

L'ONDE ÉLECTRIQUE

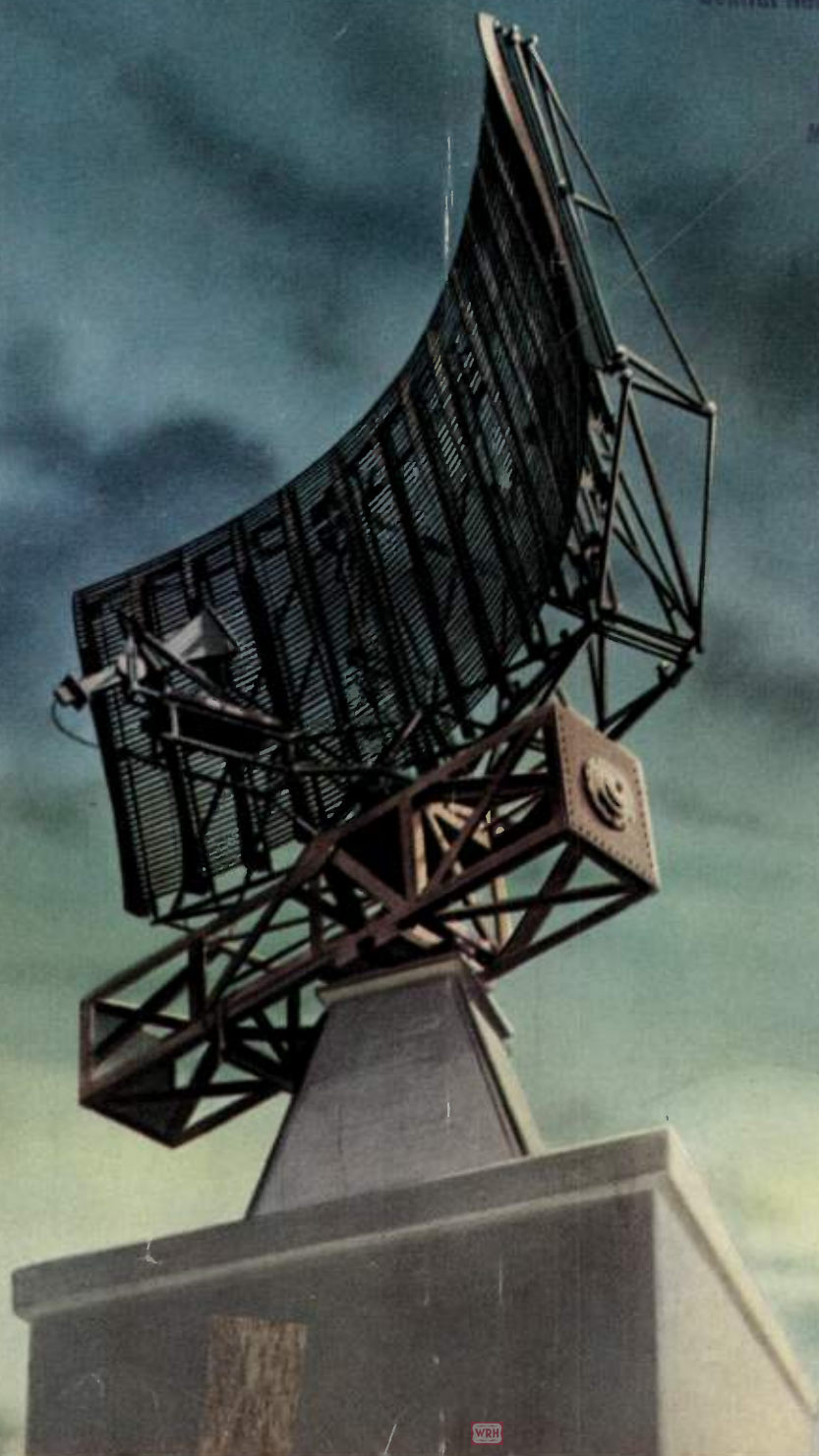
34^e ANNÉE. N° 322

JANVIER 1954

PRIX : 250 FRANCS

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

National Bureau of Standards
Central Radio Propagation Laboratory
Library
Boulder, Colorado
MAR 17 1954



NOTRE COUVERTURE

Antenne Radar de la
Compagnie Générale de
Télégraphie sans Fil.

THÉORIE ET TECHNIQUE
DES
IMPULSIONS

(1^{ère} Partie)

Library
Boulder Laboratories
National Bureau of Standards
Boulder, Colorado

NOV 18 1955

EmetHewlett

TK6540
.06

DE 1 à 3 KW

OC — OM

**POUR TRAFIC
ET RADIODIFFUSION**



**TÉLECOMMUNICATIONS
RADIOÉLECTRIQUES
et
TÉLÉPHONIQUES**



26, RUE BOYER

PARIS - XX^e — TÉLÉPHONE : + MÉNIL. 62-94

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
à partir du 1-1-54
FRANCE 2500 F
ÉTRANGER 2800 »

ÉDITIONS CHIRON
40, Rue de Seine — PARIS (6^e)
C. C. P. PARIS 53-35

Prix du
Numéro :
250 francs

Vol. XXXIV

JANVIER 1954

Numéro 322

SOMMAIRE

	Pages.
Le temps comme variable complexe	E. COLIN-CHERRY 7
La technique des impulsions et les brevets	I. INCOLLINGO. 14
Causes diverses de diaphonie dans les systèmes multiplex à impulsions	J. FAGOT. 21
Diode au germanium à espace positif	A. H. REEVES. 32
Une technique des circuits électroniques pour machine à calculer rapide	G. PIEL. 38
Repérage d'un grand nombre d'intervalles de temps élémentaires au cours d'un cycle	B. LECLERC. 47
Lignes à retard à constantes localisées	H. FEISSEL 53
Nouvelles chaînes de basculeurs utilisées dans les machines à calculer pour le comptage à base 10 et à base 12	M. BRUZAC. 59
Production de signaux de déblocage dans une machine à calculer décimale par un circuit à nombre de tubes réduit	M. BOYER. 63
Le signal minimum utilisable en réception radar et son amélioration par certains procédés de corrélation	L. GÉRARDIN. 67
Emploi des récepteurs à superréaction comme amplificateurs moyenne fréquence pour la réception des impulsions de radar	S. MARMOR. 73
Perturbation d'un oscillateur non linéaire filtré	G. CAHEN. 80
Vie de la Société	90

Sur la couverture :

Antenne radar de la Compagnie Générale de Télégraphie sans Fil

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

SOCIÉTÉ DES RADIOÉLECTRICIENS

FONDATEURS

† Général FERRIÉ, Membre de l'Institut.

† H. ABRAHAM, Professeur à la Sorbonne.
† A. BLONDEL, Membre de l'Institut.
† P. BRENOT, Directeur à la Cie Générale de T.S.F.
† J. CORNU, Chef de bataillon du Génie e. r.

† A. PÉROT, Professeur à l'Ecole Polytechnique.
† J. PARAF, Directeur de la Sté des Forces Motrices de la Vienne.
La Société des Ingénieurs Coloniaux.

SOCIÉTÉ DES RADIOÉLECTRICIENS

10, Avenue Pierre-Larousse, Malakoff (Seine)

Tél. ALESIA 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. POMFY, Inspecteur Général des P.T.T.
1926	E. BRYLINSKI, Ingénieur.
1927	† Ch. LALLEMAND, Membre de l'Institut.
1928	Ch. MAURAIN, Doyen de la Faculté des Sciences de Paris
1929	† L. LUMIÈRE, Membre de l'Institut.
1930	Ed. BELIN, Ingénieur.
1931	C. GUTTON, Membre de l'Institut.
1932	P. CAILLAUX, Conseiller d'Etat.
1933	L. BRÉGUET, Ingénieur.
1934	Ed. PICAULT, Directeur du Service de la T. S. F.
1935	† R. MESNY, Professeur à l'Ecole Supérieure d'Electricité
1936	† R. JOUAUST, Directeur du Laboratoire Central d'Electricité
1937	† F. BEDEAU, Agrégé de l'Université, Docteur ès-Sciences
1938	P. FRANCK, Ingénieur 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 et Télévision Françaises.
1952	J. de MARE, Ingénieur Conseil.

BUREAU DE LA SOCIÉTÉ

Président :

M.P. DAVID, Ingénieur en Chef à la Marine.

Président désigné pour 1954 :

M.G. RABUTEAU, Directeur Général de la Sté « Le Matériel Téléphonique »

Vice-Présidents :

R. RIGAL Ingénieur Général des Télécommunications.

R. AUBERT, Directeur Général adjoint de la S. F. R.

E. FROMY, Directeur de la Division Radioélectricité du L.C.I.E.

Secrétaire Général :

M.J. MATRAS, Ingénieur Général des Télécommunications.

Trésorier :

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

Secrétaires :

H. TESTEMALE, Ingénieur des Télécommunications.

G. ESCULIER, Ingénieur Conseil.

R. CHARLET, Ingénieur des Télécommunications.

SECTIONS D'ÉTUDES

N°	Dénomination	Présidents	Secrétaires
1	Etudes générales.	Colonel ANGOT.	M. LAPOSTOLLE.
2	Matériel radioélectr.	M. LIZON.	M. ADAM.
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.	M ^{me} 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.

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

MM. C. BEURTHÉRET, Ingénieur en Chef à la Cie Française Thomson-Houston.
 A. FLAMBAR, Ingénieur Militaire en Chef de 2^e Classe, chef de la Division D.E.M. du C.N.E.T.
 A. FROMAGEOT, Ingénieur en Chef à la Société L.T.T.
 L.J. LIBOIS, Ingénieur des Télécommunications, Service des Recherches et du Contrôle Techniques des P.T.T.
 P. LIZON, Directeur du Service Radio de la Société « Le Matériel Téléphonique ».
 R. PIRON, Ingénieur du Génie Maritime.
 M. SURDIN, Chef de la Division des Constructions Electriques au Commissariat à l'Energie Atomique.
 J. THURIN, Ingénieur des Télécommunications.
 A. BLANC-LAPIERRE, Professeur à la Faculté des Sciences d'Alger.
 L. CAHEN, ancien Ingénieur en Chef des Télécommunications.
 A. CAZALAS, Ingénieur aux Laboratoires de Télévision et Radar de la Cie pour la fabrication des Compteurs.
 P. CHAVASSE, Ingénieur en Chef des Télécommunications.

MM. A. DANZIN, Directeur de la Société « Le Condensateur Céramique »
 A. DAUPHIN, Ingénieur Militaire Principal des Télécommunications.
 J. DOCKES, Ingénieur des Télécommunications, Service des Recherches et du Contrôle Techniques des P.T.T.
 C. MERCIER, Ingénieur en Chef des Télécommunications.
 J. BOULIN, Ingénieur des Télécommunications à la Direction des Services Radioélectriques.
 F. CARDENAY, 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.

RÉSUMÉS DES ARTICLES

QUELQUES REMARQUES SUR LE TEMPS CONSIDÉRÉ COMME VARIABLE COMPLEXE, par E. COLIN-CHERRY, D. Sc. Onde Electrique de Janvier 1954 (pages 7 à 13).

Cet exposé n'a pour but que d'être une enquête académique sur les résultats obtenus en considérant le temps comme une variable complexe, dans la théorie des transformations de Fourier, Laplace, etc, particulièrement en ce qui concerne l'analyse des circuits en régime transitoire. L'auteur donne l'interprétation pratique des résultats de l'inversion de la relation courante entre le temps réel et la fréquence complexe, en se référant à la théorie des échos électriques.

CAUSES DIVERSES DE DIAPHONIE DANS LES SYSTÈMES MULTIPLEX A IMPULSIONS, par J. FAGOT, Directeur technique à la S.F.R. Onde Electrique de Janvier 1953. (pages 21 à 31).

On rappelle tout d'abord brièvement les principes généraux des systèmes multiplex à impulsions et les divers genres de modulation utilisés : d'amplitude, de durée, de position. On montre que les bruits de diaphonie proviennent de résidus d'amplitude empiétant sur les voies qui succèdent à la voie perturbatrice.

On étudie en premier lieu, la diaphonie qui, dans les amplificateurs directs du spectre vidéo des impulsions, provient de la distorsion sur la transmission des très basses fréquences, en introduisant la notion de « perturbation de diaphonie ».

On envisage ensuite la diaphonie provenant de la mauvaise transmission des composantes à haute fréquence du spectre de modulation, soit dans les amplificateurs à vidéo directe, soit dans les amplificateurs de haute fréquence modulée. On montre que les résidus sont produits par des traînages placés derrière les impulsions.

En conclusion, on peut dire que l'avantage de tous ces phénomènes est de pouvoir être décomposés. Toutes les difficultés peuvent être résolues séparément en produisant une amélioration d'ensemble. D'excellents résultats peuvent être obtenus qui donnent aux multiplex à impulsions à nombre moyen de voies des performances satisfaisant aux normes téléphoniques.

LA DIODE AU GERMANIUM « POSITIVE-GAP DIODE », par M. A. H. REEVES, Standard Telecommunication Laboratories. Onde Electrique de janvier 1954 (pages 32 à 37).

Dans ces diodes au germanium d'un type nouveau, le choix du métal de contact et une formation électrique appropriés font apparaître dans la caractéristique du courant passant une discontinuité où la résistance dynamique devient localement négative.

Il en résulte la possibilité d'introduire ces éléments dans des circuits générateurs d'impulsions, bistables, triggers, compteurs, etc ... pouvant fonctionner jusqu'à plusieurs dizaines, voire centaines de mégacycles par seconde.

UNE TECHNIQUE DE CIRCUITS ÉLECTRONIQUES POUR MACHINE A CALCULER RAPIDE, par M. PIEL, Ingénieur à la Société d'Electronique et d'Automatisme. Onde Electrique de Janvier 1954 (pages 38 à 46).

L'article développe les détails d'une technologie de machine à calculer électronique fonctionnant avec des impulsions de rythme à 800 Kc/s. Cette technologie ne fait appel qu'à un matériel classique limité — triodes à vide et diodes au germanium — et comporte un petit nombre de montages élémentaires dont les possibilités d'emploi ont été jugées satisfaisantes.

Un symbolisme schématique est développé pour chaque type de montage de sorte qu'un schéma d'organe utilisant ce symbolisme fixe intégralement le matériel utilisé.

REPÉRAGE D'UN GRAND NOMBRE D'INTERVALLES DE TEMPS ÉLÉMENTAIRES AU COURS D'UN CYCLE, par B. LECLERC, Ingénieur à la Compagnie des Machines Bull. Onde Electrique de Janvier 1954 (pages 47 à 52).

Dans l'élaboration des systèmes de transmission d'informations par impulsions codées, on éprouve très généralement le besoin de définir un intervalle de temps fixe, ou « cycle », de le subdiviser en intervalles de temps élémentaires, et de reconnaître un ou plusieurs de ces intervalles élémentaires à l'intérieur du cycle.

Ces problèmes ont été rencontrés d'une façon particulièrement typique au cours de l'étude d'une Calculatrice Electronique Arithmétique du type « Série ».

Ils ont été résolus par l'emploi d'une hiérarchie de distributeurs d'impulsions à deux, trois ou quatre directions, réalisés sous la forme de multivibrateurs généralisés, à l'aide d'étages à transformateurs interconnectés entre eux par des circuits de commutation à diodes au germanium.

Cette technique, d'une grande souplesse, conduit à des réalisations intéressantes par leur simplicité, leur stabilité et leur rendement énergétique élevé.

LIGNES A RETARD A CONSTANTES LOCALISÉES, par H. FEISSEL, Ingénieur à la Compagnie des Machines Bull. Onde Electrique de Janvier 1954 (pages 53 à 58).

Pour des raisons d'économie et de simplicité, les lignes à retard électriques à constantes localisées sont généralement réalisées en structure coaxiale, c'est-à-dire en disposant des bobines à la suite les unes des autres sur un même mandrin cylindrique.

Cette disposition fait apparaître des couplages entre bobines non adjacentes. Chaque cellule de la ligne n'est plus alors assimilable à un quadripôle, et la théorie classique des filtres ne peut s'appliquer au calcul rigoureux des caractéristiques de transmission de l'ensemble réalisé.

Une étude systématique des lignes en structure coaxiale a été conduite expérimentalement, en s'appuyant sur la théorie des filtres, non pour effectuer le calcul rigoureux des caractéristiques de transmission, mais pour orienter l'expérience.

L'effort principal a porté sur la recherche d'un retard en transmission aussi indépendant que possible de la fréquence transmise.

Il a abouti à l'élaboration d'une structure corrigée présentant, sous cet aspect, des caractéristiques sensiblement supérieures à celles que l'on obtient avec des lignes de structure classique.

PAPERS SUMMARIES

ELECTRONIC CIRCUIT TECHNIQUE FOR A HIGH SPEED RAPID CALCULATING MACHINE, by M. PIEL, *Ingénieur à la Société d'Electronique et d'Automatisme. Onde Electrique*, January 1954 (Pages 38 to 46).

The article develops the technique of an electronic calculating machine, which operates on repetitive pulses at 800 kc/s. This technique calls for only a limited choice of well established components — vacuum triodes and germanium diodes, and comprises a small number of unit assemblies designed for flexibility.

Symbols for the various units have been laid down so that a functional schematic utilising the symbols rigidly specifies the components used.

THE IDENTIFICATION OF A LARGE NUMBER OF ELEMENTARY TIME INTERVALS IN THE COURSE OF A CYCLE, by M. LECLERC, *Ingénieur à la Compagnie des Machines Bull. Onde Electrique*, January 1954 (pages 47 to 52).

In the development of pulse-code transmission systems, there is a very general need to define a fixed time interval or cycle, to subdivide it into elementary time intervals, and to recognise one or more of these elementary intervals within the cycle.

These problems were encountered in a particularly typical form during the study of an electronic arithmetical calculator of the "series" type. They were solved by using a number of 2,3 or 4 way impulse distributors, in the form of generalised multi-vibrators with the aid of transformer stages interconnected by germanium diode switching circuits.

This technique is very flexible and leads to results which are valuable because of their simplicity, stability, and high power efficiency.

LUMPED PARAMETER DELAY LINES by M. FEISSEL, *Ingénieur à la Compagnie des Machines Bull. Onde Electrique*, January 1954 (pages 53 to 58).

For reasons of economy and simplicity, lump loaded electric delay lines are usually produced in coaxial form, by arranging the coils in sequence on the same cylindrical mandril.

This disposition leads to coupling between non-adjacent coils. Each line section can no longer be simulated by a quadripole, and the classical filter theory cannot be applied to rigorous calculations of the transmission characteristics of the whole assembly.

A systematic study of lines of coaxial structure is dealt with experimentally. Filter theory is relied upon to give a pointer to the study and not to carry out rigorous calculations of the transmission characteristics.

The object in view is to derive a transmission delay as independent of frequency as possible. This has resulted in the development of a corrected structure having characteristics which are sensibly superior to those obtained with lines of the classic type.

SOME REMARKS ON TIME AS A COMPLEX VARIABLE by E. COLIN CHERRY D. Sc. *Onde Electrique*, January 1954 (pages 7 to 13).

This paper is intended only as an academic enquiry into the results of regarding time as a complex variable, in the transform theory of Fourier, Laplace, etc, especially in relation to the analysis of electrical transients. Practical interpretation is given to the results of reversing the usual real-time/complex-frequency relationship, by reference to the theory of electrical echoes.

CAUSES DIVERSES DE DIAPHONIE DANS LES SYSTÈMES MULTIPLEX A IMPULSIONS, by J. FAGOT, *Directeur technique à la S.F.R., Onde Electrique*, January 1953, (pages 21 to 31).

First the general principles are recalled of pulse multiplex systems and the various kinds of modulation employed: amplitude, time, position. It is shown that crosstalk arises from residual amplitude effects encroaching on the channels which follow the disturbing channel. In the first place consideration is given to crosstalk which, in the direct amplifiers of the pulse video spectrum, arises from distortion in the transmission of the very low frequencies, introducing the idea of "crosstalk disturbance".

Then there is considered the crosstalk arising from faulty transmission of the H. F. components of the modulation spectrum either in the direct video amplifiers, or in the modulated H.F. amplifiers. It is shown that the residual traces are caused by effects lagging behind the pulses. As a general conclusion, it can be noted that the advantage of all these phenomena is that they are capable of being decomposed. All difficulties can be resolved separately thus giving an overall improvement.

Excellent results can be secured giving to pulse multiplex with medium numbers of channels performances satisfying standard telephone requirements.

"POSITIVE GAP" GERMANIUM DIODE, by Mr. A. H. REEVES, *Standard Telecommunication Laboratories. Onde Electrique*, January 1954 (Pages 32 to 37).

In a germanium diode of a new type, the choice of contact metal and an appropriate electrical formation causes a discontinuity in the current characteristic where the dynamic resistance becomes locally negative.

As a consequence these elements can be used in circuits of various types, pulse generators, flip-flop, trigger and counter circuits, working up to tens or even hundreds of megacycles per second.

RÉSUMÉS DES ARTICLES (Suite)

NOUVELLES CHAINES DE BASCULEURS UTILISÉES DANS LES MACHINES A CALCULER POUR LE COMPTAGE A BASE 10 ET A BASE 12, par J.J. BRUZAC, Ingénieur à la Compagnie I. B. M.-France. Onde Electrique de Janvier 1954 (pages 59 à 62).

Des modifications apportées à l'interconnexion des 4 basculeurs de la décade de Potter permettent d'obtenir des chaînes de basculeurs pour le comptage à base 10 et base 12 en les dotant de la possibilité de former par une seule impulsion le complément à 9 ou à 11 du chiffre enregistré.

L'indication de l'état des basculeurs des chaînes et, partant, l'indication du chiffre enregistré se fait par lampe au néon comme dans la décade de Potter.

Toutes les chaînes sont représentées par schémas bloc et des graphiques indiquent les sauts de tension aux sorties des basculeurs pour chaque chiffre enregistré.

PRODUCTION DE SIGNAUX « GATE » DANS UNE MACHINE A CALCULER DÉCIMALE PAR UN CIRCUIT A NOMBRE DE TUBES RÉDUITS, par G. BOYER, Ingénieur à la Compagnie I. B. M.-France. Onde Electrique de Janvier 1954 (pages 63 à 66).

Dans les machines à calculer décimales, on utilise une série de signaux « Gate » à faible impédance permettant de sélectionner dans de nombreux circuits simultanément des trains de 1 à 9 impulsions.

Ces signaux peuvent être obtenus par des procédés classiques, c'est-à-dire par des basculeurs suivis de cathodes followers mais il est plus économique de les produire directement à faible impédance par des thyatron.

La principale difficulté est l'extinction de ces tubes qui doit être réalisée électroniquement ; le procédé utilisé est l'application d'une impulsion positive sur la cathode ou négative sur l'anode suivant que la charge est cathodique ou anodique.

Cette impulsion qui doit fournir aux circuits des tubes à éteindre une puissance considérable pendant le temps de désionisation est donnée par un autre thyatron qui se désamorce de lui-même aussitôt après la production de l'impulsion.

LE SIGNAL MINIMUM UTILISABLE EN RÉCEPTION RADAR ET SON AMÉLIORATION PAR CERTAINS PROCÉDÉS DE CORRÉLATION, par L. GERARDIN, Ingénieur à la Compagnie Française Thomson-Houston. Onde Electrique de Janvier 1954 (pages 67 à 72).

La réception des signaux radars, particulièrement ceux provenant d'échos isolés, se présente évidemment comme un cas particulier de la réception d'impulsions. Ce problème est, en dernière analyse, dominé par les questions de rapport signal/bruit.

Il est donc nécessaire de définir très exactement le signal minimum de façon à pouvoir, par exemple, effectuer des prévisions de portée d'appareils nouveaux.

Le critère de visibilité doit correspondre au type de transformateur physiologique adopté pour la présentation, on utilise ici celui de la déflexion moyenne, c'est-à-dire le rapport k de la valeur moyenne du signal à la valeur efficace du bruit.

Si, sur un oscilloscope type A de mesure, où le signal est permanent, le coefficient k est de -4 à -6 dB, il est de plus de $+8$ dB pour un scope. « P. P. I ».

Les expériences effectuées dans ce dernier cas confirment ces chiffres.

Pour reculer le seuil ultime, il faut exploiter mieux les informations connues a priori sur les signaux à recevoir.

Une réalisation pratique possible est sommairement décrite.

EMPLOI DES RÉCEPTEURS A SUPERRÉACTION COMME AMPLIFICATEURS MOYENNE FRÉQUENCE POUR LA RÉCEPTION DES IMPULSIONS DE RADAR, par S. MARMOR, Ingénieur, Chef du service « Réception » à la division « Détection Electromagnétique » du C. N. E. T.. Onde Electrique de Janvier 1954 (pages 73 à 79).

Après discussion des mérites respectifs des divers types de récepteurs utilisés pour l'amplification des impulsions brèves l'auteur décrit un récepteur à superréaction réalisé dans les laboratoires de la Division « Détection Electromagnétique » du Centre National d'Etudes des Télécommunications et compare les résultats obtenus en remplaçant l'amplificateur M. F. classique d'un superhétérodyne de radar par ce type de récepteur.

PERTURBATIONS D'UN OSCILLATEUR NON LINÉAIRE FILTRÉ, par G. CAHEN, Onde Electrique de Janvier 1954 (pages 80 à 89).

Soit un oscillateur constitué par un amplificateur bouclé sur lui-même au moyen d'une rétroaction, et supposé muni d'un filtre passe-bas coupant les harmoniques de la fréquence normale d'oscillation, ce qui permet de parler de gain de l'amplificateur en amplitude et en phase. L'amplificateur n'est pas linéaire, sans quoi l'amplitude de l'oscillation ne se fixerait pas à un régime déterminé. Le gain dépend donc non seulement de la pulsation ω mais aussi de l'amplitude x et, bien entendu, si le régime varie, de la dérivée logarithmique de l'amplitude $\dot{x}/x = \tau$. En régime non perturbé, le gain est 1 en amplitude et π en phase. Moyennant certaines hypothèses, on peut admettre que les gains en amplitude et en phase sont, au voisinage du régime normal, fonctions linéaires de ω , x et τ . Ceci permet d'étudier le comportement de l'oscillateur en présence d'une perturbation de faible amplitude u et de pulsation $\omega_0 + \varepsilon$, voisine de la pulsation ω_0 en régime non perturbé. On aboutit à une équation non linéaire du second ordre dont l'inconnue est l'avance de phase θ de l'oscillateur par rapport à l'excitation. Cette équation peut s'étudier géométriquement et analytiquement et permet de montrer qu'il y a accrochage de l'oscillateur sur la perturbation lorsque le rapport u/ε de l'amplitude de la perturbation à son écart de fréquence est supérieur à un certain seuil. Lorsque ce rapport est inférieur au seuil, il n'y a pas accrochage, mais glissement de la fréquence de l'oscillateur vers celle de la perturbation, glissement dont la valeur peut être calculée.

Des phénomènes analogues peuvent être observés et soumis au calcul lorsque l'on couple deux oscillateurs de fréquence légèrement différentes.

Les calculs précédents sont valables lorsque les écarts relatifs de fréquence sont petits devant les inverses des coefficients de surtension des amplificateurs.

Si cette condition n'est pas respectée, une variante des méthodes de calcul permet de retrouver la valeur du seuil d'accrochage dans le cas, assez usuel, où la fonction de gain est le produit du gain d'un amplificateur linéaire (indépendant de l'amplitude) par une fonction de l'amplitude (indépendante de la fréquence).

PAPERS SUMMARIES (Continued)

SUPER-REGENERATIVE RECEIVERS AS MEAN FREQUENCY AMPLIFIERS FOR RECEPTION OF RADAR SIGNALS, by S. MARMOR, *Ingénieur, Chef du Service "Réception" à la division "Electromagnétique" du C. N. E. T.* Onde Electrique, January 1954 (pages 73 to 79).

After a discussion on the relative merits of different types of receivers for amplifying short pulses, the author describes a superregenerative receiver designed by the laboratories of the "Electromagnetic Detection" Division of the "Centre National d'Etudes des Télécommunications" and compares the results obtained if the classical I. F. amplifier of a Radar superhetrodyne set is replaced by this type of receiver.

DISTORTIONS IN A NON LINEAR OSCILLATOR WHICH INCLUDES A FILTER, G. CAHEN. Onde Electrique, January 1954 (pages 80 to 89).

If an oscillator comprises an amplifier with feedback and a low pass filter which suppresses the harmonics of the normal frequency of oscillation, then the gain of the amplifier can be considered in terms of amplitude and phase. The gain of the amplifier is not linear, otherwise the amplitude of oscillation would not settle down to a definite value. The gain depends not only upon the pulsance ω but also upon the amplitude x , and if the conditions change then on the ratio of the derivative of the amplitude to the amplitude $\dot{x}/x = \tau$. In the undisturbed condition the gain is of amplitude 1 and phase shift zero.

Making certain assumptions it can be assumed that the amplitude and phase of the gain vector, near normal working conditions, is linear with respect to ω , x and τ . This enables the behaviour of the oscillator to be studied in the presence of a weak disturbance of amplitude u , and pulsance $\omega_0 + \epsilon$, which is close to the undisturbed pulsance ω_0 . This leads to a non linear equation of the second order in which the unknown θ is the phase shift of the oscillator with respect to the excitation.

This equation can be studied geometrically and analytically and shows that the oscillator locks to the disturbance when the ratio amplitude divided by frequency change is greater than a certain threshold. When the ratio is less than this threshold there is no locking, but the frequency of the oscillator is pulled toward that of the disturbance by a calculable amount.

Analogous effects can be observed and calculated when two oscillators of slightly different frequency are coupled.

The foregoing calculations are valid when the relative frequency displacements are small compared with the inverse of the coefficients of excess voltage of amplifier. If this condition is not met a variant of the method of calculation will yield the value of the threshold, in the quite usual case, where the function of the gain is the product of the gain of a linear amplifier (independent of amplitude) and a function of the amplitude (independent of frequency).

NEW FLIP-FLOP CHAINS USED IN CALCULATING MACHINES FOR COUNTING TO THE BASE 10 OR TO THE BASE 12, by J. J. BRUZAC, *Ingénieur à la Compagnie I. B. M. France.* Onde Electrique, January 1954 (pages 59 to 62).

Flip-Flops for counting to the base 10 or 12 can be arranged by modifications to the interconnections of the 4 flip-flops of the Potter decade arrangement. This is done by enabling them to form the complement to 9 or to 11 of the figure registered.

The condition of the flip-flops of the chains and finally the indication of the figure registered, are shown by neon lamps as in the Potter decade arrangement.

All the chains are shown in black schematic form, and the potential changes occurring at the outputs of the flip-flops, for each figure registered, are indicated by curves.

PRODUCTION OF "GATE" SIGNAL IN A DECIMAL COMPUTER WITH A REDUCED NUMBER OF TUBES, by G. BOYER, *Ingénieur à la Compagnie I. B. M. France.* Onde Electrique, January 1954 (pages 63 to 66).

In this decimal computer a series of "Gate" signals are used in a low impedance circuit so that a number of 1 to 9 impulse trains can be simultaneously selected.

These signals can be obtained in the standard manner, by a flip-flop followed by cathode followers, but it is more efficient to produce them directly at low impedance by thyratrons.

The principal difficulty is to extinguish these tubes which must be done electronically. This is carried out by applying a positive pulse to the cathode or negative pulse to the anode, depending on whether the load is connected to the cathode or the anode.

This extinguishing pulse which has to be applied to the thyatron circuit represents considerable power during the de-ionisation period, and is given by another thyatron which extinguishes itself directly after the production of the pulse.

MINIMUM USABLE SIGNAL IN RADAR RECEPTION AND ITS EASEMENT BY CORRELATION PROCESSES by L. GERARDIN, *Ingénieur à la Compagnie Française Thomson-Houston.* Onde Electrique, January 1954 (Pages 67 to 72).

Radar signal reception, particularly when individual echoes are involved, is a particular case of pulse reception. The problem is basically one of signal-to-noise ratio.

It is therefore necessary to define precisely the minimum signal, in such a manner as to specify limits for new apparatus.

The criterion of visibility must depend upon the type of display and subjective factors. Here the mean deflection is used, that is to say, the ratio K of the mean signal value to the effective (RMS) value of noise.

If, on a type A oscilloscope, where the signal is permanent, the coefficient K has a value between -4 and -6 dB. the value for a P.P.I. oscilloscope is $+8$ dB. These figures are confirmed by experience.

To improve this ultimate threshold, prior knowledge of the signal must be exploited. A practical way of doing this is concisely described.

QUELQUES REMARQUES SUR LE TEMPS CONSIDÉRÉ COMME VARIABLE COMPLEXE

PAR

E. COLIN CHERRY D. Sc.
Imperial College (Londres)

1. Introduction.

Ce compte-rendu n'a pas la prétention d'être un exposé complet ; ce n'est guère plus qu'une note brève pouvant, suivant nos espérances, suggérer des commentaires et la discussion.

Il concerne la conception du temps (tel que considéré dans l'analyse des impulsions et des transitoires) comme variable complexe et a l'intention d'examiner les conséquences de l'adoption de ce point de vue. Ces conséquences servent à faire ressortir la sagesse de la conviction de Heaviside qu'un calcul opérationnel, qui transforme les fonctions du temps en fonctions d'un opérateur p simplifierait fondamentalement l'analyse. Toutefois, les progrès récents dans l'analyse des transitoires et de la synthèse des réseaux, suggèrent qu'un nouvel examen de la possibilité de transformation entre le « temps » et la « fréquence » puisse être actuellement à conseiller.

Le lecteur doit être avisé qu'aucune nouvelle méthode pratique de l'étude des systèmes ne sera présentée ici.

Lorsque les ingénieurs électriciens étudiaient d'abord les problèmes en courant alternatif, la *fréquence* fut regardée comme une variable réelle — des ondes simplement sinusoïdales et cosinusoidales. De l'introduction, par Steinmetz, de la fréquence complexe résulta une grande réduction et une grande clarté des raisonnements mathématiques, surtout en ce qui concerne le calcul de la réponse aux transitoires. Mais l'adoption logique éventuelle du concept du *plan complexe* alla beaucoup plus loin encore tant en théorie qu'en pratique. Il est vrai de dire que l'extension de l'emploi du plan de fréquence complexe a fait apparaître beaucoup de propriétés jusqu'alors inconnues ou mal comprises de circuits. Cet emploi a clarifié le calcul opérationnel de Heaviside, il a préparé l'étude des amplificateurs à réaction et a permis l'utilisation de l'intégration des contours.

Durant ces dernières années, une application pratique directe du plan complexe a été développée sous forme de machines à calculer par analogies. De telles machines utilisent l'analogie du potentiel

entre des champs électriques représentés en deux dimensions et le plan complexe ; leur aide permet de calculer la réponse en régime permanent, en régime transitoire et autres aspects de l'analyse des réseaux linéaires, et le problème extrêmement difficile de l'analyse des réseaux peut, d'une manière considérable, être réduit à une technique expérimentale (1) (2).

Il est bien connu que les variables *fréquence* et *temps* se présentent inversement dans les relations ; les fonctions de ces variables présentant souvent une symétrie complète. La conception de la fréquence comme variable complexe ayant procuré tellement d'avantages, cet auteur s'est longtemps demandé si des avantages semblables ne pourraient découler de la conception du temps regardé comme variable complexe. Nous allons simplement examiner dans cette note quelques conséquences de cette possibilité

2. Le temps et la fréquence ; l'inversion de la transformation de Fourier.

Les ingénieurs ont été depuis longtemps familiarisés. avec la relation inverse précise existant entre le temps t et la fréquence angulaire réelle ω , dans les deux transformées de Fourier :

$$\begin{aligned} f(t) &= \int_{-\infty}^{+\infty} g(\omega) e^{j\omega t} d\omega \\ g(\omega) &= \frac{1}{2\pi} \int_{-\infty}^{+\infty} f(t) e^{-j\omega t} dt \end{aligned} \quad (1)$$

Les relations fonctionnelles exprimées ici sont conservées à un facteur constant près, si ω et t sont où que ce soit mis à la place l'un de l'autre en utilisant la transformation :

$$\begin{aligned} \omega &\longrightarrow -j \\ t &\longrightarrow \omega \end{aligned} \quad (2)$$

Brièvement, cela signifie que si une certaine forme d'ondes $f(t)$ a le spectre $g(\omega)$, une nouvelle forme d'ondes $g(-t)$ a le spectre $f(\omega)$; les représentations de forme d'ondes et de spectre sont interchangeables.

L'auteur a discuté ailleurs (3) les résultats de l'application de cette inversion de t et ω dans la théorie vectorielle élémentaire des courants alternatifs; nous obtenons alors des diagrammes vectoriels « ondes sinusoïdales » en fréquence, « spectre de temps » en un diagramme Argand en fonction de la fréquence. Virtuellement, l'ensemble de la théorie simple devient alors inversé et est interprété comme la théorie de la réponse du circuit aux *impulsions* et non aux ondes sinusoïdales.

3. Inversion possible des transformées de Laplace.

Si la fréquence variable est rendue complexe, le long du contour droit ($c + j\omega$) parallèle à l'axe $j\omega$, les équations (1) deviennent les transformées de Fouries-Mellin; s'il est possible de considérer des valeurs complexes générales, $j\omega \longrightarrow \lambda$, la relation de transformation devient alors celle des transformées de Laplace :

$$f(t) = \frac{1}{2\pi j} \int F(\lambda) e^{\lambda t} d\lambda \quad (2a)$$

$$F(\lambda) = \int_0^\infty f(t) e^{-\lambda t} dt \quad (2b)$$

Il est maintenant commode d'utiliser un nouveau symbole, F , pour la fonction de fréquence complexe.

Les fonctions intervenant alors sont $f(t)$, une fonction d'une variable réelle représentée sur un axe, et $F(\lambda)$, une fonction d'une variable complexe représentée sur un plan. D'un point de vue purement mathématique, l'extension des transformées de Fourier (1) aux transformées de Laplace aurait pu être effectuée alternativement en laissant jt devenir complexe et en conservant ω réel. Opérons ainsi mais inversons d'abord la signification des deux équations (1); écrivons $G(\omega) = 2\pi g(\omega)$ et $-\omega$ pour ω :

$$\left. \begin{aligned} G(-\omega) &= \int_{-\infty}^{+\infty} f(t) e^{i\omega t} dt \\ f(t) &= \frac{1}{2\pi} \int_{-\infty}^{+\infty} G(-\omega) e^{-i\omega t} d(-\omega) \end{aligned} \right\} \quad (3)$$

et, par analogie avec le cas précédent, laissons $jt \longrightarrow \gamma$, une variable en temps complexe,

$$\gamma = (\tau + jt)$$

$$G(-\omega) = \frac{1}{2\pi j} \int \Psi(\gamma) e^{\omega \gamma} d\gamma \quad (4a)$$

$$\Psi(\gamma) = \int_0^{-\infty} G(-\omega) e^{-\gamma \omega} d\omega \quad (4b)$$

Il faut toutefois que nous nous rendions ensuite compte que de telles équations représentent bien ce qui se passe dans des systèmes *physiques*.

Remarquez que l'équation (4b) a été écrite en $\int_0^{-\infty} d\omega$ et non $\int_{-\infty}^0 d(-\omega)$ et que l'on suppose ici que G n'est fini que pour des *fréquences négatives*; par analogie, dans l'analyse plus classique, $f(t)$ est supposé fini seulement pour les *temps positifs*, étant considéré comme une réponse transitoire débutant à $t = 0$.

On peut considérer qu'il s'agit là du « Principe de Causalité » de la réponse transitoire; on détermine par là diverses propriétés mathématiques générales des circuits. Par exemple, les parties réelle et imaginaire sont en relation, constituant une paire de transformées de Hilbert (3); là encore des conditions se présentent qui limitent la position des pôles et des zéros dans le plan de fréquence complexe (ils sont en général limités au demi-plan à gauche). Comme nous allons maintenant inverser le temps et la fréquence, nous devons, par intuition, nous attendre à quelques conditions correspondantes; c'est ainsi, qu'au début, nous avons l'intention de nous occuper uniquement de fonctions de fréquence finies d'un même côté de l'origine — le côté négatif semble plus le logique.

Nous devons ainsi restreindre l'analyse, de la manière désirée par nous, à des fonctions de temps avec le spectre négatif seul; commençons toutefois avec n'importe quelle transitoire pratique $f(t)$ et transformons la en une telle fonction. On peut montrer que cela nécessite l'addition en quadrature d'une simple transitoire relative $h(t)$, telle que $\Psi = (h + jf)$ représente la fonction désirée du temps. Examinons la relation entre $f(t)$ et $h(t)$.

La fonction $\Psi(\gamma)$ est une fonction complexe du temps :

$$\Psi(\gamma) = h(\tau, t) + j f(\tau, t) \quad (5)$$

ce qui devient sur le « vrai » axe des temps :

$$\Psi(t) = h(t) + j f(t) \quad (6)$$

où $f(t)$ est notre transitoire originale; qu'est $h(t)$ que l'on y ajoute en quadrature?

Étant en quadrature, nous pouvons écrire ainsi son spectre de Fourier, en relation avec celui de $f(t)$:

$$\left. \begin{aligned} f(t) &= \int_0^{+\infty} [a(\omega) \cos \omega t + b(\omega) \sin \omega t] d\omega \\ \text{on doit alors avoir :} \\ h(t) &= \int_0^{+\infty} [a(\omega) \sin \omega t - b(\omega) \cos \omega t] d\omega \end{aligned} \right\} \quad (7)$$

Deux fonctions comme celles-ci ont des spectres semblables, mais avec des termes de fréquence correspondante en quadrature l'un par rapport à l'autre. Des signaux tels que $\Psi(t)$ ont été utilisés par d'autres auteurs, en particulier Gabor (4) qui les a introduits dans l'analyse des réseaux en partant de la mécanique quantique. On les a ensuite mentionnés sous le terme de « signaux analytiques ».

Pour montrer qu'ils n'ont que des spectres de fréquences négatives, exprimons les sous forme exponentielle :

$$f(t) = \int_0^{+\infty} \left[\left(\frac{a - jb}{2} \right) e^{j\omega t} + \left(\frac{a + jb}{2} \right) e^{-j\omega t} \right] d\omega$$

$$h(t) = \int_0^{+\infty} \left[- \left(\frac{b + ja}{2} \right) e^{j\omega t} + \left(\frac{-b + ja}{2} \right) e^{-j\omega t} \right] d\omega$$

ce sorte que $\Psi(t) = h(t) + j f(t)$

$$= \int_0^{+\infty} [-b(\omega) + ja(\omega)] e^{-j\omega t} d\omega \quad [|\omega| > 0] \quad (8)$$

un spectre de fréquences négatives seulement, ayant des amplitudes complexes. Remarquez que ce spectre unilatéral $(-b + ja)$ est obtenu de suite directement en partant de la transitoire originale $f(t)$ par annulation de toutes les composantes positives et multiplication de toutes les composantes négatives par $2j$.

4. 4. Signaux analytiques.

Donnons une illustration pratique de cette inversion du temps et de la fréquence. La figure 1 montre un certain nombre de courbes qui apparaissent familières au lecteur. Aucune échelle horizontale n'y a

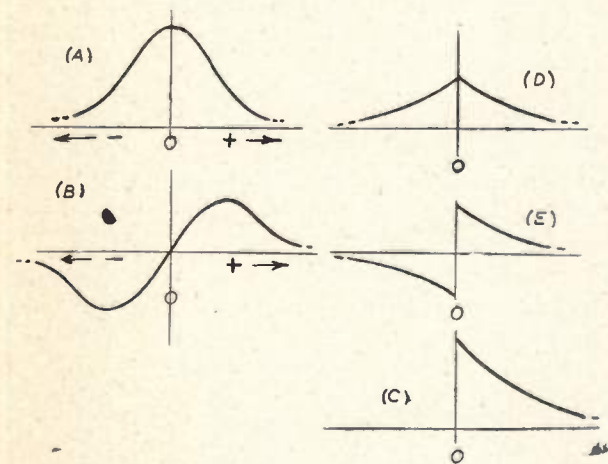


FIG. 1. — Quelques formes d'impulsions et leur spectre. L'inversion du temps et de la fréquence.

été portée ; elle peut représenter le temps ou la fréquence. Les deux courbes (A) et (B) représentent des fonctions en quadrature, leurs modules de spectre sont identiques, mais les composantes correspon-

dantes sont déphasées de $\pi/2$. Leurs équations respectives sont :

$$A(x) = \frac{1}{1 + (x/x_0)^2} ; \quad B(x) = \frac{(x/x_0)}{1 + (x/x_0)^2} \quad (9)$$

où x peut être pris soit comme t soit comme ω .

Dans le cas familier simple, $x = \omega$ et ces courbes peuvent être regardées comme représentant les parties réelle et imaginaire d'un circuit RC, à résistance-capacité, simple.

La transformée de la fonction paire (A) est la courbe exponentielle symétrique (D) ; celle de la

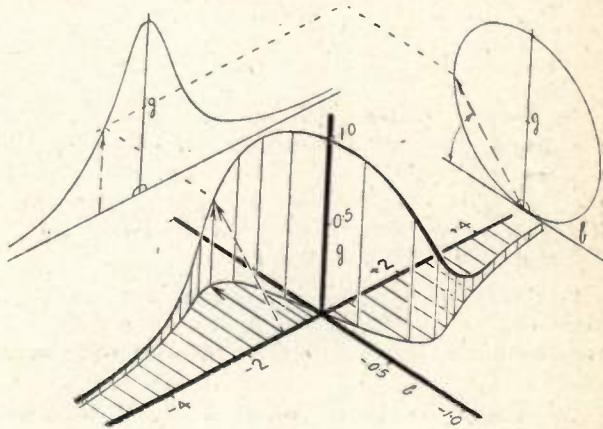


FIG. 2. — Le signal analytique.

L'addition en quadrature de deux fonctions simples pour former une courbe à trois dimensions.

fonction impaire (B) est la courbe exponentielle oblique (E). L'addition en quadrature de (A) et de (B) correspond à l'addition arithmétique de (D) et (E) correspondant à la transitoire exponentielle familière (C) pour $t > 0$ seulement. La figure 2 montre la courbe à 3 dimensions ou « diagramme de résonance solide » (7) résultant de la représentation graphique de ces fonctions en quadrature.

Si t et $-\omega$ sont partout inversés, $x = t$, et les courbes (A) et (B), deviennent les fonctions du temps en quadrature comme nos $f(t)$ et $h(t)$. Leurs spectres sont représentés par (D) et (E) [avec échelles horizontales inversées] qui, par addition arithmétique, donnent le spectre unilatéral (négatif) (C) [avec échelle inversée].

Des fonctions telles que (A) et (B) qui, indépendamment de leur symétrie paire ou impaire, ont des spectres et des transformées de Fourier identiques sont connues sous le nom de transformées de Hilbert $f(t)$ et $h(t)$ sont ainsi exprimées par les relations :

$$f(t) = \frac{1}{\pi} \int_{-\infty}^{+\infty} h(v) \frac{dv}{v - t} ; \quad h(t) = \frac{1}{\pi} \int_{-\infty}^{+\infty} f(v) \frac{dv}{v - t} \quad (10)$$

Deux points sont remarquables dans cet exemple. Premièrement, dans la figure 1, nous avons choisi pour (A) une fonction paire pure, représentant $f(t)$; c'est une simplification car $f(t)$ a, à la fois, des composantes sinusoïdale et cosinoïdale. Deuxièmement,

mement, il résulte de la non symétrie de $f(t)$ comme de $h(t)$ que le signal analytique $\Psi(t)$ a le spectre unilatéral complexe $(-b + ja)$ donné par (8) tandis que la courbe (C) de la figure 1 est une courbe simple.

Nos deux fonctions du temps et de la fréquence sont maintenant complexes et elles sont représentées par deux transformées de Fourier :

$$\left. \begin{aligned} \Psi(t) &= \frac{1}{2\pi} \int_{-\infty}^{+\infty} \Phi(\omega) e^{j\omega t} d\omega \\ \Phi(\omega) &= \dots \int_{-\infty}^{+\infty} \Psi(t) e^{-j\omega t} dt \end{aligned} \right\} \quad (11)$$

La première correspond à (8) ; par besoin de symétrie les limites ont été étendues à $-\infty$, ce qui n'apporte aucune différence, puisque $\Phi(\omega) = 0$ pour ω négatif. Nous avons également ici représenté $2\pi(-b + ja)$ par $\Phi(\omega)$ ce qui fait passer le facteur $1/2\pi$ dans la première équation.

De telles transformées de Fourier sont toutefois plus symétriques que celles de type usuel, représentées par l'équation (1) puisque Ψ et Φ sont complexes.

5. Le temps considéré comme variable complexe.

L'analyse de la dernière section a utilisé t et ω comme variables réelles (de Fourier). Par le procédé esquissé dans la section 3, nous pouvons permettre à t de devenir complexe, en utilisant la relation $\gamma = \tau + jt$. (11) se convertit alors en une paire de transformées de Laplace :

$$\left. \begin{aligned} \Psi(\gamma) &= \int_{-\infty}^{+\infty} \Phi(\omega) e^{-\gamma \omega} d\omega \\ \Phi(\omega) &= \frac{1}{2\pi j} \int \Psi(\gamma) e^{\gamma \omega} d\gamma \end{aligned} \right\} \quad (12)$$

au sujet desquelles il y a lieu d'apporter plus d'attention au choix du contour. Avec ces conventions $\gamma = (\tau, jt)$ devient le plan du *temps complexe* sur lequel on peut établir des fonctions de cette variable complexe, correspondant à (5). Nos signaux analytiques de temps doivent être considérés comme sections sur le vrai axe des temps jt .

$\Psi(t)$ est une fonction complexe ; si elle doit devenir fonction d'une variation complexe $(\tau + jt)$, les équations de Cauchy-Riemann doivent être satisfaites :

$$\frac{\partial f}{\partial \tau} = \frac{\partial h}{\partial \tau} ; \quad \frac{\partial h}{\partial t} = - \frac{\partial f}{\partial \tau} \quad (13)$$

De telles équations fournissent, en principe, les données nécessaires à l'établissement complet de la fonction complexe $\Psi(\gamma)$, débutant avec $\Psi(t)$ le signal analytique original.

Il y a lieu de remarquer que si $F(z) = u(x, y) + jv(x, y)$, n'importe quelle fonction d'une variable complexe, u et v formeront alors une paire de fonctions en quadrature le long de n'importe quel axe laissant tous les pôles d'un même côté. En conséquence, dans l'équation (12), le contour d'intégration de $\Psi(\gamma)$ doit être tel que tous les pôles de Ψ lui soient d'un même côté.

Il y a lieu de remarquer que la variable temps dans l'équation différentielle elle-même, représentant les performances physiques d'un circuit quelconque, peut être considéré comme complexe.

6. Possibilité de l'analyse pratique des transitoires.

Il semble qu'il y a deux aspects de l'application de telles notions à l'analyse des transitoires. Envisageons les successivement.

a) Nous pourrions envisager quelque réponse aux transitoires représentant de l'intérêt, $f(t)$ et construire la fonction en quadrature $h(t)$ par la transformée de Hilbert (10) et les ajouter, pour obtenir $\Psi(t) = h + jf$. On pourrait alors faire une tentative pour représenter la fonction complexe entière comme une fonction de temps complexe $\Psi(\gamma)$. Il apparaît immédiatement une difficulté qui suggère un défaut de symétrie complète entre les fonctions de temps et de fréquence : c'est-à-dire que les fonctions de fréquence (ou fonctions d'impédance) $F(\lambda)$ sont des fonctions rationnelles puisqu'elles représentent les caractéristiques de réseaux électriques linéaires à constantes localisées ; toutefois les fonctions du temps ou transitoires sont normalement des fonctions irrationnelles. La belle symétrie des transformées de Fourier réside en la possibilité d'inversion de ω et t (eq. (1)) et non en la rationalité ou autres propriétés spéciales des fonctions en transformation $g(\omega)$ et $f(t)$.

Envisageons un exemple. La réponse aux transitoires la plus fondamentale est celle d'un réseau RC, l'exponentielle simple (fig. 1 (c)) :

$$f(t) = 1 e^{-\alpha t} \quad (t \geq 0) ; \quad \text{le zéro étant ailleurs ;}$$

$$\text{De (10), on tire alors : } h(t) = \frac{1}{\pi} \int_0^{+\infty} \frac{e^{-\alpha v}}{v - t} dv$$

Avec la substitution $u = \alpha(v - t)$ la limite de l'intégrale ci-dessus devient $-\alpha t$ et l'intégrale se réduit à :

$$h(t) = \frac{1}{\pi} \int_{-\alpha t}^{\infty} \frac{e^{-u - \alpha t}}{u} du = \frac{-1}{\pi} e^{-\alpha t} Ei(\alpha t) \quad (14)$$

Le « signal analytique complet » $\Psi(t) = h(t) + j f(t)$ a été représenté sur la figure 3 comme une courbe à 3 dimensions :

$$\Psi(t) = 1 e^{-\alpha t} \left[\frac{-1}{\pi} Ei(\alpha t) + j \right] \quad (15)$$

Un examen plus approfondi de cette restriction montre que certaines réponses transitoires doivent effectivement être considérées comme la superposition d'un certain nombre de composantes distinctes ; en particulier, les circuits à caractéristiques réparties se prêtent à une telle analyse plutôt que les circuits à constantes localisées. La réponse transitoire d'un système à constantes réparties peut être regardée comme la somme des échos distincts d'un stimulus appliqué⁽¹⁾ (3). De tels systèmes peuvent peut-être être représentés par des fonctions rationnelles en termes de poles et de zéros sur un plan de temps complexe γ .

6.1. — Facteur de transposition de Heaviside. Transpositions en temps réel et complexe.

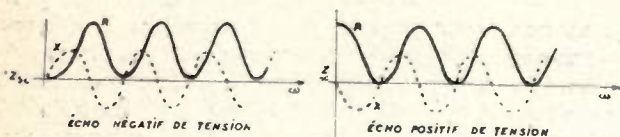
Une ligne de transmission uniforme de longueur 1, soit ouverte soit court-circuitée, provoque un écho simple, positif ou négatif, de tout transitoire appliqué à son origine.

Les caractéristiques en fréquence de l'impédance d'entrée sont de forme sinusoïdale pour le régime permanent :

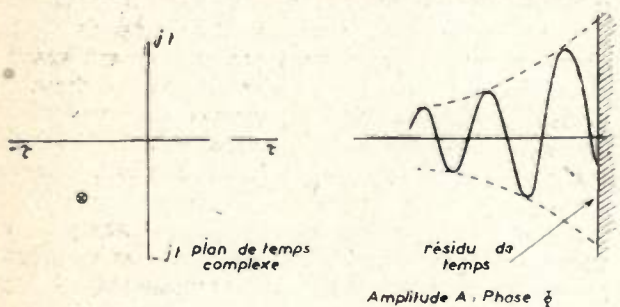
$$\left. \begin{aligned} Z_{sc}(j\omega) &= \frac{R_0}{2} (1 - \varepsilon^{-j\omega t_1}) \\ Z_{oc}(j\omega) &= \frac{R_0}{2} (1 + \varepsilon^{-j\omega t_1}) \end{aligned} \right\} \quad (20)$$

où t_1 est le retard, en temps, de l'écho.

De telles caractéristiques donnent une interprétation physique du facteur de transposition de Hea-



(a) Caractéristiques en fréquence d'une ligne ouverte et court-circuitée. — avec interprétation du facteur de transposition de Heaviside.



(b) Une paire de points conjugués sur le plan de temps complexe représente une caractéristique de fréquence exponentiellement amortie (négative).

FIG. 4. — Echos en temps réel et imaginaire ;

viside $\varepsilon^{-j\omega t_1}$. La figure 4 (a) montre une représentation de ces caractéristiques.

⁽¹⁾ Il en va de même pour un circuit à constantes localisées, mais il faut alors un nombre infini d'échos adjacents s'amalgamant en une transitoire continue (3).

On peut de suite montrer que si une caractéristique en fréquence d'une forme sinusoïdale exponentiellement amortie pouvait être exécutée correspondant à :

$$Z(j\omega) = \frac{R_0}{2} (1 + \varepsilon^{(\tau_1 + j t_1) \omega}) \quad (21)$$

cela correspondrait alors à l'établissement d'échos avec retard en temps complexe.

Remarquer que l'exposant figurant ici y est positif (eq. 21). Il en est ainsi car nous désirons appliquer une paire de signaux en quadrature $\Psi(t)$ à une telle caractéristique. De tels signaux n'ont qu'un spectre négatif ($|\omega|$ négatif), comme donné par l'équation (8). Ignorant la composante constante de Z , dans l'équation (21), la réponse de cette impédance à un tel signal sera :

$$\int_0^\infty (-b(\omega) + ja(\omega)) \varepsilon^{-j\omega(t - t_1 + j \tau_1)} \quad (22)$$

un écho de $\Psi(t)$ retardé d'un temps complexe $(t_1 - j \tau_1)$. Un tel temps, donné par l'exposant de (21), est représenté par un point unique du plan de temps complexe sur lequel nous désirons représenter le développement dans le temps des fonctions d'impédance ; qu'est-il advenu du point conjugué ?

La réponse est simple : nous nous occupons actuellement de transformées entre deux sortes de paires de fonctions en quadrature :

1° Les impédances $\Phi(\omega)$ avec parties réelle et imaginaire $a(\omega)$, $jb(\omega)$ et

2° Les réponses transitoires $\Psi(t)$ avec parties réelle et imaginaire $h(t)$, $jf(t)$.

L'équation (11) représente une telle transformation.

C'est toutefois essayer de tirer le maximum de deux mondes. Tandis que dans l'analyse normale de Fourier, nous nous occupons d'impédances complexes $a(\omega)$, $jb(\omega)$ et de réponses transitoires réelles $f(t)$, nous devrions maintenant ne nous occuper que de réponses aux transitoires complexes $h(t)$, $jf(t)$ et d'impédances réelles $a(\omega)$. Nous n'avons pas besoin à la fois de $a(\omega)$ et de $jb(\omega)$.

Dans ce cas la fonction d'impédance composante qui établit une paire d'échos conjugués est :

$$Z(j\omega) = \frac{R_0}{2} (1 + \varepsilon^{\tau_1 \omega} \cos t_1 \omega) \quad (23)$$

au lieu de (21).

Toute fonction d'impédance pratique (partie réelle) est à considérer comme la superposition linéaire de telles caractéristiques « exponentielles amorties » ; la deuxième intégrale de l'équation (11) peut être ainsi considérée.

6.2. — Résidus du plan de temps complexe.

L'analyse de la section 6.1 a été faite en termes d'intégrales de Fourier et de spectres continus

Il n'y a peut-être que peu d'intérêt à poursuivre encore cette étude académique, mais on peut peut-être remarquer quela transformée de Laplace, avec temps complexe γ , de l'équation (12), peut aussi être interprétée en termes d'échos. Dans une telle analyse, le concept de sommation des spectres continus devient la sommation d'un nombre fini de modes normaux ; si les résidus peuvent être trouvés aux poles des fonctions de transitoires $\Psi'(\gamma)$ (ou si la technique de la cuve électrolytique (1) peut être adaptée au plan de temps complexe), ces résidus représenteront les amplitudes et phases initiales (A, Φ) d'un nombre fini de caractéristiques « exponentielles amorties » dont la somme représente $Z(\omega)$; (comme dans la figure 4 (b) mais omises dans l'équation (23).

Ainsi que nous l'avons déjà vu, toutefois, l'emploi de l'analyse en temps complexe conduit à des difficultés mathématiques, en pratique, lorsque ce concept d'un nombre fini de composantes (analogues aux modes normaux) est appliqué aux caractéristiques impédance fréquence.

BIBLIOGRAPHIE

- [1] BOOTHROYD A. R., CHERRY E. COLIN, MAKAR R. « An Electrolytic Tank for the Measurement of Steady-State Response, Transient Response, and Allied Properties of Networks » *Proc. Inst. Elec. Engs.* (London), 96 Part I May 1949.
- [2] BOOTHROYD A. R. « Design of Electric Wave Filters, with the Aid of the Electrolytic Tank » *Proc. Inst. Elec. Engs.* (London) 98 Part IV 1951 (Monograph N° 8).
- [3] CHERRY E. COLIN « Pulse Response : a New Approach to Alternating Current Theory and Measurement » *Jour. I.E.E.* 92 Part III. Sept 1945.
- [4] GABOR D. « Theory of Communication », *Jour. I.E.E.* 93, Part III, nov. 1946.
- [5] WESTCOTT J. C. « The Frequency Response Method » *Trans. Soc. Inst. Tech.* (London) 4. Sept. 1952, p. 113.
- [6] OSWALD J. « Les Signaux Analytiques et Leurs Transformations ». *La Cybernétique (Editions de la Revue d'Optique Théorique et Instrumentale, 165, rue de Sèvres, 3 et 5, boulevard Pasteur, Paris, 1951.*
- [7] PROWSE W. A. « Solid Diagrams Illustrating Resonance Phenomena » *Proc. Phys. Soc.* 60 o. 131, 1948.

LA TECHNIQUE DES IMPULSIONS ET LES BREVETS

PAR

I. INCOLLINGO

doct-ing. él. (Rome), ing. radio-él. (Liège)
Institut International des Brevets, La Haye.

Le brevet joue dans l'économie moderne un rôle de plus en plus important qui se fait sentir surtout dans les techniques nouvelles.

Or, une des plus grandes contradictions dans le domaine de la propriété industrielle est le fait que, vis-à-vis d'une technique en évolution toujours croissante, les lois, les procédures, bref toute la technique du brevet à l'échelon national n'a subi que peu de changements.

Si le développement de la technique n'a plus, aujourd'hui, un caractère national, le brevet, par contre, a conservé cet aspect national et, dans chaque pays, s'inspire encore des vieux principes et des vieilles procédures qui permirent sa création. Tel qu'il est aujourd'hui, le brevet est un instrument statique : il n'apporte aucune contribution directe au développement de la technique dans le sens qu'il ne favorise pas l'orientation des techniciens sur la nouveauté au moment même où ces derniers essayent de s'orienter ou sont dans la recherche de la nouveauté.

Par contre, la technique des impulsions, comme l'électronique en général, présente par rapport aux autres techniques un aspect nouveau : elle se caractérise par une tendance à l'universalité de ses applications et à la standardisation de ses moyens. Vue d'un bureau de documentation, elle fait tache d'huile et s'étend à plusieurs, sinon à toutes les branches de l'industrie. Elle est pénétrante. Les dispositifs et les procédés que l'on trouvait en détection électromagnétique ou en télévision, se retrouvent appliqués aujourd'hui au mécanisme de contrôle et de réglage de n'importe quel montage, machine ou installation mécanique, chimique ou électrique, et telles parties d'une installation de télécommunication se retrouvent, par exemple, dans les machines à laver. La technique des impulsions se développe dans des dispositifs de base et dans ses applications. Sa tendance à l'universalité devrait nous montrer le chemin à suivre dans tous les domaines qui se rapportent à la technique en général, y comprise la technique du brevet : standardisation des méthodes de travail et mise en commun des efforts dans le but de faciliter la recherche de moyens nouveaux. Nous nous proposons d'illustrer ces exigences en considérant les trois points suivants :

1° Le brevet à l'échelon national ;

2° La recherche sur l'état de la technique pratiquée par l'Institut international et appliquée à plusieurs pays ;

3° Les difficultés qui se présentent dans l'organisation de la classification et de l'étude systématiques de la littérature concernant la technique des impulsions.

* * *

On sait que le but du brevet est de récompenser les efforts et les mérites de l'inventeur par une protection de son œuvre. Dans chaque pays, l'Etat se charge d'accorder les brevets aux ayant-droit par l'intermédiaire d'un office de la propriété industrielle. Si le principe de l'octroi du brevet et cette institution sont choses communes à tous les pays, il n'en est pas de même en ce qui concerne la procédure de l'octroi.

En effet, il y a un certain nombre de pays (France, Italie, Belgique, etc...) où l'obtention du brevet se résume, si l'on peut ainsi s'exprimer, à une simple formalité : tous ceux qui croient avoir inventé quelque chose n'ont qu'à préciser l'objet de leur invention dans une demande qui automatiquement devient un brevet après un certain laps de temps. Le brevet s'obtient par simple dépôt déclaratif.

Dans d'autres pays (Etats-Unis d'Amérique, Grande-Bretagne, Allemagne, Pays-Bas), l'Etat soumet les demandes de brevet à un examen préalable. L'office des brevets dispose à cet effet d'un personnel spécialisé et d'une importante documentation. La demande précisant l'objet de l'invention est examinée à l'aide des brevets déjà existants et de la littérature technique.

On voit donc la première différence essentielle entre les procédures de ces deux groupes de pays : l'examen préalable. Nous évitons de discuter le pour et le contre : chaque groupe comporte ses partisans et ses adversaires et les uns comme les autres ont leurs bonnes raisons.

Pays du premier groupe : à toute demande est accordé un brevet. Celui qui croit avoir réalisé une invention antérieurement à une autre n'a qu'à le prouver devant un tribunal.

¶ Pays du deuxième groupe : on constate plusieurs différences de procédures et on peut affirmer qu'il est presque impossible de trouver dans ce groupe deux pays pratiquant l'examen d'après les mêmes principes. *Que faut-il au juste examiner ?* Un point sur lequel tout le monde est d'accord, est l'examen de la *nouveauté*. C'est-à-dire : le brevet ne sera accordé que si l'objet de la demande n'est pas décrit dans la littérature précédente ou n'est pas généralement connu. Si l'objet est décrit en partie, l'invention portera sur celle qui ne l'est pas encore. Un point, par contre, sur lequel l'accord de tous les offices d'examen préalable ne se réalise pas, est celui de l'examen de la demande sur la base de la nouveauté et en même temps (voilà le point noir !) sur la base du *niveau inventif*. Imaginons, pour fixer les idées, un système multiplex constitué par des modulateurs d'amplitude, par un convertisseur (commun à tous les canaux) des impulsions modulées en amplitude en impulsions modulées en position, utilisant encore une ligne à retard comme distributeur et une impulsion de durée plus grande pour la synchronisation. Ce système a été breveté par un premier inventeur X. Imaginons qu'un deuxième inventeur Y désire avoir pour des raisons évidentes de concurrence commerciale, un brevet sur une installation possédant les mêmes caractéristiques techniques. Monsieur Y rédige alors une demande comme suit (je vous donne un exemple très banal) :

On connaît des systèmes multiplex comportant un modulateur d'amplitude dans chaque canal et un convertisseur de modulation commun à tous les canaux. J'apporte à ce système connu un perfectionnement consistant à placer cette installation sur le sommet d'une colline ou sur une tour en bois. D'où une première revendication comme suit : installation multiplex comportant toutes les caractéristiques du système précisé dans le brevet de Monsieur X, mais avec la caractéristique additionnelle que le même système est placé sur le sommet d'une colline. Ensuite une deuxième revendication avec la caractéristique additionnelle que le système est placé sur une tour en bois.

Examinons cette demande du point de vue « nouveauté pure ». Pour anticiper l'objet de la demande, il faut trouver une publication faisant explicitement mention des caractéristiques X en combinaison avec les caractéristiques additionnelles Y. Il est probable, du moins on peut le supposer, qu'aucun système X comportant les caractéristiques additionnelles de la colline ou de la tour en bois n'est décrit dans la littérature technique. Donc, Monsieur Y peut prétendre à un brevet dans les pays où l'examen de la demande est effectué sur la base de la nouveauté pure.

Mais l'examineur de l'office pratiquant l'examen préalable sur la double base de la nouveauté et du « *niveau inventif* » peut, malgré l'absence d'une antériorité complète, refuser le brevet en invoquant plusieurs arguments. Il peut alléguer que le fait de placer le système multiplex de Monsieur X sur une colline ou sur une tour en bois est à la portée de tout le monde et que la demande n'a donc aucun mérite inventif. D'après l'examineur, Monsieur X aurait pu installer son système sur un gratte-ciel, peut-être

même y en a-t-il déjà placé un, bien qu'il ne l'ait pas mentionné explicitement dans la description. L'examineur peut encore soutenir que la colline et la tour en bois d'une part, et le multiplex X de l'autre, n'ont aucun *trait commun* : la colline ou la tour en bois ne peut pas être considérée comme un perfectionnement ultérieur au système multiplex X.

Et, poussant les choses un peu plus loin, l'examineur peut encore prétendre que, même s'il y avait un mérite à placer le multiplex X sur le sommet d'une colline, et même s'il y avait bien entre eux un trait commun, la caractéristique additionnelle de placer le multiplex sur une hauteur serait susceptible de s'appliquer non seulement au multiplex X mais à *tous les systèmes multiplex*. On ne voit pas alors pourquoi les caractéristiques Y s'appliqueraient particulièrement au système X.

Pour passer à un exemple moins banal, reprenons le système multiplex X et imaginons que Monsieur Y le caractérise cette fois par le fait que l'antenne émettrice est une lentille métallique. Imaginons, comme dans la réalité, que la lentille soit connue en elle-même de même que son application comme antenne émettrice T.H.F. Supposons encore que la combinaison de cette antenne avec un système multiplex ne soit pas connue. Sous l'angle de la nouveauté pure, l'installation pourrait faire l'objet d'un brevet. Mais les arguments précédents restent valables à l'encontre de son caractère inventif. Une fois connues, en effet, la lentille et son application comme antenne, il est à la portée de tout technicien de penser à l'appliquer au système multiplex X. Ce remplacement d'antenne n'a aucun mérite inventif et n'offre aucune difficulté technique. De plus, l'antenne et le multiplex X n'ont pas de trait commun : on ne voit pas comment la lentille pourrait constituer un perfectionnement au système X ou à un système multiplex quelconque se caractérisant par un procédé de modulation des impulsions. Enfin, si la lentille peut s'appliquer au système X, elle pourrait s'appliquer également à tous les systèmes multiplex.

A cette première différence de procédure due à l'introduction du niveau inventif il faut en ajouter une autre concernant les pays se limitant à l'examen de la nouveauté pure. Le principe qui inspire, dans certains de ces pays, les examinateurs ou les juges, est qu'une seule antériorité peut être opposée à la demande de brevet ou au brevet. Pour qu'il y ait antériorité, il faut que l'identité des deux brevets soit absolue et qu'une seule antériorité présente tous les éléments de l'invention. Dans d'autres pays, par contre, on tolère l'opposition de plusieurs antériorités se rapportant chacune aux divers éléments du dispositif objet de la demande. On admet dans ces derniers la possibilité d'opposer à la demande une *mosaïque* d'antériorités. La discussion de ces deux théories nous porterait trop loin. Constatons simplement que par rapport au temps (1880) où fut émis, en Angleterre, l'avis qu'une seule antériorité était opposable, les connaissances et les moyens dont dispose l'homme du métier se sont considérablement accrus.

L'invention met donc mieux en évidence ses mérites si elle est précisée par rapport à ce qui la précède, c'est-à-dire par rapport à l'état complet de

la technique, qui ne peut ressortir aujourd'hui que de plusieurs publications.

Une dernière différence de procédure de très grande importance est celle qui concerne le respect du principe de l'*unité d'invention*, qui veut qu'un même brevet ne peut couvrir qu'une seule invention. Pour reprendre notre exemple du système multiplex : Monsieur X peut souhaiter compléter son système en y appliquant un dispositif de synchronisation capable de produire deux impulsions rapprochées, un amplificateur à contre-réaction spécial, un générateur d'ondes porteuses à cristal et une antenne comportant une lentille. Il peut vouloir le compléter encore en y appliquant un dispositif de signalisation, spécial ou non, en ajoutant par exemple dans chaque canal un dispositif qui réduise le nombre d'impulsions de canal lors de la signalisation. Un tel système multiplex pourra sans doute faire l'objet d'une demande et, après examen, d'un brevet anglais ou américain. D'après la pratique anglo-saxonne, en effet, toutes les caractéristiques essentielles de ce système, à savoir : un procédé de conversion de modulation, une méthode de synchronisation, un dispositif d'amplification, un générateur haute fréquence, une antenne, un dispositif de signalisation ou appel, *ont un trait commun* : un nouveau système multiplex de base comportant un convertisseur de modulation. Puisque ce système de base est nouveau, l'application à cette nouvelle installation de dispositifs connus ou non peut, dès lors, être considérée comme nouvelle et brevetable. Nous trouvons ainsi des brevets américains, anglais, etc... dans lesquels plusieurs caractéristiques sont décrites et revendiquées. Mais, la pratique d'autres pays est plus sévère et n'admet que l'on n'apporte dans la même demande qu'un seul perfectionnement à la fois, à l'une des parties qui composent le système. Si Monsieur X a créé un système multiplex nouveau en imaginant une nouvelle transformation de modulation, et qu'il désire en outre compléter son installation par une *nouvelle* méthode de génération de l'onde T.H.F., il ne pourra pas revendiquer cette deuxième invention dans le même brevet. Cette dernière n'apparaît ni comme une conséquence immédiate, ni comme un développement ultérieur de la méthode de transformation de modulation et pourrait donc tout aussi bien s'appliquer à un système multiplex n'ayant pas la caractéristique du système de base X.

Nombreuses sont les conséquences de cette disparité des législations, des méthodes de travail, des interprétations et, comme nous le verrons plus loin, des classifications de la littérature technique. Un brevet accordé aux Etats-Unis peut comporter plusieurs points inventifs en combinaison ou en particulier. L'objet de la même demande pourra prendre dans le texte du brevet anglais ou suédois correspondant une forme plus limitée ou présenter un autre ordre de succession des points inventifs. Car la demande sera étudiée par un autre examinateur, elle sera classée dans d'autres groupes, la recherche dans une autre classification pourra révéler des antériorités restées inconnues lors du premier examen,

l'appréciation de ces antériorités pourra être différente. Si bien que le brevet anglais ou suédois, s'il est accordé, portera dans le texte et souvent dans la substance même des discordances avec le brevet américain correspondant. Et la même chose se produisant dans le brevet allemand, danois ou hollandais, il sera souvent très difficile de reconnaître que plusieurs brevets ont la même origine. Ce qui est plus grave encore c'est qu'une demande brevetée dans un pays, peut être refusée dans un autre pour des raisons qui échappent au premier examinateur ou doit y subir des modifications dictées par des raisons purement juridiques : souvent ces modifications constituent une véritable démolition du brevet, inspirée par des motifs rédactionnels. Tout cela peut facilement mener à la disparition de l'aspect technique proprement dit. Tout se passe comme si dans chaque pays on voulait habiller l'objet technique de la demande en costume national. Or, il y a des costumes qui ne gênent pas le fonctionnement du dispositif ou du procédé : mais il y a des costumes qui les empêchent tout-à-fait de fonctionner conformément à l'idée de l'inventeur. A ces suites désastreuses du point de vue technique, viennent s'ajouter d'évidentes conséquences financières. Le dépôt de la demande dans plusieurs pays, la nécessité de suivre et combattre les observations des offices coûtent très cher. L'examen dure dans certains pays plusieurs années.

Les inconvénients de cette situation ont suscité depuis plusieurs années diverses initiatives tendant à leur porter remède par la voie de la collaboration internationale. Déjà en 1920, on avait proposé la création d'un Bureau International des Brevets. Cette tentative, et d'autres encore, d'ailleurs peu enthousiastes, avaient échoué à cause de la tâche trop ambitieuse qu'on s'était proposée : on voulait, en effet, brûler les étapes et instituer d'un coup un brevet international. Après la guerre cependant, même dans les pays qui ne pratiquaient pas l'examen préalable, s'était affirmée la nécessité de disposer d'un office de documentation susceptible de donner aux intéressés des avis sur la nouveauté des inventions. Or, l'organisation d'une documentation moderne n'est pas chose facile : elle implique la disponibilité de capitaux énormes, la formation d'un personnel spécialisé dans ce genre de travail et, même si l'on dispose de ces deux éléments, ne peut être achevée qu'au bout de cinq ans. La nécessité, donc, de documenter les intérêts amena les pays du Benelux et la France à la création de l'Institut International des Brevets de la Haye, chargé d'après l'Accord diplomatique du 6 juin 1947, de donner aux Gouvernements des Etats parties à l'Accord, des *avis motivés sur la nouveauté des inventions objet ou non de demandes de brevets*.

La mission de l'Institut International apparaît ainsi très étendue en ce sens qu'elle s'attache à toutes les phases successives du développement de la technique.

Il effectue des recherches portant sur la nouveauté d'inventions faisant l'objet de demandes de brevet.

Les frais inhérents au dépôt d'une demande de brevet dans plusieurs pays, et notamment dans ceux qui pratiquent l'examen préalable, sont, en effet, relativement élevés. Or, une convention internationale accorde au déposant d'une demande dans un premier pays un délai d'une année pour déposer dans d'autres des demandes similaires qui bénéficient de la date du dépôt prioritaire.

Pour n'engager qu'à bon escient des frais importants dans les pays étrangers, comme pour n'affronter qu'en toute connaissance les nombreuses observations de leurs offices, il est donc très avantageux de connaître, avant l'expiration du délai de priorité conventionnel, les antériorités susceptibles d'affecter la nouveauté de l'invention.

C'est ainsi que Monsieur X qui a imaginé le système multiplex précédent, pour lequel il dépose aujourd'hui (5 octobre 1953) à Paris une demande de brevet, peut attendre jusqu'au 5 octobre 1954 pour déposer également cette demande dans d'autres pays avec la même priorité française du 5 octobre 1953. Avant de s'exposer aux dépenses certainement élevées et peut-être inutiles de ces seconds dépôts, il peut ainsi mettre à profit le délai qu'on lui laisse, pour s'adresser à l'I.I.B., en vue d'obtenir un avis qui relève les antériorités se rapportant éventuellement à l'objet de sa demande.

Les recherches de l'I.I.B. concernent également la nouveauté des inventions faisant l'objet de brevets déjà délivrés. En l'absence d'un examen préalable, les demandes de brevets déposées en France, en Belgique ou en Italie, deviennent *automatiquement* des brevets. Ceux-ci, aussi bien d'ailleurs que ceux octroyés dans les pays qui pratiquent un examen limité aux seuls documents nationaux, ont, de ce fait, une *valeur commerciale assez incertaine*. Dans l'ignorance d'éventuelles antériorités, l'inventeur confiant dans la validité de son brevet, pourrait se décider à l'exploiter ou à en accorder des licences. Plus tard, le titulaire d'un brevet antérieur pourrait l'attaquer en contrefaçon devant les tribunaux et l'obliger à cesser son activité avec toutes les difficultés financières qui en découlent. De même, l'industriel qui voudrait acheter ce brevet, tiendrait au plus haut point à se mettre à l'abri d'un tel risque. Pour revenir à notre exemple, Monsieur X aurait pu breveter en France son système multiplex en 1950. Ne l'ayant pas exploité, il pourrait aujourd'hui se décider à le faire ou à céder des licences. Monsieur X ou l'industriel intéressé à l'achat du brevet, voudrait savoir si avant la date de priorité (1950) il existait des antériorités au système multiplex en question. L'I.I.B. se charge sur demande d'effectuer une recherche relative à l'objet du brevet limitée à la date de priorité en France (1950), ou à une date quelconque.

Les avis émis par l'I.I.B. sont des rapports objectifs consignants le résultat de ses recherches sans y porter aucune appréciation subjective ni jugement à l'égard de l'objet de la demande. Si Monsieur Y désire une recherche des antériorités susceptibles d'être opposées à sa demande ou à son brevet, l'avis de l'I.I.B. mentionne les brevets et les publications qui décrivent le système X. Et si dans ces publications ne

figure aucune mention relative aux caractéristiques additionnelles de la tour en bois ou de la colline, l'avis cite d'autres brevets ou publications décrivant des systèmes multiplex n'ayant pas les caractéristiques X mais comportant les caractéristiques additionnelles. L'avis peut être : « le brevet français N°... décrit un système multiplex comportant une transformation de modulations qui... » « le brevet américain N°... a trait à un système multiplex à division de fréquence, comportant une antenne directive placée sur une tour en fer ou sur une hauteur artificielle ».

Et l'avis s'arrête là en excluant toute considération touchant à la brevetabilité, au niveau inventif, etc... mais en fournissant une liste d'antériorités qui donne un tableau précis et complet de l'état de la technique.

* * *

Cette conception objective de l'étude et de l'indication des antériorités a conduit l'I.I.B. à répondre plus directement aux problèmes que pose la recherche technique elle-même en se chargeant de tous travaux documentaires spéciaux relatifs à l'état de la technique dans un secteur déterminé et concernant une construction ou un procédé de fabrication particuliers. En effet, tout technicien chargé de trouver une solution à tel ou tel problème peut exposer ce problème à l'I.I.B. qui lui fournit la documentation concernant les solutions déjà réalisées. Connaissant ces diverses solutions, le technicien est à même de choisir celle qui lui paraît la meilleure ou d'y découvrir l'inspiration qui lui permette d'en réaliser une nouvelle. C'est ainsi que Monsieur Y, chargé par un industriel d'étudier un nouveau système multiplex a intérêt, avant même de commencer ses recherches ou d'organiser son laboratoire, à demander une documentation limitée par exemple aux systèmes multiplex comportant une transformation de modulation. Ayant pris connaissance de ce qui existe dès à présent dans ce domaine en France, et dans les autres grands pays industriels, il pourra plus facilement orienter ses recherches et éviter d'aboutir après des efforts et des dépenses inutiles, à un résultat qu'il croyait nouveau mais qui ne l'était pas.

* * *

La création de l'I.I.B. comble encore une autre lacune. Imaginons qu'un industriel français veuille exploiter un nouveau procédé ou exporter un nouveau produit, breveté en France, dans un autre pays, par exemple en Belgique ou en Italie. Il veut savoir si en Belgique ou en Italie, son procédé ou son produit ne fait pas l'objet d'un brevet belge ou italien. Cela comporte une recherche dans les seuls brevets belges ou italiens. Puisqu'en Belgique, comme en Italie, l'examen préalable n'est pas pratiqué, les brevets belges ou italiens ne sont pas « pratiquement » classés. Dans ces conditions, une recherche dans ces pays est économiquement impossible. Grâce à sa classification systématique et constamment tenue à jour, l'I.I.B. se charge aujourd'hui de classer les brevets des pays signataires de l'Accord, ainsi que les brevets alle-

mands, américains, anglais et suisses. Une recherche limitée aux brevets d'un de ces pays est possible et s'effectue avec grande rapidité.

* *

Le travail de l'I.I.B. n'est plus, dès lors, un travail d'opposition mais un travail de collaboration qui facilite la recherche scientifique et la mise en valeur des inventions méritoires. La conception objective de l'état de la technique qui est la sienne, fait mieux ressortir en effet, les améliorations et les avantages offerts par les inventions nouvelles. Ce même état de la technique constitue d'autre part une source de renseignements précieux pour les chercheurs qui sont arrivés à un résultat connu, en leur fournissant une nouvelle base de départ pour l'orientation et l'aboutissement heureux de leurs futures recherches.

* *

Après avoir rapidement brossé un tableau des aspects qui différencient les procédures des offices à examen préalable, étudions le point essentiel de cet examen : la recherche. Les documents de base capables d'offrir les meilleures antériorités sont les brevets, les revues techniques, les livres scientifiques : tout cela sur une échelle internationale. Il existe aujourd'hui presque trois millions de brevets américains, plus d'un million de brevets français, 900.000 brevets allemands, 800.000 anglais, 700.000 italiens, 500.000 belges, 300.000 suisses, 80.000 hollandais, soit au total pour ces seuls pays 7 millions de brevets. Une recherche dans cet océan de documents serait chose impossible sans classification. Ici intervient encore une différence entre les offices à examen préalable, différence qui ne touche pas la procédure, cette fois, mais qui touche, chose plus importante, les résultats de la recherche. En effet, la mise à jour des antériorités dépend, pour une grande part, de la classification. Dans l'impossibilité d'illustrer toutes les théories relatives à ce sujet, nous nous référons à la classification de l'I.I.B. qui procède de la classification continentale (allemande et hollandaise) et qui a évolué en tenant compte des nouvelles tendances de la technique moderne.

L'organisation de la technique des impulsions ne peut pas être considérée comme tout-à-fait achevée. Les problèmes posés à un bureau de documentation par le développement rapide et par la pénétration de cette technique dans toutes les branches de l'industrie sont très nombreux, surtout si l'on pense que ce bureau n'a pas pu suivre progressivement cette évolution. L'Europe, préoccupée par un conflit, n'avait pratiquement aucun contact avec les pays anglo-saxons où un développement maximum allait justement s'effectuer. Même après le conflit, le secret dicté par les exigences militaires, a encore accentué ce manque de synchronisme entre le développement de la technique des impulsions et son organisation dans une classification progressive. Après la guerre, un raz de marée de demandes de brevets se rapportant à la nouvelle technique des impulsions, vint s'abattre sur les offices européens. Il fallait se mettre rapidement au courant

et créer une nouvelle classification. Or, un office de documentation moderne trouve sa raison d'exister dans l'évolution et les progrès de la technique et doit donc être préparé aux déplacements de la technique. La vieille classification allemande comportait en 1910 une seule classe 21 a^1 pour les communications électriques, c'est-à-dire pour l'électronique de jadis. Les progrès réalisés dans ce domaine demandaient déjà en 1930 la création de 4 classes séparées : 21 a^1 (télégraphie) ; 21 a^2 (appareils téléphoniques, lignes de transmission amplificateurs) ; 21 a^3 (trafic) ; 21 a^4 (radio-communications). La classe 21 a^1 comptait 233 subdivisions. Par cette vieille classification, on peut se rendre compte que l'électronique ne possédait pas encore son aspect universel. L'inventeur s'attachait en ce temps-là à perfectionner une installation télégraphique ou une partie de celle-ci. La télégraphie posait des problèmes totalement différents de ceux qui intéressaient la téléphonie. 1936 vit la création d'une 5^e classe 21 a^5 pour la télévision. Jusqu'ici donc, la classification était basée exclusivement sur la nature de la communication (signal télégraphique, téléphonique, musique, image). Les années précédant le deuxième conflit mondial montrèrent déjà la tendance de l'électronique vers l'universalité de ses applications et vers la standardisation de ses moyens. Une nouvelle classification capable de refléter cette tendance s'imposait. Déjà en 1941, l'Octrooiraad hollandais introduisait trois classes nouvelles susceptibles d'accueillir le matériel accumulé : 95 (technique des oscillations électriques) ; 96 (technique des systèmes de transmission) ; 97 (technique des communications). La classe 95 et sa division en groupes reflète la nouvelle tendance. Par des lettres on crée les groupes principaux suivants : 95 a (génération des oscillations), 95 b (modulation), 95 c (démodulation), 95 d (amplification) etc. etc... Chaque groupe principal subit alors des divisions selon les principes scientifiques et des subdivisions selon les moyens utilisés. Les groupes ayant trait à la modulation et à la démodulation sont ainsi scindés : 95 b 1 modulation en amplitude 95 b 2 modulation en fréquence 95 c 1 démodulation en amplitude 95 c 2 démodulation en fréquence. Les sous-groupes qui en dérivent sont caractérisés par les moyens : dans ce nouvel ordre d'idées, devait trouver place la nouvelle technique des impulsions. Si nous caractérisons les impulsions par le chiffre 3, on crée facilement dans la technique des oscillations électriques (classe 95) les groupes 95 a 3 (génération des impulsions), 95 b 3 (modulation des impulsions), 95 c 3 (démodulation des impulsions). Dans la classe immédiatement supérieure 96, on trouve le groupe 96 b (système multiplex). Le groupe 96 b 3 comportera le multiplex à divisions dans le temps, 96 désignant un système de transmission, b un multiplex, 3 des impulsions. La classification du multiplex 96 b 3 comporte ensuite une série de sous-groupes affectant la distribution et ses moyens, la synchronisation, la signalisation, la diaphonie, l'utilisation de moyens ou principes spéciaux (mémoire, etc...). La même constitution de groupes parallèles se retrouve pour les compteurs électroniques et pour les machines à calculer.

Un exemple typique d'application des principes généraux qui inspirent notre classification, tels que les

règles de symétries horizontale et verticale, est fourni par l'organisation de la détection électro-magnétique. La classe 97 /, qui s'y rapporte, est divisée en groupes d'après les procédés scientifiques employés. On trouve :

97 / 1 (détermination de la direction, du plan de guidage, de la vitesse, de la distance, au moyen d'oscillations électriques en général) ;

97 / 2 (détermination de la direction par repérage à minima, repérage par cardioïde de phase) ;

97 / 3 (repérage par comparaison d'amplitude) ;

97 / 4 (détermination de la direction par la méthode des champs tournants) ;

et ainsi de suite jusqu'à :

97 / 17 (repérage par échos) ;

Si on fait suivre chaque numéro caractéristique de la méthode employée par des lettres fractionnant l'installation par exemple :

b) émetteur ;

c) antenne ;

f) récepteur etc... etc...

et si on fait suivre ces dernières lettres par des numéros caractérisant les procédés de modulation employés : 2, modulation de fréquence, 3, modulation d'impulsions, ou vice-versa, on forme une série de groupes et sous-groupes symétriques dans les sens horizontal et vertical. Exemple : 97 / 17 b 3 sera un groupe comportant des émetteurs b, utilisant la méthode des échos 17 et des impulsions 3. Dans le groupe 97 / 2 (repérage à minima) on crée un sous-groupe correspondant 97 / 2 b 3 comportant des émetteurs b utilisant la méthode de repérage à minima 2 et une modulation d'impulsions 3. Par le simple changement d'un chiffre on constitue deux sous-groupes parallèles dans deux groupes différents. La symétrie horizontale concerne le même groupe en ce sens que le sous-groupe 97 / 17 / 3 sera un récepteur, méthode échos, modulation d'impulsions, tandis que 97 / 17 b 3 est l'émetteur correspondant. Deux sous-groupes dans le même groupe sont ainsi obtenus par changement d'une lettre.

Par ce système de classification, on réalise des divisions logiques et très faciles à retenir. Dans d'autres classifications par contre (voir l'américaine) aucune connexion logique n'existe entre sous-groupes. Chaque sous-groupe est numéroté systématiquement et sans aucune corrélation avec le sous-groupe des autres classes. Cette énumération appliquée à la classification de la détection électro-magnétique d'après notre système donnerait, grosso modo, la classe 97, les groupes 97 a, 97b, 97c, et les sous-groupes 97 a 1... jusqu'à 97 a 50, 97b 1... 97b 50 ; 97 c 1... 97 c 50. Les sous-groupes 97 a 32, 97b 32, 97c 32 n'auraient aucune connexion entre eux.

La technique des impulsions se trouve dans notre classification répartie dans 7 classes, 29 groupes, 82 sous-groupes. Le nombre de brevets classés dans ces différentes subdivisions dépassait à la fin du mois de juin 1953 les 12.000. Dans ces subdivisions est encore classé un nombre considérable de cartes portant le résumé des articles publiés dans 51 revues d'électronique. Dans ces 12.000 brevets ne sont pas compris des centaines de brevets concernant les dispositifs par impulsions appliqués aux diverses bran-

ches de l'industrie (servo-mécanismes, appareils de mesure en électricité, dispositifs de réglage dans les machines en général, etc..., etc...). Ces derniers brevets utilisent la technique des impulsions telle qu'elle était au moment de son apparition. Ils n'apportent en général aucune contribution ou aucun aspect nouveau à la technique des impulsions proprement dite. Ils volent, pour ainsi dire, les moyens découverts par les techniciens des impulsions pour les exploiter avec profit dans un autre domaine de la technique.

L'étendue et l'ampleur de la technique des impulsions c'est-à-dire sa pénétration et son universalité, la rendent très difficile à manier. Pensez à la télévision en couleurs utilisant un système multiplex à divisions dans le temps. Pensez que dans chaque dispositif utilisant les impulsions, on a besoin d'une synchronisation et que dans la télévision la synchronisation par impulsions est aussi vieille que la télévision elle-même. Toujours en télévision, les seuls groupes affectant la séparation des impulsions de synchronisations contiennent 2.144 brevets. L'aspect général de la technique des impulsions et sa tendance peuvent être mis en relief par un nombre croissant de brevets ou publications se rapportant à un « Commutateur électronique ». L'inventeur cherche à résoudre le problème de la génération et de la distribution qui est commun au multiplex, au radar, à la télévision, aux machines à calculer, par un seul dispositif capable en même temps de produire et de distribuer des impulsions dont on puisse faire varier facilement la fréquence, le nombre, l'espacement, la forme. L'étude de ces dispositifs très généraux devient de plus en plus laborieuse : leur compréhension implique la connaissance des problèmes qui se posent dans le multiplex, dans le radar, etc, etc... La recherche doit aussi s'étendre à ces branches. Les groupes et les sous-groupes à consulter pourraient être une vingtaine ; les brevets, les revues, les livres intéressants pourraient dépasser le millier. L'effort intellectuel est considérable : la recherche et la sélection des antécédents pourrait durer une quinzaine de jours, tout cela malgré une classification retouchée continuellement.

..

Ces difficultés qui sont inhérentes à la technique elle-même et qui tiennent au niveau très élevé qu'elle atteint aujourd'hui, sont amplifiées, dans les brevets, par la disparité des procédures nationales, avec toutes ses conséquences désastreuses sur le plan technique et financier. Une standardisation des législations et des méthodes de travail dans le domaine de la propriété industrielle n'en apparaît que d'autant plus indispensable si l'on veut que le brevet reflète la tendance fondamentale des techniques nouvelles vers l'universalité et la standardisation. Un office de brevets conçu sur le plan international comme bureau de documentation au service des chercheurs, peut et doit apporter une contribution directe au développement de ces techniques nouvelles qui, comme l'électronique, pour être jeunes, sont pleines de promesses et de possibilités. La réalisation de cette conception nouvelle serait grandement facilitée si tous les techniciens s'intéressaient davantage, non

seulement à l'obtention d'un brevet, mais à la forme, aux aspects et aux procédures d'une technique du brevet adaptée aux exigences modernes.

Un résultat de cette collaboration pourrait déjà se révéler très utile dans le domaine de la standardisation de la terminologie, du symbolisme graphique et de la schématisation. Notons en effet, pour finir, que dans la technique des impulsions rien n'a été fait jusqu'à ce jour sur ce terrain.

Terminologie : dans chaque langue nationale, on emploie quatre ou cinq expressions pour indiquer par exemple un procédé de modulation. En français, on désigne la modulation en position par les expressions: modulation en position, modulation de ou en phase, modulation de ou en temps, modulation par déplacement. La possibilité de confusion avec d'autres procédés de modulation est très grande, soit pour les examinateurs obligés de travailler dans des documents rédigés en quatre ou cinq langues, soit pour le technicien appelé à consulter des publications étrangères. D'où la nécessité d'établir dans chaque langue nationale, en parallèle avec les autres langues, une définition ou expression pour chaque procédé ou dispositif. Un exemple peut être fourni par les travaux effectués par le Bureau Central de Normalisation des Pays-Bas en matière d'expressions à employer pour les procédés de modulation des impulsions.

Symbolisme graphique : dans le multiplex à divisions de fréquence, les filtres, les modulateurs, etc... sont indiqués par des symboles graphiques. Cela facilite l'interprétation des dessins et évite la correction de ceux-ci lors du dépôt des demandes dans un pays autre que celui d'origine.

Schématisation : la compréhension des dispositifs par impulsions est très laborieuse lorsqu'on veut réaliser le procédé suivi en partant d'un schéma de détail ou d'une forme de réalisation comportant des centaines de tubes et circuits. Par contre, les phénomènes physiques mis en jeu se prêtent facilement à une schématisation de principe. Celle-ci précise rapidement le but et le procédé exploité pour sa réalisation. L'usage systématique de schémas utilisant les symboles graphiques faciliterait énormément la compréhension des documents et leur classification.

J'adresse mes bien vifs remerciements à tous ceux qui m'ont aidé dans l'établissement de cette conférence et tout spécialement Monsieur P. MILLET, Sous-Directeur de l'Institut international des Brevets ainsi que mon collègue Monsieur M. BROCHON.

CAUSES DIVERSES DE DIAPHONIE DANS LES SYSTÈMES MULTIPLEX A IMPULSIONS

PAR

J. FAGOT

Directeur technique à la S.F.R.

I. Généralités.

Le multiplexage est obtenu, comme il est connu, par succession dans le temps des impulsions qui échantillonnent chaque voie (fig. 1).

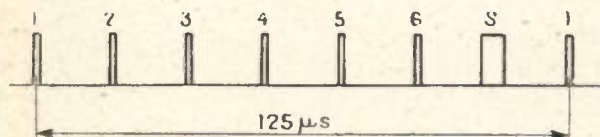


FIG. 1.

La diaphonie apparaît sous la forme d'un résidu exponentiel issu de la voie perturbatrice et modulant de façon parasite la voie perturbée (fig. 2). Cela suppose par exemple que le résidu puisse :

— faire varier l'amplitude des impulsions de la voie perturbée dans le cas de modulation d'amplitude ;

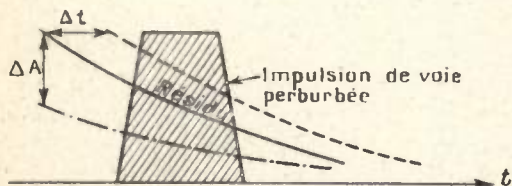


FIG. 2. — Δt , déplacement du résidu ; ΔA , variation en amplitude du résidu.

— en faire varier la durée dans le cas de modulation de durée ;

— faire varier la position des flancs dans le cas de modulation de position.

Une étude systématique des phénomènes, appuyée par l'expérience, nous conduit à classer les effets de diaphonie constatés dans les deux grandes catégories suivantes :

1° Diaphonie « basse fréquence », ou « générale » ;

2° Diaphonie « haute fréquence », ou « de traînée ».

Il convient, par ailleurs, de faire une mention spéciale des systèmes convertisseurs. Ces organes sont

utilisés pour passer d'un genre de modulation à un autre sur les impulsions déjà groupées en Multiplex. Ils peuvent, par leur principe même, être la cause d'importantes diaphonies. Nous retiendrons donc à la suite :

3° Diaphonie des dispositifs « convertisseurs ».

Je suis heureux de signaler ici l'aide précieuse que m'a apportée M.R. Casse, ingénieur à la S.F.R. qui a étudié et mis au point plusieurs équipements Multiplex et obtenu sur ceux-ci par une analyse minutieuse des phénomènes, des résultats remarquables.

2. Diaphonie basse fréquence ou générale.

2.1 Equations générales.

Les impulsions sont transmises à certains stades à travers des amplificateurs « video » qui sont des amplificateurs « résistance capacité » à large bande. Des effets de diaphonie peuvent résulter d'une transmission incorrecte des basses fréquences par ces amplificateurs.

Etudions la transmission à travers un étage amplificateur à résistances (fig. 3) d'une suite d'impulsions toutes identiques, telles qu'en produit un Multiplex en l'absence de modulation des voies.

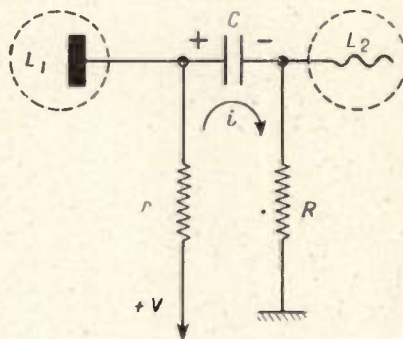


FIG. 3.

Les impulsions, supposées positives et de forme rectangulaire idéale, sont reproduites identiquement sur l'anode de la lampe L_1 . On suppose que la « bande

passante » est très grande, et que l'ensemble : C et R ne charge pas pratiquement la résistance d'anode r , beaucoup plus faible que R . La tension d'anode varie de V , tension de source (point de débit nul) à $V-A$: point de débit maximum (fig. 4). La période θ_1 correspond à la durée d'une impulsion. L'intervalle θ_2 correspond à « l'espace libre » entre deux impulsions successives. Normalement θ_2 est sensiblement plus élevé que θ_1 même dans le cas d'un grand nombre de voies.

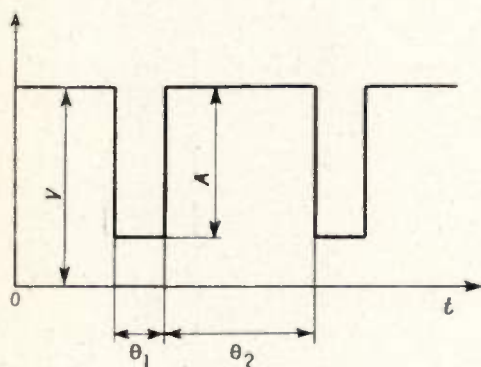


FIG. 4.

Nous nous proposons de calculer la loi de tension sur la grille de L_2 en supposant le régime d'équilibre établi. Au temps $t = 0$, pris en coïncidence avec le début d'une impulsion, on suppose que la grille est au potentiel résiduel a (tout de suite avant l'application du « front avant » de l'impulsion θ_1). A vient s'ajouter immédiatement l'amplitude $-A$ du front avant de l'impulsion, transmise intégralement à travers C .

Considérons le circuit : source d'anode V , r , C et R . Choisissons-y un sens positif de circulation de courant (fig. 3). Nous aurons pendant toute la phase θ_1 (fig. 4).

$$(1) \quad [V - A] - v_c = Ri.$$

v_c représente la tension (variable) aux bornes de C . Normalement l'armature de gauche correspond au pôle +, celle de droite au pôle -.

On écrira naturellement :

$$v_c = \frac{Q}{C} \quad \text{avec} \quad i = \frac{dQ}{dt}$$

Si le courant passait en effet dans le sens défini comme positif, il tendrait à charger le condensateur. En dérivant l'équation (1), on obtiendra :

$$-\frac{dv_c}{dt} = R \frac{di}{dt}, \quad -\frac{1}{C} \frac{dQ}{dt} = -\frac{1}{C} i = R \frac{di}{dt}.$$

L'équation du courant est donc :

$$(2) \quad CR \frac{di}{dt} + i = 0$$

Elle nous permet, par simple multiplication par R , d'écrire l'équation de la tension de grille ($v_g = Ri$)

$$(3) \quad CR \frac{dv_g}{dt} + v_g = 0.$$

La solution classique, est :

$$(4) \quad v_g = (a - A) e^{-\frac{t}{CR}}.$$

Elle est déduite de la solution générale :

$$(5) \quad v_g = K e^{-\frac{t}{CR}}.$$

en tenant compte des conditions aux limites soit :

Pour $t = 0$, v_g est égal à $(a - A)$ (conditions initiales) ;

Pour $t = \infty$, v_g est égal à 0 (il ne peut plus y avoir de courant i).

Cette première phase s'arrête en réalité à $t = \theta_1$, ce qui donne pour tension à ce moment :

$$(6) \quad v_g = (a - A) e^{-\frac{\theta_1}{CR}}.$$

La figure 5 donne l'allure des variations de v_g . On y voit que a est positif, et que pendant la phase θ_1 , la grille est négative.

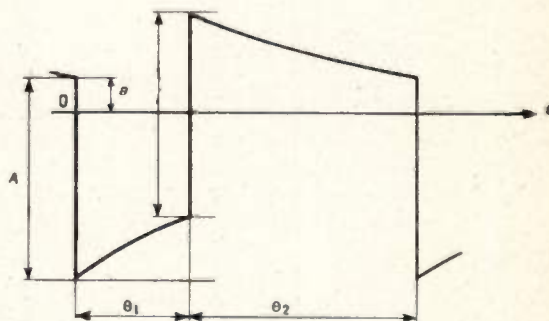


FIG. 5.

Dans la phase θ_2 l'équation (3) reste valable, et donne la même solution générale (5). La constante K est égale à la tension du départ, pour lequel on suppose maintenant $t = 0$. C'est donc la valeur de v_g donnée par (6), à laquelle on doit ajouter A , représentant le front arrière de l'impulsion (transmis intégralement au début de la phase θ_2).

Ainsi v_g prend la forme :

$$(7) \quad v_g = \left[(a - A) e^{-\frac{\theta_1}{CR}} + A \right] e^{-\frac{t}{CR}},$$

soit à la fin de cette phase θ_2 :

$$(8) \quad v_g = \left[(a - A) e^{-\frac{\theta_1}{CR}} + A \right] e^{-\frac{\theta_2}{CR}}.$$

cette valeur étant du reste celle de a .

Nous avons tous les éléments nous permettant d'étudier la diaphonie.

2.2 Diaphonie en modulation d'amplitude.

Supposons, une fois le régime d'équilibre et de non-modulation établi, qu'une des impulsions de voie subisse un accroissement d'amplitude instantané : ΔA . Quelle sera l'action produite sur les voies suivantes ?

L'équation (8) nous donne une réponse immédiate. La tension de perturbation, qui agit en amplitude sur la voie suivante est donnée par la différentielle de (8) par rapport à A , soit :

$$(9) \quad \Delta v_g = \Delta A \left[1 - e^{-\frac{\theta_1}{CR}} \right] e^{-\frac{\theta_2}{CR}}.$$

Mais CR est très grand vis-à-vis de θ_1 et θ_2 , puisqu'on cherche à réduire la diaphonie. On peut, sans inconvénient, faire les approximations suivantes dans la formule (9) :

$$e^{-\frac{\theta_1}{CR}} \approx 1 - \frac{\theta_1}{CR}, \quad e^{-\frac{\theta_2}{CR}} \approx 1 - \frac{\theta_2}{CR},$$

ce qui donne :

$$(10) \quad \Delta v_g = \Delta A \frac{\theta_1}{CR}.$$

Nous nommerons « perturbation de diaphonie » le rapport :

$$(11) \quad \frac{\Delta v_g}{\Delta A} = \frac{\theta_1}{CR}.$$

Ce résultat justifie le nom de *Diaphonie générale*. La perturbation donnée par (11) est la même sur toutes les voies, puisque le résultat donné ne contient pas θ_2 , séparation des impulsions. On le vérifierait du reste facilement.

Si nous supposons que la voie perturbatrice se trouve seule modulée à plein niveau en amplitude, nous pouvons imaginer que sa modulation est décomposée en une suite de perturbations élémentaires $\Delta A_1, \Delta A_2, \Delta A_3, \dots$ (d'amplitudes différentes les unes des autres). Il se produit dans toutes les voies des perturbations élémentaires proportionnelles, représentant les différentes ΔA affectés du coefficient :

$$\frac{\theta_1}{CR}.$$

La modulation est ainsi reproduite dans toutes les voies, en clair, au taux relatif de $\frac{\theta_1}{CR}$ qui représente ainsi le taux de diaphonie. On peut dire que la diaphonie est « linéaire », $\frac{\theta_1}{CR}$ étant constant.

Toutes les voies peuvent être modulées. Si n est le nombre des voies, les puissances de diaphonie s'ajoutent et l'amplitude de la diaphonie globale est multipliée par $\sqrt{n}$ (les puissances de diaphonies s'ajoutent). En même temps la diaphonie acquiert un

certain caractère d'inintelligibilité. On aura la formule finale à retenir :

(12) Diaphonie maximum par voie :

$$\begin{array}{ccc} N & \sqrt{n} & \frac{\theta_1}{CR} \rightarrow \text{Durée des impulsions} \\ \downarrow & \downarrow & \downarrow \\ \text{Nombre} & \text{Nombre} & \text{Constante de temps des liaisons} \\ \text{d'étages} & \text{de voies} & \end{array}$$

N étant le nombre d'étages supposés identiques mis à la suite les uns des autres.

On notera enfin qu'il existe des systèmes de compensation B.F. (constantes de temps sur la source d'anode) qui, de la même façon qu'ils améliorent la transmission des signaux rectangulaires, peuvent réduire la diaphonie.

2.3 Diaphonie en modulation de durée.

Celle-ci se fera, par exemple, par déplacement du front arrière de l'impulsion, c'est-à-dire par variation de la durée θ_1 , la somme $\theta_1 + \theta_2$ restant inchangée. Il ne pourrait y avoir aucune action sur l'emplacement des flancs de l'impulsion suivante si ces flancs étaient parfaitement verticaux. Mais la durée de

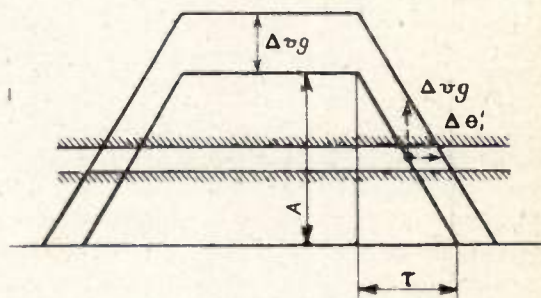


FIG. 6.

montée et de descente des flancs des impulsions présente une durée finie τ , et le mécanisme de la diaphonie est le suivant :

1° L'accroissement $\Delta \theta_1$ de θ_1 produit une amplitude résiduelle. On considérera l'équation (8) sur laquelle un accroissement brusque d'amplitude $\Delta \theta_1$ sera supposé sur θ_1 . On aura toujours :

$$\theta_1 + \theta_2 = \text{const.}, \quad \text{soit } \Delta \theta_2 = -\Delta \theta_1.$$

En prenant la différentielle de (8) par rapport à θ_2 :

$$(13) \quad \Delta v_g = -\frac{A}{CR} e^{-\frac{\theta_2}{CR}} \Delta \theta_2$$

$$(14) \quad \Delta v_g = \frac{A}{CR} e^{-\frac{\theta_2}{CR}} \Delta \theta_1.$$

Soit, étant donné les grandes constantes de temps :

$$(15) \quad \Delta v_g = A \frac{\Delta \theta_1}{CR}$$

Pour les mêmes raisons que précédemment, la diaphonie est générale (résultat indépendant de θ_2).

2° L'amplitude résiduelle fait varier la durée de l'impulsion perturbée.

On se reportera à la figure 6. Après transmission, on réalise normalement une « sélection d'amplitude » sur l'impulsion modulée en durée, c'est-à-dire que l'on agit sur une lampe saturée de façon à ne retenir que les variations d'amplitude comprises entre les limites indiquées sur la figure. A un accroissement de hauteur de Δv_g correspond ainsi un déplacement du flanc de $\Delta \theta'_1$. On voit sur la figure que :

$$(16) \quad \frac{\Delta v_g}{\Delta \theta'_1} = \frac{A}{\tau},$$

donc

$$(17) \quad \Delta \theta'_1 = \frac{\tau}{A} \Delta v_g.$$

En combinant les résultats (15) et (17) :

$$\Delta \theta'_1 = \frac{\tau}{CR} \Delta \theta_1,$$

soit

$$(18) \quad \frac{\Delta \theta'_1}{\Delta \theta_1} = \frac{\tau}{CR}.$$

C'est la « perturbation de diaphonie ». Par un raisonnement identique à ce qui a été dit précédemment, c'est aussi le taux de diaphonie (générale).

D'où la formule :

(19) Diaphonie maximum par voie :

$$\frac{N}{\text{Nombre d'étages}} \quad \frac{\sqrt{n}}{\text{Nombre de voies}} \quad \frac{\tau}{CR} \rightarrow \begin{array}{l} \text{Temps de montée} \\ \text{Constante de temps des liaisons} \end{array}$$

2.4 Diaphonie en modulation de position.

Nous devons, pour pouvoir développer les calculs, considérer ce qui se passe sur deux cycles successifs :

1° Cycle comprenant les phases θ_1, θ_2 ;

2° Cycle comprenant les phases θ'_1, θ'_2 .

En régime de non modulation on a bien entendu :

$$\theta_1 = \theta'_1 \quad \text{et} \quad \theta_2 = \theta'_2,$$

mais la différence est marquée en vue des variations à envisager.

En état de non-modulation, on arrive à la fin du premier cycle à l'amplitude résiduelle exprimée par (8) développé :

$$(20) \quad v_g = (a - A) e^{-\frac{\theta_1 + \theta_2}{CR}} + A e^{-\frac{\theta_2}{CR}}$$

et à la fin du deuxième cycle, à un résultat analogue, à condition de remplacer l'amplitude initiale a , par la valeur de v_g donnée par (20). On écrit ainsi :

$$v'_g = \left[\left((a - A) e^{-\frac{\theta_1 + \theta_2}{CR}} + A e^{-\frac{\theta_2}{CR}} \right) - A \right] \times e^{-\frac{\theta'_1 + \theta'_2}{CR}} + A e^{-\frac{\theta'_2}{CR}},$$

soit en effectuant :

$$(21) \quad v'_g = (a - A) e^{-\frac{\theta_1 + \theta_2 + \theta'_1 + \theta'_2}{CR}} + A e^{-\frac{\theta_2 + \theta'_1 + \theta'_2}{CR}} - A e^{-\frac{\theta'_1 + \theta'_2}{CR}} + A e^{-\frac{\theta'_2}{CR}}.$$

Envisageons alors une modulation de position sur l'impulsion θ'_1 considérée comme perturbatrice, et calculons le résidu diaphonique correspondant. Cette modulation peut être considérée comme se traduisant à un instant donné par un déplacement brusque de

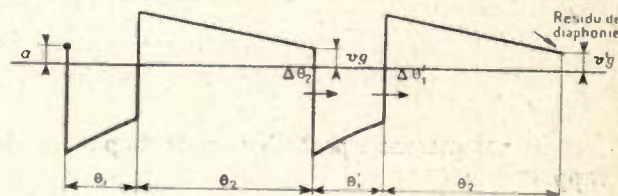


FIG. 7.

Δp de l'impulsion en bloc. Par ce déplacement, il y a déplacement du front avant et du front arrière de la même quantité. Avec nos notations :

Déplacement du front avant : cela veut dire que θ_2 est accru de $\Delta \theta_2$;

Déplacement du front arrière : cela veut dire de même que θ'_1 est accru de $\Delta \theta'_1$.

Les phénomènes étant linéaires, nous pouvons isolément supposer le déplacement de ces fronts, calculer dans chaque cas la perturbation diaphonique, et ajouter les résultats (fig. 7).

1° Front avant : on a $\Delta \theta_2$, avec

$$\theta_2 + \theta'_1 = \text{const.}$$

soit

$$\Delta \theta_2 = - \Delta \theta'_1.$$

La différentielle par rapport à θ_2 donne dans (21)

$$(22) \quad \Delta_1 v'_g = - \frac{A}{CR} e^{-\frac{\theta'_1 + \theta'_2}{CR}} \Delta \theta_2$$

2° Front arrière : on a $\Delta \theta'_1$, avec

$$\theta'_1 + \theta'_2 = \text{const.},$$

soit

$$\Delta \theta'_1 = - \Delta \theta'_2.$$

La différentielle par rapport à θ_1 donne :

$$(23) \quad \Delta_2 v'_g = \frac{A}{CR} e^{-\frac{\theta'_2}{CR}} \Delta \theta'_1.$$

On a finalement l'action globale suivante, avec :

$$\Delta \theta_2 = \Delta \theta'_1 = \Delta p,$$

$$(24) \quad \Delta v'_g = \frac{A}{CR} e^{-\frac{\theta'_2}{CR}} \left(1 - e^{-\frac{\theta'_1}{CR}}\right) \Delta p \# A \frac{\Delta p}{CR} \frac{\theta'_1}{CR}$$

(approximations identiques à celles de l'équ. (9))

Ceci chiffre le résidu en amplitude sur toutes les voies (là encore les constantes de temps sont telles que la diaphonie est générale).

Partant de ce résidu d'amplitude, on doit par un calcul calqué sur celui du paragraphe précédent, calculer la diaphonie en examinant le déplacement par rapport à la modulation utile d'un des flancs de l'impulsion. Nous considérerons dans ce cas la position d'un des flancs comme le repère utilisé pour la démodulation dans le système de modulation en position. On obtient, en utilisant le résultat exprimé par (17)

(25) Modulation de position ; taux de diaphonie :

$$N \sqrt{n} \frac{\tau}{CR} \frac{\theta_1}{CR},$$

N , nombre d'étages ;

n , nombre de voies ;

CR , constantes de temps des liaisons ;

τ , temps de montée des flancs (total) ;

θ_1 , durée des impulsions.

Par comparaison avec les systèmes précédents on voit que la diaphonie est beaucoup plus faible.

3. Diaphonie haute-fréquence ou « de Trainée ».

3.1 Généralités.

Dans un amplificateur « video », on sait que les flancs verticaux des impulsions ne sont pas parfaitement reproduits du fait de la limitation de la bande passante du côté des fréquences élevées. C'est ce qui produit en particulier un temps fini de montée τ , comme nous l'avons déjà envisagé. Une impulsion présentant la forme théorique de la figure 8 peut être considérée pour l'étude des déformations en question comme décomposée en deux « échelons-unité » de signe inverse, décalés dans le temps d'un intervalle égal à la durée de l'impulsion θ_1 (fig. 9).

La méthode est de calculer les réponses de l'amplificateur aux deux échelons unités et de les composer. L'intérêt sera alors particulièrement porté sur les « traînées » qui sont la cause de diaphonie. Voici les différentes étapes du calcul :

1° La réponse du circuit amplificateur est donnée sous une forme « imaginaire » classique pour le régime sinusoïdal (en fonction de $j\omega$).

Il suffit de remplacer $j\omega$ par p , symbole de la dérivation en « calcul symbolique » pour connaître la réponse à l'échelon unité (sous forme symbolique : $F(p)$).

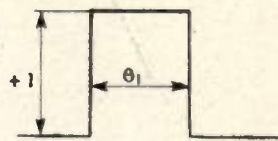


FIG. 8.

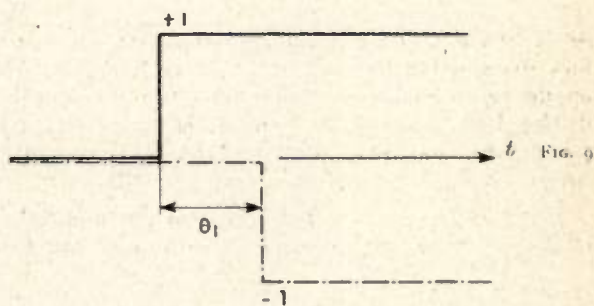


FIG. 9.

2° Il faut en déduire la réponse sous forme utilisable, c'est-à-dire sous la forme d'une fonction du temps : $F(t)$. La correspondance est obtenue en faisant appel aux règles de calcul et tableaux de correspondances du calcul symbolique. Rappelons qu'une des relations utilisées pour exprimer la loi de correspondance, est l'équation intégrale de Carson :

$$(26) \quad F(p) = p \int_0^{\infty} F(t) e^{-pt} dt$$

Dans tous les cas communément rencontrés, qui ont fait l'objet des calculs de nombreux auteurs, les réponses obtenues sont représentées sur les figures 10 et 11.

Les grandeurs définies habituellement sur ces réponses sont :

— Le temps de montée, période pendant laquelle le flanc passe de 0,1 à 0,9 de l'amplitude maximum

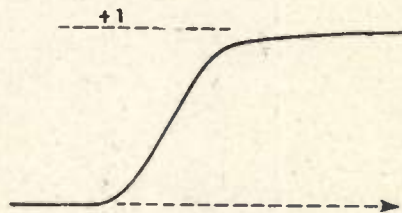


FIG. 10. — Réponse lorsqu'il n'y a pas de sur-oscillation

(pratiquement le τ utilisé précédemment) :

- l'amplitude de la sur-oscillation ;
- la période de la sur-oscillation.

On remarquera, en outre, que le palier horizontal $+1$ n'est atteint que très lentement.

3° Comme l'amplificateur est linéaire, on a le droit de déduire la réponse à une impulsion du type

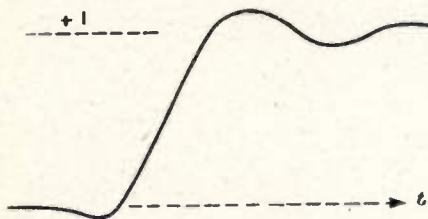


FIG. 11. — Réponse lorsqu'il y a sur-oscillation.

de la figure 8 de l'addition des réponses successives aux deux perturbations unité de la figure 9. Ainsi on devra composer, en les inversant et les décalant de θ_1 , deux réponses du type de la figure 10 ou 11. Imaginons, par exemple, des réponses suivant la figure 10. La composition donne (fig. 12) :

Les flancs sont déformés (temps de montée τ) et il existe une « traînée » qui, empiétant sur l'em-

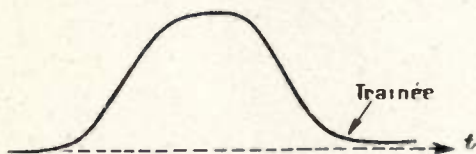


FIG. 12.

placement de la voie suivante est une cause de diaphonie. Il se présente une différence fondamentale par rapport à la diaphonie « basse fréquence » : les constantes de temps sont beaucoup plus faibles, les amortissements plus rapides, et la diaphonie n'est pas générale (nous le vérifierons par le calcul dans les prochains paragraphes).

3.2 Amplificateur à résistances. Equations de base.

Nous partons d'un amplificateur du type de la figure 3. Pour la transmission des fréquences élevées, les tensions de grille de L_2 et de plaque de L_1 sont

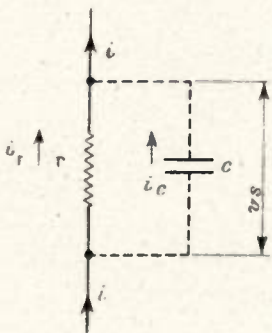


FIG. 13.

parfaitement identiques. Seul intervient l'effet de shunt de r par la somme des capacités parasites.

L'étage amplificateur peut être assimilé à une source à courant i constant (lampe pentode), dans laquelle est insérée la résistance r (d'anode), shuntée par une capacité c (toutes les capacités parasites mises en parallèle, comprenant notamment la capacité de sortie de L_1 et d'entrée de L_2). Le circuit de grille de L_2 , R , n'intervient pas comme charge résistive ($R \gg r$) (fig. 13).

Le courant d'anode i est pris comme référence : il représente à une constante près la tension d'entrée. En développant des calculs classiques, on trouve comme réponse à l'échelon unité :

$$(27) \quad F(t) = 1 - e^{-\frac{t}{cr}}$$

et à l'échelon A :

$$(28) \quad F(t) = A \left(1 - e^{-\frac{t}{cr}} \right)$$

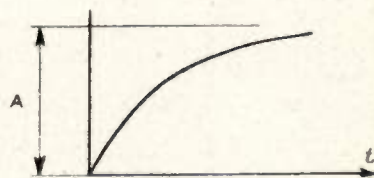


FIG. 14.

Il faut maintenant composer les réponses dues au front avant et au front arrière de l'impulsion.

— Front avant :

$$F(t) = A \left(1 - e^{-\frac{t}{cr}} \right) ;$$

— Front arrière :

$$F(t) = -A \left(1 - e^{-\frac{t-\theta_1}{cr}} \right) ;$$

— Total :

$$(29) \quad \left\{ \begin{aligned} F(t) &= A \left(-e^{-\frac{t}{cr}} + e^{-\frac{t-\theta_1}{cr}} \right) \\ F(t) &= A e^{-\frac{t}{cr}} \left(-1 + e^{\frac{\theta_1}{cr}} \right) \end{aligned} \right.$$

3.3 Diaphonie (modulation d'amplitude)

Reportons-nous à l'équation (28). Pour un temps τ sensiblement inférieur à θ_1 , durée de l'impulsion, on doit avoir pratiquement terminé la montée, c'est-à-dire que $e^{-\frac{\tau}{cr}}$ doit valoir au plus 0,1.

Ceci nous donne :

$$e^{\frac{\tau}{cr}} = 10.$$

En remplaçant τ par θ_1 qui lui est supérieur, on doit donc obtenir pour $e^{\frac{\theta_1}{cr}}$ de la formule (29) une

valeur notablement supérieure à l'unité. On pourra donc négliger -1 devant le terme $e^{-\frac{\theta_1}{cr}}$, et utiliser la formule simplifiée :

$$(30) \quad F(t) = A e^{-\frac{t-\theta_1}{cr}}$$

Pour étudier la diaphonie en modulation d'amplitude, nous supposons encore un brusque accroissement ΔA de A . La tension résiduelle varie de :

$$(31) \quad \Delta v = \Delta A e^{-\frac{t-\theta_1}{cr}}$$

et la perturbation de diaphonie vaut :

$$(32) \quad \frac{\Delta v}{\Delta A} = e^{-\frac{t-\theta_1}{cr}}$$

Par suite des constantes de temps, la perturbation n'est sensible que sur la première voie. Pour analyser exactement la diaphonie sur l'impulsion perturbée dans sa totalité, il convient de rechercher la valeur moyenne de la perturbation pendant la durée θ_1 de l'impulsion.

$$\left(\frac{\Delta v}{\Delta A}\right)_{\text{moy}} = \frac{1}{\theta_1} \int_{\theta_1}^{2\theta_1+\theta_2} e^{-\frac{t-\theta_1}{cr}} dt = \frac{e^{-\frac{\theta_1}{cr}}}{\theta_1} \int_{\theta_1}^{2\theta_1+\theta_2} e^{-\frac{t}{cr}} dt,$$

ce qui donne, tous calculs faits (on négligera au cours des calculs $e^{-\frac{\theta_1}{cr}}$ devant l'unité) :

$$(33) \quad \left(\frac{\Delta v}{\Delta A}\right)_{\text{moy}} = \frac{cr}{\theta_1} e^{-\frac{\theta_2}{cr}}$$

Comme chaque voie n'est perturbée que par la précédente le résultat est le même quel que soit le nombre n de voies en service. La perturbation de diaphonie est également le taux de diaphonie qui s'écrit :

(34) Taux de diaphonie :

$$N \frac{cr}{\theta_1} e^{-\frac{\theta_2}{cr}}$$

3.4. Modulation de durée.

Nous utilisons la relation (30)

$$F(t) = A e^{-\frac{t-\theta_1}{cr}} = A e^{-\frac{t}{cr}} e^{\frac{\theta_1}{cr}}$$

Envisageons cette fois-ci la variation brusque $\Delta \theta_1$ de θ_1 , comme il peut s'en produire par déplacement du front arrière de l'impulsion perturbatrice en modulation de durée. Il en résulte une variation brusque de l'amplitude du résidu que nous chiffrons par la différentielle :

$$\Delta v = \frac{A}{cr} e^{-\frac{t}{cr}} e^{\frac{\theta_1}{cr}} \Delta \theta_1$$

La perturbation est fonction de θ_1 . Elle est en plus dominante sur le flanc avant de l'impulsion perturbée, placée à $t = \theta_1 + \theta_2$. Sur ce flanc la perturbation vaut :

$$(35) \quad \Delta v = \frac{A}{cr} e^{-\frac{\theta_1+\theta_2}{cr}} e^{\frac{\theta_1}{cr}} \Delta \theta_1 = \frac{A}{cr} e^{-\frac{\theta_2}{cr}} \Delta \theta_1$$

Il est logique de regarder la perturbation même sur le flanc avant, l'allongement de durée qui en résulte risquant d'être décelable à la réception.

On notera que la perturbation contient θ_2 ; cela veut dire que le $\frac{\Delta v}{\Delta \theta_1}$ n'est pas le même en tous les points du cycle de modulation de l'impulsion perturbatrice. C'est la première fois que nous rencontrons ce cas. Nous y reviendrons.

Soit τ le temps de montée du flanc avant de l'impulsion perturbée (nous l'appellerons impulsion de durée θ'_1). Pour voir la répercussion du résidu d'amplitude donné par (35) sur la position de ce flanc, le raisonnement sera le même qu'au paragraphe 2.3 (voir fig. 6).

A un accroissement de hauteur de Δv correspondra un déplacement du flanc de $\Delta \theta'_1$ tel que :

$$(36) \quad \frac{\Delta v}{\Delta \theta'_1} = \frac{A}{\tau}, \quad \Delta \theta'_1 = \frac{\tau}{A} \Delta v$$

Combinant (35) et (36) on obtient :

$$\Delta \theta'_1 = \frac{\tau}{A} \frac{A}{cr} e^{-\frac{\theta_2}{cr}} \Delta \theta_1,$$

soit :

$$(37) \quad \frac{\Delta \theta'_1}{\Delta \theta_1} = \frac{\tau}{cr} e^{-\frac{\theta_2}{cr}}$$

là encore, le rapport $\frac{\Delta \theta'_1}{\Delta \theta_1}$ dépend de θ_2 . Comme θ_2 varie au cours du cycle de modulation, la diaphonie n'est plus donnée par un simple affaiblissement de la modulation perturbatrice. Même si l'on suppose au repos la voie perturbée, l'induction de diaphonie sera forte d'un côté du cycle de modulation de la voie perturbatrice, et faible de l'autre. Cela conduira à un résidu dans la voie perturbée qui pour une modulation perturbatrice sinusoïdale aurait l'allure de la figure 15. La diaphonie n'est pas linéaire. On se contentera donc ici de définir une perturbation de diaphonie que l'on s'efforcera de maintenir à une valeur acceptable pour le θ_2 minimum, c'est-à-dire pour l'intervalle dit « de garde ». Toujours pour N étages et quel que soit le nombre des voies :

$$(38) \quad \text{Diaphonie} = N \frac{\tau}{cr} e^{-\frac{\theta_2}{cr}}$$

On notera enfin que lorsque τ est uniquement produit par l'étage (cr) le rapport τ/cr vaut 2,3.

3.5. Modulation de position

Il n'y a pas de différence avec le cas précédent. Seul intervient en effet le déplacement du flanc

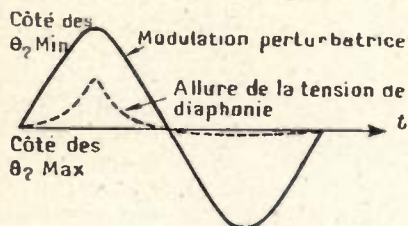


FIG. 15.

arrière de l'impulsion perturbatrice. On a le même déplacement dans la modulation de position, et les effets sont identiques.

3.6. Cas des étages en haute fréquence modulée.

L'enveloppe de modulation subit des déformations du même genre qu'en video directe, et ceci pour les mêmes causes (les capacités parasites). Ainsi un étage à circuit antirésonnant, possédant la même capacité totale c et la même résistance de charge r donne lieu sur l'enveloppe de la haute fréquence à la réponse transitoire :

$$(39) \quad F(t) = 1 - e^{-\frac{t}{2\alpha}}$$

expression identique à (27) à la différence près que la constante de temps est multipliée par 2.

4. Diaphonie des dispositifs convertisseurs.

4.1. Convertisseur « amplitude-durée ».

Le schéma utilisé est représenté par la figure 16. La lampe L_1 sert de lampe de commande. A des impulsions rectangulaires positives appliquées à la grille de L_1 et modulées en amplitude, correspon-

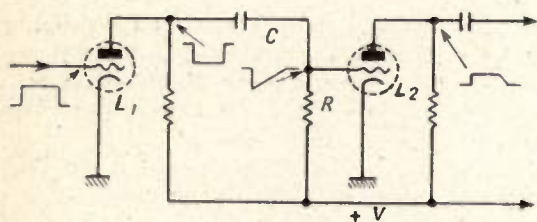


FIG. 16.

dent sur la plaque des impulsions rectangulaires négatives proportionnelles. Ces impulsions sont communiquées par l'intermédiaire d'un dispositif CR à la grille de commande L_2 . La résistance de grille de L_2 est rappelée au $+V$, tension d'anode, ce qui constitue une caractéristique essentielle du système. Mais au repos (dans l'intervalle entre les impulsions), par suite du courant de grille parcourant R , le potentiel de grille de L_2 est zéro et non $+V$.

L'état de repos est défini comme suit :

Courant de L_1 : nul ;

Tension de plaque de L_1 et armature « gauche » de C : au potentiel $+V$;

Tension de grille de L_2 et armature « droite » de C : au potentiel zéro.

Courant de L_2 : maximum (effet apparent de saturation).

Lorsqu'une impulsion positive est appliquée à la grille de L_1 , tout de suite après application du front avant, on trouve l'état suivant :

Courant de L_1 : proportionnel à l'amplitude de l'impulsion ;

Tension de plaque de L_1 et armature « gauche » de C : au potentiel $V-A$;

Tension de grille de L_2 et armature « droite » de C : au potentiel $-A$;

Courant de L_2 : nul, l'amplitude $-A$ est telle que la lampe est fortement bloquée.

Dans la brusque application du front avant de l'impulsion, la différence de potentiel aux bornes de C ne peut changer. Les tensions des sources

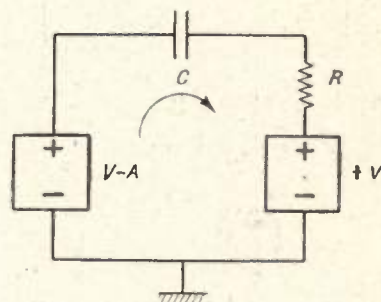


FIG. 17.

restant constantes (on suppose que la durée θ_1 de l'impulsion couvre toute cette partie du cycle), on se trouve en présence du circuit équivalent représenté sur la figure 17. Le phénomène est celui de la charge d'un condensateur dans un circuit comportant résistances et sources.

Le courant de charge du condensateur obéit à la loi générale :

$$(40) \quad i = I_0 e^{-\frac{t}{CR}}$$

et la tension aux bornes de R s'écrit en conséquence :

$$(41) \quad u_R = RI_0 e^{-\frac{t}{CR}}$$

Déterminons I_0 par les conditions initiales. Tout de suite après application du front avant, pour $t = 0$, on trouve comme force électromotrice globale en série dans le circuit :

$$V - A - V - V = -(A + V),$$

donc :

$$(42) \quad I_0 = -\frac{A + V}{R}, \quad u_R = -(A + V)e^{-\frac{t}{CR}}$$

Et en ajoutant à u_R la tension V , on obtient la tension de grille de L_2 :

$$(43) \quad u_g = V - (A + V)e^{-\frac{t}{CR}}$$

Pour $t = 0$, on retrouve $u_g = -A$. Ensuite la tension varie suivant une loi exponentielle, la valeur théorique d'équilibre pour $t \rightarrow \infty$ étant $u_g = +V$. En fait on ne peut dépasser $u_g = 0$ à cause de l'action du courant de grille.

La durée θ pendant laquelle la tension de grille de L_2 est négative est proportionnelle à l'amplitude A : c'est ce qui réalise la transformation.

Le courant d'anode de L_2 passe, avec le front raide initial, de son maximum à zéro. L_2 reste bloquée jusqu'au point où u_g atteint $-V_g$, tension de blocage de la lampe. Le courant varie alors rapidement de zéro à son maximum qu'il atteint lorsque u_g devient nul. Sur la plaque de L_2 , on recueille des impulsions positives, d'amplitude constante et de durée variable θ , proportionnelle à A .

Mais au bout du temps θ_1 arrive le second front (positif) de l'impulsion. Le courant plaque de L_1 passe alors à zéro, et sur la grille de L_2 on devrait théoriquement avoir le front A positif. Mais en réalité le développement de cette tension est fortement freiné par suite de l'action de limitation due au courant de grille. Cette limitation ne peut être chiffrée ; elle dépend des caractéristiques de la lampe L_2 . Il se produit néanmoins une pointe de tension positive sur la grille de L_2 d'amplitude αA (α assez petit). Puis, on tend, suivant une loi de forme exponentielle vers l'état d'équilibre $u_g = 0$. La constante de temps de cette décharge est mal établie, car elle fait intervenir la résistance de fuite

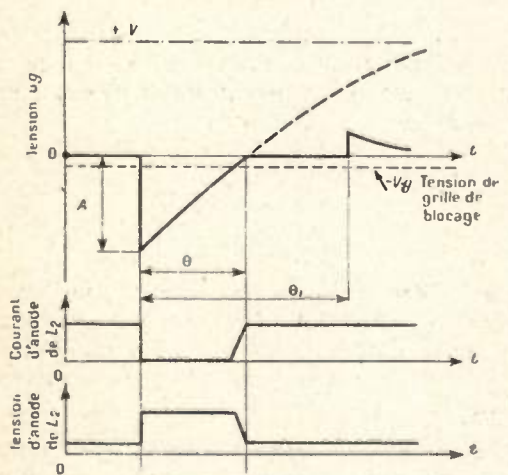


FIG. 18

de grille, peu connue, et variable du reste avec l'amplitude.

Il existe de toute façon un résidu d'amplitude qui subsiste sur l'impulsion suivante. Ce résidu

est proportionnel à A . Il agit sur l'amplitude du front avant suivant, donc finalement sur la durée de l'impulsion transformée. C'est ce résidu qui constitue la diaphonie du système convertisseur.

On a intérêt à utiliser pour C une valeur aussi faible que possible : on diminuera aussi l'importance de la diaphonie. La constante de temps CR devant permettre le fonctionnement suivant la figure 18 on ajustera donc R en conséquence.

Les constantes de temps sont telles que la diaphonie n'existe que d'une voie sur la suivante ; la diaphonie n'est pas générale, mais elle est sensiblement linéaire. Les résultats sont assez critiques et on peut être amené à utiliser deux convertisseurs : un pour les voies « impaires » l'autre pour les voies « paires ».

4.2 Convertisseurs « durée-position ».

Les impulsions modulées en durée ont, par le fait même du pré édant convertisseur, leur flanc arrière (le flanc modulé) qui présente un certain temps de montée τ . Leur flanc avant peut présenter un temps de montée différent. Comme il est de position fixe, il ne peut donner lieu à diaphonie et nous ne nous en occuperons pas.

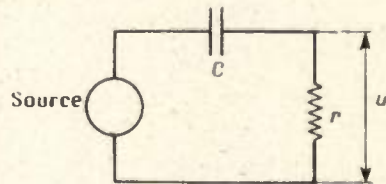


FIG. 19.

Pour nos calculs, le signal « modulé en durée » sera représenté comme issu d'un générateur (fig. 19). L'ensemble cr assure la « dérivation » nécessaire à la transformation cherchée, la tension dérivée étant recueillie aux bornes de r .

Le train d'impulsions de la source est également représenté sur la figure 20. On suppose aux impulsions une polarité négative, ceci simplement pour

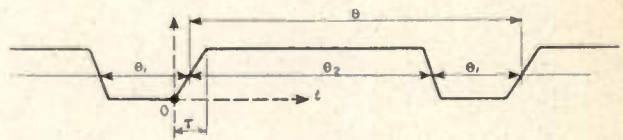


FIG. 20.

la commodité de l'écriture des calculs. Nous prenons pour $t = 0$ le point indiqué sur la figure 20, soit le début de la montée du flanc utile. Nous supposons que sur r la tension de départ est nulle ; en fait le résidu est très faible et, par suite de la linéarité du circuit, il n'y a pas d'inconvénient à le négliger.

Pendant la montée τ , la tension de source a pour expression (le maximum d'amplitude est pris égal à l'unité) :

$$(44) \quad u(t) = \frac{t}{\tau}$$

Le circuit peut être considéré comme une quadripôle très simple. Les bornes d'entrée sont les bornes du générateur, les bornes de sortie les extrémités de la résistance r .

Ce quadripôle possède une certaine réponse à l'échelon unité que les règles de calcul classiques permettent d'écrire :

$$(45) \quad F(t) = e^{-\frac{t}{cr}}$$

(réponse indicielle du réseau).

C'est la courbe du courant de charge du condensateur. La réponse à une tension $u(t)$ (qui n'est plus le signal unité) d'un réseau dont la réponse indicielle est $F(t)$, est donnée par la relation (voir un Ouvrage sur le calcul symbolique) :

$$(46) \quad r(t) = \frac{d}{dt} \int_0^t u(s) F(t-s) ds,$$

soit dans notre cas :

$$r(t) = \frac{d}{dt} \int_0^t \frac{s}{\tau} e^{-\frac{t-s}{cr}} ds = \frac{d}{dt} \frac{e^{-\frac{t}{cr}}}{\tau} \int_0^t s e^{\frac{s}{cr}} ds.$$

L'intégrale définie

$$I = \int_0^t s e^{\frac{s}{cr}} ds$$

s'écrit, tous calculs faits :

$$(47) \quad I = crt e^{\frac{t}{cr}} + c^2 r^2 \left(1 - e^{\frac{t}{cr}}\right).$$

L'expression à dériver pour avoir $r(t)$ s'écrit donc

$$\frac{crt}{\tau} - \frac{c^2 r^2}{\tau} e^{-\frac{t}{cr}} - \frac{c^2 r^2}{\tau},$$

la dérivation donnant :

$$(48) \quad r(t) = \frac{cr}{\tau} \left(1 - e^{-\frac{t}{cr}}\right).$$

C'est la tension du temps de montée, représentée sur la figure 21.

Pour obtenir une véritable dérivation, on doit adopter un faible cr vis-à-vis de τ . La figure est, par exemple, tracée avec $cr = \frac{\tau}{10}$. L'équation (48) montre que l'amplitude de l'impulsion dérivée vaut $\frac{cr}{\tau}$. Elle est d'autant plus petite que la dérivation est parfaite.

Pour la période s'étendant au delà de τ , on se trouve en présence d'un circuit comportant, en série, une source constante, un condensateur, et

une résistance. Le courant a pour expression générale :

$$i = i_0 e^{-\frac{t}{cr}}$$

et la tension aux bornes de r pour valeur :

$$(49) \quad u = u_0 e^{-\frac{t}{cr}}$$

en prenant pour $t = 0$ le début de cette deuxième phase. La tension de départ u_0 est celle d'aboutis-

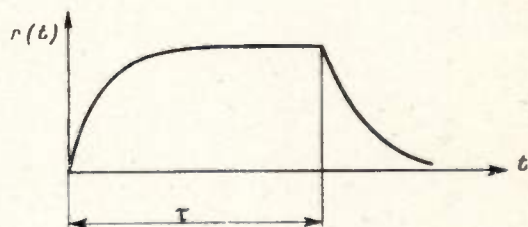


FIG. 21.

sement de la première phase, soit pratiquement si la dérivation est bien assurée :

$$u_0 = \frac{cr}{\tau}$$

Nous aurons ainsi :

$$(50) \quad u = \frac{cr}{\tau} e^{-\frac{t}{cr}}$$

Le résidu de l'impulsion dérivée est donné par (50). Il faut le considérer à la distance :

$$t = \theta = \theta_1 + \theta_2,$$

période de récurrence, c'est-à-dire là où se place l'impulsion suivante susceptible d'être perturbée. D'où le résidu :

$$(51) \quad u = \frac{cr}{\tau} e^{-\frac{\theta}{cr}}$$

Nota. — Par suite de la linéarité des circuits, la dérivation qui se place entre temps, sur le front avant de l'impulsion suivante n'a aucune action sur ce résidu, tout au moins en ce qui concerne la diaphonie.

Le résidu (51) est variable en amplitude par suite de la variation de θ . Si nous supposons un brusque accroissement $\Delta\theta$, résultant de la modulation d'une voie perturbatrice, il en résulte une brusque variation de l'amplitude du résidu, donnée par la différentielle :

$$(52) \quad \Delta u = -\frac{1}{\tau} e^{-\frac{\theta}{cr}} \Delta\theta.$$

Comme l'impulsion perturbée ne présente pas un flanc parfaitement vertical, la variation Δu produit un déplacement de ce flanc le long de l'axe des temps. C'est le phénomène déjà étudié sur la figure 6. Mais ici (voir fig. 21) les flancs de l'impulsion perturbée ne sont pas très droits. Nous considérerons le phénomène sur la figure 22 et pourrons écrire à propos de cette figure : $\frac{\Delta u}{\Delta \theta'} = \frac{du}{dt}$ de la courbe au point considéré.

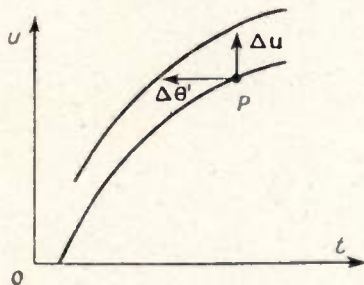


FIG. 22.

Cette dérivée, tirée de (48), s'écrit :

(53)
$$\frac{du}{dt} = \frac{1}{\tau} e^{-\frac{t}{\tau}}$$

Pour prendre une valeur défavorable, on peut considérer le point d'amplitude 0,9, pour lequel :

$$e^{-\frac{t}{\tau}} = 0,1,$$

on aura :

$$\frac{du}{dt} = \frac{0,1}{\tau} = \frac{1}{10 \tau},$$

donc

(54)
$$\frac{\Delta u}{\Delta \theta'} = \frac{1}{10 \tau}, \quad \Delta \theta' = 10 \tau \Delta u.$$

Combinant (52) et (54) :

$$\Delta \theta' = -10 e^{-\frac{\theta}{\tau}} \Delta \theta,$$

(55)
$$\text{Diaphonie } \frac{\Delta \theta'}{\Delta \theta} = -10 e^{-\frac{\theta}{\tau}}.$$

La diaphonie est donnée par (55) θ variant au cours du cycle de modulation, le $\frac{\Delta \theta'}{\Delta \theta}$ varie lui-même dans de grandes proportions. On se trouve dans le cas de diaphonie non linéaire illustrée par la figure 15.

TABEAU RÉCAPITULATIF

Type de Diaphonie	Système de Modulation	Affaiblissement de diaphonie (par étage)	Pour n voies multiplier par	Nature de Diaphonie
Générale (B.F.)	amplitude	$\theta_1 CR$	$\sqrt{n}$	linéaire
	durée	τCR	$\sqrt{n}$	linéaire
	position	$\frac{\tau}{CR} \times \frac{\theta_1}{CR}$	$\sqrt{n}$	linéaire
Trainée (H.F.)	Amplitude	$\frac{cr}{\theta_1} e^{-\frac{\theta_2}{cr}}$	1	linéaire
	durée	$\frac{\tau}{cr} e^{-\frac{\theta_2}{cr}}$	1	non linéaire
	position	$\frac{\tau}{cr} e^{-\frac{\theta_2}{cr}}$	1	non linéaire
Convertisseurs	Amplitude à durée	effet de limitation effet de constante de temps pas de formule	1	linéaire
	durée à position	$-10 e^{-\frac{\theta}{\tau}}$	1	non linéaire

BIBLIOGRAPHIE

MODULATION PAR IMPULSIONS

R. RIGAL, *Cours de Radioélectricité générale*, t. 3, livre 1 : *L'émission* : chap. V : *Modulation par impulsions* (Editions Eyrolles, Paris).
J. FAGOT, *Calcul et construction des oscillateurs, des amplificateurs de puissance et des modulateurs*, fasc. III, chap. XIII : *Emetteurs à modulation par impulsions* (Cours de l'E.S.E., Malakoff, 1952).
P. DAVID, *La réception* (Cours de l'E.S.E., compléments portant sur la modulation par impulsions).
P. Besson, *Modulation par impulsions* (vol. *Electronique, de « Techniques de l'Ingénieur »*, t. 2, 26, place Dauphine, Paris, 1^{er}, 1953).
G. POTIER, *L'utilisation des convertisseurs de modulation dans les équipements multiplex à impulsions* (*Onde électrique*, avril 1950).

CALCUL SYMBOLIQUE. RÉGIMES TRANSITOIRES

A. ANGOT, *Compléments de Mathématiques à l'usage des ingénieurs* (Editions de la Revue d'Optique, Paris).
R. POTIER et J. LAPLUME, *Le calcul symbolique et quelques applications à la physique et à l'électricité* (Hermann et Cie, Paris, 1943).
J. FAGOT, *La transmission des régimes transitoires dans les circuits radio-électriques* (*Rev. gén. Electr.*, novembre 1944).

DIAPHONIE

J. E. FLOOD, *Crosstalk in time-division-multiplex communication systems using pulse-position and pulse-length modulation* (*Proc. Inst. Electr. Eng.*, Part. IV, avril 1952).
S. MOSKOWITZ, L. DIVEN et L. FEIT, *Crosstalk considerations in time-division multiplex systems* (*Proc. Inst. Radio Eng.*, 1950, p. 1330).
H. GARDERE et J. OSWALD, *Etude de la diaphonie dans les systèmes multiplex à impulsions* (*Câbles et Transmission*, t. 2, 1948, p. 173).
NOTA : Une version plus détaillée du présent article : *Causes diverses de diaphonie dans les systèmes multiplex à impulsions* est parue dans : *Annales de Radioélectricité*, octobre 1953, Paris.

DIODE AU GERMANIUM « A ESPACE POSITIF »

Nouvel outil pour la technique des impulsions

PAR

A.H. REEVES

Standard Telecommunication Laboratories (Londres)

1. Historique.

L'origine de la diode au germanium « à espace positif » se trouve dans une série de recherches expérimentales pour déterminer les propriétés des diodes au germanium normales, recherches effectuées par les « Standard Telecommunication Laboratories » en 1950 et 1951. Au cours d'essais de modification des caractéristiques de ces diodes par traitement électrique, ou « formation » électrique, on a pu observer occasionnellement une diode ayant une propriété anormale, c'est-à-dire une résistance négative dans une certaine partie de la région conductrice de sa caractéristique. L'importance pratique possible de cette découverte fut immédiatement comprise ; d'autres études furent entreprises pour trouver le meilleur moyen pour pouvoir reproduire de telles diodes exceptionnelles, avec les caractéristiques les plus désirables.

De ce qui suit on pourra se rendre compte de l'importance probable d'une telle diode. Pour certains circuits téléphoniques à la fois à grande et à faible distance, ainsi que pour la commutation électronique dans le domaine téléphonique, de nombreux ingénieurs considèrent maintenant qu'il est économiquement désirable d'utiliser jusqu'à 1 000 canaux de conversation ou plus sur un seul circuit de transmission (fil ou radio) au moyen d'un des nouveaux systèmes à impulsions à amplitude constante existants actuellement, tels que le PTM ou le PCM.

Pour ce faire, il faut sur le PTM une fréquence de répétition d'impulsions d'au moins 6 mégacycles ; et, sur le PCM à 7 éléments, 42 mégacycles. Pour obtenir les avantages normaux d'absence de bruit et d'indépendance des conversations sur de tels systèmes, avec soit le PTM, soit le PCM, les temps de montée et de descente des impulsions doivent tous les deux ne pas être supérieurs à environ 1/100 de microseconde.

L'équipement de sortie aurait à produire et recevoir ces impulsions très brèves aux cadences élevées de répétition ci-dessus ; et il faudrait aussi fa-

briquer des émetteurs bon marché pour les transmettre. Avec le PCM l'équipement de sortie doit aussi être capable de coder et de décoder les impulsions de cette forme d'onde. Pour la commutation, des dispositifs de distribution ou des réseaux convenables devront aussi être étudiés. Les outils de base pour la presque totalité de ces besoins sont : (a) un producteur d'impulsions presque apériodiques travaillant jusqu'à une fréquence de répétition d'environ 50 Mc/s, donnant des temps de montée et de descente de l'ordre de 1/100 de microseconde ; et (b) un dispositif similaire à double stabilité apériodique, ou compteur binaire, travaillant dans des limites du même ordre. En raison de leur constante de temps relativement grande ces besoins ne peuvent pas être satisfaits par les tubes à vide poussé ou les tubes à atmosphère gazeuse existants ; il faut un outil complètement nouveau. Cependant, une diode au germanium avec une résistance de pente négative dans la direction conductrice, introduit un champ complètement nouveau de possibilités. En effet, il était connu que les capacités propres de ces diodes pouvaient être réduites au moins dans la même mesure que celle des tubes à vide, tandis que dans la direction conductrice les résistances des diodes pouvaient être de beaucoup inférieures à celles des tubes normaux. Naturellement la pente de résistance négative a donné la possibilité d'action de déclenchement.

Les essais sur des échantillons de diodes au germanium à espace positif ont dès le début donné à espérer des constantes de temps de déclenchement très réduites quand on les compare avec celles des tubes à vide et en très peu de temps des diodes furent produites en laboratoire qui atteignaient presque les limites désirées de fréquence de répétition d'impulsions, et qui respectaient pleinement les temps de montée et de descente des impulsions que nous avons indiqués plus haut.

2. Construction et caractéristiques de la diode « à espace positif ».

Une pointe, de préférence en argent contenant une faible quantité d'impureté du type N tel que l'arsenic, est mise en contact avec une surface de

(1) Qui s'appelle « formation électrique ».

germanium de type N préparé par les méthodes normales. La formation électrique se fait avec un courant anormalement élevé de l'ordre de 1 à 5 ampères, généralement en plusieurs décharges de quelques secondes chacune, jusqu'à ce qu'on observe les caractéristiques désirées de tension et de cou-

et le changement de courant correspondant peut être de 40 milliampères ou plus.

Les figures 1a à 1c montrent certaines des caractéristiques typiques, obtenues dans ce cas par alimentation à courant constant fourni par un tube pentode comme indiqué à la figure 2.

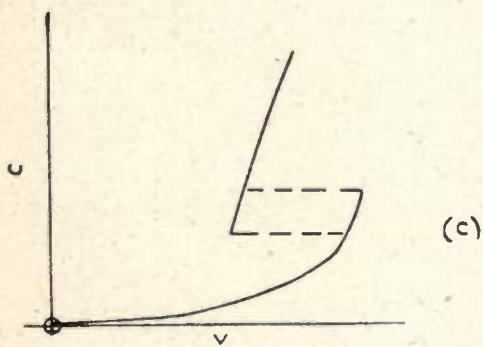
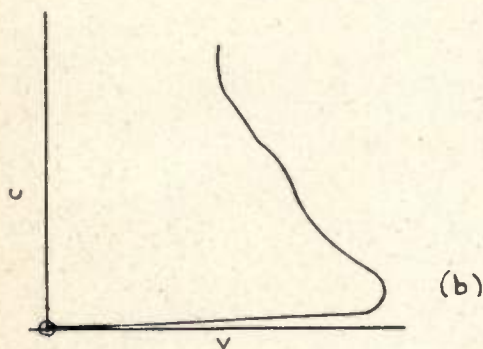
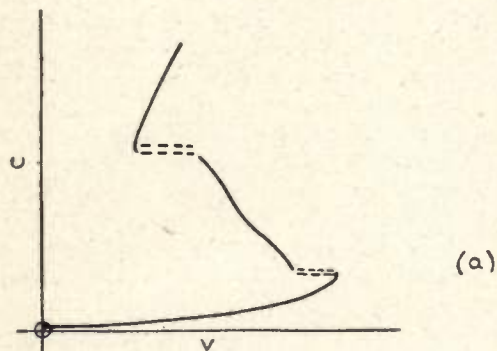


FIG. 1

rant sur un tube cathodique de contrôle par balayage à une fréquence de 50 cycles. Le courant de la formation électrique peut être continu ou alternatif.

La forme désirée de la caractéristique conductrice explique le nom d'« espace positif », car après formation correcte il y a un espace dans l'image sur le tube cathodique, indiquant à une certaine tension une double stabilité dans la direction conductrice due à une résistance négative. La différence de tension au déclenchement d'une position stable à l'autre est généralement de l'ordre de plusieurs volts

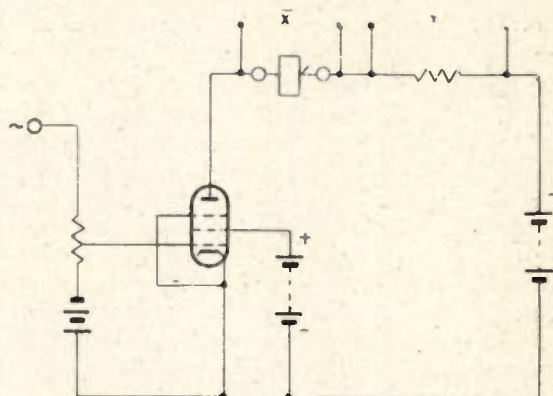


FIG. 2

La figure 3 montre les conditions pour la double stabilité d'une diode à espace positif ayant la caractéristique typique *OABC* telle qu'elle est tracée à courant constant, dans un circuit en série simple avec une batterie *E* et une résistance de charge *R*.

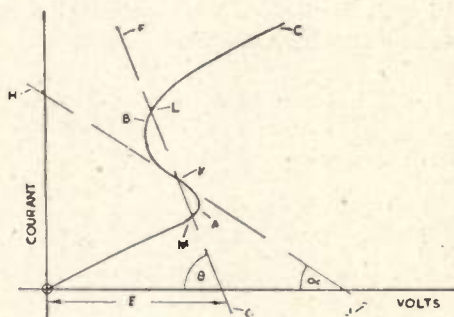


FIG. 3

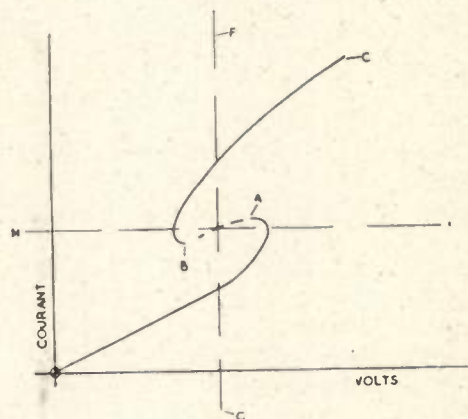


FIG. 4

La double stabilité n'est obtenue que si la ligne de charge coupe la courbe *OABC* trois fois, comme le fait la ligne droite *FG*. (La troisième intersection en

K représente, naturellement, un point d'équilibre instable entre les deux points stable L et M). Le critère pour ce résultat est que l'angle θ soit plus grand que α , où $\cot. \theta$ donne la résistance de charge R , et HJ , tangente à la courbe $OABC$ au point K , ne doit donc pas dépasser une pente limite. Il est également nécessaire qu'il y ait une valeur convenable de E pour que les trois intersections soient obtenues.

Un autre tracé typique est montré à la figure 4. Dans ce cas la ligne du tracé manque entre les points A et B , montrant que même à une résistance de pente de charge presque infinie (alimentation par pentode) une double stabilité est encore présente. L'espace AB représente un décrochage brusque durant le balayage à 50 cycles. Il faut noter que la valeur du courant en B est inférieure à celle en A . Il est évident d'après le position des lignes de charge FG et HJ que ce type de caractéristique peut montrer une double stabilité pour toutes les valeurs de résistance de charge de zéro à l'infini.

Une caractéristique de la forme $OABC$ (fig. 4) peut être expliquée en admettant qu'il y a au moins deux circuits en parallèle entre le fil de contact et le germanium, l'un présentant une caractéristique de tension/courant similaire à celle de la figure 3 et l'autre ayant effectivement une résistance à pente positive pure d'une valeur appropriée.

3. Utilisation de la diode « à espace positif » comme générateur d'impulsions brèves.

Quoiqu'avec des dispositifs spéciaux on ait obtenu des impulsions pour les essais de laboratoire d'une durée de 1/1 000 microsec. ou moins, l'appareillage

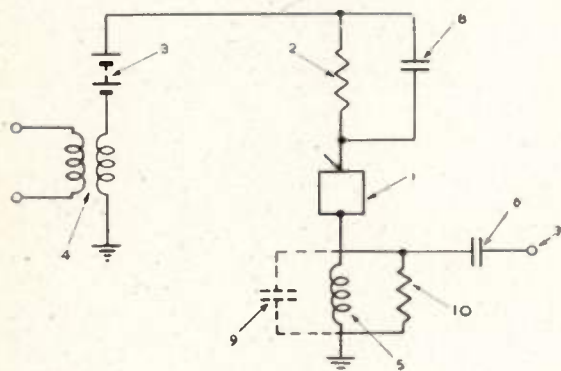


FIG. 5

est généralement assez complexe. Il est donc très clair qu'il serait avantageux d'avoir une méthode très simple et bon marché pour obtenir des impulsions jusqu'à, disons, 20 mégacycles P.R.F. (fréquence de répétition) avec des temps de montée et de descente ne dépassant pas environ 1/100 microsec. La diode « à espace positif » fournit une solution appropriée à ce problème.

Sur la figure 5 on montre un circuit très simple remplissant à peu près ces conditions. La plaque de la diode 1 à « espace positif » est reliée à la masse par la petite résistance 2, la batterie 3, et le secondaire

du transformateur 4. La base de la diode est reliée à la masse par l'impédance 5 (montrée ici comme étant une inductance en parallèle avec une capacité répartie 9) et la résistance 10. Les impulsions de sortie sont transmises à la borne 7 par le condensateur 6. Le condensateur 8 est très petit pour filtrer seulement les composantes à très haute fréquence.

En supposant que la fréquence de répétition demandée soit constante, une onde sinusoïdale de cette fréquence est appliquée au primaire du transformateur 4. La batterie 3 et la résistance 2 sont alors réglées de façon que le balayage de la caractéristique passe deux fois par période par le point de déclenchement. Avec un amortissement critique, on obtient ainsi deux impulsions très brèves à travers l'impédance 5 pour chaque période de l'onde sinusoïdale, une positive et une négative. Avec un amortissement moins critique il se produira un train amorti aux bornes 7 — un réglage qui est quelquefois utile (voir par exemple le paragraphe 6, décrivant l'utilisation de cette diode comme compteur binaire).

Avec le circuit de la figure 5 les temps de montée et de descente de l'impulsion ont pu être estimés d'une façon assez précise par la valeur des tensions d'excitation par choc produites à diverses fréquences naturelles du circuit accordé 5-9. Sur certains échantillons de diodes, le temps de montée de l'impulsion était d'environ 1/150 microsec, et le temps de descente d'environ 1/60 microsec. Cette différence entre les temps de montée et de descente (utilisée d'une façon pratique dans le circuit de compteur binaire de la section 6) peut être expliquée par l'impédance intérieure plus élevée de la diode dans la position de stabilité « hors circuit » par comparaison avec celle « en circuit ». Une fréquence de répétition d'impulsions allant jusqu'à 20 Mc/s fut obtenue avec un certain nombre d'échantillons de diodes. On a trouvé qu'une tension de sortie de plusieurs volts crête, sous 20 mA, était une valeur assez typique — correspondant à une impédance intérieure du générateur d'impulsions de l'ordre de 100 ohms.

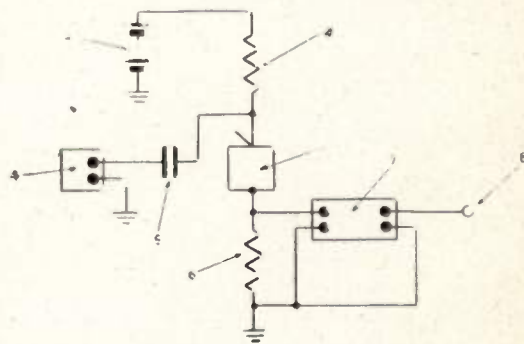


FIG. 6

Des générateurs d'impulsions basés sur les diodes au germanium à espace positif selon le circuit de la figure 5 sont maintenant utilisés par les Standard Telecommunication Laboratories ; on a trouvé que

ce sont des outils simples et très utiles pour les travaux d'expérience.

4. Utilisation de la diode à espace positif comme dispositif à double stabilité apériodique.

La figure 6 montre le premier circuit utilisé dans ce but. La pointe de la diode à espace positif 1 est reliée à la masse par la petite impédance 2 et la batterie 3, tandis que sa base est reliée à la masse à travers la petite résistance 6 (80 ohms). 4 est un générateur d'impulsions alternées positives et négatives, à des fréquences de répétition allant jusqu'à 20 Mc/s, obtenues par la déformation d'une onde sinusoïdale. La sortie de 4 était reliée à la pointe de la diode 1 par le petit condensateur 5. Le signal de sortie de la diode était observé sur un tube cathodique connecté en 8, après en avoir élevé suffisamment le niveau par un amplificateur à large bande 7.

Les résultats ont montré sur certains échantillons de diodes une succession de déclenchements effectivement apériodiques par les impulsions d'entrée à des fréquences allant jusqu'à 20 Mc/s, valeur limite

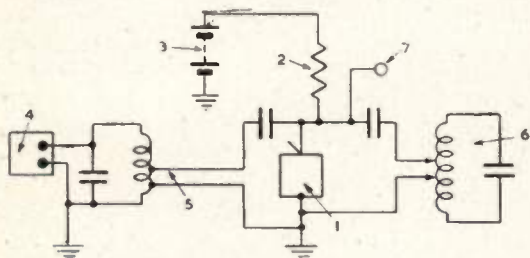


FIG. 7

du matériel d'essais. Sur d'autres échantillons de diodes le déclenchement a cessé à des fréquences de répétition légèrement inférieures.

Pour essayer les diodes à des fréquences supérieures à 20 Mc/s, on a utilisé l'arrangement de la figure 7. La base de la diode était reliée directement à masse, tandis que le fil de contact était relié à la masse

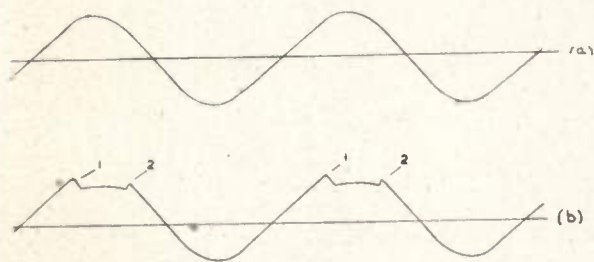


FIG. 8

par la résistance 2 et la batterie 3 comme ci-dessus. L'oscillateur à ondes sinusoïdales 4, réglable jusqu'à environ 50 Mc/s était arrangé pour appliquer une pointe de plusieurs volts à travers la diode 1, par la ligne de transmission 5 et le système d'équilibrage d'impédance 6. On observait la sortie de la diode à espace positif sur un tube cathodique spécial pour fréquences élevées connecté à la borne 7. Il n'était pas utilisé d'amplificateur à bande large ; la diode avait

été choisie comme ayant une différence de tension assez grande entre les deux points de double stabilité pour donner directement un changement dans le tracé du tube cathodique sans utiliser d'amplificateur.

Lorsque la tension de l'onde sinusoïdale appliquée à la diode était insuffisante pour causer le déclenchement, le tracé du tube cathodique en 7 était presque sinusoïdal, comme le montre la figure 8 (a). Lorsque le déclenchement de la diode se produisait on observait les dents de scie 1, 2 de la figure 8 (b) sur le tracé du tube cathodique — car le circuit syntonisé associé avec la diode était trop amorti pour que les harmoniques représentés par ces dents de scie soient supprimés.

Par ce moyen deux échantillons de diodes présentèrent effectivement des successions de déclenchements apériodiques à des fréquences de répétition allant jusqu'à 43 Mc/s.

5. Utilisation comme circuit à déclenchement à rétablissement automatique.

Comme le montre la figure 9, la diode à espace positif 1, avec la base à la masse, est reliée à la masse à l'extrémité pointe par une résistance 2 et une batterie 3. Des impulsions d'entrée redressées sont appliquées à travers la diode de façon appropriée, par exemple au moyen du transformateur d'impul-

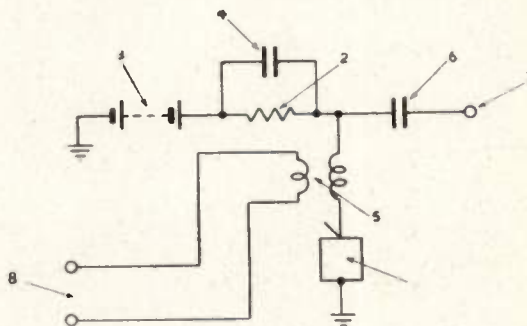


FIG. 9

sions 5 à partir des bornes 8 d'un générateur d'impulsions. Si on ajoute le condensateur 4 en parallèle avec des valeurs appropriées de 2, 4 provoqueront une oscillation de relaxation — pourvu que la batterie 3 et la résistance 2 soient telles qu'elles amènent la diode dans une région de résistance négative appropriée sur sa caractéristique. Nous changeons maintenant légèrement la tension à travers la diode pour amener cette dernière légèrement en dehors de la portion de résistance négative de la courbe ; l'oscillation cessera. Cependant, si les impulsions d'entrée sont arrangées pour balayer la caractéristique de la diode une fois de plus dans la région de résistance négative, des demi-périodes de l'oscillation de relaxation peuvent être produites à chacune de telles impulsions d'entrée, les demi-périodes d'oscillation alternées ayant lieu indépendamment des impulsions d'entrée et rétablissant les conditions initiales du circuit après un intervalle de temps réglable dépendant principalement de la valeur de la constante de temps 2-4.

Il y a ainsi une demi-période d'oscillation de relaxation libre réglée par chaque impulsion d'entrée ; cela est suivi par le début d'une demi-période de rétablissement automatique, principalement réglée par la constante de temps 2-4 ; après quoi il y a un intervalle de repos jusqu'à l'arrivée de la prochaine impulsion d'entrée.

De cette manière on peut construire un circuit de déclenchement à rétablissement automatique, stable et approprié.

6. Utilisation comme compteur binaire apériodique.

Il y a naturellement plusieurs moyens possibles d'utiliser la diode à espace positif comme compteur binaire ; peut être l'arrangement le plus simple est-il celui de la figure 10, où il ne faut qu'une seule diode pour obtenir la caractéristique binaire.

La première diode à espace positif montrée dans la figure est utilisée pour convertir l'onde de signal originale à compter, c'est-à-dire une onde ayant une forme arrondie, en une série d'impulsions brèves qui est nécessaire pour l'entrée binaire ; c'est dans son ensemble le circuit déjà montré à la figure 5. La diode à espace positif 1 a sa base reliée à la masse par une résistance 2 (100 à environ 400 ohms, dans les cas typiques). La pointe de 1 est reliée à la masse par le transformateur d'entrée 3, la résistance 4

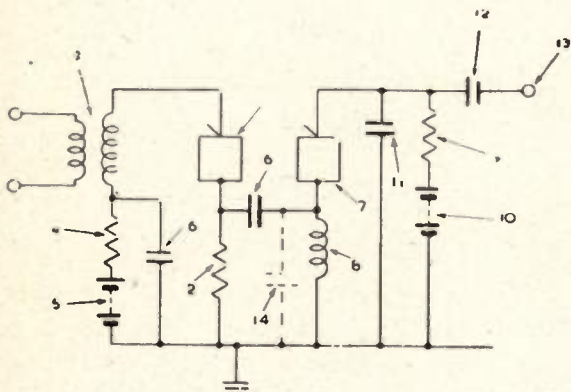


FIG. 10

et la batterie 5. Le petit condensateur 6 est utilisé comme filtre pour les impulsions brèves produites au déclenchement. Comme dans le cas de la figure 5, les formes d'ondes d'entrée arrondies (on le suppose) venant du transformateur 3, qui sont le signal à compter, sont converties en impulsions très brèves au condensateur de sortie 6. Alors ces impulsions brèves de 6 excitent par choc le circuit accordé formé par l'inductance 8 et sa capacité répartie 14, la fréquence de résonance étant dans les exemples typiques dans la région de 100 Mc/s. Le circuit de la diode 7 est complété par la résistance 9 et la batterie 10 en parallèle avec le très petit condensateur 11. La sortie du compteur binaire est obtenue aux bornes 13 à travers le condensateur 12.

On va maintenant expliquer le fonctionnement du circuit. Dans la figure 11 la courbe A est la forme

d'ondes d'entrée présumée venant du transformateur 3 (Fig. 10). La courbe B représente la forme d'onde du courant à travers la diode 1. La courbe C est la dérivée de B. On notera que les impulsions négatives 55 sont plus petites que les impulsions positives 54 ; cela est dû au fait que dans la diode à espace positif la constante de temps intérieure lors du déclenchement de « courant faible » à « courant fort » est normalement considérablement moindre que la constante de temps dans l'ordre inverse — phénomène montré dans la courbe B par une pente plus forte dans la portion 52 que dans la portion 53. Grâce au petit condensateur 6, la courbe dérivée montrée en C (Fig. 11) excite par choc le circuit accordé 8-14, donnant une forme d'onde représentée par la courbe D de la figure 11. Il y a deux trains amortis : (a) le train 66, 67, 68, dû à l'excitation par choc de l'impulsion positive 54 (la plus grande) ; et (b) le train 72 d'une amplitude relativement plus faible, comme l'impulsion d'excitation 55 est d'une valeur de pointe bien inférieure à celle de l'impulsion 54.

Les constantes de la diode 7 sont réglées de façon que le déclenchement se produise à la première pointe négative 67 du train 66, 67, 68. (Dans les conditions initiales posées le déclenchement ne peut pas se produire sur la pointe positive plus importante 68 parce que cette pointe n'est pas du signe conven-

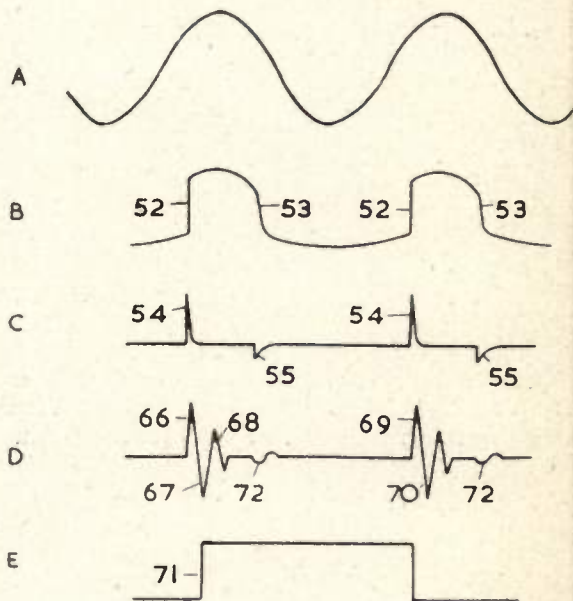


FIG. 11

ble). Le re-déclenchement à la position « hors circuit » lors de la seconde pointe positive 68 est éliminé par la constante de temps intérieure de la diode ; si la fréquence de résonance de 8-14 est de 100 Mc/s. par exemple, l'intervalle entre la pointe 67 et la pointe 68 est seulement de 1/200 microsec. beaucoup trop court pour que la diode puisse se déclencher à nouveau dans les conditions de « courant faible ». La pointe négative au début du train 72 est trop faible en amplitude pour causer le déclenchement de la diode 7 ; donc cette diode restera dans la condition de « courant fort » jusqu'à l'arrivée de la pointe

positive 69 au début du prochain train amorti 69-70. A la pointe 69, la diode 7 sera enclenchée à nouveau à sa condition initiale de « courant faible ».

Négligeant les ondulations dues aux impulsions, le courant dans la diode 7 suivra donc la courbe E de la figure 11. Ainsi il est clair qu'effectivement on puisse obtenir de cette façon une action de compteur binaire ap'riodique.

Avec certains échantillons de diodes on a pu effectuer des comptages allant jusqu'à des P.R.F. d'entrée de 15 Mc/s. Un tel outil, outre ses utilisations en télécommunication, devrait être très utile dans certains types de calculateurs électroniques à grande vitesse. etc.

7. Autres utilisations de la diode au germanium à « espace positif ».

Cette nouvelle forme de diode peut aussi être utilisée comme oscillateur H.F. dans un circuit L-C ;

une limite d'environ 300 Mc/s. a été obtenue avec certains échantillons, à une puissance de sortie de plusieurs milliwatts. La diode a aussi prouvé être un outil très utile dans un récepteur à super-régénération très compact dans cette gamme de fréquence. On remarque d'ailleurs qu'il existe toute une série de ces diodes, douées de différentes caractéristiques, et qu'on peut choisir et en fabriquer selon l'emploi envisagé. Cependant, comme le but de cet article n'est que de discuter la technique des impulsions, les autres utilisations ne seront pas décrites ici en détail.

8. Remerciements.

Je remercie la Société « Standard Telecommunication Laboratories Ltd », de Enfield, Middlesex, de m'avoir permis de publier cet article, et Monsieur R.B.W. COOKE de cette Société pour sa coopération dans le travail expérimental.

UNE TECHNIQUE DE CIRCUITS ÉLECTRONIQUES POUR MACHINE A CALCULER RAPIDE

PAR

M. G. PIEL

*Ingénieur à la Société d'Electronique
et d'Automatisme*

INTRODUCTION.

Cette note a pour objet d'exposer dans ses principaux détails une technologie que nous utilisons pour réaliser une machine à calculer électronique, fonctionnant en numération binaire série.

Cette technologie combine, moyennant un petit nombre de règles fixes, l'emploi de triodes et diodes. Elle répond au problème de rapidité qui nous est imposé et qui nous amène à définir pour les impulsions de rythme, une fréquence de 0,8 Mc. En effet, connaissant la charge de calculs à faire effectuer à la machine en un temps donné, il nous a fallu choisir un schéma fonction de la fréquence des impulsions de rythme. Un schéma plus compliqué ne nous aurait pas permis de diminuer beaucoup cette fréquence et n'aurait donc pas abouti à la solution la plus économique. La durée d'un nombre de seize chiffres est donc fixée à $16 \times 1,25 \mu s = 20 \mu s$, l'intervalle de temps par chiffre étant de $1,25 \mu s$. L'usage pratiquement général de liaisons continues nous a permis de fonctionner en « impulsions jointives », c'est-à-dire que deux impulsions consécutives de même sens n'en forment qu'une de durée $2,5 \mu s$.

Ceci correspond à limiter au plus juste la fréquence la plus élevée à transmettre. Nous verrons dans cet exposé comment, à partir de circuits de base simples, moyennant quelques extensions, nous obtenons des schémas notablement allégés.

1. ASPECT GÉNÉRAL DE LA TECHNOLOGIE CHOISIE.

L'usage de triodes est à la base de cette technologie. Toutefois, l'emploi exclusif de triodes aboutirait à des schémas très chargés en tubes et nous avons dû choisir une solution mixte avec triodes et diodes (diodes au germanium). De même l'emploi exclusif de liaisons continues dans des montages à triodes et diodes, s'il permet de réaliser n'importe quel montage logique, a pu être abandonné dans certains cas où l'usage de transformateurs et de liaisons capacitives (ces dernières, très rares) est rendu possible et se

montre particulièrement avantageux. Nous allons décrire dans les chapitres suivant, les montages typiques utilisables.

2. CIRCUITS LOGIQUES A TRIODES.

2.1. — Triode inverseuse.

La figure 1a donne une vue de principe d'un amplificateur à courant continu élémentaire, sensiblement tel que nous l'utilisons et la figure 1b le symbole que nous emploierons par la suite dans les schémas logiques.

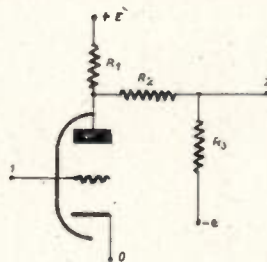


FIG. 1 a



FIG. 1 b

Nous supposons, tout d'abord, que la triode étant bloquée, le potentiel de sortie, sur la borne 2, déterminé par les résistances R_1, R_2, R_3 est nul par rapport au potentiel de la cathode pris comme référence.

(Ceci n'est pas rigoureux, en réalité, et nous admettons un écrêtage par courant de grille, écrêtage que nous négligerons en premier lieu pour simplifier notre exposé, Cf chapitre 5, Caractéristiques technologiques).

Enfin, lorsque la borne 1 est au potentiel zéro, nous supposons que le gain du tube est tel que la tension de sortie soit capable de bloquer la triode d'un montage identique placé à la suite.

Au potentiel zéro nous ferons correspondre la valeur 1 d'un signal, au sens de l'algèbre de Boole, et au potentiel $-v$, la valeur 0. On aura $-v \leq -v_c$, v_c étant le recul de grille maximum admis avec une marge de garantie fixe.

A un signal X appliqué à la borne 1 du montage de la figure 1 correspondra un signal $\bar{X}$ sur la borne de sortie 2 ($\bar{X} + X = 1$, $X\bar{X} = 0$)

L'utilisation de n triodes montées en parallèle sur la même charge constituée par les résistances R_1 , R_2 , R_3 , les anodes étant connectées au même point (entre

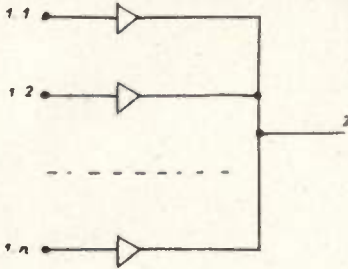


FIG. 2.

R_1 et R_2) est donnée en exemple par la figure 2. Si X_1 , X_2 , ..., X_n sont les signaux appliqués aux bornes d'entrée 1.1., 1.2., ..., 1.n le signal de sortie Y , sur la borne 2, est

$$Y = \bar{X}_1 \cdot \bar{X}_2 \dots \bar{X}_n$$

En effet, ce signal est 1 (tension nulle) seulement lorsque toutes les triodes sont bloquées, les signaux d'entrée étant tous 0.

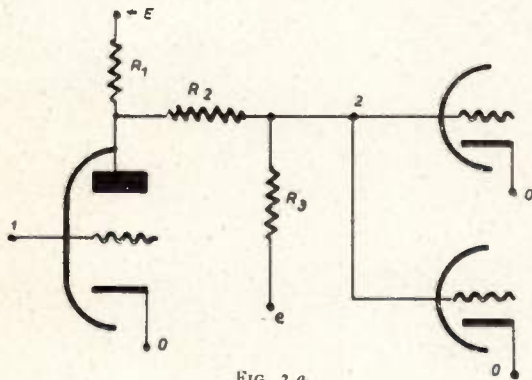


FIG. 3 a

Les formules fondamentales

$$\bar{X}_1 \cdot \bar{X}_2 \dots \bar{X}_p = \bar{X}_1 + X_2 + \dots + X_p \quad (1)$$

$$\text{et } \bar{X}_1 \cdot X_2 \dots X_p = \bar{X}_1 + \bar{X}_2 + \dots + \bar{X}_p \quad (2)$$

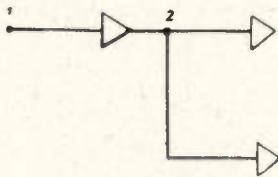


FIG. 3 b

nous permettent de concevoir n'importe quel schéma logique constitué avec des triodes montées comme nous l'avons indiqué.

La figure 3a et la figure 3b correspondante donnent le cas d'attaque de deux grilles à partir d'une même sortie (2), ceci pour montrer qu'un seul pont de résistances R_1 , R_2 , R_3 est utilisé en ce point. Nous verrons plus loin un tableau figurant différentes sortes de schémas combinant deux signaux (chapitre 4).

2.2. — Triode cathodyne.

La figure 4a, représente une triode cathodyne supposée telle que le signal appliqué sur la barre 1 soit pratiquement transmis sans changement à la borne 2. La figure 4b donne le symbole que nous adopterons à la place de la figure 4 a.

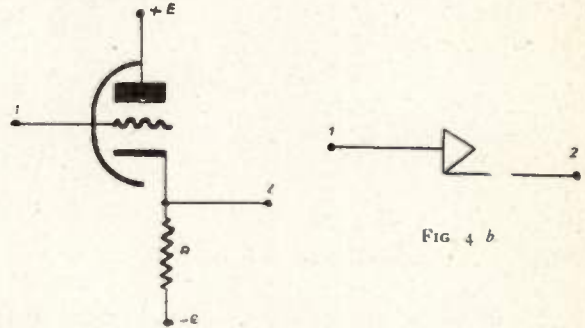


FIG. 4 b

FIG. 4 a.

La cathodyne sera retenue par nous en raison de ses propriétés de séparatrice et d'abaisseur d'impédance c'est-à-dire qu'elle nous servira aux liaisons entre chassis à l'attaque en parallèle de plusieurs grilles (plus de 2) et également dans certains montages utilisant des diodes, montages décrits plus loin.

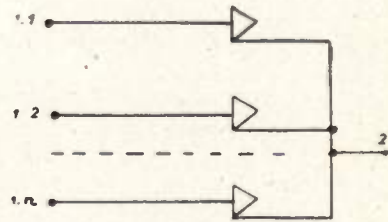


FIG. 5.

Un exemple de mise en parallèle de plusieurs cathodynes, les cathodes ayant en commun la charge R est donné par la figure symbolique 5.

L'opération logique sur les signaux d'entrée X_1 , X_2 , ..., X_n appliqués aux bornes 1.1, 1.2, ..., 1.n représentée par le signal de sortie Y est

$$Y = X_1 + X_2 + \dots + X_n,$$

le signal de sortie Y n'étant 0 que si tous les signaux d'entrée sont 0.

L'emploi de ce dernier type de montage est très rare dans notre technologie ; nous en verrons la raison plus loin.

2.3. — Triode commandée par la cathode.

Dans un montage tel que celui de la figure 1 la triode peut recevoir un signal sur sa cathode de sorte que lorsque la tension de cathode sera nulle, l'équation $Y = \bar{X}$ sera vérifiée, tandis que lorsqu'elle sera positive et supérieure au recul de grille v_c , Y sera 1 quelque soit X . Si Z désigne le signal de cathode, avec comme conventions $Z = 0$, si la tension de cathode est nulle,

$Z = 1$ si la tension de cathode est positive, la fonction de sortie Y correspondra à l'équation $\bar{Y} = X\bar{Z}$ ou $Y = \bar{X} + Z$, comme on peut le vérifier à l'aide du tableau comprenant toutes les combinaisons possibles entre X et Z :

X	Z	Y
0	0	1
1	0	0
0	1	1
1	1	1

La figure 6a montre la disposition adoptée pour permettre le contrôle par la cathode d'une triode avec

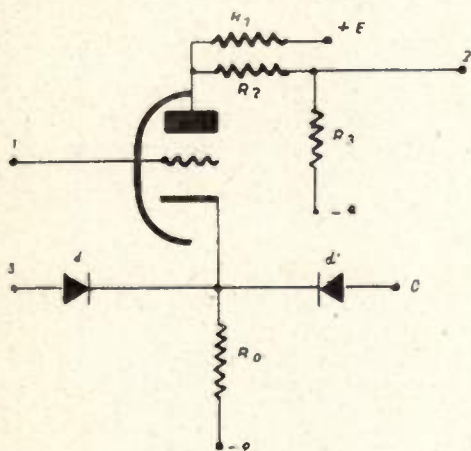


FIG. 6 a.



FIG. 6 b.

un signal dont les deux valeurs 0 et 1 correspondent respectivement aux deux tensions V_0 et V_1 telles que

$$V_0 < 0, \quad V_1 > V_c,$$

Ce signal est appliqué sur la borne 3 et contrôle le potentiel de cathode par l'intermédiaire d'une diode à cristal d . Une diode à cristal d' et une résistance R_0 montée comme l'indique la figure, assurent la butée à zéro du potentiel de cathode, quelque soit l'état de la triode, lorsque le potentiel appliqué sur la borne 3 est négatif.

La figure 5 b est une représentation symbolique de la figure 5 a.

Les tensions applicables sur grilles, pouvant prendre deux valeurs 0, $-v$, nous désignerons les signaux de grille par l'expression : *signaux G*.

Les tensions permettant le contrôle de cathode pouvant prendre deux valeurs $V_0 \leq 0$, $V_1 \geq V_c$, nous

désignerons les signaux correspondants par l'expression : *signaux K*.

Les signaux K dans notre machine sont :

- les impulsions de rythme.
- les impulsions de chiffre,
- certaines impulsions cycliques ou acycliques, de durée bien définie.

Ces signaux K sont toujours sortis sur transformateurs. Les impulsions de rythme sont celles à partir desquelles toute cadence de fonctionnement de la machine découle, leur fréquence est 0,8 Mc comme nous l'avons dit.

Les impulsions de chiffre ont une durée de $1,25\mu s$ et sont au nombre de 16, chacune de ces impulsions est décalée par rapport à la précédente de $1,25\mu s$.

3. — CIRCUITS A DIODES.

3.1. — Circuits de cathode.

Nous venons de voir apparaître au paragraphe précédent l'emploi de diodes au germanium dans le circuit de cathode d'une triode.

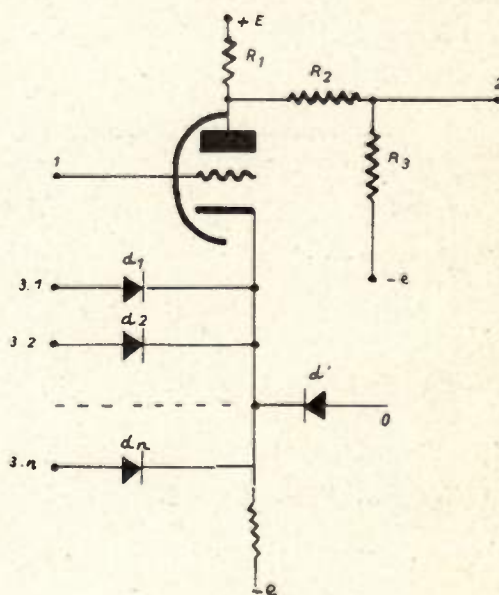


FIG. 7.

La figure 7 correspond au cas d'une application multiple d'impulsions en commande de cathode.

Le signal résultant des signaux $Z_1, Z_2, \dots, Z_n$ appliqués sur les bornes 3.1, 3.2, ..., 3.n est sur la cathode :

$$Z = Z_1 + Z_2 + \dots + Z_n$$

Si le signal de grille est X , le signal de sortie (borne 2) sera

$$Y = \bar{X} + Z_1 + Z_2 + \dots + Z_n$$

Ce genre de montage peut être avantageux dans

certain cas ; il correspond à l'emploi dans le circuit de cathode d'un « mélangeur » classique à diodes.

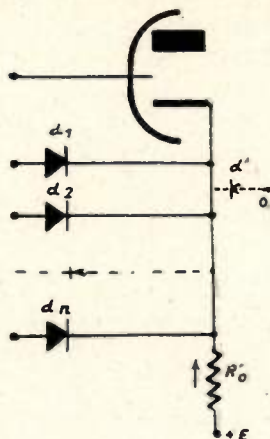


FIG. 8.

Nous remarquons cependant que le montage correspondant à un « conditionneur » obtenu en inversant le sens des diodes $d_1, d_2 \dots d_n$ (fig. 8) ne pourrait utiliser la diode de butée d'(inversée) sans difficulté technologique (mise en court-circuit des sources de signaux K pour la partie négative de ces signaux) et que, cette diode supprimée, la butée devrait se faire malgré la transmission du courant de grille à travers les sources de signaux K .

Nous admettons qu'un tel genre de butée sera exclu (en raison des difficultés que nous ne pouvons

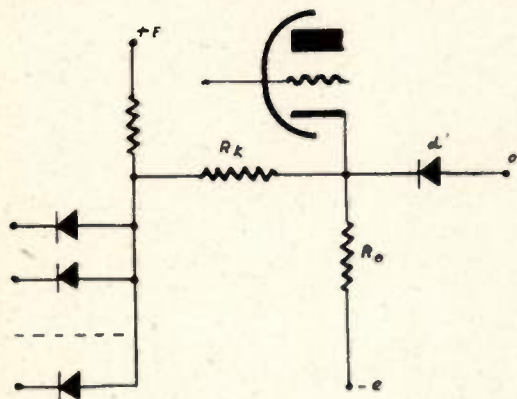


FIG. 9.

développer complètement ici) et que les signaux K pourront toujours devenir nettement négatifs.

D'ailleurs, on doit pour ces signaux, tolérer une certaine variation d'amplitude suivant la charge représentée par les circuits qu'ils attaquent. D'autre part, pour simplifier le schéma de la machine et la technologie, nous avons décidé de ne pas utiliser plus de deux standards de signaux en admettant en outre que les signaux K peuvent servir également comme signaux G :

$$V_0 \leq -V_c, \quad V_1 \geq +V_c$$

(Cf § 3,2)

Ceci dit, le montage correspondant à un « conditionneur » dans le circuit de cathode pourrait alors être celui de la figure 9.

Le rendement énergétique à prévoir pour ce montage est sans doute inférieur à celui du montage de la figure 7.

Le fait que les signaux K peuvent être sortis sur transformateur en direct, ou en inverse, correspond à des possibilités telles que nous pourrions éviter ce dernier type de montage. (Cf. chapitre 4).

Ceci d'ailleurs nous permet de simplifier l'aspect de la technologie.

3.2. — Circuits de grille.

Dire que les signaux K peuvent être utilisés comme signaux G , signifie qu'ils peuvent être appliqués à des grilles.

La figure 10 montre le schéma correspondant à l'attaque d'une triode par un signal K moyennant l'usage d'une résistance de grille en série, R_g .

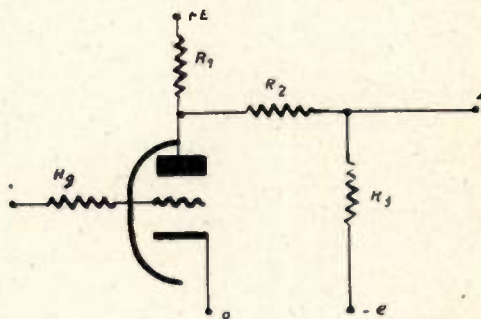


FIG. 10.

La figure 11 montre la possibilité d'utiliser un « conditionneur » dans le circuit de grille d'une triode, les signaux appliqués aux bornes 1.1, 1.2, ... 1.n étant des signaux K ,

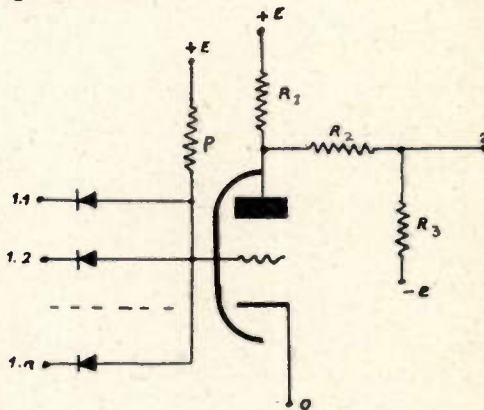


FIG. 11.

On remarquera que dans ce cas, la butée du signal de grille est réalisée par l'apparition du courant de grille mais ne s'effectue pas à travers les sources des signaux K .

Le montage utilisant les diodes en sens inverse, montage qui réaliserait un mélangeur (fig. 12) nécessiterait l'emploi de la résistance R_g .

Le montage de la figure 12 pourra être utilisé.

Le conditionneur de la figure 11 pourra être utilisé entre deux triodes comme le montrent la figure 13 a et la figure symbolique correspondante 13 b.

Il en serait de même pour le mélangeur de la figure 12.

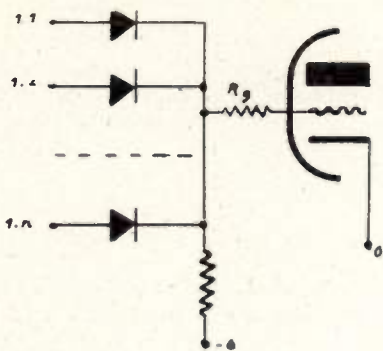


FIG. 12.

Ceci nous donne un grand nombre de possibilités logiques. (Cf. chapitre 4).

Toutefois, dans le cas de la figure 13, nous nous limiterons à l'emploi de 2 ou 3 diodes compte tenu

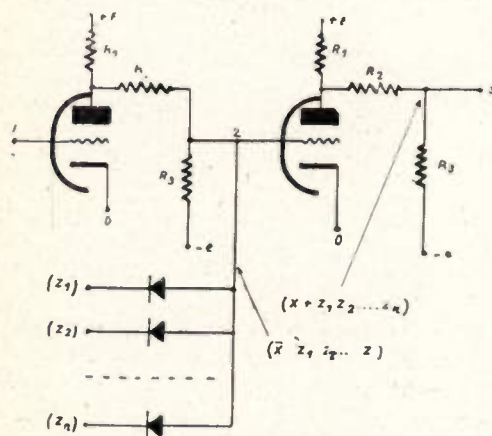


FIG. 13.

de l'impédance du pont de résistance vis à vis des résistances inverses des diodes.

Un résultat intéressant nous est fourni par la substitution possible aux signaux K utilisés ici dans ces circuits de grille, avec des diodes, par des signaux sortant de cathodynes.

Il s'impose alors de distinguer ces derniers signaux, du type G , puisqu'ils offrent une possibilité supplémentaire sur les signaux normaux de grille.

Nous les appellerons : *signaux GK*.

Le schéma de la figure 13 pourra alors être retenu dans le cas où les signaux appliqués aux diodes seront du type GK . (limitation à 2 ou 3 diodes). Il serait plus délicat de prévoir l'usage de mélangeurs tels que celui de la figure 12 (placé en attaque directe comme sur la figure 12 a entre 2 triodes) avec des signaux GK eu égard à l'impédance de sortie des cathodynes à et la résistance inverse des cristaux pour des raisons faciles à analyser, aussi avons-nous dû éliminer ces cas d'emploi.

Dans le cas de la figure 13 en supposant $X = 0$ il faudra qu'une quelconque des cathodynes, par exemple celle délivrant le signal Z_1 , soit susceptible d'abaisser suffisamment le potentiel du point 2 lorsque l'on

aura $Z_1 = 0$, pour que le signal en ce point puisse alors être considéré comme étant 0.

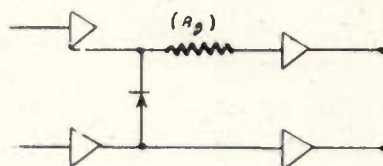


FIG. 14.

Nous avons choisi nos éléments de circuits (résistances R, R_1, R_2, R_3) pour que cette possibilité nous soit offerte, en notant toutefois qu'à toute cathodyne ne pourra correspondre qu'une seule diode.

Ceci nous ne empêche pas d'ailleurs d'avoir des schémas du genre de celui de la figure 14 dans lequel une cathodyne attaque à la fois une grille (avec résistance série Rg) et une diode.

4. — SCHEMAS LOGIQUES ÉQUIVALENTS.

Le tableau I montre dans chaque colonne les circuits équivalents correspondant à un même signal de sortie.

Dans ce tableau, nous avons représenté sur la première ligne les circuits avec des triodes inverseuses ne recevant que des signaux G .

Dans le reste du tableau, nous avons volontairement éliminé pour chaque colonne, les schémas qui nécessitent un nombre de triodes supérieur à celui du schéma de la première ligne et nous avons supposé que les signaux K pouvaient être pris directs ou inverses. Pour simplifier, nous avons également éliminé tout montage comportant un « mélangeur » dans le circuit de grille.

Les schémas marqués d'une astérisque sont ceux qui ne se terminent pas par une triode inverseuse. Ces schémas ne devront être suivis que par une telle triode inverseuse, s'ils sont intéressés par des signaux GK en provenance d'une cathodyne. Dans le cas où ils sont intéressés par des signaux K , ils pourront être suivis d'une inverseuse ou d'une cathodyne.

Nous avons en effet admis que deux cathodynes ne devront jamais être à la suite l'une de l'autre par mesure de sécurité et qu'il en sera de même si une cathodyne est associée à cette diode : derrière la diode on ne devra retrouver qu'une triode inverseuse. (cas du circuit $X_1 X_2$ de la deuxième ligne, par exemple.).

C'est d'ailleurs pour ce fait que les montages de la dixième ligne utilisant des cathodynes pour réaliser des fonctions logiques prennent un faible intérêt. Il est facile de voir, grâce à ce tableau, que les ressources de la technologie choisie ont été nettement étendues par l'emploi de diodes et par la définition de deux standards de signaux.

Le cas de la fonction $\bar{X}_1 + \bar{X}_2$ réalisée avec des signaux G est typique : le montage à cinq triodes de la première ligne est avantageusement remplacé par un montage utilisant deux cathodynes (transformateurs de signaux G en signaux GK) et le montage de

TYPE DE MONTAGE	TYPE DE SIGNAL	SIGNALS DE SORTIE					
		X_1 X_2	$X_1 + X_2$	$\bar{X}_1 + X_2$	$\bar{X}_1 + \bar{X}_2$	$X_1 X_2$	$\bar{X}_1 X_2$
AVEC INVERSEUSES ET DIODES	1 G, K, GK G, K, GK						
	2 G, K, GK K, GK						
	3 K, GK K, GK						
	4 K, GK K, GK						
	5 G, K, GK K INVERSE						
	6 K INVERSE G, K, GK						
	7 K INVERSE K, GK						
	8 K, GK K INVERSE						
	9 K INVERSE K INVERSE						
	10 G, K, GK G, K, GK						
INVERSEUSES ET TRIODES A COMMANDE DE CATHODE	11 K G, K, GK						
	12 G, K, GK K						
	13 G, K, GK K INVERSE						
	14 K INVERSE G, K, GK						

TABLEAU 1. — Schémas équivalents.

la quatrième ligne, 3^e colonne, comme le montre la figure 15.

Si un des signaux X_1 ou X_2 est un signal K on pourra finalement aboutir au montage d'une triode unique à commande de cathode (schéma de la treizième ligne — troisième colonne).

Les exemples d'application que nous donnons au chapitre 6, ne seront autres que des combinaisons de schémas pris dans ce tableau.

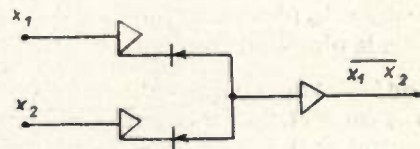


FIG. 15.

5. — CARACTÉRISTIQUES TECHNOLOGIQUES.

Les tubes utilisés sont des doubles triodes à cathodes séparées, 12 A T 7.

Au schéma de la figure 1 correspondent les valeurs suivantes :

$$R_1 = R_2 = R_3 = 10.000 \text{ ohms.}$$

$$+ E = 62 \text{ volts ;}$$

$$- e = - 28 \text{ volts.}$$

Le gain du tube est tel que, bloqué on a normalement + 2 volts sur la borne 2 et, passant (grille au potentiel zéro), on a - 5 volts sur cette même borne, ceci pour un tube de caractéristiques moyennes.

Pour un même tube $V_c = 1$ volt à - 1,5 volt. On a donc un écrêtage de 2 volts par courant de grille et de 3,5 à 4 volts par cut-off. Le signal utile étant ainsi de l'ordre de 1/7 à moins de 1/4 du signal total (7 volts). On voit que la marge de sécurité choisie est acceptable. Elle permet d'ailleurs d'utiliser des résistances à $\pm 5\%$ à condition que celles-ci restent dans leurs tolérances.

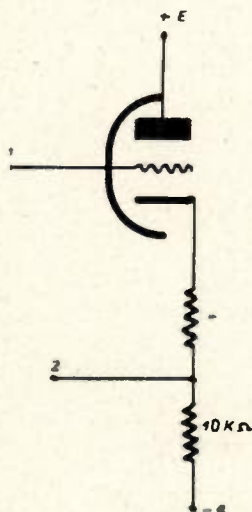


FIG. 16.

Lorsque plusieurs grilles doivent être attaquées à partir du point 2 (figure 1) des résistances de 4,7 k sont intercalées entre ce point et les grilles en question.

Il en est de même lorsque l'on attaque une grille avec un signal K ou GK. (Fig. 10, $R_g = 4,7 \text{ K}$).

La cathodyne dont le schéma est fourni par la figure 16 utilise une résistance r (de l'ordre de 300 ohms) dont le but est d'obtenir du signal sortant que sa tension la plus haute soit sensiblement égale à la tension la plus haute du signal de grille.

Dans un montage tel que celui de la figure 13, si le signal X est 0 et tous les signaux Z égaux à 1, les diodes sont alors polarisées en « tension inverse » de 2 volts au plus, que les signaux Z soient du type K ou du type KG.

Le potentiel de grille du deuxième tube est alors en butée par le courant de grille qui s'écoule principalement par le pont de résistances.

Les diodes utilisées sont des diodes au germanium, type O A 50. Les tubes 12 A T 7 sont chauffés en 12,6 V pour réduire l'importance des courants de chauffage à distribuer dans la machine.

Les signaux du type K (impulsions récurrentes généralement fournies par les dispositifs de synchronisation de la machine) sont fournis par montages à transformateur et tubes penthodes P L 82 (sortie sur une centaine d'ohms), ou tubes 12 A T 7 (sortie sur 500 à 1.000 ohms).

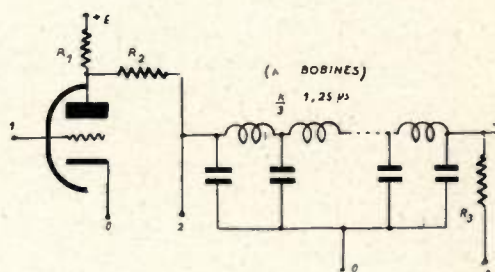


FIG. 17.

L'alimentation générale est constituée par une source de tension de 90 volts, non stabilisée. Un diviseur potentiométrique, asservi électroniquement fixera le potentiel zéro de façon à réaliser la proportionnalité de e (28 volts) à E (62 volts).

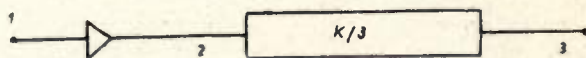


FIG. 18.

Le montage normal à triode (fig. 1) consomme environ 3 mA, tube bloqué, et 4,4 mA, tube passant (2,1 mA dans le tube). Dans ce dernier cas la dissipation anodique est de 0,038 watts, soit environ cinquante fois inférieure à la dissipation maximum permise pour ce tube.

Les lignes de retard utilisées sont des lignes à constantes localisées, à couplages inductifs.

Le retard par section (1 bobine) est de l'ordre de $\frac{1,25}{3} \mu\text{s}$ c'est-à-dire qu'il faut 3 sections pour contenir un chiffre.

La figure 17 montre le montage d'une ligne après une triode tel que nous l'utilisons.

La figure 18 est la forme symbolisée de la figure 17.

La régénération de phase (Cf. chapitre 6) peut se faire après chaque tronçon de 23 sections, donnant approximativement 9,4 μs de retard, le régénérateur complétant le retard de 10 μs .

6. — EXEMPLES D'APPLICATION.

6.1. — Régénérateur.

Le régénérateur est indispensable dans toute machine du type série et principalement dans les boucles de mémoires. Il a ici pour but de rectifier la position dans le temps des fronts de signaux.

Il est constitué par deux triodes (Fig. 19 a) à commande par la cathode : T_1 et T_2 , attaquées par la grille respectivement par X et $\bar{X}$, X étant le signal à régénérer.

Deux triodes T_3 et T_4 constituent une bascule. Le signal de commande des cathode I_R est commun

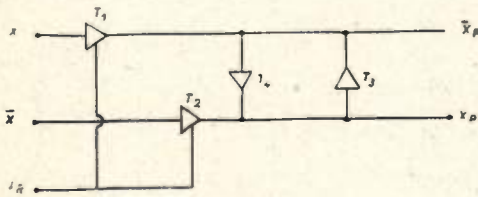


FIG. 19 a.

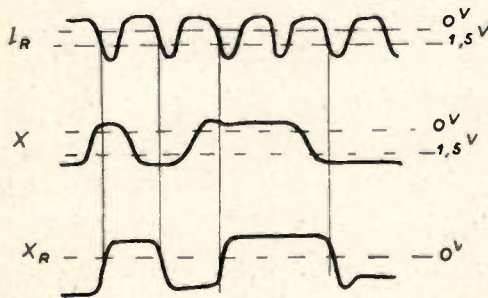


FIG. 19 b.

et représenté sur la figure 19 b diagramme, ce sont les « impulsions de rythme » à 0,8 Mc.

Sur le diagramme on a représenté également un signal X à régénérer et le signal régénéré X_R . A l'instant initial du diagramme, on a supposé la triode T_1 bloquée et la triode T_3 passante.

La bascule changera d'état sur le front négatif de l'impulsion I_R chaque fois qu'à cet instant le signal X a une nouvelle valeur par rapport à l'instant qui a précédé.

Le signal X_R est donc en retard en moyenne de $\frac{1,25}{2} \mu s$ sur le signal X .

Dans le cas de ce régénérateur des capacités de 40 pF shuntent les résistances R_2 (Cf. Fig. 1) des triodes T_3 et T_4 afin d'accroître la rapidité de basculement.

Le fonctionnement de ce régénérateur est extrêmement souple en ce sens qu'il est susceptible de régénérer des signaux dont les valeurs 1 et 0 correspondent à des tensions à peine nulles et de l'ordre de -1 volt respectivement et des impulsions dont la durée peut être réduite à moins de 0,7 μs .

6.2. — Additionneur.

L'additionneur représenté (Fig. 20) est du type « two half-adders » c'est-à-dire qu'il est constitué par deux circuits identiques triodes (1, 2, 3, 4, 5 et triodes 6, 7, 8, 9, 10) réalisant chacun les fonctions ET et OU exclusif (P_1, P_2 et S_1, S_2) le premier entre les signaux X_1 et X_2 , le second entre les signaux X'_1 et X'_2 .

On a supposé que les deux codes binaires série, à additionner étaient disponibles sous la forme de signaux du type KG :

$\bar{X}_1$ sur la borne 1 et

X_2 sur la borne 2.

Sur la figure on distingue la boucle classique de report (partie supérieure du schéma) incorporant une ligne de retard θ (retard de 1,25 μs). Il est à noter que cette boucle ne nécessite pas l'emploi de régénérateur.

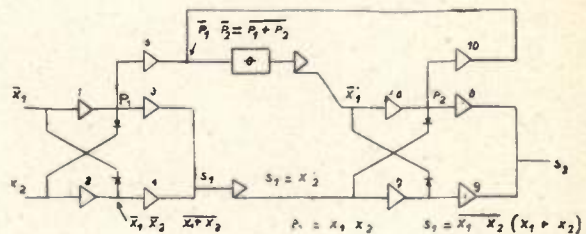


FIG. 20.

Le schéma comporte 12 triodes soit 6 tubes du type 12 A T 7.

6.3. — Registre à décalage, pas à pas.

Si on place à la suite les uns des autres une série de régénérateurs, en ayant soin d'incorporer entre

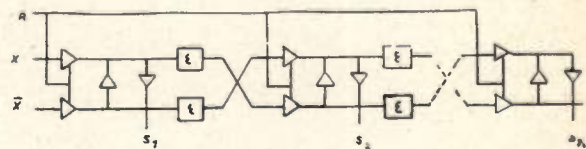


FIG. 21.

eux des éléments de retard ϵ (environ 0,5 μs) comme le montre la figure 21, on constitue ainsi un registre

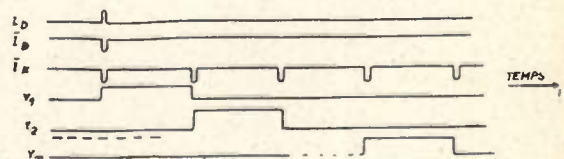


FIG. 22.

dans lequel un code peut se décaler en permanence, si les impulsions I_R sont permanentes.

Si l'on dispose d'un signal D du type K , figurant une impulsion de durée $N\theta$ (N entier, $\theta = 1,25 \mu s$), négative, ce signal appliqué en parallèle avec I_R sur les cathodes (montage du type de celui de la figure 7) le code emmagasiné par l'entrée (X et $\bar{X}$) pourra être décalé de N places vers la droite pendant la durée du signal D .

Si au lieu du signal I_R nous disposons d'une impulsion de chiffre $\bar{I}_K$ et si les éléments ϵ ont un retard

de $1,25 \mu s$, nous pourrions obtenir grâce à une impulsion de déclenchement auxiliaire I_D et $\bar{I}_D$ (fig. 22) appliquée à la place de X et $\bar{X}$, une suite de signaux rectangulaires de durée $20 \mu s$ chacun (I_K étant à la cadence $50 K_c$) sur les sorties $S_1, S_2 \dots S_n$ des bascules.

Les signaux représentés en $Y_1, Y_2, \dots Y_n$ sont en succession jointive dans le temps et peuvent être utilisés dans la machine à contrôler des transferts de codes de 16 chiffres.

REPÉRAGE D'UN GRAND NOMBRE D'INTERVALLES DE TEMPS ÉLÉMENTAIRES AU COURS D'UN CYCLE

PAR

B. LECLERC

Chef de laboratoire à la Compagnie des Machines Bull

I. — INTRODUCTION.

Dans les problèmes de transmission d'informations par impulsions, on est fréquemment conduit à considérer une « information élémentaire » caractérisée par la présence ou l'absence d'une impulsion sur un conducteur dans un intervalle de temps élémentaire.

On affecte alors un nombre défini d'intervalles de temps élémentaires consécutifs à la représentation quantitative d'une valeur instantanée de la grandeur considérée. La transmission prend ainsi un caractère cyclique ; l'on est même parfois conduit à distinguer plusieurs hiérarchies de « cycles », un « grand cycle » contenant un nombre déterminé de « petits cycles ».

Cette configuration a été rencontrée d'une façon particulièrement typique au cours de l'étude d'une Calculatrice Electronique arithmétique du type « série » travaillant sur des nombres décimaux de 12 chiffres.

Dans une calculatrice de ce type, on traduit en effet chaque chiffre par un code contenu dans une suite de 4 intervalles de temps élémentaires ; un nombre de 12 chiffres comporte alors 12 groupes de 4 intervalles élémentaires. En outre, dans l'organisation qui va être décrite, on a été conduit à créer un « grand cycle » dont la durée est égale à 8 fois le temps d'écoulement d'un nombre de 12 chiffres, soit $4 \times 12 \times 8 = 384$ intervalles de temps élémentaires.

On peut ainsi reconnaître trois hiérarchies de cycles dont la définition et le découpage ont été confiés à des distributeurs d'impulsions hiérarchisés, ou Rythmeurs ⁽¹⁾.

Après avoir passé en revue les différents signaux émis par ces distributeurs d'impulsions, nous décrirons la conception technologique des générateurs utilisés, qui s'apparentent à des multivibrateurs généralisés et font un large usage de transformateurs d'impulsions et de circuits de commutation à diodes au germanium.

(1) Nous appelons ici *rythmeurs* ce que la littérature anglo-saxonne désigne sous le nom de « Timing Circuits ». Ce terme ne nous semble pas avoir, jusqu'ici, reçu de traduction satisfaisante, les expressions « circuits d'horloge », « circuits de définition du temps » étant d'un maniement lourd, et impropres à toute déclinaison.

II. — LES RYTHMES ET LEURS FONCTIONS.

1. Rythmes « a » (fig. 1).

La période de chacun des trains d'impulsions a_1, a_2, a_3, a_4 , est égale à un intervalle élémentaire (désigné par θ dans ce qui suit).

Ces quatre trains sont décalés l'un par rapport à l'autre d'un quart de période, et les largeurs d'impulsions sont telles qu'il y ait recouvrement entre la

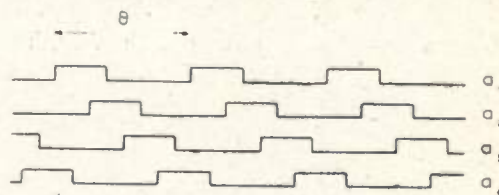


FIG. 1. — Rythmes « a ». Signaux théoriques.

fin des impulsions appartenant à un train et le début des impulsions appartenant au train suivant.

Ces signaux sont utilisés pour piloter des régénérateurs d'impulsions dont la fonction est d'assurer la

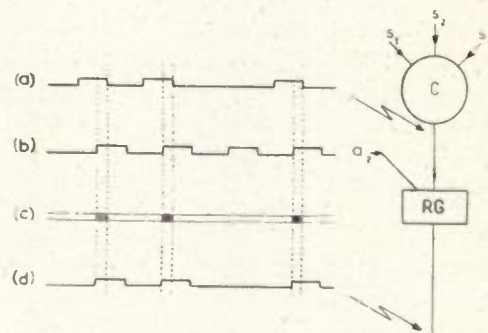


FIG. 2. — Utilisation des rythmes « a ».

Les circuits de commutation C combinent entre eux les signaux S_1, S_2 et S_3 , et délivrent le signal résultant (a) sur l'entrée du régénérateur RG . Celui-ci reçoit par ailleurs le train récurrent de rythmes « a_2 », détecte les coïncidences entre les signaux des lignes (a) et (b) et délivre à partir de chaque coïncidence (c) , une impulsion qui se prolonge jusqu'à disparition de l'impulsion de rythme « a_2 » correspondante. Les durées de fronts du signal (a) , non représentées sur la figure, peuvent atteindre les deux tiers de l'impulsion théorique sans altérer le fonctionnement, grâce au décalage du signal pilote (b) par rapport au signal à régénérer (a) .

propagation des impulsions codées à travers des circuits de commutation et des lignes à retard (fig. 2).

On peut noter que la brièveté des instants de lecture par rapport à l'intervalle de temps élémentaire permet de régénérer sans ambiguïté des signaux dont les temps de front sont relativement longs. On parvient à ce résultat en utilisant sur le régénérateur un circuit de réaction qui prolonge l'impulsion engendrée au-delà de l'instant de lecture.

On verra comment une technique tout à fait analogue est utilisée dans le générateur d'impulsions lui-même pour engendrer les rythmes.

Une utilisation typique des rythmes « a » est d'assurer une propagation synchrone des impulsions dans les différentes « boucles de mémoire » du calculateur.

En effet, les impulsions codées caractéristiques d'un nombre sont conservées en propagation dans un circuit bouclé comportant des éléments de retard et des régénérateurs d'impulsions pilotés par des signaux de rythme « a ».

2. Rythmes « b » (fig. 3).

Le signal b_1 permet de distinguer des intervalles de temps pairs et des intervalles de temps impairs. Pour des raisons de commodité, on fait délivrer par le générateur le signal complémentaire $\bar{b}_1$.

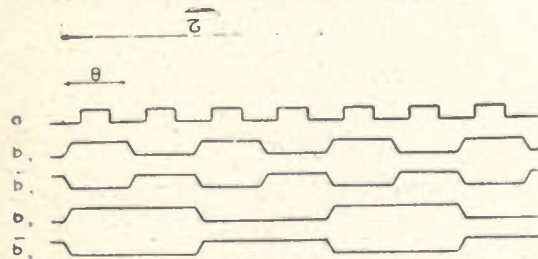


FIG. 3. — Rythmes « b ».

b_1 est positif dans les intervalles de temps pairs, $\bar{b}_1$ est positif dans les intervalles de temps impairs. La période du signal b_2 est double de celle du signal b_1 .

En combinant les signaux b_1 , $\bar{b}_1$, b_2 , $\bar{b}_2$, on saura distinguer un intervalle quelconque dans une suite de quatre intervalles élémentaires.

Cette possibilité est utilisée dans la calculatrice pour distinguer les poids arithmétiques 1, 2, 4 et 8 constitutifs du code d'un chiffre décimal.

La durée des fronts des rythmes « b » est telle qu'une combinaison de ces signaux peut être lue, par un régénérateur piloté en rythme « a », aux instants de recouvrement ($a_1 a_2$), ($a_2 a_3$), ($a_3 a_4$), mais non pas ($a_4 a_1$) qui apparaît dans une période d'indétermination.

3. Rythmes « d » (fig. 4).

Le code représentatif d'un chiffre décimal occupant 4 intervalles, les rythmes « d » ont été créés pour repérer 12 ordres décimaux dans une suite de 48 intervalles.

Les signaux d_1 , $\bar{d}_1$, d_2 , $\bar{d}_2$, sont homologues aux signaux b déjà vus, et permettent le repérage d'un ordre décimal parmi 4.

Chacun des signaux d_3 , d_4 , ou $\bar{d}_5$ est affecté au repérage d'un groupe de 4 ordres décimaux.

Les rythmes « d » sont utilisés notamment dans les circuits de transposition espace-temps et temps-espace qui assurent la liaison des boucles de mémoire, où les informations se présentent en succession dans

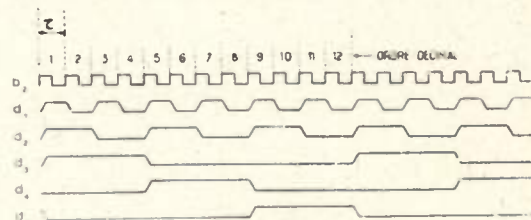


FIG. 4. — Rythmes « d ».

le temps, avec des circuits d'introduction et d'extraction « parallèles », où l'on affecte un conducteur à chaque ordre décimal.

Ce sont les mêmes problèmes de transposition qui se posent dans les dispositifs de transmission multiplex à sélection de temps.

La transposition espace-temps s'effectue dans une matrice de redresseurs dont chaque ligne est affectée à un ordre décimal et dont les colonnes sont alimentées par les rythmes « d » (fig. 5).

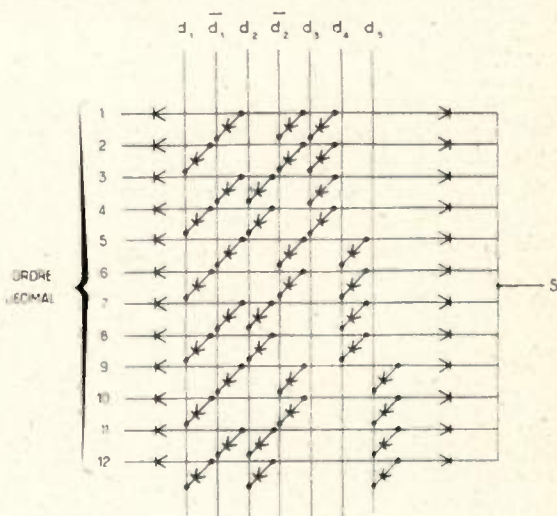


FIG. 5. — Pyramide de transposition Espace-Temps.

Une méthode réciproque est utilisée pour effectuer la transposition temps-espace.

4. Rythmes « e » (fig. 6).

Les rythmes « e », en combinaison avec les signaux déjà vus, permettent la définition d'un grand « cycle » dont la période est égale à 8 fois le temps d'écoulement d'un nombre de 12 chiffres.

Un signal « e_0 » ou « $\bar{e}_0$ » est défini positif pendant le même temps qu'un signal d_3 , d_4 ou $\bar{d}_5$.

On voit, d'après le diagramme, que les signaux e_0 , $\bar{e}_0$ sont alternativement établis en face de chacun des signaux d_3 , d_4 , ou $\bar{d}_5$. Les 6 combinaisons que l'on peut réaliser avec ces signaux définissent par conséquent des positions distinctes dans un cycle dont la durée est 3 fois la période e_0 .

Ce résultat est équivalent à celui qu'on aurait pu obtenir dans un distributeur hiérarchisé de structure classique, avec des signaux e_0 , $\overline{e_0}$ de durée trois fois plus grande.

Le phénomène de précession qui permet d'utiliser des signaux de même durée apparaît chaque fois que les nombres de sorties dans chaque groupe de générateurs sont premiers entre eux, comme c'est ici le cas pour 2 et 3.

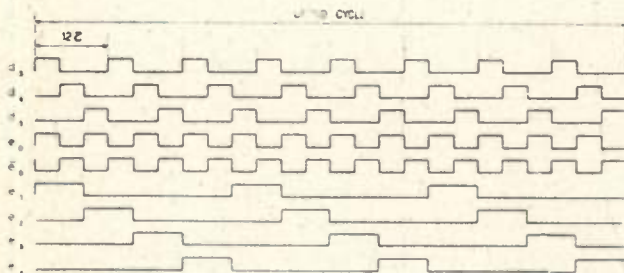


FIG. 6. — Rythmes e_0 . — Signaux théoriques.

De la même façon, l'association aux signaux précédents des signaux e_1 , e_2 , e_3 , e_4 dont la durée est égale à une période e_0 définit un cycle de durée $3 \times 2 \times 4 = 24$ durées e_0 . En effet, l'association de e_0 , $\overline{e_0}$ avec e_1 , e_2 , e_3 , e_4 définit 8 intervalles de temps dans une période e_1 , comme le ferait un générateur à 8 sorties.

L'association considérée est donc équivalente à l'association de 2 groupes comportant respectivement 8 et 3 sorties, et le principe énoncé ci-dessus s'applique puisque les nombres 8 et 3 sont premiers entre eux.

Cette faculté que l'on a de réduire la durée des signaux de hiérarchie élevée est avantageuse lorsque, comme c'est le cas ici, les générateurs utilisent des transformateurs d'impulsions.

III. — TECHNOLOGIE DES RYTHMEURS.

1. Généralités.

Les étages générateurs de rythmes sont apparentés à des multivibrateurs généralisés.

Comme nous l'avons dit, ils ont été conçus pour être incorporés à une Calculatrice Electronique arithmétique.

Pour une raison d'homogénéité, ils utilisent les mêmes organes fondamentaux qui ont été choisis pour la calculatrice proprement dite.

En particulier, les commutations y sont réalisées par des circuits comportant principalement des redresseurs au germanium.

De tels circuits peuvent toujours se ramener à deux circuits élémentaires, illustrés par la figure 7 (a) et (b), qui apparaissent en groupements plus ou moins complexes.

Supposons que les tensions sur chacune des entrées a_1 , b_1 , c_1 , a_2 , b_2 , c_2 , ne puissent prendre que l'une des deux valeurs 0 et -10 volts.

Si le signal affirmatif est 0 volts, dans le cas de la figure 7 (a) on peut dire que l'on recueille un signal à la sortie s_1 , si les signaux a_1 et b_1 ... et n_1 sont simultanément présents.

Dans le cas de la figure 7 (b), on peut dire que l'on recueille un signal à la sortie s_2 si l'un des signaux a_2 ou b_2 ... ou n_2 est présent, ou si plusieurs d'entre eux sont présents simultanément.

Dans le premier cas, on réalise la combinaison logique « et ». Nous dirons alors que nous avons affaire à un conditionneur, car l'action d'un signal apparaissant sur l'une des entrées est conditionnée par la présence simultanée de signaux convenables sur les autres entrées.

Dans l'autre cas, on réalise la combinaison logique du « ou » inclusif. Nous dirons alors que nous avons affaire à un mélangeur, car ce circuit est souvent utilisé pour regrouper des conditionnements logiques créés par ailleurs.

Il est évident que si, au contraire, on choisit comme affirmatif le niveau -10 volts, les fonctions logiques des deux circuits ci-dessus sont inversées, le circuit figure 7 (a) devenant mélangeur, et le circuit figure 7 (b) devenant conditionneur.

L'utilisation, dans les rythmeurs et dans la calculatrice proprement dite, de circuits de commutation à diodes au germanium conduit à délivrer les signaux de rythme avec une amplitude faible, mais un débit élevé.

L'emploi de transformateurs d'impulsions permet d'adapter la basse impédance de l'utilisation aux circuits plaque des tubes à vide qui délivrent la puissance.

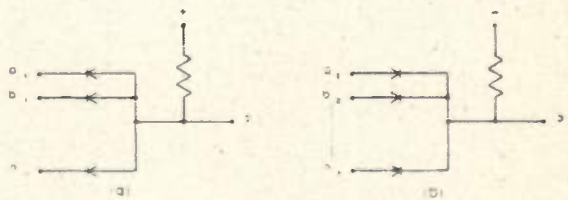


FIG. 7. — Conditionneur et mélangeur.

Il permet en outre, en réalisant des transformateurs à plusieurs secondaires, de disposer du signal fourni par un étage dans les deux polarités et sous plusieurs niveaux distincts, et laisse toute latitude dans le choix des composantes continues.

2. Schéma de principe d'un étage type.

Tous les étages générateurs sont établis suivant un même schéma type, et seuls le dimensionnement des organes et l'interconnexion des étages distinguent les uns des autres les générateurs appartenant aux différents groupes.

Le principe de fonctionnement de l'étage type est le suivant :

Un signal pilote, commun à tous les étages d'un même groupe, apparaît chaque fois qu'un changement d'état doit affecter l'un quelconque des étages de ce groupe (ou plusieurs d'entre eux simultanément).

Ce signal pilote n'a d'action sur un étage donné que pour une combinaison donnée d'états sur l'ensemble des étages du groupe.

Cette action se traduit, soit par une « commande de début » dont l'effet est d'amorcer la génération d'un signal sur l'étage considéré, précédemment au repos,

soit par une « commande de fin » qui provoque le retour au repos de cet étage.

Entre la « commande de début » et la « commande de fin », l'étage est maintenu au travail par un circuit de réaction.

Dans la pratique, le signal pilote est subdivisé en deux signaux distincts, de polarités inversées.

Le signal de polarité positive est appliqué en 1 (a) sur le schéma de la figure 8, et conditionne la commande de début. Le signal pilote de polarité négative est appliqué en 1 (b) et conditionne la commande de fin.

L'interdépendance des étages est réalisée grâce aux entrées 2 (a, b... n) et 3 (a, b... n) qui conditionnent respectivement les signaux pilotes appliqués en 1 (a) et 1 (b).

Le circuit de réaction est réalisé par un secondaire du transformateur de sortie, basé sur une tension continue, qui attaque le circuit grille par l'intermédiaire d'une résistance R .

Le signal engendré est délivré dans les deux polarités, en s et s, par deux autres secondaires du même transformateur.

Le circuit utilisé pour regrouper la commande de début, la commande de fin et le signal de réaction sur le point 6 du schéma exige évidemment que le potentiel du point 5 soit à tout moment supérieur ou égal au potentiel du point 4. On doit tenir compte de ce fait dans la formation des commandes de début et de fin. En particulier, ces commandes ne peuvent jamais coexister.

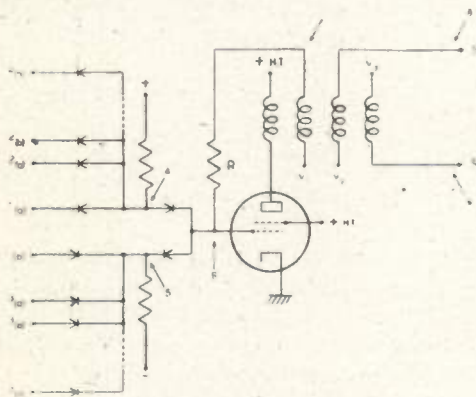


Fig. 8. — Etage type.

Dans la pratique, la puissance demandée à chaque étage générateur conduit à utiliser, dans le primaire du transformateur de sortie, un ou plusieurs tubes dont le recul de grille est important.

Or, la recherche d'une bonne sécurité de fonctionnement des circuits de commutation à diodes au germanium exige l'utilisation, dans ces circuits, de signaux d'amplitude faible. L'on est donc conduit à prévoir pour chaque étage un tube pilote à faible recul de grille qui attaque le ou les tubes de sortie par l'intermédiaire d'un transformateur.

3. Interconnexion des étages.

Rythmes « a » :

Le groupe de 4 étages qui engendrent les rythmes « a » est piloté par un signal en forme de créneau régulier dont la fréquence (1 Mc/s) est quatre fois plus

élevée que celle des signaux délivrés par chacun des étages (fig. 9).

Dans ce cas particulier, par suite de l'égalité de durée entre alternances positives et négatives du signal pilote et en raison du recouvrement recherché entre les différents rythmes « a », le même signal est appliqué, sur chaque étage, à la fois sur l'entrée 1 (a) et sur l'entrée 1 (b).

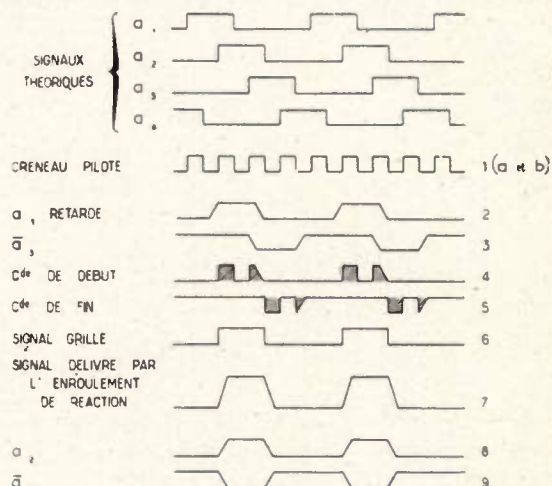


Fig. 9. — Génération des rythmes « a ».

Les chiffres à droite de la figure renvoient aux points correspondants du schéma fig. 8.

Dans la formation de la « commande de début », le créneau pilote est conditionné par le signal de polarité positive issu de l'étage précédent, et convenablement retardé dans une ligne à retard à constantes localisées.

La « commande de fin » est fournie par le créneau pilote, conditionné par le signal de polarité négative, désigné par a_3 sur la figure 9, issu de l'étage suivant.

On peut noter, d'après les lignes 4 et 5 de la figure 9, que les signaux « commande de début » et « commande de fin » comportent des résidus de commutation en dehors de leur portion utile. Mais l'effet de ces résidus est nul, car ils apparaissent dans une position telle qu'ils ne font que confirmer un état existant.

Si les conditions d'interconnexion qui viennent d'être décrites étaient appliquées à l'ensemble des quatre étages du groupe, celui-ci se maintiendrait en fonctionnement s'il était convenablement démarré. Toutefois, lors de la mise sous tension, chaque étage demandant, pour démarrer, la présence d'un signal issu de l'étage précédent, on risquerait, selon le mode d'établissement des tensions d'alimentation, soit de voir l'ensemble des quatre étages rester au repos, soit de voir l'ensemble démarrer sur un faux mode, c'est-à-dire en délivrant des signaux différents de ce qui vient d'être décrit.

On résout cette difficulté en modifiant la définition de la commande de début pour l'un des quatre étages.

Pour cet étage particulier, le créneau sera conditionné non plus par la sortie positive retardée de l'étage précédent, mais simultanément par les sorties négatives des deux étages suivants. Ce double conditionnement est réalisé en utilisant les entrées 2 (a) et 2 (b) de la figure 9.

Groupes à 2 sorties.

Chaque groupe dit « à 2 sorties » comporte un seul étage délivrant deux signaux complémentaires caractérisés par l'égalité de durée des alternances positives et négatives.

En règle générale, les commandes de début et de fin d'un groupe à 2 sorties sont obtenues en conditionnant les signaux pilotes par l'état de l'unique étage qui constitue ce groupe. Cette action de l'étage sur lui-même a lieu en connectant sa sortie $\bar{s}$ (fig. 8) aux conditionneurs formant les commandes de début

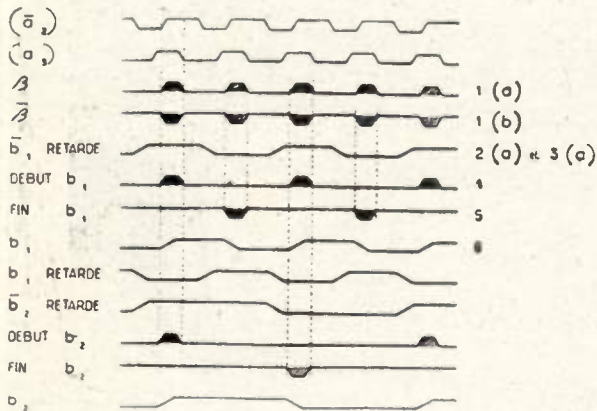


FIG. 10. — Génération des rythmes « a ».

et de fin, à travers un élément de retard dont le retard est au moins égal au temps d'établissement ou de coupure du signal engendré par l'étage.

On voit par exemple, sur la figure 10 et sur le tableau de la figure 11, que le signal de début du rythme b_1 est obtenu lorsque le signal pilote β et le signal b_1 retardé sont simultanément positifs.

RYTHMES	SIGNAL PILOTE DEBUT	SIGNAL PILOTE FIN	CONDITION DEBUT : DOIVENT ETRE SIMULT POSITIFS	CONDITION FIN : DOIVENT ETRE SIMULT NEGATIFS
a_1	CRENEAU IM e_0	CRENEAU IM e_0	$\bar{a}_2, \bar{a}_3$	$\bar{a}_2$
a_2	..	..	$\bar{a}_2, R$	$\bar{a}_3$
a_3	..	..	$\bar{a}_2, R$	$\bar{a}_4$
a_4	..	..	$\bar{a}_3, R$	$\bar{a}_4$
b_1	β	$\bar{\beta}$	$\bar{b}_1, R$	$\bar{b}_1, R$
b_2	..	..	$\bar{b}_1, R, \bar{b}_2, R$	$b_1, R, \bar{b}_2, R$
d_1	b_2	$\bar{b}_2$	$\bar{d}_1, R, \bar{d}_2, R$	$\bar{d}_1, R, \bar{d}_2, R$
d_2	..	..	$\bar{d}_1, R, \bar{d}_2, R$	$d_1, R, \bar{d}_2, R$
e_0	..	..	$\bar{d}_1, R, \bar{d}_2, R, \bar{e}_0, R$	$d_1, R, \bar{d}_2, R, \bar{e}_0, R$
d_3	δ	$\bar{\delta}$	$\bar{d}_3, R, \bar{d}_4, R$	$\bar{d}_3, R$
d_4	..	..	d_3, R	d_3, R
d_5	..	..	d_4, R	d_4, R
e_1	γ	$\bar{\gamma}$	$\bar{e}_1, R, \bar{e}_2, R, \bar{e}_3, R$	$\bar{e}_1, R$
e_2	..	..	e_1, R	e_1, R
e_3	..	..	e_2, R	e_2, R
e_4	..	..	e_3, R	e_3, R

R = "RETARDE"

$$\beta = a_2 a_3 \quad \delta = b_2 d_1 R d_2 R \quad \gamma = \delta e_0$$

FIG. 11. — Une notation telle que : $\delta = b_2 d_1 R d_2 R$ signifie que δ est affirmatif lorsque les signaux composants b_2 , $d_1 R$ et $d_2 R$ sont simultanément affirmatifs.

De même, le signal de fin est obtenu lorsque le signal pilote β et le signal b_1 retardé sont simultanément négatifs.

Dans certains cas, un même signal pilote est utilisé pour deux ou plusieurs groupes à 2 sorties, de hiérarchies différentes.

On complète alors les conditionnements de début et de fin de chaque étage en faisant intervenir sur un groupe donné le ou les groupes de hiérarchie inférieure dépendant du même signal pilote. C'est ainsi que $\bar{d}_1 R$ et $\bar{d}_2 R$ interviennent, avec $\bar{e}_0 R$, dans la condition de début du signal e_0 .

Groupes à plus de deux sorties.

Un groupe à n sorties comporte n étages dont chacun délivre un signal de durée égale à un $n^{\text{ème}}$ de période. Les différents signaux délivrés sont décalés, d'un étage au suivant, d'une quantité égale à la durée du signal, et, par conséquent, ne se recouvrent pas.

En toute rigueur, il suffirait, pour obtenir n sorties, de disposer de $n-1$ étages, la dernière sortie étant obtenue en mettant en évidence l'absence simultanée de signal sur tous les étages du groupe.

Toutefois, cette relation n'est généralement pas mise à profit dans les groupes à plus de deux sorties où elle est difficilement exploitable.

Nous ne nous étendons pas sur les interconnexions réalisées. Le tableau de la figure 11 sera facilement compris d'après ce qui précède.

Notons simplement un artifice analogue à celui qui a été décrit pour les Rythmes « a », et qui consiste, pour assurer un démarrage correct lors de la mise sous tension, à conditionner le signal de début d'un étage dans chaque groupe par l'absence simultanée de signaux sur les autres étages du groupe.

IV. — PERFORMANCES.

Des générateurs répondant aux principes qui viennent d'être exposés on fait l'objet d'une réalisation, illustrée par la figure 12.

Les propriétés générales de cet ensemble vont être décrites ci-dessous :

1. Caractéristiques des signaux délivrés.

Nous allons donner ici les caractéristiques des Rythmes « a », et celles des Rythmes « e », qui se situent aux deux extrémités de la hiérarchie de signaux décrite au paragraphe II.

Les Rythmes « a » sont émis à la cadence de 250 Kc/s. Les impulsions engendrées ont une durée de 1,5 microseconde, et un temps de front de 0,15 microseconde.

La puissance utile de crête délivrée par chacun des 4 étages est de 17 watts.

Chaque étage comporte deux sorties de niveaux différents en polarité positive, et une sortie en polarité négative.

Les Rythmes « e » sont émis à une cadence très voisine de 2 Kc/s. Les impulsions qui les constituent ont une durée de 128 microsecondes, et un temps de front de 6 microsecondes.

La puissance utile de crête délivrée par chaque étage est de 8 watts.

Chaque étage comporte trois sorties de niveaux différents en polarité positive, et une sortie en polarité négative.

2. Stabilité.

Influence de la fréquence :

A charge nominale des étages, la fréquence du créneau de base peut être modifiée du simple au double, soit environ $\pm 35\%$ de part et d'autre de la fréquence nominale, sans altérer le fonctionnement de l'ensemble.

Influence du vieillissement des tubes.

La charge nominale des étages, qui n'est jamais dépassée dans la pratique, a été choisie telle que le débit anodique des tubes qui lui correspond soit égal à la moitié du débit nominal résultant de l'examen des caractéristiques théoriques de ces tubes.

Dans ces conditions, le vieillissement des tubes est sans influence sur le fonctionnement de l'ensemble.

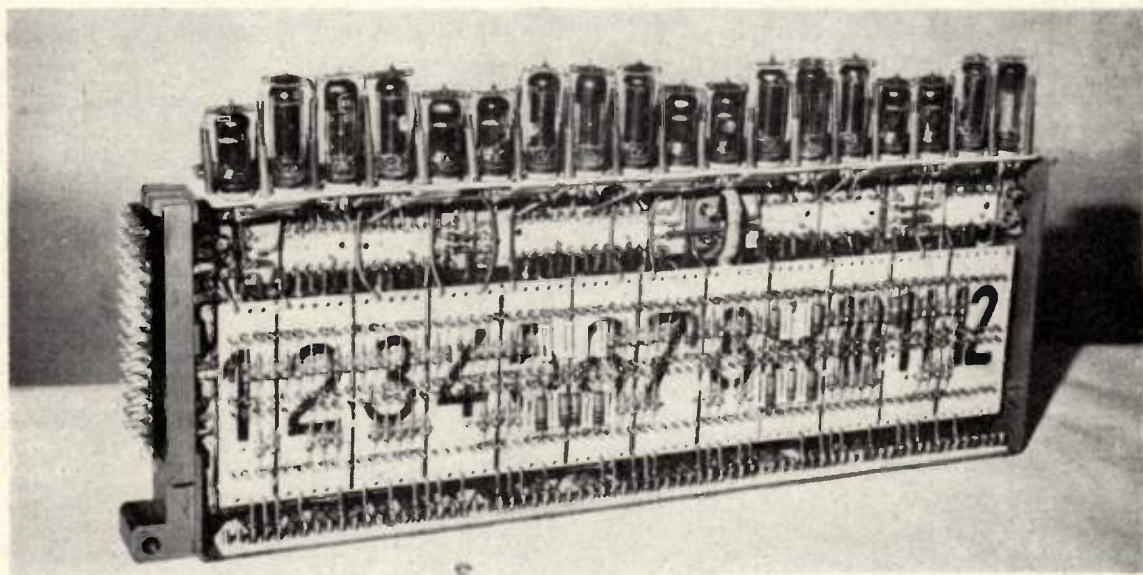


FIG. 12. — Aspect d'un châssis.

Les groupes (d_3, d_4, d_5) et (e_1, e_2, e_3, e_4) sont réunis sur ce châssis. On distingue les transformateurs d'impulsions à proximité immédiate des supports de lampes. La partie centrale du châssis est occupée par la plaque isolante qui supporte les circuits de commutation. Les résistances relatives à ces circuits sont à plat sur la tranche inférieure du châssis. Les lignes à retard, disposées derrière la plaque centrale ne sont pas visibles sur la figure.

Les limites sont dues à la présence de lignes à retard, dimensionnées pour la fréquence nominale. Un déplacement excessif de la fréquence de fonctionnement entraîne un décalage des signaux transmis par ces lignes, et une altération des commandes de début et de fin.

Influence des tensions d'alimentation :

Si l'on agit sur la tension du secteur, la plage de variation admissible est pratiquement limitée par les conditions de chauffage des tubes, et il n'est pas recommandé de dépasser les valeurs extrêmes admises par les constructeurs dans ce domaine, soit $\pm 10\%$.

Si la tension de chauffage est maintenue constante, toutes les autres tensions variant d'une façon homothétique, des variations beaucoup plus importantes sont admissibles, et le fonctionnement reste correct au quart de la tension nominale. Du côté des tensions élevées, la limite est définie par la dissipation maximum admissible sur les électrodes des tubes et dans les résistances des circuits de commutation.

On est toujours maître des variations indépendantes de certaines tensions par rapport à d'autres. Les circuits de commutation des rythmeurs dont les performances sont décrites ici ont été calculés en supposant ces variations limitées à $\pm 5\%$. Il était évidemment possible de prendre toute autre hypothèse, mais ce chiffre s'est révélé surabondant dans tous les cas pratiques d'emploi.

dans de très larges limites de variation des caractéristiques.

3. Rendement énergétique.

Le rendement énergétique est défini comme le rapport de la puissance utile délivrée par un étage chargé au nominal, à la somme des puissances appliquées aux différents éléments de l'étage. Cette dernière quantité inclut en particulier la puissance de chauffage des tubes, la puissance consommée par le tube pilote, et la puissance dissipée dans les résistances incorporées aux circuits de commutation. Elle tient compte également du rendement du transformateur de sortie.

Dans ces conditions, le rendement énergétique global d'un étage de Rythme « a » est de 35 %.

Si l'on exclut de ce bilan la puissance de chauffage des tubes, le rendement énergétique est de 50 %.

4. Matériel mis en œuvre.

L'ensemble des 16 étages générateurs et de l'oscillateur pilote comporte un total de 53 tubes, 250 diodes, 17 éléments de retard et 33 transformateurs d'impulsions.

Grâce à l'emploi de transformateurs à plusieurs secondaires, cet ensemble présente 51 sorties distinctes, la plupart des signaux étant délivrés dans les deux polarités et sous plusieurs niveaux.

Il faut noter que la faculté de disposer d'un signal sous divers niveaux permet, en choisissant pour chaque utilisation le niveau le mieux adapté, d'augmenter dans de larges proportions le rendement énergétique d'exploitation de l'étage qui délivre ce signal, et se traduit finalement par une économie de matériel et de puissance.

5. Résultats d'exploitation.

Les Rythmeurs dont nous venons de décrire les propriétés ont été incorporés à une Calculatrice Electronique actuellement construite en série, et dont plus de 30 exemplaires sont actuellement en service.

Ces circuits s'avèrent d'un fonctionnement très sûr, et les rares défauts constatés en exploitation sont dus pour la plupart à des accidents survenant aux lampes (ruptures filaments, court-circuits internes). Ces défauts apparaissent d'ailleurs, en général, dans les 200 premières heures de fonctionnement.

Il n'est pas actuellement possible de donner d'indication sur l'influence du vieillissement des organes, car les 5 000 heures de fonctionnement du premier en date de ces générateurs sont insuffisantes pour provoquer des variations de caractéristiques susceptibles d'intervenir sur le fonctionnement de l'ensemble.

Les performances étudiées au paragraphe IV permettent de situer ce type de générateur par rapport à d'autres solutions d'un emploi plus classique, et font en particulier ressortir son intérêt lorsque le problème posé exige l'emploi de distributeurs d'impulsions capables de délivrer des puissances relativement élevées, ou comportant de très nombreuses sorties.

L'exposé a été centré sur une application particulière, mais les principes de base qui s'en dégagent ont un caractère général indépendant, dans une large mesure, de la forme des signaux délivrés, du nombre de sortie des distributeurs, de la fréquence de base et de la puissance utile.

LIGNES A RETARD A CONSTANTES LOCALISÉES

PAR

H. FEISSEL

*Ingénieur de Recherches
à la Compagnie des Machines Bull*

1. INTRODUCTION :

Une ligne à retard à constantes localisées est généralement apparentée à un filtre passe bas.

Historiquement, les premières réalisations ont d'abord utilisé des structures passe bas prototypes. Dans de telles structures, la relation phase fréquence est assez peu linéaire et, par conséquent, le retard n'est pas constant dans la bande passante.

Des résultats meilleurs ont été obtenus avec des structures dérivées en m , mais la réalisation matérielle de lignes comportant des sections dérivées en m conduit à un encombrement relativement important car les différentes sections doivent être éloignées les unes des autres suffisamment pour éliminer les couplages parasites.

Pour des considérations d'économie et d'encombrement évidentes, il est intéressant de placer toutes les bobines constitutives d'une ligne à retard sur un même mandrin cylindrique (disposition coaxiale).

Dans ce cas, en raison des mutuelles entre bobines non adjacentes, la ligne n'est plus assimilable à une suite de quadripôles élémentaires, comme c'était le

cas pour les structures passe bas dérivées en m . Dans ces conditions, le calcul rigoureux des caractéristiques de transmission devient extrêmement ardu, et l'on est conduit à étudier ces structures en effectuant des relevés expérimentaux.

Dans la recherche d'un dimensionnement optimum, on sera parfois tenté, pour guider l'expérience, de faire des hypothèses simplificatrices permettant d'appliquer la théorie des filtres, mais dans chaque cas particulier, la validité de ces hypothèses devra être éprouvée a posteriori en confrontant les résultats du calcul simplifié avec les résultats expérimentaux.

2. APPAREILLAGE DE MESURE ET DÉPOUILLEMENT DES RÉSULTATS.

Pour effectuer des relevés expérimentaux du retard et de l'atténuation en fonction de la fréquence, un montage suivant figure 1 a été utilisé :

Un générateur sinusoïdal à fréquence variable attaque simultanément un atténuateur étalonné et la ligne à étudier, terminée de part et d'autre sur

son impédance caractéristique ; les deux amplificateurs qui suivent attaquent chacun une paire de plaques d'un tube cathodique et l'on note la suite des fréquences pour lesquelles le déphasage dans la ligne est égal à $K\pi$.

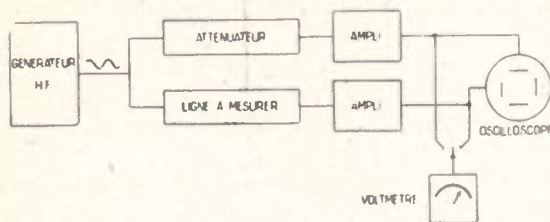


FIG. 1. — Méthode de relevé des caractéristiques retard-fréquence et atténuation-fréquence.

Le nombre de points de mesure est d'autant plus grand que la ligne à relever comporte plus de sections.

Soit F une fréquence pour laquelle le déphasage est $K\pi$. Le retard à cette fréquence est :

$$\tau = \frac{\varphi}{\omega} = \frac{K\pi}{2\pi F} = \frac{K}{2F}$$

On obtient ainsi des points d'une courbe définie jusqu'au voisinage de la fréquence de coupure, où les mesures deviennent difficiles en raison de la grande atténuation.

La similitude des réponses en amplitude et en phase des deux amplificateurs doit être soigneusement ajustée et contrôlée, ce que l'on effectue en alimentant leurs entrées en parallèle.

L'atténuateur doit introduire un déphasage négligeable.

Pour avoir une précision de l'ordre de ± 1 pour mille dans la mesure des fréquences, on utilise un oscillateur à fréquence variable contrôlé par battements avec une fréquence de référence.

Un voltmètre commutable sur l'une ou l'autre des sorties des amplificateurs permet de repérer l'égalité des niveaux que l'on obtient en agissant sur l'atténuateur.

On peut ainsi mesurer l'atténuation dans la ligne en fonction de la fréquence.

A partir des relevés effectués sur l'appareillage qui vient d'être décrit, on trace des courbes en coordonnées réduites, en rapportant le retard de la structure et les fréquences transmises au retard τ_0 à la fréquence zéro.

On trouve donc en abscisses des fréquences réduites $\tau_0 F$, et en ordonnées des temps réduits $\frac{\tau}{\tau_0}$.

Le retard τ_0 n'est pas directement mesurable. Il est défini par extrapolation comme la limite du retard lorsque la fréquence tend vers zéro.

Certains auteurs rapportent les coordonnées réduites à la fréquence de coupure F_c , mais cette notation, commode lorsque l'on présente des courbes calculées, n'a pas été adoptée ici, car la définition expérimentale de F_c par extrapolation manque de précision.

Pour comparer entre elles les performances de diverses structures, il reste à choisir une longueur unitaire pour chaque type de structure étudié.

Dans le cas d'une ligne idéale, si l'on considère une longueur unitaire telle que le déphasage qu'elle introduit à la fréquence de coupure est égal à 2π , la courbe retard-fréquence qui lui est relative, tracée dans les coordonnées réduites qui viennent d'être définies, sera un segment de droite d'ordonnées unité et de longueur unité.

Les courbes des filtres passe bas dérivés en m , tracées pour différentes valeurs de m en rapportant le retard à deux cellules consécutives sont limitées à la fréquence de coupure qui correspond à un déphasage de π par cellule. Les points limites de ces courbes sont situés sur l'hyperbole $y = \frac{1}{x}$ qui passe par le point de coordonnées 1/1.

Pour obtenir des résultats comparables dans les relevés effectués sur des structures coaxiales, les courbes de retard sont rapportées à une portion de ligne élémentaire comprenant deux bobines.

3. CARACTÉRISTIQUES DE STRUCTURES CONNUES.

La figure 2 présente un réseau de courbes retard-fréquence théoriques relatives à des passe-bas dérivés en m . Les retards ont été calculés en $\frac{\varphi}{\omega}$, et les courbes tracées en coordonnées réduites comme il a été dit précédemment.

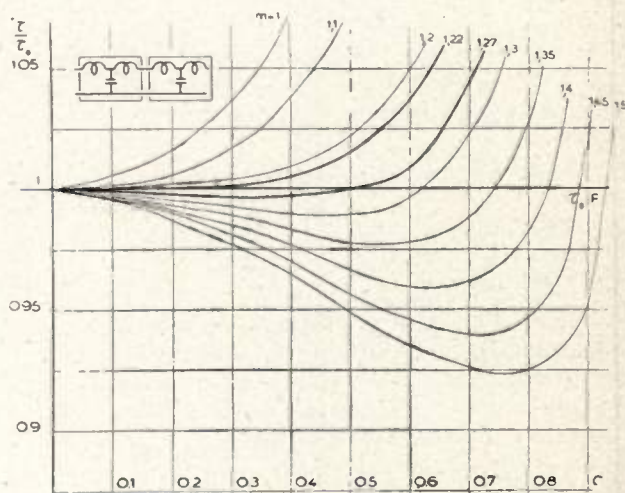


FIG. 2. — Courbes retard-fréquence des passe-bas dérivés en m .

La figure 3 présente un réseau relevé expérimentalement sur des lignes comportant des bobines coaxiales et pour des coefficients de couplage entre bobines adjacentes $K = \frac{M}{L}$ compris entre 5 et 20 %.

Ce réseau a été complété sur la figure par la courbe calculée du passe bas prototype, qui correspond à un couplage nul entre bobines.

L'examen comparé des figures 2 et 3 fait apparaître des différences de nature profondes entre les deux réseaux de courbes.

Le retard augmente aux fréquences moyennes pour des valeurs progressivement croissantes de α .

Pour des valeurs de α voisines de l'optimum, on voit que les courbes sont en moyenne descendantes pour des valeurs trop fortes du couplage, et montantes dans le cas contraire.

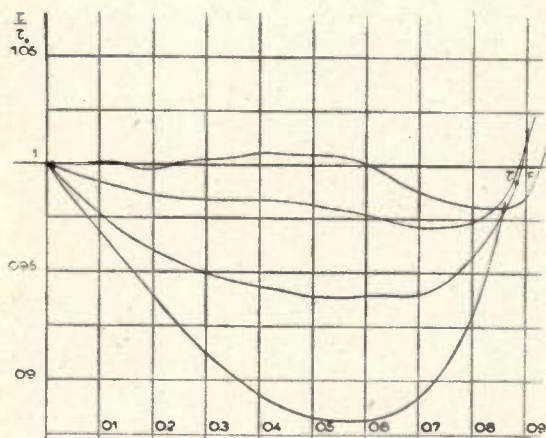


FIG. 10. — Courbes retard-fréquence de structures coaxiales corrigées pour diverses valeurs de α . Couplage trop élevé (20 %).

La connaissance de cette allure générale des courbes permet de définir dans quel sens il faut agir pour améliorer les performances d'une ligne dont on a relevé la caractéristique retard-fréquence.

Dans le choix d'une caractéristique retard-fréquence, on s'attachera à obtenir un retard très constant pour les fréquences les plus élevées dans la bande à transmettre et l'on pourra être moins sévère

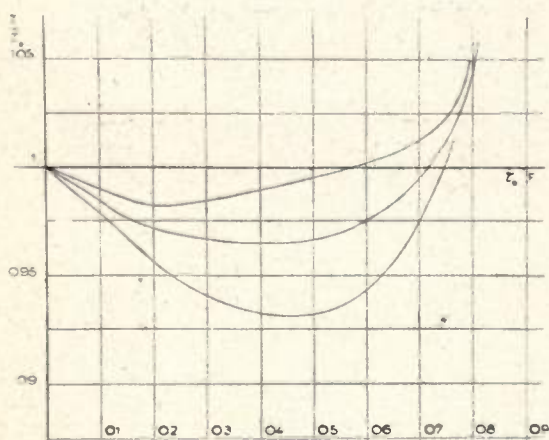


FIG. 11. — Courbes retard-fréquence de structures coaxiales corrigées pour diverses valeurs de α . Couplage trop faible (15 %).

vers les fréquences basses car, d'une part, le retard associé à un déphasage donné est inversement proportionnel à la fréquence et, d'autre part, la distortion introduite par un déphasage petit sur une composante donnée d'un signal est proportionnelle à ce déphasage.

Inévitablement, les courbes retard-fréquence sont très incurvées au voisinage de la fréquence de coupure. Si l'atténuation n'est pas déjà grande à la fréquence pour laquelle le retard commence à varier appréciablement, il en résulte l'apparition d'oscillations

transitoires dans la réponse de la ligne à l'échelon unité.

Pour éliminer ces oscillations, il a été trouvé avantageux d'associer à une structure comportant un grand nombre de sections de fréquence de coupure F_c quelques sections de même nature et de fréquence de coupure $0,8 F_c$ par exemple.

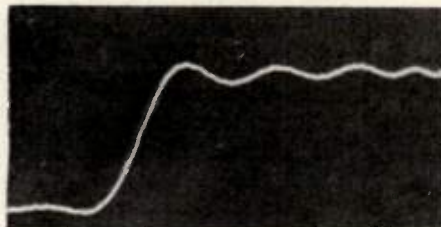


FIG. 12. — Transmission de l'échelon unité à travers 20 cellules d'une ligne corrigée.

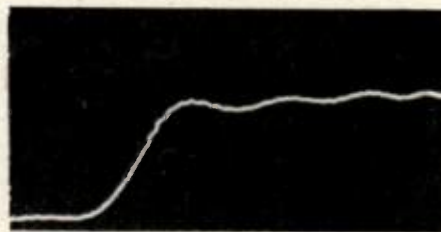


FIG. 13. — Transmission de l'échelon unité à travers 40 cellules d'une ligne corrigée.

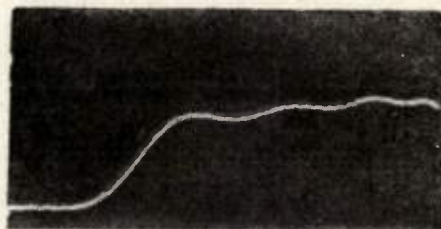


FIG. 14. — Transmission de l'échelon unité à travers 56 cellules d'une ligne corrigée.

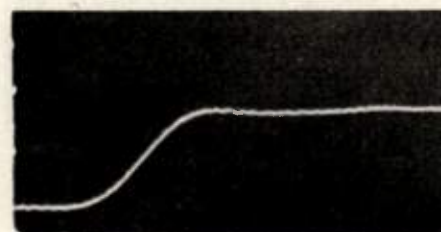


FIG. 15. — Transmission de l'échelon unité à travers 59 cellules, dont 3 cellules à la fréquence de coupure $0,8 F_c$.

L'adjonction de ces quelques sections altère très peu la courbe retard-fréquence de l'ensemble, tout en déplaçant la courbe atténuation-fréquence d'une façon appréciable vers les fréquences basses.

Les figures 12 à 15 représentent la transmission d'un échelon unité à travers diverses longueurs de la ligne corrigée dont la courbe retard fréquence est représentée sur les figures 7 et 8.

Les principales caractéristiques des éléments de cette ligne sont les suivants :

— Coefficient de correction : $\alpha = \frac{C_2}{C_m} = 0,1$

— Coefficient de couplage entre deux bobines adjacentes :

$$K = \frac{M}{L} = 18\%$$

Terminaison des extrémités suivant figure 9a.

Les différentes vues ont été prises avec la même vitesse de balayage, en ajustant convenablement l'instant de déclenchement.

On peut remarquer que la durée de front ne varie pratiquement pas le long de la ligne.

Le retard total d'une ligne comportant 60 bobines est de 30 durées de front, soit un retard d'une durée de front pour deux bobines.

En fait, les coefficients α et K sont à ajuster dans chaque cas particulier au voisinage des valeurs qui ont été indiquées ci-dessus.

Dans le cas relativement extrême d'une structure comportant 12 bobines et terminée, à une extrémité suivant figure 7b, à l'autre extrémité suivant figure 7c, on a trouvé une réponse optimum pour

$$\alpha = 0,14 \text{ et } K = 20 \%$$

Les figures 16 et 17 présentent à titre de comparaison, la transmission à travers ligne non corrigée pour $K = 12 \%$.

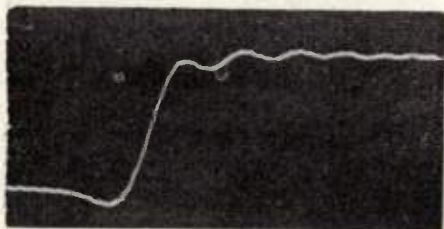


FIG. 16. — Transmission de l'échelon unité à travers 20 cellules d'une ligne non corrigée.

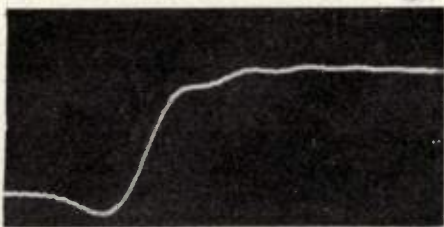


FIG. 17. — Transmission de l'échelon unité à travers 40 cellules d'une ligne non corrigée.

On remarquera que, si le front apparent est comparable à celui des figures 12 à 15, il est ici précédé et suivi d'une oscillation parasite dont l'amplitude croît avec la longueur de ligne considérée. Pour que deux fronts transmis successivement ne réagissent pas l'un sur l'autre, on devra donc les espacer en tenant compte de ces oscillations.

4. Choix de L et C_m .

Pour calculer les éléments constitutifs d'une ligne dans un cas pratique, on dispose généralement des données :

τ = retard total

δ = durée de front à transmettre

Z_c = impédance caractéristique

Les éléments à déterminer sont :

n = nombre de bobines

C_m = (défini précédemment)

L = self inductance d'une bobine

On peut prendre $n = 2 \frac{\tau}{\delta}$, et pour une ligne suffisamment longue, on trouve :

$$C_m = \frac{\tau}{nZ_c}$$

$$L = \frac{\tau Z_c}{\beta n}$$

où β est un coefficient, fonction des couplages entre bobines et généralement compris entre 1,4 et 1,5.

Si l'on veut tenir compte de l'influence des extrémités on a, pour définir les capacités, une relation plus exacte :

$$\Sigma C_s = \frac{\tau}{Z_c}$$

où l'on désigne par ΣC_s la somme de tous les condensateurs qui ne sont pas des condensateurs de correction C_2 .

CONCLUSION.

D'après ce qui vient d'être vu et, notamment, d'après les figures 7 et 8, il apparaît avec évidence

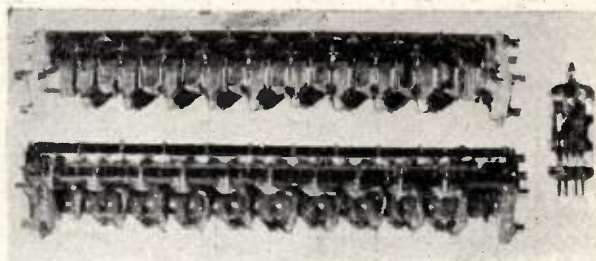


FIG. 18. — Aspect d'une ligne corrigée.

Les condensateurs à la masse sont sous la plaquette support. Les condensateurs de correction leur sont diamétralement opposés par rapport aux bobines.

que dans une ligne à retard, il est intéressant de placer des condensateurs de correction lorsqu'on utilise des bobines en disposition coaxiale.

Cette disposition se prête bien à une fabrication en série et plusieurs milliers de lignes de divers types ont été construits sur ce principe pour équiper des calculatrices électroniques dans lesquelles elles assurent la fonction de mémoire et délivrent des retards opératoires.

BIBLIOGRAPHIE

- P. DAVID. Les filtres électriques, généralités. Paris, Gauthier Villars 1952.
 M.J.E. GOLAY. The Ideal Low-Pass Filter in the form of a Dispersionless Lag Line. *Proceedings of the I.R.E. and Waves and Electrons.*, march 1946.
 H.E. KALLMANN. Equalized Delay Lines. *Proceedings of the I.R.E. and Waves and Electrons.* septembre 1946

NOUVELLES CHAINES DE BASCULEURS UTILISÉES DANS LES MACHINES A CALCULER POUR LE COMPTAGE A BASE 10 ET A BASE 12

PAR

J. J. BRUZAC
Compagnie I.B.M. — France

1°) GÉNÉRALITÉS

Nombre de dispositifs de comptage d'impulsions utilisent des chaînes de basculeurs montés en diviseurs de fréquence.

Généralement ces chaînes portent le nom de décades du fait qu'elles servent à compter des nombres à base 10.

c'est-à-dire de ramener à l'état normal tous les basculeurs lors de l'enregistrement de la dixième impulsion (Fig. 2).

Ce dispositif est constitué par une triode de blocage dont la grille reçoit une impulsion provenant du basculeur 8 au moment où ce dernier reçoit la dixième impulsion.

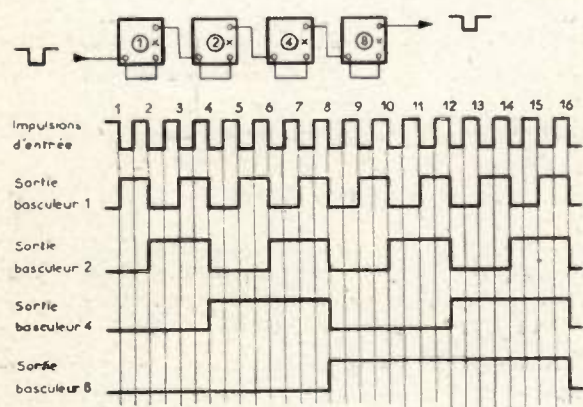


FIG. 1. — Chaîne de 4 basculeurs montés en diviseurs de fréquence

L'une des décades la plus connue est celle de Potter et ses dérivées utilisant une chaîne de 4 basculeurs et un tube de blocage.

Les quatre basculeurs sont appelés 1, 2, 4, 8 suivant la valeur propre d'enregistrement de chacun d'eux.

La chaîne de basculeurs représentée figure 1 et ne comportant pas de tube de blocage, permet le comptage à base 16. En effet, comme le montre le diagramme des tensions des sorties des basculeurs de la figure 1, le nombre de combinaisons des états de quatre basculeurs est de $2^4 = 16$.

Les croix sur les figures indiquent les triodes conductrices lorsque les basculeurs sont à l'état normal.

Pour constituer une décade, il faut donc adjoindre à cette chaîne un dispositif constitué par un tube de blocage permettant d'arrêter le comptage à 9,

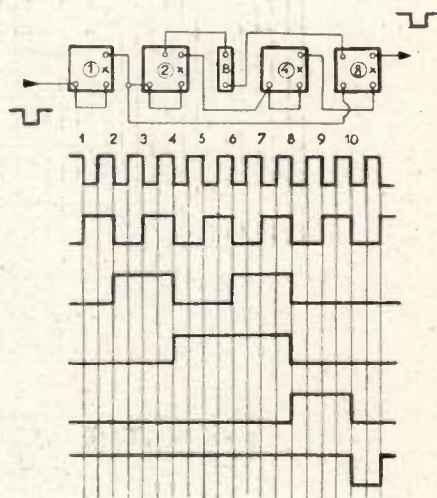


FIG. 2. — Décade du type 1-2-4-8.

La plaque de ce tube, étant reliée directement à la résistance de charge de la triode conductrice à l'état normal du basculeur 2, abaisse la tension plaque de cette dernière triode, ce qui empêche la mutation du basculeur 2.

Le retour à l'état normal du basculeur 8 produit une impulsion de report enregistrable dans une chaîne similaire représentant l'ordre décimal immédiatement supérieur soit directement, soit par l'intermédiaire d'un basculeur.

Il est en effet nécessaire, dans les machines du type parallèle dans lesquelles l'entrée de tous les chiffres composant un nombre se fait simultanément dans des décades séparées, de différer l'impulsion de report.

Cette impulsion peut se produire à n'importe quel moment de l'enregistrement, donc pendant qu'un

train d'impulsions est enregistré dans la décade devant recevoir le report. Elle est par conséquent enregistrée dans un basculeur, qui, remis à l'état normal après le train d'impulsions de comptage délivre lui-même une impulsion qui effectue le report.

En partant de la décade de Potter, il est possible d'obtenir des chaînes pour le comptage à bases 10 et 12 et leur donner une possibilité nouvelle très importante qui sera traitée plus loin.

2^o COMPTAGE A BASE 10.

2.1. Décade du type 1-2-2-4 (fig. 3).

En intervertissant l'attaque des grilles du basculeur 8 d'une décade de Potter, on obtient une décade dans laquelle on provoque la mutation du basculeur 8 à la deuxième impulsion, tout en empêchant par

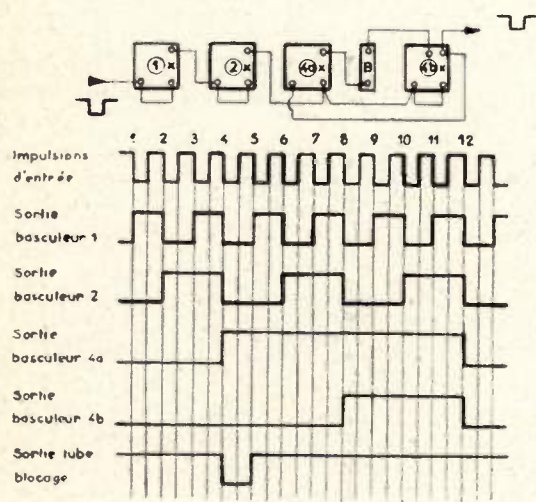


FIG. 3. — Décade du type 1-2-4-8 avec attaque inversée du basculeur 8.

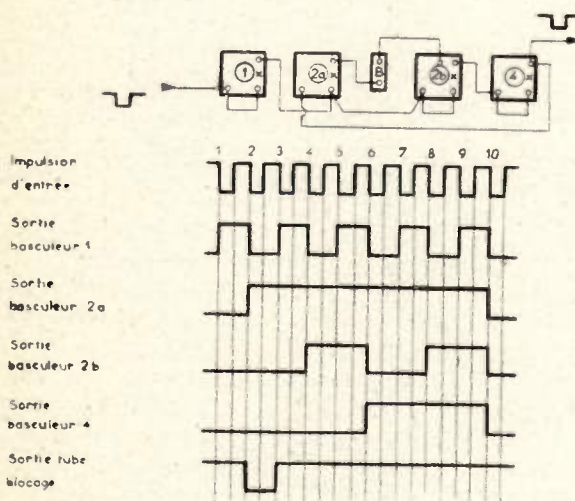


FIG. 4. — Représentation logique d'une décade du type 1-2-2-4 (1-2-4-8 modifié).

le tube de blocage, toute action sur le basculeur 2. Le basculeur 8 inversé (8i) reste à l'état basculé jusqu'à la dixième impulsion.

A partir de la quatrième impulsion, le basculeur reprend sa fonction normale.

La figure 4 montre la représentation logique de cette décade que nous pourrions appeler 1-2-2-4 suivant la valeur propre d'enregistrement de chacun des basculeurs le constituant.

2.2. Décade du type 1-1-2-5 (fig. 5).

La décade 1-2-2-4 de la figure 4 est représentée avec son tube de blocage entre le deuxième et troisième basculeur, le troisième basculeur contrôlant la mutation du deuxième.

En mettant ces deux basculeurs et leur tube de blocage en tête de la chaîne on obtient une décade

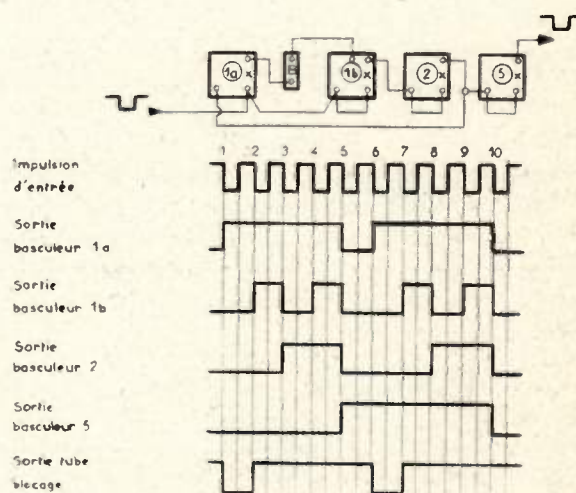


FIG. 5. — Décade du type 1-1-2-5.

dans laquelle la première impulsion change l'état du basculeur 1-a, ce dernier empêchant par le tube de blocage la mutation du basculeur 1-b.

Par contre la deuxième impulsion, sans effet sur le basculeur 1-a, provoque la mutation du basculeur 1-b. Les basculeurs 1-b, 2 et 5 sont montés normalement en diviseurs de fréquence.

La cinquième impulsion, tout en provoquant la mutation du basculeur 5, remet à l'état normal le basculeur 1-a et la chaîne, sauf pour le basculeur 5 à l'état excité, se retrouve dans les mêmes conditions qu'au moment de l'enregistrement de la première impulsion.

3^o COMPTAGE A BASE 12.

3.1. Duodécade du type 1-2-4-4 (fig. 6).

Le tube de blocage de la chaîne 1-2-2-4 étant inséré entre les troisième et quatrième basculeurs la chaîne obtenue peut avoir 12 positions distinctes.

En effet, les basculeurs 1-2 et 4-a étant montés en diviseurs de fréquence, la quatrième impulsion provoque la mutation du basculeur 4-a et, ce dernier,

par l'intermédiaire du tube de blocage, empêche la mutation du basculeur 4-b.

A la huitième impulsion, par contre, le basculeur 4-b passe à l'état excité, le basculeur 4-a restant à l'état excité.

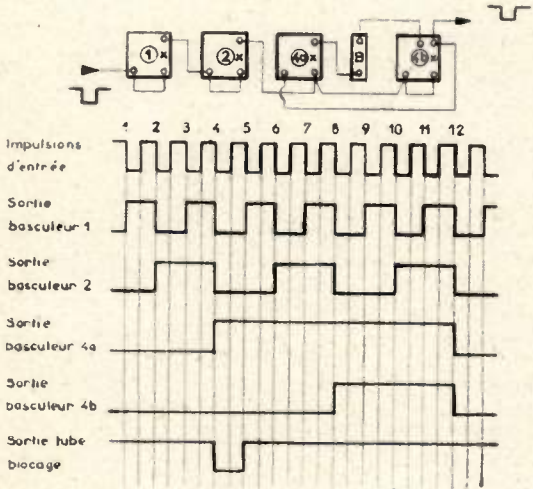


FIG. 6. — Duodécade du type 1-2-4-4-

Ce n'est qu'à la douzième impulsion que le basculeur 4-b, revenant à son état normal, remet à l'état normal le basculeur 4-a.

3.2. Duodécade du type 1-2-2-6 (fig. 7).

Cette duodécade ne diffère de la décade 1-2-2-4 que par l'attaque de la grille gauche du basculeur 2-a. En effet, cette grille étant attaquée par la sortie du basculeur 2-b au lieu de la sortie du basculeur 4, c'est la sixième impulsion qui provoque la mutation des basculeurs 2-a et 6.

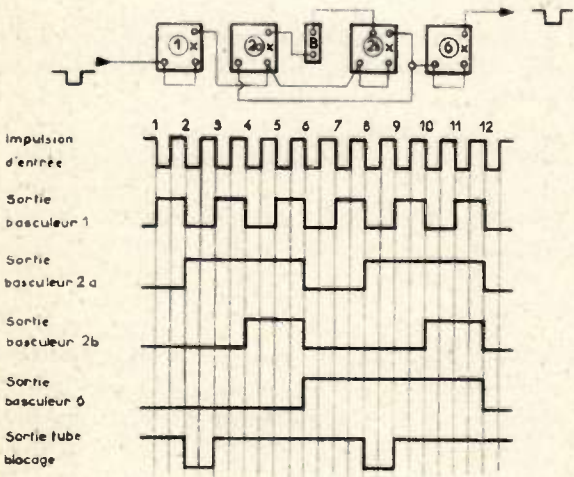


FIG. 7. — Duodécade du type 1-2-2-6.

A partir de la septième impulsion les basculeurs 1, 2-a et 2-b se retrouvent dans le même état que lors de l'enregistrement de la première impulsion.

10) AUTRES CHAINES DE COMPTAGE.

En modifiant l'interconnexion des 4 basculeurs de la décade de Potter on peut obtenir d'autres types de chaînes comme par exemple la duodécade 1-1-3-6, les chaînes 1-1-3-3 et 1-1-2-4 pour le comptage à base 9, etc.

50) FORMATION DES COMPLÉMENTS.

NOMBRE D'IMPULSIONS ENREGISTRÉES	BASCULEURS			
	1	2 a	2 b	4
0	0	0	0	0
1	1	0	0	0
2	0	1	0	0
3	1	1	0	0
4	0	1	1	0
5	1	1	1	0
6	0	1	0	1
7	1	1	0	1
8	0	1	1	1
9	1	1	1	1

FIG. 8. — Chaîne type 1-2-2-4.

NOMBRE D'IMPULSIONS ENREGISTRÉES	BASCULEURS			
	1 a	1 b	2	5
0	0	0	0	0
1	1	0	0	0
2	1	1	0	0
3	1	0	1	0
4	1	1	1	0
5	0	0	0	1
6	1	0	0	1
7	1	1	0	1
8	1	0	1	1
9	1	1	1	1

FIG. 9. — Chaîne type 1-1-2-5.

NOMBRE D'IMPULSIONS ENREGISTRÉES	BASCULEURS			
	1	1	4 a	4 b
0	0	0	0	0
1	1	0	0	0
2	0	1	0	0
3	1	1	0	0
4	0	0	1	0
5	1	0	1	0
6	0	1	1	0
7	1	1	1	0
8	0	0	1	1
9	1	0	1	1
10	0	1	1	1
11	1	1	1	1

FIG. 10. — Chaîne type 1-2-4-4-

NOMBRE D'IMPULSIONS ENREGISTRÉES	BASCULEURS			
	1	2 a	2 b	6
0	0	0	0	0
1	1	0	0	0
2	0	1	0	0
3	1	1	0	0
4	0	1	1	0
5	1	1	1	0
6	0	0	0	1
7	1	0	0	1
8	0	1	0	1
9	1	1	0	1
10	0	1	1	1
11	1	1	1	1

FIG. 11. — Chaîne type 1-2-2-6.

Désignons par 0 l'état normal d'un basculeur et par 1 son état excité. Les figures 8, 9, 10, et 11 montrent les états des différents basculeurs des chaînes décrites ci-dessus en fonction du nombre d'impulsions enregistrées.

Si nous considérons (Fig. 12) une décade du type 1-2-2-4 ayant enregistré 5, nous voyons que les basculeurs 1, 2-a et 2-b sont à l'état excité, le basculeur 4 étant à l'état normal.

Le complément de 5 étant 4 en base 10, nous constatons que pour 4, le basculeur 4 est à l'état excité, les basculeurs 1, 2-a et 2-b étant à l'état normal.

Pour obtenir le complément d'un chiffre enregistré dans une décade 1-2-2-4, il suffit donc d'inverser l'état de chaque basculeur de la chaîne.

De même dans le cas d'une duodécade du type 1-2-4-4, si le chiffre 6 y est enregistré, les basculeurs 2 et 4-b sont à l'état excité. Le complément en base 12 de 6 étant 5 et l'enregistrement de 5 étant caractérisé par l'état excité des basculeurs 1 et 4-a, l'inversion de l'état de chaque basculeur de la chaîne suffit pour obtenir ce complément.

Il en est de même pour toutes les chaînes décrites ci-dessus.

ETAT DU BASCULEUR POUR :	BASCULEURS			
	1	2 a	2 b	4
Chiffre en clair	1	1	1	0
Chiffre en complément	0	0	0	1

ETAT DU BASCULEUR POUR :	BASCULEURS			
	1	2	4 a	4 b
Chiffre en clair	0	1	0	1
Chiffre en complément	1	0	1	0

FIG. 12. — Conversion clair, complément

Par ce procédé, la soustraction est très simplifiée et se résume à additionner le nombre à soustraire au complément du premier nombre.

Prenons par exemple l'opération :

A-B-C-D-

dans laquelle $Q = 342$ $B = 29$ $C = 13$ $D = 193$.

Le complément à 324 étant 675, on ajoute successivement à ce dernier les nombres 29-13 et 193.

$$\begin{array}{r}
 +675 \\
 29 \\
 \hline
 +704 \\
 13 \\
 \hline
 +717 \\
 193 \\
 \hline
 910
 \end{array}$$

A la fin de l'opération, en formant le complément du résultat obtenu, nous avons le résultat en clair soit 89.

Il est évident que dans une suite d'opérations comportant des additions et des soustractions, il est nécessaire de former le complément du résultat à chaque changement de signe.

Prenons par exemple l'opération :

A-B-C-D-E

où $A = 1528$ $B = 128$ $C = 53$ $D = 21$ $E = 264$

Complément de 1528 = 8471.

$$\begin{array}{r}
 8471 \\
 + 128 \\
 \hline
 8599 \\
 + 53 \\
 \hline
 8652
 \end{array}$$

Complément de 8652 = 1347

$$\begin{array}{r}
 1347 \\
 + 21 \\
 \hline
 1368 \\
 + 264 \\
 \hline
 1632
 \end{array}$$

La formation du complément n'a pas lieu d'être effectué si la dernière opération est une addition.

La mutation de tous les basculeurs d'une chaîne est provoquée soit par une très brève impulsion positive sur la cathode ou négative sur les plaques.

Cette mutation étant très rapide ne présente qu'un inconvénient restreint, si on le compare à l'avantage obtenu dans la facilité du déroulement des calculs.

6°) CONTROLE DES NOMBRES ENREGISTRÉS.

Sur toutes les chaînes décrites ci-dessus, on peut monter, comme sur la décade de Potter, des indicateurs constitués par des lampes au néon qui indiquent l'état de chaque basculeur des chaînes.

7° TUBES DE BLOCAGE.

Les dispositifs de blocage, décrits comme étant constitués par des triodes, peuvent être remplacés par un classique interrupteur à diode, dont les circuits sont incorporés dans le câblage, permettant une économie appréciable de place.

PRODUCTION DE SIGNAUX DE DÉBLOCAGE DANS UNE MACHINE A CALCULER DÉCIMALE PAR UN CIRCUIT A NOMBRE DE TUBES RÉDUIT

PAR
G. BOYER

Compagnie I.B.M. France

Les machines à calculer électroniques sont généralement équipées d'un grand nombre de tubes qui représentent une part importante du prix total.

La recherche permanente de prix de revient plus bas conduit très souvent à des circuits économiques présentant cependant une sécurité de fonctionnement parfaite, condition indispensable dans ce domaine.

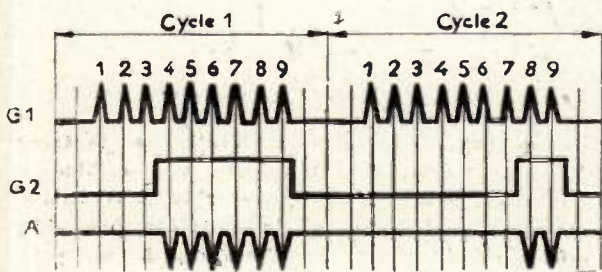


FIG. 1.

Le dispositif décrit ci-dessous illustre cette évolution. Il utilise environ le tiers des tubes qui seraient nécessaires pour obtenir les mêmes résultats par des méthodes plus classiques.

Dans un calculateur électronique toute opération est composée d'un certain nombre de périodes élémentaires que l'on peut appeler « cycles ». Durant chaque cycle d'une machine travaillant en système décimal, chaque compteur peut recevoir un nombre d'impulsions variable entre 0 et 9 et il est nécessaire d'interposer entre les compteurs et la source d'impulsions un élément capable de n'en laisser passer que le nombre voulu : ceci peut être fait d'une façon très classique en utilisant un tube à deux électrodes de commande — une pentode par exemple.

A chaque cycle 9 impulsions sont appliquées à l'une des grilles, l'autre reçoit un signal « gate » dont la durée détermine le nombre d'impulsions obtenues sur l'anode.

Dans l'exemple de la figure 1, le compteur recevra 6 impulsions au cours du premier cycle et 2 au cours du deuxième.

Les différents compteurs pouvant recevoir des chiffres différents au cours du même cycle, il est

utile de disposer des 9 signaux représentés figure 2. Il peut également arriver que l'un d'eux soit utilisé par la totalité des compteurs (si tous doivent recevoir le même chiffre) et sur certaines machines où

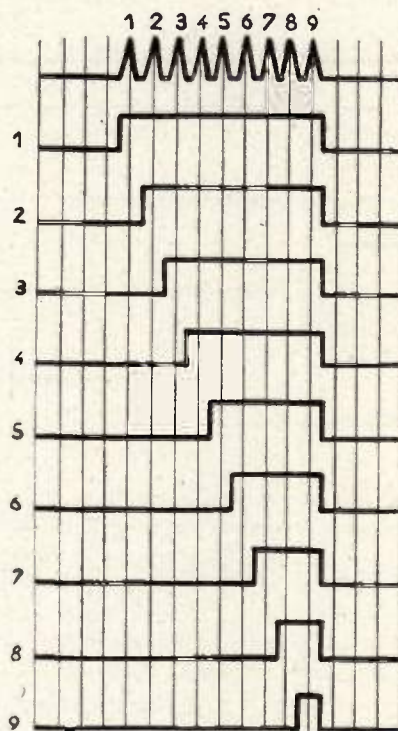


FIG. 2.

les connexions sont longues cela impose d'obtenir ces signaux sous faible impédance.

Deux circuits permettant de résoudre le problème posé vont être décrits ; les signaux « gate » y sont obtenus par des basculeurs déclenchés successivement et ramenés ensemble à leur état primitif (restauration).

Il est indispensable de les faire suivre de « cathodes followers ». Tous deux conduisent sensiblement au même nombre de tubes.

1° 9 basculeurs $B_1, B_2, \dots, B_9$ montés comme l'indique la fig. 3 reçoivent d'un côté les 9 impulsions

marquant le début des « gate » et se commandent en chaîne par l'autre entrée.

Supposons les tous dans l'état où tendent à les mettre les 9 impulsions, sauf B_1 .

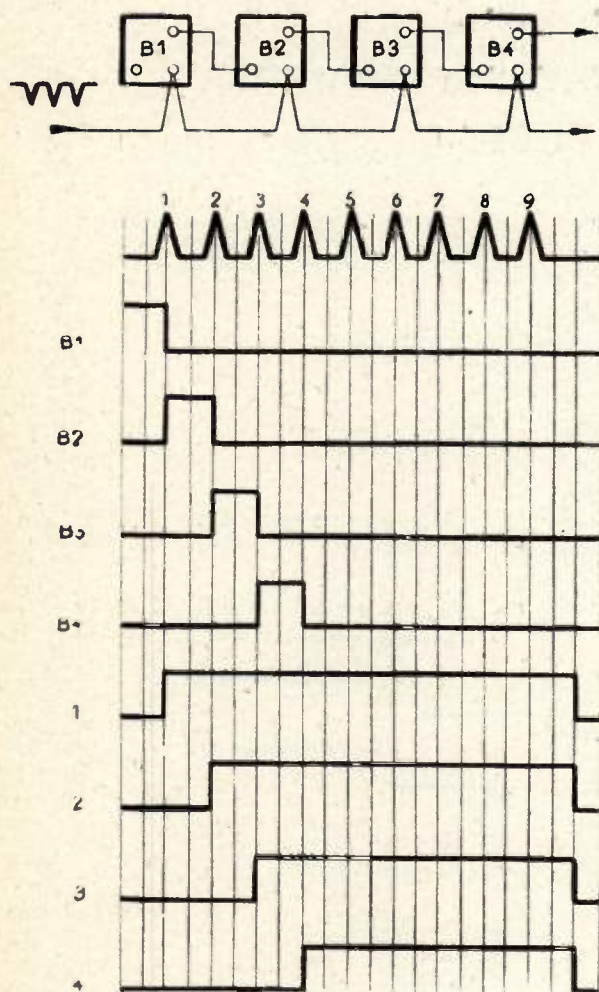


FIG. 3.

Le premier top ne peut que basculer B_1 ; les tops suivants sont sans action. B_1 en fonctionnant entraîne B_2 que le top 2 ramène à son état primitif, etc...

On obtient ainsi les signaux de la figure 3 dont les flancs négatifs peuvent déclencher les 9 basculeurs donnant les signaux désirés ; le dispositif est complété par 9 « cathodes followers ».

Un basculeur étant couramment réalisé avec un tube double triode, cet ensemble demande 27 tubes sans tenir compte des circuits auxiliaires.

2° Il est possible également de n'utiliser que les 9 basculeurs délivrant les signaux recherchés à condition d'insérer entre chacun d'eux un système de « gate » provoquant chaque déclenchement avec une impulsion de retard (fig. 4).

Un tube à deux grilles de contrôle commande chaque basculeur et reçoit d'une part les impulsions 1, 3, 5, 7, 9 ou 2, 4, 6, 8 et d'autre part le signal délivré par le basculeur précédent.

La première impulsion recueillie sur l'anode de G_3 , par exemple, est bien l'impulsion 3 qui doit déclencher le basculeur B_3 .

La division des impulsions d'attaque en paires et

impaires n'est pas indispensable. On peut insérer dans la liaison entre un basculeur et la grille qu'il commande, une constante de temps qui empêche le déblocage à ce moment même, mais l'autorise pour l'impulsion suivante.

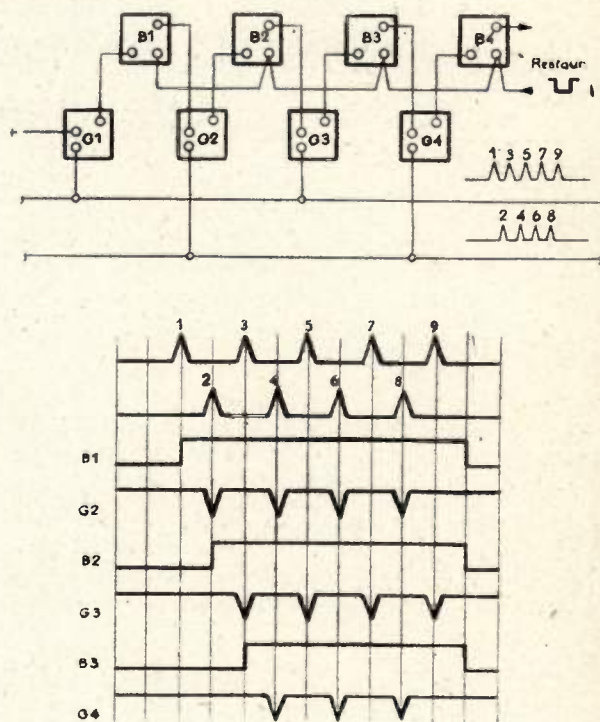


FIG. 4.

Bien qu'apparemment plus simple, ce système conduit au même nombre de tubes que le premier.

Le montage à thyatron qui va être décrit maintenant s'inspire directement de ce dernier procédé ; les tubes sont des 2D.21 — tube miniature qui possède deux électrodes de commande.

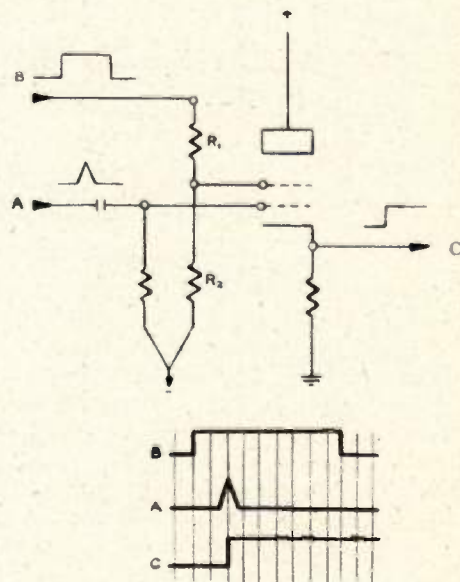


FIG. 5.

Soit l'un d'eux monté suivant le schéma de la figure 5 la charge est disposée entre masse et cathode

et par conséquent, une tension positive existe sur cette électrode quand le thyatron est conducteur.

Pour obtenir cette conduction, il faut deux conditions simultanées :

- une impulsion positive en A
- une tension positive en B .

Le pont R_1, R_2 est calculé pour que la tension nécessaire en B soit égale à la tension cathodique pendant la conduction.

Le dispositif comprend 9 éléments identiques disposés en chaîne suivant le schéma de la figure 6.

T_1, T_3, T_5, T_7, T_9 reçoivent les impulsions 1, 3, 5, 7, 9
 T_2, T_4, T_6, T_8 reçoivent les impulsions 2, 4, 6, 8.

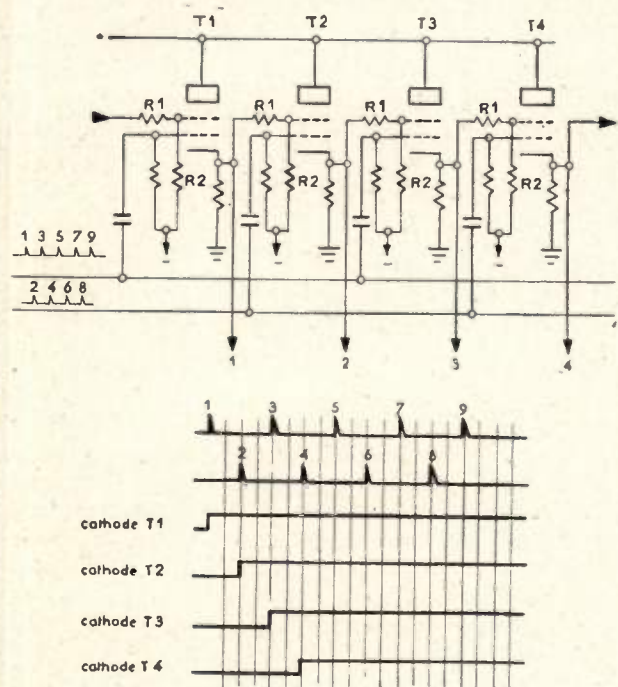


FIG. 6.

D'autre part, chacun des ponts diviseurs fixant le potentiel de la deuxième électrode de commande reçoit la tension cathodique du thyatron précédent, le pont de T_1 étant supposé initialement à un potentiel positif suffisant.

A la première impulsion T_1 est le seul thyatron qui peut s'amorcer puisque les deuxièmes grilles de tous les autres sont bloquées, mais l'amorçage de T_1 prépare celui de T_2 qui a lieu à l'impulsion suivante.

La conduction de T_2 commence donc à 2 et prépare celle de T_3 etc...

Comme dans le procédé précédent la division des impulsions en paires et en impaires n'est pas indispensable à condition de retarder le moment où le potentiel de la deuxième électrode de commande permet l'amorçage du tube par une constante de temps.

Le 2D.21 pouvant débiter un courant moyen de 100 mA il est possible de choisir une valeur de résistance assez faible comme charge cathodique ; les signaux « gate » désirés sont donc obtenus directement à faible impédance sur les cathodes de T_1 ,

$T_2 \dots T_9$. En fait la tension anode-cathode des thyatrons étant constante quelle que soit l'intensité débitée, la charge cathodique peut varier dans de grandes proportions sans que cela modifie l'amplitude du signal.

Avec les montages à basculeurs décrits plus haut, il n'y avait pas de difficultés à rétablir l'état primitif après les 9 impulsions, chacun d'eux n'exigeant pour cela qu'une énergie minime. Par contre, le désamorçage des 9 thyatrons débitant sur des charges choisies volontairement assez faibles pose un problème beaucoup plus complexe puisqu'il ne peut être question de couper mécaniquement le circuit anodique.

On peut désamorcer un thyatron par une impulsion négative sur la plaque ou positive sur la cathode, suivant que la charge est anodique ou cathodique ; mais il faut que cette impulsion fournisse au circuit une énergie supérieure à celle qui est contrôlée par le thyatron pendant le temps nécessaire à la déionisation.

En effet, pendant sa durée cette impulsion représente une source d'énergie en parallèle sur la source normale. Si elle est suffisante, l'intensité traversant le tube tombe à 0, ce qui provoque la déionisation aussi sûrement qu'une coupure mécanique du circuit.

Une telle impulsion est difficile à obtenir d'un tube à vide, mais peut être aisément produite par un thyatron du même type (le 2D.21 peut accepter des pointes de courant considérables).

Le montage de la figure 7 rappelle le dispositif utilisé dans les relaxateurs utilisés — de moins en

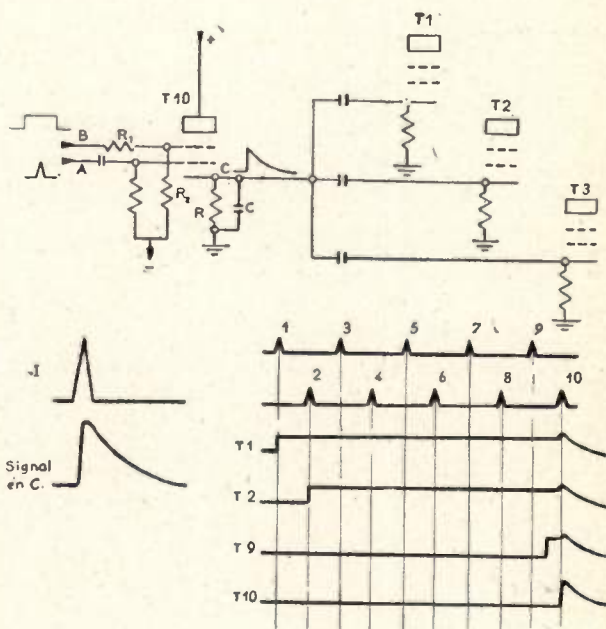


FIG. 7.

moins, il est vrai — pour la production des tensions de balayage des oscilloscopes.

Il suffit d'ailleurs de supprimer la polarisation négative de grille pour que la relaxation, se produise mais ici la courbe observée est l'exponentielle de décharge de C presque entière. En relâché ou en déclen-

ché, le tube se désamorce immédiatement après la conduction si le produit RC est suffisamment élevé.

Il s'agit ici d'un élément T_{10} dont les circuits de commande sont semblables à ceux des autres éléments $T_1, T_2 \dots T_9$ et qui se monte à la suite de la chaîne décrite plus haut. L'amorçage se produit de la même manière à l'impulsion 10 (rajoutée au groupe 2, 4, 6, 8).

Le désamorçage est immédiat et provoque celui de $T_1, T_2 \dots T_9$, grâce à un couplage par condensateurs entre la cathode de T_{10} et chacune des autres cathodes : la constante de temps de cette liaison doit être supérieure au temps de déionisation qui est de l'ordre de $50 \mu s$. En réalité, certaines précautions doivent être prises : C étant considérable, une valeur trop forte des condensateurs de liaison peut provoquer l'extinction de $T_1, T_2 \dots T_9$ juste après leur amorçage.

Des réactions ont lieu par l'intermédiaire des condensateurs de liaison mais sont rendues négligeables par la valeur élevée de C .

La vitesse de fonctionnement de ce circuit est limitée par la valeur minimum que l'on peut donner à RC . En effet, la tension de cathode de T_{10} revient à 0 par une exponentielle de constante de temps RC . Tant qu'elle est sensiblement positive un réamorçage est difficile et d'autre part, l'impulsion obtenue est réduite et ne suffit plus pour déioniser $T_1, T_2 \dots T_9$.

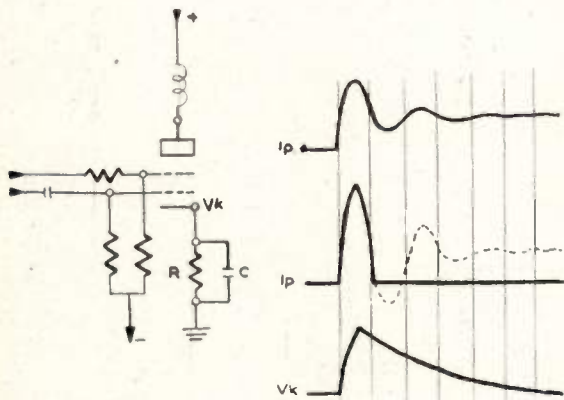


FIG. 8.

Au dessous de $RC = 1 \text{ ms}$, le désamorçage de T_{10} devient défectueux.

Il est cependant possible d'obtenir une extinction satisfaisante de T_{10} au dessous de cette valeur en insérant une self dans le circuit anodique (figure 9). Juste après la conduction, on peut admettre que L et C forment un circuit oscillant qui vient d'être excité et qui, par conséquent résonne. Si la surtension du circuit est insuffisante, on observe des oscillations amorties et le tube reste conducteur car l'intensité ne devient jamais nulle.

Si la surtension est bonne, la déionisation a lieu au cours de la première période dès que l'intensité a tendance à s'inverser.

Ce procédé est intéressant, puisqu'il permet à peu de frais, une augmentation de vitesse sensible mais

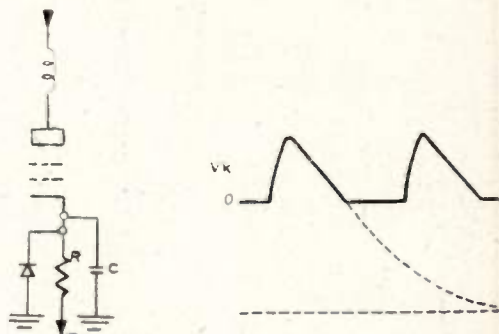


FIG. 9.

on est vite limité dans son application. En effet, il a été dit que la valeur minimum des condensateurs de liaison était fixée par le temps de déionisation. Ces condensateurs déterminent à leur tour la valeur de C qui doit être de valeur beaucoup plus élevée pour éviter les réactions. Le seul moyen de diminuer RC est donc de diminuer R .

Or les deux conditions R faible et bonne surtension sont contradictoires.

Pratiquement un compromis est possible, la self L devant évidemment être très peu résistante.

Enfin pour augmenter encore la vitesse, le circuit cathodique peut être modifié comme l'indique la figure 9. Après désamorçage, la cathode a tendance à atteindre la tension de l'alimentation négative et atteint pratiquement le potentiel 0 par une droite au lieu d'une exponentielle. Un redresseur l'empêche de descendre au-dessous, dès ce moment, un nouvel amorçage est possible.

LE SIGNAL MINIMUM UTILISABLE EN RECEPTION RADAR ET SON AMELIORATION PAR CERTAINS PROCÉDÉS DE CORRELATION

PAR

L. GERARDIN

Compagnie Française Thomson-Houston

INTRODUCTION

La réception des signaux radar, particulièrement ceux provenant d'échos isolés, se présente comme un cas particulier de la réception d'impulsions, à ceci près que la conservation de la forme des signaux a peu d'importance. Ce problème est, en dernière analyse, dominé par la question de rapport signal/bruit. Celle-ci prend, dans ce cas particulier, une importance toute spéciale.

En effet, dans un système de télécommunication, on se garde toujours une certaine marge de sécurité en rapport signal minimum reçu /bruit à la réception. Le radar, lui, fonctionne toujours à la limite ultime possible de détectabilité des signaux reçus, puisque sa portée est conditionnée étroitement par le signal minimum utilisable par l'observateur.

Il est donc nécessaire de définir très exactement ce signal, de façon à pouvoir, par exemple, effectuer des prévisions de portée d'appareils nouveaux. Les procédés, quels qu'ils soient, qui permettront la diminution de la puissance de ce signal nécessaire, se présentent, à priori, comme très intéressants.

PARTIE I

SIGNAL MINIMUM UTILISABLE

I. — CRITÈRE DE VISIBILITÉ ADOPTÉ :

Dans le cas de la détection d'échos lointains, dont on ignore à priori les coordonnées, on utilise pratiquement toujours une présentation P.P.I. (ou analogues, « B » par exemple), avec modulation d'intensité d'un tube cathodique.

Il est donc logique d'admettre comme critère de visibilité le rapport d'amplitude du signal à l'amplitude de bruit :

$$\frac{\text{amplitude signal}}{\text{amplitude bruit}} = k \quad (1)$$

Il s'agit de déterminer par le calcul la valeur de k nécessaire pour avoir un signal minimum utilisable, et de confronter avec l'expérience les résultats obtenus.

Notons tout de suite que pour des signaux d'impulsions mélangés à du bruit, (1) sera employée sous la forme :

$$\frac{\text{valeur moyenne signal} + \text{variance signal}}{\text{variance bruit (en l'absence de signal)}} = k \quad (2)$$

définition qui a déjà été employée dans un certain nombre de publications sur ce sujet.

II. — DÉFINITION DE L'OBSERVATEUR :

Avant d'entreprendre le calcul du rapport k en fonction des caractéristiques du signal, il faut définir de façon mathématique l'opérateur radar.

Le problème de réception se présente à celui-ci comme un problème de choix simple : à tel endroit de l'espace (correspondant univoquement à un point du système de représentation), y a-t-il ou non un obstacle ? en langage statistique, l'opérateur essaye donc les deux hypothèses suivantes : du bruit seul est présent, ou il y a du bruit et un signal. Ne pouvant effectuer qu'un certain nombre d'observations, la décision de l'opérateur peut être affectée de l'une ou l'autre des erreurs suivantes : prendre du bruit pour un signal, alors que seul du bruit est présent, ou bien : considérer un signal comme du bruit, alors qu'il y a effectivement, un signal. Les probabilités respectives de ces erreurs sont notées : « p » et « q ». Lorsque l'on aura décrit la façon dont l'opérateur utilise les données pour faire son choix, le processus d'observation sera totalement défini analytiquement.

Trois types principaux d'opérateurs ont été utilisés [7], dits : NEYMAN-PEARSON, Idéal ou Séquentiel.

L'opérateur NEYMAN-PEARSON, défini comme suit : la probabilité « p » étant fixée, ainsi que le nombre d'observations « n », la probabilité « q » doit être

minima (ou vice-versa), semble correspondre assez exactement à l'opérateur radar usuel. La différence entre les trois types d'observateurs mathématiques reste, d'ailleurs, de l'ordre des erreurs de mesure possibles (deux décibels environ).

S'il est toujours possible de résoudre formellement le problème du calcul des variations, posé par l'exigence de rendre une probabilité minima, il est assez délicat d'aller jusqu'aux solutions numériques. On se contentera donc, dans ce qui suit, de se fixer $(1 - q)$ probabilité de succès, du signal étant présent, de le voir, quitte à calculer ensuite et à voir si p a une valeur suffisamment faible pour permettre une exploitation correcte sans fausses alarmes exagérées ; L'opérateur va travailler sur une différence de brillance (la forme n'influant pratiquement pas pour les échos lointains et de petites dimensions). La différence minima de brillance nécessaire pour discriminer un signal uniforme d'un fond uniforme a été déterminée par HOPKINSON [1]. Dans le cas d'un fond de bruit, il est nécessaire de considérer non la luminosité moyenne du fond, mais sa luminosité maxima, c'est-à-dire la brillance des plus brillants spots de bruit. La brillante moyenne de l'écho devra être « h » fois cette quantité. Les courbes d'HOPKINSON prises dans le cas d'un signal juste reconnaissable, permettent de donner la valeur 2 à ce coefficient « h ». Le choix de cette valeur, qui influe sur toute la suite des calculs, est évidemment, le point de la théorie le plus délicat et celui qui risque de présenter le plus d'arbitraire.

III. — PROBABILITÉ DE SUCCÈS $(1 - q)$:

Le choix de la probabilité de succès est assez arbitraire également. On désire, par exemple, le système P.P.I. ayant effectué un premier et un second tour mécaniquement, être quasi certain, à 1 % près, par exemple, qu'un spot apparu à la même place correspond bien à un signal. Ceci correspond donc à une probabilité de succès $(1 - q)$ de 0,9.

IV. — DÉFINITION DU TUBE CATHODIQUE.

Il reste à définir mathématiquement le tube cathodique. Le signal composite (bruit + signal) à la sortie du récepteur a une certaine distribution d'am-

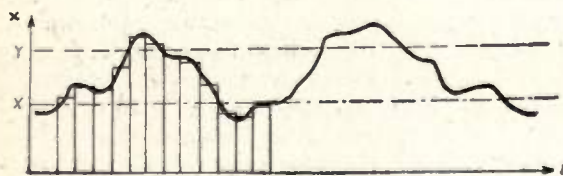


FIG. 1. — Signal d'attaque du P.P.I.

plitude x . Tout signal d'amplitude égale ou inférieure à X (niveau de polarisation, ne sera pas perçu. Tout signal d'amplitude égale ou supérieure à Y (niveau d'écrêtage), donnera la brillance maxima (fig. 1). En premières approximations, on supposera que les niveaux X et Y sont très voisins, c'est-à-dire que le tube fonctionne en tout ou rien. Le signal bruité peut, avec une bonne approximation, se décom-

poser, pour faciliter les calculs, en un certain nombre d'impulsions rectangulaires non corrélées entre elles. La largeur de ces impulsions est environ $1/2 B$, avec « B » bande passante du récepteur moyenne fréquence. Chaque spot lumineux comprend un certain nombre de telles impulsions élémentaires ; sa brillance sera la résultante de l'excitation de toutes impulsions, qu'elles se produisent en un cycle de balayage ou dans quelques cycles adjacents parmi les « n » considérés.

V. — CALCUL DE LA CONSTANTE « K ».

Les définitions précédentes une fois posées, il est possible de calculer la constante « k ». Le problème a été traité par A.W. Ross [6]. La procédure suivie est en résumé la suivante. Le type de détecteur étant déterminé (linéaire ou quadratique), on connaît les fonctions de distribution d'amplitude du bruit seul et d'un bruit mélangé à du signal continu.

S'étant fixé le niveau de polarisation, on connaît la probabilité pour qu'une impulsion élémentaire de bruit dépasse ce niveau et apparaisse sur l'écran. Chaque spot comprenant un certain nombre de telles impulsions élémentaires (nombre fonction des vitesses de balayages électriques et mécaniques et du diamètre du spot), on déduit la probabilité qu'à l'intensité optique d'un spot de bruit d'égalier ou de dépasser une certaine valeur.

On peut effectuer les mêmes calculs pour le signal bruité, et déterminer ainsi, connaissant « k » (critère de détection) et le rapport signal/bruit, la probabilité de détection du signal.

En répétant ces calculs, on dresse des tables de résultats qui se mettent plus pratiquement sous forme de courbes. Un exemple en est donné dans l'article cité pour un détecteur linéaire. On y voit que si l'on dispose de dix signaux il faut un rapport signal/bruit de 8,2 dB, pour assurer la probabilité de succès de 0,9. Une diminution de signal de 2 dB n'assure plus qu'une probabilité de 0,3.

VI. — EXPÉRIENCES EFFECTUÉES.

Ces résultats ont été confirmés par des expériences effectuées avec un radar 10 cm, et un pupitre

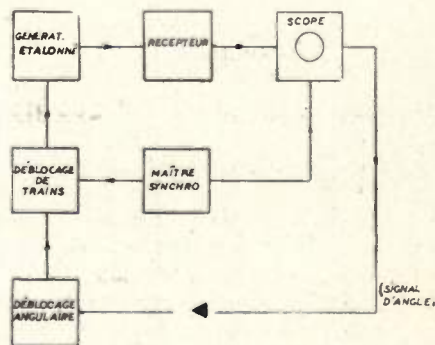


FIG. 2. — Dispositif de mesure du signal minimum utilisable.

P.P.I. à tube à 12 pouces. Le montage d'essais était le suivant (fig. 2). Le récepteur radar était attaqué par un générateur de signal étalonné (impulsions

hyperfréquences) synchronisé. Un maître synchronisateur déclenchait en permanence le balayage du pupitre d'observations et, par trains, le générateur. Ces trains comportaient de 30 à 4 impulsions. Ils se produisaient soit à un angle donné sur le scope, soit n'importe où, grâce à un dispositif de déblocage mécanique du système générateur de trains, lié au système de rotation du P.P.I.

Les impulsions utilisées étaient de une microseconde, la vitesse de rotation de 6 tours/minute, la durée des balayages de 25 ou 150 km. Cette dernière quantité causait une variation de 1 dB environ. On note une différence moyenne de 1 dB entre le cas où l'opérateur sait à peu près la direction dans laquelle va apparaître le signal (cas qui se présente pratiquement pour une calibration), et le cas où l'opérateur ignore totalement où va apparaître l'écho (cas d'une exploitation réelle).

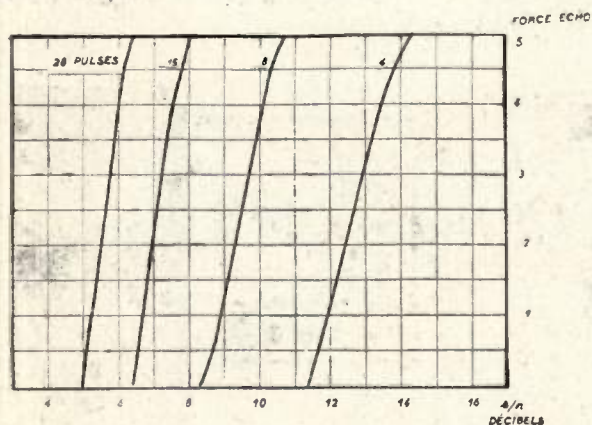


Fig. 3. — Courbes expérimentales du signal minimum utilisable.

La figure 3 résume les résultats de mesure sur balayage 150 km, dans le cas où l'on connaît la direction de l'écho. On voit la variation rapide de la force de l'écho pour moins de 2 dB de variation de signal (le détecteur n'était plus linéaire, mais quadratique). La non-adaptation de la largeur de l'impulsion au diamètre du spot causait une perte de 1 dB supplémentaire par rapport aux courbes publiées par A.W. Ross.

On constate la variation très rapide de la force du signal pour moins de 2 dB de variation de signal. Ceci est dû au fait que le détecteur n'était plus linéaire dans la région utilisée, mais presque quadratique.

VII. — CONCLUSIONS.

On peut déduire de tous ces résultats une conclusion très importante. En conditions normales, sur un radar P.P.I., le signal minimum utilisable est le signal tangentiel. Le signal minimum classique, mesuré en type « A » avec un signal permanent, n'a guère de signification, en particulier lorsque l'on veut calculer la portée en appliquant l'équation du radar.

De plus, puisqu'à une différence de puissance de signal de 2 dB correspond une forte variation de la probabilité de détection, la largeur de faisceau d'antenne utilisée pour le calcul du nombre d'impulsions sera la largeur à 1 dB, c'est-à-dire environ moitié de la largeur classique à 3 dB (avec une forme de diagramme de rayonnement gaussienne).

PARTIE II

AMÉLIORATION DU SIGNAL MINIMUM UTILISABLE

I. — POSITION DU PROBLÈME.

Le signal minimum utilisable ayant été ainsi bien défini (pratiquement le signal tangentiel), on peut chercher comment l'améliorer. Ce résultat s'obtient d'abord en diminuant le bruit, c'est-à-dire en réalisant des circuits donnant le minimum de facteur de bruit global. La limite dans cette voie est actuellement pratiquement atteinte.

Il faut alors chercher à exploiter toutes les caractéristiques connues du signal pour mieux le différencier du bruit de fond.

Un récepteur radar normal, de 2 Mc/s de bande par exemple, pour fixer les idées, peut amplifier toutes les impulsions de largeur supérieure à 0,5 microseconde, quelle que soit leur cadence de répétition.

Il y a là un véritable gaspillage de bande, puisqu'en fait, les impulsions qui nous intéressent sont identiques aux impulsions émises qui, en particulier, ont même durée et même fréquence de répétition. D'un autre point de vue, on peut dire que l'écho provenant d'un objectif défini (fixe ou mobile) a lieu, au cours des périodes successives, au même instant de la période, avec une configuration d'énergie renvoyée identique (tout au moins pour quelques périodes).

L'information d'instant d'apparition de l'écho peut être exploitée facilement, puisqu'il est possible de la mesurer de façon très précise, même après passage des signaux dans les circuits de réception. Le dispositif de mesure devra être doué d'une certaine « mémoire ».

Celle-ci peut être statique (c'est-à-dire obtenue par l'emploi d'un tube à mémoire) ou dynamique (avec une ligne à mercure). C'est ce dernier mode de réalisation qui a été choisi. Il offre l'avantage d'une facilité plus grande de reproductibilité du dispositif de mémoire.

II. — FILTRE LINÉAIRE — THÉORIE SOMMAIRE.

Le dispositif le plus simple à réaliser, pour exploiter la donnée temps privilégiée, consiste dans l'addition linéaire simple des signaux correspondants à des périodes successives.

On peut se rendre compte du gain procuré. Si l'on ajoute les signaux d'échos de n récurrences successives, ils s'additionnent en amplitude, et l'amplitude du signal résultant est n fois l'amplitude d'un des signaux, supposés tous égaux, (ce qui a toujours lieu, l'amplitude de fluctuation des échos, même mobiles, étant lente).

Au contraire, les bruits correspondants étant des signaux aléatoires indépendants ne s'ajoutent qu'en variance et l'amplitude du signal du bruit résultant n'est que $\sqrt{n}$ fois l'amplitude d'un des bruits.

Le gain en rapport signal/bruit est donc $\sqrt{n}$.

Ces considérations sont très élémentaires. En fait, il faut tenir compte des faits suivants : d'abord de la présence du second détecteur qui, aux faibles va-

leurs du rapport signal/bruit moyenne fréquence modifie ce rapport.

De plus, les signaux ne sont pas stationnaires ; l'antenne illumine l'objectif pendant un temps très court (correspondant à l'envoi de 10 ou 15 impulsions par exemple) le bruit, par contre, est présent en permanence. Il est donc nécessaire que la somme des bruits tende rapidement vers une limite, donc que la mémoire du système soit finie.

Enfin, en liaison très étroite avec la durée de la mémoire, on a une constante de temps d'établissement du signal dans le dispositif d'addition. Les signaux en sortie n'auront pas tous la même amplitude et si l'on observe le résultat sur un écran type P.P.I., par exemple, on aura encore une diminution de l'efficacité, puisque l'on a vu que le P.P.I. est très sensible aux variations de signal autour du signal minimum détectable.

La théorie complète en vidéo avec nombre de périodes de bruit égal au nombre de périodes de signal, a déjà été faite en utilisant pour le rapport signal/bruit le critère de DWORK, et les théories classiques de prédiction de WIENER [11]. On détermine ainsi la fonction de transfert optimale à donner au dispositif linéaire d'amélioration.

$$H(p) = \frac{Cte}{p} (1 - e^{-p\tau}) \frac{1 - e^{-npT}}{1 - e^{-pT}} \quad (3)$$

avec τ durée des impulsions de signal, et T récurrence des impulsions.

Ce filtre peut se réaliser avec une ligne à mercure de retard T , bouclée sur elle-même par un circuit de réaction de gain égal à 1, en série avec un circuit de soustraction analogue à celui d'un radar à élimination d'échos fixes de retard nT .

On remarque de suite qu'il faudra asservir deux lignes à mercure l'une sur l'autre, de plus l'une (retard nT) est de réalisation mécanique presque impossible, et l'autre doit fonctionner avec un gain de réaction de 1, c'est-à-dire à la limite ultime de l'accrochage.

Pratiquement, on supprime le circuit de soustraction, et on rend la mémoire finie en diminuant le gain « g » de boucle. La valeur optimum de celui-ci est définie par le nombre d'impulsions avec lequel on travaille. Il est voisin de 0,8 ou 0,9, en pratique. On a, dans ce cas, comme fonction de transfert du système :

$$H'(p) = \frac{1}{1 - ge^{-pT}} \quad (4)$$

et l'amélioration possible en type « A » est de 4 à 5 db

III. — RÉALISATION DU FILTRE LINÉAIRE.

Fondamentalement, le dispositif comprend (fig. 4) : un ensemble d'addition, une ligne à mercure et un circuit de soustraction. Bien entendu, il y a de nombreuses annexes : la ligne à mercure fonctionnant sur une certaine fréquence porteuse (10 MHz dans le cas actuel), nécessite un modulateur qui transforme les signaux vidéo en signaux moyenne fré-

quence modulée. L'atténuation de la ligne est compensée par un amplificateur. Enfin, le radar est piloté par la ligne elle-même, d'où circuits spéciaux à cet effet.

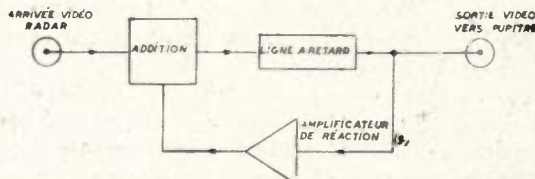


FIG. 4. — Principe du filtre linéaire à mémoire finie.

L'appareil se présente en une baie (fig. 5). De bas en haut, on distingue : l'alimentation, le panneau intégrateur proprement dit le panneau de commandes et mesures, et un panneau annexe de générateur et synchroscope de contrôle.

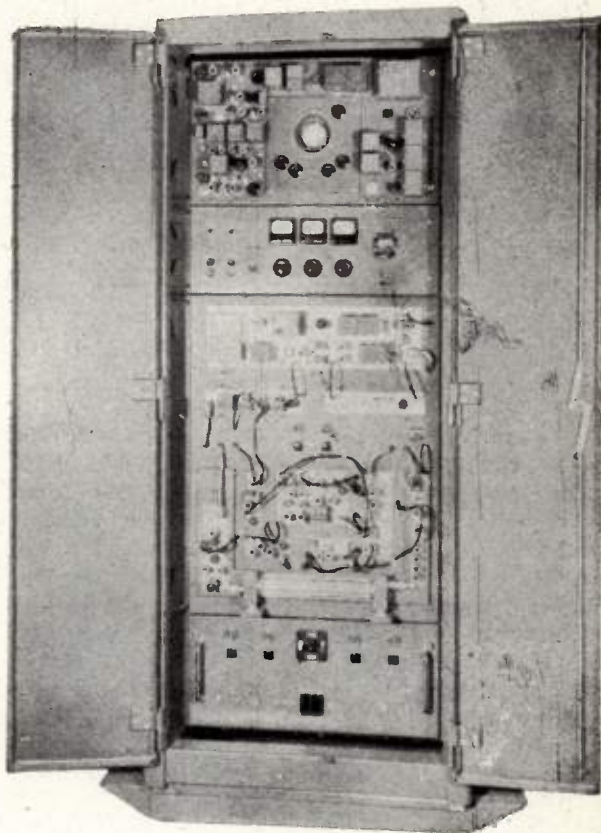


FIG. 5. — Réalisation du filtre linéaire.

Les résultats obtenus ont confirmé la théorie, soit au cours de mesures statiques au générateur, soit sur échos réels, par les soins de l'Administration.

On voit sur les figures 6 et 7 quelques résultats obtenus au cours d'essais préliminaires en usine.

IV. — FILTRES NON LINÉAIRES.

Pour intéressants que soient les résultats obtenus, il est naturel de chercher s'il est possible de faire mieux. A priori, un système qui mettrait en jeu des moments d'ordre plus élevé des distributions d'amplitude du signal étudié, doit conduire à une meil-

leure exploitation de l'information, donc à une plus grande discrimination du signal par rapport au bruit.

Une théorie générale d'optimisation de tels filtres non linéaires a été présentée récemment par L.A. ZADEH [17]. L'optimisation consiste classiquement à rendre minima la moyenne quadratique de la différence entre le signal sorti réellement du filtre et le signal désiré. Mais même dans un cas très simple, le signal d'impulsion stationnaire (analogue à un écho radar) mélangé à du bruit gaussien stationnaire, il est impossible de résoudre analytiquement l'équa-

Le signal étant supposé non fluctuant :

$$s_1 = s_2 = s' \quad (8)$$

la valeur moyenne de signal est s'^2 . Sa variance est la variance du terme $s' (n_1 + n_2)$. Les deux bruits n_1 et n_2 étant gaussiens et indépendants (par hypothèse), la variance de la somme est la somme des variances, soit $2 s' \sigma$.

Le bruit, en l'absence de signal, $n_1 n_2$ est caractérisé par :

$$\overline{n_1 n_2} = 0. \quad (\overline{n_1 n_2})^2 = 4 \sigma^4. \quad \text{Variance } (n_1 n_2) = 2 \sigma^2. \quad (9)$$

L'équation (2) s'écrit donc dans ce cas :

$$\frac{s'}{\sigma} + \frac{s'_2}{2 \sigma^2} = k \quad (10)$$

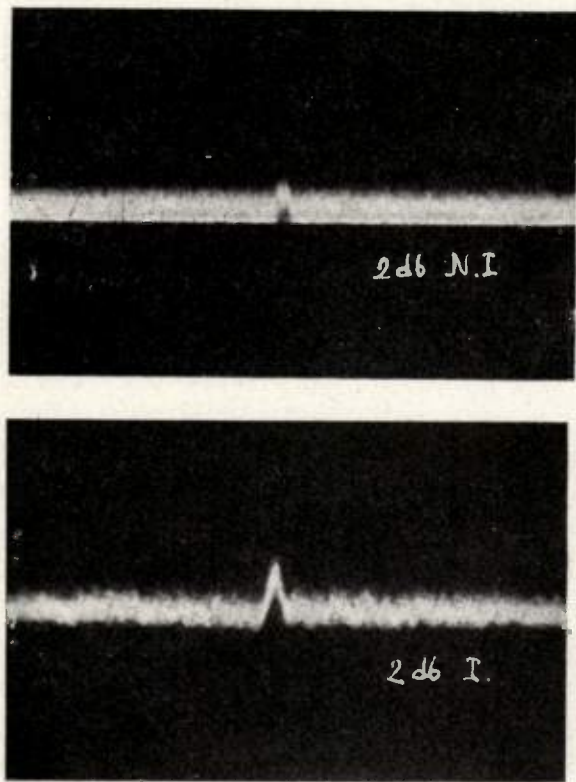


FIG. 6. — Amélioration du filtre linéaire (type « A »).

tion intégrale d'optimisation pour un filtre de classe T1 (c'est-à-dire mettant en jeu la valeur du signal à deux instants différents).

On peut essayer de se rendre compte, sur ce cas simple, du gain que peut procurer une opération non linéaire. La plus simple est la multiplication du signal à l'instant t_1 par le signal à l'instant t_2 . Le signal composite considéré est constitué par du bruit gaussien $n(t)$, superposé à un signal continu d'amplitude « s » (ou zéro).

Le bruit est défini classiquement par :

$$\overline{n(t)} = 0. \quad \overline{n(t)^2} = \sigma^2. \quad \text{Variance } n(t) = \sigma. \quad (5)$$

À l'entrée du dispositif de multiplication, le rapport signal-bruit, défini suivant (2) est :

$$\frac{s}{\sigma} = k. \quad (6)$$

À la sortie du dispositif de multiplication, le signal est donné par :

$$(s_1 + n_1) (s_2 + n_2) = s_1 s_2 + s_1 n_2 + s_2 n_1 + n_1 n_2 \quad (7)$$

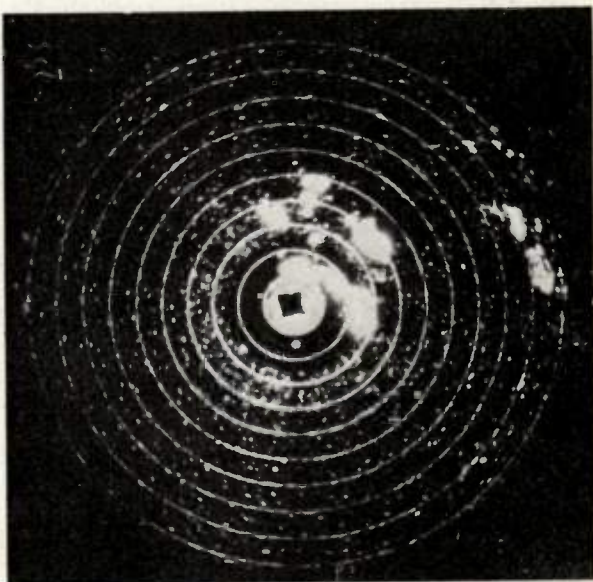
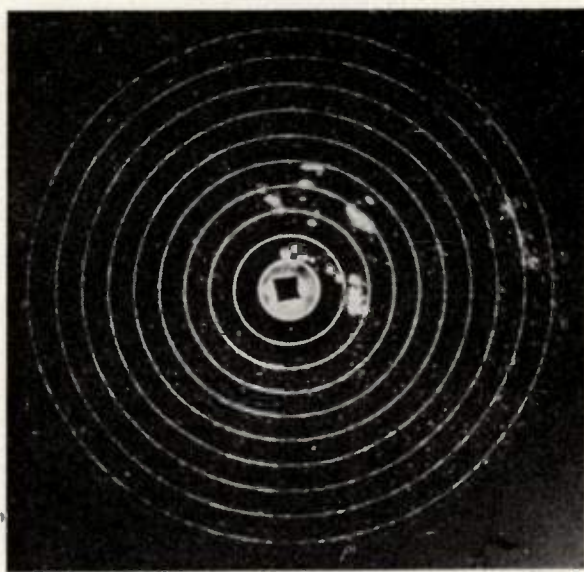


FIG. 7a et b. — Amélioration du filtre linéaire (P.P.I.).

On peut, k étant connu, calculer le nouveau rapport $\frac{s'}{\sigma}$ et, donc, le gain par rapport au rapport signal/bruit initial $\frac{s}{\sigma}$.

k ou $\frac{s}{\sigma}$	0,5	1	4	10
$\frac{s'}{\sigma}$	0,414	0,732	2	3,57
gain dB	1,6	2,7	6	9

On voit l'importance, pour le problème du filtre non linéaire, d'une bonne définition du coefficient

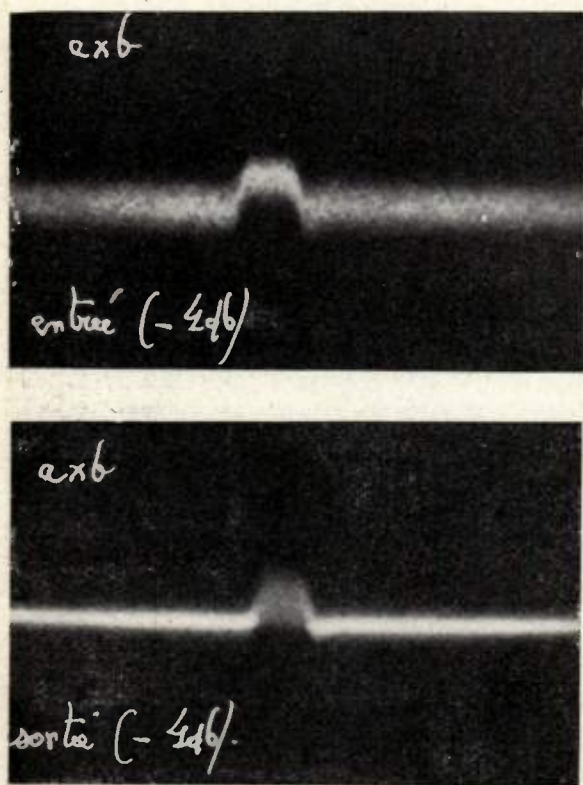


FIG. 8. — Amélioration due à une multiplication (type « A »).

« k » correspondant au signal minimum utilisable sur P.P.I., puisque l'amélioration réalisée est une fonction non linéaire de ce coefficient « k » (ce qui n'était pas le cas pour le filtre linéaire).

La figure 8 montre l'amélioration obtenue expérimentalement par une multiplication. Les résultats sont en bon accord avec la théorie simple exposée ci-dessus.

CONCLUSION

On peut considérer, en conclusion, que les systèmes d'amélioration de signal minimum utilisable ont quitté le domaine du laboratoire pour celui de l'application pratique. Il convient de remercier le Service Technique des Télécommunications de l'Air, sur l'initiative duquel a été entreprise la réalisation du filtre linéaire et tout spécialement MM. NIFFELS et DESCAMPS, avec qui ont eu lieu de nombreux échanges de vues sur tous les points envisagés dans cet exposé.

BIBLIOGRAPHIE

I. — Signal minimum utilisable :

- [1] R.G. HOPKINSON : Visibility of Cathode Ray Tube traces in Radar. *J.I.E.E.* part. III, A. n° 5, 1946, page 795.
- [2] R. PAYNE SCOTT : Visibility of Small Echoes on P.P.I. Displays, *P.I.R.E.*, février 1948, page 180-196.
- [3] E. PARKER and P.R. WALLIS : Three dimensional Cathode Ray Tube Displays. *J.I.E.E.*, part. III, page 383-387. 1948.
- [4] A. W. ROSS : Visibility of Radar Echoes. Analysis of Intensity Modulated Displays. *Wireless Eng.*, Mars 1951, page 79-92.
- [5] R.E. SPENCER : The Detection of Pulses Signal near the Noise threshold. *Journal Brit. Inst. Rad. Eng.*, octobre 1951. Pages 435-455.
- [6] I.L. DAVIES : On determining the Presence of Signals in Noise. *Proc. I.E.E.*, part. III, page 45-51, Mars 1952.
- [7] D. MIDDLETON : Statistical Criteria for the Detection of Pulses Carriers in Noise. *Journal of Applied Physics*. avril 1953. pages 371-378 et 379-391.

II. — Filtres linéaires

- [8] W.J. CUNNINGHAM, J.C. MAY et J.K. SKALNIK. Intégration Noise Reducer for Radar. *Electronics*, septembre 1950. pages 76-78.
- [9] J.V. HARRINGTON et T.F. ROGERS. Signal to Noise Improvement through Integration in a Storage Tube. *P.I.R.E.* octobre 1950. pages 1197-1203.
- [10] L.A. ZADEH. On the Theory of Filtration of Signals. *Zeitschrift für Angewandte Math. and Physic* Vol. III, Fas. 2, (15 Mars 1952) pages 149-156.
- [11] L.A. ZADEH et J.R. RAGAZZINI. Optimum Filters for the Detection of Signals in Noise. *P.I.R.E.*, octobre 1952, pages 1223-1231.
- [12] L. ROBIN. Etude du Signal formé d'Impulsions Rectangulaires Périodiques Superposées à un Bruit Thermique. L'ensemble étant échantillonné, recherche de l'échantillonnage Optimum Eventuel. *Ann. Télécom.* avril 1953, pages 127-130.

III. — Filtres non linéaires :

- [13] Y.W. LEE, T.P. CHEATHAM et J.B. WIENER : Application of Correlation Analysis to the Detection of Periodic Signals in Noise. *P.I.R.E.*, octobre 1950, pages 1165-1172.
- [14] JEAN ICOLE et J. OUDIN : Analyse Temporelle et Filtrage. *Ann. Télécomm.* février 1952, pages 99-108.
- [15] L. ROBIN : Fonction de Corrélation Propre et Spectre de la Densité de Puissance de Bruit Thermique Échantillonné. Filtrage de Signaux Périodiques Simples dans ce Bruit. *Ann. Télécom.* septembre 1952, pages 375-387.
- [16] THOMAS G. SLATTERY : The Detection of a Sine Wave in the Presence of Noise by the Use of a non Linear Filter. *P.I.R.E.*, octobre 1952, pages 1232-1236.
- [17] L.A. ZADEH : Optimum non Linear Filter. *Journal of appl. Physic.* avril 1953, pages 396-404.

EMPLOI DES RÉCEPTEURS A SUPERRÉACTION COMME AMPLIFICATEURS MOYENNE FRÉQUENCE POUR LA RÉCEPTION DES IMPULSIONS DE RADAR

PAR

Serge MARMOR

*Ingénieur, chef du service « Réception »
à la Division « Détection Electromagnétique » du CNET*

GÉNÉRALITÉS.

Pour recevoir et amplifier des impulsions brèves (de l'ordre de 0,5 à 5 μ s) provenant d'un radar on dispose de 3 types classiques de récepteurs :

- 1) Le récepteur « cristal-vidéo » ⁽¹⁾
- 2) Le superhétérodyne
- 3) L'amplificateur à superréaction ⁽²⁾.

Chacun de ces appareils présente des avantages et des inconvénients. Le récepteur à détection directe (cristal-vidéo) est, ainsi que son nom l'indique, composé d'un détecteur à cristal qui reçoit et détecte l'énergie HF, suivi d'un amplificateur « vidéo fréquence ». Ce type d'appareil est séduisant par la simplicité de sa construction et la modicité de son prix de revient. Par contre au point de vue du souffle ce récepteur est franchement mauvais ; pour une bande de 2 Mc/s on ne peut guère espérer déceler des signaux dont la puissance est inférieure à 10^{-9} watt. La sélectivité est déterminée par le circuit d'antenne en amont du cristal. Lorsque le nombre d'étages amplificateurs dépasse 4, on est gêné par les dépassements inverses (undershoots) qui désensibilisent le récepteur pendant des dizaines de microsecondes après le passage d'un signal rectangulaire puissant. De plus, on est obligé de prendre des précautions pour éviter les effets microphoniques. Ces récepteurs ont trouvé principalement leur emploi pour des balises.

Le récepteur superhétérodyne est le plus répandu de tous. Sa popularité est justifiée par de nombreuses considérations.

Le facteur de bruit est excellent. Pour une bande passante de 2 Mc/s, on peut percevoir des signaux dont la puissance avoisine 10^{-13} watt à 3.000 Mc/s. La sélectivité peut être ajustée selon les besoins à des valeurs très différentes, car elle est pratiquement déterminée par les caractéristiques des circuits MF.

Le temps de désensibilisation après le passage d'une impulsion puissante peut être ramené à des valeurs très acceptables.

Ce type d'appareil se prête à la construction en série, à condition que la maquette soit bien étudiée du point de vue des réactions possibles entre l'entrée et la sortie.

Par contre, le superhétérodyne amplifie certaines fréquences indésirables (ondes images, battements entre harmoniques du signal et de l'oscillateur local etc). Le nombre de lampes utilisées pour construire un superhétérodyne de radar ayant un gain de 120 dB et une bande de 2 Mc/s est de 10 et même davantage. Pour la réception des ondes centimétriques, l'oscillateur local est fréquemment constitué par une lampe chère (p. ex. un klystron) nécessitant l'emploi d'alimentations stabilisées.

Enfin, la maquette doit être très soigneusement étudiée pour éviter les réactions gênantes entre la sortie et l'entrée du récepteur. Un poste auquel on n'aurait pas accordé suffisamment de soins à ce point de vue a une bande passante qui se déforme quand on change le point de fixation de la connexion de masse ; parfois on assiste à des accrochages intempestifs.

Tous les constructeurs de récepteurs à gros gain connaissent ce genre de difficultés et la patience qu'il faut déployer pour les surmonter.

La superréaction, dont nous allons parler maintenant, a pendant longtemps suscité beaucoup de méfiance à cause de son fonctionnement capricieux. Le premier article qui traite de ce sujet date de 1922 (Armstrong). A cette époque, on construisait quelques récepteurs de ce type destinés à l'écoute des ondes métriques pour lesquelles la technique des autres types de récepteurs n'était pas au point. Puis, vers 1928 et 1935, les articles de David Frinck, Scroggie, Ataka, expliquèrent les divers modes de fonctionnement des amplificateurs à superréaction, mais il fallut attendre la seconde guerre mondiale pour qu'on en entreprit avec succès la fabrication en série sur une grande échelle.

⁽¹⁾ Récepteurs à détection directe, par S. MARMOR. *Annales des Télécommunications*. Tome 5, n° 7, juillet 1950.

⁽²⁾ Récepteurs à superréaction pour réception des impulsions brèves, par S. MARMOR. *Annales des Télécommunications*. Tome 6, n° 6, juin 1951.

Du point de vue du souffle, les récepteurs à super-réaction sont intermédiaires entre les superhétérodynes et les cristal-vidéo. On peut descendre à une puissance d'entrée de l'ordre de 10^{-11} à 10^{-12} watt. On distingue 3 principaux modes de fonctionnement :

- 1) Le mode linéaire ;
- 2) Le mode logarithmique ;
- 3) A autoblocage.

Pour les applications qui nous intéressent seul le mode linéaire peut convenir en raison des fréquences de découpage élevées qu'il tolère.

PRINCIPE DE FONCTIONNEMENT.

Les récepteurs à superréaction sont essentiellement composés d'une lampe amplificatrice à réaction dans laquelle, grâce à un oscillateur auxiliaire dit de « découpage » ou de « quench », la condition d'amorçage des oscillations libres est périodiquement satisfaite ou non satisfaite.

La croissance et la décroissance de ces oscillations libres se font suivant une loi exponentielle. Lorsqu'on fait travailler la lampe en « mode linéaire », on n'atteint pas la région de saturation pour laquelle l'amplitude de la tension oscillante est constante. L'amplitude de la première oscillation est proportionnelle à la tension (du signal ou du souffle) présente sur la grille au moment où les conditions d'amorçage sont satisfaites. La relation qui relie les amplitudes des $(n + 1)^{ème}$ et première oscillations est

$$\frac{V_{n+1}}{V_1} = e^{n \cdot Q} \quad (1)$$

Q étant la surtension du circuit accordé pendant la période envisagée.

La tension oscillante qui apparaît aux bornes du circuit accordé est ensuite détectée et amplifiée en vidéo-fréquence.

La formule (1) indique que le gain d'un étage à superréaction augmente rapidement avec le nombre n d'oscillations comprises dans la période d'entretien.

On ne peut donc pas augmenter exagérément la fréquence de l'oscillateur de découpage sans diminuer sérieusement le gain de l'étage.

Après avoir atteint une amplitude maximum à la fin de la période d'entretien des oscillations libres, ces dernières décroissent exponentiellement ; cette période de l'amortissement est un temps mort qu'il convient de raccourcir le plus possible.

Il faut que l'amplitude des oscillations libres à la fin de la période d'extinction soit négligeable devant la tension du signal ou du souffle présente dans le C.O. au moment du démarrage du cycle de croissance suivant, afin que ce soit cette tension qui détermine l'amplitude de la $(n + 1)^{ème}$ oscillation de la formule (1). On obtient ainsi l'indépendance de phase entre deux périodes oscillatoires successives complètes (croissance et décroissance) ; le relevé de l'amplification en fonction de la fréquence a l'aspect, pour ce mode de fonctionnement, d'une courbe de résonance normale, en forme de cloche (fig. 1).

Au contraire, quand l'extinction des oscillations libres est insuffisante, il y a cohérence de phase entre deux cycles consécutifs complets et la tension de

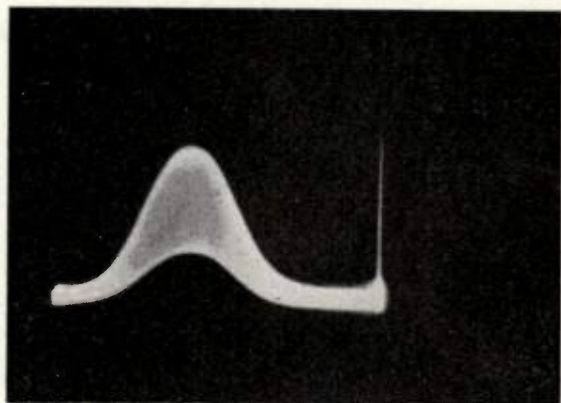


FIG. 1.

l'oscillateur de découpage module l'enveloppe de l'onde HF, faisant apparaître des bandes latérales ayant un espacement égal à la fréquence f_q de cet oscillateur. La courbe d'amplification devient irrégulière : elle présente des creux et des protubérances pointues au sommet desquelles, uniquement, la sensibilité est bonne. (fig. 2 et 3).

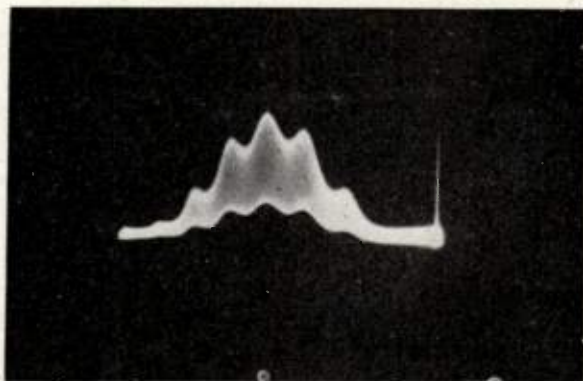


FIG. 2.

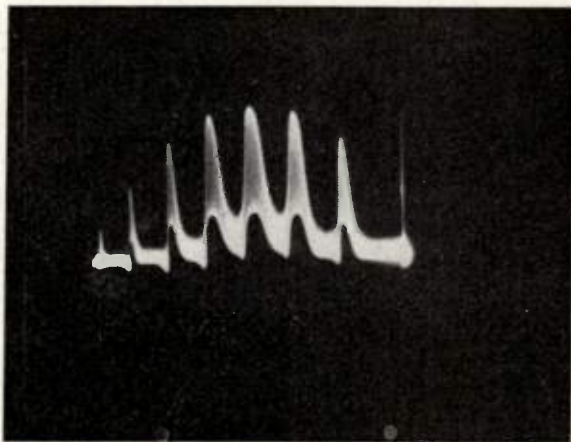


FIG. 3.

CARACTÉRISTIQUES SPÉCIALES REQUISES DES RÉCEPTEURS À SUPERRÉACTION DESTINÉS À L'AMPLIFICATION DES IMPULSIONS BRÈVES.

1. Relation entre la fréquence de découpage f_q et la durée τ de l'impulsion.

Si l'on désire recevoir toutes les impulsions émises par le radar à une fréquence de récurrence f_r , il faut que pendant la durée de chacune de ces impulsions

La largeur de bande de l'amplificateur vidéo doit être suffisante pour amplifier cette tension. Ainsi, si 1 cycle oscillatoire complet dure 1 micro seconde, l'amplificateur vidéo devra passer 1 Mc/s environ. L'allongement ou le raccourcissement de la durée τ de l'impulsion n'influe pas sur la largeur de la bande vidéo mais sur le pourcentage d'impulsions effectivement reproduites.

Dans les anciens récepteurs à superréaction desti-

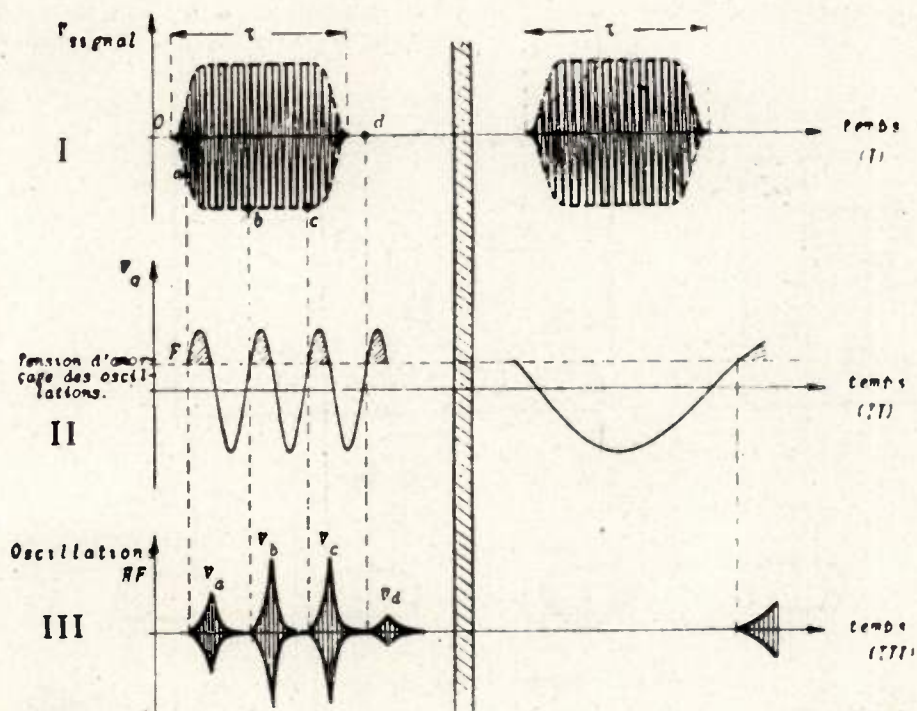


FIG. 4

on ait au moins un cycle oscillatoire complet ; si τ est la durée de l'impulsion, la fréquence minimum de « quench » est donc égale à

$$f_{q \text{ min}} = \frac{1}{\tau} \quad (2)$$

Ainsi, pour ne pas perdre d'impulsions ayant une durée de 1 microseconde, la fréquence de l'oscillateur de découpage sera au moins égale à 1 Mc/s.

La figure 4 indique la réponse du récepteur pour une impulsion de même durée quand $f_q > \frac{1}{\tau}$ (à gauche) et $f_q < \frac{1}{\tau}$ (à droite).

2. Largeur de bande de l'amplificateur vidéo-fréquence

Chaque paquet d'impulsions détermine, après détection et passage dans un circuit intégrateur à constante de temps convenablement choisie, une tension dont le contour peut être grossièrement assimilé à un triangle dont l'amplitude est proportionnelle au bruit ou au signal existant dans le C.O. au moment du déclenchement des oscillations libres.

nés à l'écoute radiophonique, la fréquence de découpage f_q était de l'ordre de 30 à 40 Kc/s. Elle était inaudible, et de ce fait, peu gênante. On se contentait d'un filtrage rudimentaire pour empêcher cette tension de « quench » de parvenir jusqu'au haut-parleur. La figure 5 donne le schéma d'un de ces récepteurs, muni d'une CAG. La tension de découpage est envoyée, après détection, dans une lampe amplificatrice dont le circuit plaque comporte un transformateur T_1 accordé sur la fréquence f_q . La tension que ce dernier délivre est appliquée à la grille de l'oscillateur HF après détection et passage dans un circuit RC dont la constante de temps est égale à 0,1 seconde.

Pour empêcher la tension de découpage V_q d'atteindre le haut-parleur on a prévu un simple filtrage réalisé par une self de choc et un condensateur de 1 000 pF.

Ces dispositions sommaires seraient insuffisantes pour un récepteur de radar attaquant un oscillographe cathodique dont la trace horizontale doit être débarrassée des ondulations indésirables à la fréquence f_q .

Dans l'exemple que nous avons cité, la fréquence de découpage $f_q = 1 \text{ MC/s}$ tombe en plein dans la bande passante de l'amplificateur vidéo et l'élimination de V_q pose un problème délicat à résoudre.

- 1 » détectrice
- 1 » 1^{re} vidéo-fréquence
- 1 » cathode suiveuse
- 1 » 2^e vidéo-fréquence.

Total : 11 tubes.

Bande à 3 dB = 2,5 Mc/s

Facteur de bruit : 4 dB.

2) Récepteur à superréaction CNET.

Ce récepteur comprend

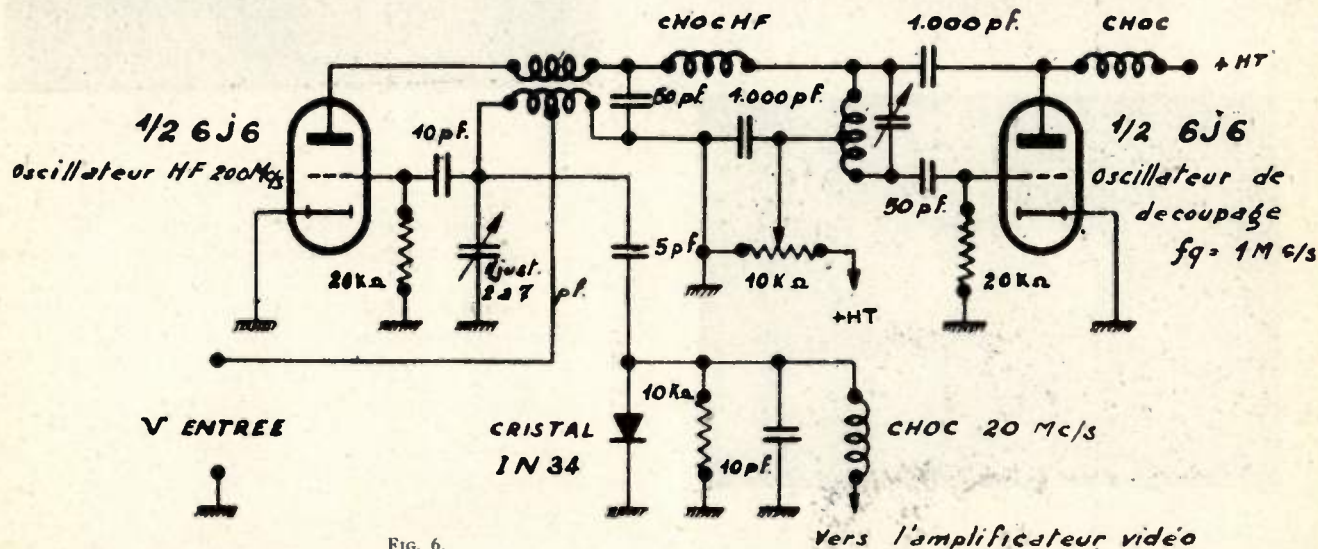


FIG. 6.

1 lampe double 6J6 fonctionnant en :

- oscillateur HF à superréaction ; $f_{HF} = 200 \text{ Mc/s}$
- oscillateur de découpage ; $f_q = 1 \text{ Mc/s}$.

3 lampes amplificatrices vidéo-fréquence (2: EF42 + 1: EL 83).

Total : 4 lampes.

Bande passante à 3 db : 3 Mc/s.

Alimentation : 6,3 v ; 50 pps — pour le chauffage
+ 200 v ; 50 mA — pour la H.T.

La lampe de sortie attaque la plaque Y_1 du tube à rayons cathodiques.

Comparaison des deux récepteurs.

a) Courte distance.

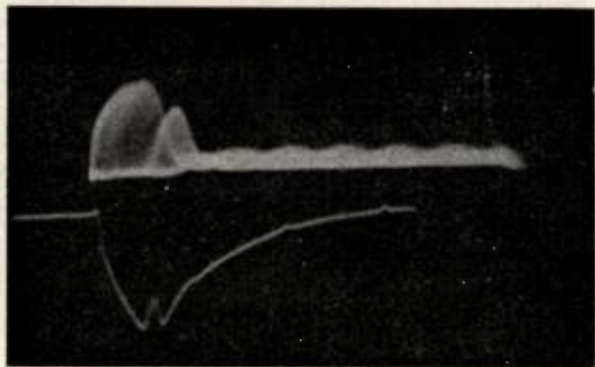


FIG. 7.

Type d'indicateur : Scope A.

Objectif et Distance : Pylône métallique à 110 mètres.

Remarques : L'image du haut provient du récepteur à superréaction. On remarquera l'image de l'écho visible à l'intérieur de l'impulsion de l'émetteur et les ondulations rémanentes de V_q (durée entre deux maxima : 1 micro seconde).

L'image du bas provient du récepteur à MF classique. On peut y distinguer les « pips » du marqueur distants entre eux de 2 microsecondes.

b) Moyenne distance.

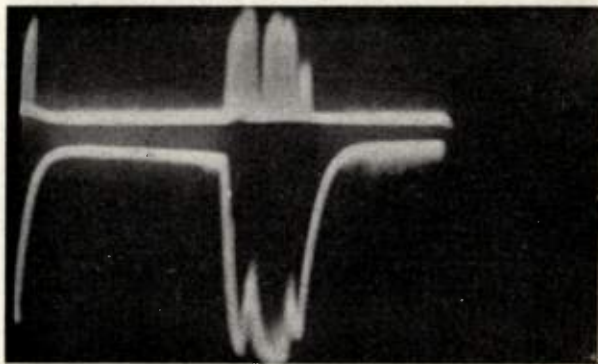


FIG. 8.

Type d'indicateur : Scope A ; Durée de l'impulsion : $\tau = 2 \text{ microsecondes}$.

Objectif et distance : Cap Ferrat distant de 7 kilomètres.

Remarques : L'image du haut provient du récepteur à superréaction. La luminosité du TRC est très poussée (pour la photographie) et la trace en est épaissie. On remarquera la réapparition presque immédiate d'une herbe clairsemée après le passage de l'impulsion de l'émetteur. L'image de l'écho est pleine et lumineuse.

L'image du bas provient de l'amplificateur MF classique. La trace se soulève au passage de l'écho. Le récepteur est désensibilisé pendant un certain temps après le passage de l'impulsion de l'émetteur et l'écho réapparaît franchement sous sa forme dense habituelle après la disparition de l'écho (soit environ 45 microsecondes après l'émission).

Type d'indicateur : P.P.I., Durée de l'impulsion : $\tau = 0.5$ microseconde.

Objectif et distance : Cap Ferrat (7 Km) et Baie des Anges à Nice (15 kilomètres).



FIG. 9.



FIG. 10.

Remarques :

La fig. 9 correspond au récepteur à superréaction.

La fig. 10 correspond au récepteur classique.

On remarquera l'absence d'échos de mer au large du Cap Ferrat sur la fig. 9. Les échos situés à 15 kilomètres sont moins fournis que sur la fig. 10. Le fond noir de la fig. 9 est parsemé de taches plus claires dues au souffle (peu gênantes).

c) Grandes distances.

Type d'indicateur : Scope A.

Objectif et distance : Cap Camarat à 85 kilomètres.

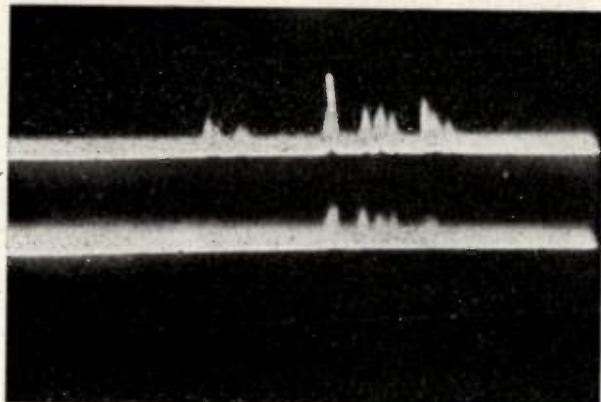


FIG. 11.

Remarques :

La trace du haut est relative au récepteur classique.

La trace du bas concerne le récepteur à superréaction.

Aux grandes distances, la supériorité du superhétérodyne devient manifeste.

II. — INFLUENCE DE LA DURÉE τ DES IMPULSIONS.

a) Moyennes distances :

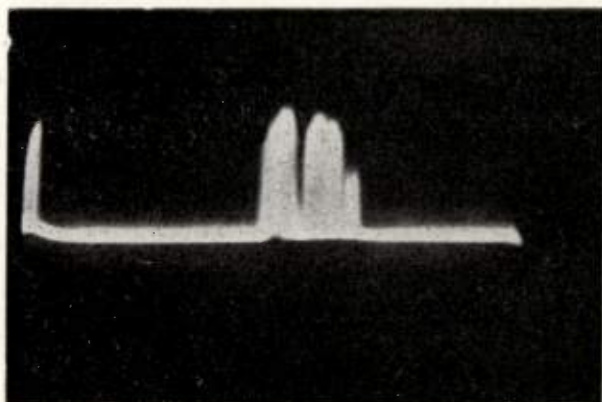


FIG. 12.

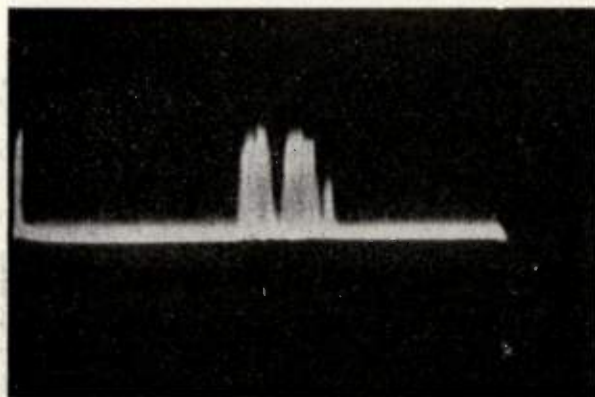


FIG. 13.

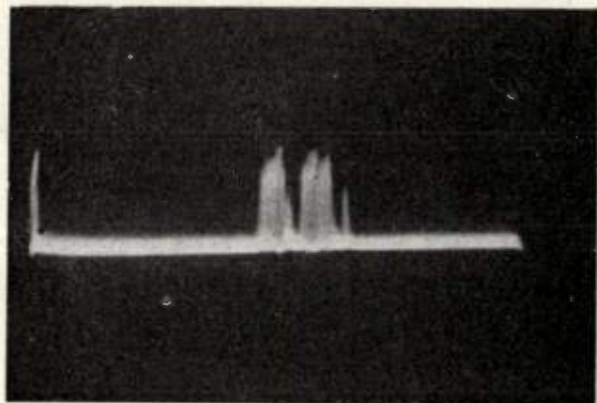


FIG. 14

Type d'indicateur : Scope A.

Distance et objectif : Cap Ferrat à 7 kilomètres.

Durée de l'impulsion :

τ microseconde	N° Figure
2	12
1	13
0,5	14

Remarques : On constate, ainsi que prévu, une diminution du nombre d'impulsions reçues quand τ

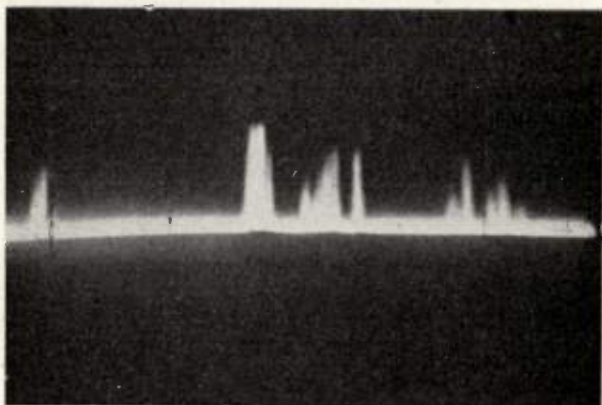


FIG. 15

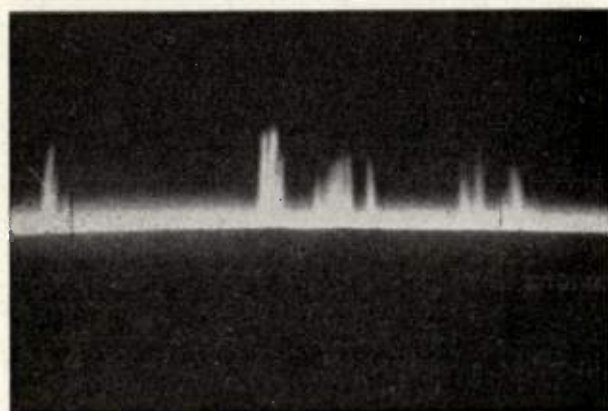


FIG. 16

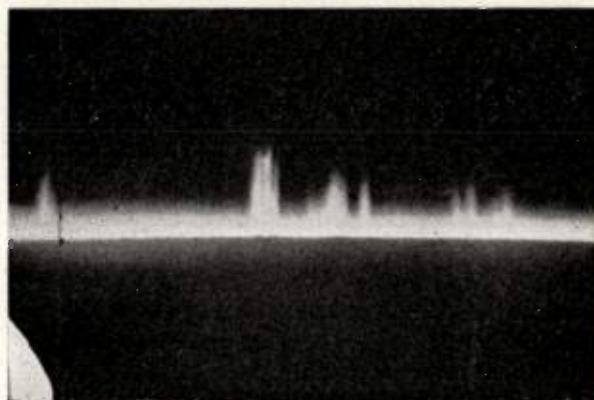


FIG. 17

diminue. L'écho sur la droite de l'image est un phare. Quand τ diminue, il se détache de plus en plus nettement.

b) Grandes distances.

Type d'indicateur : Scope A.

Distance et objectif :

De la gauche vers la droite, les échos représentent : Le cap Ferrat, le Cap d'Antibes, les Iles de Lérins, l'Estérel. L'échelle est de 60 kilomètres.

Durée de l'impulsion :

τ microseconde	N° figure
2	15
1	16
0,5	17

Les échos s'affaiblissent quand τ diminue à cause de la diminution du nombre d'impulsions reproduites.

*
* * *

CONCLUSIONS :

Le récepteur à superréaction est moins sensible que le superhétérodyne classique. Il peut dans certains cas, quand on ne recherche pas une grande précision et une très longue portée, fournir des résultats satisfaisants ; l'économie de tubes (4 au lieu de 11 dans le cas du radar étudié) est considérable ; la puissance nécessaire pour l'alimenter est beaucoup moins importante. La facilité avec laquelle l'oscillateur HF à superréaction peut être mué en émetteur peut être mise à profit pour la construction de transceivers. Le récepteur à superréaction est très stable car il ne contient pas de chaînes d'amplification sur la même fréquence et que le phénomène de réaction et d'accrochage HF y est nécessaire pour le fonctionnement de l'appareil au lieu de présenter un caractère indésirable comme dans les autres types de récepteurs.

PERTURBATIONS D'UN OSCILLATEUR NON-LINÉAIRE FILTRÉ

PAR

Gilbert CAHEN

Etablissement des équations de perturbation d'un circuit bouclé comprenant un amplificateur non linéaire et un filtre passe-bas coupant les harmoniques de la fréquence fondamentale d'oscillation. Entraînement par une perturbation sinusoïdale ; traînage de fréquence ; seuil de synchronisation ; relation entre fréquence et amplitude. Couplage de deux oscillateurs, seuil de synchronisation.

1. — INTRODUCTION.

1,1 Historique.

DUTILH, [1], puis KOCHENBURGER [2] ont montré que, parmi les circuits bouclés non linéaires, une classe importante comprenait ceux dont, en un point de la boucle d'asservissement, la grandeur physique soumise à l'oscillation a une variation pratiquement sinusoïdale en fonction du temps, les harmoniques de l'oscillation fondamentale pouvant être négligées.

LOEB [3], [4] et [5] a prouvé la fécondité de cette notion en définissant une fonction de transfert générale du circuit ouvert — dont le bouclage provoque l'entrée du système en oscillation — et en montrant que la connaissance de cette fonction de transfert permet de définir un critère général d'accrochage des oscillations — dont le critère de NYQUIST est un cas particulier — et un critère de stabilité des oscillations correspondantes. Dans ce qui suit, nous allons appliquer ces mêmes hypothèses à l'étude de l'entraînement des oscillateurs non linéaires filtrés.

1,2 — La définition mathématique précise d'un circuit amplificateur, non linéaire, filtré, présente des difficultés évidentes. Nous allons adopter la suivante :

Appliquons d'abord, à l'entrée du circuit, une excitation sinusoïdale d'affixe angulaire ($\omega t + \theta$) :

$$1,2 - 1 \quad X = x \cos (\omega t + \theta)$$

l'amplitude x , la pulsation ω et la phase θ étant constantes. Grâce au filtrage, on recueillera, à la sortie, une oscillation également sinusoïdale, la « réponse », que nous écrirons :

$$1,2-2 \quad Y = xg \cos (\omega t + \theta + \varphi)$$

Ainsi, l'action de l'amplificateur se traduit par une modification de l'amplitude, qui de x passe à xg , et de phase, qui de θ passe à $\theta + \varphi$.

g et φ sont, chacun, des fonctions de x et de ω .

Cette définition permet [3] de définir les conditions d'oscillation stationnaire du système bouclé ; elle ne suffit pas pour l'étude des régimes transitoires ni des perturbations.

1,3 — Si nous appliquons, maintenant, à l'entrée, une excitation encore donnée par la formule 1,2-1, mais dans laquelle x et θ sont des fonctions du temps, nous pourrions encore appliquer la formule 1,2-2 ; mais, dans ce cas, y et φ seront des fonctions du temps qui ne seront pas reliées de façon simple à x et ω .

Supposons, cependant, que θ varie lentement avec le temps, en ce sens que sa dérivée première reste petite devant ω ; que, d'autre part, x varie lentement avec le temps, en ce sens que l'on peut encore assimiler X à une fonction sinusoïdale modulée d'amplitude x et de pulsation apparente

$$\omega + \dot{\theta} = \bar{\omega}$$

1,4 — Dans le cas d'un amplificateur linéaire, g et φ dépendent, en toute rigueur, comme l'ont montré CARSON et FRY [11], de la pulsation $\bar{\omega}$, du décrement ou variation logarithmique de l'amplitude, $\tau = \frac{\dot{x}}{x}$, et de toutes leurs dérivées successives.

Cependant, au voisinage d'un régime d'équilibre caractérisé par une pulsation ω_0 , autrement dit, si $\frac{\omega - \omega_0}{\omega_0}$ et $\frac{\tau}{\omega_0}$ sont petits, et plus précisément si les produits de ces quantités par le coefficient de sur-

(¹) Communication présentée en séance de section à la Société Française des Electriciens et à la Société des Radioélectriciens, le 15 juin 1953.

tension Q du circuit peuvent être négligés devant l'unité, on est en droit de ne pas tenir compte des dérivées de $\bar{\omega}$ et de τ . L'étude de cette question, déjà amorcée dans deux de nos notes antérieures [9] et [12] (cette dernière publiée en collaboration avec M. LOEB), est reportée en annexe.

Nous admettrons donc, à titre d'approximation, et sous la réserve précédente, que g et φ sont, dans un amplificateur linéaire, fonctions uniquement de $\bar{\omega}$ et de τ . On peut même ajouter que $g \exp(j\varphi)$ est alors fonction analytique de $\tau + j\omega$, ce qui entraîne les relations :

$$1,4-1 \left\{ \begin{array}{l} \frac{1}{g} \frac{\partial g}{\partial \tau} = \frac{\partial \varphi}{\partial \bar{\omega}} \\ \frac{1}{g} \frac{\partial g}{\partial \bar{\omega}} = - \frac{\partial \varphi}{\partial \tau} \end{array} \right.$$

1,5 — Revenons au cas des amplificateurs non linéaires. Nous savons que, dans ce cas, le gain g et la phase φ dépendent de l'amplitude x de l'oscillation à l'entrée. Comme précédemment, ils dépendront aussi, bien entendu, de ω et de τ , mais cette dépendance ne sera plus analytique [12], ce qui nous interdit de parler de gain complexe $g \exp(j\varphi)$, puisque les conditions 1,4-1 ne sont pas remplies.

Ici encore, g et φ dépendent certainement, dans le cas général, de l'ensemble de l'« histoire » de x et de $\bar{\omega}$ mais, comme au paragraphe 1,4, nous admettrons, à titre d'hypothèse fondamentale, que l'on peut se contenter de prendre en considération les dépendances de g et de φ , par rapport à x , ω et τ .

Rappelons que la condition de validité de cette hypothèse, indiquée au paragraphe 1,4 et justifiée en annexe, est que $Q \frac{\omega - \omega_0}{\omega^0}$ et $\frac{Q\tau}{\omega_0}$ soient assez petits devant l'unité pour que leurs carrés soient négligeables.

Notons que la validité de notre hypothèse ne suppose pas que le taux d'harmoniques soit nul absolument ; il suffit que les variations possibles des proportions des divers harmoniques soient assez réduites pour n'entraîner, par le jeu de leurs battements, que des variations négligeables de l'amplitude et de la phase de l'oscillation fondamentale.

1,6 — Au voisinage du régime normal d'oscillation, caractérisé par une amplitude x_0 , une pulsation ω_0 , un gain $g_0 = 1$ et une phase $\varphi = \pi$, comme nous le verrons au paragraphe 2,1, le gain g et la phase φ ne dépendent pas explicitement des autres dérivées de g et de φ , de telle sorte que l'on puisse écrire :

$$1,6-1 \left\{ \begin{array}{l} \delta g = \frac{\partial g}{\partial x} \delta x + \frac{\partial g}{\partial \bar{\omega}} \delta \bar{\omega} + \frac{\partial g}{\partial \tau} \delta \tau \\ \delta \varphi = \frac{\partial \varphi}{\partial x} \delta x + \frac{\partial \varphi}{\partial \bar{\omega}} \delta \bar{\omega} + \frac{\partial \varphi}{\partial \tau} \delta \tau \end{array} \right.$$

Pour la commodité des calculs, nous adopterons, en outre, les notations suivantes :

$$1,6-2 \left\{ \begin{array}{ll} \frac{\partial g}{\partial x} = a \sin \alpha & \frac{\partial \varphi}{\partial x} = a \cos \alpha \\ \frac{\partial g}{\partial \bar{\omega}} = b \sin \beta & \frac{\partial \varphi}{\partial \bar{\omega}} = b \cos \beta \\ \frac{\partial g}{\partial \tau} = x_0 c \sin \gamma & \frac{\partial \varphi}{\partial \tau} = x_0 c \cos \gamma \end{array} \right.$$

les huit quantités x_0 , ω_0 , a , b , c , α , β , γ suffiront à caractériser physiquement l'amplificateur ; nous verrons, au paragraphe 5, comment mesurer les six dernières.

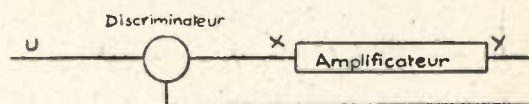
On trouve que la condition d'analyticité, représentée par les conditions 1,4-1 s'écrit avec les notations 1,6-2.

$$1,6-3 \left\{ \begin{array}{l} \gamma - \beta = \frac{\pi}{2} \\ b = x_0 c \end{array} \right.$$

Cette double condition n'est, en général, pas remplie dans un système non linéaire ; la première l'est, cependant, très souvent.

2. — LES ÉQUATIONS DE PERTURBATION DU SYSTÈME BOUCLÉ.

2,1 — Si nous bouclons l'amplificateur A sur un discriminateur D (fig. 1), nous aurons le schéma d'un asservissement de x à u .



Nous admettrons que :

$$2,1-1 \quad x = u - y$$

En l'absence de perturbation U , le système se mettra à osciller, à condition que Y soit égal en module et opposé en phase à X .

Les conditions d'oscillation sont donc :

$$2,1-2 \quad y_0 = x_0$$

$$2,1-3 \quad \varphi_0 = \pi$$

A ce moment, le système a une pulsation d'oscillation ω_0 et une amplitude d'oscillation x_0 (conditions d'équilibre).

L'objet de la présente étude est le comportement du système en présence d'une excitation.

$$2,14 \quad U = u x_0 \cos \omega t$$

dont la pulsation est voisine de la pulsation d'équilibre ω_0 et dont l'amplitude est petite par rapport

à l'amplitude d'équilibre x_0 . Plus précisément (cf. paragraphe 1,5).

$$2,1-5 \quad \left\{ \begin{array}{l} \omega - \omega_0 = \varepsilon \ll \frac{\omega_0}{Q} \end{array} \right.$$

En outre, nous bornerons notre étude au cas où

$$2,1-6 \quad u \ll 1$$

2,2 — Nous ne pouvons, a priori, connaître le comportement de l'amplificateur A en présence de la somme algébrique de U et de Y , puisque le principe de superposition n'est pas applicable. Nous allons tourner la difficulté, en admettant qu'à l'entrée de A nous avons une oscillation de pulsation $\omega_0 + \varepsilon$, dont la phase θ et l'amplitude $x_0 + \xi$ peuvent être fonctions du temps ; nous admettrons, également, que ξ reste petit devant l'unité si u et ε sont petits.

Dans ce cas, d'après l'hypothèse faite aux paragraphes 1,3 et 1,6, on aura à la sortie de A une oscillation Y , d'amplitude $(x_0 + \xi)(1 + \delta g)$, de pulsation $\omega + \varepsilon$, et de phase $\pi + \delta \varphi$, et nous pourrions écrire l'équation de bouclage 2,1-1 comme suit :

$$2,2-1 \quad (1 + \delta g)(x_0 + \xi) \cos [(\omega_0 + \varepsilon)t + \theta + \pi + \delta \varphi] = u x_0 \cos (\omega_0 + \varepsilon)t - (x_0 + \xi) \cos [(\omega_0 + \varepsilon)t + \theta]$$

En égalant séparément les termes en $\cos (\omega_0 + \varepsilon)t$ et $\sin (\omega_0 + \varepsilon)t$ et en négligeant les termes du second ordre, nous obtenons, après développement :

$$2,2-2 \quad \left\{ \begin{array}{l} -\delta g \cos \theta + \delta \varphi \sin \theta = u \\ -\delta g \sin \theta - \delta \varphi \cos \theta = 0 \end{array} \right.$$

d'où

$$2,2-3 \quad \left\{ \begin{array}{l} \delta g = -u \cos \theta \\ \delta \varphi = u \sin \theta \end{array} \right.$$

Nous remplacerons δg et $\delta \varphi$ par leurs valeurs 1,6-1, compte tenu des notations 1,6-2, en remarquant que

$$\delta x = \xi, \quad \delta \omega = \varepsilon + \dot{\theta}, \quad \text{et} \quad \tau = \frac{\ddot{\xi}}{x_0}$$

Nous trouvons :

$$2,2-4 \quad \left\{ \begin{array}{l} -u \cos \theta = a \sin \alpha \cdot \xi + b \sin \beta \cdot (\varepsilon + \dot{\theta}) + c \sin \gamma \cdot \ddot{\xi} \\ u \sin \theta = a \cos \alpha \cdot \xi + b \cos \beta \cdot (\varepsilon + \dot{\theta}) + c \cos \gamma \cdot \ddot{\xi} \end{array} \right.$$

Nous extraierons de ces équations la valeur de ξ et celle de $\ddot{\xi}$

$$2,2-5 \quad \left\{ \begin{array}{l} \xi = \frac{u \cos (\theta - \gamma) - b (\varepsilon + \dot{\theta}) \sin (\gamma - \beta)}{a \sin (\gamma - \alpha)} \\ \ddot{\xi} = \frac{u \cos (\theta - \alpha) + b (\varepsilon + \dot{\theta}) \sin (\alpha - \beta)}{c \sin (\gamma - \alpha)} \end{array} \right.$$

Nous dériverons la première relation et égalerons à la seconde, pour avoir, en définitive, (cf nos communications [13] et [14]) :

$$2,2-6 \quad \left\{ \begin{array}{l} \frac{c \sin (\gamma - \beta)}{a \sin (\alpha - \beta)} \ddot{\theta} + \left[1 + \frac{uc \sin (\theta - \gamma)}{ab \sin (\alpha - \beta)} \right] \dot{\theta} \\ + \varepsilon \left[1 - \frac{u \cos (\theta - \alpha)}{b \sin (\alpha - \beta)} \right] = 0 \end{array} \right.$$

Notons tout de suite que le coefficient de $\dot{\theta}$ est forcément positif si, en l'absence d'excitation extérieure ($u = \varepsilon = 0$) le régime d'oscillation est stable : on retrouve la cond tion de LOEB [5], sous une forme différente.

3. — ETUDE DE L'ENTRAÎNEMENT DE L'OSCILLATEUR.

3,1 — L'équation 2,2-12 entre dans la catégorie des équations non linéaires du second ordre, à variable indépendante non explicite, et linéaires par rapport aux dérivées, auxquelles nous avons consacré une étude spéciale ([8] et [15]). Posons

$$3,1-1 \quad h = \frac{u}{b \varepsilon \sin (\alpha - \beta)}$$

$$3,1-2 \quad k = \varepsilon h \frac{c}{a} = \frac{ab \sin (\alpha - \beta)}{u c} = m \lambda$$

$$3,1-3 \quad m = \varepsilon \frac{c \sin (\gamma - \beta)}{a \sin (\alpha - \beta)}$$

$$3,1-4 \quad \theta - \alpha = \psi$$

$$3,1-5 \quad \gamma - \alpha = \mu$$

$$3,1-6 \quad \varepsilon t = t_1$$

$$3,1-7 \quad \int_0^{\psi} (1 + k \sin (\psi - \mu)) d\psi = q = \psi - k \cos (\psi - \mu)$$

$$3,1-8 \quad m (1 - h \cos \psi) = r$$

$$3,1-9 \quad m \dot{\psi} = p = s - q$$

(la dérivée étant prise par rapport au temps t_1)

L'équation 2,2-12 peut revêtir les deux formes suivantes :

$$3,1-10 \quad \left\{ \begin{array}{l} m \dot{\psi} + \left[1 + k \sin (\psi - \mu) \right] \dot{\psi} \\ + (1 - h \cos \psi) = 0 \end{array} \right.$$

$$3,1-11 \quad \left\{ \begin{array}{l} \frac{ds}{d\psi} = \frac{q - s}{r} \end{array} \right.$$

La première forme est utile pour l'étude analytique.

La deuxième forme se prête à l'emploi de la méthode géométrique, dérivée de la méthode de LIENARD, que nous avons étudiée dans les notes précitées [8] et [15] : traçons, comme dans ces mémoires, une courbe $Q(\psi)$ et une courbe $R(\psi)$ définies com-

me suit : on porte (figure 2) q en ordonnées pour les deux courbes, et en abscisses, ψ pour la première, $\psi - r$ pour la seconde, de telle sorte que $\overline{OP} = \psi$, $PQ = q$, $\overline{RQ} = r$.

Soit S un point d'ordonnée $s = p + q$. La condition 3,11 exprime que la normale, en S , à la courbe intégrale, passe par R .

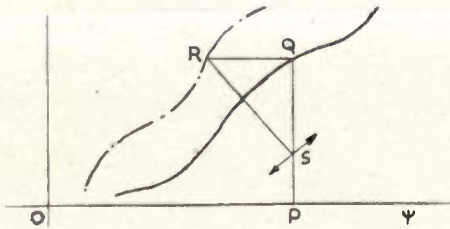


FIG. 2

3,2 — Selon la forme des fonctions r et q , nous aurons diverses classes de solutions.

Nous limiterons ici notre étude aux cas pratiquement utiles, compte tenu des hypothèses qui ont conduit à l'équation fondamentale 2,2-12.

En particulier, nous savons que le coefficient k , du fait de la présence de u au numérateur, est petit devant l'unité (hypothèse 2,1-6) ; cela ne cesserait d'être vrai que si le coefficient de non-linéarité, a , devenait très faible. Il en résulte que la courbe q (ψ) sera une sinusoïde en axes obliques, très aplatie et s'écartant peu de la droite à 45°. La périodicité est 2π , tant pour la courbe q que pour la courbe r .

Nous allons étudier successivement les différents cas, caractérisés par les valeurs de h , donc de l'amplitude u de la perturbation.

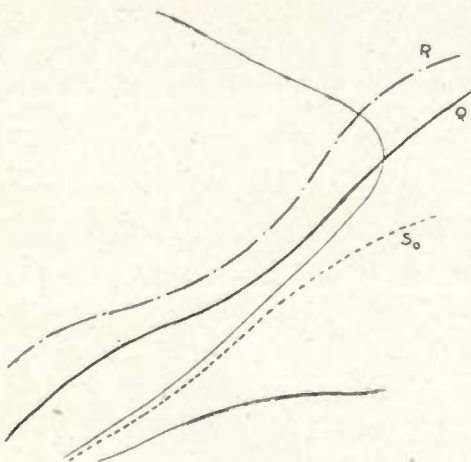


FIG. 3

3,3. — Cas où h est inférieur à l'unité.

3,31 — Les courbes (Q) et (R) ne se rencontrent pas. Comme il a été démontré dans le précédent mémoire [15] il existe une trajectoire limite (S_0), de direction générale parallèle à celle des courbes (Q) et (R) (figure 3). (Cycle limite de deuxième espèce).

* La courbe (Q) est en traits pleins, la courbe (R) en traits mixtes.

3,32 — La quantité $p = m \frac{d\psi}{dt}$, qui représente la distance verticale de la courbe (S) à la courbe (Q), oscillera autour d'une valeur constante, une fois atteinte — pratiquement — la trajectoire limite.

La phase ψ aura donc une variation séculaire et une variation périodique. A la variation séculaire correspondra un écart de fréquence moyen entre la fréquence de la perturbation et celle de l'oscillation principale ; à la variation périodique correspondra une modulation en fréquence de cette oscillation principale, qui entraînera à son tour — par le jeu de la formule 2,2-11 — une modulation en amplitude : celle-ci équivaut à un battement entre l'oscillateur et sa perturbation.

3,33 — Nous allons rechercher une solution approximative tenant compte du fait que ε est petit, ou, plus précisément que $m = \varepsilon \frac{c \sin(\gamma - \beta)}{a \sin(\alpha - \beta)}$ est petit devant l'unité.

L'équation 3-1-2 peut encore s'écrire, en tenant compte de 3,1-3 :

$$3,33-1 \quad -\dot{\psi} = 1 - h \cos \psi + m [\ddot{\psi} + \lambda \dot{\psi} \sin(\psi - \mu)]$$

m étant, pour les faibles écarts entre ω_0 et la pulsation du perturbateur ($\varepsilon \ll \omega_0$), petit par rapport à l'unité. La solution de première approximation s'obtient en faisant $m = 0$.

$$3,33-2 \quad -\dot{\psi}_1 = 1 - h \cos \psi$$

Dérivons cette expression :

$$3,33-3 \quad -\ddot{\psi}_1 = h \sin \psi$$

Nous substituerons ces deux expressions dans le facteur de m , entre crochets, de la formule 3,33-2, et obtiendrons la solution de 2^e approximation.

$$3,33-4 \quad -\dot{\psi}_2 = 1 - h \cos \psi + m \left[-h \sin \psi - \lambda \sin(\psi - \mu) (1 - h \cos \psi) \right]$$

Les termes manquants sont du 2^e ordre en m . Le temps nécessaire pour que retrouve la même valeur, à 2π près, vaut

$$3,33-5 \quad T = \frac{1}{\varepsilon} \int_0^{2\pi} \frac{d\psi}{1 - h \cos \psi} = \frac{1}{\varepsilon} \int_0^{2\pi} \frac{d\psi}{(1 - h \cos \psi) \left[1 - m \left\{ \lambda \sin(\psi - \mu) + \frac{1 - h \cos \psi}{h \sin \psi} \right\} \right]}$$

et, en négligeant les termes en m^2 :

$$3,33-6 \quad T = \frac{1}{\varepsilon} \int_0^{2\pi} \frac{d\psi}{1 - h \cos \psi} \left[1 + m \lambda \sin(\psi - \mu) + \frac{m h \sin \psi}{1 - h \cos \psi} \right]$$

$$3,33-7 \quad T = \frac{m \lambda \sin \mu}{h \varepsilon} \int_0^{2\pi} d\psi + \frac{1}{\varepsilon} \left(1 - \frac{m \lambda \sin \mu}{h} \right) \int_0^{2\pi} \frac{d\psi}{1 - h \cos \psi} + \frac{1}{\varepsilon} \int_0^{2\pi} \frac{d\psi_1}{1 - h \cos \psi} \frac{m \lambda \sin \psi \cos \mu}{1 - h \cos \psi} + \frac{1}{\varepsilon} \int_0^{2\pi} \frac{m h \sin \psi d\psi}{(1 - h \cos \psi)^2}$$

Les deux dernières intégrales s'annulent et l'on trouve, en définitive :

$$3,33-8 \quad T = \frac{2\pi}{\varepsilon} \frac{1 + \frac{\varepsilon c}{a} \sin(\alpha - \gamma) (1 - \sqrt{1 - h^2})}{\sqrt{1 - h^2}}$$

La pulsation correspondante est :

$$3,33-9 \quad \delta \omega = \frac{-\varepsilon \sqrt{1 - h^2}}{1 + \frac{\varepsilon c}{a} \sin(\alpha - \gamma) (1 - \sqrt{1 - h^2})}$$

Comme la pulsation de la perturbation est $\omega_0 + \varepsilon$ il en résulte pour la pulsation de l'oscillateur perturbé une valeur :

$$3,33-10 \quad \omega = \omega_0 + \varepsilon \left[\frac{(1 - \sqrt{1 - h^2}) \left(1 + \frac{\varepsilon c}{a} \sin(\alpha - \gamma) \right)}{1 + \frac{\varepsilon c}{a} \sin(\alpha - \gamma) (1 - \sqrt{1 - h^2})} \right]$$

le terme entre crochets représente le trainage de fréquence bien connu. Il représente une fraction de ε qui tend à s'annuler si h est nul, et à atteindre l'unité si h tend vers 1.

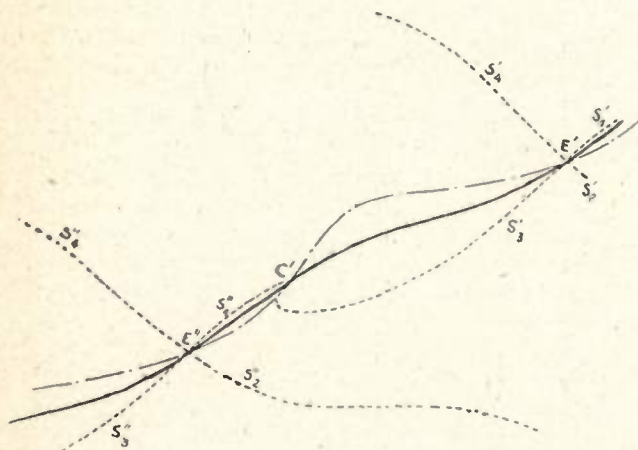


FIG. 4

3,4 Cas où h est supérieur à l'unité (Fig. 4).

3,41 — La courbe (R) rencontre la courbe (Q), en une série de points ; ceux-ci sont des points singuliers, définis par :

$$3,41-1 \quad r = m \varepsilon (1 - h \cos \psi) = 0$$

En ces points, on a :

$$3,41-2 \quad \cos \psi = \frac{1}{h}$$

$$3,41-3 \quad \frac{dq}{d\psi} = 1 + k \sin(\psi - \delta)$$

$$3,41-4 \quad \frac{dr}{d\psi} = \pm m \varepsilon \sqrt{h^2 - 1}$$

Cette quantité restera généralement petite devant l'unité ; l'équation caractéristique de 3,1-10, au voisinage des points singuliers, aura donc son discriminant positif. Les points singuliers seront, dès lors, alternativement des cols et des nœuds stables, selon les définitions de notre précédent mémoire [15].

La seule disposition possible, avec cette hypothèse, est celle de la figure 4.

3,42 — La variation d'amplitude ξ donnée par 2,2-5, est alors définie par la formule :

$$3,42-1 \quad A \xi = u \sin(\theta - \alpha) - b \varepsilon \cos(\alpha - \beta)$$

la formule 3,41-2 peut, avec ces notations, s'écrire :

$$3,42-2 \quad 0 = u \cos(\theta - \alpha) - b \varepsilon \sin(\alpha - \beta)$$

On élimine aisément θ entre ces deux opérations, pour aboutir à la relation :

$$3,42-3 \quad u^2 = a^2 \xi^2 + 2ab\xi\varepsilon \cos(\alpha - \beta) + b^2\varepsilon^2$$

Si ε est donné, cette équation définit une hyperbole ; en portant en abscisse $\frac{u}{b\varepsilon}$ et en ordonnées $\frac{a\xi}{b\varepsilon}$, l'hyperbole sera équilatère (fig. 5).

Nous avons montré, par ailleurs, en collaboration avec M. LOEB [9] que ce régime correspond à celui de la synchronisation de l'oscillateur avec la perturbation : la condition de cette synchronisation est l'existence d'une racine à l'équation 3,43-2, ce qui exige que le rapport de l'amplitude relative u de la perturbation à son écart de fréquence soit supérieur à un seuil qui vaut, avec les notations adoptées ici, $b \sin(\alpha - \beta)$.

3,43- Au voisinage immédiat du nœud stable, nous savons pouvoir remplacer 3,1-10 par l'équation :

$$3,43-1 \quad m\ddot{\psi} + (1 + k \sin(\psi_0 - \delta))\dot{\psi} + \varepsilon h \sin \psi_0 \psi = 0$$

Rappelons les conditions de stabilité déjà données : les trois coefficients de cette équation linéaire du second ordre à coefficients constants doivent être de même signe, à savoir (puisque k étant petit, le second terme est forcément positif) :

$$3,43-2 \quad m = \frac{c \sin(\gamma - \beta)}{a \sin(\gamma - \beta)} > 0$$

$$3,43-3 \quad \varepsilon h \sin \psi_0 = \frac{u \sin(\theta_0 - \alpha)}{b \sin(\alpha - \beta)} = \frac{a\xi + b\varepsilon \cos(\alpha - \beta)}{b \sin(\alpha - \beta)} > 0$$

Rappelons que, pour les systèmes analytiques, $\gamma - \beta = \frac{\pi}{2}$. La condition 3,43-2 est alors équivalente

à la condition donnée par M. LÆB dans son étude [5].

Même dans les systèmes non analytiques, $\sin(\gamma - \beta)$ restera généralement positif ; il devra donc en être de même pour $\sin(\alpha - \beta)$. La condition 3,43-3 s'écrit alors :

$$3,43-4 \quad \varepsilon \left(\frac{a\xi}{b\varepsilon} + \cos(\alpha - \beta) \right) > 0$$

et montre que le point figuratif de la figure 5 doit être sur une demi-branche supérieure, ou inférieure, selon que ε est positif ou négatif : d'autre part, b et u étant positifs par définition, on devra utiliser la branche de gauche ou la branche de droite, selon que ε est positif ou négatif. Il en résulte que seules les demi-branches renforcées sur la figure doivent être considérées comme stables.

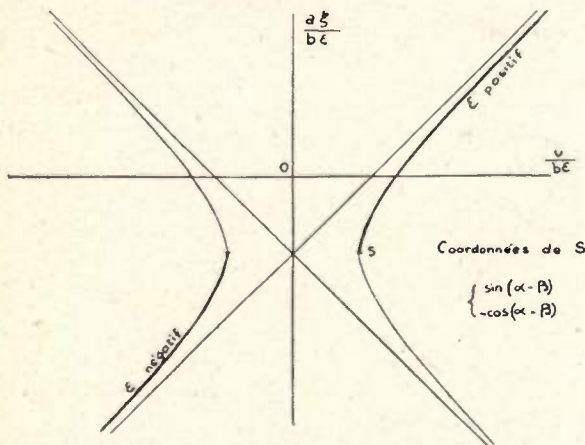


Fig. 5

3,44- Toujours au voisinage du centre d'équilibre, on notera que le coefficient d'amortissement est donné, en valeur absolue, par l'expression :

3,44-1

$$\frac{1 + k \sin(\psi_0 - \delta)}{m} = \frac{a \sin(\alpha - \beta)}{c \sin(\gamma - \beta)} \left(1 + \frac{uc \sin(\theta_0 - \gamma)}{ab \sin(\alpha - \beta)} \right)$$

Cette question d'amortissement est étudiée, plus en détail, dans le mémoire [12].

4. — DÉTERMINATION EXPÉRIMENTALE DES CARACTÉRISTIQUES.

4,1- Comme nous l'avons vu, nous avons besoin de connaître, au voisinage du régime d'équilibre λ_0 , ω_0 pour lequel $y_0 = x_0$ et $\varphi_0 = \pi$, les six dérivées partielles de y et de φ par rapport à l'amplitude x , à la pulsation ω et la dérivée logarithmique à l'amplitude, $\frac{x}{x} = \tau$.

4,2- La mesure des dérivées partielles par rapport à x et à ω est aisée : si nous appliquons à l'entrée de

l'amplificateur, en circuit ouvert, une oscillation sinusoïdale :

$$X = x \cos \omega t$$

nous recueillons à la sortie une oscillation sinusoïdale

$$Y = y \cos(\omega t + \varphi)$$

et il suffit de faire varier d'abord x , ω conservant la valeur ω_0 , puis ω en maintenant à x la valeur x_0 ; la mesure de y et de φ dans les divers cas nous fournira les valeurs de u , x , b , et β et nous donnera en même temps des indications sur le domaine de linéarité.

4,3- Pour avoir c et γ , il nous faut connaître les dérivées partielles par rapport à τ . A cet effet, plutôt que d'appliquer une excitation du type $\exp(-\tau t)$, nous appliquerons une excitation modulée :

$$4,3-1 \quad X = x_0 \cos \omega_0 t (1 + u \cos \mu t)$$

la pulsation μ étant faible par rapport à ω , et le taux de modulation restant petit.

Autrement dit, l'amplitude sera :

$$4,3-2 \quad x = x_0 (1 + u \cos \mu t)$$

On aura alors :

$$4,3-3 \quad y = y_0 + y'_x \lambda_0 u \cos \mu t - y'_\tau u \mu \sin \mu t$$

$$4,3-4 \quad \varphi = \pi + \varphi'_x x_0 u \cos \mu t - \varphi'_\tau u \mu \sin \mu t$$

Ou, compte tenu de nos notations précédentes :

$$4,3-5 \quad \delta y = (1 + a x_0 \sin \alpha) x_0 u \cos \mu t - x_0 c u \mu \sin \gamma \sin \mu t$$

$$4,3-6 \quad \delta \varphi = a x_0 u \cos \alpha \cos \mu t - x_0 c u \mu \cos \gamma \sin \mu t$$

Si l'on peut commodément mesurer l'amplitude et la phase de l'oscillation à la pulsation μ , des grandeurs y et de φ , (φ est la phase à la pulsation ω), on en déduira les valeurs de c et de γ .

On ne pourra, parfois, mesurer commodément que l'amplitude de l'oscillation de y et de φ , à savoir :

4,3-7

$$\left| \delta y \right| = u x_0 \left[(1 + a x_0 \sin \alpha)^2 + x_0^2 \mu^2 c^2 \sin^2 \gamma \right]^{\frac{1}{2}}$$

$$4,3-8 \quad \left| \delta \varphi \right| = u x_0 \left[a^2 \cos^2 \alpha + \mu^2 c^2 \cos^2 \gamma \right]^{\frac{1}{2}}$$

Connaissant, par l'expérience précédente, a et α , on pourra en déduire les valeurs de c et de $\tan \gamma$; le fait que, dans le cas général, γ sera voisin de $\beta + \pi/2$ indiquera la détermination à prendre pour γ parmi les deux valeurs possibles.

5. — PERTURBATIONS D'UN OSCILLATEUR AU VOISINAGE D'UNE FRÉQUENCE DE RÉSONANCE. SEUIL DE SYNCHRONISATION.

Au voisinage d'une résonance, la condition posée au paragraphe 1,5 limite les écarts relatifs de fréquence, susceptibles d'être étudiés par les méthodes précédentes, à des valeurs faibles par rapport à l'in-

verse du coefficient de surtension de l'amplificateur filtrant. Pour peu que celui-ci soit important, le champ d'application de la méthode se trouve très diminué.

Rappelons ici que l'origine de cette limitation se trouve dans la complexité du calcul de la variation du gain et de la phase en régime transitoire ; mais une autre cause de défaillance de la méthode se trouve dans le fait que les formules de développement 1,6-1 sont linéaires par rapport à $\delta\omega$ et, par suite, ne sont pas valables au voisinage d'un maximum de gain.

Si nous nous limitons, toutefois, au cas de l'étude des régimes d'entraînement, qui sont stationnaires, nous pourrions négliger le premier phénomène et essayer d'exploiter les conditions d'entraînement dans un cas au moins (déjà étudié dans la note [1] de M. DUTILH, et dans nos études [9] et [12]) celui où le gain g_{L} de l'amplificateur est le produit d'une fonction $f(x)$ de la seule amplitude, par le gain $\Gamma(j\omega)$ d'un amplificateur linéaire. Comme le régime est permanent, les restrictions rappelées au paragraphe 1,5, concernant la nonanalyticité du gain g_{L} , n'ont plus d'objet.

5,1 — Ainsi, nous pourrions poser :

$$5,1-1 \quad G(x, \omega) = g e^{j\tau} = f(x) \Gamma(j\omega)$$

et, comme nous nous trouvons au voisinage d'une résonance, nous pourrions l'écrire sous la forme

$$5,1-2 \quad \Gamma(j\omega) = - \frac{1 + jQ \left(\frac{\omega_0}{\Omega} - \frac{\Omega}{\omega_0} \right)}{1 + jQ \left(\frac{\omega}{\Omega} - \frac{\Omega}{\omega} \right)}$$

Dans cette formule, Q est le coefficient de surtension ; Ω est la pulsation de résonance ; ω_0 est, comme précédemment, la pulsation de régime en l'absence d'entraînement ; ω la pulsation de la perturbation, c'est-à-dire du régime entraîné. La fonction $f(x)$ est le facteur de non-linéarité, égal à l'unité pour l'amplitude x_0 du régime non entraîné. Il est commode d'utiliser, en outre, les notations suivantes :

$$5,1-3 \quad \left\{ \begin{array}{ll} \delta = \omega_0 - \Omega & \varepsilon = \omega - \omega_0 \\ b = \frac{2Q}{\Omega} & z = \frac{x - x_0}{x_0} \\ a = -x_0 \frac{df}{dx} \end{array} \right.$$

Dans ces conditions, et sous réserve que ε et δ restent petits devant ω_0 (mais non forcément devant $\frac{\omega_0}{Q}$), nous aurons

$$5,1-4 \quad \Gamma(j\omega) = \frac{1 + j b \delta}{1 + j b (\varepsilon + \delta)}$$

5,2 — Le raisonnement se conduit ensuite comme précédemment.

Si l'on applique à l'amplificateur une excitation

$$5,2-1 \quad X = x e^{j(\omega t + \theta)}$$

on recueille, à la sortie, une réponse :

$$5,2-2 \quad Y = x f(x) \Gamma(j\omega) e^{j(\omega t + \theta)}$$

La condition de bouclage sur la perturbation :

$$5,2-3 \quad U = u x_0 e^{j\omega t}$$

s'exprime comme suit :

$$5,2-4 \quad x e^{j\theta} + x f(x) \Gamma(j\omega) e^{j\theta} = u x_0$$

ceci s'écrit, compte tenu des notations 5,1-3 :

$$5,2-5 \quad x_0 (1 + z) \left[1 - (1 - az) \frac{1 + j b \delta}{1 + j b (\varepsilon + \delta)} \right] e^{j\theta} = u x_0$$

ou encore :

$$5,2-6 \quad x_0 (1 + z) \frac{az + j b (\varepsilon + az \delta)}{1 + j b (\varepsilon + \delta)} e^{j\theta} = u x_0$$

Comme la phase θ nous importe peu, il suffira d'égaliser les modules des deux membres pour obtenir :

$$5,2-7 \quad u^2 = (1 + z)^2 \frac{a^2 z^2 + b^2 (\varepsilon + az \delta)^2}{1 + b^2 (\varepsilon + \delta)^2}$$

Pour trouver le seuil, il suffit de chercher la valeur z_0 qui rend u minimum, donc d'annuler la dérivée du second membre.

Comme z_0 est petit par hypothèse, sa valeur approximative est

$$5,2-8 \quad z_0 = - \frac{b^2 \varepsilon (\varepsilon + a \delta)}{b^2 (\varepsilon^2 + 4a \delta \varepsilon + a^2 \delta^2) + a^2}$$

Nous n'écrirons pas la valeur du seuil, que l'on obtient en substituant la valeur ci-dessus dans 5,2-7.

6. — ENTRAÎNEMENT MUTUEL DE DEUX OSCILLATIONS COUPLÉES.

6,1 — Considérons (figure 6) deux oscillateurs, 1 et 2, non-linéaires, filtrés au sens précédent, constitués par des circuits bouclés, connectés par un organe de couplage C .

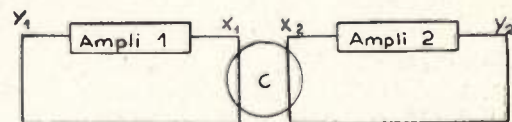


FIG. 6

Nous affecterons respectivement aux grandeurs relatives aux deux oscillateurs des indices 1 et 2, les paramètres sans indices caractérisant des éléments communs, savoir λ , un paramètre de couplage, supposé petit, et $\varepsilon = \omega_1 - \omega_2$ l'écart entre les pulsations ω_1 et ω_2 auxquelles oscillent respectivement les circuits lorsqu'ils sont découplés ($|\varepsilon| \ll \omega_1$)

Les équations des circuits, lorsqu'il y a couplage, sont :

$$6,1-1 \quad \begin{cases} X_1 + Y_1 - \lambda p_2 x_2 = 0 \\ X_2 + Y_2 - \lambda p_1 x_1 = 0 \end{cases}$$

Nous posons :

$$6,1-2 \quad \begin{cases} X_1 = x_1 \cos(\omega_1 t + \theta_1) \\ X_2 = x_2 \cos(\omega_2 t + \theta_2) \end{cases}$$

les phases θ_1 et θ_2 pouvant être variables, dans les conditions que nous avons discutées au début du présent mémoire. Ceci nous conduit à préciser que les équations 6,1-1 sont symboliques, car les coefficients de couplage λp_1 et λp_2 (d'ailleurs égaux si le couplage satisfait au principe de réciprocité) peuvent introduire respectivement des avances de phase η_1 , et η_2 , de telle sorte que $\lambda p_2 X_2$, par exemple, s'écrirait $\lambda p_2 X_2 \cos(\omega_2 t + \theta_2 + \eta_2)$

Nous posons, d'autre part,

$$6,1-3 \quad \zeta = (\omega_1 t + \theta_1) - (\omega_2 t + \theta_2) = \varepsilon t + \theta_1 - \theta_2$$

ce qui permet d'exprimer la différence de « phase » entre X_1 et $\lambda p_2 X_2$, à la pulsation ω_1 , par $\zeta - \eta_2$ et, de même, la différence de phase entre X_2 et $\lambda p_1 X_1$ à la pulsation ω_2 , par $-\zeta - \eta_1$.

6,2 — Plaçons nous d'abord dans le cas où l'écart relatif des fréquences est petit devant l'inverse du coefficient de surtension, de telle sorte que les approximations des paragraphes 1,5 et suivants soient valables. En développant les équations 6,1-1 comme nous l'avons procédé au paragraphe 2, nous trouverons :

$$6,2-11 \quad \begin{cases} \delta g_1 = -\lambda p_1 x_1 \cos(\zeta - \eta_1) & \delta \varphi_1 = \lambda p_1 x_2 \sin(-\zeta \eta_1) \\ \delta g_2 = -\lambda p_2 x_2 \cos(\zeta + \eta_2) & \delta \varphi_2 = -\lambda p_2 x_1 \sin(\zeta + \eta_1) \end{cases}$$

Nous pouvons utiliser les équations 2,2-15 en remplaçant ε par 0, θ par $\zeta - \eta_1$ et $-(\zeta + \eta_2)$ respectivement, et u par $\lambda p_1 x_2$ et $\lambda p_2 x_1$ respectivement pour les deux boucles. On obtient, dès lors, les équations suivantes :

$$6,2-2 \quad \begin{cases} c_1 \sin(\gamma_1 - \beta_1) \ddot{\theta}_1 + a_1 \sin(\alpha_1 - \beta_1) \dot{\theta}_1 \\ = \frac{\lambda p_2 x_2}{b_1 x_1} \left[c_1 \cos(\zeta - \eta_1 - \gamma_1) - a_1 \zeta \sin(\zeta - \eta_1 - \alpha_1) \right] \\ c_2 \sin(\gamma_2 - \beta_2) \ddot{\theta}_2 + a_2 \sin(\alpha_2 - \beta_2) \dot{\theta}_2 \\ = \frac{\lambda p_1 x_1}{b_2 x_2} \left[c_2 \cos(\zeta + \eta_2 + \gamma_2) - a_2 \zeta \sin(\zeta + \eta_2 + \alpha) \right] \end{cases}$$

que nous compléterons par

$$6,2-3 \quad \varepsilon + \dot{\theta}_1 - \dot{\theta}_2 = \zeta$$

En éliminant θ_1 et θ_2 entre ces trois équations, on aboutit à une équation différentielle du 3^e ordre en ζ , que nous n'écrirons pas ; nous nous contenterons de chercher s'il existe une solution synchronisée. Dans ces conditions, $\ddot{\theta}_1$ et $\ddot{\theta}_2$ sont nuls ; $\dot{\theta}_1$, $\dot{\theta}_2$ et

ζ sont constants. Ceci impose, à son tour, que ζ soit constant, donc $\dot{\zeta}$ nul.

On aura alors, en tirant $\dot{\theta}_1$ et $\dot{\theta}_2$ de ces équations et en retranchant le premier du second :

$$6,2-4 \quad \varepsilon = \lambda \left[\frac{p_1 x_1 c_2 \cos(Z + \eta_2 + \gamma_2)}{a_2 b_2 x_2 \sin(\alpha_2 - \beta_2)} - \frac{p_2 x_2 \cos(Z - \eta_1 - \gamma_1)}{a_1 b_1 x_1 \sin(\alpha_1 - \beta_1)} \right]$$

que l'on pourra écrire sous la forme :

$$6,2-5 \quad \varepsilon = \lambda r \cos(Z + \rho)$$

L'existence d'une solution stable sera alors subordonnée à la condition

$$6,2-6 \quad \lambda > \frac{r}{|\varepsilon|}$$

qui définit, comme précédemment, un seuil de synchronisation.

6,3 — Plaçons-nous maintenant, comme nous l'avons fait au paragraphe 5, dans le cas où les deux oscillateurs fonctionnent, au voisinage de leurs fréquences de résonance, avec un coefficient de surtension élevé, et où leur gain est décomposable en produit d'une fonction de l'amplitude seule pour le gain complexe d'un amplificateur linéaire.

Ici, encore, les hypothèses qui ont servi de départ à la discussion du paragraphe 6,2 doivent être modifiées pour étudier les cas où l'écart relatif des pulsations naturelles des deux oscillateurs cesse d'être petit devant les inverses des coefficients de surtension.

Nous adopterons, pour chacun des oscillateurs, les notations 5,1-3, sauf à leur affecter les indices 1 et 2 respectivement. Les gains des amplificateurs linéaires sont donnés par la formule 5,1-4. Nous désignerons les coefficients de couplage par $\lambda p_1 e^{j\psi_1}$ et $\lambda p_2 e^{j\psi_2}$. Dès lors, compte tenu de la formule 5,2-6, les équations de couplage 6,1-1 s'écrivent :

$$6,3-1 \quad \begin{cases} e^{j\theta} x_1 (1 + z_1) \frac{a_1 z_1 + j b_1 (\varepsilon_1 + a_1 z_1 \delta_1)}{1 + j b_1 (\varepsilon_1 + \delta_1)} \\ = \lambda p_2 e^{j\psi_2} x_2 (1 + z_2) \\ e^{-j\theta} x_2 (1 + z_2) \frac{a_2 z_2 + j b_2 (\varepsilon_2 + a_2 z_2 \delta_2)}{1 + j b_2 (\varepsilon_2 + \delta_2)} \\ = \lambda p_1 e^{j\psi_1} x_1 (1 + z_1) \end{cases}$$

θ est la phase relative de la première oscillation par rapport à la seconde, à la pulsation commune ω du régime synchronisé.

Les deux équations complexes 6,3-1 représentent quatre équations réelles qui, avec l'équation des fréquences :

$$6,3-2 \quad \varepsilon = \omega_1 - \omega_2 = \varepsilon_2 - \varepsilon_1$$

doivent permettre le calcul des cinq inconnues :

$z_1, z_2, \varepsilon_1, \varepsilon_2$ et θ .

En éliminant θ , les quatre équations en z_1, z_2, ε_1 et ε_2 sont, si a_1 et a_2 sont réels :

6,3-3

$$\left\{ \begin{aligned} x_1^2 (1 + z_1)^2 [a_1^2 z_1^2 + b_1^2 (\varepsilon_1 + a_1 z_1 \delta_1)] \\ = \lambda^2 p_2^2 x_2^2 (1 + z_2)^2 [1 + b_1^2 (\varepsilon_1 + \delta_1)^2] \\ x_2^2 (1 + z_2)^2 [a_2^2 z_2^2 + b_2^2 (\varepsilon_2 + a_2 z_2 \delta_2)^2] \\ = \lambda^2 p_1^2 x_1^2 (1 + z_1)^2 [1 + b_2^2 (\varepsilon_2 + \delta_2)^2] \\ \text{arc tg } \frac{b_1 (\varepsilon_1 + a_1 z_1 \delta_1)}{a_1 z_1} - \text{arc tg } b_1 (\varepsilon_1 + \delta_1) - \psi_1 \\ + \text{arc tg } \frac{b_2 (\varepsilon_2 + a_2 z_2 \delta_2)}{a_2 z_2} - \text{arc tg } b_2 (\varepsilon_2 + \delta_2) - \psi_2 = 0 \\ \varepsilon_2 - \varepsilon_1 = \varepsilon \end{aligned} \right.$$

La valeur minimum du coefficient de couplage λ , qui rendra ces équations compatibles, définira le seuil de synchronisation.

Nous n'étudierons, plus en détail, qu'un cas particulier, celui où les deux oscillateurs sont identiques. (On peut alors poser $p_1 = p_2 = 1$).

Les équations 6,3-3 se ramènent alors à :

$$6,3-4 \left\{ \begin{aligned} \lambda^2 &= \frac{a^2 z^2 + b^2 (\varepsilon + a z \delta)^2}{1 + b^2 (\varepsilon + \delta)^2} \\ \frac{b \varepsilon (1 - a z)}{a z + b^2 (\varepsilon + \delta)(\varepsilon + a z \delta)} + \text{tg } \psi &= 0 \end{aligned} \right.$$

Ces formules permettent notamment de calculer la dérive de fréquence qui résulte du couplage, et qui est donnée par la formule

$$6,3-6 \quad b \varepsilon = \frac{\lambda \sin \psi (1 + b^2 \delta^2)}{1 - \lambda (\cos \psi + b \delta \sin \psi)}$$

Elle est assortie d'une variation d'amplitude

$$6,3-7 \quad z = \frac{b \varepsilon}{a} \cdot \frac{\cotg \psi - b (\varepsilon + \delta)}{1 + b \varepsilon \cotg \psi + b^2 \delta (\varepsilon + \delta)}$$

On sait, qu'en principe, le couplage introduit un déplacement de la fréquence d'oscillation. La formule 6,3-6 montre que c'est exact en général, notamment dans le cas où le couplage est purement inductif ($\psi = \pi/2$), avec cette circonstance particulière que le déplacement de fréquence s'annule avec ψ , autrement dit lorsque le couplage n'introduit pas de déphasage : dans ce cas, il n'y a variation que de l'amplitude, donnée par la formule :

$$6,3-8 \quad a z = \pm \lambda$$

Le signe sera négatif ou positif, selon le sens du couplage.

7-7. — CONCLUSIONS.

Les systèmes non linéaires, lorsqu'ils dépendent d'équations différentielles du deuxième ordre, ont

déjà fait l'objet de nombreuses études ; la généralisation de l'emploi de dispositifs électriques ou électromécaniques complexes comprenant parfois des relais rend souhaitable que l'on puisse avoir une idée du comportement de ces dispositifs, alors même qu'aucune forme ne peut être attribuée, a priori, à l'équation différentielle — ou à l'équation fonctionnelle — qui les régit.

Le problème a pu être abordé dans un cas très particulier, mais très important pour l'Ingénieur, à savoir celui où la présence d'inerties (électriques ou mécaniques) rend pratiquement sinusoïdale la réponse du système, malgré sa non-linéarité : on a, successivement, étudié les variations de régime du système en présence d'une faible perturbation imposée, et le couplage de deux oscillateurs.

Deux variantes sont à considérer :

Ou bien le système a une courbe de réponse en fréquence relativement aplatie au voisinage du régime d'oscillation non perturbé. Dans ce cas, on peut étudier le régime perturbé, en admettant que les petites variations du gain et de la phase de l'amplificateur sont linéaires, en fonction des trois caractéristiques de l'oscillation (amplitude, fréquence, dérivée logarithmique de l'amplitude).

Ou bien on se trouve au voisinage d'une pointe de résonance de l'amplificateur, ce qui conduit à définir ses caractéristiques uniquement par sa fréquence de résonance, par sa fréquence naturelle d'oscillation en circuit bouclé, par son coefficient de surtension et par son « coefficient de non linéarité » (dérivée du gain par rapport à l'amplitude).

On retrouve des résultats classiques :

— Existence d'un seuil de synchronisation, soit que l'on applique à un oscillateur une perturbation de fréquence imposée, soit que l'on couple deux oscillateurs dont les fréquences propres et les caractéristiques sont différentes ;

— Existence d'un traînage de fréquence, si l'amplitude de perturbation, ou le coefficient de couplage (dans le cas où l'on a deux oscillateurs) est inférieur au seuil.

Mais l'intérêt de la présente méthode est qu'elle ne suppose pas connue l'équation différentielle ou fonctionnelle complète du ou des oscillateurs : nous n'avons fait intervenir dans nos calculs que les caractéristiques mesurables des amplificateurs filtrés au voisinage du régime d'oscillation spontané des systèmes bouclés, ce qui permet d'étendre les conclusions pratiques de l'étude à des systèmes électriques ou mécaniques complexes dont l'étude analytique rigoureuse se heurterait à des difficultés considérables.

Nous tenons, en terminant, à rappeler que les travaux de M. LOEB sont à l'origine de la présente étude et que de nombreuses et fructueuses discussions avec lui, qui ont trouvé leur expression dans des publications communes rappelées dans la bibliographie, ont contribué à l'avancement de la question.

BIBLIOGRAPHIE

- [1] DUTILH, Théorie des servo-mécanismes non linéaires, *Radio* (Fr.), mai 1950, n° 5, pp. 1-7.
- [2] KOCHENBURGER, Analysing contactor servomechanisms by frequency response methods *Electrical Engineering* (U.S.A.) août 1950, 69, n° 8, pp. 687-692.
- [3] LOEB, Un critérium général de stabilité des servomécanismes liés de phénomènes héréditaires (*C.R. Ac. Sc. Fr.*, t. 233, pp. 344-345, séance du 30/7/51).
- [4] LOEB, Analyse harmonique des servomécanismes non-linéaires (*C.R. Ac. Sci. Fr.* t. 233, pp. 733-735, séance du 1^{er} octobre 1951).
- [5] LOEB, Phénomènes héréditaires dans les servomécanismes ; un critérium général de stabilité. (*Annales des Télécommunications*, décembre 1951, t. 6, p. 346).
- [6] VAN DER POL, Oscillations sinusoïdales et de relaxation (*Onde Electrique*, 1930, t. IX, pp. 245 et 293).
- [7] LIENARD, Etude des oscillations entretenues (*Revue Générale d'Electricité*, 26 mai et 2 juin 1928, t. XXIII, pp. 901 et 946).
- [8] CAHEN, Etude topologique de certaines équations différentielles non-linéaires (*C.R. Ac. Sci. Fr.*, t. 235, p. 1003, 3 novembre 1952).
- [9] CAHEN et LOEB, Entrainement des oscillateurs non-linéaires filtrés (*Annales des Télécommunications*, 7, 10 octobre 1952, p. 411).
- [10] LOEB, Transitoires dans les servo-mécanismes non-linéaires filtrés (*Annales des Télécommunications*, 7, 10 octobre 1952, nota p. 408).
- [11] CARSON et FRY, Variable Frequency electric circuit theory with application to the theory of frequency modulation (*Bell Syst. Tech. J.*, U.S.A., octobre 1937, 16, n° 4, p. 513).
- [12] CAHEN et LOEB, Calcul de l'amortissement dans les oscillateurs filtrés (*Annales des Télécommunications*, 8 mars 1953, p. 97).
- [13] CAHEN, Perturbations des oscillateurs filtrés (*C.R. Ac. Sci. Fr.*, t. 235, p. 1.614, 22 décembre 1952).
- [14] CAHEN, Perturbations des oscillateurs filtrés II (*C.R. Ac. Sci. Fr.*, t. 236, p. 356, 26 janvier 1953).
- [15] CAHEN, Intégration géométrique d'équations différentielles non-linéaires du second ordre, avec second membre. (*Société Française des Electriciens — Société des Radioélectriciens — Séance commune* du 15 juin 1953).

ANNEXE

RAPPEL DES CONDITIONS
DE FONCTIONNEMENT D'UN SYSTÈME
LINÉAIRE
EN RÉGIME OSCILLATOIRE
NON « PSEUDO-STATIONNAIRE »

Appliquons à un réseau linéaire, qui peut être actif, caractérisé par un gain complexe $G(p)$, (avec $p = \tau + j\omega$) une excitation :

$$A-1 \quad X = e^{j(\omega_0 t + \theta)} = K e^{j\omega_0 t}$$

θ et K sont supposés variables et, éventuellement, complexes, ce qui permet de représenter par la partie imaginaire de θ le logarithme, changé de signe, de l'amplitude.

CARSON et FRY ont montré [11] que la réponse du réseau linéaire à l'excitation X est donnée par

$$A-2 \quad Y = \left[G_0 K + \sum_{n=1}^{\infty} \frac{1}{n!} \frac{d^n G_0}{d p^n} \frac{d^n K}{d t^n} \right] e^{j\omega_0 t}$$

l'indice 0 indiquant que G , et ses dérivées, sont calculés par $p = j\omega_0$. CARSON et FRY en ont déduit la valeur suivante du gain :

$$A-3 \quad G = G_0 + \sum_{n=1}^{\infty} \frac{d^n G_0}{d p^n} D_n$$

les coefficients D_n étant donnés par les formules de récurrence symboliques suivantes :

$$A-4 \quad \left\{ \begin{array}{l} D_1 = j\dot{\theta} \\ D_n = \frac{1}{n} \left(j\dot{\theta} + \frac{d}{dt} \right) D_{n-1} \end{array} \right.$$

Dans tout ce qui suit, nous supposons que l'écart entre la fréquence « instantanée » $\bar{\omega}$ et la fréquence de régime ω_0 reste petit devant cette dernière ; sous une autre forme, la dérivée logarithmique de l'amplitude, τ et la vitesse de variation de la phase, $\dot{\theta}$ sont petits devant ω_0 ; nous allons examiner, dans le cas d'un système linéaire, quelle simplification cette hypothèse nous permet d'apporter à la formule A-3.

D'une manière générale, on sait que le gain d'un amplificateur sera une fraction décroissante de la fréquence, dès que celle-ci devient élevée ; le gain complexe, s'il peut s'exprimer par une fraction rationnelle, aura donc un dénominateur de degré supérieur au numérateur.

Faisons d'abord abstraction des cas de résonance ; les dérivées $\frac{d^n G_0}{d p^n}$ seront donc des fractions de degré inférieur (algébriquement) à $-n$, de l'ordre de grandeur de ω_0^{-n} .

Or, les coefficients D_n sont de l'ordre de grandeur de $\dot{\theta}^n$. Le terme de rang n de A-3 sera donc de l'ordre de grandeur de $\left(\frac{\omega_0}{\dot{\theta}} \right)^n$.

Il pourrait n'en pas être de même dans le cas d'une résonance. Considérons donc un circuit résonant, le gain unité d , de coefficient de surtension Q . Le tableau ci-dessous donne les valeurs des premiers coefficients de la formule A-3. Des résultats analogues seraient obtenus avec d'autres expressions du gain :

$G(p)$	$\frac{p}{\omega_0 Q (1 + Q^{-1} \omega_0^{-1} p + \omega_0^{-2} p^2)}$
$G(\omega_0)$	1
$\frac{d G_0}{d p}$	$-2 \frac{\omega}{Q}$
$\frac{d^2 G_0}{d p^2}$	$8 \left(\frac{Q}{\omega_0} \right)^2 \left(1 - j \frac{Q^{-1}}{1} \right)$

Et ainsi de suite, le terme de rang n de la formule A-3 étant de l'ordre de $\left(\frac{Q \dot{\theta}}{\omega_0} \right)^n$. Nous pourrions donc négliger les termes de rang élevé, toutes les fois que $\frac{\dot{\theta}}{\omega_0}$ est petit devant Q^{-1} .

Sans prétendre à la rigueur de cette démonstration, nous retiendrons la règle précédente pour définir le domaine de l'applicabilité de nos hypothèses.

VIE DE LA SOCIÉTÉ

RÉUNION DU BUREAU

Le bureau s'est réuni le mercredi 16 décembre 1953 sous la présidence de M. P. DAVID, *Président de la Société des Radioélectriciens*.

Etaient présents :

MM. AUBERT, BOUTHILLON, BUREAU, de MARE, FROMY, MATRAS, PICAULT, RABUTEAU, TESTEMALE.

Etaient excusés :

MM. CABESSA, LIBOIS, RIGAL.

Au cours de la séance les principaux points suivants ont été examinés :

- 1^o Organisation du congrès sur les procédés d'enregistrement sonore et leur extension à l'enregistrement des informations ;
- 2^o Renouvellement du bureau et du conseil ;
- 3^o Relations avec l'Association des Ingénieurs Electroniciens (en présence de M. RAYMOND, président de cette association et convoqué spécialement).

RÉUNIONS EN SORBONNE

Réunion du samedi 28 novembre 1953.

Au cours de cette séance présidée par M. P. DAVID, M. TANTER, *Chef de Service du Laboratoire Central des Télécommunications* fit un exposé sur le Récepteur L.C.T. de radar à élimination des échos sur obstacles fixes ».

M. TANTER indique tout d'abord que le Récepteur L.C.T. de radar à élimination des échos sur obstacles fixes est un récepteur de radar 10 cm comportant essentiellement des circuits assurant la discrimination entre les échos d'obstacles fixes et mobiles et des circuits de soustraction dont l'élément de retard est constitué par une ligne à retard à propagation d'ultra-sons dans le mercure.

Un des premiers modèles de ce récepteur, actuellement construit en série a été livré à l'Administration Française, qui a procédé à des essais complets.

Le conférencier expose ensuite les résultats expérimentaux qui sont conformes aux prévisions théoriques.

Cette conférence a été illustrée par des projections.

Le texte de la communication de M. TANTER paraîtra dans un prochain numéro de l'Onde Electrique.

Réunion du samedi 5 décembre 1953

Cette séance présidée par M. P. DAVID était consacrée à un exposé de M. Charles J. HIRSCH, *Ingénieur en Chef du Laboratoire de recherches de HAZELTINE ELECTRONICS CORPORATION* sur les principes du système N.T.C.S. de télévision en couleurs.

M. HIRSCH rappelle tout d'abord que le NATIONAL TELEVISION SYSTEM COMMITTEE (N.T.S.C.) vient de soumettre à l'approbation du Gouvernement des Etats Unis un nouveau système compatible de Télévision en couleurs, dont l'élaboration est le résultat de trois années de collaboration volontaire des meilleurs ingénieurs des Etats-Unis.

Le conférencier, qui a participé activement à ces travaux expose ensuite les principes suivis :

- transmission séparée et simultanée de l'information en noir et blanc (signal de *Luminance*) et de l'information de couleur à *Luminance* constante (signal de *Chrominance*) ;
- invisibilité du signal de *Chrominance* sur les récepteurs en noir et blanc ;
- séparation par détection synchrone de la double information, contenue dans le signal de *Chrominance* ;
- possibilité de réduire considérablement la bande passante affectée à ce signal,

Cette conférence a été illustrée par des projections.

ACTIVITÉ DES SECTIONS

Première section. — Etudes générales.

Le groupe de Mathématiques appliquées à la Radioélectricité avait organisé le mercredi 15 novembre sous la présidence de M. l'Ingénieur Militaire en Chef ANGOT une séance, au cours de laquelle M. R. CAZENAVE, *Docteur ès Sciences, Ingénieur de la Société L.T.T.* fit un exposé sur :

« Intégrales et Fonctions elliptiques usuelles »

Dans la première partie de sa communication, Mr. CAZENAVE donne un aperçu d'ensemble sur la théorie des fonctions elliptiques vu du point historique et rappelle :

- l'origine des intégrales elliptiques des deux premières espèces,
- l'intégrale elliptique de troisième espèce de Legendre,
- le développement de la théorie des fonctions elliptiques et *théa* de JACOBI,
- le passage à la fonction elliptique de WEIERSTRASS.

La deuxième partie est relative à la représentation géométrique dans le plan des intégrales elliptiques de Legendre de première et seconde espèces.

Ces deux intégrales représentent l'aire en coordonnées polaires d'un secteur de courbe relatif à une quartique pour l'intégrale de première espèce, à une biquadratique (inverse de la quartique par rapport au cercle-unité) pour l'intégrale de seconde espèce.

Enfin, dans la troisième partie, le conférencier donne des indications sur la représentation géométrique de $sn u$, $cn u$, $dn u$ et u sur la sphère unité.

Dans un triangle sphérique rectiligne de côtés $a, b, c = \pi/2$ et d'angles A, B, C on a $sn u = \sin a$, $cn u = \cos a$ et $dn u = \cos A$, tandis que u est l'arc de spirale de Seiffert partant du point-origine de l'équateur.

Sixième section. — Electronique.

Au cours de la séance du lundi 14 décembre 1953, présidée par M. CAZALAS, M. N'GUYEN-THIEN-CHI *Chef de laboratoire à la Compagnie Générale de Télégraphie sans Fil, Ingénieur Conseil à la Compagnie Industrielle des Métaux Electroniques* fit un exposé sur le « Molybdène ».

Le conférencier développe successivement les points suivants :

- Généralités sur les matériaux électroniques dont le molybdène est l'un des plus importants ;
- Elaboration du molybdène *ductile* par les techniques spéciales de la Métallurgie des Poudres : analogie avec celle du tungstène ;
- Bref aperçu sur le molybdène fondu à l'arc ;
- Propriétés du molybdène (physiques, mécaniques, thermiques, électriques, etc...) ;
- Emplois du molybdène en électronique (électrode, scelléments verre-métal, alliages $W - Mo$, etc...) ;
- Quelques autres applications du molybdène (contacts électriques, alliages réfractaires, etc...).

L'exposé fut illustré par des projections.

Septième section (Documentation) et deuxième section (Matériel radioélectrique).

Ces deux sections avaient organisé en commun le lundi 7 décembre, une séance sous la présidence de M. L. CAHEN pour la

7^e section et de M. P. LIZON pour la 2^e section, au cours de laquelle M. J. SUCHET Ingénieur à la Société Philips fit une mise au point technique et bibliographique sur :

Les produits céramiques utilisés en Radioélectricité.

Le conférencier rappelle d'abord que l'usage croissant en Radioélectricité de matériaux élaborés par une technique spéciale dite « technique céramique » a suscité durant les dernières années la publication d'un grand nombre de travaux et d'études sur les pièces détachées dont ils permettent la réalisation. Ces articles restreignent toutefois leur point de vue soit aux céramiques pour condensateurs, soit aux céramiques magnétiques, soit aux semi-conducteurs. M. SUCHET s'attache ici à montrer les relations et les analogies existant entre les compositions chimiques, les structures moléculaires et les propriétés électriques ou magnétiques des différentes classes de céramiques. Il donne un tableau des diverses pièces détachées céramiques et passe en revue leurs principaux domaines d'utilisation. Un grand nombre de références bibliographiques postérieures à 1950 sont citées.

Cette mise au point technique et bibliographique paraîtra dans un prochain numéro de l'Onde Electrique.

Outre les communications sur les sujets énumérés ci-dessus, des conférences de mise au point sur des sujets cristallographiques d'intérêt général sont prévues.

Expositions. — Une exposition de matériel et d'ouvrages scientifiques aura lieu pendant la durée du Congrès.

Excursions. — Des excursions et visites seront organisées à l'intention des personnes accompagnant les congressistes.

Colloques. — Deux colloques sont prévus sur les sujets suivants :

1. Localisation de l'atome hydrogène et liaison hydrogène.
2. Mécanisme des changements de phases dans les cristaux.

Les colloques se tiendront le 29 et le 30 juillet ; des séances préliminaires auront lieu entre le 21 et le 28 juillet.

— Pour tous renseignements s'adresser au Secrétaire du Congrès :

A. J. ROSE. *Laboratoire de Minéralogie*, 1, rue Victor-Cousin, Paris-5^e.

CONGRÈS INTERNATIONAL SUR LES PROCÉDÉS D'ENREGISTREMENT SONORE ET LEUR EXTENSION A L'ENREGISTREMENT DES INFORMATIONS

¶ Ainsi que nous l'avons annoncé dans nos précédents numéros le Congrès sur les procédés d'enregistrement Sonore et leur extension à l'enregistrement des Informations se tiendra du 5 au 10 Avril 1954 à la Maison de la Chimie, 28, rue Saint-Dominique, Paris (7^e).

Participeront à ce congrès, 70 conférenciers, 250 auditeurs et 30 exposants environ parmi lesquels les représentants les plus qualifiés de l'Industrie Cinématographique et Radioélectrique, du Centre National de la Cinématographie, du Groupement des Acousticiens de Langue Française (GALF), de la Radiodiffusion-Télévision Française, ainsi que de nombreuses personnalités étrangères.

Pendant toute la durée du Congrès et pendant la journée du Dimanche 11 Avril, l'exposition de matériel installée dans le hall de la Maison de la Chimie sera ouverte au public.

Pour tous renseignements complémentaires s'adresser Société des Radioélectriciens, 10, avenue Pierre-Larousse à Malakoff (Seine), téléphone ALE. 04-16.

INFORMATIONS

60^e anniversaire de l'Ecole Supérieure d'Electricité (E. S. E.)

Notre président M. DAVID a reçu une convocation pour assister aux séances du Comité d'Organisation des Cérémonies de commémoration du 60^e anniversaire de l'E.S.E.

Troisième Congrès International de Cristallographie.

L'Union Internationale de Cristallographie, 1, rue Victor Cousin, Paris (5^e) organise le troisième congrès de cristallographie à la Faculté des Sciences de Paris du 21 au 28 juillet 1954.

Programme : Les sujets suivants seront traités :

1. Appareillage et techniques,
2. Progrès récents dans la détermination des structures,
3. Structures des minéraux, y compris les minéraux synthétiques et les céramiques,
4. Structures des métaux et alliages,
5. Structures non organiques,
6. Structures organiques,
7. Structures des protéines et structures des composés analogues,
8. Ordre, désordre et déformation dans les structures cristallines,
9. Liquides et cristaux liquides,
10. Verres,
11. Transformations thermiques,
12. Diffusion en dehors des réflexions sélectives,
13. Croissance des cristaux,
14. Diffraction des neutrons,
15. Divers.

DOCUMENTATION

La bibliothèque de notre Société a reçu un ouvrage édité par MASSON et Cie, sur « La Théorie des Fonctions Aléatoires » (Applications à divers phénomènes de fluctuation) par :

M. A. BLANC LAPIERRE. — Professeur de Physique Théorique à la Faculté des Sciences d'Alger,

M. Robert FORTET. — Professeur à la Faculté des Sciences de Paris, chargé d'un cours de calcul des Probabilités.

avec un chapitre de :
M. J. KAMPÉ DE FÉRIET. — Professeur à la Faculté des Sciences de Lille.

NOUVEAUX MEMBRES

MM.	Présentés par MM.
ABADIE Maurice, Ingénieur en Chef à la Société Alsacienne de Constructions Mécaniques	MATRAS et TESTEMALE.
ANTIER Jean, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.
ASCH Georges, élève à l'E.S.E. (Division Radio)	DAUPHIN et GAUSOT.
BENNETOT Michel de, Ingénieur de l'Ecole Navale à la Cie Gle de T. S. F.	WARNECKE et GUÉNARD.
BEZIE Alexandre, Ingénieur en Chef des Télécommunications	MATRAS et TESTEMALE.
BIGARD Jean, Ingénieur E. C. P., Dr Industriel à la Cie Gle de T. S. F.	WARNECKE et GUÉNARD.
BLAMOUTIER Michel, Ingénieur E.P.C.I.	MANUEL et GUYOT.
BOUTHINON Guy, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.
BOUYEURE Jacques, élève à l'Ecole Centrale de T.S.F.	QUINET et CHRÉTIEN
CAHEN Olivier, Ingénieur au Central National d'Etudes des télécommunications	LAPOSTOLLE et PICQUENDAR.
CONAN Robert, Capitaine de l'Armée de l'Air	COUDEYRE et BURON.
DACOS Fernand, Professeur à l'Institut Montefiore (Belgique)	DAVID et MATRAS.
DESPOITS Gérard, élève à l'E.C.T.S.F.	QUINET et CHRÉTIEN
DIDIER André, Professeur au Conservatoire National des Arts et Métiers	MATRAS et TESTEMALE.
DUCAMUS Jean, Ingénieur des télécommunications Licencié ès-Sciences	RIGAL et MATRAS.
DUCOT Claude, Ingénieur E. P. Chef de Dépt au Laboratoire Electronique de Physique appliquée	ANDRIEUX et CAYZAC.
FEREC Roger, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.

MM.	présentés par MM.	MM.	présentés par MM.
GENIEUX Pierre, Ingénieur Radio diplômé du Centre d'Etudes Techniques de Paris	VALLÉE et LAVAL-LÉE.	MERLET Lucien, Ingénieur en Chef des Télécommunications (R.T.F.)	MATRAS et TESTEMALE.
GILLET Auguste, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.	MOGLIATTO Marcel, élève à l'Ecole Centrale de T.S.F.	QUINET et CHRÉTIEN
GIRAUD Jacques, Ingénieur à la Radio Industrie .	CAMERINI et ZERMIZOGLOU.	OGET Christian, élève à l'E.C.T.S.F.	QUINET et CHRÉTIEN
GLAVANY René, Agent Technique Principal au C.N.E.T.	POINCELOT et BURGEAT.	PETIT Jacques, élève à l'E.C.T.S.F.	QUINET et CHRÉTIEN
GONNEVILLE Olivier de, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.	PIERRE Jacques, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.
GRUMEL Henri, Ingénieur en Chef à la S.F.R. ...	AUBERT et FAGOT.	PORTIER William, élève à l'Ecole Centrale de T.S.F.	QUINET et GRANSARD.
GUILBAUD Georges, Ingénieur E.P.C.I. au Département Recherches Electroniques de la Cie Gle de T.S.F.	WARNECKE et GUÉNARD.	QUEVRIN Janic, élève à l'E.C.T.S.F.	QUINET et CHRÉTIEN.
IATROPOULOS Evangelos, Ingénieur Civil de l'Ecole Nationale Supérieure des Télécommunications	POTIER et FUHRMANN.	RAYBAUD Jean, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.
ITALO INCOLLINGO, Docteur Ingénieur électricien à l'Institut International des Brevets de la Haye	SÉBÉO et BOUTHILLON.	ROUS Jean, élève à l'E.C.T.S.F.	QUINET et CHRÉTIEN
LE BLAINVAUX Armand, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.	SARRETTE Bernard, Ingénieur Radio E.S.E.	DAUPHIN et GAUSSOT
LE COZ Gérard, élève à l'Institut Polytechnique de Grenoble	BENOIT et GRANSARD.	SCHIRMER Jean, Ingénieur à la Cie Gle de T.S.F. .	THIEN-CHI et DUGAS
LESTEL Jacques, Ingénieur au C.N.E.T.	LAPOSTOLLE et PICQUENDAR.	SCHWEITZER Jean, Ingénieur E.P.C.I. à la S.F.R. (Départ. Lampes)	MARTINOFF et PRÉVOST.
		SEUROT Max, élève à l'E.C.T.S.F.	CHRÉTIEN et QUINET
		THEZARD François, élève à l'E.C.T.S.F.	QUINET et CHRÉTIEN
		THILLIEZ Jacques, élève à l'E.S.E. (Division Radio)	VARRET et DAUPHIN
		TON-THAT So, élève à l'E.C.T.S.F.	QUINET et CHRÉTIEN
		TRENTINIAN Jacques de, élève de l'Ecole Supérieure d'Electricité (Division Radio)	DAUPHIN et GAUSSOT.
		TOUSSAINT Maurice, Ingénieur à la S.F.R.	AUBERT et FAGOT.