures à l'unité.

Par définition :

$$\begin{aligned} (A20) \quad K &= 1 + 4s \sum_{n=1}^{\infty} \left[\frac{1}{\sqrt{s^2 + n^2}} - \frac{1}{\sqrt{(2s)^2 + n^2}} \right] = 1 + \frac{4s}{\sqrt{1 + s^2}} \\ &\quad - \frac{4s}{\sqrt{1 + 4s^2}} + 4s \sum_{n=2}^{\infty} \left[\frac{1}{\sqrt{s^2 + n^2}} - \frac{1}{\sqrt{(2s)^2 + n^2}} \right]. \end{aligned}$$

Comme s est inférieur à 1, on peut développer $1/\sqrt{s^2 + n^2}$ par la formule du binôme

$$\begin{aligned} (A21) \quad \frac{1}{\sqrt{s^2 + n^2}} &= \frac{1}{n} \left(1 + \frac{s^2}{n^2} \right)^{-\frac{1}{2}} \\ &= \frac{1}{n} - \frac{1}{2} \frac{s^2}{n^3} + \frac{1.3}{2^2 2!} \frac{s^4}{n^5} - \frac{1.3.5}{2^3 3!} \frac{s^6}{n^7} + \dots + \\ &+ (-1)^q \frac{1.3.5 \dots (2q-1)}{2^q q!} \frac{s^{2q}}{n^{2q+1}} + \dots \end{aligned}$$

Le second terme de la somme fournit un développement analogue valable également pour $s \leq 1$, puisque dans la somme $n \geq 2$, et donc $2s \leq n$. Il suffit de remplacer partout s par $2s$. On obtient ainsi

$$\frac{1}{\sqrt{(2s)^2 + n^2}} = \sum_{n=1}^{\infty} (-1)^q \frac{1.3.5 \dots (2q-1)}{2^q q!} \frac{(2s)^{2q}}{n^{2q+1}}.$$

Introduisons les fonctions de Riemann

$$(A23) \quad \zeta(u) = \sum_{n=1}^{\infty} \frac{1}{n^u}.$$

On voit que

$$\sum_{n=2}^{\infty} \frac{1}{n^{2q+1}} = \zeta(2q+1) - 1.$$

Nous obtenons finalement le développement en série entière de la somme ordonnée suivant les puissances croissantes de s :

$$(A24) \quad \sum_{n=2}^{\infty} \left[\frac{1}{\sqrt{s^2 + n^2}} - \frac{1}{\sqrt{(2s)^2 + n^2}} \right] \\ = \frac{s^2}{2} (2^2 - 1) [\zeta(3) - 1] \\ - s^4 \frac{1.3}{2^2 2!} (2^4 - 1) [\zeta(5) - 1] + s^6 \frac{1.3.5}{2^3 3!} (2^6 - 1) \\ [\zeta(7) - 1] + \dots + (-1)^{q-1} s^{2q} \frac{1.3.5 \dots (2q-1)}{2^q q!} \\ (2^{2q} - 1) [\zeta(2q+1) - 1] + \dots$$

La série étant alternée, l'erreur est inférieure au premier terme négligé. De plus, les coefficients $\zeta(2q+1) - 1$ décroissent très rapidement comme en témoigne la liste ci-dessous extraite de la table de Jahnke-Emde :

$$\begin{array}{ll} \zeta(3) - 1 = 0,202 & \zeta(5) - 1 = 0,0369 \\ \zeta(7) - 1 = 0,00835 & \zeta(9) - 1 = 0,002008 \\ \zeta(11) - 1 = 0,000494 & \zeta(13) - 1 = 0,0001227 \\ \zeta(15) - 1 = 0,0000306 & \zeta(17) - 1 = 0,00000764 \end{array}$$

Au delà, la fonction $\zeta(2q+1) - 1$ se réduit pratiquement à $\frac{1}{2^{2q-1}}$, et chacun des termes devient égal au quart du précédent.

Le calcul de la fonction G s'effectue de façon analogue.

Les termes $2^{2q} - 1$ de la formule (A24) sont remplacés par des termes de la forme $2^q - 1$. La convergence est plus rapide pour une même valeur de s et a lieu jusqu'à $s = \sqrt{2}$.

On peut étendre le domaine de convergence en isolant dans la somme le terme $n = 2$ comme nous l'avons fait pour $n = 1$.

APPENDICE IV

Influence d'une dimension limitée suivant la coordonnée x .

Les calculs développés dans le présent article reposaient sur l'hypothèse d'un échantillon indéfiniment étendu suivant l'axe des x (axe de l'alignement des pointes d'injection de courant). On peut se demander quel est l'ordre de grandeur de l'erreur commise lorsque cette hypothèse n'est pas fondée, par exemple lorsque la surface de l'échantillon se rapproche de la forme carrée.

On se rend bien compte intuitivement que cette erreur est petite, car la densité des lignes de courant est faible dans les régions extérieures à l'intervalle des pointes (voir figure 12), de sorte qu'on ne perturbe que très faiblement leur répartition dans cet

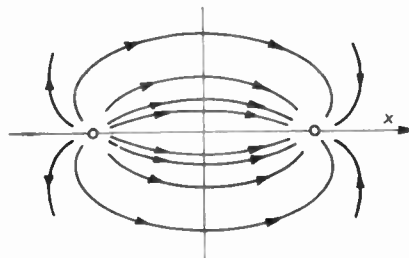


FIG. 12.

intervalle lorsqu'on limite l'échantillon suivant la direction Ox ,

Il est facile de généraliser les formules déjà établies. Il suffit d'ajouter, dans les formules donnant le potentiel, la contribution des images des points — sources par rapport aux nouveaux plans-limites $x = +x_0$ et $x = -x_0$. L'examen de la figure 13 montre que les abscisses des nouveaux points-images de la source A sont données par la formule générale :

$$(A25) \quad \alpha_q = 2q x_0 + (-1)^q l,$$

q prenant toute valeur entière de $-\infty$ à $+\infty$, zéro exclu. Le terme $q = 0$ redonne d'ailleurs le terme général qui seul avait été pris en considération.

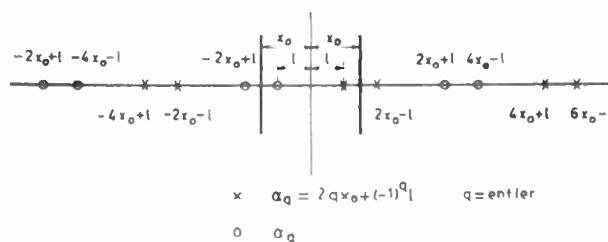


FIG. 13.

Les images de la source B s'obtiennent en prenant les valeurs opposées à α_q , puisque B est symétrique de A par rapport à l'origine.

Bornons-nous au cas des échantillons d'épaisseur faible étudiés par la méthode des pointes alignées.

Dans le cas d'un échantillon indéfini suivant la direction Ox , la tension recueillie est donnée par la formule (14) que nous retranscrivons ci-après :

$$(14) \quad v = \frac{\rho I}{\pi z_0} \log \frac{\text{Sh } \pi \frac{l+a}{2y_0}}{\text{Sh } \pi \frac{l-a}{2y_0}}.$$

On en déduit l'expression de v dans le cas de l'échantillon limité :

$$(A26) \quad v = \frac{\rho I}{\pi z_0} \left[\text{Log} \frac{Sh \pi \frac{l+a}{2y_0}}{Sh \pi \frac{l-a}{2y_0}} + \sum_{q \neq 0} \text{Log} Sh \frac{\pi}{2y_0} (\alpha_q + a) - \sum_{q \neq 0} \text{Log} Sh \frac{\pi}{2y_0} (\alpha_q - a) \right].$$

Ces deux derniers termes correctifs étant faibles, il nous suffit d'en calculer une valeur approchée. L'argument des fonctions hyperboliques est grand en valeur absolue.

Nous avons rigoureusement :

$$\begin{aligned} Sh \frac{\pi}{2y_0} (\alpha_q + a) &= \frac{1}{2} \left[e^{\frac{\pi}{2y_0} (\alpha_q + a)} - e^{-\frac{\pi}{2y_0} (\alpha_q + a)} \right] \\ &= \frac{1}{2} e^{\frac{\pi}{2y_0} (\alpha_q + a)} \left[1 - e^{-\frac{\pi}{y_0} (\alpha_q + a)} \right]. \end{aligned}$$

et sensiblement, pour $q > 0$ ($\alpha_q > 0$), donc $e^{-\frac{\pi}{y_0} (\alpha_q + a)} \ll 1$,

$$(A27) \quad \begin{aligned} \text{Log} Sh \frac{\pi}{2y_0} (\alpha_q + a) \\ \approx \text{Log} \frac{1}{2} + \frac{\pi}{2y_0} (\alpha_q + a) - e^{-\frac{\pi}{y_0} (\alpha_q + a)}. \end{aligned}$$

En retranchant, comme indiqué par la formule (A26), la contribution des termes en $-a$, nous obtenons :

$$(A28) \quad \frac{\pi a}{y_0} - e^{-\frac{\pi a}{y_0}} \left[e^{\frac{\pi a}{y_0}} - e^{-\frac{\pi a}{y_0}} \right] = \frac{\pi a}{y_0} + 2 e^{-\frac{\pi a}{y_0}} Sh \frac{\pi a}{y_0}.$$

Pour les valeurs négatives de q , nous obtenons par un calcul analogue la formule homologue de (A27)

$$(A29) \quad \begin{aligned} \text{Log} Sh \frac{\pi}{2y_0} (\alpha_q + a) \\ \approx \text{Log} \left(-\frac{1}{2} \right) - \frac{\pi}{2y_0} (\alpha_q + a) - e^{\frac{\pi}{y_0} (\alpha_q + a)}, \end{aligned}$$

puis la formule homologue de (A28)

$$(A30) \quad -\frac{\pi a}{y_0} - e^{\frac{\pi a}{y_0}} \left[e^{\frac{\pi a}{y_0}} - e^{-\frac{\pi a}{y_0}} \right] = -\frac{\pi a}{y_0} - 2 e^{\frac{\pi a}{y_0}} Sh \frac{\pi a}{y_0}.$$

Il faut maintenant sommer les expressions (A28) et (A30) sur toutes les valeurs entières, respectivement positives et négatives, de q , le terme $q = 0$ étant exclu. Remarquons dès maintenant que le terme $\pi a/y_0$ disparaîtra dans l'addition de (A28) et (A30) et nous ne nous en occuperons plus.

Il faut encore distinguer entre les valeurs paires et impaires de q :

1° q pair positif. $q = 2k, k = 1, 2, 3, \dots$

$$\begin{aligned} \alpha_q &= 4kx_0 + l, \\ \sum_q e^{-\frac{\pi \alpha_q}{y_0}} &= \sum_k e^{-\frac{4\pi kx_0}{y_0} - \frac{\pi l}{y_0}} = e^{-\frac{\pi l}{y_0}} \sum_k e^{-\frac{4\pi kx_0}{y_0}}. \end{aligned}$$

Cette dernière somme est celle d'une progression géométrique de raison $e^{-\frac{4\pi x_0}{y_0}}$ et dont le premier terme est également $e^{-\frac{4\pi x_0}{y_0}}$. Elle a donc pour valeur $e^{-\frac{\pi l}{y_0}} / \left[1 - e^{-\frac{4\pi x_0}{y_0}} \right]$

$$(A31) \quad \sum_q e^{-\frac{\pi \alpha_q}{y_0}} = e^{-\frac{\pi l}{y_0}} e^{-\frac{4\pi x_0}{y_0}} / \left[1 - e^{-\frac{4\pi x_0}{y_0}} \right] = \frac{e^{-\frac{\pi l}{y_0}}}{e^{\frac{4\pi x_0}{y_0}} - 1}.$$

2° q impair positif $q = 2k - 1, k = 1, 2, 3, \dots$

$$(A32) \quad \begin{aligned} \alpha_q &= 4kx_0 - 2x_0 - l, \\ \sum_q e^{-\frac{\pi \alpha_q}{y_0}} &= e^{\frac{\pi}{y_0} (2x_0 + l)} \sum_k e^{-\frac{4\pi kx_0}{y_0}} = \frac{e^{\frac{2\pi x_0}{y_0}} e^{\frac{\pi l}{y_0}}}{e^{\frac{4\pi x_0}{y_0}} - 1} \end{aligned}$$

3° q pair négatif $q = -2k, k = 1, 2, 3, \dots$

$$(A33) \quad \begin{aligned} \alpha_q &= -4kx_0 + l, \\ \sum_q e^{\frac{\pi \alpha_q}{y_0}} &= e^{\frac{\pi l}{y_0}} \sum_k e^{-\frac{4\pi kx_0}{y_0}} = \frac{e^{\frac{\pi l}{y_0}}}{e^{\frac{4\pi x_0}{y_0}} - 1}. \end{aligned}$$

4° q impair négatif $q = -2k + 1, k = 1, 2, 3, \dots$

$$(A34) \quad \begin{aligned} \alpha_q &= -4kx_0 + 2x_0 - l, \\ \sum_q e^{\frac{\pi \alpha_q}{y_0}} &= e^{\frac{\pi}{y_0} (2x_0 - l)} \sum_k e^{-\frac{4\pi kx_0}{y_0}} = \frac{e^{\frac{2\pi x_0}{y_0}} e^{-\frac{\pi l}{y_0}}}{e^{\frac{4\pi x_0}{y_0}} - 1}. \end{aligned}$$

Additionnons les expressions (A31), (A32), (A33), (A34), avec les signes imposés par (A28) et (A30) :

$$(A35) \quad \begin{aligned} 2 Sh \frac{\pi a}{y_0} \left[e^{\frac{\pi l}{y_0}} + e^{\frac{2\pi x_0}{y_0}} e^{\frac{\pi l}{y_0}} - e^{\frac{\pi l}{y_0}} - e^{\frac{2\pi x_0}{y_0}} e^{-\frac{\pi l}{y_0}} \right] \frac{1}{e^{\frac{4\pi x_0}{y_0}} - 1} \\ = 2 Sh \frac{\pi a}{y_0} \frac{\left(e^{\frac{\pi l}{y_0}} - e^{-\frac{\pi l}{y_0}} \right) \left(e^{\frac{2\pi x_0}{y_0}} - 1 \right)}{e^{\frac{4\pi x_0}{y_0}} - 1} = \frac{4 Sh \frac{\pi a}{y_0} \cdot Sh \frac{\pi l}{y_0}}{e^{\frac{2\pi x_0}{y_0}} + 1}. \end{aligned}$$

D'où finalement l'expression de la valeur corrigée de v :

(A36)

$$v = \frac{\rho I}{\pi z_0} \left[\text{Log} \frac{\text{Sh } \pi \frac{l+a}{2y_0}}{\text{Sh } \pi \frac{l-a}{2y_0}} + \frac{4 \text{Sh } \frac{\pi a}{y_0} \cdot \text{Sh } \frac{\pi l}{y_0}}{e^{\frac{2\pi x_0}{y_0}} + 1} \right].$$

Cas particulier des pointes équidistantes ($l = 3 a$) :
(A37)

$$v = \frac{\rho I}{\pi z_0} \left[\text{Log } 2 \text{Ch } \frac{\pi a}{y_0} + \frac{4 \text{Sh } \frac{\pi a}{y_0} \cdot \text{Sh } \frac{3\pi a}{y_0}}{e^{\frac{2\pi x_0}{y_0}} + 1} \right].$$

L'erreur relative commise en considérant l'échantillon comme indéfini suivant la direction Ox est :

$$\varepsilon = \frac{1}{e^{\frac{2\pi x_0}{y_0}} + 1} - \frac{4 \text{Sh } \frac{\pi a}{y_0} \cdot \text{Sh } \frac{3\pi a}{y_0}}{\text{Log } 2 \text{Ch } \frac{\pi a}{y_0}}$$

Le calcul montre que pour $x_0 = y_0 = 2l = 6a$,

$\varepsilon = 1,1 \%$. Pour $x_0 = y_0 = l = 3 a$, cas limite où les pointes reposent sur les bords de l'échantillon, $\varepsilon = 9 \%$.

On peut donc, avec une bonne approximation, considérer l'échantillon comme indéfiniment étendu suivant l'axe d'alignement des pointes. Un calcul analogue pourrait être effectué sans trop de difficulté dans le cas des pointes en carré et dans le cas de la mesure de l'effet Hall. Il est hors de doute qu'on aboutirait à la même conclusion.

L'échantillon étudié se présente souvent sous forme de pastille approximativement circulaire. L'étude de la répartition des lignes de courant n'est guère accessible au calcul. Toutefois, il est intuitif que la répartition des lignes de courant dans la section circulaire diffère peu de la répartition dans le carré circonscrit, et celle-ci ne se modifie guère, nous venons de le voir, lorsqu'on augmente indéfiniment la dimension suivant l'axe Ox .

On en conclut que les formules établies dans cet article sont pratiquement applicables aux pastilles circulaires, le paramètre y_0 désignant le rayon de la pastille. Cette remarque permet une extension considérable du champ d'application de ces formules.

ONDES ÉLECTROMAGNÉTIQUES POLARISÉES ELLIPTIQUES ET CIRCULAIRES⁽¹⁾

PAR

Maurice BOUX
Docteur ès Sciences

Généralités

En optique, une source lumineuse monochromatique est pratiquement constituée par un grand nombre d'oscillateurs électromagnétiques élémentaires qui vibrent avec des phases et des polarisations distribuées au hasard. À partir d'une certaine distance de la source, le champ électromagnétique est perpendiculaire à la direction de propagation ; son énergie reste constante en moyenne, mais le champ électrique fluctue très rapidement en intensité et en direction autour de la direction de propagation. On dit qu'on a un faisceau de lumière non polarisée. Si sur le trajet d'un faisceau de lumière parallèle on dispose un cristal de tourmaline convenablement taillé qu'on appelle nicol polariseur, ce cristal a la propriété de ne laisser passer que la composante du champ électrique située dans un plan donné. Le faisceau à la sortie du nicol est formé de lumière polarisée rectiligne.

On le vérifie en plaçant sur le faisceau un second nicol, dit nicol analyseur, qui laisse passer le faisceau si son plan de polarisation est parallèle au plan de polarisation du premier nicol, et qui éteint progressivement le faisceau si on fait tourner ce nicol jusqu'à ce que son plan de polarisation devienne perpendiculaire au premier. Si on interpose entre les deux nicols un cristal biréfringent convenablement taillé et convenablement orienté, l'onde qui sortira sera polarisée elliptique ou exceptionnellement circulaire, et on pourra le vérifier en faisant tourner le nicol analyseur autour de l'axe du faisceau (fig. 1).

⁽¹⁾ Conférence faite à la Société Française des Radioélectriciens, le 21 Mai 1954.

Dans le cas des ondes électromagnétiques de la radioélectricité, on n'a habituellement comme source qu'un oscillateur unique qui excite une antenne de forme bien déterminée. Contrairement à ce qui se passe en optique, cette antenne rayonne donc une vibration électromagnétique qui, en chaque point, a une forme et une phase bien déterminée. Tant qu'on n'a pas recherché une forme particulière de la vibration, on a construit les antennes de la façon la plus commode possible, et on a abouti à des formes qui donnent généralement une polarisation rectiligne. Il en est ainsi pour les doubles ou les cadres usuels dans les liaisons de trafic. Aux hyperfréquences, on a utilisé des dipôles, des fentes dans les guides ou des cornets, et on a obtenu ainsi de la polarisation rectiligne.

Pour avoir la meilleure utilisation possible, il suffit de disposer à la réception une antenne à polarisation rectiligne placée au mieux dans le champ incident qui provient de l'émetteur, c'est-à-dire telle que le champ qu'elle pourrait émettre soit parallèle à celui qu'elle reçoit.

Mais quand on a eu à réaliser soit des liaisons de trafic, soit des liaisons d'interrogateur-répondeur avec les avions, on a dû tenir compte du fait que l'antenne de l'avion est en général fixe par rapport à l'avion et peut être loin de la position optimum pour la liaison. Il est donc indiqué dans certains cas d'augmenter la sécurité de la liaison en imposant soit à l'antenne de réception, soit à l'antenne d'émission d'être polarisée elliptique. Cette nécessité est encore plus impérieuse lorsqu'il s'agit de fusées radioguidées qui non seulement évoluent dans l'espace à la manière des avions, mais encore peuvent tourner sur elles-mêmes autour de leur axe.

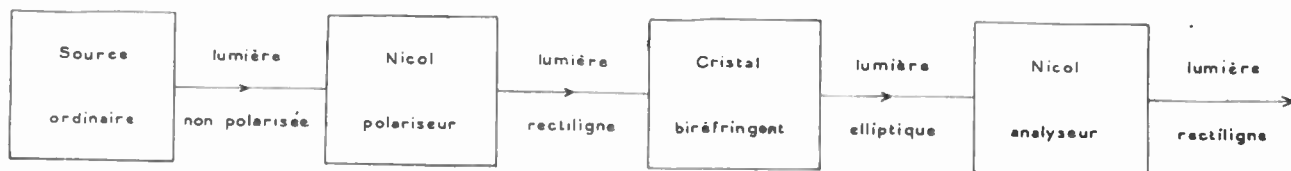


FIG. 1. — Polarisation et analyse d'un faisceau lumineux

D'un point de vue différent, l'écho d'un objectif de radar a une intensité différente suivant la polarisation de l'onde du radar, et il peut être recommandé de choisir pour le radar une polarisation appropriée aux objectifs qu'il cherche et pour éliminer ou en tous cas diminuer les échos inutiles ou nuisibles. Ainsi on a réussi, en choisissant une polarisation circulaire à diminuer sensiblement les échos dus à la pluie, et d'augmenter par là-même le contraste entre les échos dus à un avion qui serait dans un nuage de pluie et les gouttes de pluie elles-mêmes. Cela est dû au fait que si on envoie sur un objet sphérique (goutte d'eau) une onde polarisée circulaire droite par exemple, cet objet renvoie une onde polarisée circulaire gauche, et que l'antenne du radar ne peut recevoir qu'une onde circulaire de même sens que celle qu'elle a émise. Ce phénomène est encore très intéressant dans le cas des câbles hertziens; si les antennes d'émission et de réception sont faites pour fonctionner en polarisation circulaire, l'antenne de réception recevra bien le faisceau direct, mais sera peu sensible aux ondes qui proviendraient des réflexions sur le sol. Si le faisceau de liaison est assez près du sol, au-dessus de forêts par exemple, ou assez près de la mer, les effets du vent qui agite les arbres ou qui crée les vagues ne causeront que peu de fluctuations à la liaison. Le principe de l'élimination de certaines réflexions par l'emploi de la polarisation circulaire est mentionné dans l'ouvrage de L.-N. RIDENOUR « *Radar System Engineering* » parag. 3.10 in fine.

Ce n'est pas que dans les faisceaux rayonnés qu'on trouve une polarisation elliptique ou circulaire. Elle existe aussi aux hyperfréquences à l'intérieur des cavités ou guides d'ondes, soit qu'on forme intentionnellement cette polarisation elliptique dans un but déterminé, soit qu'elle soit naturellement associée à un phénomène relativement simple.

Nous nous proposons dans cet exposé de passer en revue la littérature qui traite de la polarisation elliptique et circulaire. Cette question n'est pas traitée avec l'ampleur qu'elle mériterait dans les ouvrages, et il faudra en rechercher les divers aspects dans un certain nombre d'articles. Les titres de ces articles et les quelques explications que nous en donnerons permettront, j'espère, de se faire une idée de l'état actuel de la question.

Nous commencerons par les articles qui traitent de la polarisation elliptique à l'intérieur des cavités ou lignes de transmissions et dont par suite les applications sont limitées aux hyperfréquences et nous poursuivrons par les articles qui étudient les phénomènes liés au rayonnement; nous terminerons par quelques articles relatifs aux études théoriques.

Dans le courant de cet exposé nous indiquerons avec précision les références des publications que nous aurons pu étudier. Nous indiquerons dans la Bibliographie les auteurs et les titres d'autres publications que nous n'avons pu lire, mais qui

peuvent intéresser les lecteurs susceptibles de se les procurer.

Avant de citer les articles plus techniques, empruntons à celui de A. KASTLER « *La polarimétrie dans le domaine des ondes hertziennes* » paru dans le *Supplemento al Nuovo Cimento* Vol. IX, Série IX, N° 3, pp. 315-321, l'historique de la suite des premiers articles scientifiques sur la question.

HERTZ dans son travail fondamental (1889) avait mis en évidence que les ondes hertziennes obtenues à partir d'un doublet sont polarisées; une grille placée sur le faisceau a un effet comparable à celui d'un cristal de tourmaline sur un faisceau lumineux polarisé. A. RICHI a mis en évidence les propriétés de la polarisation dans la réflexion et la réfraction (*L'ottica delle oscillazioni elettriche* 1897). Il a obtenu une vibration circulaire avec un parallélépipède de FRESNEL en paraffine. RICHI et MACK en 1895 ont mis en évidence la biréfringence du bois, tandis que GARBASSO, LEBEDEV et BOSE ont étudié celle des cristaux (1896). Ils ont atteint la longueur d'onde de 6 mm et ont fait des nicols avec des monocristaux de soufre. GARBASSO a réalisé, avec des lames clivées de gypse, des lames demi-onde et quart d'onde. BOSE a mis en évidence des propriétés de polarisation rotatoire en travaillant sur des cristaux de quartz traversés par les ondes dans la direction de l'axe optique. Plus tard LINDMAN (1924) a étudié la polarisation rotatoire avec des modèles macroscopiques formés de boules placées d'une manière dissymétrique et avec des hélices métalliques.

Pour trouver des études dirigées vers les applications il faut arriver aux années de la guerre 1939-45 et nous allons les passer en revue ci-après.

2. — Polarisation elliptique dans les cavités et guides d'ondes.

Dans le cas des fréquences métriques et inférieures, les lignes de transmission sont habituellement soit des paires bifilaires nues ou blindées, soit des lignes coaxiales. Les lignes de champ électrique ou magnétique sont dans des plans perpendiculaires aux conducteurs; et ces champs, d'intensité variable avec le temps à l'intérieur de chaque période, ont une direction fixe: ils sont polarisés rectilignes (fig. 2). Il en est d'ailleurs

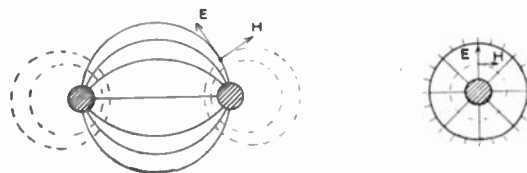


FIG. 2. — Champ électrique (trait plein) et champ magnétique (trait pointillé) a) d'une ligne bifilaire, b) d'un coaxial.

de même pour les vecteurs densités de courant sur les conducteurs.

Le phénomène est assez différent dans le cas de la transmission par guide d'onde (fig. 3a); dans le cas simple du guide d'onde rectangulaire par-

couru par une onde H_{01} , les lignes de champ électrique sont parallèles au petit côté du rectangle de section, et ce champ est bien polarisé rectiligne, mais la structure des lignes de champ magnétique est différente. A un instant donné ces lignes forment des familles de boucles (fig. 3b) situées

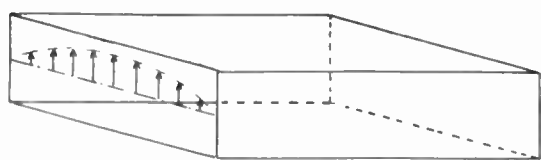


FIG. 3 a. — Champ électrique dans un guide rectangulaire pour une onde $T E_{01}$.

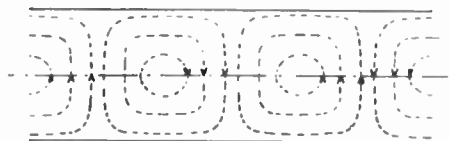


FIG. 3 b. — Champ magnétique d'une onde progressive $T E_{01}$ dans un guide rectangulaire.

dans des plans parallèles à la face large du guide d'onde, et la périodicité de ces boucles le long du guide d'onde est une demi-longueur d'onde guidée. Mais toute la figure formée par ces boucles se déplace avec le temps avec une vitesse égale à la vitesse de phase dans le guide. Ainsi si on

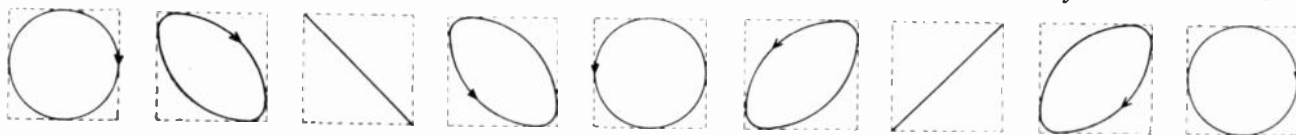


FIG. 6. — Variation de la polarisation d'une onde elliptique le long d'un guide à section voisine d'un carré.

examine le champ magnétique en un point M fixe, son support tourne et fait un tour par période. Une analyse simple montre que ce champ est polarisé elliptique. Si on étudie la polarisation en un point M variable dans un plan parallèle au grand côté depuis l'axe du guide jusqu'au petit côté, la polarisation du champ magnétique en M

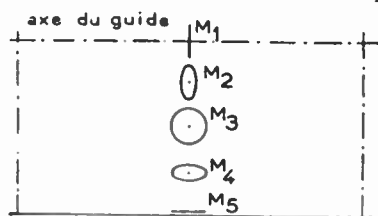


FIG. 4. — Variation de l'ellipse de polarisation du champ magnétique d'une onde progressive $T E_{10}$ dans un guide rectangulaire.

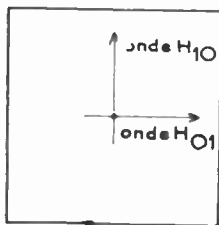


FIG. 5. — Décomposition en deux ondes rectilignes perpendiculaires d'une onde polarisée elliptique dans un guide carré.

est d'abord rectiligne perpendiculaire à l'axe, puis elliptique de grand axe perpendiculaire à l'axe du guide; elle passe par une polarisation circulaire, puis devient elliptique de grand axe parallèle à l'axe du guide, puis devient rectiligne parallèle à l'axe du guide. (fig. 4).

La polarisation du champ magnétique dans le guide d'onde rectangulaire est due au déphasage

que présentent ses deux composantes suivant l'axe du guide et sa perpendiculaire.

On peut obtenir au moyen d'un guide carré (fig. 5) un champ électrique polarisé elliptique (ou en particulier circulaire) en faisant parcourir le guide par deux ondes H_{01} et H_{10} analogues à celles du guide précédent, mais chacune parallèle à un côté du carré, et déphasées de 90° l'une par rapport à l'autre. Si le guide n'a pas une section exactement carrée, les deux longueurs d'ondes guidées ne seront pas les mêmes et à mesure qu'on se déplace le long du guide le déphasage variera; si on est parti de deux champs égaux, correspondant à une polarisation circulaire, le cercle initial se déforme peu à peu suivant les courbes de la fig. 6.

La flèche indique dans quel sens tourne le champ.

On voit par cet exemple toute l'importance que prendront les déphaseurs dans la génération et l'étude des ondes à polarisation elliptique, et on devine les précautions qu'il faudra prendre pour conserver à une onde une polarisation déterminée.

Aux fréquences centimétriques, on a utilisé ce guide carré pour faire de la polarisation circulaire. Un cornet basé sur ce principe et destiné à rayonner une onde polarisée circulaire est décrit dans l'ouvrage « *Very High Frequency Techniques* » Vol. 1, 1^{re} édition parag. 6.5, qui a été rédigé par le *Radio Research Laboratory* de l'Université de

HARVARD (Mc Grawhill 1947). Mais on a trouvé souvent plus pratique de travailler avec un guide circulaire fonctionnant dans le mode fondamental H_{11} (fig. 7). Pour cela on décompose par la pensée le vecteur champ électrique en deux composantes rectangulaires suivant deux axes dont il est la bissectrice, soient E_1 et E_2 , et on retarde de 90°

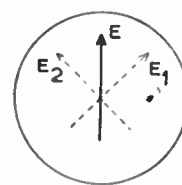


FIG. 7. — Décomposition d'une onde rectiligne E dans un guide circulaire en deux ondes rectilignes perpendiculaires entre elles E_1 et E_2

la phase de E_1 par exemple par rapport à E_2 . Les dispositifs utilisés sont soit une lame de diélectrique placée suivant le plan diamétral du guide, soit une ailette métallique, soit des tiges métalliques (fig. 8). Quel que soit le dispositif employé, le guide a une structure différente pour les ondes E_1 et E_2 dans la section qui contient le déphaseur, et les longueurs d'onde guidées sont différentes. On donne au déphaseur une longueur convenable pour que le déphasage soit de 90° . Dans le cas où le dispositif de déphasage est formé de deux tiges

métalliques, l'espacement des deux tiges est d'un quart de longueur d'onde guidée, et on agit sur le diamètre des tiges pour obtenir le déphasage voulu. Ce dernier dispositif a sur les deux précédents l'avantage d'une construction mécanique simple, mais il est plus sensible à la fréquence; pour augmenter sa largeur de bande, on remplace le dispositif à deux tiges par un dispositif à trois ou quatre tiges espacées d'un quart de longueur d'onde guidée et dont on a déterminé convenablement les diamètres.

On trouvera dans l'article de J.-F. RAMSAY, *Circular polarization for C. W. Radar* paru dans MARCONI REVIEW, Octobre 1952, des croquis assez clairs et des indications qui permettent de se faire une idée de ces dispositifs et de guider dans l'étude de leur bande passante. Cet article est le texte d'une conférence faite à Londres en Juin 1950 au

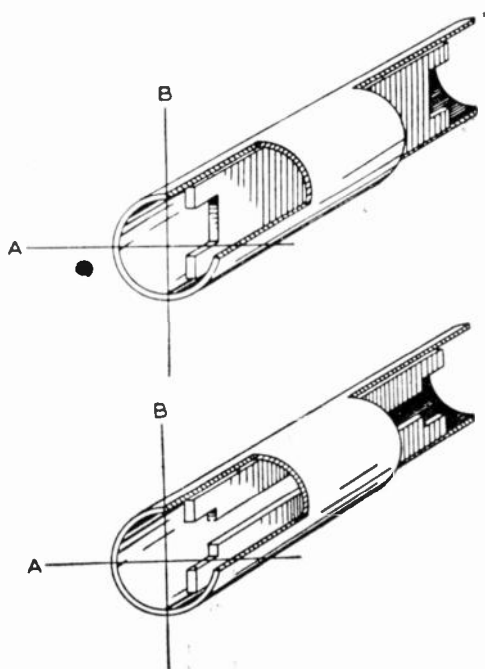


FIG. 8. — Dispositifs de déphasage d'une composante d'une onde qui se propage dans un guide circulaire (a) lame de diélectrique, (b) ailette métallique (Figures extraites de l'article de G.A. Fox « An adjustable wave guide phase changer » P.I.R.E., Vol. 35, pp. 1685-1698 Déc. 1947).

Congrès des Aériens Centimétriques pour les radars de navigation de la Marine. L'auteur indique qu'en 1945-46 une antenne de radar à polarisation circulaire avait été construite pour un radar à ondes entretenues sur la bande de 10 cm. La polarisation circulaire servait de dispositif TR pour découpler l'émission de la réception.

On utilisait le fait que l'onde circulaire, droite par exemple, réfléchi sur un plan perpendiculaire à la direction d'incidence se transforme en une onde circulaire gauche. Ainsi l'onde de réception ne peut pas revenir vers l'émetteur. On place sur son trajet un dispositif qui la dirige vers le récepteur. Malheureusement les objectifs renvoyaient une polarisation qui a une ellipticité sensible, et une partie du signal d'écho est perdue dans la voie de l'émetteur.

Un article de A. GARDNER FOX (1) « *An adjustable waveguide phase changer* » dans les P.I.R.E. de Décembre 1947 décrit les dispositifs du même genre utilisés comme déphaseurs; il donne pas mal de détails sur leur fonctionnement et leur construction et mentionne un joint tournant basé sur ce principe.

Les dispositifs décrits ci-dessus peuvent être utilisés simplement pour obtenir la polarisation circulaire destinée à l'alimentation d'une antenne, mais aussi, on peut en jumeler deux pour faire tourner d'un angle arbitraire le plan de polarisation dans le guide : On place comme sur la figure 9 un « polariseur » et un « analyseur ». L'analyseur reçoit l'onde circulaire et la transforme en

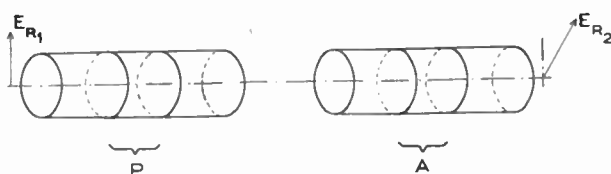


FIG. 9. — Dispositif servant à faire tourner le plan de polarisation d'une onde rectiligne.

onde rectiligne. En faisant tourner l'analyseur d'un certain angle, l'onde rectiligne de sortie tourne du même angle.

3. — Gyrateurs.

Le dispositif précédent a en commun avec tous les dispositifs électriques passifs habituels d'être réciproque, c'est-à-dire que si on renvoie par la droite du dispositif l'onde rectiligne qui en était sortie, les rôles des analyseurs et des polariseurs seront inversés et on retrouvera en sortie à gauche l'onde d'où on était parti. Autrement dit, le dispositif fait tourner en sens inverse le vecteur E_{R2} pour l'amener sur E_{R1} . On a imaginé récemment un dispositif appelé « gyrateur » qui n'est pas réciproque : quel que soit le sens dans lequel se présente une onde rectiligne le dispositif la fait tourner dans le même sens.

Si on a réalisé un « gyrateur » qui fait tourner de 45° le plan de polarisation (fig. 12), et si on le place entre deux guides rectangulaires à 45° l'un de l'autre, l'onde directe qui va du guide G_1 au guide G_2 arrivera sur G_2 avec la bonne polarisation pour être transmise, mais l'onde réfléchi qui vient de G_2 tournera encore de 45° dans le gyrateur, mais dans un sens tel qu'elle arrivera sur G_1 avec son champ électrique parallèle au grand côté; elle n'y pourra donc entrer. Le gyrateur pourra ainsi servir de dispositif « à sens unique » et permettra en particulier d'isoler un générateur de sa charge.

Le gyrateur est constitué par un bloc de ferrite convenablement taillé placé dans un guide circulaire et autour de ce guide se trouve un solénoïde qui produit un champ magnétique intense dans

(1) Bell Telephone Laboratories Red Bank N.J.

l'axe du guide. Si le courant qui passe dans l'enroulement est fixe, ou si le solénoïde est remplacé par un aimant permanent, ou si le bloc de ferrite est lui-même aimanté, on obtient un gyrateur tel que nous l'avons sommairement décrit. Mais en modulant le courant du solénoïde, on module par cela même l'onde hyperfréquence qui passe dans le gyrateur.

Du point de vue de la théorie des circuits, un gyrateur est un quadripôle passif qui n'est pas réciproque (fig. 10), c'est-à-dire que l'on a entre les courants qui entrent à gauche et à droite et les tensions les relations :

$$V_1 = -S i_2 \quad V_2 = S i_1$$

C'est TELLEGEN dans son numéro « *The gyrator, a new circuit element* » paru dans les *Philippis Research Reports Vol 3, Avril 1948 pp. 81-101*, qui

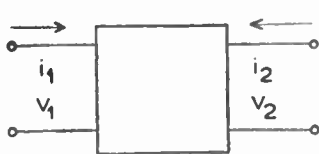
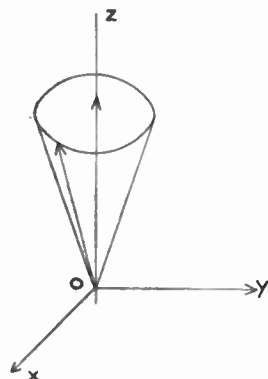


FIG. 10. — Quadripôle.

FIG. 11. — Cône de précession du moment magnétique d'un électron.



a montré comment cet élément permet la simplification de certains problèmes de circuits, ainsi que de nouvelles combinaisons intéressantes. Divers auteurs ont montré la possibilité de l'existence des gyrators, mais on n'a obtenu des résultats susceptibles d'applications pratiques qu'au moment où on a introduit les ferrites dans le domaine des hyperfréquences.

Le fonctionnement de ces éléments de circuit est basé sur la résonance magnétique des électrons dans les ferrites. Un électron peut pour cette théorie être considéré comme une sphère chargée négativement qui tourne sur elle-même. Elle est ainsi douée d'un moment angulaire, et cette rotation produit un moment magnétique colinéaire au moment angulaire qui dépend de la charge de l'électron, de son diamètre et de sa vitesse de rotation.

Si on applique à un bloc de ferrite, substance qui possède des électrons relativement libres un champ magnétique assez intense, tous ces moments magnétiques finissent après amortissement par être orientés dans le champ. Si on superpose au champ intense fixe dirigé, par exemple suivant Oz, un champ magnétique alternatif dirigé suivant un axe perpendiculaire Oy, on montre que, après amortissement, le moment magnétique finit par décrire un cône de révolution, dont l'ouverture dépendra de l'intensité et de la fréquence du champ alternatif, et il induira donc dans la direction Ox perpendiculaire au plan yOz un flux magnétique déphasé de 90° par rapport au champ alternatif.

A partir de là, on montre que si on place un bloc de ferrite dans un guide d'onde circulaire parcouru par une onde du mode H_{11} de champ magnétique $\vec{h}_1$, et situé dans un champ magnétique intense $\vec{H}$ dirigé suivant son axe, l'onde H_{11} peut être décomposée en deux ondes circulaires l'une « droite » l'autre « gauche » qui se propagent avec des vitesses différentes.

On appellera « droite » l'onde circulaire qui tourne dans le sens du courant du solénoïde qui produirait le champ $\vec{H}$, et « gauche » celle qui tourne en sens inverse. Après la traversée du bloc de ferrite, les deux ondes circulaires se recomposent suivant une onde $\vec{h}_2$ qui aura tourné d'un certain angle θ par rapport à $\vec{h}_1$. Si maintenant on changeait le sens de propagation et si on envoyait l'onde $\vec{h}_2$ à partir de la droite, les vitesses de propagation des ondes « droite » et « gauche » sont les mêmes, et l'onde résultante $\vec{h}_1$ aura tourné dans le même sens. Elle fera avec $\vec{h}_1$ l'angle 2θ . Le système n'est donc pas réciproque.

Nous venons de décrire ce qui se passe lorsque la fréquence de l'onde envoyée dans le guide est éloignée de la fréquence de résonance magnétique : les deux ondes « droite » et « gauche » sont atténuées assez peu et à peu près de la même quantité. Mais lorsque la fréquence se rapproche de la fréquence de résonance, l'onde « droite » est atténuée très fortement alors que l'onde « gauche »

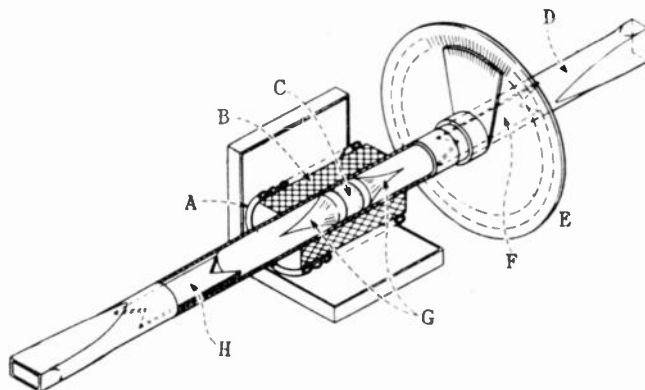


FIG. 12. — Schéma d'un gyrateur dans son montage d'essai (Figure extraite de l'article C.L. Hogan « The ferromagnetic Faraday effect at microwave frequencies and its applications. The microwave gyrator ». B.S.T.J. Vol. 31, n° 1.

a. Serpentin de refroidissement ; b. Solénoïde ; c. Ferrite ; d. Section tournante ; e. Disque repère fixe ; f. Lame radiale pour absorber les ondes polarisées verticalement ; g. Transitions progressives pour réduire les réflexions ; h. Lame radiale pour absorber les ondes polarisées horizontalement.

l'est très peu ; on obtient un dispositif qui produit une onde polarisée circulaire ; bien qu'il absorbe à peu près la moitié de l'énergie qui lui est appliquée, c'est un dispositif qui présente de l'intérêt.

Il est hors du cadre de cet exposé de faire la théorie complète des gyrateurs ; on trouve dans l'article de C. L. HOGAN « *The Ferromagnetic*

Faraday effect at Microwave Frequencies and its applications : The Microwave Gyrator » paru dans le B. S. T. J. de Janvier 1952, des renseignements détaillés et des calculs concernant le « gyrateur » et des combinaisons de gyrateurs que cet auteur appelle le *circulateur* et le *circulateur à polarisation*. Des symboles sont introduits pour ces nouveaux éléments de circuit.

En liaison avec les propriétés précédentes citons encore deux articles très clairs, plus physiques que techniques de Karl K. DARROW parus dans les B. S. T. J. de Janvier 1953 et Mars 1953 « *Magnetic resonance : Part. I — Nuclear Magnetic Resonance* » et « *Magnetic resonance; Part II — Magnetic Resonance of Electrons* » où l'auteur étudie le phénomène général de la résonance magnétique et traite en particulier dans son deuxième article le cas de la résonance ferromagnétique.

4. — Réseaux polarisants.

Lorsqu'on veut faire rayonner une antenne en polarisation circulaire, on peut se contenter de faire passer le faisceau polarisé rectiligne qui émerge de l'antenne à travers un réseau de lames parallèles formant guides d'ondes et inclinées à 45°

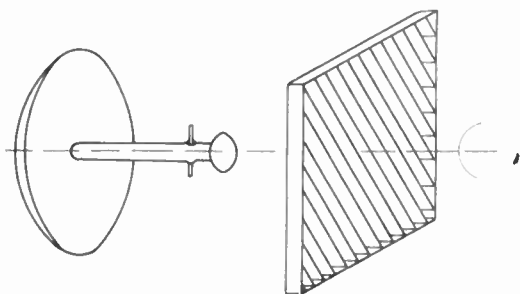


FIG. 13. — Réseau polarisant circulaire placé en avant d'une antenne

sur le plan de polarisation. Le champ rectiligne se décompose en deux champs égaux, l'un perpendiculaire aux lames et l'autre parallèle. Le premier passe pratiquement sans déformation et la largeur et l'espacement des lames parallèles sont tels que la différence entre la longueur d'onde guidée pour

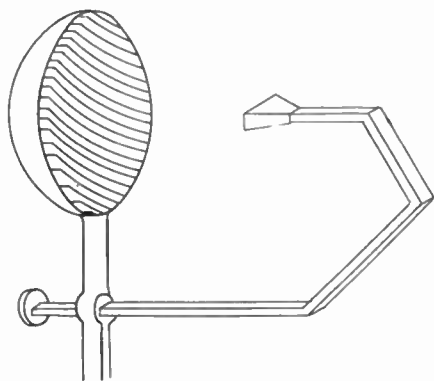


FIG. 14. — Réseau polarisant circulaire placé devant un réflecteur.

le second champ et la longueur d'onde dans l'air cause un déphasage de 90° entre ces deux champs. L'onde qui sort du réseau est donc polarisée circulaire dans l'axe du faisceau rayonné (fig. 13).

Au lieu de placer le réseau en avant de l'antenne, il est plus commode souvent, au point de vue mécanique de le placer contre le réflecteur de l'antenne. Le déphasage causé par le réseau de lames doit alors être la moitié du précédent car le faisceau traverse deux fois ce réseau (fig. 14).

On peut obtenir le même résultat avec un cornet qui émet en polarisation rectiligne et une lentille électromagnétique à lames à 45° de la polarisation (fig. 15).

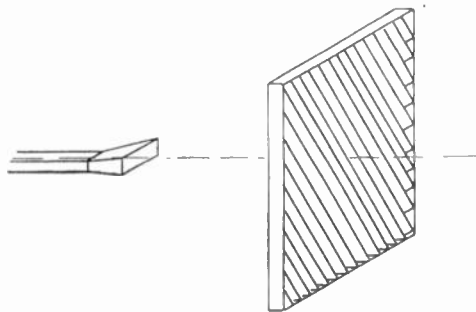


FIG. 15. — Réseau polarisant placé devant un cornet

L'article de J. F. RAMSAY déjà cité « *Circular polarization for C. W. Radar* » donne des détails sur ces dispositifs et des dispositifs analogues.

5. — Antennes en hélice.

Les dispositifs précédents ont l'avantage de pouvoir donner des faisceaux fins en polarisation circulaire, mais ce sont en général des dispositifs encombrants. Lorsqu'on recherche un faisceau relativement large, on peut utiliser une antenne de construction relativement assez facile, l'antenne en hélice. C'est une antenne qui est formée d'un conducteur enroulé en hélice. La fig. 16 indique

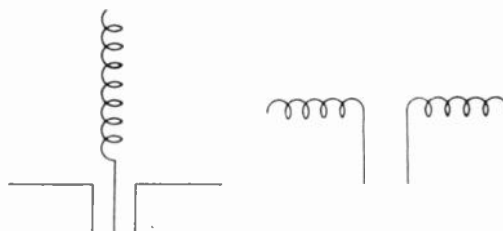


FIG. 16 a). FIG. 16 b).
Hélices alimentées ; a) en bout ; b) en son centre

deux modes d'excitation utilisés : dans le premier une extrémité de l'hélice est reliée au conducteur central d'un coaxial, alors que le conducteur extérieur du coaxial s'épanouit sous forme de plateau ; dans le second, l'hélice est attaquée en son milieu par une ligne bifilaire. Les paramètres qui interviennent dans une antenne en hélice sont le diamètre des spires et leur espacement, ainsi que le nombre de spires. On se rend compte que suivant les valeurs relatives des paramètres l'hélice peut se confondre pratiquement avec une boucle si le diamètre de la spire reste constant, alors que l'espacement devient nul ou bien vient se confondre avec une antenne rectiligne si c'est le diamètre de la spire qui devient nul

sans que l'espacement le devienne. Aussi l'hélice peut-elle rayonner de façons très différentes suivant les valeurs relatives de paramètres précédents et de la longueur d'onde. Sans entrer dans le détail de la recherche de tous les modes possibles de rayonnement, la fig. 17 indique les formes remarquables des modes de rayonnement qu'on peut appeler fondamentaux; en (a) l'hélice rayonne un faisceau unique dont le maximum est dirigé suivant l'axe de l'hélice : c'est le mode *axial*. En (b) le faisceau a un trou dans la direction de l'axe, et

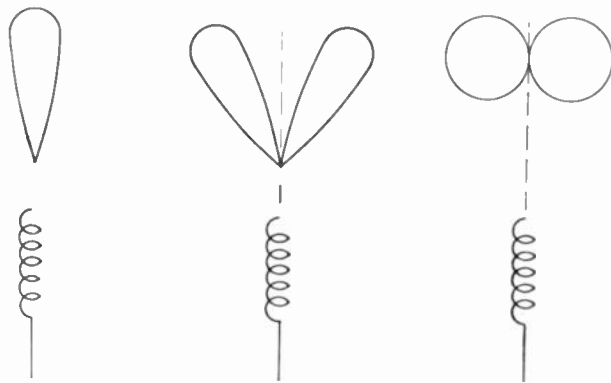


Fig. 17 a.

Fig. 17 b.

Fig. 17 c

Schémas des modes de rayonnement d'une hélice : a) mode axial, b) mode conique, c) mode normal

un maximum réparti sur un cône : c'est le mode *cônique*. En (c) le rayonnement est à peu près celui d'un dipôle, c'est-à-dire que le faisceau est omnidirectionnel avec un maximum dans le plan perpendiculaire à l'axe de l'hélice et un zéro dans la direction de l'axe : c'est ce qu'on appelle le mode *normal*. La polarisation émise par ce type d'antenne est en général elliptique. Dans le mode axial, elle arrive à être sensiblement circulaire dans la direction de l'axe si l'antenne a un assez grand nombre de spires.

Dans le mode *normal*, en réglant convenablement les divers paramètres, on arrive à avoir une antenne omnidirectionnelle autour de l'axe de l'hélice qui rayonne en polarisation à peu près circulaire dans toutes les directions.

C'est une antenne de ce dernier type que décrit H. A. WHEELER (1) dans son article « *A helical antenna for circular polarization* » paru dans les P. I. R. E. de Décembre 1947. Il fait une théorie du fonctionnement de cette antenne en considérant une spire hélicoïdale comme équivalente à une boucle plane située dans le plan de section droite du cylindre et un radiateur rectiligne parallèle à l'axe. En choisissant convenablement les paramètres, les radiateurs rectilignes présentent un déphasage de 90° avec les spires, et le rayonnement est polarisé circulaire. La théorie utilisée n'est valable avec une bonne approximation que pour les spires voisines du point d'attaque; mais elle convenait au problème de WHEELER qui était celui d'une antenne pour les ondes métriques; il pou-

vait avoir un espacement assez petit vis-à-vis de la longueur d'onde pour utiliser un assez grand nombre de spires. Citons ici l'avantage complémentaire des antennes en hélice; deux antennes en hélices, enroulées en sens inverses sont très découplées et peuvent être mises sans inconvénient assez près l'une de l'autre.

Le mode conique ne semble pas pour l'instant présenter beaucoup d'intérêt pour les applications et ne semble pas avoir été étudié en détail. Au contraire le mode *axial* présente un gros intérêt aux hyperfréquences. J. D. KRAUS dans son article « *Helical Beam Antennas* » paru dans Electronics d'Avril 1947 décrit une antenne en hélice pour la bande de 10 cm. Il donne là quelques diagrammes de rayonnement et quelques renseignements relatifs à la construction, mais ses calculs et ses mesures sont reportés avec beaucoup plus de détails dans deux autres articles : J. D. KRAUS et J. C. WILLIAMSON, dans le premier article « *Characteristics of Helical Antennas radiating in the axial mode* », paru dans le Journal of Applied Physics de Mai 1953, ont étudié théoriquement et expérimentalement les diagrammes de rayonnement d'une antenne en hélice dans le mode axial entre 250 et 500 Mc/s.

Une théorie montre qu'on peut considérer la distribution sur l'hélice comme la somme de quatre ondes (fig. 18) : une onde 1 atténuée exponentiellement et une onde 2 progressive qui se propagent en direction de l'extrémité libre de l'hélice, et deux ondes 3 et 4 en sens inverse respectivement atténuée exponentiellement et progressive.

Dans la région centrale de l'hélice, on ne trouve pratiquement que les ondes 2 et 4, et on peut définir un taux d'onde stationnaire.

$$= \frac{I_2 + I_4}{I_2 - I_4}$$

Par raison de symétrie, si une hélice d'un grand nombre de tours est parcourue par une onde progressive, elle rayonne, sur l'axe, de la polarisation circulaire. Si la polarisation n'est pas circulaire,

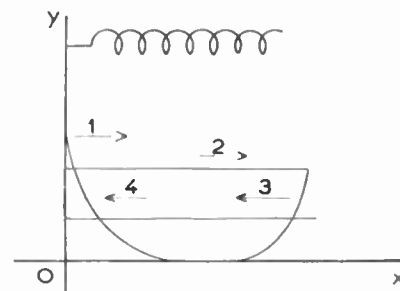


Fig. 18. — Les quatre ondes d'une hélice

c'est qu'il existe d'autres ondes. Dans l'expérimentation rapportée dans l'article précédent, les auteurs ont constaté un diagramme très irrégulier à l'extrémité inférieure de la gamme (250 Mc/s) et ce diagramme s'améliore progressivement quand la fréquence augmente; à 450 Mc/s les irrégularités sont très légères.

(1) Ingénieur-Conseil Physicien, Great Neck New-York.

Les résultats de cette expérimentation sont que, pour une hélice qui fonctionne dans le mode axial, on peut obtenir un bon faisceau dirigé suivant l'axe avec une largeur de bande assez grande; la largeur du faisceau est d'environ 45° à demi-puissance et la largeur de bande est d'environ 20 pour cent.

Le deuxième article auquel je faisais allusion est celui de O. J. GLASSER et J. D. KRAUS intitulé « *Measured Impedances of Helical Beam Antennas* » et paru dans le *Journal of Applied Physics* de Février 1948, où les auteurs étudient l'impédance d'entrée des mêmes hélices. Ils concluent qu'aux fréquences trop basses pour le fonctionnement en mode axial, l'impédance varie très rapidement avec la fréquence, mais que dans la gamme qui correspond au mode axial, l'impédance est relativement constante. Pour des hélices qui ont une longueur physique donnée, la variation d'impédance avec la fréquence décroît quand on augmente le pas de l'hélice. L'effet du nombre de spires est faible aux fréquences du mode axial dès que ce nombre dépasse un certain minimum. Pour les hélices étudiées, la résistance moyenne se situe aux environs de 130 ohms.

Les travaux de cette équipe ont été effectués à l'Université de l'Etat d'Ohio.

On trouve encore une antenne destinée à la polarisation circulaire pour ondes métriques et qui a une certaine parenté avec les hélices de WHEELER dans l'article de G. H. BROWN et O. M. WOODWARD Jr. paru dans la *R.C.A. Review* de Juin 1947. Cette antenne est formée de quatre dipôles placés en carré, mais inclinés. Des détails pratiques de construction sont indiqués, et les résultats expérimentaux sont donnés.

Les antennes en hélice sur ondes métriques dont nous venons de parler ont surtout été utilisées pour les liaisons de trafic avec les avions. Les antennes à réseau de lames sur ondes centimétriques permettent d'obtenir un faisceau fin et sont utilisées soit sur les radars pour diminuer l'influence des échos dus à la pluie, soit pour les câbles hertziens en vue de réduire les fluctuations provenant des échos du sol. Les antennes en hélice sur ondes centimétriques sont des antennes d'encombrement réduit et par suite de faisceaux larges; il semble que leur emploi soit plus particulièrement conseillé sur les avions et les fusées.

Nous n'avons trouvé dans la littérature aucune mention des *antennes diélectriques* pour la polarisation circulaire. Il semble qu'elles peuvent jouer le même rôle que les antennes en hélice aux fréquences centimétriques.

6. — Autres applications de la polarisation elliptique.

Sans entrer dans le détail, signalons que les services de la météorologie tant en France qu'à l'étranger ont cherché à utiliser les radars pour étudier les diverses formations nuageuses.

Signalons entre autres à ce sujet deux articles de N. R. LABRUM (1) « *The Scattering of Radio Waves by Meteorological Particles* » et « *Some Experiments on Centimeter-Wavelength Scattering by Small Obstacles* » parus dans le *Journal of Applied Physics* de Décembre 1952 où l'auteur étudie en particulier la polarisation des ondes diffractées par les particules atmosphériques ce qui doit permettre de distinguer dans les nuages les zones où les particules sont des gouttes de pluie, de la neige ou de grêle.

Mentionnons aussi dans cet ordre d'idées un article de J. BROU, A. VASSY et E. VASSY « *Echos Anormaux observés avec le radar du Musée Océanographique de Monaco* » paru dans la revue « *Geofisica pura et applicata* Volume 25 p. 71-82 (1953) » qui mentionne des échos obtenus sur atmosphère sans nuage apparent.

Cet article ne mentionne d'ailleurs pas l'usage de la polarisation circulaire.

Ces diverses applications ont donc conduit à étudier soigneusement la polarisation elliptique en général et à faire des mesures dans ce sens. D'autre part on a dû étendre aux ondes polarisées elliptiques les calculs mathématiques qui avaient été faits surtout en vue de la polarisation rectiligne.

7. — Les mesures sur la polarisation elliptique.

Les mesures sur les ondes polarisées elliptiques impliquent des mesures d'amplitude ou de puissance lorsqu'il s'agit de mesurer le rapport des axes de l'ellipse et des mesures de phase pour étudier le déphasage entre deux champs composants qui peuvent éventuellement donner une polarisation elliptique.

On trouve dans les P. I. R. E. de Mai 1951 un article de J. I. BOHNERT qui donne la description d'un appareillage pour effectuer la mesure du rapport des axes de l'ellipse et la mesure du gain. Les méthodes classiques d'étude des antennes s'appliquent aux antennes à polarisation elliptique à condition de prendre les précautions voulues. Citons pour ces méthodes les deux ouvrages de la collections du M. I. T. : C. G. MONTGOMERY, « *Technique of Microwave Measurements* », Chapitre 15. et S. SILVER, « *Microwave Antenna Theory and Design* », Chapitres 15 et 16. Malheureusement ces ouvrages ne tiennent guère compte de l'importance que peut prendre l'étude des antennes à polarisation elliptique.

On trouve des photographies très instructives de l'écran panoramique d'un radar sur 1 300 Mc/s en polarisation circulaire (fig. 19 et 20) dans un article de W. D. WHITE « *Circular Radar Cuts Rain Clutter* », paru dans la revue *Electronics* de Mars 1954. L'élimination des échos de pluie grâce à la polarisation circulaire, même à une longueur d'onde aussi basse que 25 cm est extrêmement nette. L'auteur donne les chiffres suivants : par

(1) Division of Radiophysics, Commonwealth Scientific and Industrial Research Organization, Sydney (Australie).

rapport à la polarisation rectiligne la polarisation circulaire donne une atténuation de 15 à 30 dB pour les échos de pluie, alors que l'écho d'un avion est atténué de 6 à 8 dB. Les meilleurs contrastes avion-pluie s'obtiennent quand l'avion est à un

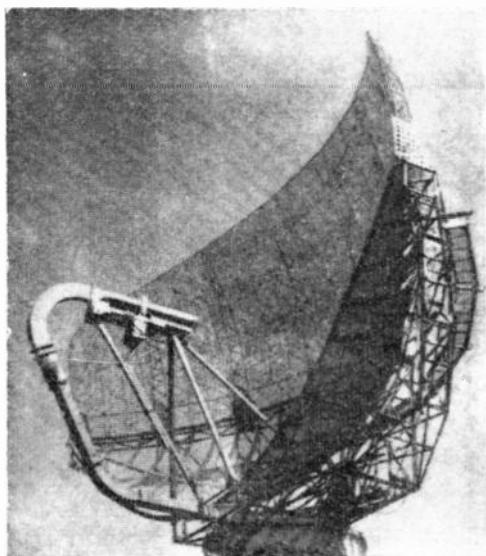


FIG. 19. Antenne en polarisation circulaire expérimentale pour 1300 Mc/s avec laquelle a été prise la photographie de la figure 20 (Figure extraite de l'article "Circular radar cuts rain clutter", Electronics, Mars 1954).

site élevé. Pour les faibles sites, l'amélioration est moindre, car une partie du faisceau réfléchi sur la pluie arrive au radar après avoir touché le sol, ce qui inverse encore la polarisation. On trouve encore dans le même article une photographie

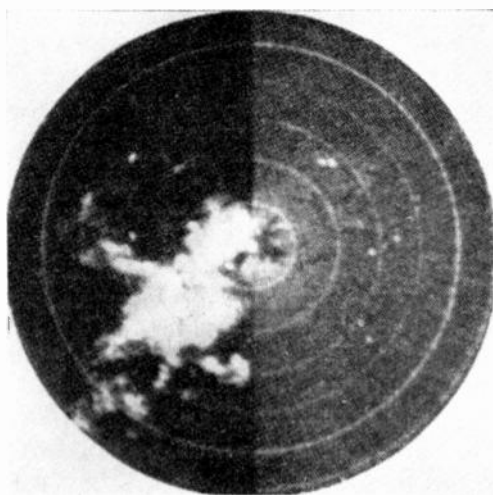


FIG. 20. — L'indicateur panoramique montre un gros écho de pluie (à gauche) et le résultat obtenu avec la polarisation circulaire (à droite) (Figure extraite de l'article "Circular radar cuts rain clutter", Electronics, Mars 1954).

intéressante qui montre l'effet combiné de l'élimination des échos fixes avec la polarisation circulaire destinée à l'élimination des échos de pluie.

8. — Théorie de la polarisation elliptique.

De nombreux articles ont paru sur la question de la polarisation elliptique et en particulier sur la question délicate de savoir quelle énergie peut recevoir d'une onde polarisée elliptique quelconque une antenne de réception apte à émettre une autre onde polarisée elliptique. Les calculs ne sont pas particulièrement compliqués, mais les antennes sont liées soit à un émetteur soit à un récepteur, et il est délicat de donner des formules générales qui s'appliquent à tous les cas. Aux fréquences radio, les antennes sont en général couplées directement à l'émetteur ou au récepteur et les quantités accessibles sont la tension aux bornes et le courant dans l'antenne. Aux hyperfréquences, l'antenne est généralement reliée au récepteur ou à l'émetteur par une ligne de transmission bifilaire, coaxiale ou en guide d'onde et les quantités accessibles sont la puissance émise ou reçue et les impédances soit de l'antenne, soit du récepteur. Et ces impédances sont connues au moyen de mesures de taux d'onde stationnaire et de phase.

On trouve dans l'ouvrage de S. A. SCHELKUNOFF « *Electromagnetic Waves* », plusieurs passages où il est parlé de la polarisation elliptique. Les formules générales dérivées des équations de Maxwell sont valables en polarisation elliptique, mais souvent les résultats n'ont pas été établis dans ce sens avec les détails qui seraient actuellement utiles. On trouve encore une courte étude de la réception des ondes polarisées elliptiques dans l'ouvrage de S. A. SCHELKUNOFF et H. T. FRIS « *Antennas Theory and Practice* » pp. 388-393.

Il semble qu'un des premiers articles qui aient essayé de donner une vue générale sur la question dans le but d'utiliser les rayonnements polarisés elliptiques soit constitué par le travail d'un groupe de quatre auteurs, paru dans les P. I. R. E. de Mai 1951 et intitulé « *Techniques for Handling Elliptically Polarized Waves with Special reference to antennas* ». L'introduction est de H. G. BOOKER (1). La première partie, de V. H. RUMSEY (2), a pour titre « *Transmission between elliptically polarized antennas* » et commence par une étude de la représentation des ondes elliptiques au moyen des diagrammes analogues aux diagrammes d'impédances. Elle donne des formules relatives à une liaison entre deux antennes très éloignées et ayant des polarisations quelconques. La deuxième partie est de G. A. DESCHAMPS (3) et son titre est « *Geometrical representation of the polarisation of a plane electromagnetic wave* ». Elle indique des relations entre les divers paramètres de l'onde elliptique et donne en particulier une formule pour la hauteur effective d'une antenne dans une direction donnée. La troisième partie intitulée « *Elliptically Polarized Waves and Antennas* » est de M. L. KALES (1). Partant d'une représentation de l'onde elliptique en vecteurs complexes, cet auteur établit une for-

(1) Cornell University Ithaca (New-York).

(2) Ohio State University, Columbus, Ohio.

(3) Federal Telecommunications Laboratories, Nutley (New-Jersey).

mule donnant la puissance reçue par une antenne quelconque en présence d'une onde elliptique. La quatrième partie est l'article déjà cité de J. I. BOHNERT (1) sur l'appareillage de mesure pour la polarisation elliptique.

Dans le même numéro des P.I.R.E. et à la suite du travail précédent, citons l'article de M. G. MORGAN (2) et W. R. EVAND (2) « *Synthesis and Analysis of Elliptic Polarization Loci in Terms of Space-Quadrature Sinusoidal Components* » qui étudie comment on peut décomposer une onde elliptique sur deux ou trois axes de coordonnées, et inversement la construire à partir de ses composantes.

Les travaux précédents ont eu pour précurseurs les chercheurs qui ont vu l'intérêt de l'étude des ondes elliptiques et qui ont déjà donné des formules utiles. Citons quatre articles par ordre chronologique. L'article de W. SICHAK et S. MILAZZO (3). « *Antennas for circular polarization* » dans les P. I. R. E. d'Août 1948, indique qu'une onde polarisée circulaire se réfléchit sur une surface métallique avec un sens de polarisation inversé et que cette onde n'est pas reçue par l'antenne qui l'a émise, et il décrit une antenne formée d'une boucle destinée à émettre et à recevoir de la polarisation circulaire. L'article de YUNG CHING YEH « *The receiving power of a receiving antenna and the criteria for its design* » paru aux P. I. R. E. de Février 1949 donne une formule de réception pour une onde polarisée elliptique quand on connaît la polarisation de l'onde incidente. Il s'intéresse surtout aux fréquences radio et parle de l'effet polarisant des réflexions sur l'ionosphère et du champ magnétique terrestre. L'article de G. SINCLAIR « *The transmission and reception of elliptically polarized waves* » paru aux P. I. R. E. de Février 1950 donne une formule qui permet de calculer la tension aux bornes d'une antenne de réception.

Enfin nous citerons l'article de E. ROUBINE (4) paru dans l'Onde Electrique de Juin 1950 et intitulé « *Les propriétés directives des antennes de réception* ». L'auteur essaie d'adapter la formule du théorème de réciprocité des antennes au cas où l'émission et la réception se font en polarisation elliptique et donne une relation entre la surface effective et le gain.

Nous n'exposerons pas ici plus à fond ces études théoriques, il n'est pas possible de les résumer et d'en faire une critique en quelques lignes. Disons seulement que les auteurs sont d'accord lorsqu'il s'agit de représenter une onde polarisée elliptique, et d'écrire à un facteur près la relation entre les ondes que peut émettre et recevoir une même antenne.

Mais ce facteur qui se calcule de façon assez simple sous certaines hypothèses que les auteurs font explicitement ou implicitement est d'un calcul plus compliqué si on veut le déterminer pour le cas général, où l'antenne est alimentée par une ligne de transmission terminée sur une charge non obligatoirement adaptée.

Conclusion.

Nous avons essayé non de passer en revue les résultats tant théoriques qu'expérimentaux relatifs à l'étude et à l'emploi de la polarisation elliptique ou circulaire, mais plutôt de donner une idée des problèmes qu'elle pose et des applications qu'on peut lui trouver et d'indiquer une série d'articles qui traitent de la question et auxquels les chercheurs désireux de se documenter plus à fond pourront se reporter. Ces articles comportent aussi en général chacun une bibliographie qui ne nous a pas toujours été accessible et à laquelle il serait sans doute intéressant de se reporter aussi.

Nous avons d'ailleurs été amenés à parler d'applications qui ne se rapportent peut-être pas obligatoirement à la polarisation elliptique, mais qui sont basées sur des travaux connexes : il s'agit en particulier des déphaseurs et des gyrateurs.

A l'intérieur des circuits hyperfréquences, les principales applications actuelles de la polarisation circulaire sont certains joints tournants et certains dispositifs de T. R. sur les détails desquels nous n'avons pas pu insister. Nous avons décrit des types d'antennes à polarisation circulaire soit à faisceaux fins, soit à faisceaux larges. Il semble que, dans la technique du radar, il serait intéressant d'étudier des antennes à balayage conique en polarisation circulaire qui auraient un double avantage : atténuer les échos de pluie ou du sol et peut-être opposer une plus grande difficulté au brouillage en ce sens qu'il est plus difficile de trouver un plan de référence dans le faisceau.

Nous avons donné quelques résultats concernant l'élimination des échos de pluie parus dans la littérature, mais nous pensons ultérieurement pouvoir exposer (1) certains résultats plus personnels, en particulier sur les antennes. Indiquons cependant qu'une conférence est prévue à la 5^e Section de la Société des Radioélectriciens au mois de Juin où M. PIRCHER, de la Compagnie Française Thomson-Houston, doit exposer certains résultats théoriques et expérimentaux qu'il a obtenus (2).

Il semble que la question de l'étude et des applications de la polarisation elliptique ou circulaire soit tout à fait d'actualité et qu'elle mérite qu'on consacre des efforts tant à la partie purement théorique qu'à l'expérimentation de ses possibilités et au développement des applications.

(1) Naval Research Laboratory, Washington.

(2) Thayer School of Engineering, Hanover (N.M.).

(3) Federal Telecommunication Laboratories.

(4) Maître de Conférences à la Faculté des Sciences de Lille.

(1) Depuis cette conférence, l'auteur a publié un article théorique qu'on trouvera à la bibliographie.

(2) Le texte de cette communication paraîtra dans un prochain numéro de l'Onde Electrique.

BIBLIOGRAPHIE

GÉNÉRALITÉS, DÉPHASEURS ET RÉSEAUX POLARISANTS.

- Very High Frequency Techniques, Vol. I, 1^{re} édition pp. 153-161.
- FOX G.A. — An Adjustable waveguide phase changer. *P.I.R.E.* Vol. 35, pp. 1685-1698, Décembre 1947.
- RAMSAY J.F. — Circular polarization for C.W. Radar. *Marconi Review*, 2^e quart. 1952, pp. 71-72.
- RIDENOUR L.N. — Radar System Engineering. Vol. I de la Collection du MIT, 1^{re} édition, pp. 84-85 (Mac Graw-Hill).

RÉSONANCE MAGNÉTIQUE ET GYRATEURS.

- KASTLER A. — La polarimétrie dans le domaine des ondes hertziennes. *Supplemento al Nuovo Cimento*. Vol. IX, Série IX, n° 3, pp. 315-321.
- HOGAN C.L. — The Ferromagnetic Faraday Effect at Microwave frequencies and its applications. The Microwave gyrator. *B.S.T.J.* Vol. 31, n° 1, Janvier 1952, pp. 1-31.
- DARROW K.K. — Magnetic Resonance I Nuclear Magnetic Resonance, II Magnetic resonance of Electrons. *B.S.T.J.* Vol. 32, n° 1, Janvier 1953, pp. 74-99 et n° 3, Mars 1953, pp. 384-405.
- SURDUTS A. — L'effet Faraday dans les conducteurs et semi-conducteurs. *C.R. Ac. Sci. Fr.*, 9 Mars 1953, 1953, 236 n° 10, pp. 1005-1007.
- BELJERS H.G. et SNOCK J.L. — Gyromagnetic Phenomena occurring within Ferrites. *Phys. Rev.*, 11 May 1950, pp. 313-322.
- YAGER W.A., GALT J.K., MERRITT F.R., WOOD E.A. — Ferromagnetic resonance in nickel ferrite. *Phys. Rev.* 80 pp. 744-748 (1950).
- TELLEGEN B.D.H. — The gyrator, a new electric element. *Philippis Research Reports*. Vol. 3, Avril 1948, pp. 81-101.
- BLOCH F. — Nuclear induction. *Phys. Review*, Vol. 70, n° 7 et 8, 1^{er} et 15 Oct. 1946, pp. 460-474.
- BLOCH F., HANSEN W.W. et PACKARD M. — The Nuclear induction experiment. *Phys Review*, Vol. 70, n° 7 et 8, 1^{er} et 15 Oct. 1946, pp. 474-485.
- GOLDSTEIN L., LAMBERT M et HENEY J. — Magneto-optics of an electron gas with guided microwaves. *Phys. Rev.*, Vol. 82, n° 6, 15 Juin 1951, pp. 956-57.
- WICHER E.R. — The influence of Magnetic fields upon the propagation of electro-magnetic waves in artificial dielectrics. *J. of Applied physics*. Vol. 22, n° 11 Novembre 1951, pp. 1327-1329.

ANTENNES EN HÉLICE.

- KRAUS J.D. — Helical Beam Antenna for circular Polarization. *Electronics*, Avril 1947, pp. 109-111.
- KRAUS J.D. et WILLIAMSON J.C. — Characteristics of helical antennas radiating in the axial mode. *J. of Applied Physics*, Vol. 19, Janvier 1948 pp. 87-96.
- GLASSER O.J. et KRAUS J.D. — Measured Impedances of helical beam antennas. *J. of Applied Physics*, Vol. 19, Février 1948, pp. 193-197.
- CHATTERJEE J.S. — Radiation of a conical helix. *J. of Applied Physics*. Vol. 24, n° 5, Mai 1953, pp. 550-559.
- KRAUS J.D. — Helical Beam Antennas for wide band applications. *P.I.R.E.* Vol. 46, n° 10, Oct. 1948, pp. 1236-1242.
- WHEELER H.A. — A helical Antenna for Circular Polarization. *P.I.R.E.* Vol. 35, n° 12, Décembre 1947, pp. 1484-1488.
- BROWN G.H. et WOODWARD O.M. Jr. — Circularly polarized omnidirectional antenna. *R.C.A. Review*, Vol. 8, n° 2, Juin 1947, pp. 259-269.

MÉTÉOROLOGIE.

- LABRUM N.R. — Some experiments on Centimeter Wavelength scattering by small obstacles. *J. of Applied Physics*. Vol. 23, n° 12 Décembre 1952, pp. 1320-1323.
- LABRUM N.R. — Scattering radio waves by meteorological Particles. *J. of Applied Phys.* Vol. 23, n° 12, Décembre 1952, pp. 1324-1330.
- BOWEN E.G. — Radar observations of rain and their relations to the mechanism of rain formation. *J. of Atmospheric and Terrestrial Physics*, 1951, 1, p. 125.
- HOOPER N.E.N. et KIPPAX A. — Radar echoes from meteorological precipitations. *P.I.R.E.* 1950, 97, Part. I, p. 89.
- HOOD A.D. — Quantitative measurements at three and ten centimeters of radar echoes intensities from precipitation. *National Research Council of Canada Report E.R.A.*, 180.
- RYDE J.W. — Attenuation and Radar echoes produced at centimeter wavelengths by various meteorological phenomena, publié dans *Meteorological factors in radiowaves propagation*. Physical Society and Royal Meteorological Society, London, 1948.
- KIELY D.G. — Rain clutter Measurements with C.W. Radar Systems operating in the 8 mm wavelength band. *P.I.R.E.* Vol. 101 Part. 111, n° 70, Mars 1954, pp. 101-108.
- BROC J., VASSY A et VASSY E. — Échos anormaux observés avec le radar du musée océanographique de Monaco.

MESURES ET THÉORIE.

- MONTGOMMERY C.G. — Microwave measurements. Vol. II de la collection du MIT. 1^{re} édition Ch. 15, Mc Graw Hill. Edition française Chiron.
- SILVER S. — Microwave antennas theory and design. Vol. 12 de la collection du MIT. 1^{re} édition. Ch. 15, Mc Graw Hill.
- WHITE W.D. — Circular radar cuts Rain clutter. *Electronics*, Mars 1954.
- SCHIELKUNOFF S.A. — Electromagnetic Waves. Van Nostrand, New-York.
- YUNG CHING YEH. — The receiving power of a receiving antenna and the criteria for its design. *P.I.R.E.* Vol. 37, Février 1949, pp. 145-58.
- SINCLAIR G. — The Transmission and reception of elliptically polarized waves. *P.I.R.E.* Vol. 38, pp. 148-151, Février 1950.
- SICHAKE W. et MILAZZO S. — Antennas for Circular polarization. *P.I.R.E.* Vol. 36, pp. 997-1002, Août 1948.
- ROUBINE E. — Les propriétés directives des antennes de réception. *L'Onde Electrique*. n° 279, Juin 1950, pp. 259-266.
- BOOKER N.G., RUMSAY V.H., DESCHAMPS G.A., KALES M.L., BOHNERT J.L. — Technics for handling elliptically polarized waves with special reference to antennas. Introduction, I Transmission between elliptically polarized antennas, II Geometrical representation of the polarization of a plane electromagnetic wave, III Elliptically polarized waves and antennas, IV Measurements on elliptically polarized antennas. *P.I.R.E.* Vol. 39, n° 5, Mai 1951, pp. 533-552.
- MORGAN M.G. et EVANS W.R. — Synthesis and Analysis of elliptic polarization loci in terms of space quadrature components. *P.I.R.E.* Vol. 39, n° 5, Mai 1951, pp. 552-556.
- ROUBINE E. — Antenne de réception. Division radioélectrique et électronique. Ecole Supérieure d'Electricité.
- ROUBINE E. — Cours de lignes HF et d'antennes. Ecole Supérieure d'Electricité. 1948.
- DESCHAMPS G.A. — New Chart for the solution of transmission-livre and polarization problems. *Electrical Communication*, Vol. 30, N° 3, Sept 1953, pp. 247-254 et aussi *Trans I.R.E.*, Prof. group on microwave theory and techniques. Vol. 1, mars 1953, pp. 5-13.
- DESCHAMPS G.A. — Geometrical Representation of the polarisation of a plane electromagnetic wave. *Proc. I.R.E.*, Vol. 39, mai 1951, pp. 540-544.
- BOUX M. — La polarisation elliptique du rayonnement électromagnétique. *Ann. des Télécommunications*, T. 9, N° 1, oct-nov.-déc. 1954.

UTILISATION DE LA MESURE DE PHASES DANS UN SYSTÈME DE GUIDAGE ET DE LOCALISATION DES MOBILES

PAR

J. ZAKHEIM

Ingénieur E.S.E. et I.E.T.

Chef de Groupe de Recherches à l'O.N.E.R.A.

I. — INTRODUCTION.

Le dispositif connu sous le nom « Radio-alignement ONERA » est destiné essentiellement à la détermination de la direction dans laquelle se trouve un mobile. Le fonctionnement est basé sur la mesure des variations des phases des ondes émises par un émetteur porté par le mobile et reçues au sol par un ensemble de récepteurs convenablement disposés. Le matériel embarqué est constitué par un simple émetteur d'ondes entretenues de faible puissance. Son poids très réduit autorise son emploi même sur des engins de très faibles dimensions. L'appareillage au sol peut être indifféremment utilisé soit pour l'indication au pilote du mobile de son orientation par rapport à une direction de référence, soit, aussi bien, au maintien automatique d'un engin dans un plan de tir déterminé.

II. — PRINCIPE.

Soit $XOYZ$ un trièdre de référence (fig. 1), M la position du mobile, A et B deux récepteurs placés sur l'axe des X symétriquement par rapport à O . Désignons par ψ l'angle formé par la droite OM et l'axe des X .

Si la distance OM est grande par rapport à la distance AB , le déphasage des ondes arrivant aux récepteurs A et B peut s'écrire sous la forme :

$$1) \quad \varphi = 2\pi \cdot \frac{AB}{\lambda} \cdot \cos \psi ;$$

avec λ la longueur d'onde émise par le mobile M .

La relation entre le déphasage φ et l'angle ψ

fait intervenir le rapport $2\pi \cdot \overline{AB}/\lambda$ appelé facteur de sensibilité du système.

Si on a $AB < \lambda/2$ le facteur de sensibilité est faible et au maximum égal à π . Par contre, et cela est essentiel, la mesure de φ est liée sans ambiguïté à la direction ψ car φ ne sort jamais des limites $\pm \pi$.

En prenant une distance $\overline{AB}$ grande par rapport à la longueur d'onde le facteur de sensibilité augmente. La précision du système augmente aussi et peut, comme l'a d'ailleurs démontré l'expérience,

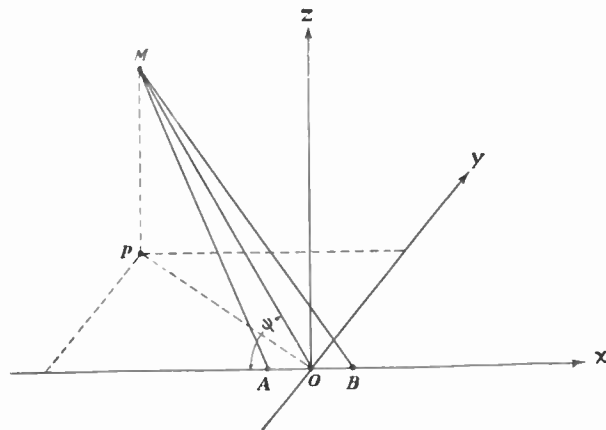


FIG. 1

dépasser la valeur pratiquement nécessaire dans le cas d'un guidage. Mais on se heurte alors à une difficulté dans la détermination de la valeur exacte de ψ . En effet si $\overline{AB} \gg \lambda$ on a une série des valeurs des angles ψ définis par :

$$2) \quad \psi_n = \arccos \left[\frac{\varphi + n \cdot 2\pi}{2\pi \overline{AB}/\lambda} \right] ;$$

qui pour des valeurs de n : $n_1 = 0$; $n_2 = 1$; $n_3 = 2$ etc. donnent des valeurs de déphasage φ_n :

$$\varphi_1 = \varphi ; \varphi_2 = \varphi + 2\pi ; \varphi_3 = \varphi + 4\pi ; \text{etc.}$$

entre lesquelles le phasemètre est absolument incapable d'opérer une distinction, son organe mobile reprenant la même position angulaire φ pour toutes ses valeurs de ψ_n .

Le Radio-alignement combine l'utilisation de deux bases de mesure. Une base, dite « base longue », correspond à une distance $\overline{AB} \gg \lambda$. L'autre base, dite « base courte » a ses antennes écartées seulement de $\lambda/2$. Elle sert à lever l'ambiguïté et à verrouiller les indications du phasemètre de la « base longue » sur une valeur correcte du nombre des tours complets.

Soulignons immédiatement les particularités suivantes du système.

a) Indépendance de la continuité de l'émission :

Une interruption momentanée de l'émission de l'émetteur mobile (par masquage d'antenne par exemple) est sans effet sur le fonctionnement du Radio-Alignement, celui-ci redonne des indications correctes dès la réapparition de l'émission.

b) Nature des angles mesurés :

L'angle mesuré est, par principe, situé dans le plan oblique AMB . Cette caractéristique permet, comme nous le verrons plus loin, de déterminer, à l'aide d'un nombre convenable des radio alignements, les trois coordonnées du mobile.

La mesure des différences des phases en haute-fréquence étant délicate, on tourne la difficulté en faisant battre les ondes arrivant aux récepteurs A et B avec un émetteur hétérodyne commun. On obtient ainsi à la sortie des récepteurs des signaux basse-fréquence reproduisant fidèlement les variations des déphasages haute-fréquence. La mesure proprement dite est effectuée à l'aide d'un phasemètre asservi réalisé de la manière suivante :

Les signaux basse fréquence provenant des récepteurs A et B (fig. 2) sont envoyés dans un discriminateur des phases D suivi d'un amplificateur

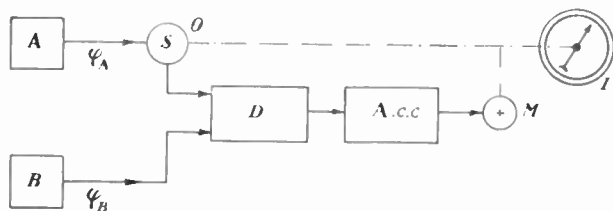


FIG. 2

de courant continu A_{cc} alimentant le moteur M . Le signal du récepteur A passe, avant d'atteindre le discriminateur, par un déphaseur constitué par un selsyn S entraîné par le moteur M . L'équilibre du

système correspondant à un déphasage de $\pi/2$ entre les tensions alimentant le discriminateur, si 0 désigne le déphasage introduit par le selsyn, nous aurons la relation :

$$\varphi_A + 0 - \varphi_B = \pi/2 ;$$

d'où

$$3) \quad (\varphi_A - \varphi_B) = \frac{\pi}{2} - 0 ;$$

Par conséquent, à chaque variation du déphasage $(\varphi_A - \varphi_B)$ correspondra une modification équivalente dans la position angulaire 0 du selsyn, position indiquée par l'index I qui lui est accouplé.

Les deux phasemètres, affectés respectivement à la « base longue » et à la « base courte » ainsi que les différents potentiomètres qu'ils entraînent dans leur rotation, sont dans notre réalisation assemblés dans un appareil unique représenté sur la fig. 3.

Nous allons passer en revue les différents résultats que l'on peut obtenir avec ce dispositif suivant le nombre d'ensembles « radio-alignement » utilisés simultanément.

III. — RADIO-ALIGNEMENT SIMPLE.

Le radio-alignement simple est essentiellement destiné à un guidage d'un mobile dans un plan vertical. En effet le plan YOZ (fig. 1) est caractérisé par un déphasage nul entre les signaux reçus aux deux récepteurs A et B et ceci indépendamment du rapport $\frac{MP}{OM}$. Le phasemètre mesurant l'angle φ

est accouplé à un dispositif produisant des signaux de guidage. Ces signaux, fonction du sens et de la grandeur de l'écart du mobile à partir du plan, peuvent être :

1° soit utilisés, après une transmission par une V.H.P. classique, pour indiquer au pilote du mobile sous une forme auditive ou visuelle sa position par rapport au plan de référence.

2° soit injectés directement dans la chaîne de télécommande du mobile, l'asservissant ainsi au plan de référence.

En faisant passer le plan YOZ par une piste d'atterrissage ou par l'axe d'une base de mesure, il devient possible de prendre un avion en charge à une distance considérable et de l'amener exactement sur l'axe désiré sans que les conditions de visibilité entrent en ligne de compte.

Dans le cas des signaux auditifs, il est avantageux d'utiliser un système auquel les pilotes sont déjà habitués. C'est ainsi que nous utilisons des signaux composés des lettres N (trait-point) et A (point-trait) s'imbriquant les uns dans les autres. L'intensité avec laquelle chaque lettre est émise dépend de la position du mobile par rapport à l'alignement. Tant que le mobile se trouve dans le plan de guidage les deux lettres sont émises avec la même intensité,

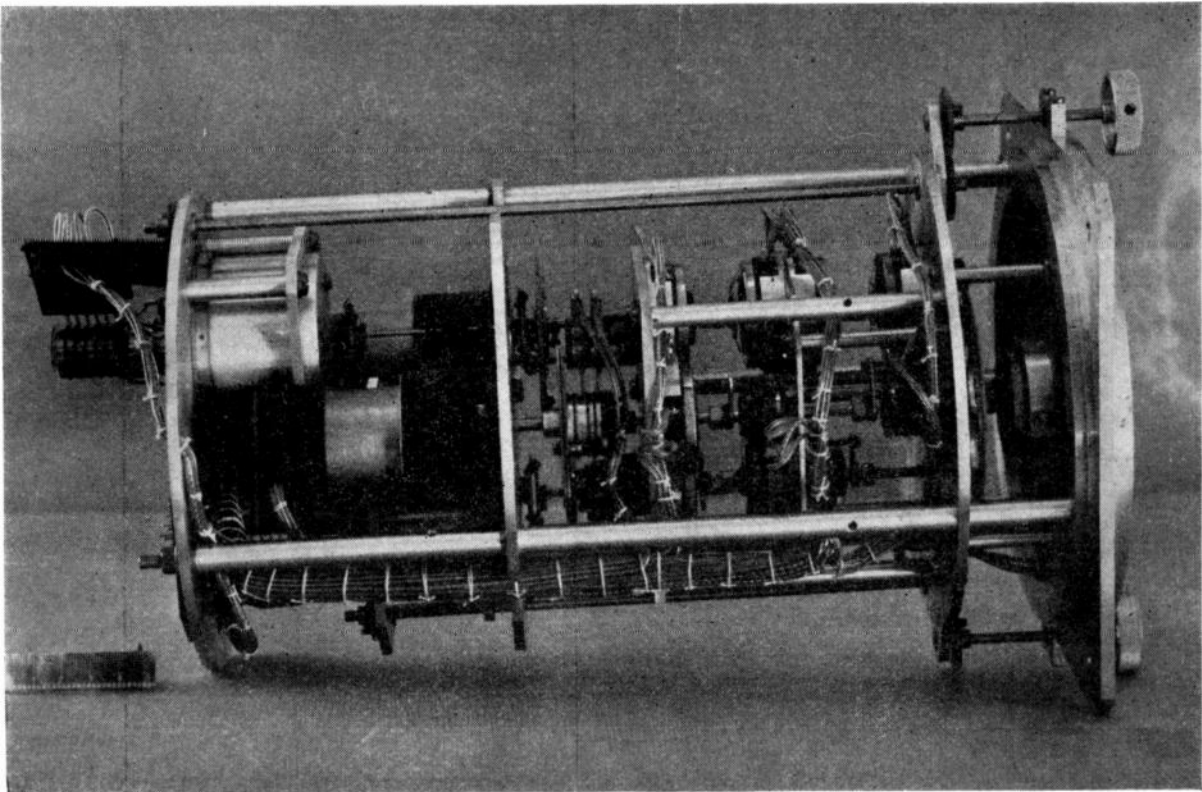


Fig. 3

le pilote entend donc un son continu. De part et d'autre du plan l'une ou l'autre des lettres s'estompe progressivement. A partir d'un certain angle d'écart une seule lettre est émise, l'autre étant bloquée.

Les signaux de guidage sont obtenus de la manière suivante :

Sur un tambour en matière isolante, dont le déve-

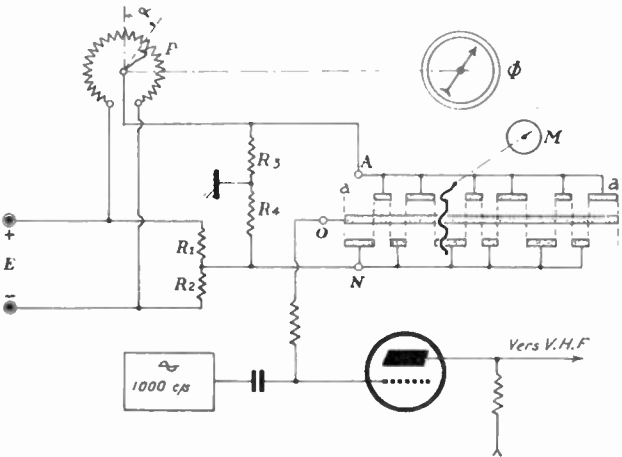


Fig. 4

loppement est indiqué sur la fig. (4), sont disposées deux rangées de plots courts et longs convenablement groupés. Une série des plots aboutit à la borne A, l'autre série à la borne N. Un balai triple, dont le mouvement est assuré par le moteur M, établit le contact entre la barre collectrice 0 et les dif-

férents plots. La barre 0 est reliée à la grille d'un tube L à pente variable.

Les deux séries des plots sont alimentées à partir d'un pont constitué par les résistance R_1 et R_2 et le potentiomètre P calé sur l'arbre du phasemètre.

Les tensions apparaissant aux bornes A et N. sont ainsi rendues variables en fonction de l'indication du phasemètre Φ , quand la tension sur l'une des bornes monte, la tension sur l'autre borne descend et vice versa. La polarisation du tube L se trouve donc modulée à travers le dispositif décrit par les tensions en A et N, et ceci alternativement à la cadence de la manipulation des lettres A et N, enchevêtrées entre elles.

Le tube L étant intercalé entre un générateur à 1 000 p/s et le départ vers les étages de modulation d'un émetteur VHF, on assure ainsi l'envoi vers le pilote du mobile des signaux de correction de route.

La largeur du faisceau de guidage, c'est-à-dire l'écart angulaire du mobile par rapport au plan de guidage provoquant le passage de l'émission de la lettre A seule à l'émission de la lettre N seule est donnée par la relation :

$$\theta b = \frac{2\theta_0}{\pi} \cdot \frac{\lambda}{D} \cdot \frac{V_{gb}}{E} :$$

avec :

λ/D = rapport de la longueur d'onde à la longueur de la base,

θ_0 = l'angle couvert par l'enroulement du potentiomètre P ,

E = la tension de la source alimentant le pont,

V_{gb} = la polarisation du cut-off du tube L .

Dans notre réalisation l'angle θ_0 est variable entre les valeurs 1 et 3 degrés.

En comparaison avec les beams de guidage usuels nous signalerons les points suivants :

1° Le faisceau de guidage du radio alignement est plusieurs fois plus étroit que les beams classiques, le guidage gagne ainsi énormément en précision.

2° Le radio alignement ONERA peut, tout en conservant ses antennes fixes, décaler son orientation d'environ $\pm 75^\circ$ par rapport à l'orientation principale. Cette particularité permet de prendre en charge un avion se trouvant pratiquement dans n'importe quel azimut et de le « tirer » vers le terrain.

3° Dans les beams classiques le sol ignore la position de l'avion par rapport au plan de guidage.

Par contre notre dispositif renseigne continuellement les opérateurs au sol sur la direction dans laquelle se trouve le mobile.

La transmission des signaux de guidage vers l'avion se fait à l'aide de la VHF ordinaire faisant partie de l'équipement de bord. Dans ces conditions, l'équipement supplémentaire se compose uniquement

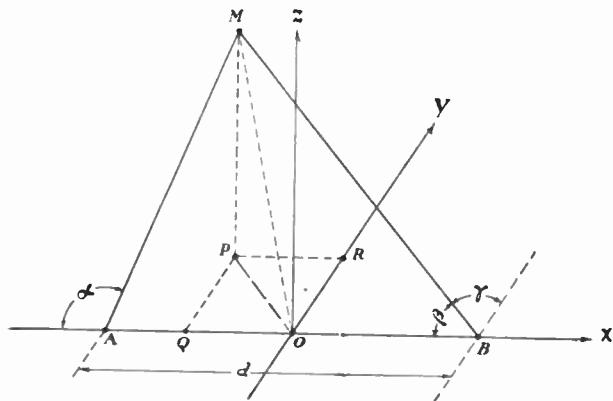


FIG. 5

de l'émetteur de localisation, dont le poids, alimentations comprises, ne dépasse pas 2 kg.

Les portées réalisées avec un équipement de ce genre dépassent 80 kilomètres à une altitude de vol de 2 000 mètres.

VI. — RADIO-ALIGNEMENT DOUBLE.

L'utilisation simultanée de deux dispositifs radio-alignement implantés suivant la fig. 5 aux points A et B distants d'une quantité connue d permet de connaître à chaque instant l'abscisse du mobile. L'ensemble se prête donc très bien à la mesure de vitesses sur bases en vol horizontal.

Dans l'hypothèse que la force et l'orientation du vent restent constantes pendant la durée de l'expérience, l'appareillage peut fournir aussi bien les vitesses aérodynamiques moyennes que les vitesses instantanées.

Citons également, comme autre exemple d'utilisation d'un radio-alignement double; les larguers des maquettes à partir d'un avion porteur évoluant à haute altitude. Dans ce cas on utilise généralement

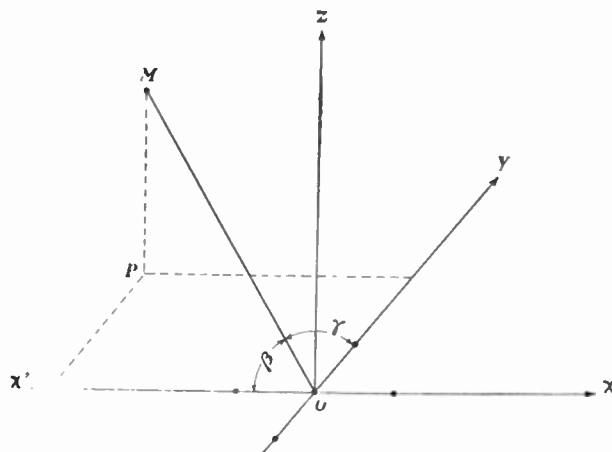


FIG. 6

un ensemble des cinéthéodolites pour le relevé de la trajectoire de la maquette. Il est nécessaire, pour la bonne précision des mesures, que l'avion se présente suivant un axe prescrit et que le larguer soit effectué à la verticale d'un point bien déterminé.

Le radio-alignement double résout le problème. Les deux ensembles sont groupés suivant la fig. 6. On mesure essentiellement les angles :

$$\beta = \widehat{MOX'} \text{ et } \gamma = \widehat{MOY},$$

ce qui détermine l'orientation de OM . L'axe des X étant orienté suivant l'axe du larguer, les indications du phasemètre mesurant γ sont utilisées pour le guidage de l'avion de manière à le maintenir toujours dans le plan XOZ . Si d'autre part, on connaît avec assez de précision l'altitude du vol, la mesure de l'angle β donne immédiatement l'abscisse de l'avion. On connaît ainsi à chaque instant la position de l'avion et on peut commander le larguer à la verticale d'un point bien déterminé. En particulier en faisant coïncider la verticale du larguer avec l'axe OZ , la connaissance exacte de l'altitude de l'avion devient superflue.

On dispose actuellement à l'ONERA d'un double radio-alignement monté dans un véhicule tous terrains et absolument autonome. La fig. 7 donne une vue intérieure de cette installation.

V. — RADIO-ALIGNEMENT TRIPLE.

L'exploitation simultanée de trois radio-alignements convenablement groupés permet de déter-

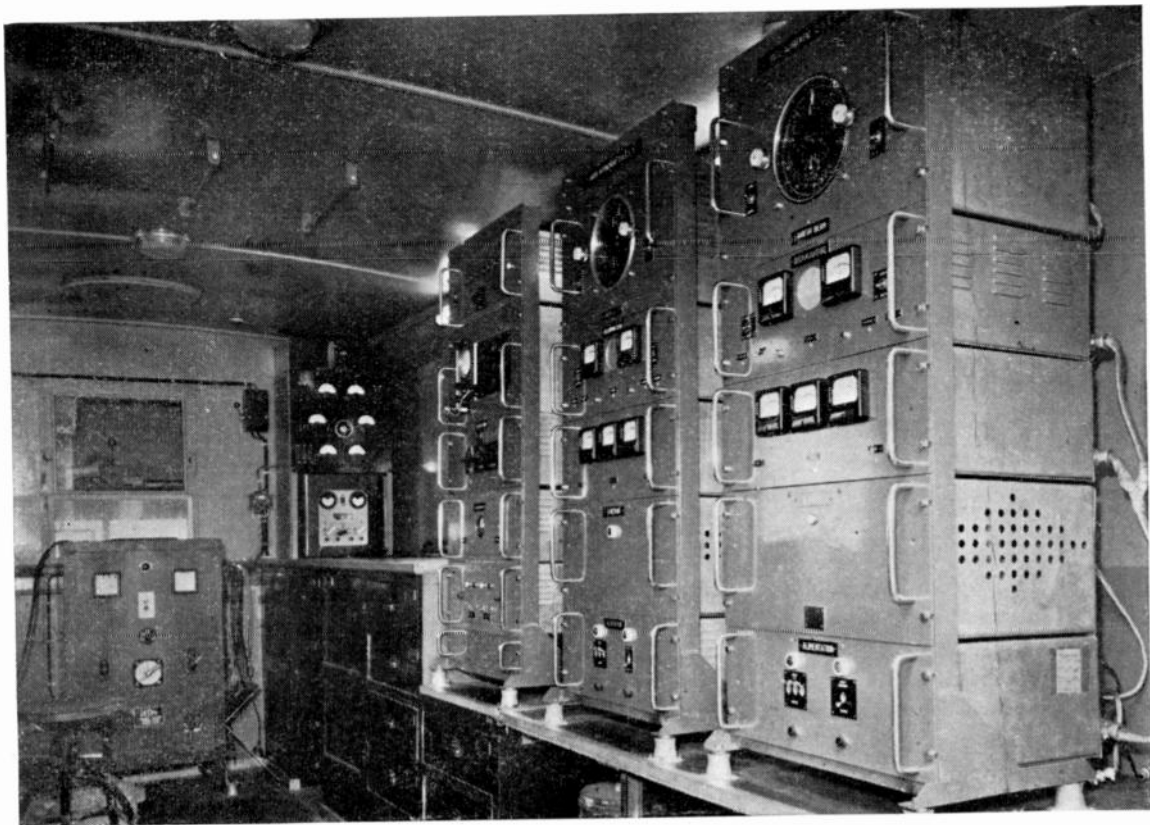


Fig. 7

miner sans ambiguïté les coordonnées d'un mobile donc également de relever une trajectoire.

En effet, si au point *B* (voir fig. 5) on place un ensemble double mesurant les angles β et γ , nous disposons des éléments d, α, β, γ suffisants pour la détermination des coordonnées *X, Y* et *Z* du point *M* où se trouve le mobile.

Il est évident que les deux points *A* et *B* doivent être reliés à un point central de mesure soit par fil soit par radio. Cette liaison devra permettre une transmission continue et précise des valeurs des angles α, β et γ qui seront soit enregistrées en vue d'un dépouillement ultérieur, soit injectées directement dans un ordinateur fournissant les valeurs des coordonnées.

MULTIPLIEURS A DÉCOUPAGE TEMPOREL

PAR

Madame M. LILAMAND

Ingénieur à la Société d'Electronique
et d'Automatisme

1. — INTRODUCTION.

Dans les machines à calculer analogiques, la première forme d'opération non linéaire est la multiplication. La S. E. A. a retenu trois types de multiplieurs : le multiplieur à servomécanisme, le multiplieur à diodes mettant en œuvre un traducteur de fonction parabolique, et enfin le multiplieur à modulation ou à découpage temporel dont il va être question. Celui-ci a l'intérêt d'être plus précis que le multiplieur à servomécanisme, d'être distributif, et de nécessiter beaucoup moins de matériel que le multiplieur à diodes.

L'idée de base consiste dans le produit d'une modulation de durée et d'une modulation d'amplitude sous la forme proposée en particulier par GOLDBERG :

L'amplitude d'un signal rectangulaire est fonction linéaire d'un des facteurs du produit, la durée dans le temps étant fonction de l'autre ; de telle sorte que la valeur moyenne dans le temps du signal soit proportionnelle au produit.

Les multiplieurs basés sur ce principe ⁽¹⁾ avaient jusque là une précision insuffisante ou une complexité trop grande. Nous avons recherché à la S. E. A. une solution à la fois économique et précise à laquelle nous sommes arrivés, par la mise au point d'un commutateur de précision ⁽²⁾ et par la compensation des capacités parasites des lampes.

Nous avons pu ainsi obtenir une précision de 0,1 % pour une fréquence moyenne de découpage de 1,5 kHz.

2. — PRINCIPE.

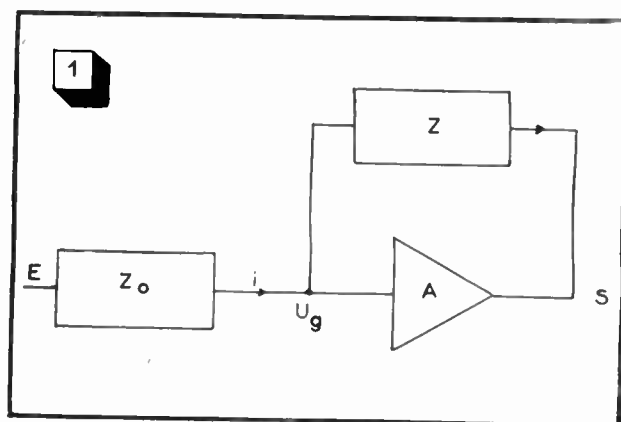
Rappelons d'abord ce qu'on entend par déphaseur et intégrateur en analogie.

(1) Voir bibliographie.

(2) Brevets P. V. 658 091, du 13 novembre 1953 : « Perfectionnements à la commutation électronique » (François-Henri RAYMOND),

P. V. 659.013, du 28 novembre 1953 : « Perfectionnements apportés aux commutateurs électroniques », (François-Henri RAYMOND Monique Berthe LEJET).

Si nous nous reportons à la figure 1, A est un amplificateur de gain infini et d'admittance nulle, u_g est la tension de grille de l'amplificateur maintenue nulle par la contre-réaction.



Ecrivons que le courant dans l'impédance d'entrée Z_0 est égal au courant dans l'impédance de bouclage Z_1

$$\frac{E}{Z_0} = - \frac{S}{Z_1}$$

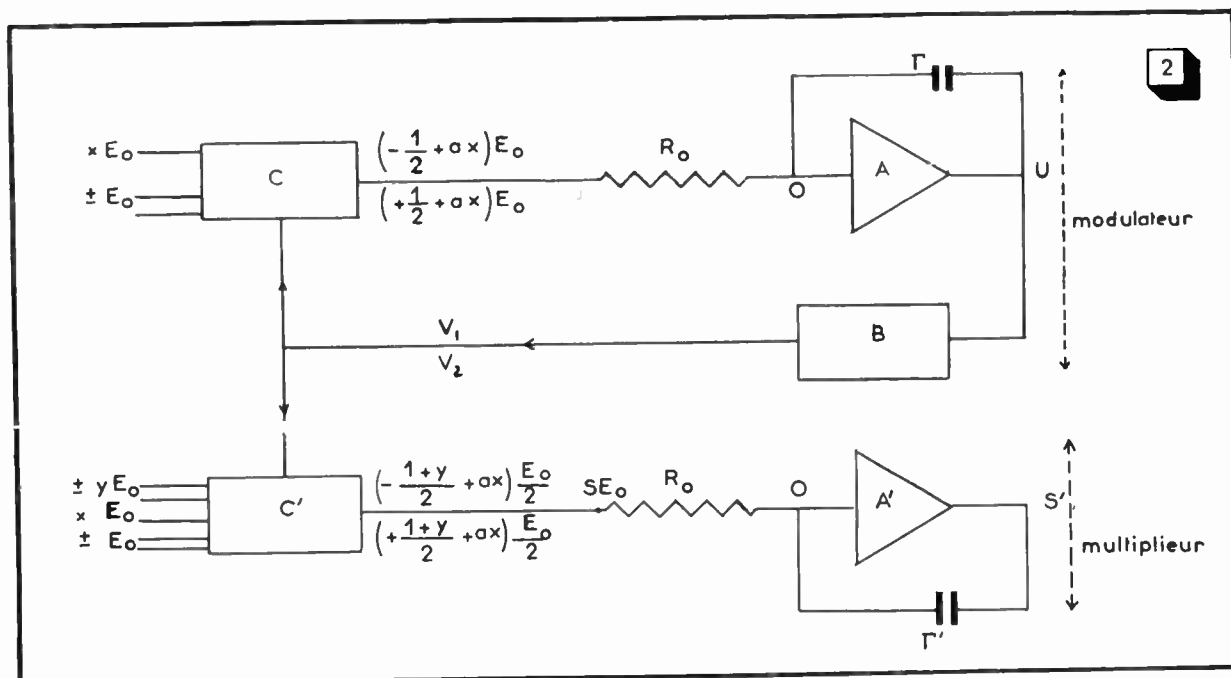
$$S = - \frac{Z_1}{Z_0} E$$

$$\text{Si } \begin{cases} Z_0 = R_0 \\ Z_1 = R_1 \end{cases}$$

nous avons un déphaseur $S = - \frac{R_1}{R_0} E$

$$\text{Si } \begin{cases} Z_0 = R_0 \\ Z_1 = \frac{1}{Cp} \end{cases}$$

nous avons un intégrateur $S = - \frac{1}{R_0 C} \int E dt$



Le schéma de principe indiqué figure 2 ne correspond pas exactement à la réalité, mais il permet de comprendre plus facilement ce qui se passe.

E_0 est une tension de référence fixe ($E_0 = 100$ Volts).

x et y sont les deux nombres (compris entre -1 et $+1$) à multiplier.

C et C' sont deux commutateurs à deux positions commandés en synchronisme par un basculeur B (en réalité B comprend un organe sensible, le basculeur proprement dit, et un organe de puissance).

Le fonctionnement est tel que la tension de sortie U de l'intégrateur oscille entre deux limites U_1 et U_2 fixées d'avance par la bascule et qui sont successivement atteintes.

B agit précisément lorsque :

$$U = U_1 \quad \text{et} \quad U = U_2$$

Pour fixer les idées :

dans une première phase :

B délivre la tension de commande : $V_c = V_1$

C délivre la tension $\left(-\frac{1}{2} + ax\right) E_0$ $0 < a < \frac{1}{2}$

C' délivre la tension $\left(-\frac{1+y}{2} + ax\right) \frac{E_0}{2}$

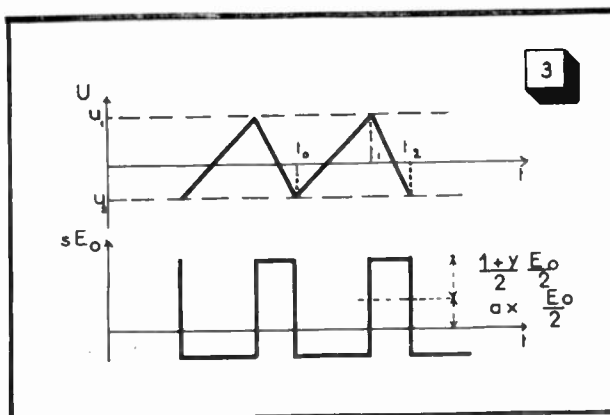
dans une deuxième phase :

B délivre la tension de commande : $V_c = V_2$

C délivre la tension $\left(+\frac{1}{2} + ax\right) E_0$

C' délivre la tension $\left(+\frac{1+y}{2} + ax\right) \frac{E_0}{2}$

Considérons le multiplieur en fonctionnement. A l'instant $t_0 + 0$, $U = U_2$, B vient de basculer et C délivre $\left(-\frac{1}{2} + ax\right) E_0$, c'est le début d'une première phase (fig. 3).



Première phase :

$$U = U_2 + \frac{1}{R_0 \Gamma} \int_{t_0}^t \left(\frac{1}{2} - ax\right) E_0 \, dt \quad \text{du} \quad t \geq t_0$$

U croît (car $\frac{1}{2} - ax > 0$). Soit t_1 l'époque laquelle $U = U_1$ nous avons :

$$(1) \quad U_1 = U_2 + \frac{1}{R_0 \Gamma} \int_{t_0}^{t_1} \left(\frac{1}{2} - ax\right) E_0 \, dt$$

Au temps t_1 une nouvelle commutation a lieu et le système entre dans sa deuxième phase.

Deuxième phase :

$$U = U_1 - \frac{1}{R_0 \Gamma} \int_{t_1}^t \left(\frac{1}{2} + ax \right) E_0 \, du \, t \geq t_1$$

U décroît avec t (car $\left(\frac{1}{2} + ax \right) \geq 0$) jusqu'à l'époque t_2 à laquelle $U = U_2$. Nous avons :

$$(2) \quad U_2 = U_1 - \frac{1}{R_0 \Gamma} \int_{t_1}^{t_2} \left(\frac{1}{2} + ax \right) E_0 \, du$$

B agit alors et le cycle de commutations recommence.

Les relations (1) et (2) s'écrivent, en posant

$$\lambda = R_0 \Gamma \cdot \frac{U_1 - U_2}{E_0}$$

$$(1') \quad \lambda = \int_{t_0}^{t_1} \left(\frac{1}{2} - ax \right) du = \left(\frac{1}{2} - ax_m^1 \right) T_1$$

$$(2') \quad \lambda = \int_{t_1}^{t_2} \left(\frac{1}{2} + ax \right) du = \left(\frac{1}{2} + ax_m^2 \right) T_2$$

x_m^1 et x_m^2 étant les valeurs moyennes respectives de $x(t)$ sur les intervalles (t_0, t_1) et (t_1, t_2) .

La valeur moyenne de la tension $s E_0$ à la sortie du commutateur C au cours du cycle de commutations (t_0, t_2) est $s_m E_0$ telle que :

$$\begin{aligned} s_m &= \frac{1}{2} \frac{1}{T_1 + T_2} \left[\int_{t_0}^{t_1} \left(-\frac{1+y}{2} + ax \right) du \right. \\ &\quad \left. + \int_{t_1}^{t_2} \left(\frac{1+y}{2} + ax \right) du \right] \\ &= \frac{1}{2} \frac{\left(-\frac{1+y_m^1}{2} + ax_m^1 \right) T_1 + \left(\frac{1+y_m^2}{2} + ax_m^2 \right) T_2}{T_1 + T_2} \end{aligned}$$

Les relations (1') et (2') donnent :

$$\frac{\frac{1}{2} - ax_m^1}{T_2} = \frac{\frac{1}{2} + ax_m^2}{T_1} = \frac{1 + a(x_m^2 - x_m^1)}{T_1 + T_2}$$

$$(3) \quad \frac{T_1}{T_1 + T_2} = \frac{\frac{1}{2} + ax_m^2}{1 + a(x_m^2 - x_m^1)}$$

$$(4) \quad \frac{T_2}{T_1 + T_2} = \frac{\frac{1}{2} - ax_m^1}{1 + a(x_m^2 - x_m^1)}$$

d'où :

$$s_m = \frac{1}{2} \frac{1 - a \frac{x_m^2 y_m^1 + x_m^1 y_m^2}{2} + \frac{y_m^2 - y_m^1}{4}}{1 + a(x_m^2 - x_m^1)}$$

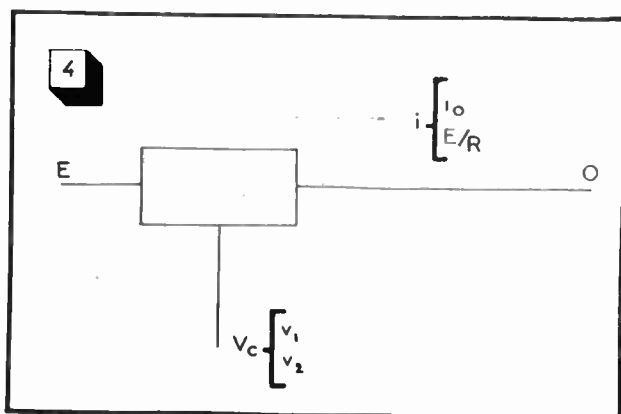
Si x et y varient peu durant un cycle de commutations, on aura (ce que nous supposons par la suite), :

$$s_m = -\frac{a}{2} xy$$

3. — RÉALISATION.

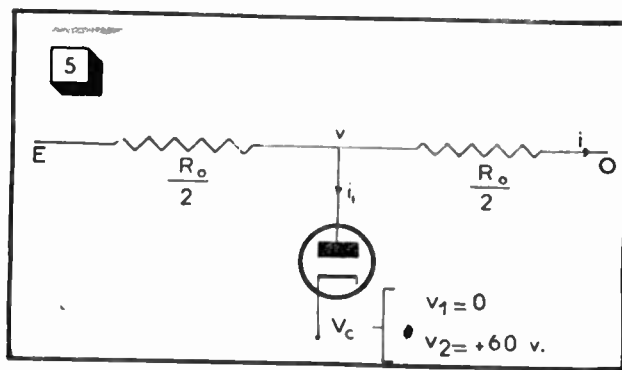
Le multiplieur comprend deux parties :

Le modulateur (partie supérieure de la figure 2)



où se produit la modulation de durée et le multiplieur proprement dit (partie inférieure de la figure 2) où se produit la modulation d'amplitude.

L'élément de base du modulateur et du multiplieur est le commutateur.



3.1. — Commutateur à diodes.

Le fonctionnement du commutateur schématisé figure 4 est le suivant :

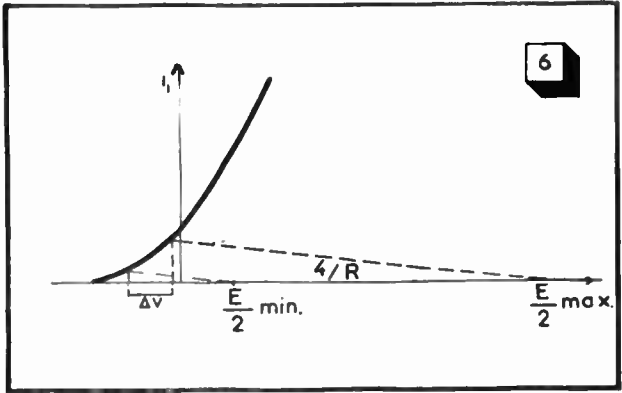
À l'entrée du commutateur, nous avons une tension positive E variant de $E_{\max}$ à $E_{\min}$ (10 à 100 Volts), à la sortie une tension nulle et un courant i tel que :

$$i = i_0 \text{ indépendant de } E, \quad \text{quand } V_c = V_1$$
$$i = \frac{E}{R} \text{ proportionnel à } E, \quad \text{quand } V_c = V_2$$

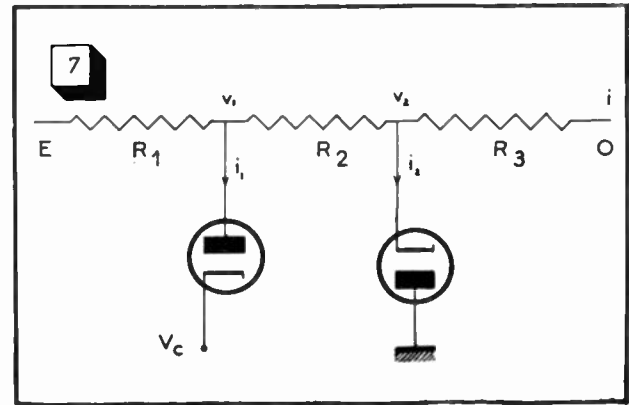
Le système le plus simple qui vient à l'idée est celui indiqué sur la figure 5.

$$i = 0 \quad \text{quand } V_c = 0 \text{ (diode conductrice)}$$
$$i = \frac{E}{R_0} \quad \text{quand } V_c = + 60 \text{ V (diode bloquée)}.$$

En réalité, quand la diode est conductrice, v n'est pas rigoureusement nul; il subsiste une tension dite de déchet qui dépend de E .

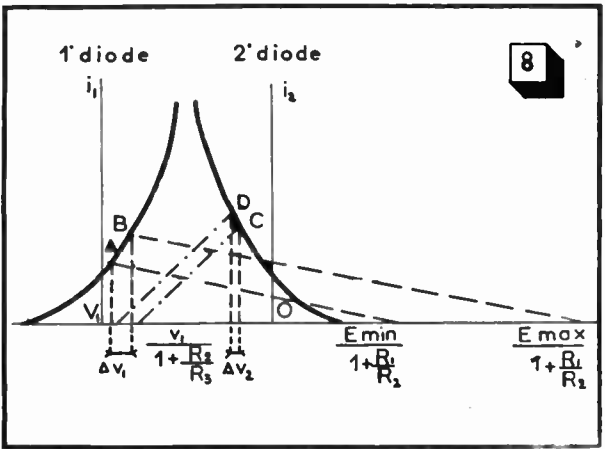


Nous voyons sur la caractéristique de la diode, figure 6, la variation Δv de v quand E varie de E_{min} à E_{max} .
Si la caractéristique était linéaire dans cette région, cette variation serait compensable, mais

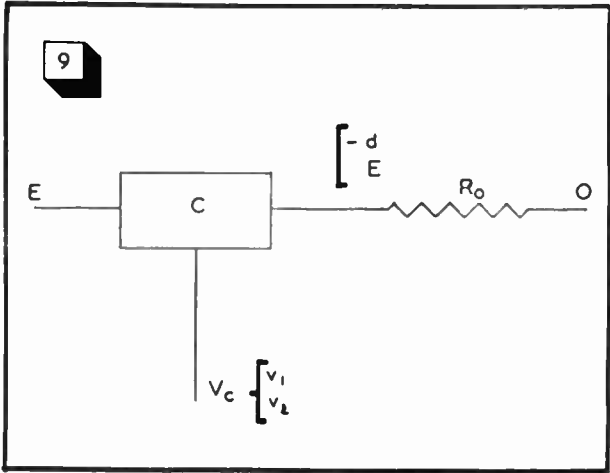


ce n'est pas le cas; et Δv atteint 250 mV dans les conditions suivantes :

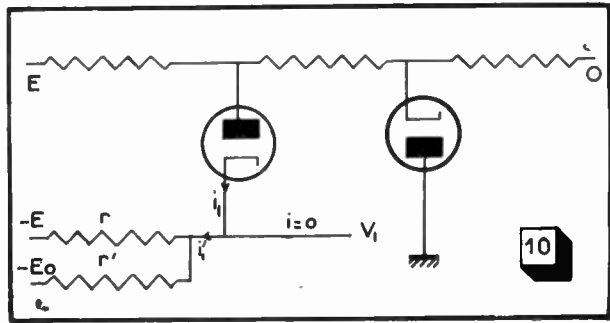
$$\left\{ \begin{array}{l} \text{diode } EB\ 41 : \text{chauffage } 6\text{ V.} \\ R_0 = 1\text{ M}\Omega \\ 10\text{ V} < E < 100\text{ V.} \end{array} \right.$$



Le commutateur finalement retenu après un examen que nous croyons complet de la question, est celui de la figure 7.



La première diode est polarisée négativement ($V_c = V_1 = - 11\text{ V}$) ce qui crée dans la seconde diode (quand les diodes sont conductrices) un courant i_2 constant indépendant de E . Le fonctionnement se voit sur le graphique, figure 8.



La variation Δv_2 de v_2 reste alors inférieure à 1 mV. dans les conditions suivantes :

$$\left\{ \begin{array}{l} \text{Diode } EB\ 41 : \text{chauffage } 6\text{ V.} \\ 10\text{ V} < E < 100\text{ V} \quad v_1 = - 11\text{ V} \\ R_1 = 0,5\text{ M}\Omega \quad R_2 = 0,1\text{ M}\Omega \quad R_3 = 0,4\text{ M}\Omega \end{array} \right.$$

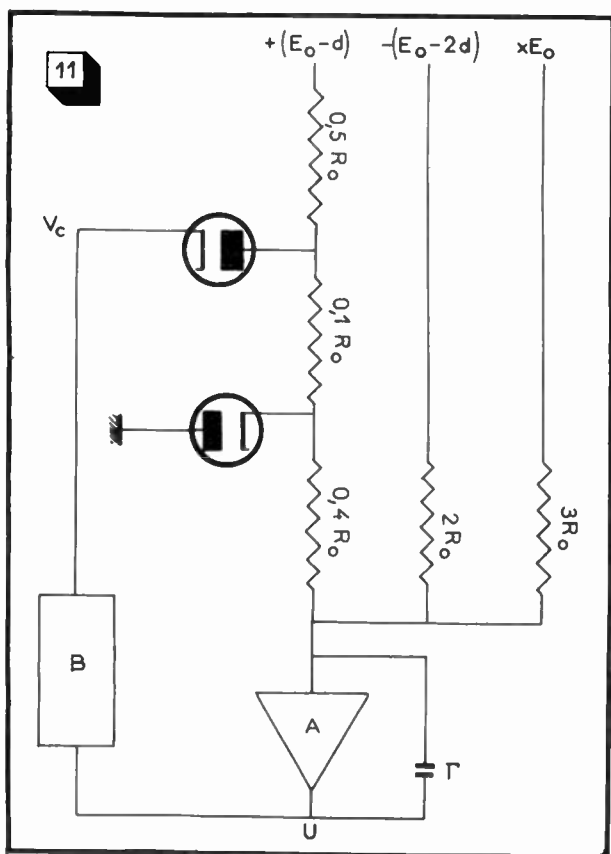
En résumé, le commutateur se comportera (fig. 9) comme une résistance $R_0 = R_1 + R_2 + R_3$ à laquelle serait appliquée une tension égale à :

$$\left\{ \begin{array}{l} -d \text{ quand } V_c = V_1 \text{ (diodes conductrices)} \\ E \text{ quand } V_c = V_2 \text{ (diodes bloquées)} \end{array} \right.$$

avec par définition :

$$-d \cdot \frac{R_3}{R_1 + R_2 + R_3} = v_2 \quad \left(\begin{array}{l} \text{le signe « - » a été} \\ \text{choisi parce que en gé-} \\ \text{néral } v_2 \text{ est négatif} \end{array} \right)$$

Tout ceci suppose que la tension v_1 reste constante, quel que soit le courant i_1 qui passe dans la diode.



En fait, la source de tension V_1 (sortie cathodyne d'un étage de puissance) présente une résistance interne non négligeable (5 kΩ) aussi faut-il, par le dispositif, figure 10, faire passer dans la source en même temps que le courant i_1 un courant i'_1 égal à i_1 et de signe contraire de telle sorte que le courant total soit sensiblement nul.

3.2. — Modulateur (1).

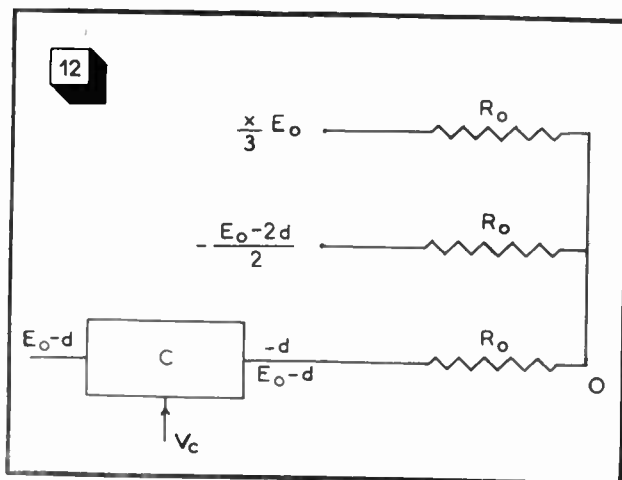
Le schéma du modulateur est indiqué figure 11.

Pendant le temps T_1 les diodes sont conductrices et la tension d'entrée rapportée à une résistance R_0 est (voir figure 11) :

(1) Brevet P. V. 659.014, du 28 novembre 1953 : « Opérateur pour le calcul électronique analogique notamment multiplieur » (F. H. RAYMOND).

$$\left\{ \begin{array}{l} \frac{x E_0}{3} \text{ pour la première branche} \\ \frac{-E_0 - 2d}{2} \text{ pour la seconde branche} \\ -d \text{ pour la troisième branche} \end{array} \right.$$

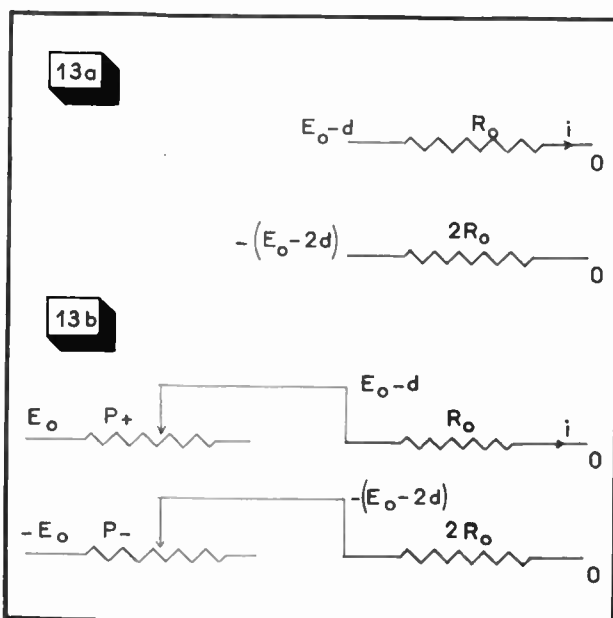
$$= \left(-\frac{1}{2} + \frac{x}{3} \right) E_0 \text{ au total.}$$



Pendant le temps T_2 les diodes sont bloquées et la tension d'entrée rapportée à une résistance R_0 est (voir la figure 12) :

$$\left\{ \begin{array}{l} \frac{x E_0}{3} \text{ pour la première branche} \\ \frac{-E_0 - 2d}{2} \text{ pour la seconde branche} \\ +E_0 - d \text{ pour la troisième branche} \end{array} \right.$$

$$= \left(+\frac{1}{2} + \frac{x}{3} \right) E_0 \text{ au total.}$$



Pratiquement, pour obtenir les tensions $E_0 - d$ et $-(E_0 - 2d)$ d étant positif, le schéma, figure 13 a est remplacé par le schéma équivalent, figure 13 b, après réglage adéquat des potentiomètres P_+ et P_- pour avoir le même courant.

Si maintenant nous nous rapportons au principe du paragraphe 2, en faisant :

$$\left\{ \begin{array}{l} x'_m = x^2_m = x \\ a = \frac{1}{3} \end{array} \right.$$

Les formules (3) et (4) deviennent :

$$(3') \quad \frac{T_1}{T_1 + T_2} = \frac{1}{2} + \frac{x}{3}$$

$$(4') \quad \frac{T_2}{T_1 + T_2} = \frac{1}{2} + \frac{x}{3}$$

La fréquence de découpage F se calcule à partir de (1') et (2'), ce qui donne :

$$F = \frac{1}{T_1 + T_2} = \frac{E_0}{(U_1 - U_2) R_0 \Gamma} \left(\frac{1}{4} - \frac{x^2}{9} \right)$$

Nous pouvons modifier cette fréquence, en faisant varier Γ .

Le rapport $a = \frac{1}{3}$ est choisi afin que :

$$\left(\frac{T_1}{T_2} \right)_{\min.} \geq \frac{1}{5} \quad \text{et} \quad \frac{F_{\min.}}{F_{\max.}} \geq \frac{1}{2}$$

3.3. — Multiplieur.

Le schéma du multiplieur est indiqué figure 14.

Calculons la tension d'entrée moyenne sur la première branche E_m rapportée à une résistance R_0 (voir figure 15 a et 15 b).

$$\left\{ \begin{array}{l} \text{Temps } T_1 \text{ diodes conductrices entrée} = -d' \\ \text{Temps } T_2 \text{ diodes bloquées entrée} = \frac{1+y}{2} E_0 \end{array} \right.$$

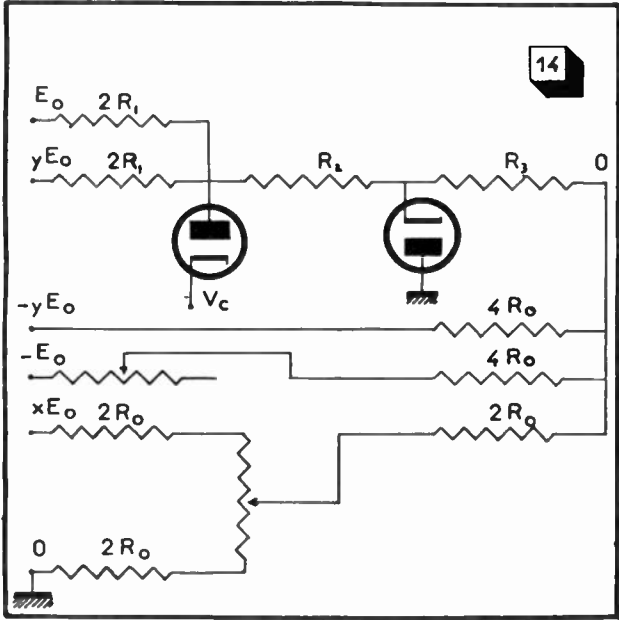
entrée moyenne :

$$E_m = \frac{-d' T_1 + \frac{1+y}{2} E_0 T_2}{T_1 + T_2}$$

En utilisant les relations (3') et (4'), paragraphe 3.2. Nous trouvons :

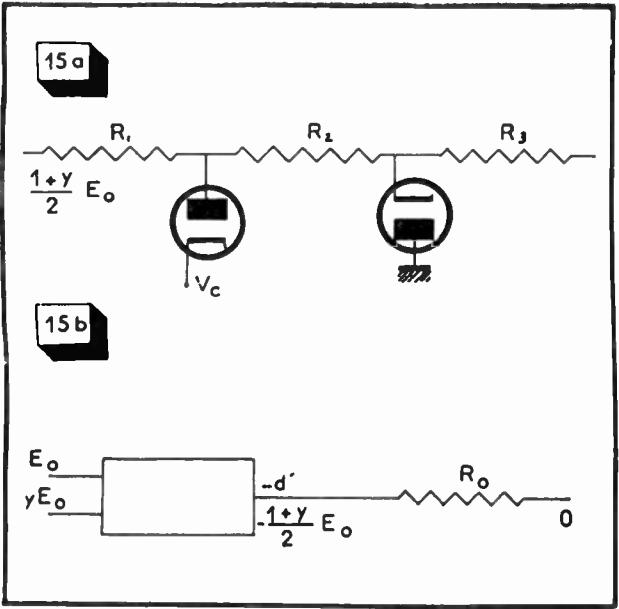
$$E_m = -\frac{xy}{6} E_0 + \frac{y}{4} E_0 - x \left(\frac{1}{6} + \frac{d'}{3 E_0} \right) E_0 + \left(\frac{1}{4} - \frac{d'}{E_0} \right) E_0$$

Il suffira donc d'ajouter trois autres entrées non commutées afin de compenser tous les termes différents de $-\frac{xy}{6} E_0$ (comme pour le modulateur, nous utilisons des potentiomètres pour tenir compte

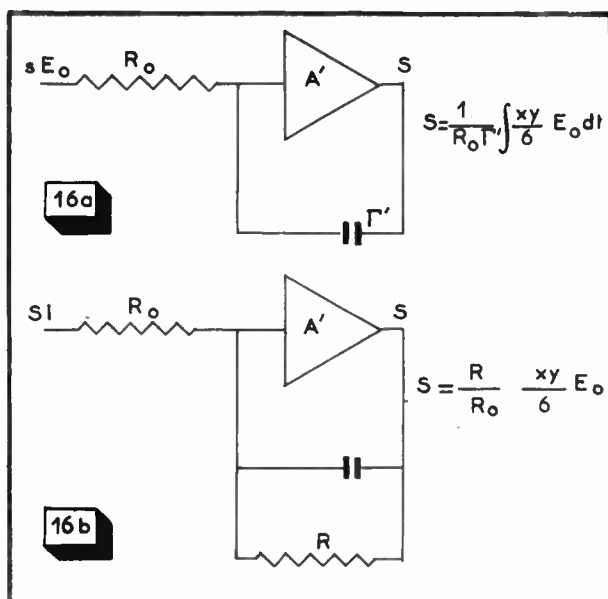


des tensions en d'). En faisant l'approximation $d' = 0$ nous retrouvons le schéma de principe indiqué figure 12.

Toujours figure 2, suivant que l'amplificateur A' est bouclé par une capacité Γ' (fig. 16 a) ou par une



résistance R (fig. 16 b) en parallèle avec une capacité suffisamment grande pour constituer un filtre, nous aurons les tensions suivantes à la sortie S du multiplieur :

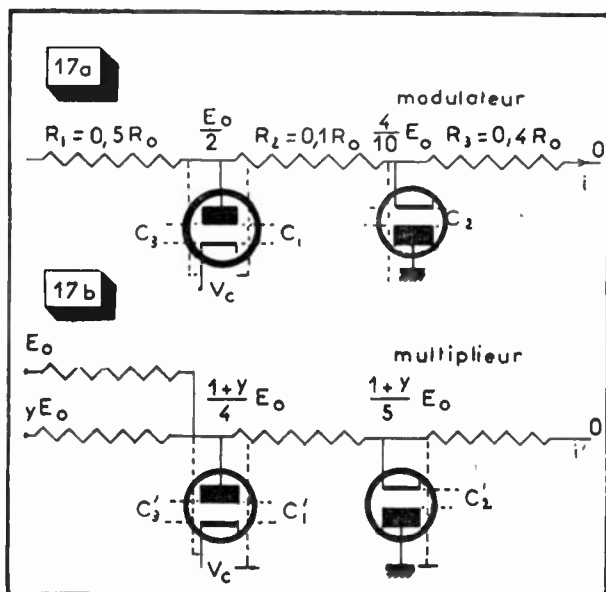


3.4. — Compensation des capacités parasites.

Les capacités parasites sont de l'ordre de 5 pF. Elles produisent des erreurs importantes (un à deux centièmes), elles n'interviennent que parce que nous avons des éléments non linéaires (les diodes) et elles proviennent presque uniquement de ces diodes. Nous voyons figure 17, les capacités parasites C_1 C_2 C_3 du modulateur C'_1 C'_2 C'_3 du multiplieur.

3.41. — Influence de C_1 C_2 , C'_1 C'_2 .

Il est facile de montrer que si les capacités C_1 C_2 sont égales aux capacités C'_1 C'_2 c'est-à-dire si les deux commutateurs ont même coefficients de transfert, les erreurs introduites par ces capacités se compensent.



3.42. — Influence de C_3 et C'_3 .

La capacité C_3 prend la charge :

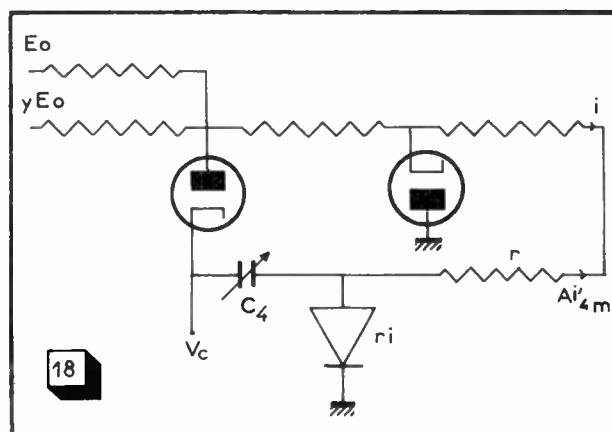
$$Q = C_3 \left(V_2 - \frac{E_0}{2} \right)$$

lorsque $V_c = V_a = +70$ V et se décharge dans la diode quand $V_c = V_i = -11$ V.

Pendant la charge de C_3 , le courant dans la résistance R_3 est non pas $i = E_0/R_0$ mais $i + \Delta i_3$. L'erreur moyenne en courant pendant une période est :

$$\Delta i_{3m} = \frac{1}{T} \frac{R_1}{R_1 + R_2 + R_3} Q$$

$$= + \frac{1}{T} \frac{R_1}{R_1 + R_2 + R_3} C_3 \left(V_2 - \frac{E_0}{2} \right) \text{ (expression positive)}$$



C_3 a même influence qu'une capacité négative $-\Delta C_2$ placée en parallèle avec C_2 et qui introduit une erreur :

$$\Delta i_{3m} = \frac{1}{T} \frac{R_1 + R_2}{R_1 + R_2 + R_3} (-\Delta C_2) \left(0 - \frac{E_0}{4} \right)$$

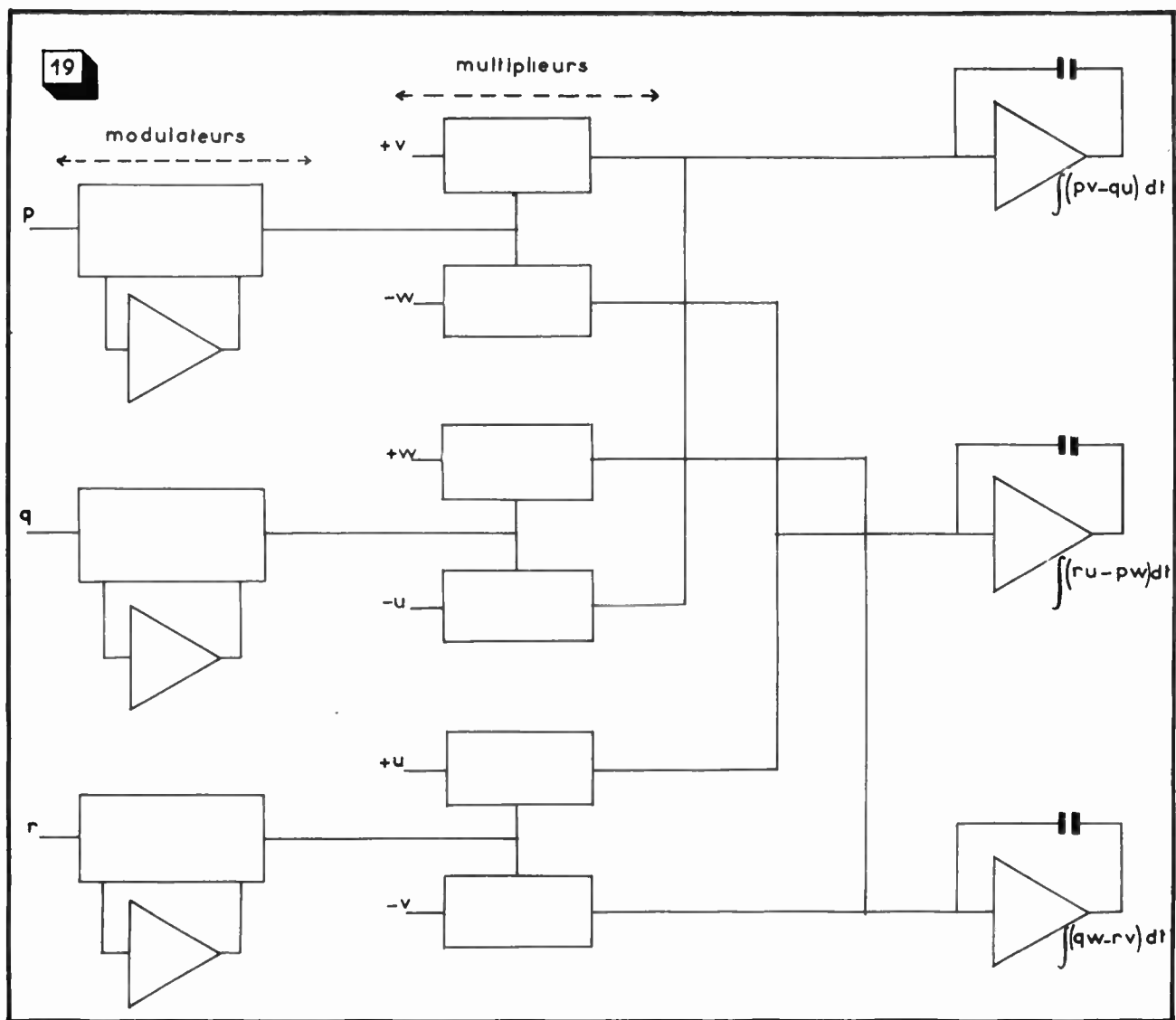
De même la capacité C'_3 introduit une erreur $\Delta i'_{3m}$ telle que :

$$\Delta i'_{3m} = \frac{1}{T} \frac{R_1}{R_1 + R_2 + R_3} C'_3 \left(V_2 - \frac{1+y}{4} E_0 \right)$$

Le terme en $\frac{1+y}{4} E_0$ est identique à celui apportée par une capacité $\Delta C'_2$ en parallèle avec C'_2 qui introduit une erreur :

$$\Delta i'_{3m} = \frac{1}{T} \frac{R_1 + R_2}{R_1 + R_2 + R_3} \Delta C'_2 \left(0 - \frac{1+y}{5} E_0 \right)$$

En ajoutant une capacité ajustable C_0 en parallèle avec C_2 nous pourrions compenser les erreurs



introduites par C'_1 , C'_2 , $\Delta C'_2$ par celles introduites par C_1 , C_2 , $-\Delta C_2$ et C_0 .

Il reste à compenser le terme en V_2 de $\Delta i'_{3m}$.

Cela est réalisé par la capacité C_4 du schéma figure 18 qui introduit une erreur de signe contraire :

$$\Delta i'_{4m} = \frac{1}{T} C_4 (V_1 - V_2) \frac{ri}{r_i + r}$$

4. — EXEMPLE D'UTILISATION D'UN MULTIPLIEUR A MODULATION.

Une propriété essentielle du multiplieur est la « distributivité ». Cette propriété est particulièrement utilisée pour calculer des produits vectoriels. En mécanique du vol, par exemple, il faut obtenir l'intégrale des composantes de produits vectoriels de la forme :

$\vec{w} \wedge \vec{V}$ (dans l'équation de la quantité de mouvement).

$\vec{w} \wedge \vec{H}$ (dans l'équation du moment cinétique).

Si les composantes pqr du vecteur $\vec{w}$ alimentent trois modulateurs, la tension de commande venant de chaque modulateur pourra commander autant de multiplieurs que nécessaire.

Nous voyons, figure 19, comment réaliser le produit $\vec{w} \wedge \vec{V}$ c'est-à-dire les composantes :

$$\begin{cases} qw - rv \\ ru - pw \\ pv - qu \end{cases}$$

Avec les mêmes modulateurs, mais trois autres multiplieurs, nous pourrions réaliser de même $\vec{w} \wedge \vec{H}$.

Il faudra donc au total pour effectuer 12 produits :

- 3 amplificateurs et 3 modulateurs,
- 3 amplificateurs et 6 multiplieurs (pour $\vec{w} \wedge \vec{H}$),
- 3 amplificateurs et 6 multiplieurs (pour $\vec{w} \wedge \vec{V}$).

soit 9 amplificateurs, 3 modulateurs, 12 multiplieurs d'où grosse économie d'amplificateurs et de modulateurs.

5. — RÉSULTATS OBTENUS : CONCLUSION.

L'amplificateur A' (fig. 16 b) étant bouclé en gain de 6 afin d'avoir :

$$S = x y E_0$$

La précision est :

$$\frac{\Delta S}{E_0} = \frac{1}{1\,000}$$

Elle est limitée, par la précision des résistances, et par la stabilité des lampes (variation des caractéristiques des diodes due au vieillissement ou aux

chocs). Le multiplieur réalisé par GOLDBERG aux Etats-Unis est rendu indépendant des caractéristiques des tubes par l'utilisation d'un asservissement en courant et non en tension, mais il a l'inconvénient de nécessiter 6 amplificateurs.

C'est pourquoi, quitte à diminuer légèrement la précision, nous avons recherché à la S. E. A. une solution plus simple ne nécessitant que deux amplificateurs.

6. — BIBLIOGRAPHIE.

BAUM R. V. AND MORRILL : « A Time Division Multiplier for a General Purpose Electronic Differential Analyser », Paper presented to the National Convention on the I.R.E. 1951.

GOLDBERG EDWIN A : « A High Accuracy Time Division Multiplier », Cyclone Symposium 1952, *R.C.A. Review*, Vol. XIII, September 1952, page 265.

C.D. MORRILL AND R.I. BAUM : « A Stabilized Electronic Multiplier », *Transactions of the I.E.R. (P. G. E. C. — 1)* December 1952.

CORRECTEUR AUTOMATIQUE DE FRÉQUENCE POUR ÉMETTEUR DE TÉLÉCOMMUNICATIONS SUR ONDES CENTIMÉTRIQUES

PAR

J. CAYZAC

*Ingénieur aux Laboratoires d'Electronique
et de Physique appliquées*

INTRODUCTION.

a) Généralités.

Dans le domaine des télécommunications hyperfréquence, les exigences sur la stabilité de la fréquence des émetteurs sont suffisamment strictes pour poser parfois un problème délicat à résoudre. Pour éviter les risques d'interférences indésirables entre les émissions, une précision de l'ordre de quelque 10^{-4} peut être rendue nécessaire, notamment dans les liaisons à plusieurs tronçons qui utilisent différentes fréquences groupées dans une même bande.

La technique actuelle des tubes dans certains cas permet de réaliser des émetteurs à modulation de fréquence d'une puissance de quelques watts ou quelques dizaines de watts ne comportant qu'un seul tube hyperfréquence auto-oscillateur.

Le fait d'être aisément modulable en fréquence est lié pour l'oscillateur à une instabilité en fonction de ses tensions d'alimentation. Citons à titre d'exemple le tube à réflexions multiples dont la puissance SHF est de 10 watts à 3 400 Mc/s et qui nécessite une tension d'anode de 3 000 volts; la pente en fréquence, exprimée en kilocycles par volt, de cette électrode est de 100 à 200; cette considération montre que pour garantir en exploitation une stabilité suffisante, une alimentation de très haute précision serait nécessaire. D'autre part, la variation des sources de tensions n'est pas la seule cause de dérive; il est donc plus rentable et plus sûr d'utiliser des alimentations stabilisées d'un type courant et de demander à un correcteur automatique de fréquence de réduire de façon globale la dérive due aux différents éléments.

b) Conditions d'utilisation du correcteur automatique de fréquence.

Notons tout d'abord que l'analyse du fonctionnement des dispositifs dont nous parlerons dans le

présent article ne s'applique en toute rigueur qu'aux émissions à modulation d'amplitude ou aux émissions à modulation de fréquence dans lesquelles les composantes spectrales sont symétriquement réparties autour d'une fréquence fixe. Cela n'empêche pas qu'ils puissent être souvent employés efficacement aussi dans des cas où le spectre est dissymétrique (onde modulée en fréquence par des signaux de télévision, par exemple), bien qu'alors il se produise en général un léger décalage de la fréquence centrale lors de l'application du signal modulateur.

Dans le cas mentionné ci-dessus le porteur n'apparaît dans aucun circuit dissocié de la bande de fréquences chargée des informations à transmettre; le correcteur automatique devra donc délivrer une tension de correction qui sera fonction de l'écart existant entre une fréquence de référence et par exemple le milieu de la bande de fréquences du signal émis. Cela implique l'emploi d'un circuit discriminateur dont la bande d'utilisation soit calculée en fonction de l'excursion de fréquence de l'émetteur et dont le temps de réponse soit suffisamment lent pour ne pas restituer les composantes de fréquences les plus basses de la modulation transmise.

L'énergie SHF consommée par le correcteur automatique ne doit naturellement pas représenter une fraction substantielle de la puissance de sortie de l'émetteur.

Le couplage des circuits de correction de fréquence au feeder de l'aérien ne doit pas introduire de réflexions notables dans celui-ci; tout élément sélectif risquerait de créer des distorsions dans la modulation.

DESCRIPTION D'UN PROCÉDÉ DE CORRECTION AUTOMATIQUE DE FRÉQUENCE.

Le procédé qui consiste à comparer la fréquence centrale de l'émetteur avec celle d'un oscillateur à

cristal convenablement multipliée remplit parfaitement les conditions électriques demandées ci-dessus.

Nous nous proposons toutefois de décrire un autre procédé, qui, bien que moins précis, est suffisant lorsque la stabilité requise est de l'ordre de

de la caractéristique sera prise comme fréquence de référence dans le dispositif d'asservissement et elle doit être particulièrement stable ; sa stabilité est fonction d'une part de celle de l'accord des cavités et d'autre part de celle des deux détecteurs.

Si la sensibilité de l'un des deux cristaux se trouve

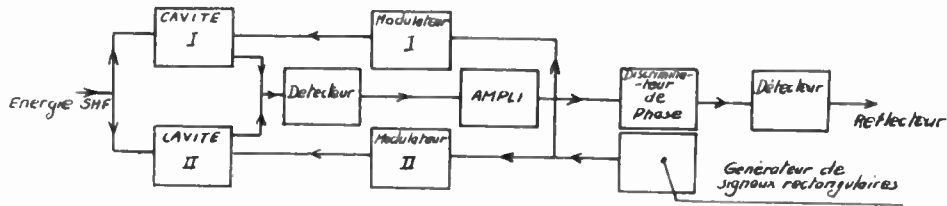


FIG. 1

10^{-4} à 10^{-6} ; il présente en outre un certain nombre d'avantages concernant la réalisation, notamment lorsque la partie hyperfréquence doit se présenter sous la forme d'une tête séparée du reste de l'équipement. Dans ce cas, le correcteur de fréquence dont il est question ne fait appel, sur la tête hyperfréquence, à aucun circuit contenant des tubes électroniques, ni à aucun élément nécessitant un entretien quelconque ; les circuits alors contenus dans cette tête sont peu encombrants et peuvent être reliés sans inconvénient à la seconde partie du correcteur de fréquence par plusieurs dizaines de mètres de câble. Cette seconde partie ne contient que des tubes de réception ordinaires.

a) Principe de fonctionnement.

Un circuit discriminateur rappelant le système Travis est utilisé : 2 cavités résonnantes sont accordées sur deux fréquences distinctes dont la différence est fonction de la largeur de la caractéristique désirée.

Il serait naturellement possible en détectant la

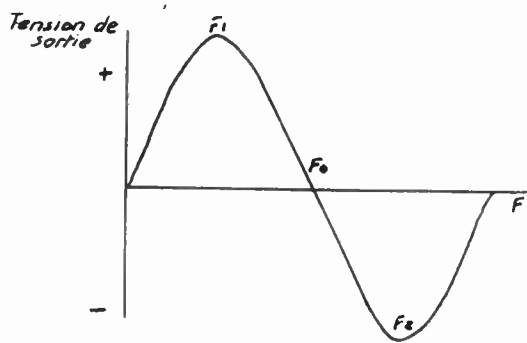


FIG. 2

tension UHF de chaque cavité à l'aide de deux détecteurs à cristal de polarités inverses d'obtenir une tension de sortie après mélange qui varie en fonction de la fréquence selon la loi représentée figure 2.

La fréquence F_0 qui est la fréquence centrale

modifiée pour une raison quelconque, F_0 change de valeur et devient F'_0 comme le montre la figure 3. La courbe n'est d'ailleurs plus symétrique par rapport à F'_0 et si le centre de la bande de fréquence du signal appliqué coïncide avec F'_0 la tension de

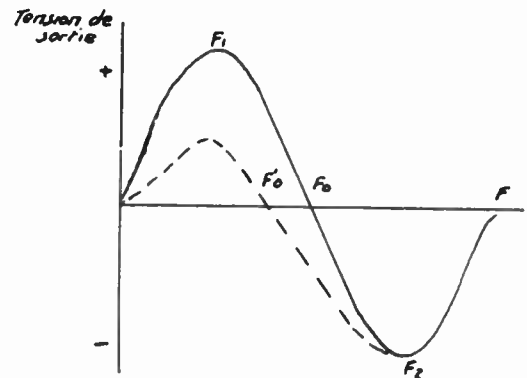


FIG. 3

sortie n'est plus nulle comme dans le fonctionnement normal. Ce second effet produit malheureusement un écart qui s'ajoute à celui qui causait déjà le déplacement de F_0 . On voit donc l'importance que peut avoir la stabilité des détecteurs, importance d'autant plus grande que l'écart entre les fréquences d'accord F_1 et F_2 des deux cavités est plus grand.

Le temps de réponse demandé au discriminateur est relativement très long ; des constantes de temps de l'ordre de la seconde sont parfaitement admissibles puisque la dérive à corriger est très lente ; c'est pourquoi une solution permettant d'éliminer la condition d'équilibrage de deux détecteurs en maintenant la stabilité du système a pu être utilisée : elle consiste à n'employer qu'un seul cristal branché alternativement à la sortie de l'une ou de l'autre cavité. Cette commutation dans le temps a lieu à très basse fréquence, à 50 périodes par exemple. L'information d'écart de la fréquence par rapport à F_0 est traduite de la façon suivante :

1° Si la fréquence centrale F du spectre du signal incident est comprise entre F_0 et F_1 ou se trouve au-delà de F_1 , le courant à la sortie du cristal

est plus élevé lorsqu'il est branché sur la cavité I que lorsqu'il est branché sur la cavité II (fig. 4 a).

2° Si F est comprise entre F_0 et F_2 ou se trouve au-delà de F_2 , c'est la cavité II qui donne le courant le plus élevé (fig. 4 b).

3° Si $F = F_0$ les courants sont égaux pour les deux cavités comme le montre la figure 4 c.

En résumé l'écart de fréquence entre F et F_0 est traduit par la variation d'amplitude des courants (fig. 4) ; le signe de cet écart est donné par la phase du signal rectangulaire d'amplitude relativement à la phase de commutation (fig. 4 a et 4 b).

L'égalité des fréquences F et F_0 se traduit par l'absence de signal alternatif.

Ce dernier point a une importance capitale en ce qui concerne l'altération possible de la valeur de la fréquence de référence par les circuits d'interprétation.

- Ces circuits doivent être sensibles :
- 1° à la phase du signal, laquelle devra déterminer la polarité de la tension de sortie ;
 - 2° à l'amplitude du signal.
- On ne transmet du signal de la figure 4 que la

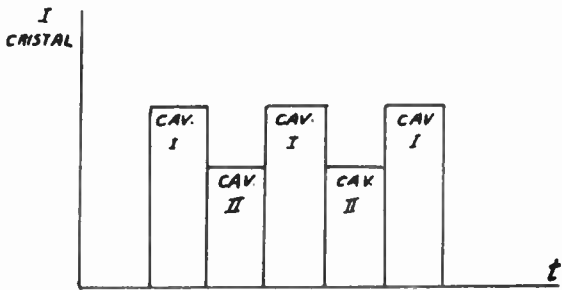


FIG. 4 a

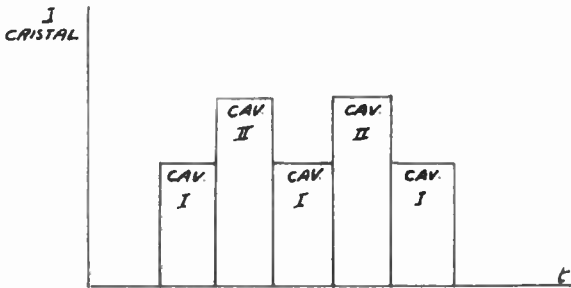


FIG. 4 b

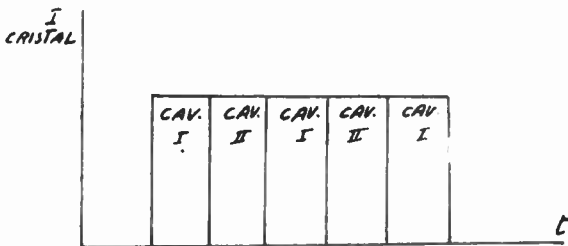


FIG. 4 c

composante alternative ; par conséquent l'amplificateur utilisé est du type à courant alternatif. Une variation des caractéristiques de cet ampli-

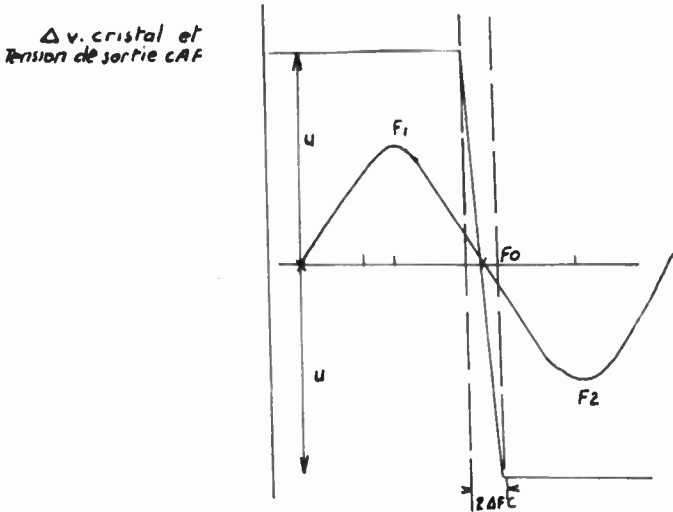


FIG. 5

cateur peut modifier les rapports d'amplitude des signaux, ce qui se traduira, soit par une variation de la pente du discriminateur, soit par une déformation, le seul point de la caractéristique qui ne peut être modifié est le point F_0 caractérisé par l'absence d'information transmise à l'amplificateur.

Voyons dans quelle mesure intervient la stabilité de l'amplificateur à courant alternatif :

La valeur limite de la tension de sortie du discriminateur est égale à $\pm u$ (fig. 5), cette valeur limite est calculée en fonction des caractéristiques du tube hyperfréquence utilisé.

La variation maximum de fréquence que pourra effectuer l'oscillateur stabilisé sans que le système de régulation sorte de sa zone de fonctionnement est donnée par :

$$\pm \Delta F = \frac{u}{GS}$$

où S est la pente $\frac{dV}{dF}$ du circuit hyperfréquence,

V la tension alternative à la sortie du cristal, G le gain des circuits d'interprétation et d'amplification.

Il apparaît donc que si G diminue, la plage de régulation s'élargit et la précision du correcteur est réduite dans les mêmes proportions que G , ce qui se traduit également par une réduction du coefficient de stabilisation.

Le coefficient de stabilisation défini par :

$$\frac{dF}{df} = \frac{\text{variation de fréquence sans correcteur}}{\text{variation de fréquence avec correcteur}}$$

(ces variations étant naturellement dues à la même cause, une modification de tension d'alimentation du tube oscillateur par exemple), est égal à : $SG\alpha$.

étant la pente en Mc/s par volt de l'électrode de contrôle du tube hyperfréquence.

Par exemple, si l'écart de fréquence maximum Δf est de $10^{-4} F_0$ pour un gain normal de l'amplificateur, et si la perte de gain en cours de service est de 6 dB, Δf max. devient égal à $2 \cdot 10^{-4} F_0$, mais le centre de la plage de régulation ne se trouve pas décalé. Si le système de correction est calculé avec une marge de sécurité suffisante, la tension correspondant à Δf max. n'est jamais atteinte et la tension de sortie reste en général, sauf pendant la période initiale d'échauffement, au voisinage de 0 ; la variation de gain de l'amplificateur ne se traduit donc que par un écart de fréquence bien inférieur à $p \Delta f$ max., p étant le rapport entre le gain normal et le gain, accidentellement réduit en cours de service, de l'amplificateur. Naturellement cette considération n'est valable que par le fait que la fréquence centrale, définie par l'absence de signal, ne peut pas être altérée par une modification des caractéristiques des circuits d'interprétation.

Précisons que ce procédé est particulièrement intéressant dans le cas où la caractéristique hyperfréquence doit couvrir une bande de fréquences beaucoup plus large que l'écart autorisé par la précision demandée, ce qui est vrai en général dans les systèmes de télécommunications en ondes décimétriques ou centimétriques.

b) Exemple de réalisation.

Le procédé décrit ci-dessus a été utilisé pour asseoir un tube dont la puissance de sortie est de 10 watts ; la fréquence F_0 de travail est de 3 407,5 Mc/s, l'excursion de fréquence due à la modulation de ± 5 Mc/s ; l'électrode de contrôle possède une pente de 50 kc/s par volt ; la plage de tension de contrôle est de + 100 V à - 100 V par rapport au potentiel de la masse.

1° Partie hyperfréquence du correcteur automatique de fréquence.

La largeur de bande requise est de l'ordre de 13 Mc/s, on choisit donc les fréquences d'accord des cavités égales à $F_1 = F_0 - 7,5$ Mc/s et $F_2 = F_0 + 7,5$ Mc/s.

Voyons quelle est la valeur de la surtension à donner aux cavités pour que la courbe de discrimination soit linéaire de part et d'autre de F_0 .

Les équations donnant les tensions de sortie des deux cavités sont :

$$V_1 = \frac{V_0}{\sqrt{1 + 4Q^2 \left(\frac{\Delta f_1}{F_1} \right)^2}} \quad V_2 = \frac{V_0}{\sqrt{1 + 4Q^2 \left(\frac{\Delta f_2}{F_2} \right)^2}}$$

V_1 et V_2 sont les tensions de sortie des deux cavités pour une fréquence :

$$F = F_1 + \Delta f_1 = F_2 + \Delta f_2$$

V_0 est la tension de sortie de l'une ou de l'autre des cavités à sa fréquence d'accord et Q le coefficient de surtension que nous supposons égal pour les deux cavités.

La tension de sortie après combinaison des deux tensions est, à un facteur constant près :

$$y = V_1 - V_2 = \frac{V_0}{\sqrt{1 + 4Q^2 \left(\frac{\Delta f_1}{F_1} \right)^2}} - \frac{V_0}{\sqrt{1 + 4Q^2 \left(\frac{\Delta f_2}{F_2} \right)^2}}$$

En posant :

$$x = 2Q \left(\frac{F - F_0}{F_0} \right) a = 2Q \left(\frac{F_0 - F_1}{F_0} \right) = 2Q \left(\frac{F_2 - F_0}{F_0} \right)$$

et en tenant compte de ce que : $F_1 \approx F_2 \approx F_0$

$$\text{on peut écrire : } y = \frac{V_0}{\sqrt{1 + (x + a)^2}} - \frac{V_0}{\sqrt{1 + (x - a)^2}}$$

$$\text{qui sera linéaire autour de } x = 0 \text{ si } \left(\frac{d^3 y}{d^3 x} \right)_{x=0} = 0$$

ce qui a lieu pour : $a = \sqrt{1,5}$

$$\text{c'est-à-dire pour : } Q = \frac{F_0 \sqrt{1,5}}{2(F_0 - F_1)}$$

comme : $F_0 = 3407,5$ Mcs et

$$F_0 - F_1 = 7,5 \text{ Mcs}$$

$$\text{on trouve : } Q = \frac{3407,5 \times 1,225}{2 \times 7,5} = 278$$

Cette valeur du coefficient de surtension commune Q est aisée à obtenir. Pratiquement on a réalisé les deux cavités dans un même bloc métallique placé au-dessus du guide canalisant l'énergie. De ce guide partent deux lignes coaxiales auxiliaires très faiblement couplées, disposées dans un même plan perpendiculaire au sens de circulation de l'énergie. Chacune de ces lignes excite l'une des deux cavités. Cette disposition rend les niveaux relatifs indépendants du taux d'onde stationnaire existant dans le feeder de l'aérien.

Les cavités travaillent en mode TE 011, comme l'indique la figure 6, elles sont attaquées par le petit côté, donc par l'intermédiaire d'une boucle, dont on a rendu l'orientation réglable pour doser aisément l'énergie injectée.

Dans la cloison qui sépare les deux cavités est pratiquée une ouverture de petite dimension permettant de coupler un cristal détecteur unique simultanément aux deux cavités ; ce dispositif est schématisé figure 7.

Le cristal détecte en permanence une tension qui est fonction de la somme des niveaux SHF sortant des cavités. La commutation à 50 cycles dont il a été question précédemment doit donc avoir lieu avant

ce système de couplage à 3 directions ; c'est par l'action directe d'un modulateur à absorption sur chaque cavité qu'a été réalisée la commutation (1).

Les deux modulateurs fonctionnent en opposition

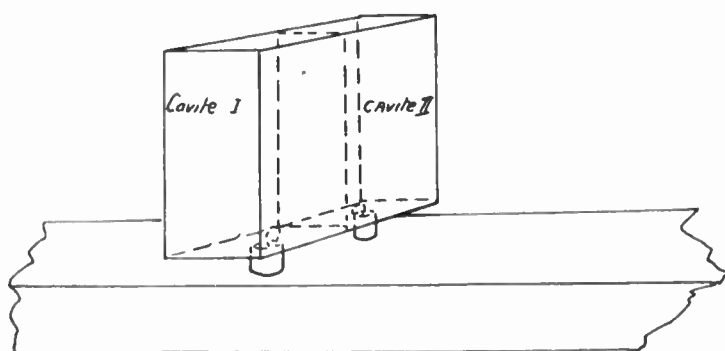


Fig. 6

de phase en sorte que l'une des cavités est complètement amortie et ne délivre pratiquement plus aucune énergie au cristal tandis que l'autre possède le coefficient de surtension prévu par le calcul précédent (environ 300) ; dans la phase de commutation suivante, l'inverse se produit.

Les modulateurs à absorption utilisent les propriétés du ferroxcube : une ligne coaxiale couplée d'un côté à la cavité à moduler est terminée à son extrémité opposée par une boucle traversée par un barreau de ferroxcube. Ce matériau possède un angle de perte qui est fonction du champ magnétique dans lequel il est plongé en sorte que si on place le barreau dans une bobine parcourue par un courant variable,

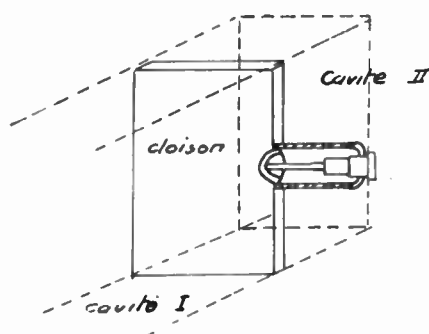


Fig. 7

on modifie le coefficient de surtension de la cavité au rythme de la pulsation de ce courant. Notons que l'adaptation du modulateur à la cavité est réalisée par un transformateur coaxial à deux sections

(1) Les modulateurs à absorption utilisés dans ce dispositif de correction automatique de fréquence sont une application de l'effet signalé par H. G. BELJERS, W. J. van de LINDT et J. J. WENT dans *Journal of Applied Physics*, décembre 1951.

Rappelons que pour des fréquences comprises entre 100 et 10 000 Mc/s certaines ferrites présentent des pertes magnétiques susceptibles d'être réduites de façon notable par l'application d'un champ magnétique relativement faible et dirigé convenablement. Si on introduit une barre de cette ferrite dans une cavité résonnante couplée à un guide d'onde parcouru par l'énergie hyperfréquence, une modulation de cette énergie à la sortie du guide d'onde peut être obtenue en soumettant la barre de ferrite à un champ magnétique variable.

quart-d'onde, disposé sur la ligne mentionnée plus haut. L'allure de la courbe donnant le coefficient de surtension d'une cavité en fonction du nombre d'ampères-tours parcourant la bobine modulatrice est reproduite figure 8.

On voit que la valeur de Q est indépendante du signe du courant.

Il sera donc nécessaire d'appliquer à la bobine modulatrice un signal non symétrique par rapport au niveau 0 ; la réalisation de cette condition présente des inconvénients étant donné que la bobine doit nécessairement être excitée à travers un transformateur d'adaptation ; il est beaucoup plus simple de « polariser » magnétiquement le ferroxcube par un champ magnétique permanent, ce qui situe le point de travail en A ou en B sur la courbe de la figure 8, selon le sens de l'aimant ; les signaux de commutation peuvent dans ce cas être symétriques par rapport au courant 0. Un autre avantage de

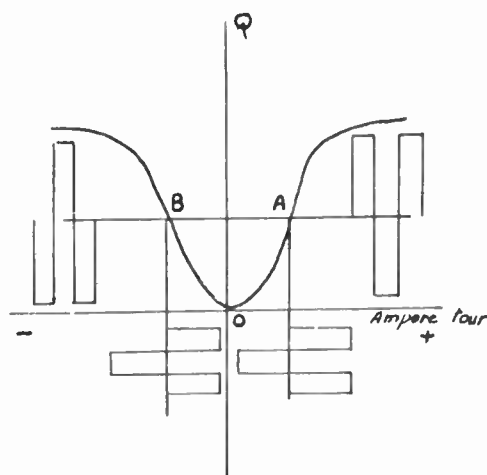


Fig. 8

cette combinaison est de permettre un fonctionnement en opposition de phase des deux modulateurs, excités par le même courant modulant, par simple inversion de signe du champ de polarisation, comme le montre la figure 8 (point de fonctionnement A ou B).

La résolution du problème capital de la stabilité du coefficient de surtension des cavités pendant la période active de la commutation se trouve facilitée par le fait qu'à partir d'une certaine valeur de champ magnétique l'angle de perte du ferroxcube ne varie plus ; il suffit donc de prévoir un courant de commutation surabondant dans la bobine de modulation pour que le fonctionnement ne soit plus tributaire des instabilités du signal modulateur. Ces instabilités ne pourraient d'ailleurs pas avoir de très graves conséquences puisque nous venons de voir que les bobines des modulateurs sont montées en série et parcourues par le même courant ; l'effet serait donc symétrique et n'agirait que sur la pente de la caractéristique globale, sans déplacer la fréquence de référence F_0 . Le schéma de principe du système de modulation est donné figure 9.

Il faut noter que toute modification des éléments couplés aux cavités (transformateurs, lignes et boucles) réagit légèrement sur les fréquences de résonance ; pour éviter les variations de fréquence dues à la température, il a paru plus simple d'enfermer l'ensemble dans une enceinte munie d'un thermostat que d'essayer de réaliser des compensations.

après amplification le tube dont la tension d'écran sera en phase avec cette tension rectangulaire donnera un courant tant que l'écran restera positif ; l'autre restera bloqué, son potentiel d'écran étant négatif pendant toute la durée de la partie positive du signal de grille.

Aux bornes de la résistance de charge du premier

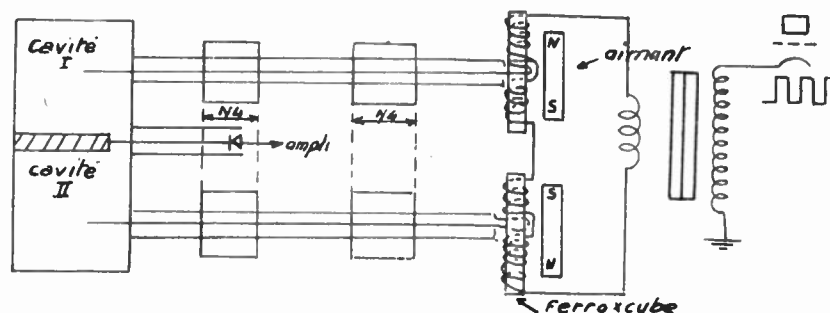


FIG. 9

2° Partie basse fréquence du correcteur automatique de fréquence.

Le discriminateur de phase a pour rôle, comme nous l'avons vu, d'identifier quelle est la cavité qui fournit le niveau le plus élevé ; en fonction de cette indication, il doit aiguiller le signal rectangulaire dont l'amplitude est fonction de l'écart de fréquence sur l'un ou l'autre des deux détecteurs de polarités inverses qui fournissent la tension de correction au réflecteur du tube hyperfréquence.

Le principe de fonctionnement du discriminateur de phase est le suivant :

On applique, sur les écrans de deux tubes pentodes montés comme l'indique la figure 10 des signaux de formes identiques mais dont les phases diffèrent de

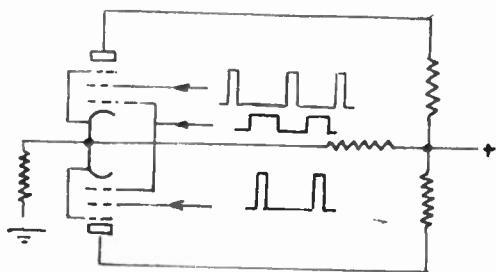


FIG. 10

180° ; ces signaux sont engendrés à partir du courant de modulation du commutateur à ferrocube ; la durée de la partie positive a été réduite de moitié, en partie au début et en partie à la fin des impulsions positives comme l'indique la figure 10.

D'autre part, quelle que soit la valeur instantanée de la tension de l'écran des pentodes, la polarisation des grilles de commande est telle que les deux tubes sont maintenus à la coupure ; si dès lors on applique sur les grilles la tension rectangulaire du cristal

tube on peut recueillir une tension de forme rectangulaire dont l'amplitude varie en fonction du niveau du signal de grille. Aucune tension alternative n'existe dans ce cas aux bornes de la résistance d'anode du second tube. Un détecteur, respectivement D_1 ou D_2 , est disposé dans chaque circuit

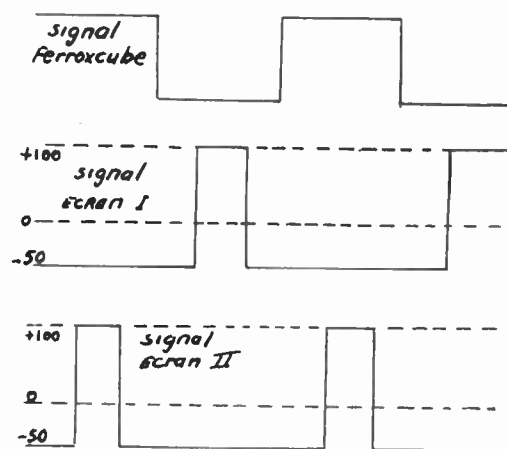


FIG. 11

d'anode. Un additeur effectue la somme des tensions détectées comme l'indique le schéma de la figure 12. Pour ne pas perdre la moitié du signal alternatif, des diodes de restauration de zéro, R_1 et R_2 ont été utilisées.

Le circuit additeur est constitué par deux tubes à charge cathodique, de manière que la tension de sortie ne risque pas d'être modifiée par le débit variable du réflecteur. Les grilles de ces deux tubes sont directement connectées à la sortie des détecteurs lesquels donnent des tensions de même signe ; la cathode C de l'un des deux tubes est reliée à la masse, la cathode A de l'autre tube est reliée au réflecteur, l'alimentation étant bien entendu isolée. Si le détecteur D_1 fournit une tension alors que le détecteur D_2 n'est pas excité, la tension de sortie

qui est envoyée au réflecteur, égale à la somme algébrique des tensions $V_o - V_c$ et $V_a - V_b$ (voir figure 12) est évidemment négative. Si au contraire c'est le détecteur D_2 qui fournit une tension, le point A et par suite le réflecteur se trouveront à un potentiel positif par rapport à la masse. Si les détecteurs D_1 et D_2 ne donnent pas de courant, la tension de sortie est nulle, à condition toutefois que l'équilibrage des courants dans les deux tubes de sortie soit parfait : un essai sur des tubes usagés a

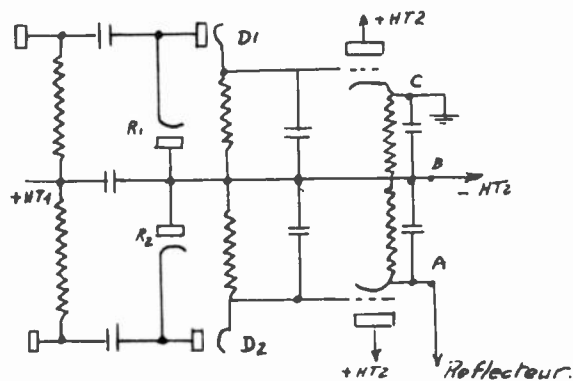


FIG. 12

montré que l'on pouvait s'attendre, si on utilise des tubes à caractéristiques peu poussées, à des déséquilibres de $\pm 0,5$ V au maximum.

Le gain de la chaîne est, comme il sera vu plus loin, de 10^5 donc le déséquilibre de $\pm 0,5$ V équivaut à un signal détecté dans le cristal de sortie des cavités égal à $\frac{500 \text{ mV}}{10^5} = 0,5$ microvolt.

c) Coefficient de stabilisation.

Soit ± 250 kc/s l'écart de fréquence moyenne tolérable, par rapport à la fréquence centrale fixée par les cavités.

Soit ± 100 volts la zone des tensions de correction que l'on s'autorise à appliquer au réflecteur : cette zone est limitée par les caractéristiques du tube, par rapport auxquelles une grande marge de sécurité a été prise.

Si la sensibilité du cristal est telle qu'un écart de 250 kc/s provoque à la sortie de ce cristal, compte tenu des couplages des cavités et de la puissance émise l'apparition d'une tension de 1 mV, le gain de la chaîne amplificatrice totale se trouve fixé à $\frac{100 \text{ V}}{1 \text{ mV}} = 10^5$

La pente fréquence/tension du réflecteur détermine alors à la fois le coefficient de stabilisation et la plage des dérives naturelles de fréquence que le système pourra corriger sans sortir de la zone de tension permise.

Supposons que cette pente soit 50 kc/s/volt, on obtiendra un gain de boucle (fréquence/fréquence) de 20, selon le schéma de la chaîne :

$$250 \text{ kc/s} \longrightarrow 1 \text{ mV} \longrightarrow 100 \text{ V} \longrightarrow 5 \text{ Mc/s}$$

Le coefficient de stabilisation sera alors de 20 (théoriquement 21). En pratique pour conserver un coefficient de stabilisation convenable même lorsque la puissance émise ou la sensibilité du cristal diminueront, il est préférable de donner à la chaîne amplificatrice un gain 1,5 fois plus fort ; on évitera, si on le désire, que la tension de correction dépasse ± 100 volts en déclenchant un dispositif d'alarme (ou même de correction discontinue sur la tension d'une autre électrode) lorsque cette limite sera atteinte ; on réduira par là la plage de dérive naturelle à la zone ± 100 V tout en conservant (à pleine puissance et à pleine sensibilité de cristal) un coefficient de stabilisation 1,5 fois plus fort.

Comme le gain du discriminateur de phase est en fait de l'ordre de 60 (la saturation n'intervenant pas avant que la tension de sortie atteigne ± 150 V), il suffit d'une tension de $\pm 2,5$ V à l'entrée de ce circuit, en sorte que le gain des amplificateurs à intercaler entre le cristal et l'entrée du discrimi-

nateur de phase est seulement de $\frac{1,5 \cdot 10^5}{60} = 2500$.

APPENDICE I

Il est parfois commode de pouvoir changer rapidement la longueur d'onde d'un émetteur sans que cette opération nécessite un réglage compliqué.

Si le tube hyperfréquence est capable de fonctionner dans une bande de fréquence assez large, le correcteur automatique de fréquence tel que nous l'avons décrit est établi pour une valeur fixe.

Pour rendre réglable la fréquence de travail du discriminateur, il est nécessaire de pouvoir décaler simultanément d'une même quantité les accords des cavités.

Les coefficients de surtension de ces dernières ainsi que les niveaux relatifs d'énergie qu'elles sont susceptibles de délivrer au cristal ne doivent pas varier de façon notable.

Avec un couple de deux volumes résonnants munis chacun d'une vis d'accord, plongeant dans un sens parallèle à celui du champ électrique, et un dispositif de commande unique de deux déplacements de ces deux vis, on a pu obtenir un alignement qui se conserve de façon satisfaisante dans une bande de fréquence de 150 Mc/s ; notons que dans ces limites, pour un décalage donné de la fréquence incidente par rapport à la fréquence centrale de la caractéristique de discrimination, la tension de sortie du cristal varie dans un rapport de 1 à 2,5.

Le coefficient de régulation et la zone d'écarts possibles par rapport à la fréquence de référence varient dans la même proportion. Dans la plupart des cas, cela n'a pas grande importance ; toutefois, si on désire rattraper cette variation de pente, il suffit d'agir sur le couplage des cavités au guide d'onde qui conduit l'énergie SHF.

Si on fait abstraction de cette dernière considération, une seule commande suffit pour caler instantanément l'émetteur sur une fréquence quelconque incluse dans une bande de 150 Mc/s.

La figure 13 montre la partie hyperfréquence du correcteur munie de son dispositif de commande unique.

APPENDICE II

Nous avons envisagé dans tout ce qui précède le cas d'un correcteur automatique de fréquence purement électronique ; notons qu'avec un tube hyperfréquence muni d'une commande mécanique d'accord, on peut, moyennant une importante

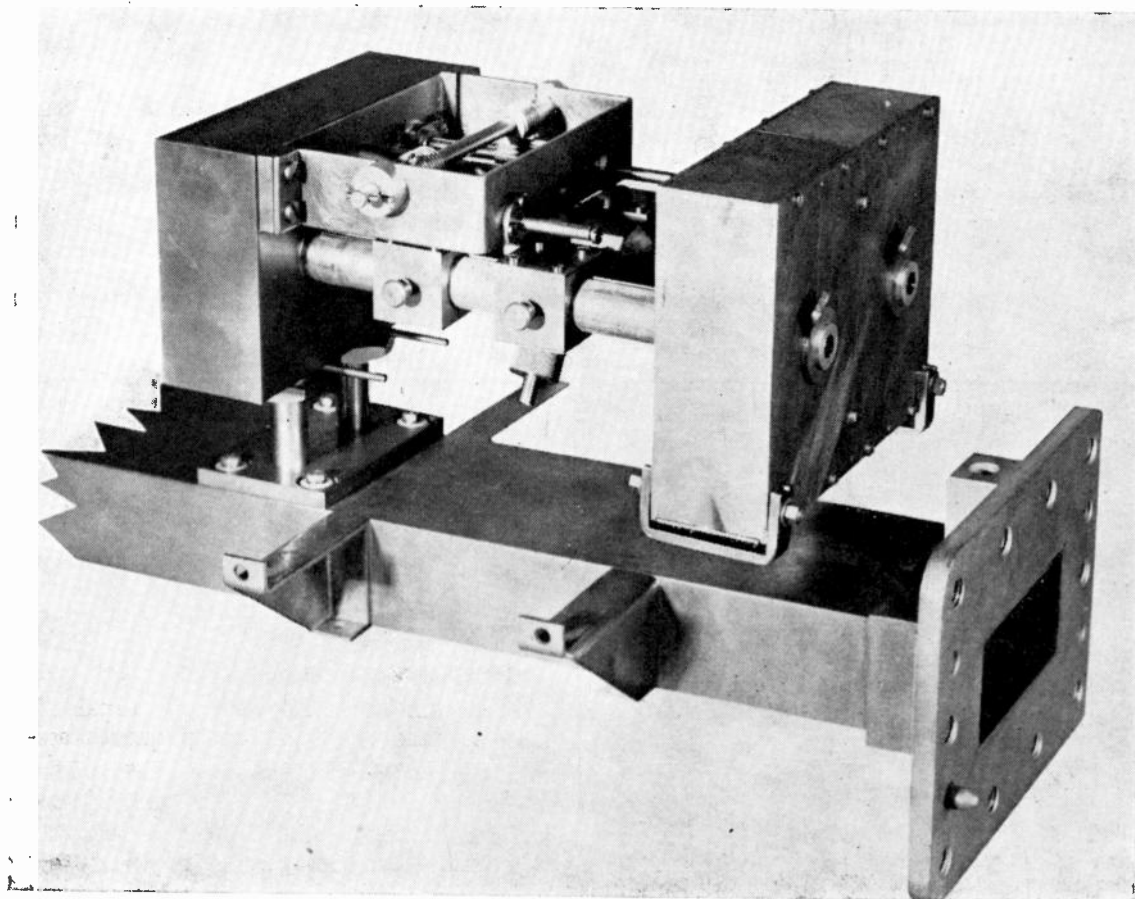


FIG. 13

Il a été mentionné ci-dessus que l'ensemble constitué par le correcteur automatique de fréquence était placé dans une enceinte dont la température était commandée par un thermostat. Pour que le réglage de longueur d'onde soit valable sans risque de modification ultérieure, il doit être effectué à la température de fonctionnement qui est de 65 degrés.

A cet effet a été prévu un dispositif tel qu'il ne soit pas nécessaire d'ouvrir l'enceinte à température constante. La commande unique est assurée par un arbre isolant qui traverse la paroi de cette enceinte ; les joints sont étudiés de façon à introduire un minimum de pertes thermiques.

simplification de la partie basse fréquence, adjoindre un moteur d'asservissement à ce dispositif.

La tension qui prend naissance à la sortie du cristal détecteur du bloc hyperfréquence est variable en amplitude et en phase en fonction de l'importance et du signe de l'écart de fréquence par rapport à une fréquence de référence. Cette tension est donc, après une amplification convenable, directement applicable à l'un des enroulements d'un moteur diphasé d'asservissement dont le sens de rotation est réversible ; l'autre enroulement est alimenté en permanence par un courant alternatif à 50 périodes d'amplitude et de phase constantes.

LA MODULATION D'ESPACEMENT D'IMPULSIONS

PAR

G. X. POTIER

Ingénieur à la Société Le Matériel Téléphonique

L'objet de la présente communication est d'indiquer les avantages et inconvénients de la modulation d'espacement d'impulsions et, plus particulièrement, de montrer comment certains défauts peuvent être éliminés à l'aide de dispositifs enregistreurs utilisés tant à l'émission qu'à la réception.

1. — Définition de la modulation d'espacement.

Afin de préciser ce que l'on entend par modulation d'espacement, examinons, dans un train d'impulsions de liaison multiplex, les positions respectivement occupées par l'impulsion de synchronisation définissant le début d'un train d'impulsions et les deux premières impulsions de voie. Le procédé de modulation le plus couramment utilisé consiste à faire varier linéairement, en fonction du signal de modulation, la position des impulsions de voie. Pour chacune des impulsions, cette position est définie à partir de l'impulsion de synchronisation prise comme origine des temps.

Supposons que nous utilisions un tel procédé de transmission et que seule la première impulsion de voie soit modulée ; sa position variera au rythme de la modulation, tandis que la position de la seconde impulsion restera immuable. Le temps séparant les deux impulsions sera proportionnel à la valeur instantanée du signal de modulation. Si, maintenant, la seconde impulsion est à son tour modulée, le temps séparant les deux impulsions sera fonction de la différence des valeurs instantanées des deux signaux de modulation.

La modulation d'espacement consiste à faire en sorte que le temps séparant deux impulsions consécutives soit proportionnel à la valeur instantanée d'un seul signal de modulation lorsque les deux impulsions sont simultanément modulées par deux signaux distincts.

Pour satisfaire une telle condition, il est nécessaire que le temps séparant les deux impulsions reste constant lorsque seule la première d'entre elles est modulée. En d'autres termes, il faut que le signal de modulation de la première impulsion provoque non seulement le déplacement de celle-ci mais également un déplacement identique de toutes les impulsions qui lui font suite.

Si nous utilisons la modulation d'espacement dans un équipement multiplex, la position d'une impulsion de rang « n » sera déterminée, en l'absence de modulation sur la voie correspondante, par la somme des valeurs instantanées des signaux présents dans les $(n-1)$ voies précédentes. C'est à partir de cette position que le signal appliqué à la voie « n » peut provoquer un déplacement de l'impulsion correspondante et de toutes celles qui la suivent.

2. — Avantages de la modulation d'espacement.

Malgré la différence signalée ci-dessus, la modulation d'espacement s'apparente étroitement à la modulation de position et il ne semble pas qu'elle puisse jouir de propriétés dont cette dernière serait dépourvue.

Divers auteurs ont cependant, pour des raisons différentes que nous allons maintenant examiner, envisagé l'utilisation de ce procédé de modulation.

2.1 Diaphonie par la voie précédente.

On sait que la limitation de la bande passante d'un circuit transmettant des impulsions provoque un traînage de celles-ci. Au moment de l'apparition d'une impulsion, il subsiste un résidu de l'impulsion précédente provoquant une modification de la forme de l'impulsion transmise. Lorsque l'impulsion précédente est modulée, la valeur du résidu varie et il en résulte un changement de la forme de l'impulsion au rythme de la modulation de l'impulsion précédente. Ce phénomène bien connu se traduit dans les liaisons multiplex à impulsions par une diaphonie dite de traînage qui, pour un espacement donné entre les impulsions, est d'autant plus importante que la bande passa

Le premier brevet décrivant la modulation d'espacement est, à notre connaissance, celui déposé par MM. REEVES et CHATTERJEA (1) qui imaginèrent ce procédé de modulation pour réduire la diaphonie de traînage.

Les auteurs font très justement remarquer que la déformation d'une impulsion par celle qui l'a précédée n'est pas gênante en elle-même ; seules

les modifications de cette déformation créent la diaphonie. La diaphonie de traînage peut par conséquent être supprimée si la déformation apportée par une impulsion modulée est constante. Cette condition est satisfaite lorsque le temps séparant deux impulsions consécutives reste constant en présence d'une modulation appliquée à la première d'entre elles. Un procédé de modulation présentant cette particularité fournit obligatoirement des impulsions dont l'espacement est proportionnel à la valeur instantanée d'un signal de modulation.

L'utilisation de la modulation d'espacement permet donc, pour un taux de diaphonie donné, de réduire la bande passante nécessaire à une valeur que ne permet pas d'atteindre la modulation de position.

On peut faire remarquer que la diaphonie de traînage réapparaît lorsque la seconde d'une paire d'impulsions est modulée. Le trouble ainsi créé est cependant peu gênant car il ne fait que provoquer une distorsion du signal transmis qui est bien inférieure à celle habituellement admise. Ce défaut sera examinée de façon plus détaillée dans un prochain paragraphe.

2.2 Quantité d'information transmise.

Les impulsions modulées en position subissent habituellement, sous l'influence du signal de modulation, un déplacement de part et d'autre de la position occupée en l'absence de modulation. Le temps séparant deux impulsions consécutives doit par conséquent être tel que des déplacements opposés des deux impulsions puissent se produire simultanément sans que les impulsions risquent d'être confondues.

Avec la modulation d'espacement, il est inutile de réserver entre deux impulsions le temps nécessaire au retard de la première impulsion puisque celui-ci provoque un retard identique de la seconde impulsion, évitant ainsi tout risque de confusion.

Examinons maintenant la position occupée par la dernière impulsion de voie d'une liaison multiplex en absence de signal de modulation sur l'ensemble des voies. Lorsque l'on utilise la modulation de position, le temps disponible entre cette impulsion et l'impulsion de synchronisation suivante permet un déplacement égal au déplacement maximum auquel un signal de modulation peut donner naissance.

Si l'on utilise la modulation d'espacement, tout retard d'une impulsion de voie est transmis aux suivantes et la dernière impulsion subit un retard égal à la somme de ceux provoqués par les divers signaux de modulation. Dans ces conditions, il semble normal que le temps disponible avant l'impulsion de synchronisation soit égal à celui prévu en modulation de position multiplié par le nombre d'impulsions. S'il en est ainsi, les deux procédés de modulation transmettent rigoureusement la même quantité d'information.

Cependant, dans les conditions normales d'utilisation d'un équipement multiplex, il n'y a aucune relation de phase entre les divers signaux de modulation. La probabilité d'un retard maximum subi

simultanément par toutes les impulsions est pratiquement nulle et le temps précédant l'impulsion de synchronisation indiqué plus haut possède une valeur inutilement grande. Sa réduction permet d'accroître la quantité d'information transmise soit en augmentant le nombre de voies, soit en augmentant le déplacement des impulsions et par conséquent le rapport signal/bruit.

En utilisant la modulation d'espacement telle qu'elle a été décrite jusqu'ici, le temps séparant deux impulsions consécutives est au moins égal à la moitié de celui existant en modulation de position et toute augmentation du déplacement maximum des impulsions nécessitera une augmentation identique du temps séparant les impulsions en l'absence de modulation. L'amélioration du rapport signal/bruit restera, de ce fait, inférieure à 3 dB.

Une augmentation plus importante du rapport signal/bruit peut être obtenue lorsque le temps séparant deux impulsions consécutives est, en l'absence de modulation, réduit à une valeur juste suffisante pour assurer le fonctionnement correct des sélecteurs. Ceci exige que le signal de modulation ne puisse que retarder l'apparition de l'impulsion. Ce procédé de modulation, dit à composante continue incorporée, ne présente pas de difficulté et a déjà été envisagé pour d'autres applications [2].

La combinaison de la modulation d'espacement et de la transmission avec composante continue incorporée permet de transmettre une quantité d'information sensiblement supérieure à celle transmise en utilisant la modulation de position.

3. — Inconvénients de la modulation d'espacement.

Après avoir examiné les avantages procurés par la modulation d'espacement, il convient de s'assurer que son utilisation ne soulève pas des difficultés la rendant pratiquement inutilisable.

3.1. — Sélection des différentes voies de transmission.

La première remarque que l'on peut faire concerne la sélection des impulsions dans l'équipement d'arrivée. Habituellement cette sélection s'effectue en deux temps de la manière suivante : un premier sélecteur isole d'abord l'impulsion de synchronisation, puis celle-ci décalée dans le temps, d'une valeur convenable, permet d'effectuer la sélection des diverses impulsions de voie. Cette méthode parfaitement adaptée à la modulation de position ne peut être utilisée avec la modulation d'espacement.

Lorsque l'on utilise la modulation d'espacement on ignore a priori la position qui sera occupée par l'impulsion correspondant à une voie déterminée. La sélection doit donc reposer sur une autre caractéristique. On ne peut songer à différencier entre elles les impulsions de voie par des amplitudes ou durées différentes. Pratiquement la seule méthode utilisable consiste à s'intéresser au rang occupé par la voie et à effectuer la sélection à l'aide d'un compteur d'impulsions. Le compteur doit posséder autant de sorties qu'il y a de canaux. Le plus souvent chacune d'entre

elles transmet une impulsion dont la durée est égale au temps séparant deux impulsions consécutives occupant un rang déterminé et toujours le même.

Le compteur électronique permet donc d'effectuer la sélection de chaque canal de transmission et transforme la modulation d'espacement en modulation de durée. Cette transformation est avantageuse car les impulsions modulées en durée possèdent une composante à la fréquence du signal de modulation beaucoup plus importante que les impulsions modulées en position ou en espacement. Une telle transformation est d'ailleurs couramment utilisée dans les liaisons par impulsions modulées en position.

3.2. — Bruit d'origine diaphonique.

Habituellement la reconstitution du signal s'effectue en appliquant les impulsions modulées en durée à un filtre passe-bas qui isole la composante basse fréquence. Dans le cas présent quelques difficultés surgissent du fait que les impulsions sont non seulement modulées en durée par le signal désiré mais également modulées en position par l'ensemble des signaux appliqués aux voies précédentes. Cette modulation de position possède une composante basse-fréquence qui sera isolée par le filtre au même titre que le signal désiré et donnera naissance à une diaphonie qu'il convient d'éliminer.

La composante à la fréquence de modulation F_m existant dans le spectre d'impulsions d'amplitude A et de durée t , modulées en position et subissant un déplacement td a une amplitude :

$$ABF = A t F_m 2 \pi \frac{td}{T_R}$$

en appelant T_R la période de répétition des impulsions. En modulation de durée, si l'on désigne par td_1 le déplacement maximum du front modulé ; l'amplitude du signal BF est $ABF = A \frac{td_1}{T_R}$. La diaphonie inséparable d'un filtrage direct des impulsions modulées simultanément en durée et en position est donnée par le rapport des amplitudes ci-dessus. Sa valeur est :

$$D = \frac{td_1}{td} \frac{1}{2 \pi F_m t}$$

Si nous considérons la dernière voie d'un équipement transmettant 60 communications, nous aurions approximativement : $\frac{td_1}{td} = 3$ et la durée t de l'impulsion en l'absence de modulation ne serait pas sensiblement inférieure à $1 \mu s$. Dans ces conditions, pour une fréquence de modulation de 3 000 c/s l'écart diaphonique serait de l'ordre de 25 dB.

Pratiquement le déplacement de l'avant-dernière impulsion n'est pas proportionnel à un signal de modulation mais à la somme de ceux appliqués aux différentes voies. Le signal perturbateur se présente donc sous la forme d'une diaphonie incompréhensible, comparable à un bruit.

Les tolérances admises pour ce type de diaphonie sont moins sévères que celles données pour la diaphonie compréhensible. La valeur de l'écart diaphonique indiquée ci-dessus est cependant trop importante et rendrait la modulation d'espacement inutilisable si, ainsi que nous le montrerons dans un paragraphe suivant, il n'était pas possible de supprimer le défaut indiqué ici.

3.3. — Distorsion de traînage.

Nous avons déjà fait remarquer au paragraphe 2.1 qu'une impulsion modulée était affectée d'une déformation due à la présence de l'impulsion précédente.

Un trouble ne se manifestant qu'en présence d'un signal de modulation ne peut être assimilé ni à une diaphonie ni à un bruit, mais à une distorsion d'un type particulier. Cette façon de la présenter est ici d'autant plus justifiée que le signal indésirable étant provoqué par le déplacement de l'impulsion transmettant le signal se modifiera au rythme de ce dernier.

Le passage de la modulation de position à la modulation d'espacement transforme la diaphonie de traînage en une distorsion de traînage. La valeur de cette dernière peut sans inconvénient être sensiblement plus grande que celle de la diaphonie. Une diaphonie de 40 dB est en effet inacceptable tandis qu'une distorsion de même valeur (1 %) se combinant avec celle existant normalement dans un canal de transmission n'affecte pratiquement pas la qualité de la transmission.

4. — Possibilités offertes par la modulation d'espacement.

Pour préciser les améliorations procurées par la modulation d'espacement, il est commode d'examiner comment les possibilités données par un équipement à modulation de position seraient modifiées par l'utilisation de la modulation d'espacement.

Partant d'un équipement à 24 voies, dont les caractéristiques sont approximativement les mêmes chez tous les constructeurs, nous chercherons tout d'abord dans quel rapport il est possible de réduire la bande passante nécessaire à la transmission en conservant le même nombre de voies, la même puissance moyenne et le même rapport signal/bruit.

Après avoir montré la réduction de bande passante qui peut être obtenue pour transmettre un nombre de voies donné, nous déterminerons l'accroissement possible du nombre de voies sans modification de la bande passante, de la puissance moyenne ni du rapport signal/bruit de chaque voie.

4.1. — Réduction de la bande passante nécessaire à la transmission d'un nombre de voies téléphoniques donné.

Soit un équipement multiplex téléphonique à 24 voies utilisant 25 impulsions de $0,5 \mu s$ de durée se déplaçant de $\pm 1,5 \mu s$. La période de répétition

est de 125 μ s et le temps de garde entre deux impulsions consécutives est de 1,5 μ s.

Le temps réservé à la transmission des impulsions et à leur temps de garde est de $25 \times (1,5 + 0,5) \mu$ s = 50 μ s. Ce temps est jugé nécessaire pour effectuer les opérations permettant l'identification des différentes voies. Le temps restant disponible, soit $25 \times 2 \times 1,5 = 75 \mu$ s est réparti entre les impulsions afin de permettre le déplacement de celles-ci, c'est-à-dire la transmission de l'information.

C'est ici qu'apparaît la différence essentielle existant entre la modulation de position et la modulation d'espacement.

Lorsque l'on utilise la modulation de position, le temps disponible est également réparti une fois pour toutes entre les impulsions.

En modulation d'espacement, le temps disponible est réparti uniquement entre les voies actives et proportionnellement à la valeur instantanée du signal de modulation. Le rapport du temps total nécessaire à la transmission de N voies au temps nécessaire pour une seule voie est égal au facteur de charge du signal multiplex utilisé dans la technique des courants porteurs.

On sait que le facteur de charge est le rapport de la valeur maximum du signal composite formé de la somme des signaux vocaux provenant de N lignes téléphoniques à la valeur maximum du signal provenant d'une seule ligne.

Les valeurs suivantes du facteur de charge peuvent être déduites de l'article de B.D. HOLBROOK et J. T. DIXON [3] :

Nombre de voies	Facteur de charge
12	1,9
24	2
48	2,2
60	2,4
120	2,65

Examinons maintenant les conséquences de la division de la bande passante par un coefficient k . La durée des fronts des impulsions sera multipliée par k , ce qui réduira le rapport signal/bruit ; pour redonner à celui-ci sa valeur première, il faudra multiplier le déplacement des impulsions par k . D'autre part, pour se trouver dans les mêmes conditions de fonctionnement du récepteur, la durée des impulsions et le temps de garde devront être multipliés par k . Afin que cette augmentation de la durée ne provoque pas une augmentation de la puissance moyenne, il faudra diviser par k la puissance transmise pendant les impulsions. Pour éviter que cette réduction de puissance diminue le rapport signal/bruit, il faudra à nouveau multiplier le déplacement des impulsions par un facteur qui, cette fois sera $\sqrt{k}$.

En résumé, si nous divisons par k la bande passante, le temps disponible réservé à la modulation de l'ensemble des impulsions deviendra $(125 - 50 k)$

μ s et, pour conserver le même rapport signal/bruit, il faudra que le déplacement individuel maximum de chaque impulsion soit porté à : $(3 k^{3/2}) \mu$ s, c'est-à-dire que le temps disponible pour la modulation de l'ensemble des impulsions par le nouveau procédé soit $2 \times 3 \times k^{3/2} \mu$ s = $(125 - 50 k) \mu$ s qui donne environ $k = 2,1$. La modulation d'espacement permet donc de diviser par 2,1 la bande passante nécessaire à la transmission de 24 voies téléphoniques.

4.2. — Augmentation du nombre de voies transmises.

La diminution de la bande passante permet d'effectuer des liaisons identiques à celles assurées actuellement en évitant un gaspillage de fréquences porteuses.

La recherche de l'accroissement du nombre de voies montre que la modulation d'espacement permet la transmission d'un nombre de voies qui ne peut pratiquement pas être envisagé en utilisant la modulation de position. Ces possibilités nouvelles ont été indiquées par MM. GLOESS et LABOIS [4].

Nous calculerons tout d'abord le nombre de voies qu'il est possible de transmettre en conservant le temps de garde des équipements actuels (1,5 μ s), puis réduirons celui-ci à une valeur plus faible mais restant compatible avec la bande passante et les possibilités de sélection.

Si, conservant la bande passante actuelle, nous multiplions le nombre de voies par n , il faudra diviser par n la puissance émise pendant les impulsions afin de conserver la même puissance totale émise. Dans ces conditions, le rapport signal/bruit se trouvera réduit et pour lui redonner sa valeur primitive il faudra multiplier le déplacement des impulsions par $\sqrt{n}$ bien que le temps disponible ne soit plus que $(125 - 50 n) \mu$ s.

Cette condition peut être satisfaite pour $n =$ légèrement supérieur à 2, grâce au fait que la valeur maximum de la somme de 48 signaux vocaux est égale à 2,2 fois la valeur maximum d'un seul signal vocal.

La modulation d'espacement appliquée aux équipements à impulsions couramment utilisés permet donc, soit de diviser par deux la bande passante utilisée, soit de doubler le nombre de communications tout en conservant le même rapport signal/bruit et la même puissance moyenne. En d'autres termes, nous dirons que l'utilisation de la modulation d'espacement permet de doubler la quantité d'information transmise sans modification de la capacité du canal utilisé.

Il est intéressant de calculer le nombre maximum de voies que la modulation d'espacement permettrait de transmettre sans accroissement de la bande passante et de la puissance moyenne actuellement utilisées pour la transmission de 24 voies téléphoniques. Le temps de garde minimum devant exister entre deux impulsions est égal à la durée de celles-ci, soit dans le cas présent 0,5 μ s. La réduction du temps de garde permettrait d'envisager la transmission de 96 voies.

Le nombre de communications étant multiplié par 4, il faudrait diviser la puissance de crête dans le même rapport et multiplier le déplacement par deux pour conserver le même rapport signal/bruit. Le déplacement de l'ensemble des impulsions étant égal, pour 96 voies, à 2,6 fois celui d'une seule impulsion, il est nécessaire de disposer de $3.2.2,6 = 15,6 \mu s$. Le temps disponible étant de $125 - 96 = 29 \mu s$ il est possible de conserver le rapport signal/bruit obtenu actuellement avec les liaisons 24 voies.

5. — Exemples de modulateur et démodulateur d'espacement d'impulsions.

On peut considérer que chacune des impulsions provenant d'un équipement multiplex pour lequel on utilise la modulation d'espacement a subi une double modulation de position. La première, provoquée par le déplacement des impulsions précédentes constitue une modulation indésirable à éliminer ; la seconde due à la présence d'un signal doit seule servir à reconstituer le signal de modulation original.

Ainsi qu'il a été expliqué au paragraphe 3.1, l'utilisation d'un compteur électronique pour effectuer la sélection des impulsions permet de transformer la modulation de position, provoquée par le signal à transmettre, en une modulation de durée tandis que la modulation par les voies précédentes continue à se présenter sous la forme d'une modulation de position.

Ceci constitue une première étape permettant de différencier les deux modulations et présente de surcroît l'avantage de réduire l'importance de la perturbation apportée par la modulation indésirable. Cependant, ainsi qu'il a été indiqué dans le paragraphe 3.2, le trouble résiduel est encore trop important.

Le dispositif de démodulation décrit ci-dessous [5] permet d'éliminer le bruit d'origine diaphonique. Le fonctionnement d'un modulateur évitant certains défauts pouvant accompagner la modulation d'espacement est également examiné.

5.1. — Démodulateur.

Le bruit d'origine diaphonique est provoqué par la modulation de position dont sont affectées les impulsions modulées en durée, disponibles en sortie, du compteur électronique de sélection. Sa suppression exige que l'on applique au filtre passe-bas de démodulation des impulsions ne subissant pas une telle modulation, c'est-à-dire équidistantes. Ces impulsions doivent pouvoir être obtenues à partir de celles fournies par le compteur-sélecteur.

Le principe suivant peut être utilisé. Après avoir créé une tension dont la valeur est proportionnelle à la durée des impulsions modulées, on conserve cette tension qui sera transmise au filtre par l'intermédiaire d'une porte électronique pendant la présence d'une impulsion de commande. Les impulsions de commande étant équidistantes, il en sera de même de celles appliquées au filtre. Cette façon de

procéder transforme des impulsions modulées simultanément en durée et en position en impulsions modulées uniquement en amplitude.

Le dispositif d'enregistrement utilisé doit prendre connaissance de l'information en un temps qui ne doit pas être supérieur à une microseconde. Il doit conserver le souvenir de la tension enregistrée, avec une précision supérieure à 1/1 000, pendant une durée de l'ordre de cent microsecondes. Le souvenir de la tension enregistrée doit pouvoir être effacé en quelques microsecondes. Un tel dispositif ne peut être qu'électronique. Le fait qu'il en existe un exemplaire par voie oblige à concevoir un appareil aussi simple que possible.

Un mode d'enregistrement pouvant être utilisé consiste à charger un condensateur à l'aide d'un courant constant pendant la durée des impulsions modulées en durée. La tension disponible aux bornes d'un condensateur C après une charge de durée t

par un courant I sera : $V = \frac{I}{C} t$; elle sera donc pro-

portionnelle à la valeur instantanée du signal de modulation. Après la transmission de l'information au filtre, des impulsions de commande sont appliquées à un circuit qui provoque la décharge rapide du condensateur.

Les impulsions assurant la commande de la porte électronique et la décharge de la capacité d'enregistrement sont obtenues à partir des impulsions

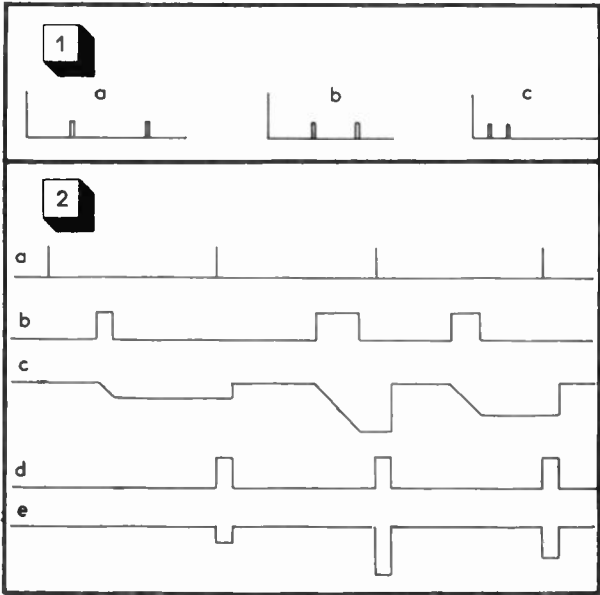


FIG. 1. — Positions occupées par les deux premières impulsions de voies en absence de modulation.

- a) modulation de position.
- b) modulation d'espacement avec déplacement bidirectionnel.
- c) modulation d'espacement avec déplacement unidirectionnel.

FIG. 2. — Suppression de la modulation de position indésirable.

- a) Positions occupées par les impulsions de synchronisation.
- b) Impulsions modulées simultanément en position et en durée fournies par le compteur-sélecteur.
- c) Tension aux bornes de la capacité d'enregistrement.
- d) Impulsions de commande de la porte électronique.
- e) Impulsions modulées en amplitude transmises au filtre de démodulation.

de synchronisation qui sont les seules se répétant d'une façon rigoureusement périodique.

La forme des signaux obtenus en différents points du démodulateur sont reproduits sur la figure 2 qui facilite la compréhension du fonctionnement de celui-ci.

5.2 Modulateur.

Il convient maintenant de s'assurer que le procédé utilisé pour éliminer la diaphonie ne donne pas naissance à d'autres défauts. Le taux de distorsion, en particulier, peut être affecté par les opérations de démodulation. La méthode préconisée ci-dessus ne provoquera qu'un retard constant, et par conséquent sans importance, si les déplacements des impulsions provoqués par les modulateurs sont proportionnels aux valeurs instantanées prises par le signal de modulation à des instants équidistants. Bien que cette condition soit satisfaite dans les équipements où les impulsions occupent toujours la même position en absence de modulation, il n'en est pas forcément de même ici.

Le procédé de modulation le plus simple consiste en effet à appliquer au modulateur l'impulsion provenant de la voie précédente, laquelle définit l'instant de référence pour la voie considérée, en même temps que le signal à transmettre. En opérant de cette manière le déplacement de l'impulsion est proportionnel à la valeur prise par le signal de modulation au moment où se présente l'impulsion de voie précédente, c'est-à-dire à des instants non équidistants. Une transmission effectuée dans de telles conditions sera inévitablement accompagnée de perturbations d'origine diaphonique. On remarquera toutefois que le signal indésirable ainsi créé ne se manifeste qu'en présence de modulation, il s'apparente donc à une distorsion ou au bruit de quantification des équipements à impulsions codées.

La suppression d'un tel défaut ne s'avère pas indispensable ; cependant, si l'on désire l'éliminer, le procédé d'enregistrement utilisé à la démodulation permet d'atteindre ce résultat. Cette fois la valeur enregistrée sera celle du signal à transmettre, l'opération sera effectuée à des instants équidistants coïncidant avec les impulsions de synchronisation. Chaque grandeur enregistrée sera utilisée par le modulateur à un instant défini par l'impulsion de

voie précédente. L'effacement d'une grandeur enregistrée s'effectuera un court instant avant l'apparition de l'impulsion de synchronisation suivante.

6. — Conclusion.

Pour résumer ce qui précède concernant la modulation d'espacement, nous dirons que :

1° Elle permet d'effectuer des transmissions par impulsions en utilisant de façon rationnelle le temps disponible pour l'ensemble des communications.

2° La quantité d'informations transmise par ce procédé est multipliée par un facteur qui, dans les conditions habituelles de transmission, est au moins égal à deux.

3° Le principe même de la transmission fait apparaître un bruit d'origine diaphonique dont l'importance est telle qu'il risque de rendre ce procédé de transmission inexploitable.

4° L'utilisation d'enregistreurs permet de supprimer ce défaut.

Cette dernière précaution étant prise, la modulation d'espacement permet de faire bénéficier les liaisons par impulsions de l'effet statistique qui semblait jusqu'ici l'apanage des équipements multiplex à courants porteurs.

Elle permet d'envisager la construction d'équipements à impulsions transmettant une centaine de communications avec une qualité identique à celle obtenue actuellement avec un nombre de voies nettement plus faible et ceci sans accroissement sensible de la bande passante utilisée.

BIBLIOGRAPHIE

1. — P. K. CHATTERJEA et A. H. REEVES. — Brevet français 1.014.487 du 20 Mars 1947.
2. — BEATTY et SCULLY. — Brevet français 870.474 du 28 Février 1941.
3. — B. D. HOLBROOK et J. T. DIXON. — Load rating theory for multi-channel — *B.S.T.J.* — Octobre 1939.
4. — P. F. GLOESS et L. J. LIBOIS. — Brevet français 1.007.764 du 27 Mars 1948.
5. — G. X. POTIER. — Brevet français 1.045.582 du 7 février 1950.
6. — E. FITCH. — The Spectrum of modulated pulses — *JIEE*, Vol. 84, Partie III A, N° 13.

CIRCUIT DÉCLENCHEUR FERRORESONNANT PARALLÈLE A RÉACTION⁽¹⁾

PAR

J. GARCIA SANTESMASES

*Dr. Sc., Ing. E. S. E., Professeur à l'Université
de Madrid et Directeur de « l'Instituto
de Electricidad del C. S. I. C. »*

ET

M. RODRIGUEZ VIDAL

*Docteur ès-Sciences « Departamento
« Instituto de Electricidad del Consejo Superior
de Investigaciones Científicas », Madrid*

INTRODUCTION.

En partant du phénomène de la ferrorésonance en parallèle (bobine avec noyau de fer et condensateur en parallèle) l'un de nous a réalisé un circuit déclencheur ferrorésonnant qui a été l'objet d'une communication présentée au « Symposium on Automatic Digital Computation, National Physical Laboratory, 1953 » (Teddington, Angleterre) et qui a été l'objet aussi de quelques travaux [1]. Dans ce circuit on obtient une zone de deux tensions stables, lorsqu'il est alimenté par une source d'intensité constante, et on obtient des rythmes d'impulsions assez élevés, de l'ordre de 200 kc/s., avec l'emploi de ferrites [2].

Dans la présente communication nous nous proposons de décrire un nouveau circuit déclencheur dans lequel, si on applique bien la ferrorésonance en parallèle, on suit un chemin distinct de celui utilisé dans nos travaux antérieurs, puisque nous partons initialement d'un amplificateur magnétique à réaction.

La possibilité d'utilisation du noyau saturable, monté en forme d'amplificateur magnétique à réaction, comme circuit bi-stable, a été premièrement mis en évidence par A. S. FITZGÉRALD dans l'année 1949 [3] : FITZGÉRALD observa la bi-stabilité en étudiant expérimentalement la caractéristique d'un amplificateur magnétique à un seul noyau, avec

réaction extérieure, en fonction du nombre de spires de l'enroulement à réaction.

Lorsque l'enroulement à courant alternatif d'un amplificateur magnétique est shunté par un condensateur, les caractéristiques de l'amplificateur présentent une déformation qui peut être utilisée pour l'obtention d'une nouvelle zone de deux intensités stables au moyen d'une réalimentation convenable, lorsqu'on alimente par une source de tension constante.

De cette façon l'obtention d'un circuit avec deux zones bi-stables est possible, et au moyen d'un choix convenable des paramètres de ce circuit elles peuvent arriver à coïncider en donnant lieu à un circuit avec trois états stables.

CARACTÉRISTIQUES DU CIRCUIT FERRORESONNANT PARALLÈLE AVEC EXCITATION INDÉPENDANTE.

Soit un circuit ferrorésonnant parallèle tel qu'il est indiqué schématiquement sur la fig. 1 a).

Si nous remplaçons la self-inductance L avec noyau de fer par une self-inductance pure (dépendant de l'excitation $e = n_e I_e$ et de l'intensité qui traverse la bobine) en série avec une résistance égale à la résistance ohmique de la bobine et en parallèle avec une autre résistance r_1 représentative des pertes totales dans le fer et une capacité aussi en parallèle, englobée dans la C_1 , il nous reste un circuit équivalent indiqué par la fig. 1 b).

L'étude analytique du circuit complet se compli-

(1) Communication présentée au Congrès sur les Procédés d'Enregistrement le 8 avril 1954.

que extraordinairement étant donnée la difficulté de représenter d'une manière appropriée la self-induction L et la résistance r_1 par des fonctions de l'intensité d'excitation I_e et du courant I_L qui traverse la bobine.

Pour nous donner une idée qualitative du comportement du circuit, supposons-le dans sa forme idéale la plus simple : résistance ohmique de la bobine nulle ; résistance r_1 infinie (pertes dans le fer négligeables) et résistance de charge R nulle.

Lorsque dans ces conditions, on alimente le circuit avec une source de tension, l'intensité I_L qui traverse la bobine en fonction de l'intensité d'excitation, nous est donnée par une courbe de la forme indiquée par la $I_L(e)$ sur la fig. 2 a. L'intensité du courant à travers le condensateur en parallèle, due à la même tension sera constante et égale à ωVC_1 ; et l'intensité, I , totale nous est donnée par valeur absolue de la différence des deux intensités I_L et I_C car celles-ci sont en opposition de phase et on obtient la courbe $I_1(e)$.

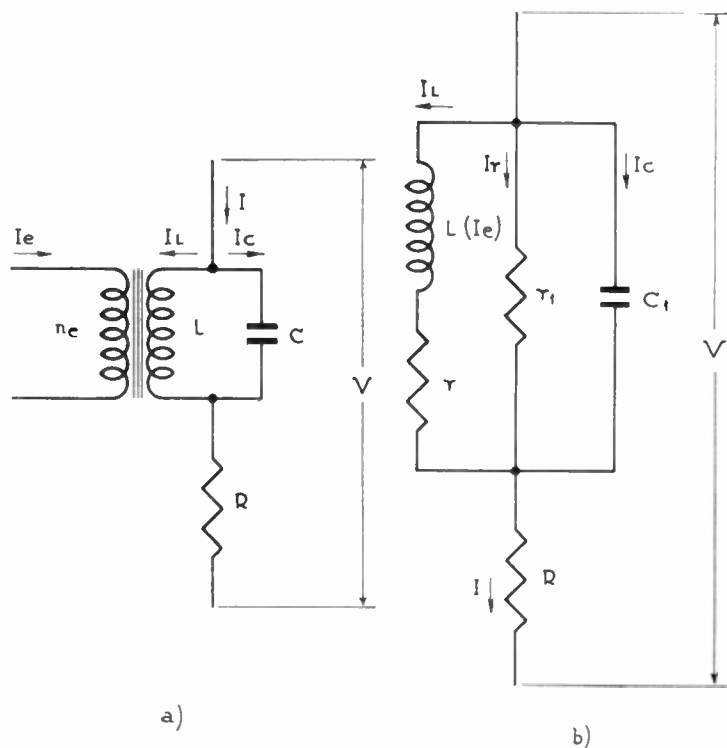


FIG. 1.

L'effet de la résistance r_1 , représentative des pertes dans le noyau, est d'ajouter à cette intensité I_L , une valeur $I_r = V/r_1$, en quadrature avec elle : et on obtient pour l'intensité totale :

$$I = (I_L^2 + I_r^2)^{\frac{1}{2}}$$

et si en première approximation nous supposons r_1 constante, nous obtenons comme représentation

de l'intensité totale I , la courbe indiquée sur la figure.

Si la capacité est suffisamment élevée pour que la droite $I = \omega VC_1$ coupe la courbe $I_L(e)$ on vérifie que : initialement, pour une excitation nulle, le courant a une composante capacitive ; à mesure que l'excitation augmente, le courant inductif va en augmentant, de sorte que le courant total diminue, jusqu'à ce que, au moment de la résonance, le courant inductif devienne égal et opposé au capacitif, ce pourquoi le courant total est uniquement résistif à ce moment et faible, étant en phase avec la tension d'alimentation. A partir de ce moment, le courant total a une composante inductive croissante avec l'excitation.

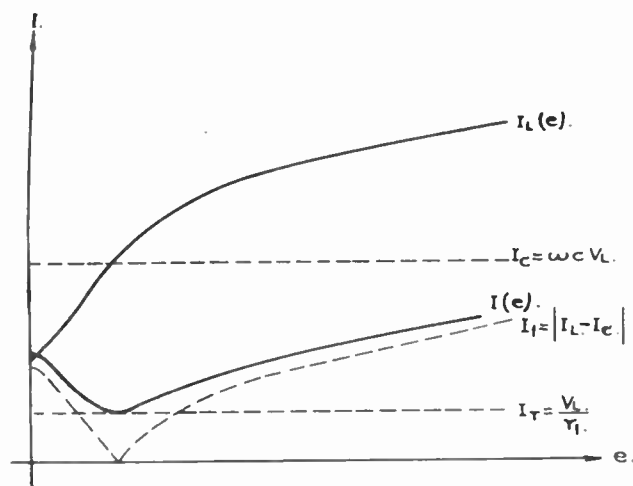


FIG. 2 a.

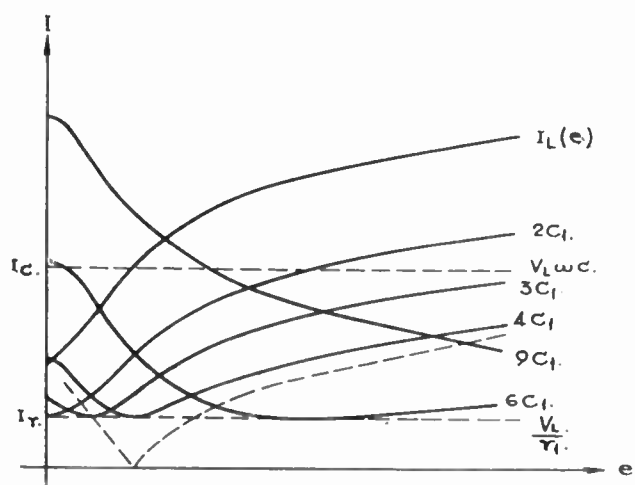


FIG. 2 b.

A mesure que la capacité augmente, le minimum d'intensité (résonance) s'obtient à des intensités d'excitation plus grandes, jusqu'à ce que pour des valeurs très grandes de celle-ci, on cesse d'obtenir la résonance. Sur la figure 2 b, on représente les graphiques obtenus au moyen de la construction antérieure, pour diverses capacités, qui sont qualitativement en accord avec les résultats expérimentaux (fig. 3).

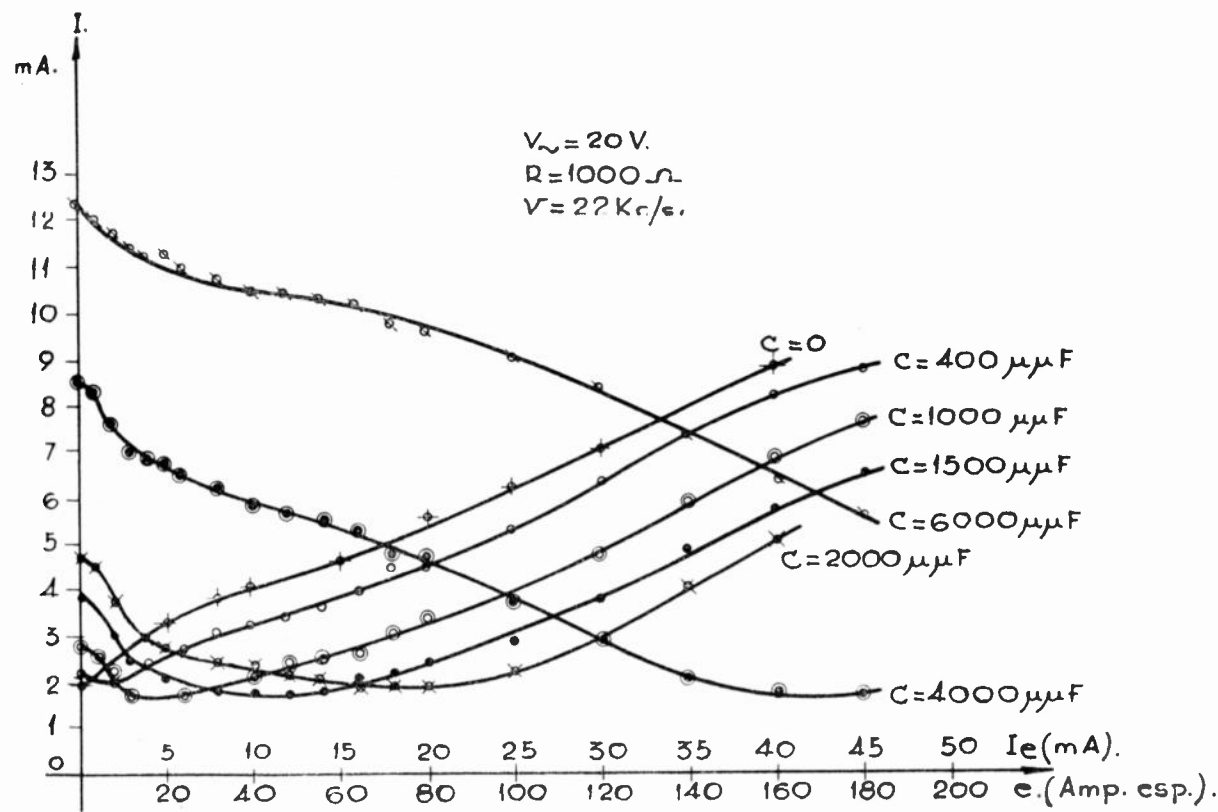


Fig. 3

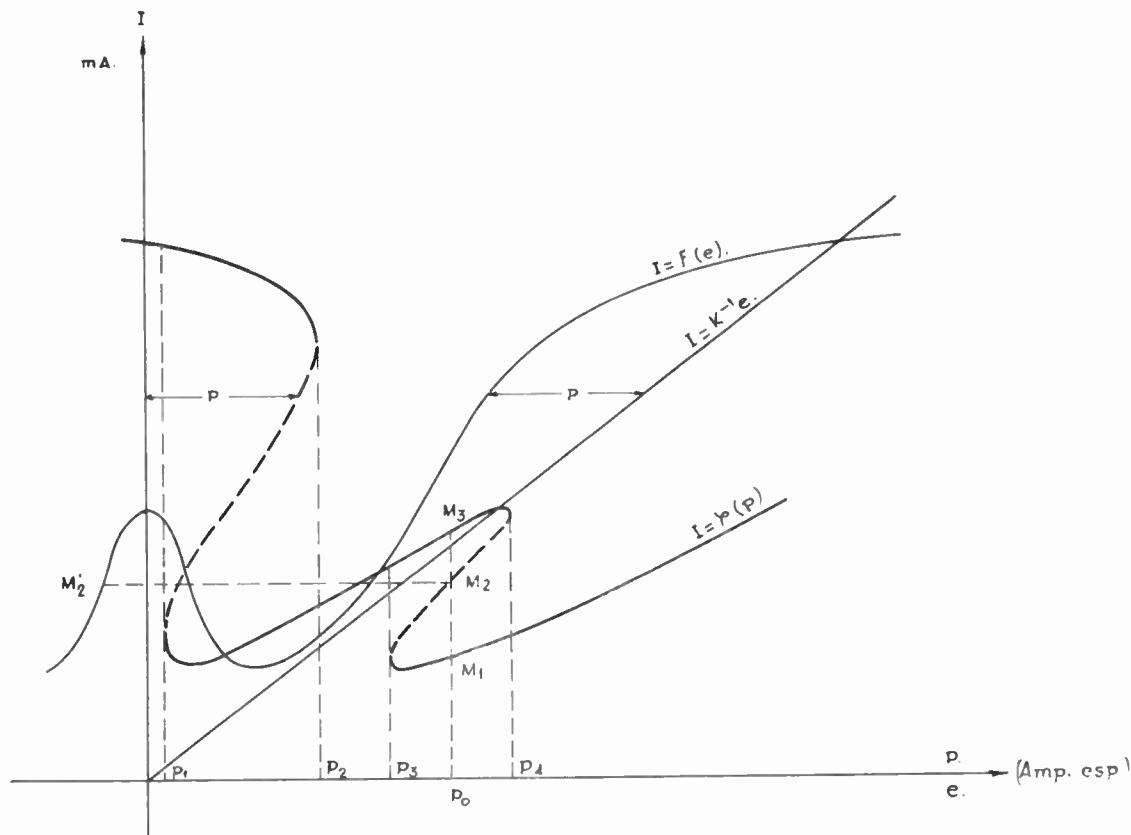


Fig. 4.

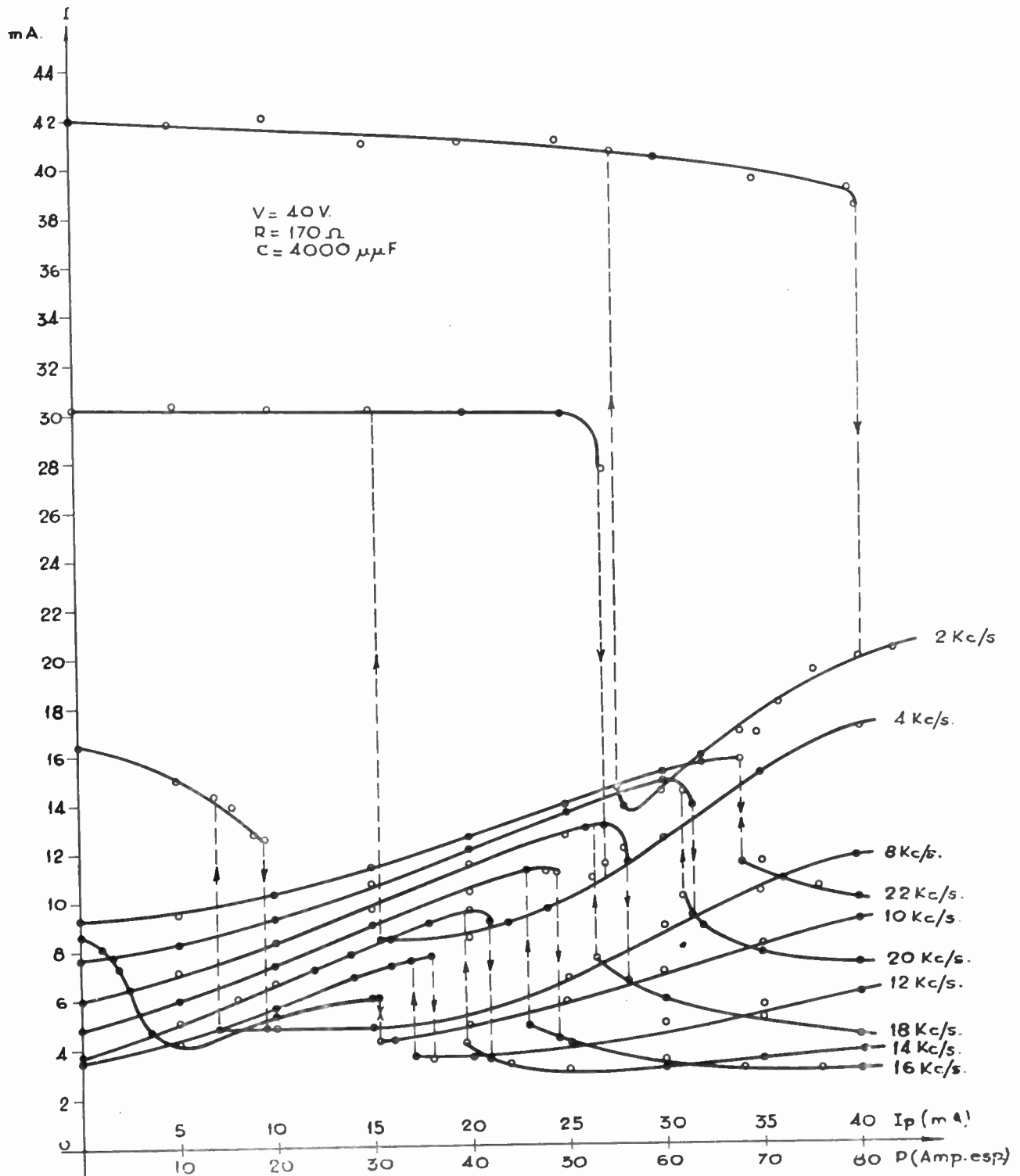


FIG. 6.

nulle on commence à noter les changements de forme des caractéristiques jusqu'à ce qu'on atteigne la seconde zone bi-stable. Cette seconde zone de bi-stabilité, due à la ferrorésonance du circuit alternatif, est obtenue chaque fois pour des intensités de polarisation plus grandes à mesure que nous augmentons la fréquence; il existe une fréquence pour laquelle on obtient les meilleures con-

ditions de travail autant pour la grandeur de la chute du courant que pour son ampleur; fréquence qui peut être facilement déplacée, sans faire plus que modifier la capacité du condensateur du circuit ferrorésonnant. Pour des fréquences au-dessus de celle-là, la zone bi-stable devient plus étroite, sans que la grandeur de la chute se modifie beau-

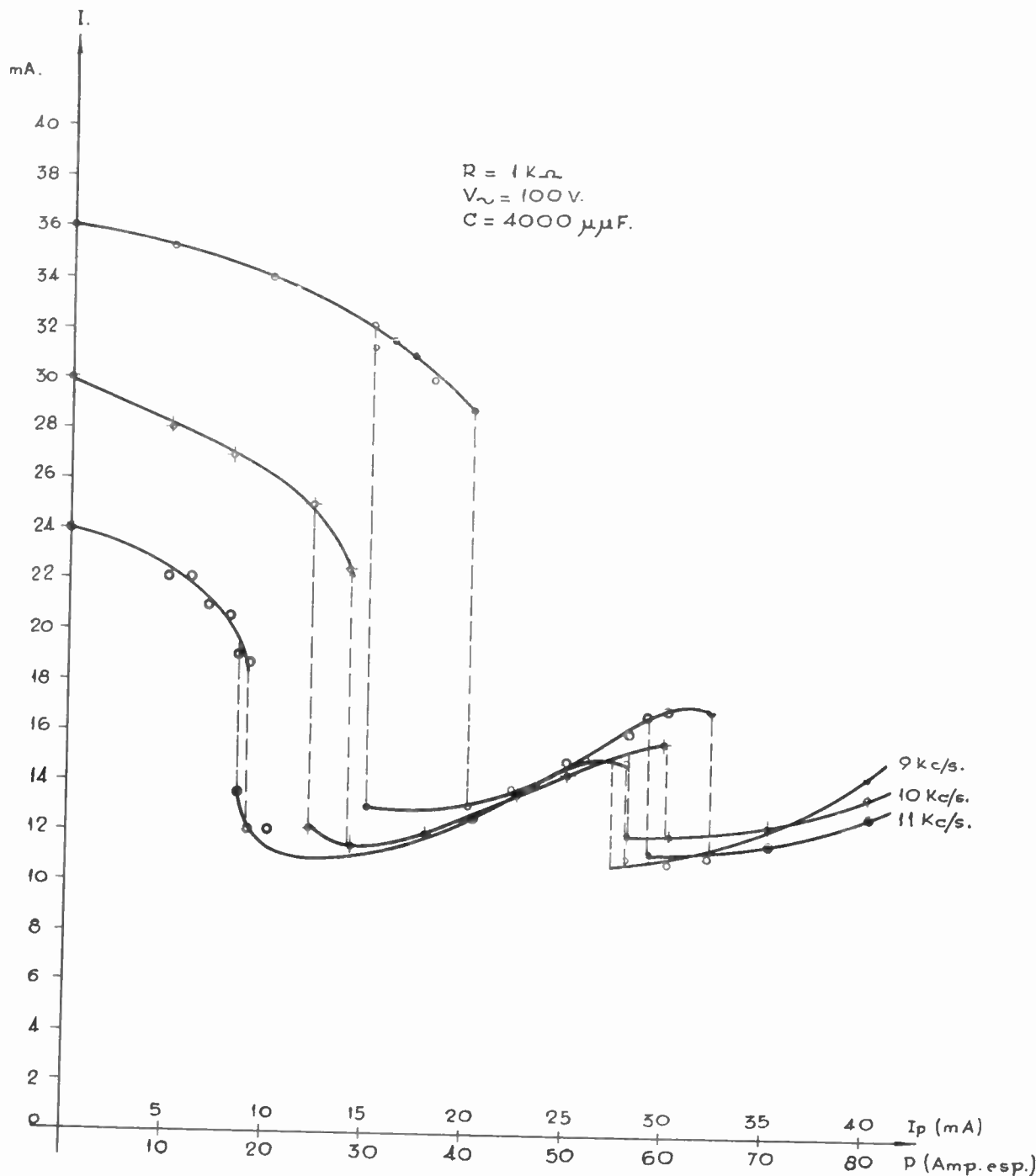


Fig. 7.

coup, jusqu'à ce qu'elle finisse par disparaître totalement.

Tout ce qui a été exposé est résumé dans les graphiques de la fig. 6 obtenus pour une résistance de charge de $200\text{ }\Omega$ et une capacité dans le circuit résonnant de $4\,000\text{ }\mu\mu\text{F}$, la tension d'alimentation étant de 40 volts efficaces.

Dans les graphiques de la fig. 7, on représente les deux zones bi-stables obtenues simultanément pour une même fréquence, quand la tension d'alimentation est élevée à 100 volts. Il est possible au moyen d'un choix convenable de la fréquence et avec une tension suffisamment élevée, de faire que les deux zones de bi-stabilité présentent un inter-

valle de polarisation commun, dans lequel le circuit a trois états stables.

b). — Influence de la capacité.

Sur la figure 8 sont résumés les résultats expérimentaux obtenus pour des capacités différentes, lorsque la tension d'alimentation est de 80 volts efficaces; la fréquence est de 22 Kc/s. et la résistance de charge de $1\text{ k}\Omega$.

Ces résultats sont d'accord avec les courbes qualitatives de la fig. 2 b; et on en déduit déjà l'existence d'une capacité optima, qui dans notre cas est approximativement de $1\,750\text{ }\mu\mu\text{F}$.

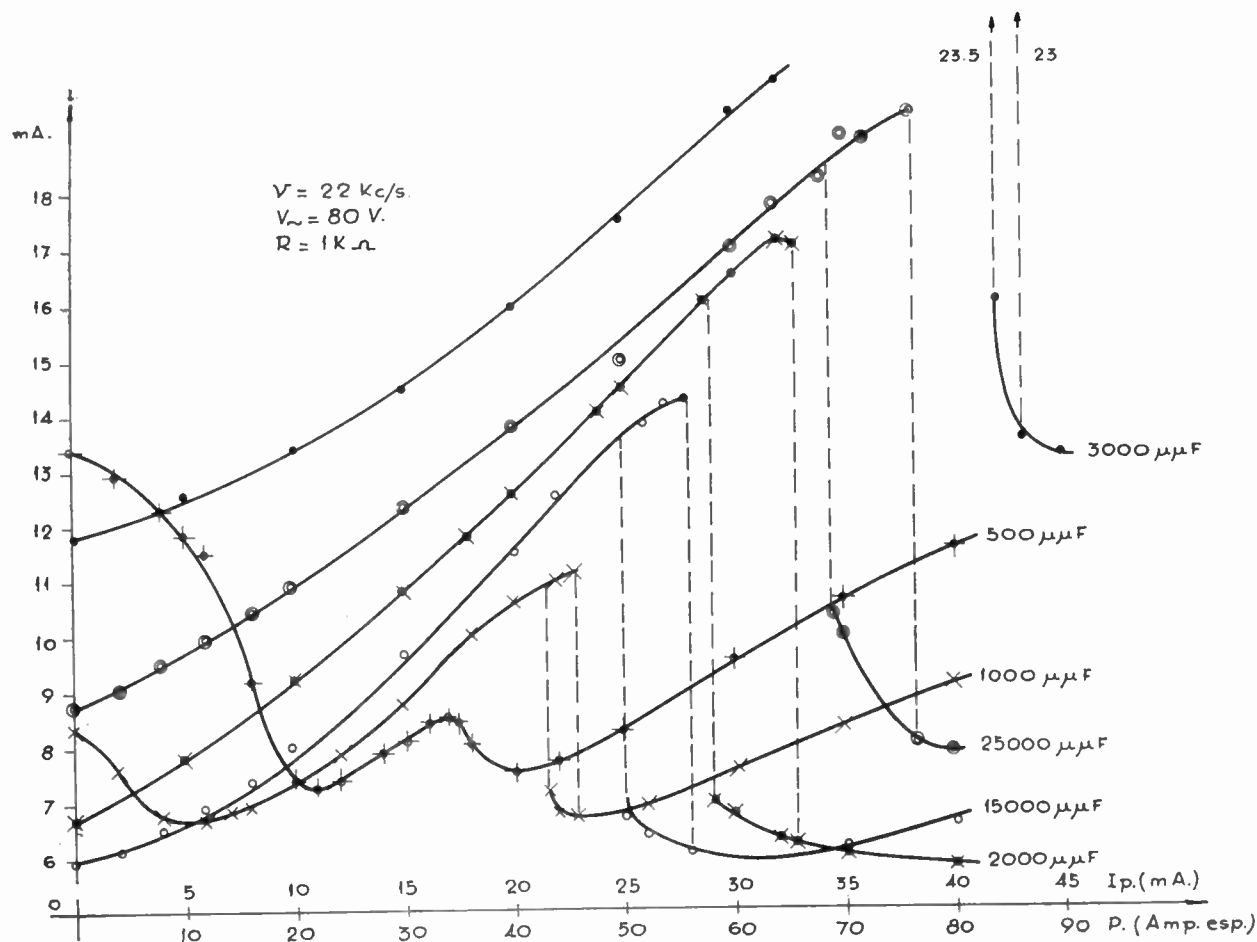


FIG. 8.

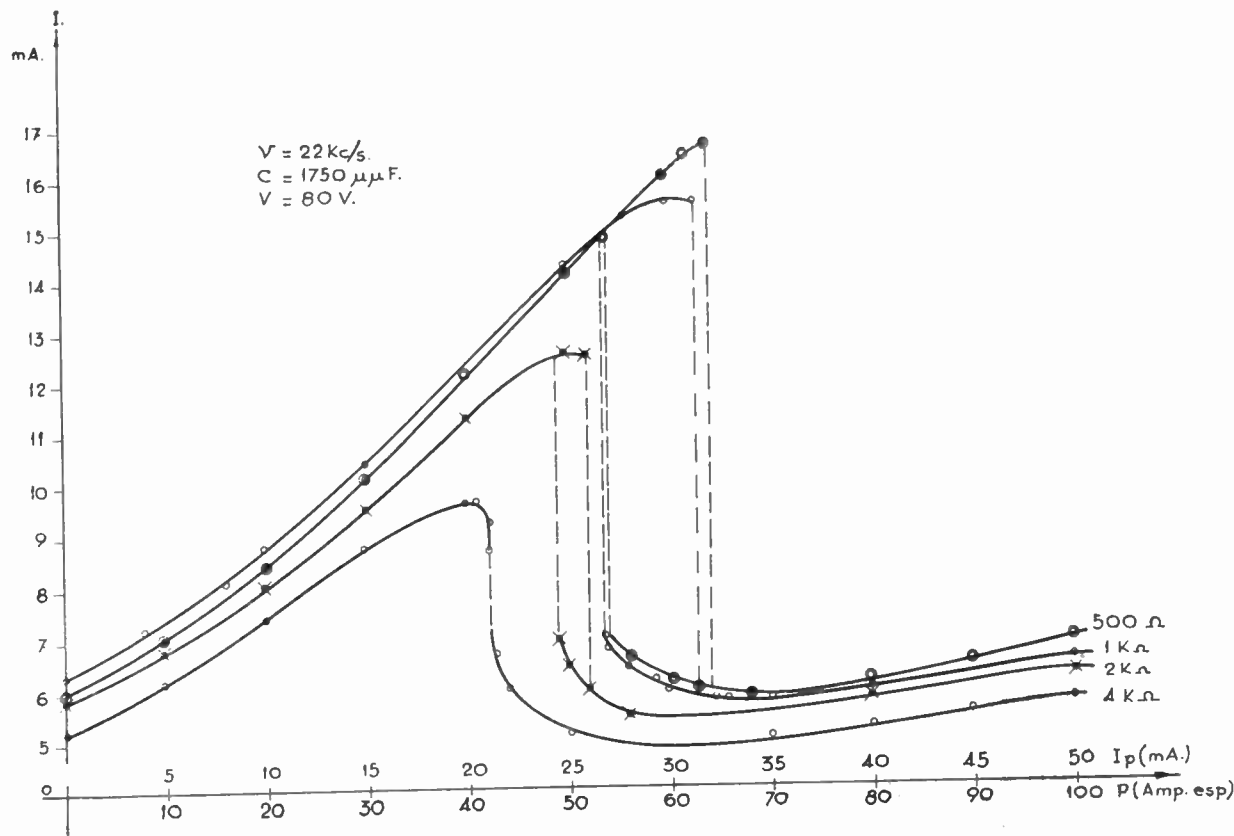


FIG. 9.

Fig. 11.

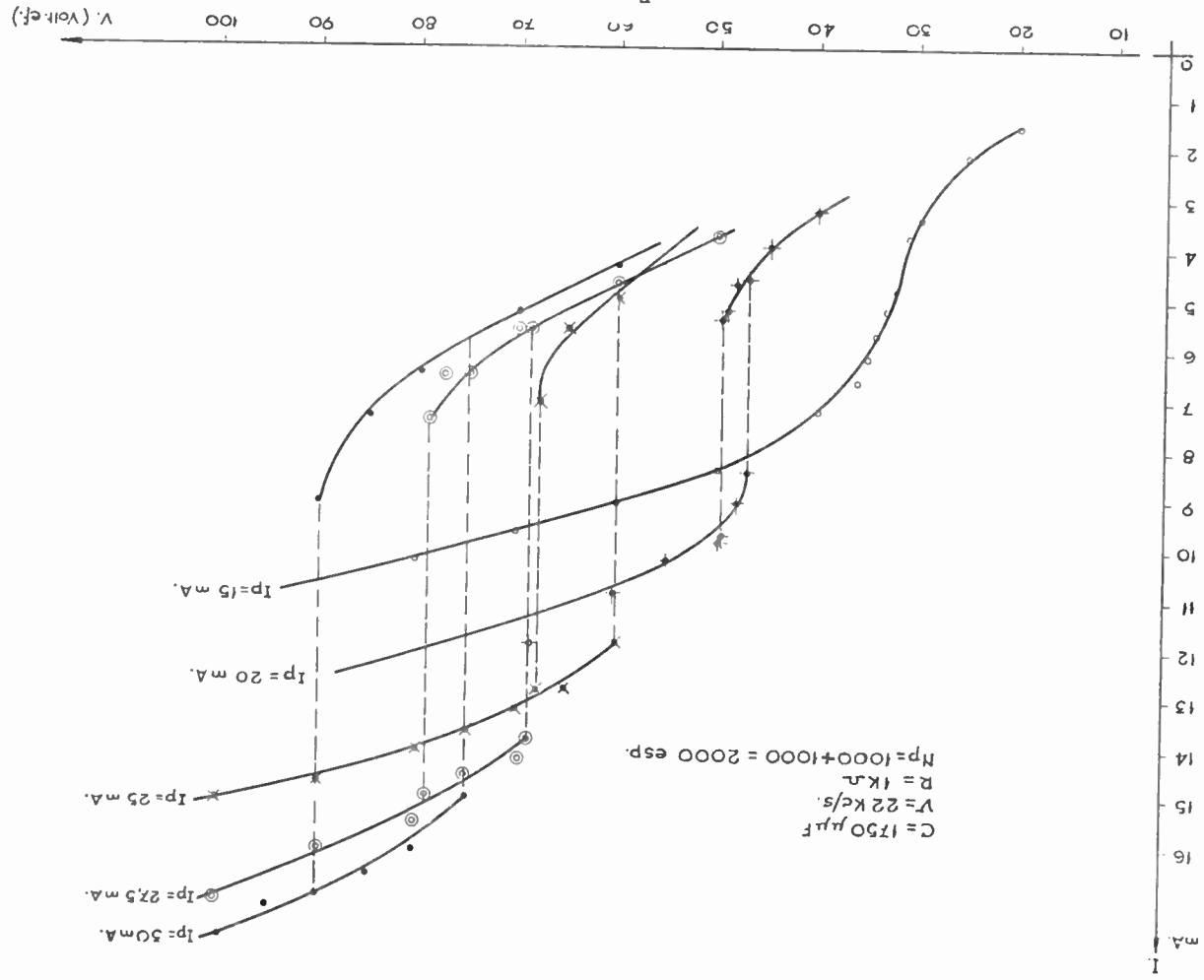
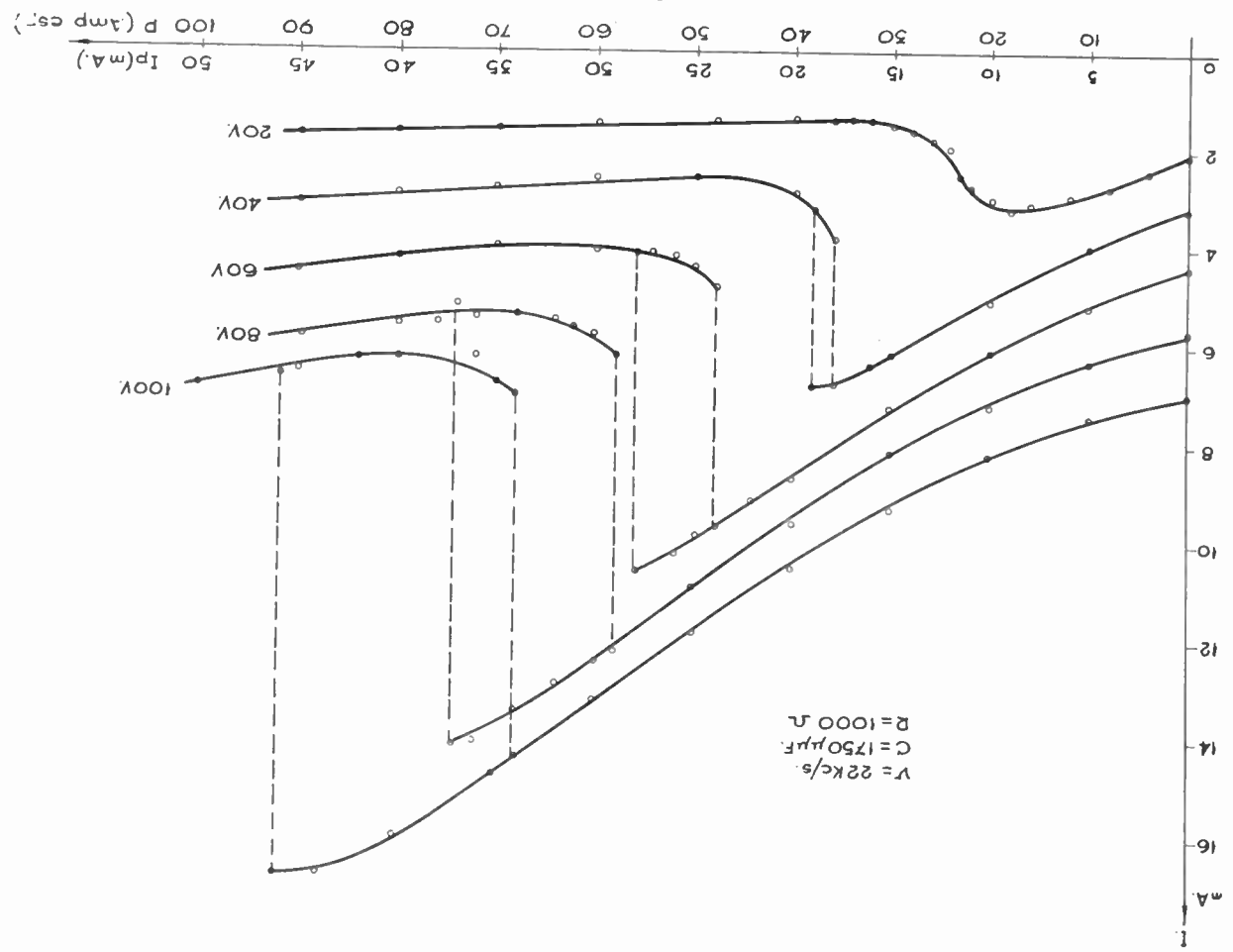


Fig. 10.



c). — Influence de la résistance de charge.

La zone bi-stable se modifie très peu pour les différentes résistances de charge, chaque fois que celles-ci ne sont pas excessives (pour notre cas expérimental cela se produit pour des résistances inférieures à $1\text{ k}\Omega$). (fig. 9).

A mesure qu'on augmente la résistance de charge, au-dessus de cette valeur, la zone bi-stable se réduit et la grandeur de la chute diminue, en même temps qu'elle se déplace vers les intensités de polarisation plus petites; jusqu'à ce qu'elle finisse par disparaître pour des résistances supérieures à $4\text{ k}\Omega$. Cette valeur de charge maxima peut être prédite à partir des caractéristiques du circuit sans réaction.

d). — Influence de la tension d'alimentation.

Dans les graphiques de la fig. 10, on représente comment varie la zone bi-stable quand le circuit est alimenté par des tensions différentes, les autres paramètres du circuit demeurant constants ($v = 22\text{ ke/s}$, $C = 1\,750\text{ }\mu\text{F}$; et $R = 1\,000\text{ }\Omega$).

On doit observer que la tension influe notablement davantage sur la largeur de la zone bi-stable (laquelle augmente à mesure qu'augmente la tension) que dans la grandeur de la chute laquelle augmente aussi.

CARACTÉRISTIQUES. I - V

La zone des intensités stables peut être obtenue, non seulement quand on varie la polarisation, tous les autres paramètres demeurant constants, mais aussi par la variation de n'importe lequel de ceux-ci; pour lesquels uniquement peut offrir un intérêt pratique le cas où varie la tension d'alimentation.

Sur la fig. 11 sont représentées les zones bi-stables, obtenues par variation de la tension d'alimentation, pour diverses intensités de polarisation.

CONCLUSIONS.

Nous avons décrit les fondements et les premières expériences d'un circuit déclencheur ferroresonant parallèle autoexcité dont les caractéristiques peuvent se résumer comme suit :

a) Dans ce circuit on obtient une zone de deux intensités stables, pour une même intensité de polarisation, quand on l'alimente avec une source de tension, contrairement à ce qui se produit dans le circuit ferroresonant parallèle avec excitation indépendante, alimenté par une source d'intensité, où l'on obtient deux tensions stables;

b) La fréquence où l'on obtient la zone bi-stable peut facilement se déplacer en modifiant la capacité en parallèle du circuit ferroresonant;

c) Il peut se faire que cette nouvelle zone de bi-stabilité coexiste avec celle de Fitzgerald pour une même fréquence, et même que chacune ait un intervalle de polarisation commun dans lequel le circuit présentera trois états stables;

d) Le passage de l'une à l'autre branche de la zone bi-stable peut s'obtenir non seulement par variation de l'intensité de polarisation mais aussi en modifiant un quelconque des autres paramètres du circuit (capacité, fréquence, etc.). Le passage offrant uniquement de l'intérêt étant celui obtenu par variation de la tension d'alimentation.

BIBLIOGRAPHIE

1. J. GARCIA SANTESMAS. Progress Report of the Harvard University n° 22 (Cambridge Mass., 1952). *Anales de Física y Química*, vol. 48, A 171 y 355, 1952. Symposium on Automatical Digital Computation, N.P.L. (Teddington, 1953)
2. J. GARCIA SANTESMAS.
 M. RODRIGUEZ VIDAL.
 J. J. SANCHEZ RODRIGUEZ.
 } *Anales de Física y Química*,
 50.A, 47, 1954
3. A. S. FITZGERALD, J. Frank. Inst. 247, 233, 1949.

VIENT DE PARAÎTRE

L'INGÉNIEUR DU SON

en RADIODIFFUSION
CINÉMA
TÉLÉVISION

PAR
V. JEAN-LOUIS
PRÉFACE
de **M. le Général LESCHI**
Directeur des Services techniques
de la Radiodiffusion - Télévision Française

ACOUSTIQUE PSYCHOTECHNIQUE. — Le son et l'oreille, les microphones, l'acoustique architecturale.

LA PRISE DE SON. — L'espace sonore, les emplacements microphoniques, la dynamique sonore, le mixage, les systèmes d'enregistrement optiques et magnétiques.

L'INGÉNIEUR. — Ses fonctions, le concours pour le recrutement d'ingénieurs du son (studio d'essai), l'examen écrit, les tests sonores, conclusion, bibliographie.

UN VOLUME de 296 pages 15 X 24 cm, broché ou relié pleine toile grenat, titres bronze.

Prix : broché, 2 700 F + poste 70 F = 2 770 F ; relié 3 000 F + poste 70 F = 3 070 F

Frais de recommandation en plus s'il y a lieu.

ÉDITIONS CHIRON — 40, rue de Seine — PARIS

Compte chèques postaux PARIS 53.35