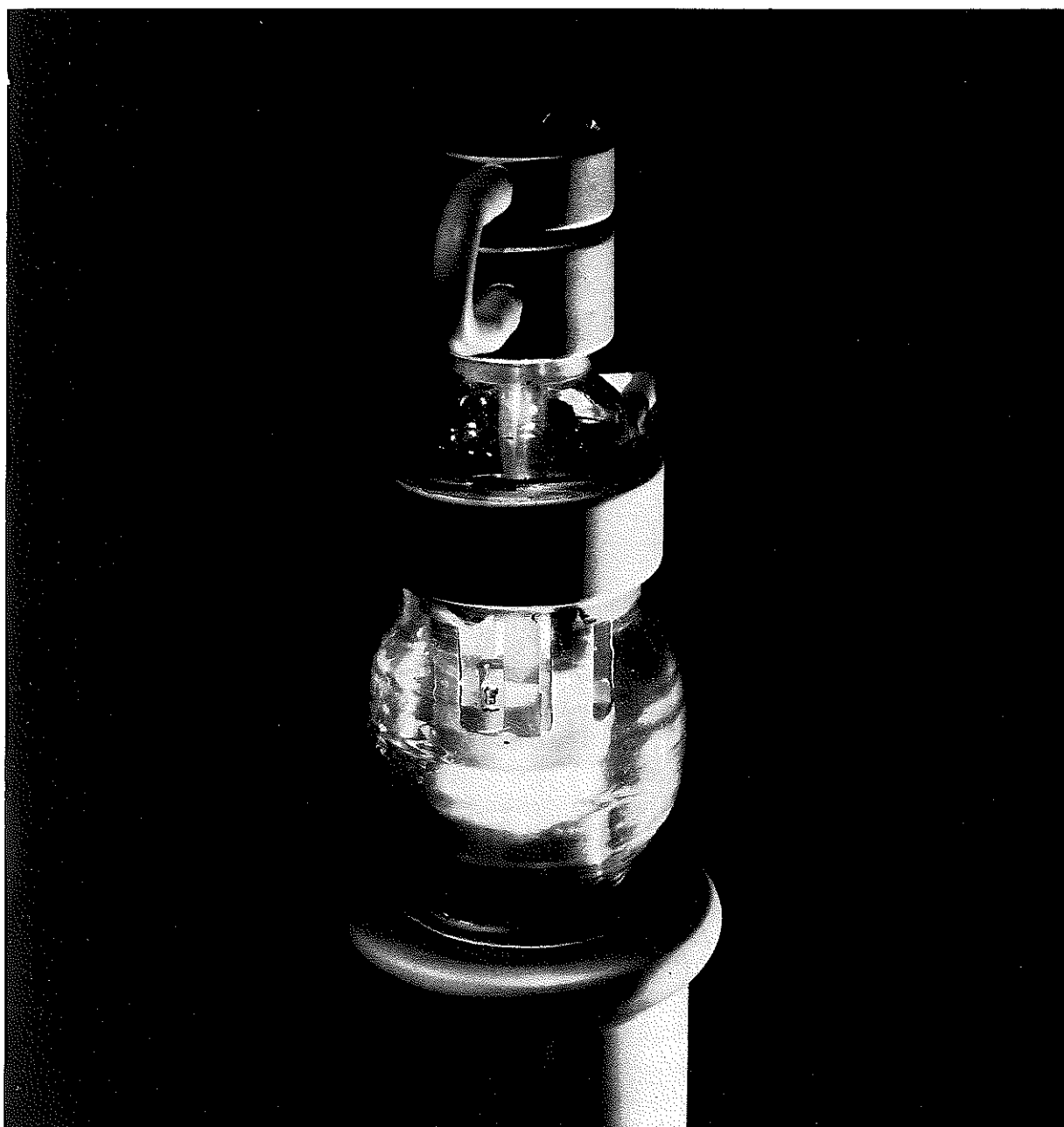


REVUE BROWN BOVERI



Lampe d'émission démontable Brown Boveri de grande puissance.



Numéro spécial: Haute fréquence.

RÉALISATIONS NOUVELLES

dans tous les domaines de la Haute Fréquence

POSTES ÉMETTEURS

à grande puissance, de radiodiffusion et de radiocommunication, équipés de lampes démontables Brown Boveri.

GÉNÉRATEURS

haute fréquence pour buts scientifiques et industriels, à ondes longues et à ondes très courtes.

PETITS APPAREILS

émetteurs-récepteurs pour l'armée, la police, la défense aérienne, etc.

ÉMETTEURS

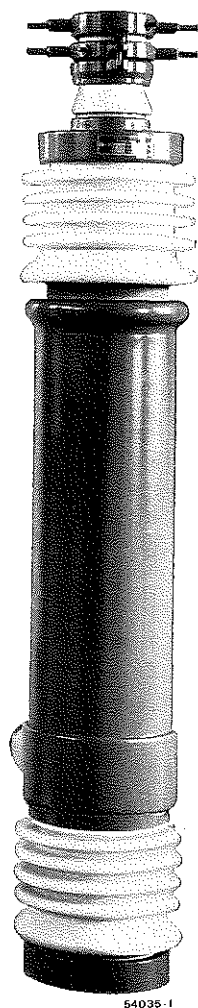
pour le service de sécurité de la navigation aérienne et du trafic maritime.

POSTES fixes et de bord.

TÉLÉCOMMANDE, TÉLÉMESURE ET TÉLÉRÉGLAGE

par courants porteurs, sur fil ou sans fil.

Ce matériel moderne BROWN BOVERI est le fruit de nouvelles recherches et d'un travail de précision réputé.



REVUE BROWN BOVERI

PUBLIÉE PAR LA SOCIÉTÉ ANONYME BROWN, BOVERI & C^{IE} A BADEN (SUISSE)

XXVIII^{me} ANNÉE

DÉCEMBRE 1941

N° 12

La Revue Brown Boveri paraît mensuellement. La reproduction d'articles ou d'illustrations est permise, à condition de citer leur provenance.
Prix de l'abonnement annuel pour la Suisse Fr. 10.—, prix du fascicule pour la Suisse Fr. 1.—, port et emballage non compris.

SOMMAIRE:

	Pages		Pages
Introduction	387	La radio dans l'aviation	414
Lampes démontables de grandes puissances	389	Modulation en fréquence	417
Générateurs d'ondes ultra-courtes	393	Etude théorique du couplage de deux circuits oscillants par l'intermédiaire d'une ligne de transmission	423
Procédés de brouillage automatiques de la parole	397	Nouveaux résultats des recherches en physique atomique	436
Installations modernes de radio pour la police	409	Postes émetteurs de radiodiffusion	446
Équipements portatifs de radio pour l'armée	413		

INTRODUCTION.

La découverte par de Forest et von Lieben des lampes à vide très poussé à grilles polarisées, a donné, avant le conflit mondial de 1914 déjà, un nouvel essor à la technique de la haute fréquence, édifiée sur les idées et les travaux fondamentaux de Maxwell, Hertz et Marconi. Les événements historiques avaient alors accéléré son développement et à la fin de la guerre, la technique était suffisamment perfectionnée pour permettre l'utilisation dans l'économie de paix des améliorations trouvées. Déjà, les progrès réalisés dans le domaine de la télégraphie et de la téléphonie commerciale avaient fortement influencé la période de l'après guerre. Mais cette influence n'est pas comparable à l'action profonde de la radiodiffusion sur la manière de vivre et la mentalité de la majorité des êtres humains. Spécialement à l'époque actuelle, qui apporte de nouveau de grands bouleversements politiques et militaires, on ne peut sous-estimer l'importance du fait que des millions d'hommes sont pour ainsi dire instantanément renseignés sur tous les événements importants et entendent même directement la voix des hommes qui dirigent leur pays.

La Suisse n'a pas pris une part importante au développement de cette technique qui avait même conduit à des domaines encore inexplorés, mais elle s'est contentée de fabriquer des objets ayant un écoulement suffisant dans le pays. Qui

ne connaît pas l'histoire de la technique dans notre pays, attribuerait sans doute cette attitude à l'exiguïté du territoire et à la modicité de ses moyens en hommes et en matières premières. Mais il faudrait alors oublier complètement les travaux remarquables de notre pays dans le domaine du courant fort. Des hommes comme Burgin, Thury et C. E. L. Brown, pour ne citer que les plus marquants, ne se sont pas laissés arrêter par ces conditions défavorables et ont tout mis au service du développement de ce nouveau domaine leur renom, leur travail et leurs moyens financiers. Ils n'ont pas seulement remporté des succès personnels, mais ont surtout rendus des services inestimables à notre pays. Pourquoi, plus tard, n'en fut-il pas de même dans le domaine de la haute fréquence? Parce qu'à côté de l'absence d'une activité maritime, cette technique s'est développée pendant la guerre mondiale et les années suivantes. Nous admettons aussi que les chefs de notre industrie avaient été découragés par les difficultés économiques, toutes différentes de celles rencontrées précédemment, de rechercher une activité nouvelle. Mais notre peuple aussi, dans son ensemble, avait perdu la force créatrice qui doit toujours l'animer et ne se trouvait pas après la guerre dans une disposition d'esprit favorable. Exige que simultanément les conditions de travail

soient améliorées et le travail personnel réduit n'a jamais conduit au succès dans aucun domaine. Nous sommes persuadés que cette mentalité a maintenant disparu et espérons avec confiance qu'il en sera ainsi pour longtemps.

Si nous avons décidé de compenser, dans la mesure de nos forces, la négligence passée, il ne peut plus être question de travaux de créateurs comme ceux des maîtres cités, car toute cette technique est maintenant trop avancée. Il est toutefois indubitable qu'on est encore loin d'avoir épuisé toutes les possibilités de perfectionnement; il existe encore bien des lacunes à remplir, par exemple, dans le domaine des ondes centimétriques. Du reste on n'atteint, en général, jamais un état définitif dans la technique. Chacun est appelé à participer à ces perfectionnements, quels que soient ses moyens, s'il attaque un problème avec courage et possède de claires notions de physique. Ce numéro de notre Revue doit donner quelques renseignements sur la voie que nous avons déjà parcourue et sur celle que nous comptons encore suivre.

Débutant dans un domaine qui nous était entièrement nouveau, deux voies nous étaient simultanément ouvertes. D'une part, nous devions chercher à obtenir assez rapidement des commandes pour prendre contact avec les besoins du marché et pour former notre personnel, composé en grande partie de jeunes gens. D'autre part, nous voulions, conformément aux principes que nous appliquons dans les autres domaines de notre activité, contribuer au perfectionnement de cette technique, en nous attaquant à des problèmes nouveaux. Lors de l'application de ce programme, nous nous sommes aperçu que ces deux activités ne pouvaient pas être nettement séparées. Nous voulons toutefois les considérer ici séparément.

Les postes d'émission offraient un domaine approprié pour nos premières armes, la fabrication des récepteurs ayant déjà été entreprise par quelques maisons du pays. Mentionnons ici l'émetteur à portée limitée de Kloten, en service depuis 1939 et servant au guidage des avions. Les équipements émetteurs-récepteurs que nous avons construits pour l'armée suisse, nous ont donné l'oc-

casion d'introduire des améliorations intéressantes par rapport aux appareils existants spécialement en ce qui concerne la puissance en fonction du poids. La commande de l'installation de radio de la police de Zurich nous offrit une autre occasion de développer une technique nouvelle. Mentionnons nos essais pour assurer une liaison bilatérale de postes mobiles se déplaçant en ville, en utilisant simultanément des ondes moyennes et ultra-courtes. L'utilisation de ce système à deux ondes fut possible grâce à la modulation en fréquence qui n'avait encore été que peu utilisée en Europe. Un travail intensif fut nécessaire pour créer le récepteur portatif dont le dispositif d'appel était une des principales difficultés.

C'est surtout dans la construction des lampes que nous avons cherché de nouvelles voies. Des lampes à grande puissance démontables, dont la cathode peut être remplacée, sont plus économiques pour les grands postes émetteurs que les lampes scellées qui deviennent inutilisables lorsque la cathode est brûlée. Ces lampes ont aussi une propriété intéressante, qui manque aux lampes scellées: celle de permettre une adaptation ultérieure de leur caractéristique à des conditions spéciales de fonctionnement.

Nous avons entrepris la construction des générateurs d'ondes décimétriques et centimétriques. Ainsi qu'on le sait les triodes normales ne conviennent pas dans ce cas. Les difficultés rencontrées jusqu'à présent dans la production de ces ondes ont empêché le développement de la technique qui s'y rapporte. Ces ondes offrent des possibilités d'utilisation très intéressantes pour la transmission de renseignements ou pour la sécurité des transmissions militaires. Nous sommes aussi actifs dans ce domaine.

Parmi les autres travaux de notre laboratoire de recherches, mentionnons encore pour finir les appareils de brouillage pour la téléphonie sans fil, quoiqu'ils n'aient pas de rapport direct avec la haute fréquence proprement dite. Nous croyons que la téléphonie sans fil et, dans certains cas, aussi celle par fil seront susceptibles de nouvelles utilisations si l'on arrive à rendre la parole émise incompréhensible à l'utilisateur d'un

récepteur normal. De tels appareils sont déjà utilisés dans la téléphonie transocéanique, mais ils sont très lourds et ne garantissent pas un secret de communication absolu. Nos travaux nous ont déjà conduit à la construction d'appareils portatifs garantissant un secret équivalent à celui des télégrammes chiffrés. Nous n'avons pas voulu nous arrêter à ce résultat; après de nouvelles études, nous avons posé les bases de la construction d'un nouvel appareil, garantissant le secret d'une façon bien supérieure. Le problème difficile de la synchronisation de l'émetteur et du récepteur avait déjà été résolu sur les anciens appareils. C'est pourquoi nous sommes persuadés de pouvoir mener à bonne fin la réalisation définitive d'un «chiffreur» répondant aux conditions posées.

Nous avons pour terminer la grande satisfaction de pouvoir exprimer ici notre vive recon-

naissance au Professeur Dr P. Scherrer de l'Ecole Polytechnique Fédérale de Zurich pour l'article «Nouveaux résultats des recherches en physique atomique» qu'il a eu l'amabilité de nous adresser comme collaboration à ce numéro de notre Revue. Après avoir construit, comme premier travail de notre département «Haute Fréquence», l'émetteur à ondes courtes de 40 kW et 20 m de longueur d'onde, alimentant le cyclotron de l'Institut de Physique de l'E. P. F., il nous a paru indiqué de faire connaître à nos lecteurs l'importance du cyclotron pour la dissociation des atomes. Le Professeur Dr Scherrer ne s'est pas borné à réaliser ce désir, mais a pénétré fort avant dans le domaine des recherches les plus récentes en physique et en physiologie. Nous sommes reconnaissants à l'auteur de son exposé qui intéressera certainement tous nos lecteurs. (MS 811) Th. Boveri. (J. C.)

LAMPES DÉMONTABLES DE GRANDES PUISSANCES.

Indice décimal 621.396.615.16

La lampe d'émission est l'élément le plus important d'un émetteur. L'auteur décrit dans cet article une nouvelle construction de lampes en métal et matières céramiques, qui sont très robustes et économiques et ont des avantages au point de vue physique. Elles peuvent être utilisées pour diverses puissances, tensions et fréquences. Ce sont les premières lampes de grande puissance, construites en Suisse.

1^o INTRODUCTION.

Les travaux dans ce domaine tout nouveau ont été entrepris au début de 1937. Nous nous étions alors fixé pour but d'équiper nos postes d'émission de lampes de notre fabrication. Il nous fallut tout d'abord choisir entre la construction de lampes scellées normales en verre et métal, sans pompe et celle de lampes démontables robustes en métal et matières céramiques, raccordées en permanence à une pompe à vide élevé. Nous nous sommes décidés pour les lampes démontables. Ce choix nous donna la possibilité d'utiliser notre expérience de la technique du vide, acquise dans la construction des mutateurs et d'arriver en un court délai à construire des lampes d'émission utilisables en se servant de l'outillage existant. Nous sommes même arrivés à faire des essais, couronnés de succès dès l'automne 1939, avec deux lampes de 150 kW, au poste émetteur de Beromünster. Ces lampes servirent même un jour à une brève émission officielle, sans entraîner de perturbations. Cette dernière épreuve

fut rendue difficile, car il fallait adapter exactement les caractéristiques de nos lampes à celles des lampes scellées de l'installation.

C'est aussi pour des raisons économiques que nous avons été conduits à adopter la construction démontable. Les lampes d'émission scellées ont en effet une durée limitée à quelques milliers d'heures, car la cathode se brûle peu à peu. Le remplacement d'une lampe de grande puissance est toujours coûteux et se reproduit périodiquement; une réparation n'est pas économique. Notre construction démontable permet par contre de changer facilement et rapidement la cathode. Les frais ne s'élèvent alors qu'à quelques pour-cent du prix d'achat des lampes, et l'on peut compter sur une durée illimitée de celles-ci.

Nous n'avons toutefois pas perdu de vue la construction sans pompe. Nous avons maintenant trouvé une nouvelle solution qui permet de remplacer les cathodes défectueuses sans toucher aux joints entre les parties métalliques et l'isolation. Il s'agit d'une construction partiellement démontable en métal et en matières céramiques.

La construction de lampes démontables n'est pas nouvelle. En Angleterre, on essayait déjà avant 1920 d'élever, de cette façon à quelques kW, la puissance

unitaire des types en verre qui n'était que d'environ 100 W. Cet essai fut sans succès, car la pompe à vapeur de mercure, que l'on utilisait, avec sa pression minimum de service de 10^{-3} mm de hauteur de colonne de Hg correspondant à la pression des vapeurs de mercure à la température ambiante, était tout à fait insuffisante. Les moyens connus pour abaisser la pression de vapeur à une valeur admissible sont beaucoup trop compliqués. Avec la pression mentionnée ci-dessus la caractéristique de commande ne suit pas, à cause de la ionisation des gaz, la loi de la charge spatiale $I_a = k \cdot U^{\frac{3}{2}}$, mais une autre loi qui entraîne une forte distorsion, rendant les lampes inutilisables pour la téléphonie. Ce n'est que pour des pressions inférieures à 10^{-5} mm de hauteur de colonne de Hg qu'une lampe d'émission fonctionne parfaitement. Quelques années plus tard, les recherches furent reprises en France, en utilisant une pompe moléculaire qui fournit une pression absolue inférieure à 10^{-5} mm de hauteur de colonne de Hg, mais avec un faible débit. Ces essais permirent la construction de lampes d'émission utilisables jusqu'à une dissipation anodique de 10 kW. Le perfectionnement de ces lampes resta arrêté pendant quelques années. On avait, en effet, appris à construire des lampes scellées de grandes puissances et atteint en 1935 une dissipation anodique de 300 kW. Quelques années auparavant la construction des lampes démontables avait reçu une nouvelle impulsion par la découverte de la pompe à diffusion à vapeur d'huile, qui avait une pression de 10^{-5} mm de hauteur de colonne de Hg. On construisit à l'étranger à titre d'essai des lampes d'émission jusqu'à 500 kW de dissipation anodique. Nous sommes actuellement à même, grâce à l'expérience acquise, de construire des lampes d'émission démontables pour toutes puissances, tensions et fréquences que l'on rencontre en pratique.

II° CONSTRUCTION DES LAMPES D'ÉMISSION DÉMONTABLES.

Les figures 1a et 1b qui représentent une lampe démontable et une lampe scellée, permettent, par comparaison, de voir la différence de leur construction. Il s'agit ici de la triode connue avec l'anode cylindrique (1) en cuivre, la grille de commande hélicoïdale (2) en molybdène, et la cathode (3) formée de plusieurs filaments en tungstène. Les anodes sont refroidies par un liquide, la grille et la cathode, en revanche, abandonnent leur chaleur à l'ambiance, surtout par radiation. La cathode ou la grille sont reliées aux brides refroidies (4) et (5). Les joints verre-métal, de la construction 1b, supportent des températures supérieures à 100° , et ne demandent en général aucun refroidissement. Le refroidissement de toutes les brides des lampes démontables est, en revanche, nécessaire afin que la

température des joints entre les pièces métalliques et celles en matière céramique soit inférieure au maximum admissible de 50°C . A l'ouverture (6), au bas de la lampe 1a, est raccordée la pompe à vide élevé, qui maintient continuellement la pression. Les lampes sont des appareils qui doivent satisfaire à toute une série de

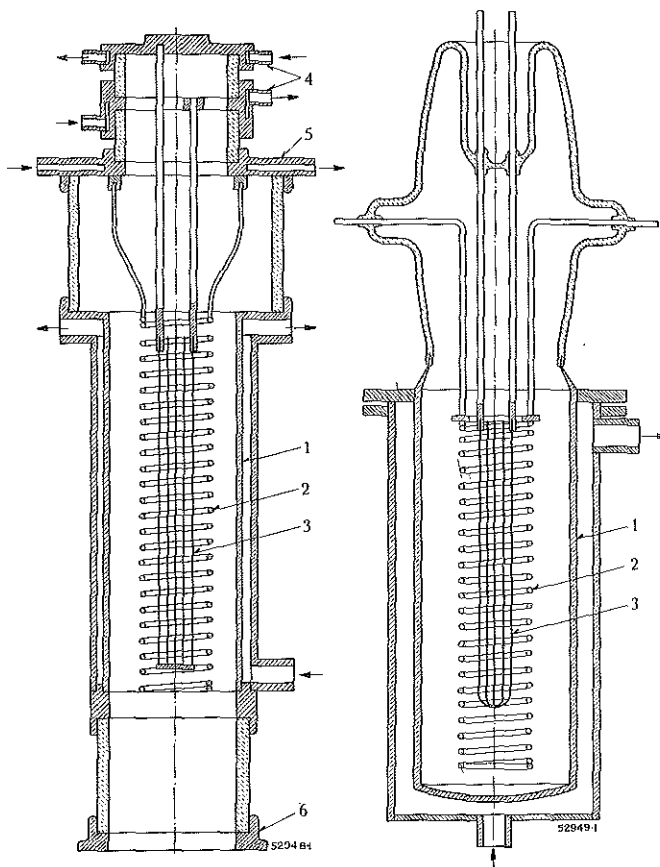


Fig. 1a. — Lampe d'émission démontable, refroidie par eau.

Fig. 1b. — Lampe d'émission scellée, refroidie par eau.

Cette comparaison montre que la construction des lampes facilement démontables en métal-matière céramique, est robuste. Le poids des lampes est à peine plus élevé car l'anode, la pièce importante, peut être plus fortement chargée.

- | | | |
|-----------------------------|-----------------------|-------------------------|
| 1 = Anode. | 3 = Cathode. | 5 = Bride de la grille. |
| 2 = Grille de commande. | 4 = Bride de cathode. | 6 = Bride de la pompe. |
| → = Eau de refroidissement. | | |

conditions. Elles travaillent avec des pressions très basses, inférieures à 10^{-5} mm de hauteur de colonne de Hg. La fréquence de fonctionnement peut atteindre quelques centaines de méga-cycles. Les pièces en matière céramique doivent être parfaitement étanches au vide le plus poussé et présenter de très faibles pertes diélectriques, et pour les très hautes fréquences supporter également des températures de plusieurs centaines de degrés (le quartz satisfait à ces conditions). La température de fonctionnement de la cathode varie entre 2200 et 2300°C . La densité d'énergie W/cm^2 à l'anode est environ cinq fois plus élevée que celle des éléments évaporateurs d'une chaudière Velox. L'énergie électrique amenée à l'anode est transformée

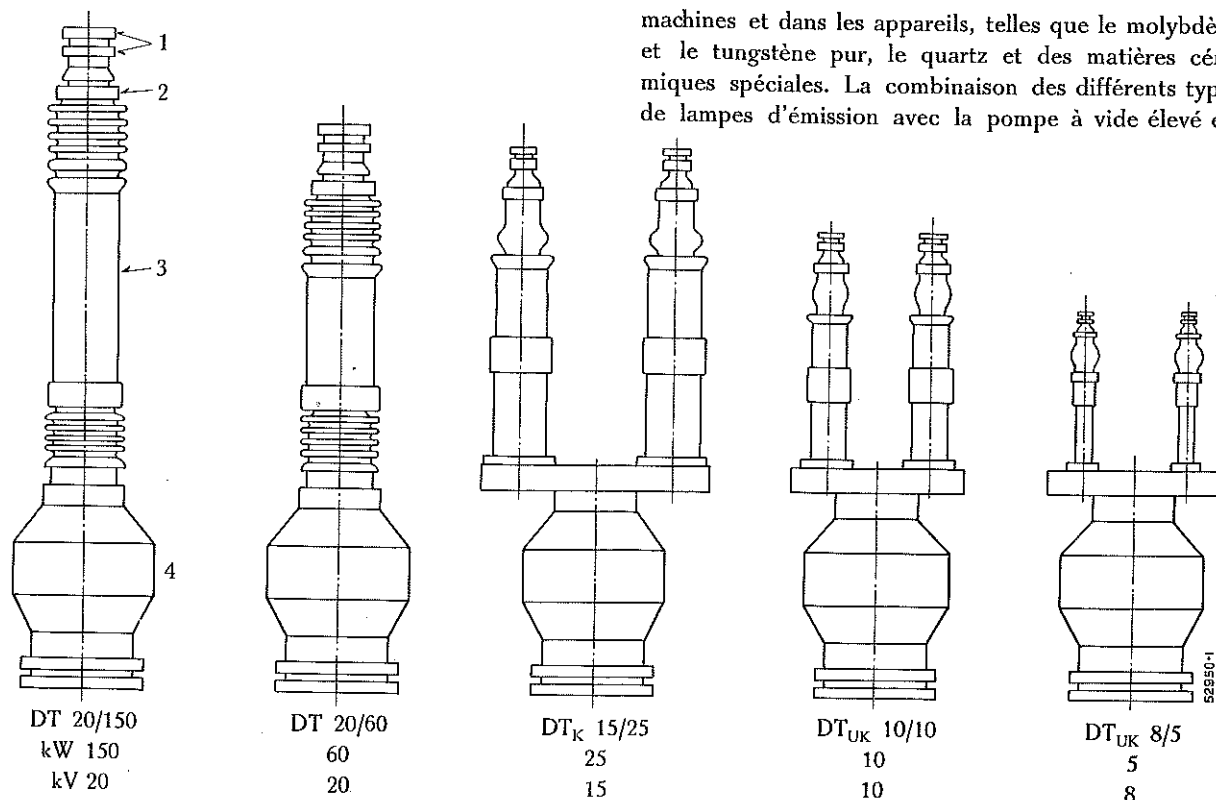


Fig. 2. — Lampe d'émission refroidie par eau.

1 = Borne de la cathode.

2 = Borne de la grille.

3 = Anode.

4 = Pompe à vide élevé.

La pompe à vide élevé sert de socle aux lampes. Celles-ci peuvent, selon la puissance, être montées sur la pompe, seules ou par deux.

en chaleur. Les tensions de fonctionnement peuvent atteindre environ 50 kV en haute fréquence et quelques centaines de kV à la fréquence industrielle.

Pour la construction, on utilise des matières qui ne sont pas ou presque pas utilisées dans les grandes

machines et dans les appareils, telles que le molybdène et le tungstène pur, le quartz et des matières céramiques spéciales. La combinaison des différents types de lampes d'émission avec la pompe à vide élevé est

montrée dans la figure 2. La pompe sert de socle à la lampe. La construction démontable permet d'élever presque à volonté la puissance et la tension.¹⁾ La figure 3 montre un type de lampe pour faible puissance, refroidie à l'air, et la figure 4 une lampe scellée en métal et matière céramique, et partiellement démontable. Cette dernière peut à choix être refroidie à l'air ou par un liquide, selon la

¹⁾ Revue Brown Boveri 1941 p. 319.

Fig. 3. — Lampe d'émission démontable refroidie par air.

- | | |
|---------------------------------------|--|
| 1 = Anode. | 5 = Borne de cathode avec réfrigérant. |
| 2 = Réfrigérant d'anode. | 6 = Pompe à vide élevée. |
| 3 = Réfrigérant de la bride d'anode. | 7 = Amenée d'air. |
| 4 = Borne de grille avec réfrigérant. | 8 = Ventilateur. |

Ce type refroidi à l'air convient particulièrement bien pour les faibles puissances et où le refroidissement par eau n'entre pas en ligne de compte.

Fig. 4. — Lampe d'émission scellée et partiellement démontable en métal et matière céramique pour refroidissement par liquide.

- | | |
|--------------------------|-----------------------|
| 1 = Queue. | 4 = Anode. |
| 2 = Borne de la cathode. | 5, 6, 7 = Isolateurs. |
| 3 = Borne de la grille. | |

Une cathode brûlée est facile à remplacer dans cette construction, lorsque la lampe a été coupée en x-x. La lampe est ensuite soudée, vidée de son air et scellée.

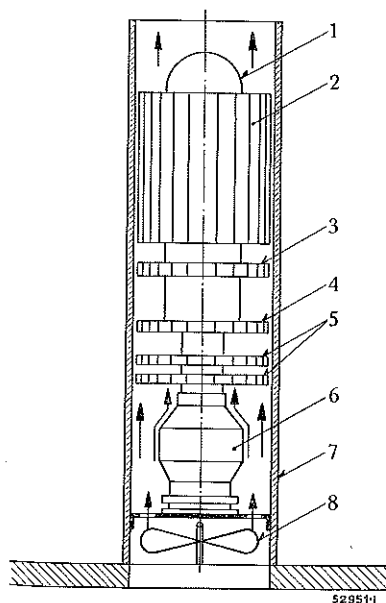


Fig. 3.

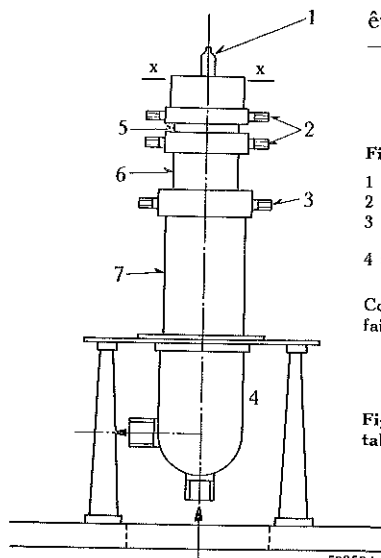


Fig. 4.

puissance demandée et les conditions locales de refroidissement. La particularité de ce type est qu'une cathode brûlée et même une grille défectueuse peuvent être changées facilement comme dans une lampe démontable. Il faut alors couper la tête métallique des lampes en x—x, remplacer la cathode ou la grille, souder un nouveau couvercle, évacuer les gaz et fermer. Les joints fondus entre les pièces métalliques et celles en matière céramique ne doivent plus être touchés. Cette opération ne peut naturellement être faite qu'en usine et peut être souvent répétée. Dans les lampes scellées en métal-verre, il faut couper le verre et après le remplacement de la cathode refermer en soudant le verre. Cette opération ne peut être effectuée qu'une ou deux fois pour les lampes de faibles puissances, pour les grandes puissances il est préférable le plus souvent de remplacer la lampe entièrement.

Le figure 5 montre une série de lampes d'émission pour des puissances de 5 à 150 kW, des tensions de 8 à 20 kV et des gammes d'ondes de 1 à 100 m et plus.

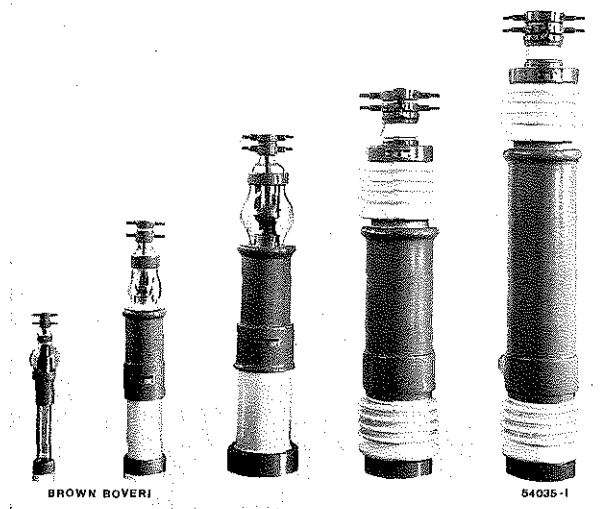


Fig. 5. — Lampes d'émission démontables.

Construites pour une dissipation anodique de 5 à 150 kW,
une tension anodique de 8—20 kV,
une puissance utile maximum de 10—500 kW,
une longueur d'onde de 1—100 m.

Cette série de lampes est déjà en vente.

La figure 6 montre une lampe d'émission, type DT 20/150, avec sa pompe. Un émetteur Brown Boveri d'une puissance utile de 100 à 200 kW, avec deux lampes de ce type, est depuis plusieurs mois en service avec de bons résultats, dans un poste de radio-diffusion. La lampe du type DT_K 15/25 fonctionne d'une façon très satisfaisante dans l'installation du cyclotron de l'Ecole Polytechnique Fédérale. Deux lampes montées en push-pull donnent une puissance utile supérieure à 40 kW pour une longueur d'onde de 18 m. Lors d'essais avec le type DT_K 10/10 en push-pull sur une onde de 5 m on a atteint une puissance utile de 10 à 15 kW.

III^e AVANTAGES DES LAMPES DÉMONTABLES.

Les lampes démontables ont, comparées aux lampes scellées en verre et métal, toute une série d'avantages tant économiques que techniques. L'avantage le plus intéressant est celui que nous avons déjà cité: la possibilité de changer facilement et rapidement les cathodes brûlées ou défectueuses. Ce changement peut par exemple s'effectuer avec un peu d'habitude en moins de 30 à 40 min pour une lampe de 150 kW. Avant la

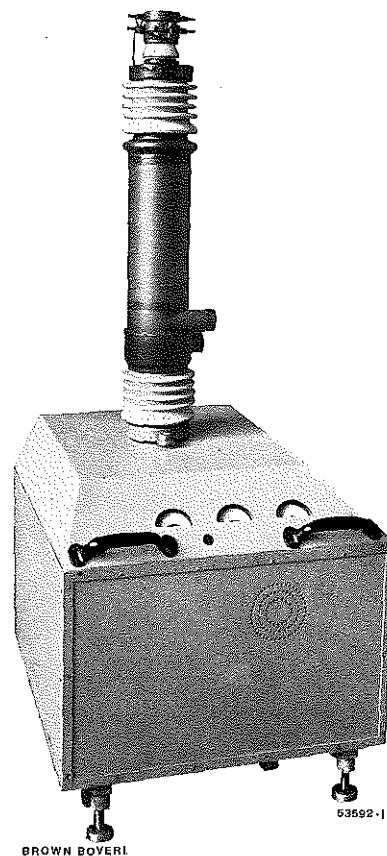


Fig. 6. — Lampe démontable de 150 kW.

Cathode interchangeable, d'où durée illimitée.

La pompe à vide poussé, celle à vide préliminaire et leurs accessoires sont montés dans le coffret métallique. Le tout est mobile et d'un faible encombrement.

mise en service il faut encore attendre deux à trois heures jusqu'à ce que la nouvelle cathode soit dégazée. Cette dernière période peut être ramenée à 15 ou 30 min si l'on garde sous vide des cathodes dégazées en réserve. Il sera judicieux de ne jamais attendre la rupture de la cathode mais d'entreprendre le remplacement après quelques milliers d'heures de service pendant un arrêt particulièrement long. Les installations importantes doivent toujours avoir des lampes de réserve prêtes à fonctionner. La possibilité de changer les cathodes est un avantage économique important, qui a d'autant plus de poids que la puissance des lampes est élevée.

Les avantages techniques ne sont pas moins importants. Les lampes démontables supportent des surcharges continues d'environ 50 % de la puissance nominale, ce qui est intéressant pour augmenter la puissance d'un émetteur, ou lorsque l'appareil sert à la production d'énergie à haute fréquence dans l'industrie.

La construction robuste a en outre comme conséquence que la grille et les connexions supportent des courants à haute fréquence qui sont inadmissibles dans les lampes scellées. Au début on croyait que la construction des lampes démontables ne se justifiait que pour des puissances supérieures à 50 kW. Mais au cours des perfectionnements on constata avec étonnement que cette construction était aussi intéressante pour les petits types ayant une dissipation anodique de 5 kW et tout particulièrement pour les ondes ultra-courtes.

Nous avons mis au point des constructions, qui ont de faibles capacités et de faibles inductivités et permettent donc d'atteindre des puissances utiles relativement élevées avec les ondes courtes.

Une autre propriété remarquable est la possibilité de modifier la caractéristique de la lampe dans certaines limites, ce qui est utile lors de la mise au point d'un émetteur. On peut rendre identiques les caractéristiques de deux lampes travaillant en push-pull. On peut aussi choisir la caractéristique de grille selon le but à remplir par la lampe, ou la capacité de charge de l'étage préliminaire. La courbe du courant de grille en fonction de la tension de grille peut être ascendante ou descendante et avoir des valeurs positives ou négatives. Pour les lampes scellées de construction rigide ce réglage n'est pas possible, et l'on doit choisir deux lampes de caractéristiques semblables. En effet de par la fabrication il y a toujours dans la position respective

Lampes d'émission démontables avec refroidissement par eau et pompe à vide élevé à fonctionnement continu.

	Type des lampes				
	DT _{UK} 8/5	DT _{UK} 10/10	DT _K 15/25	DT 20/60	DT 20/150
Dissipation anodique admissible kW	5	10	25	60	150
Tension d'anode kV	8	10	15	20	20
Puissance utile maximum kW	10 ²⁾	20 ²⁾	50 ²⁾	120 ²⁾	300 ²⁾
Pour une longueur d'onde de m	6	8	12	50	70
Longueur d'onde minimum m	0,9	2	3	4	6
Tension de chauffage . . V	7	10	12	16	32
Courant de chauffage . . A	75	145	145	430	430
Courant de saturation . A	4	8	10	50	100

¹⁾ Par lampe. ²⁾ Pour 2 lampes.

de l'anode, de la cathode et de la grille de légers écarts, qui modifient la caractéristique.

Notre système avec pompe à vide élevé et pompe préliminaire peut paraître un peu compliqué et encombrant, mais le groupe de pompes avec le dispositif de mesure du vide a donné d'excellents résultats. Quelques pompes moléculaires ont déjà fonctionné sans perturbation, pendant plus de deux ans en service continu. La pompe préliminaire ne travaille que par intermittence.

Nous sommes actuellement en mesure de fournir des lampes d'émission démontables, pour tous usages, puissances, tensions et fréquences (voir le tableau ci-dessus). Sous peu, nous pourrions aussi livrer des lampes scellées de petites puissances, à refroidissement par air ou par liquide.

(MS 812)

A. Gaudenzi. (J. C.)

GÉNÉRATEURS D'ONDES ULTRA-COURTES.

Indice décimal 621.396.615.14
621.385.029.6

La modulation de la vitesse des électrons en mouvement est le moyen d'obtenir les ondes les plus courtes. Le flux d'électrons, homogène à l'origine, sera coupé en paquets qui céderont leur énergie dans le champ alternatif d'un condensateur. L'application de ce principe conduit, d'une part, au «transator» avec système résonnant qui produit en même temps la réaction et qui sera traversé par le flux électronique, d'autre part, au «turbator», où un courant circulaire d'électrons servira à l'excitation d'un résonateur d'espace.

1^o INTRODUCTION.

L'extension de la gamme d'ondes aux ondes décimétriques et centimétriques représente un progrès important pour la technique. La possibilité de concentrer tout simplement ces ondes à l'aide des miroirs concaves permet l'emploi d'ondes dirigées, rendant possible un grand nombre de nouvelles applications. Comme la

longueur des ondes est encore relativement grande par rapport aux molécules de l'air et aux particules en suspension dans l'atmosphère, ces ondes ne possèdent pas les grands désavantages de la propagation des rayons lumineux, soit l'absorption et la dispersion.

A l'encontre des ondes courtes, dont l'importance pour la radiotechnique fut reconnue seulement lors de la découverte de leur grande portée, la possibilité d'application des ondes dirigées fut d'avance bien déterminée. Cependant, le générateur pour des ondes si courtes manquait. Deux exemples de leur vaste domaine d'application peuvent être cités, ce sont la détermination de la distance et de la direction, ainsi que la

technique des communications sans fil, à canaux multiples. Celle-ci peut atteindre un espace particulièrement étendu dans les pays plats qui ne disposent pas encore d'un grand réseau téléphonique. Il est même possible qu'elle supplante largement, étant plus économique, la technique actuelle des câbles à longue distance.

Bien entendu, les plus grandes exigences seront posées pour un tel système à canaux multiples quant à la caractéristique linéaire de modulation de l'onde porteuse qui permet théoriquement, en raison de sa fréquence élevée, la modulation simultanée de plusieurs centaines de canaux. Un problème reste irrésolu à l'heure actuelle, celui de savoir jusqu'où la réalisation pourra être faite dans ce sens.

Le développement des tubes électroniques a atteint aujourd'hui un très haut degré qui répond à toutes les exigences même à celles de la génération des ondes métriques de la télévision. Les lampes pour un mètre et moins sont caractérisées par des constructions spéciales, avec de petites distances entre les électrodes et des entrées courtes, ayant entre elles de faibles capacités. Il en est ainsi, par exemple, des lampes «gland» qui constituent une petite merveille de la technique. Ces lampes travaillent encore avec une commande pratiquement instantanée du courant par la tension de grille, c'est-à-dire que le temps de passage des électrons (temps de transit) est encore négligeable par rapport à la durée d'oscillation. Cependant, si le temps de transit était comparable à la durée d'oscillation, il se produirait un déphasage entre la tension de grille et le flux d'électrons commandé, qui serait mis en évidence par un aplatissement de la pente de la caractéristique. Dans un montage à réaction, ceci signifie une augmentation des difficultés des conditions initiales d'oscillation et une forte réduction du rendement. Pour ces raisons, les distances entre électrodes sont de $1/10$ mm et moins dans une lampe «gland» pour des ondes de 60 cm et une tension de fonctionnement d'environ 200 V.

II^o PRINCIPE DES GÉNÉRATEURS D'ONDES ULTRA-COURTES.

Si l'on veut passer aux ondes encore plus courtes (inférieures à 50 cm), on devra baser la commande du flux d'électrons sur un principe complètement nouveau. Au lieu de commander directement l'amplitude du flux d'électrons au moyen de la tension de grille, des mécanismes sélectifs déterminés produiront, dans le flux d'électrons, homogène du début, des concentrations et des raréfactions, qui seront entraînées à la vitesse des électrons. L'essentiel consiste en ce que ces concentrations produisent par influence, suivant les lois connues, un courant dans un condensateur couplé à un circuit oscillant, courant qui a pour conséquence de produire un champ freinant les paquets d'électrons. Par là, l'énergie du flux homogène d'électrons sera transformée précisément, en utilisant le temps de transit, en énergie alternative, car la tension entre armatures du condensateur pendant le passage des paquets d'électrons ne changera pas de signe, ou changera de façon

que, l'action de freinage du champ prédomine celle d'accélération.

Ce principe général, s'il n'est pas encore exprimé positivement partout, sert de base à tous les générateurs à temps de transit, présentant apparemment les formes les plus diverses. Depuis très longtemps, il appartient au domaine de la technique de la haute fréquence, et l'activité intense des laboratoires scientifiques et industriels de la haute fréquence consiste à rechercher à ce sujet l'application judicieuse de ce principe, c'est-à-dire la construction de générateurs d'ondes décimétriques et centimétriques stables et puissants.

La définition ci-dessus des oscillations créées par la modulation de vitesse serait cependant incomplète si trois caractères plus généraux et positifs n'étaient encore cités :

- 1^o L'oscillation doit s'entretenir, respectivement s'établir d'elle-même dans de tels générateurs. *Le principe du montage à réaction* est également pris comme base ici, de même que pour les tubes normaux. Par ce montage réactif, une discontinuité est créée en un certain point, dans le flux homogène électronique, discontinuité qui produit l'excitation de l'oscillation. Cette réaction est liée, dans la plupart des cas, au mécanisme du mouvement des électrons et ce montage à réaction n'est alors pas du tout visible.
- 2^o La génération d'ondes aussi courtes exige que le condensateur de couplage, mentionné plus haut, fasse partie du circuit oscillant. Tous ces organes sont placés, la plupart du temps, dans le tube électronique lui-même.
- 3^o Ces circuits oscillants sont actuellement réalisés sous forme de résonateurs d'espace. Ceux-ci sont des corps métalliques vides qui empêchent un rayonnement et se distinguent par un facteur de surtension particulièrement élevé (jusqu'à 10^6 , alors que les circuits oscillants normaux ont environ 10^2). Ce facteur, favorable au maintien d'une tension alternative élevée, est d'une grande importance pour les générateurs précités où les électrons doivent être freinés lors d'un passage qui ne se fait qu'une fois. Les résonateurs d'espace ont également l'avantage technique de permettre des constructions complètement métalliques, où le corps métallique du résonateur forme en même temps la chambre de décharge électronique.

On pourra se représenter les exigences devant être envisagées pour la précision du mécanisme de triage, lorsqu'on ne considère pas la longueur d'onde, mais le cycle qui agit sur le temps de transit. La fréquence correspondant à une onde de 3 cm et qui s'élève à 10^{10} cycles, mesurée par oscillographie pendant une seconde exigerait une longueur de film de 10 000 km pour une longueur d'un mm par cycle!

Dans ce qui suit, on décrira quelques réalisations de conceptions différentes, choisies parmi de nombreuses formes de construction et qui se détachent de celles connues jusqu'ici par des caractéristiques avantageuses.

III⁰ FORMES D'EXÉCUTION ET PROPRIÉTÉS.1⁰ Le «transator».

La figure 1 montre un générateur pour une longueur d'ondes de 25 cm et une puissance en haute fréquence de 1,2 watt. Le mécanisme de la production d'oscillations est particulièrement visible ici et peut être éclairci par cet exemple. Un flux d'électrons s'écoule dans la direction axiale de la cathode, située dans la partie inférieure du tube à travers un système d'optique électronique qui rend possible en même temps la modulation. Le circuit oscillant se compose d'un système de Lecher court-circuité, avec les deux grilles doubles, disposées selon la figure 2. Le calcul montre que des oscillations stationnaires peuvent se former par couplage avec le flux d'électrons¹⁾. La tension haute fréquence,

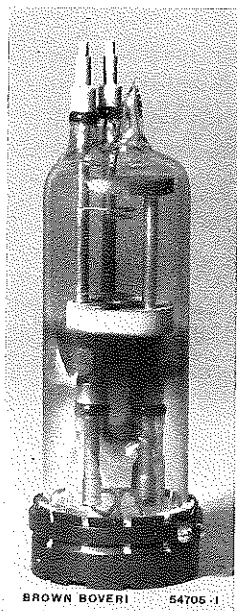


Fig. 1. — «Transator» pour ondes décimétriques.

La cathode est logée dans un support céramique d'un nouveau genre. La perforation de l'électrode d'accélération et des quatre grilles ont pour but une bonne concentration du flux électronique.

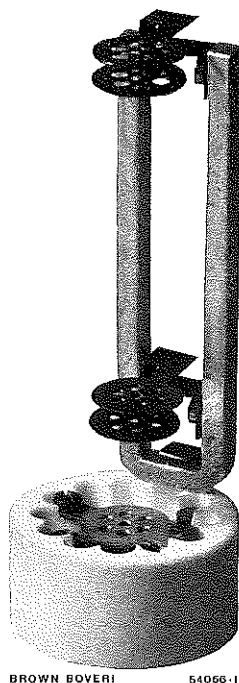


Fig. 2. — Système oscillant du «transator» selon figure 1.

Il est constitué en système de Lecher pour quart d'onde. La construction simple et robuste favorise également la réduction des pertes.

aux bornes de la grille double inférieure (électrode de commande), module la vitesse des électrons, mais pas le courant. A la sortie de la double grille inférieure, les électrons sont donc animés de vitesses différentes; ils se réunissent dans l'espace entre les deux paires de grilles et la concentration désirée est obtenue entre les deux grilles supérieures (électrode

de couplage) qui représentent la capacité de couplage du résonateur. Par réaction, une partie de la tension induite aux bornes de cette capacité sera conduite à l'électrode de commande, ce qui donne naissance à l'oscillation et l'entretient.

La réaction est bien perceptible dans ce cas. L'énergie de haute fréquence sera sortie par l'intermédiaire d'un système de Lecher, couplé directement au circuit de résonance. L'anode, située à la partie supérieure de la lampe, a pour but de recueillir les électrons. La combinaison de ces éléments en une unité organique a l'avantage d'être plus simple que la combinaison de deux résonateurs accordés, séparés et couplés entre eux par réaction.

La figure 3 montre un générateur construit d'après le même principe. Le circuit de résonance est constitué ici par un système de Lecher concentrique, fermé des deux côtés, et traversé transversalement par le flux d'électrons. Il forme un résonateur d'espace spécial, dont l'onde propre est égale au double de la longueur du cylindre. Les surfaces de pénétration du flux d'électrons dans les doubles cylindres du résonateur forment l'électrode de commande et les deux surfaces de sortie, l'électrode de couplage. La longueur de l'espace de transit est déterminé par le diamètre du tube intérieur. La prise d'énergie se produit également par couplage direct sur le circuit de résonance. La combinaison: électrode de commande, espace de transit, électrode de couplage et réaction, forme, avec le circuit de résonance, une construction très simple. En outre, vu sa simplicité géométrique et aucune perte par rayonnement n'entrant en ligne de compte, ce générateur peut être calculé complètement au point de vue quantitatif²⁾. Les rayons des deux cylindres con-

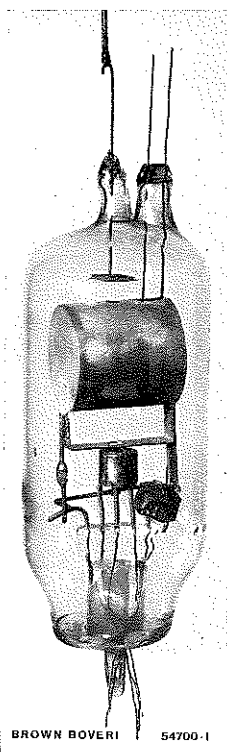


Fig. 3. — «Transator» avec résonateur cylindrique.

Le résonateur, constitué comme système Lecher concentrique, a une longueur de 3 cm, correspondant à une longueur d'ondes de 6 cm.

centriques sont parfaitement déterminés par la tension d'accélération et la longueur d'onde. Ensuite, on peut calculer la résistance de perte du système; elle est de $6 \cdot 10^4$ ohms pour une onde de 6 cm. Il en résulte un courant électronique de 14 mA, pour la mise en oscillation, lorsque le couplage du circuit de charge est pratiquement nul. Dans ce cas, la fréquence varie de 0,1 % lorsque la vitesse varie de 1 % (ce qui correspond à une variation de tension de 2 %). Pour

¹⁾ F. Lüdi, Helv. Phys. Acta Vol. XIII, Fascic. Sec. (1940) page 122.

F. Lüdi, Helv. Phys. Acta Vol. XIII, Fascic. Sext. (1940) page 498.

²⁾ F. Lüdi, à paraître prochainement dans les «Helv. Phys. Acta».

une onde de 6 cm, cette variation de fréquence correspond à 5.10^5 cycles. Une telle variation, relativement grande, se présente en principe pour tous les générateurs à temps de transit auto-excités. Elle est déterminée par le changement du temps de transit des groupes d'électrons, auquel correspond toujours un changement de la phase de réaction.

Les deux générateurs mentionnés sont caractérisés par le mouvement transversal des électrons dans un système oscillant comprenant un résonateur unique. C'est pourquoi ces formes de construction sont désignées par le terme de «transator».

2° Le «turbator».

Le résonateur¹⁾, représenté sur la figure 4, appartient à un générateur réalisé d'après une autre idée. L'émission d'électrons a lieu dans l'axe du système. Un courant circulaire d'électrons sera produit par l'action combinée d'un champ électrique continu radial et d'un champ magnétique axial. Les segments sont disposés alternativement sur les côtés du résonateur d'espace cylindrique, de telle façon qu'une tension alternative apparaît entre les segments voisins. L'action sur les électrons est plus compliquée que dans l'exemple précédent; elle se laisse également ex-

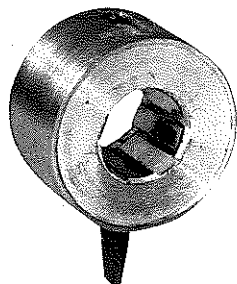


Fig. 4. — Système oscillant d'un turbator.

Le résonateur est constitué par un métal à point de fusion élevé et bien dégazé.

pliquer théoriquement²⁾. Le processus est à peu près le suivant: Le champ alternatif entre les segments crée, avec le champ magnétique constant, une accélération ou un ralentissement dans la direction tangentielle, des électrons émis à des temps différents; il en résulte un effet de sélection. Il se produit par là, dans le courant circulaire, homogène au début, des concentrations et des raréfactions. Ces groupes d'électrons livrent alors graduellement leur énergie, tirée du champ électrique continu, au champ alternatif du résonateur d'espace, contrairement aux deux générateurs mentionnés en premier lieu. La réaction, quoiqu'elle ne soit pas apparente, existe également ici. L'électrode de commande et celle de couplage sont formées par les mêmes segments.

¹⁾ F. Lüdi, rapport de session de la Sté Suisse de Physique des 7 et 8 septembre 1941.

²⁾ F. Lüdi, *Helv. Phys. Acta*, Vol. XIII, Fasc. Sext. (1940), page 77.

Dans son fonctionnement, ce générateur a une certaine ressemblance avec un alternateur. Le résonateur d'espace avec segments fixant la fréquence, est analogue au stator, tandis que la rotation des paquets d'électrons correspond au mouvement du rotor. Vu ces caractéristiques, ce type de construction a été désigné par le nom de «turbator».

Ce genre de construction des plus simples a les avantages suivants: Le flux circulaire d'électrons supprime l'optique électronique, nécessaire dans le «transator».

La longueur d'ondes est fixée d'une façon précise par le résonateur d'espace, soit par son diamètre et la capacité entre les segments. On sait que, dans les premiers magnétrons, l'onde fondamentale n'est pas excitée seule, mais qu'une série d'ondes voisines gêne la transmission. La liaison organique des segments avec le résonateur d'espace est une très bonne solution, aussi bien à l'égard de la dissipation anodique que de la possibilité d'influencer l'émission. Une exécution robuste, complètement en métal, peut également être réalisée.

Une première exécution, selon figure 4, a produit une puissance en haute fréquence d'environ 1 watt, pour une longueur d'ondes de 10,2 cm, mesurée au moyen d'un résonateur d'espace réglable, avec un écart absolu de moins de 0,1 mm. Le tube oscille d'une façon stable, sans qu'il soit nécessaire de stabiliser la source d'énergie et sans que la cathode ne soit détruite par le retour d'électrons. Une modulation jusqu'à près de 100 % est possible avec une électrode de commande appropriée, sans variations de fréquence sensibles. Le champ magnétique ne s'élève alors qu'à 400—500 gauss.

Ce champ faible permet l'utilisation d'aimants permanents légers, ce qui est de grande importance pour la construction des appareils.

Les générateurs connus depuis longtemps sous le nom de «tubes à champ freinant» et «magnétrons», employés pour les ondes courtes, utilisent un autre mécanisme de sélection pour grouper les électrons. Dans ce cas, les électrons dits de «fausse phase» sont écartés du flux par une électrode spéciale. Ce procédé ne réussit qu'imparfaitement. C'est pourquoi, seules de petites puissances peuvent être atteintes avec des rendements limités.

On a pourtant réussi à obtenir dans ce domaine des conditions favorables, en combinant le groupement en phase des électrons avec les résonateurs d'espace. Ce n'est que par les recherches les plus minutieuses de tous les éléments servant à la production d'oscillations que ces avantages sont devenus possibles, ainsi que les progrès signalés plus haut.

(MS 813)

Dr F. Lüdi. (Ma.)

PROCÉDÉS DE BROUILLAGE AUTOMATIQUES DE LA PAROLE.

Indice décimal 621.396.47

Le secret de la conversation n'a pas seulement une grande importance pour la transmission de renseignements militaires, mais aussi pour les conversations commerciales. Le secret est surtout important pour la transmission multiple sans fil, qui présente encore des possibilités inespérées de perfectionnement.

Nous avons entrepris, depuis longtemps, l'étude du brouillage de la parole et le perfectionnement d'appareils pour le brouillage automatique de transmissions téléphoniques, ou télégraphiques.

Nous décrivons dans cet article les méthodes fondamentales de brouillage ainsi que celles que nous avons mises au point pour augmenter le secret, et nous les comparons entre elles.

I^o INTRODUCTION.

Les transmissions sans fil peuvent, comme on le sait, être reçues, dans bien des cas, par des indiscrets. Même dans les transmissions par câbles ou par lignes aériennes, il existe des possibilités d'écoute. C'est pourquoi les télégrammes transmis sont souvent chiffrés. Il y a aussi des procédés de brouillage automatiques des messages parlés. Le récepteur, auquel le message est destiné, est muni d'un dispositif ajusté sur le programme de brouillage convenu et qui reproduit le langage clair initial.

Un moyen connu de rendre difficile à des tiers la réception de signaux transmis par sans fil, consiste à modifier continuellement ou brusquement l'onde porteuse selon un programme convenu. Toutefois, si la bande de l'amplificateur est suffisamment large ou s'il y a un dispositif d'accord automatique rapide, une réception utilisable de telles émissions est possible sans connaître le programme d'accord. De tels procédés ne sont pas appréciés à cause de la grande largeur de bandes nécessaires. D'autres propositions pour augmenter le secret en modulant la fréquence porteuse ne conviennent pas pour un système de brouillage à usages multiples, car les messages brouillés sont souvent transmis par des lignes téléphoniques dont la bande de fréquence est limitée. Dans de nombreux cas, par exemple montage d'appareils de brouillage dans des centraux téléphoniques ou même chez un abonné au téléphone, il ne faut pas avoir à s'inquiéter, si la communication est transmise par téléphone sans fil sur une partie de la distance. Cela entraîne le brouillage des messages à basse fréquence, sans élargissement de la bande de fréquence nécessaire.

Dans les méthodes connues de chiffage des télégrammes on écrit côte à côte les signes employés dans le même ordre mais en décalant les séries, et l'on remplace les signes d'une série par ceux de l'autre (substitution). Fréquemment on change aussi l'ordre

des signes (transposition). Ces procédés peuvent être réalisés automatiquement par des appareils appropriés.

Il est plus difficile de brouiller les messages parlés et de rétablir ensuite le message initial, car les mots et les phrases parlés ne sont pas composés de signes simples et séparés qui, comme les lettres et les chiffres, conviennent au chiffage; les vibrations de la parole sont un phénomène continu, qui est formé, pour correspondre aux multiples expressions et aux nuances diverses, d'un très grand nombre de différentes formes d'ondes, d'amplitudes, de fréquence et de durée diverses.

II^o COMPOSITION DE LA PAROLE ET POSSIBILITÉS DE BROUILLAGE.

Les voyelles et les consonnes ont des caractères tout différents provenant de la façon différente de les prononcer. Les voyelles «v» se composent d'une onde fondamentale et de nombreux harmoniques. Elles se différencient les unes des autres par leur fréquence fondamentale $f_0 = \frac{\omega_0}{2\pi}$ et par les amplitudes a_k de leurs composantes :

$$v = V(t) \cdot \sum_{k=1}^n a_k(t) \cos \left[k \int \omega_0(t) dt + \varphi_k \right]$$

La fonction $V(t)$ est l'expression des interruptions dues aux coupures de la phrase et aux consonnes. Les voyelles réparties dans la parole sont donc représentées par un mélange d'oscillations de fréquence et d'amplitude modulées agissant par intermittences et ayant une sonorité.

Les consonnes présentent par contre un spectre continu d'amplitudes sans oscillation fondamentale prépondérante, c'est-à-dire, qu'elles sont un bruit.

$$w = W(t) \int_0^{\omega_m} \left\{ b(\omega) \cdot \cos(\omega t) + c(\omega) \cdot \sin(\omega t) \right\} d\omega$$

Par un brouillage automatique de la parole, ces oscillations complexes sont modifiées selon un système convenu. Il faut cependant remarquer que les plus petites particularités du langage initial sont entendues par une oreille exercée. Cela peut suffire pour comprendre ou deviner quelques mots. Un auditeur, auquel le message n'est pas destiné, peut aussi essayer par variations systématiques des signaux brouillés de retrouver tout ou partie du procédé de brouillage.

Après le déchiffrement de quelques passages du texte, le reste va plus facilement, car on peut faire peu à peu de nouvelles suppositions. Si la clef reste la même, c'est-à-dire, si la relation entre le langage clair et le langage chiffré est toujours la même, il serait facile de la déterminer après avoir déchiffré un court passage d'un message. C'est pourquoi il faut constamment modifier le brouillage selon un programme établi; il faut aussi que le message puisse être rétabli d'une façon satisfaisante.

En considérant toutes ces difficultés et exigences on comprend que la plupart des brouillages automa-

III^e PROCÉDÉ DE SUBSTITUTION.

Un procédé de brouillage connu est le décalage, d'une valeur convenue, de la fréquence des vibrations de la parole initiale. Cette substitution de fréquence est obtenue par modulation avec des fréquences auxiliaires constantes ou variables qui, en général, sont choisies en dehors de la bande de fréquence de la parole.

Une simple modulation des vibrations x de la parole par une fréquence auxiliaire g_1 produit à la sortie du modulateur M deux mélanges de fréquence, l'un étant la somme de la fréquence auxiliaire avec les

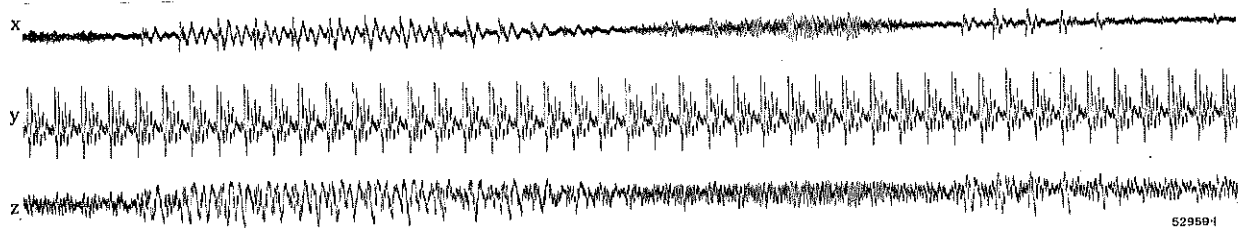


Fig. 1. — Procédé utilisant un signal perturbateur.

Par la superposition d'un signal perturbateur y au signal primitif x , on obtient le signal perturbé z . Ce brouillage est insuffisant, car on retrouve l'allure de x dans z .

tiques de la parole soient techniquement irréalisables ou n'assurent pas un secret suffisant.

Le procédé d'addition, qui superpose les oscillations d'un signal perturbateur à celles du message, est tout à fait insuffisant. Dans la figure 1 on donne en x une portion de l'oscillogramme du mot «Gesellschaft» et en y l'oscillation perturbatrice, la voyelle a . Dans l'oscillogramme du signal résultant z on retrouve encore bien des particularités de x . La séparation est encore plus facile pour une oreille exercée. Même si l'amplitude perturbatrice est un multiple de l'amplitude de la parole, certains mots peuvent quand même être compris. L'inévitable distorsion de phase et d'amplitude produite par la transmission modifie tellement le signal perturbateur superposé que la reconstitution du langage correct à la réception, par addition d'une onde de signe opposé, n'est pas possible.

Les procédés, qui brouillent le message en l'émettant par un dispositif à caractéristique non linéaire sont aussi insuffisants. Le message primitif est rétabli en utilisant, du côté récepteur, aussi un dispositif à caractéristique non linéaire. Avec ce système le message brouillé n'est jamais complètement incompréhensible, et une reconstitution satisfaisante est impossible à la réception à cause de l'inévitable déformation de phase et d'amplitude due à la transmission. Les mêmes objections peuvent être formulées pour de nombreux procédés maintenant le secret des messages par distorsion de phases ou d'amplitude, par interruption ou par mélange avec d'autres oscillations.

fréquences composant la parole, l'autre la différence. Sur la figure 2 la somme des fréquences est éliminée par un filtre passe-bas TP , si bien que seule la différence z_a est transmise, et les fréquences qui la composent sont comprises dans la bande de fréquence de la parole, si la fréquence auxiliaire correspond à la fréquence maximum de la parole. A la réception, on effectue, avec un dispositif semblable, la modulation de l'oscillation transmise par la fréquence auxiliaire g_1 , ce qui reconstitue le langage clair. Cette inversion de fréquence, souvent utilisée, ne répond pas aux exigences posées car la fréquence auxiliaire nécessaire au déchiffrement peut facilement être déterminée par des essais.

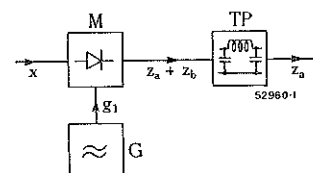


Fig. 2. — Procédé d'inversion.

Par la formation du produit de modulation du message primitif x et d'une fréquence auxiliaire g_1 , on produit le signal inversé z_a . Ce brouillage est insuffisant, car la fréquence auxiliaire utilisée g_1 est facilement trouvée par des essais.

Souvent les fréquences de la parole, décalées par une première modulation, sont modulées une seconde fois avec une fréquence auxiliaire g_2 . Dans les dispositifs de la figure 3, la différence de fréquence de la première modulation est supprimée par un filtre

passé-haut HP et la somme des fréquences de la seconde modulation par un filtre passe-bas TP. Si bien qu'il reste un signal z_c dont les composantes sont celles du message, décalées de la différence des deux

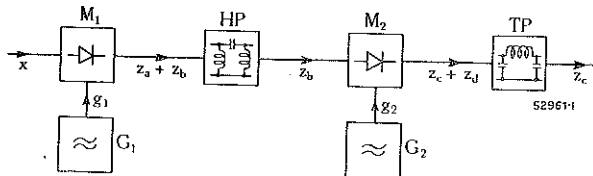


Fig. 3. — Décalage de la fréquence.

Par une double modulation du signal primitif y avec les fréquences auxiliaires g_1 et g_2 , on produit le message brouillé z_c avec composante de fréquence décalée. Le brouillage est insuffisant, car le décalage de fréquence est facilement trouvé par des essais.

fréquences auxiliaires. Même ce décalage ne répond pas aux conditions posées, car le décalage nécessaire pour reformer le langage clair est facilement déterminé par des essais.

Pour augmenter le secret, le spectre de fréquences de la parole est réparti en n bandes x_1 jusqu'à x_n qui pour la transmission ne sont pas décalées de la même valeur. De cette façon la bande x_m de la parole est décalée dans une bande x_k , aussi comprise dans la bande de fréquence de la parole. La clef de ce décalage de fréquence est donnée par la liaison entre m et k :

m	k
1	3
2	4
3	1
4	5
5	2

Elle peut être caractérisée par le nombre S qui est obtenu par une permutation de la succession normale 1, 2, 3, ... n . Dans notre exemple $S = 34152$.

Ce brouillage par substitution est obtenu à l'aide du dispositif de la figure 4. La première bande de fréquence x_1 est transmise par le filtre de bande F_1 au modulateur M_1 qui la module avec la fréquence auxiliaire g_1 issue du générateur G_1 . La somme des fréquences y_1 tombe dans la bande de fréquence que laisse passer le filtre passe-bande BP, et est transmise au modulateur N_3 . La différence de fréquence qui s'y forme constitue la troisième bande de fréquence z_3 du signal chiffré, car seules les différences peuvent traverser le filtre passe-bas TP. Les autres bandes de fréquences du message sont décalées de la même façon. Les bandes de fréquences traversant les filtres $F_1, F_2, \dots, F_5$, BP et les fréquences auxiliaires $g_1, g_2, \dots, g_5$ émises par les générateurs G_1 et G_2 ont par exemple les valeurs suivantes:

F_1	200	750 cycles	g_1	3850 cycles
F_2	750	1300 »	g_2	4400 »
F_3	1300	1850 »	g_3	4950 »
F_4	1850	2400 »	g_4	5500 »
F_5	2400	2950 »	g_5	6050 »
BP	3100	3650 »		

Les filtres d'entrée $F_1 \dots F_5$ peuvent, dans certains cas, être supprimés, si les fréquences que doivent bloquer les filtres tombent, après la première modulation, dans la zone bloquée par le filtre BP. Le même dispositif sert à rétablir le message initial par un nouveau décalage des fréquences.

Dans la figure 5, le premier oscillogramme représente le message non chiffré, le mot «Gesellschaft».

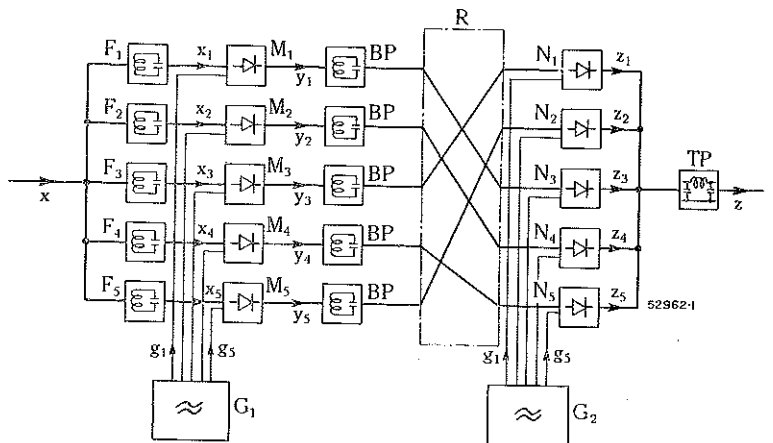


Fig. 4. — Brouillage par substitution.

Par une double modulation avec les fréquences auxiliaires $g_1 \dots g_5$ et une permutation des canaux, en R, on obtient le message brouillé z dont les bandes de fréquence sont décalées de valeurs différentes. Le déchiffrement par des indiscrets est difficile, car un rétablissement correct doit être fréquemment recherché.

En dessous sont reproduits séparément les spectres des cinq bandes de fréquences $x_1 \dots x_5$ avant le brouillage et $z_1 \dots z_5 = x_{31} \dots x_{15}$ après brouillage avec la clef $S = 34152$, le premier indice représentant le rang occupé avant le brouillage et le deuxième indice celui occupé après le brouillage.

Comme la clef S peut être assez facilement déterminée par des essais, elle est modifiée dans les nouveaux dispositifs après un certain temps t_0 . C'est pour cela que l'on utilise le permutateur R qui échange les cinq bandes de fréquence selon un programme convenu. Ce programme est donné par la succession des clefs $S_1 = 34152$, $S_2 = 51342$, $S_3 = 53421$ comme sur le tableau suivant:

m	k_1	k_2	k_3	k_4	k_5	k_6
1	3	5	5	1	3	4
2	4	1	3	4	2	5
3	1	3	4	5	4	2
4	5	4	2	2	1	3
5	2	2	1	3	5	1

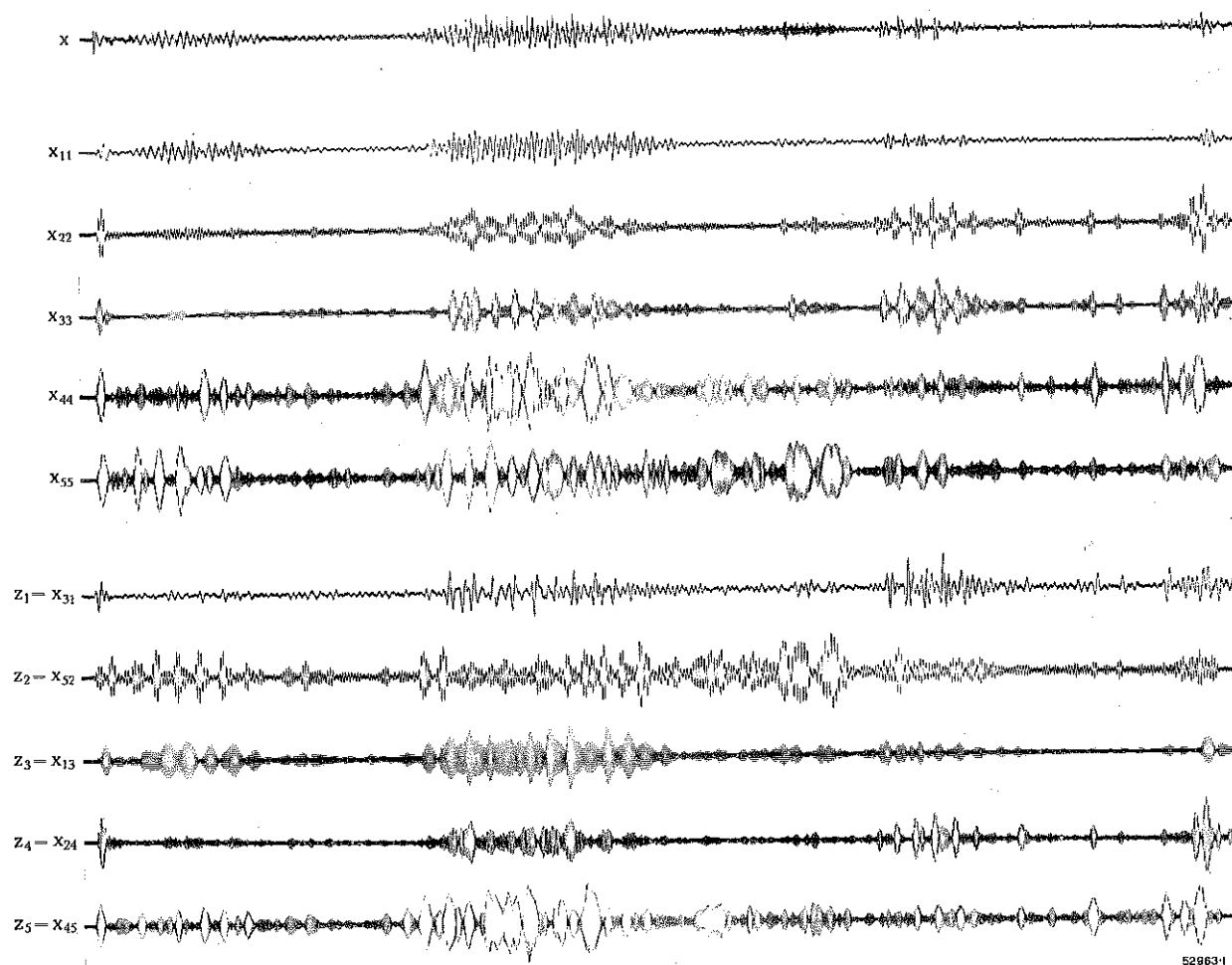


Fig. 5. — Brouillage par substitution.

x = Oscillogramme du mot non brouillé « Gesellschaft ». $x_{11} \dots x_{55}$ = Oscillogrammes des cinq bandes de fréquence du mot « Gesellschaft ». $z_1 \dots z_5$ = Oscillogrammes des bandes de fréquence permutées dont le message brouillé est formé. Des essais de décalage de fréquence permettent de fixer clairement que, par exemple, la bande de fréquence z_4 est formée de la deuxième bande de fréquence du message, x_{22} . Le secret est ainsi limité.

La bande x_1 du spectre apparaît successivement aux bandes 3, 5, 5, 1 ce qui donne dans le message brouillé les bandes x_{13} , x_{15} , x_{15} , x_{11} , etc.

Le message x à transmettre est subdivisé en éléments $a, b, c, \dots$, d'une durée t_0 . Le spectre dans la première bande est également divisé en éléments $a_1, b_1, c_1, \dots$, celui de la deuxième bande en a_2, b_2, c_2 , etc.

Le message peut être représenté de la façon suivante :

$$x = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \end{pmatrix} = \begin{pmatrix} a_1 & b_1 & c_1 & d_1 & e_1 & f_1 & g_1 & h_1 & i_1 & k_1 & \dots \\ a_2 & b_2 & c_2 & d_2 & e_2 & f_2 & g_2 & h_2 & i_2 & k_2 & \dots \\ a_3 & b_3 & c_3 & d_3 & e_3 & f_3 & g_3 & h_3 & i_3 & k_3 & \dots \\ a_4 & b_4 & c_4 & d_4 & e_4 & f_4 & g_4 & h_4 & i_4 & k_4 & \dots \\ a_5 & b_5 & c_5 & d_5 & e_5 & f_5 & g_5 & h_5 & i_5 & k_5 & \dots \end{pmatrix}$$

A chaque ligne on retrouve tous les termes faisant partie d'une des bandes de fréquences, c_2 , par exemple, est le troisième élément de la deuxième bande de fréquence.

Avec le programme de brouillage donné, le signal transmis z est représenté par :

$$z = \begin{pmatrix} z_1 \\ z_2 \\ z_3 \\ z_4 \\ z_5 \end{pmatrix} = \begin{pmatrix} a_{31} & b_{21} & c_{51} & d_{11} & e_{41} & f_{31} & g_{11} & h_{31} & i_{11} & k_{21} & \dots \\ a_{52} & b_{52} & c_{42} & d_{42} & e_{22} & f_{32} & g_{52} & h_{12} & i_{52} & k_{32} & \dots \\ a_{13} & b_{33} & c_{23} & d_{53} & e_{13} & f_{13} & g_{23} & h_{53} & i_{23} & k_{13} & \dots \\ a_{24} & b_{14} & c_{34} & d_{24} & e_{34} & f_{14} & g_{14} & h_{14} & i_{34} & k_{54} & \dots \\ a_{15} & b_{15} & c_{15} & d_{35} & e_{55} & f_{25} & g_{35} & h_{25} & i_{45} & k_{45} & \dots \end{pmatrix}$$

Dans chaque ligne on retrouve tous les éléments d'une bande de fréquence du message brouillé. Le premier indice indique de quelle bande de fréquence du message original cet élément provient et le deuxième dans quelle bande de fréquence l'élément a été décalé.

L'intervalle entre deux changements de clef, c'est-à-dire, la durée t_0 d'un élément est, en général, d'une seconde ou plus. Le déchiffrement du message reçu et enregistré est facile à un indiscret. On prend dans une bande de fréquence le spectre d'un

élément qui se reproduit souvent, et on le décale en fréquence d'une valeur variable jusqu'à ce que disparaisse à l'écoute cette impression de pas naturel provoqué par un décalage incorrect. On recherchera, de la même façon, le décalage des autres bandes de fréquence de cet élément et de tous les autres éléments, ce qui permettra de rétablir le langage clair. Ce procédé mène au but avec une rapidité surprenante, lorsqu'on est entraîné. Toutefois le temps employé augmente lorsque la durée de chaque élément diminue, car la détermination du décalage correct est difficile pour des éléments courts.

Par un procédé de synchronisation spécial nous avons pu réduire la durée des éléments à 0,05 s, ce qui permet d'améliorer notablement l'efficacité de ce brouillage par substitution. Les oscillogrammes obtenus

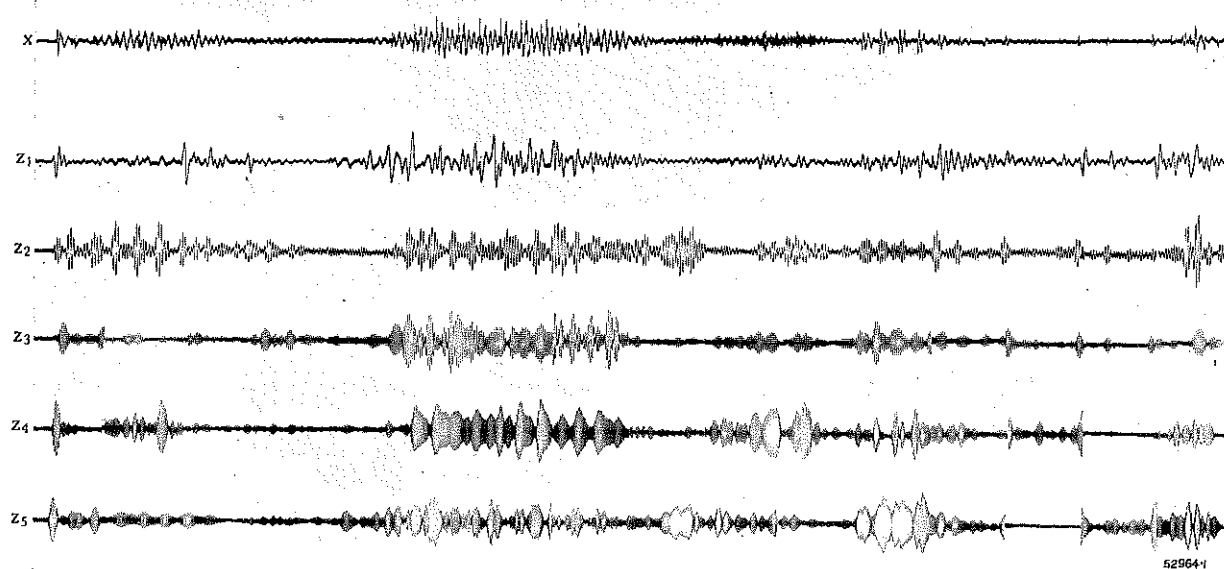


Fig. 6. — Brouillage par substitution.

x = Oscillogramme du mot non brouillé «Gesellschaft». $z_1 \dots z_5$ = Oscillogrammes des bandes de fréquence perméutées de x .
Par la permutation des bandes de fréquence, le rétablissement du message est rendu plus difficile et une notable amélioration du secret est obtenue.

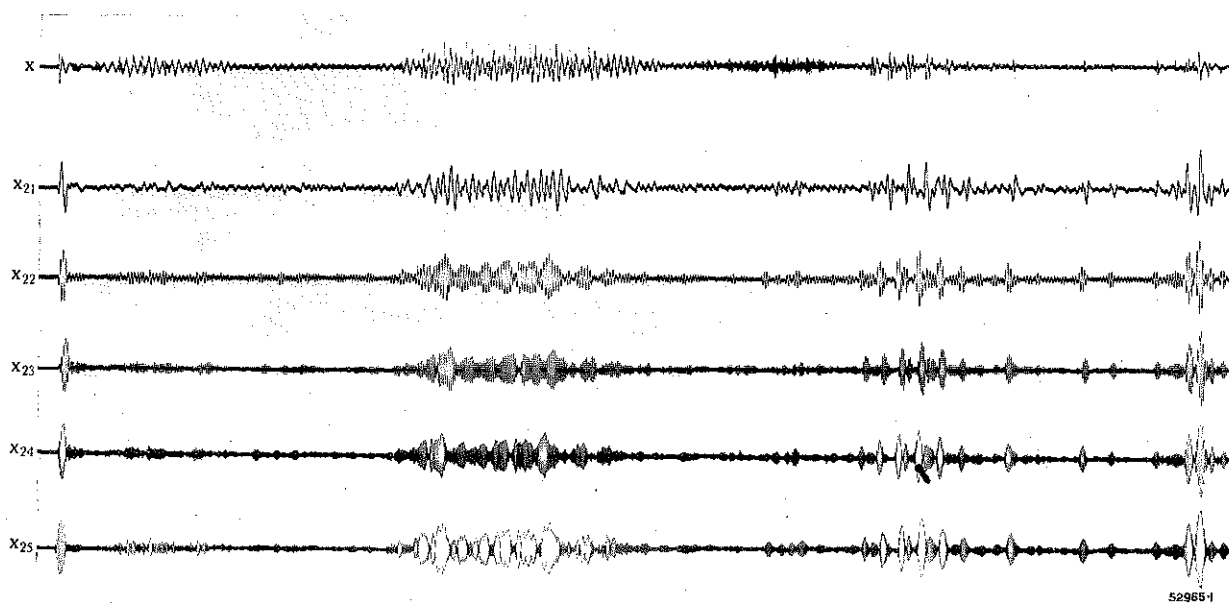


Fig. 7. — Décalage des bandes de fréquence.

x = Oscillogramme du mot non brouillé «Gesellschaft». $x_{21} \dots x_{25}$ = Deuxième bande de fréquence de x après décalage dans les premières jusqu'aux cinquièmes bandes de fréquence.
On voit nettement que la courbe d'amplitude n'est pas modifiée par ce décalage de fréquence. Le message brouillé par cette substitution de fréquence peut donc être retrouvé par la recherche du rythme de l'énergie.

avec des éléments d'une durée de 0,03 s sont donnés dans la figure 6. Le déchiffrement est aussi rendu plus difficile si l'on augmente le nombre des bandes de fréquences en réduisant simultanément leur largeur; mais cela augmente le nombre des filtres et modulateurs nécessaires. L'emploi de fréquences auxiliaires de modulation g , à variation continue, permet de remplacer le décalage brusque des fréquences par une variation continue.

La courbe, en fonction du temps, de l'amplitude moyenne $X_k = \sqrt{x_k^2}$ de chaque bande de fréquences n'est pas modifiée par le décalage de fréquence. Cela est par exemple visible sur la figure 7, qui reproduit le spectre de la deuxième bande de fréquence x_2 du mot «Gesellschaft» pour plusieurs décalages. Le rythme de l'émission reste le même, quelle que soit la durée des éléments, et le diagramme d'amplitude $Z_0(t) = \sqrt{z_1^2 + z_2^2 + \dots}$ de tout le message chiffré z correspond au diagramme d'amplitude $X_0(t) = \sqrt{x_1^2 + x_2^2 + \dots}$ du message clair. La fréquence fondamentale caractéristique ω_0 des voyelles brouillées peut en outre être facilement déterminée, car elle est la différence des composantes des fréquences décalées. C'est pourquoi on peut aussi, sans peine, dessiner le diagramme mélodique $\omega_0(t)$ qui donne la courbe de la tonalité des fréquences fondamentales en fonction du temps. Le déchiffrement du message chiffré se fera par comparaison des diagrammes d'amplitude $Z_0(t)$ et mélodique $\omega_0(t)$ avec des diagrammes connus de mots et de phrases. Bien que ces diagrammes dépendent de la prononciation, de l'accent, etc., un emploi systématique de cette méthode et quelques exercices permettent le déchiffrement de messages brouillés par substitution.

Nous pouvons, par un moyen relativement simple, compliquer le déchiffrement, en affaiblissant inégalement les différentes bandes de fréquence du message brouillé avant de l'émettre. Ce moyen modifie le diagramme d'amplitude qui ne peut plus servir au déchiffrement. Ces variations peuvent être supprimées facilement au récepteur, auquel le message est destiné, par une réduction correspondante de l'amortissement selon un programme convenu.

IV^e MÉTHODE DE TRANSPOSITION.

Les éléments du message initial $a, b, c, \dots$ d'une durée t_0 peuvent être décalés dans le temps de durées inégales. D'après une clef $T = 41532$, le premier élément prend alors la quatrième place, le second la première, etc. si bien qu'à la série initiale x correspond la série z :

$$\begin{aligned} x &= a b c d e \\ z &= b e d a c \end{aligned}$$

La réalisation de cette méthode de transposition est obtenue en retardant l'émission des éléments, de temps inégaux, mais qui diffèrent l'un de l'autre d'un multiple entier de t_0 . Comme des retards ne peuvent pas être négatifs dans la clef indiquée $T = 4\bar{1}5\bar{3}2$, les chiffres surmontés d'un trait seront augmentés de cinq unités d'où l'on obtient $T^* = 46587$ et l'on obtient alors le message chiffré suivant:

$$z^* = \dots a c b e d$$

Dans le dispositif de la figure 8, les éléments du message sont inscrits par aimantation variable sur une bande d'acier L se déplaçant dans le sens de la flèche, par des têtes d'inscription $K_1 \dots K_5$. La tête de lecture K_6 lit les éléments dans l'ordre où ils sont inscrits et la bobine de soufflage K_7 efface le message. Pour le brouillage avec la clef T revenant périodiquement,

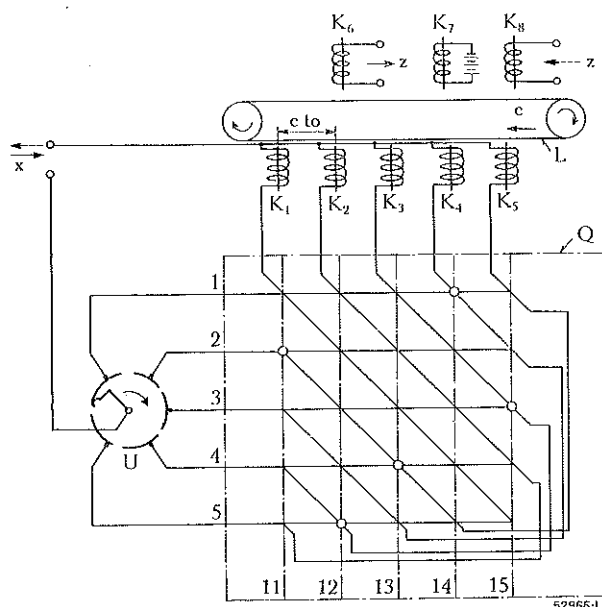


Fig. 8. — Brouillage par transposition.

Le message non brouillé x est transmis par le commutateur U et le tableau Q , alternativement aux différentes têtes d'inscription $K_1 \dots K_5$, de façon que l'inscription des éléments sur la bande d'acier L se fasse dans un ordre différent. La lecture du signal brouillé se fait par la tête K_6 et le soufflage par K_7 . Pour le débrouillage, le passage du message se fait dans le sens inverse (flèches pointillées).

les oscillations du message sont envoyés successivement dans les conducteurs horizontaux du tableau de brouillage Q par le commutateur U tournant dans le sens indiqué par la flèche; les conducteurs en diagonale de ce tableau sont reliés aux têtes d'inscription. Les liaisons entre ces conducteurs, indiquées par de petits cercles, correspondent à la clef T ; le premier conducteur horizontal est relié à son quatrième croisement avec une diagonale, le deuxième avec le premier croisement, etc. Le premier élément du message est inscrit, pour la position initiale indiquée du commutateur, en quatrième position sur la bande d'acier. Lorsque le commutateur a passé sur le deuxième

contact et que la bande a avancé de la longueur d'un élément entre deux têtes, le deuxième élément s'inscrit en sixième position. Après un tour du commutateur les cinq premiers éléments sont inscrits sur la bande dans l'ordre du message brouillé z et sont lus par la tête K_6 . Si le commutateur et la bande continuent à se déplacer, le même processus de brouillage se reproduit périodiquement.

Le brouillage doit être choisi tel que jamais deux éléments ne s'inscrivent l'un sur l'autre et qu'il n'y ait pas de vide sur la bande. Le message brouillé est alors sans lacune et sans recouvrement. L'étude d'une phase de l'inscription montre que cette condition est remplie si sur chaque conducteur horizontal 1...5 et sur chaque verticale 11...15, il n'y a qu'un point de liaison. Cela mène à la conclusion que la clef, comme dans la méthode de substitution, est obtenue par une permutation d'une série continue de chiffre 1, 2, 3...n.

Avec un même appareillage à la réception, une nouvelle permutation reproduit le message original. La clef utilisée T_2 est une permutation de la série de chiffre 1, 2, 3...n. d'après la clef d'émission T . A la clef d'émission $T = 41\ 532$ ou $T^* = 46\ 587$ correspond la clef de réception $T_e = 25\ 413$ ou $T_e^* = 25\ 468$. Par le brouillage avec la clef d'émission T^* , la succession d'éléments $x = a\ b\ c\ d\ e$ devient le signal émis y , après débrouillage avec la clef T_e^* à la réception on obtient le signal x_e , qui est reculé de cinq éléments par rapport à x :

$$\begin{array}{l} x = a\ b\ c\ d\ e \\ y = \dots a\ c \quad b\ e\ d \\ x_e = \dots \quad a\ b\ c\ d\ e \end{array}$$

Le signal brouillé est, pour le déchiffrement, écrit par la tête K_3 . Les éléments lus par les têtes $K_1 \dots K_5$ et repris par le commutateur U rendent le message primitif, si les clefs des deux appareils se correspondent.

Avec la plupart des clefs, certaines têtes d'inscription ou de lecture sont employées plusieurs fois, d'autres pas du tout. La clef de la figure 8 laisse par exemple les têtes K_1 et K_2 inemployées. Il est donc

possible de supprimer quelques têtes en acceptant une réduction du nombre de combinaisons.

L'indiscret peut, en essayant diverses clefs, chercher à reconstituer un message compréhensible avec les signaux reçus. Avec la méthode de transposition décrite, ce procédé conduit rapidement au but à cause de la répétition périodique rapide du même programme de brouillage.

Une augmentation suffisante de la période n'est pas possible avec un seul commutateur à cause du nombre réduit de positions de commutation. On peut aussi prévoir au poste émetteur plusieurs têtes de lecture mais en nombre différent de celui des têtes d'inscription; elles sont aussi reliées à un tableau de brouillage et à un commutateur. Dix têtes d'émission et onze têtes de lecture donnent avec ce brouillage en cascade une période de 110 éléments dans le message émis.

Dans la figure 9 on a reproduit un fragment du spectre du mot «Gesellschaft» et au-dessous le message intermédiaire y brouillé avec la clef $T_s = 46\ 587$; un nouveau brouillage à la lecture selon la clef $T_a = 437\ 658$ donne le message z , brouillé en cascade.

Parmi les nombreuses possibilités d'augmenter la période, nous citerons encore l'emploi de bandes perforées. La figure 10 montre le principe d'un dispositif pour l'inscription brouillée selon un programme sur bande perforée. Un trou O_{nm} placée à l'intersection de la $n^{\text{ème}}$ ligne transversale avec la $m^{\text{ème}}$ ligne longitudinale de la bande E , correspond à la $m^{\text{ème}}$ tête d'inscription et au $n^{\text{ème}}$ élément à transmettre. Comme jusqu'au moment de l'inscription de cet élément la bande d'acier a progressé de $n-1$ éléments, l'inscription aura lieu à la $(n+m-1)^{\text{ème}}$ place de la bande d'acier. La succession des inscriptions peut être déterminée à l'aide du schéma à droite de la figure 10. En traçant une ligne horizontale à partir du $n^{\text{ème}}$ élément du message non brouillé x jusqu'au trou O_{nm} et de là une ligne oblique, on atteint le $(n+m-1)^{\text{ème}}$ élément du message brouillé z . Les vides et les recouvrements sont évités, si sur chaque ligne horizontale et sur chaque ligne oblique il ne se trouve qu'un trou.

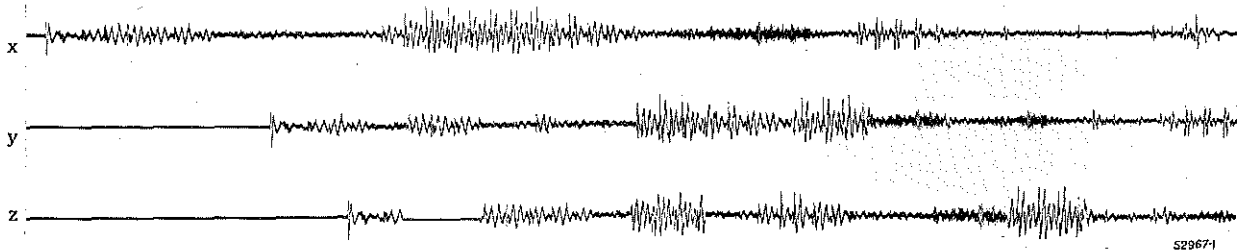


Fig. 9. — Brouillage par translation.

x = Oscillogramme du mot non brouillé «Gesellschaft».

y = Message intermédiaire après la première permutation des éléments.

z = Message brouillé après la deuxième permutation des éléments.

Le message brouillé est incompréhensible à cause de la permutation de tous les éléments. Un déchiffrement est cependant possible par tâtonnement.

Pour la commande d'un grand nombre de têtes d'inscription on peut aussi faire plusieurs trous dans la même ligne horizontale, si l'on dispose d'un combinateur électrique faisant correspondre à chaque combinaison de trous la commande d'une tête. Pour

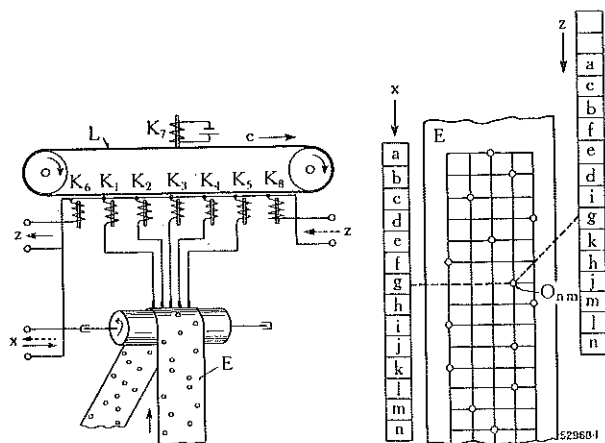


Fig. 10. — Brouillage par transposition.

x = Message clair. —> = Direction à l'émission.
z = Message brouillé.> = Direction à la réception.

La permutation des éléments du message peut être modifiée par le couplage des têtes d'inscription ou de lecture $K_1 \dots K_5$ à l'aide d'une bande perforée E. On obtient ainsi un programme de brouillage de grande période, ce qui complique énormément le déchiffrement.

une inscription correcte, les conditions sont les mêmes que précédemment si l'on considère une bande ayant autant de lignes longitudinales qu'il y a de têtes et en représentant une combinaison de trous par un seul sur la ligne correspondant à la tête qu'elle commande.

La durée d'une syllabe de la langue allemande parlée normalement varie de 0,05 à 0,5 s. La durée d'un élément t_0 doit être plus courte que la durée d'une syllabe et le décalage moyen d'un élément plus long. En pratique on utilise les valeurs suivantes:

Durée de l'élément = 0,05 s

Retard maximum = 0,8 s.

En partant de ces valeurs, nous avons construit des appareils de brouillage qui répondent parfaitement aux conditions posées en ce qui concerne l'exactitude de la synchronisation et la compréhension du message remis au clair, même pour des conditions défavorables de transmission. La courbe d'amplitude qui reste la même avec la méthode de substitution normale est complètement transformée par la permutation des éléments de la méthode de transposition. Un déchiffreur indiscret ne peut donc plus utiliser le système de l'énergie. Si le décalage des éléments est insuffisant, on peut toutefois se faire une idée du contenu du message en recherchant les voyelles

d'un groupe d'éléments et en les rangeant dans un ordre correct qui permet de reconnaître les mots. Une clef de déchiffrement choisie arbitrairement ramène souvent plusieurs éléments à peu près dans l'ordre primitif. Avec un peu d'exercice il est possible de deviner des mots du message, après quelques essais de déchiffrement avec des clefs différentes.

Les possibilités de déchiffrement diminuent donc rapidement avec l'augmentation du décalage des éléments. Mais cette augmentation accroît la durée de la transmission de l'entrée de l'appareil de brouillage jusqu'à la sortie de l'appareil de déchiffrement, ce qui rend plus difficile les communications bilatérales et réduit les possibilités d'emploi.

V⁰ COMBINAISON DES DEUX MÉTHODES DE BROUILLAGE.

Aussi bien la méthode de substitution des bandes de fréquences que celle de transposition ne satisfont pas toujours à la condition posée concernant le secret. Avec la méthode de substitution ordinaire, un déchiffrement est possible en se servant des diagrammes d'amplitude et de mélodie, et les éléments chiffrés suivant la méthode de transposition peuvent aussi être remis dans l'ordre correct pour certains passages, si leur décalage est court. Dans les deux cas, le déchiffrement par celui qui n'a pas la clef demande une grande routine et un appareillage approprié, ainsi que beaucoup de temps.

Pour tous les cas où le secret doit être conservé, même vis-à-vis d'indiscrets bien équipés, nous avons mis au point un procédé qui supprime les défauts des deux méthodes de brouillage mentionnées. Du signal initial x dont les éléments étaient dans l'ordre a b c ..., etc., et que l'on a partagé par exemple en cinq bandes de fréquences $x_1 \dots x_5$, on forme un signal z dont les composantes sont changées de fréquence et décalées dans le temps:

$$z = \begin{pmatrix} z_1 \\ z_2 \\ z_3 \\ z_4 \\ z_5 \end{pmatrix} = \begin{pmatrix} . & . & a_{11} & b_{31} & c_{21} & e_{21} & f_{31} & h_{11} & g_{11} & g_{31} & . & i_{11} & . \\ . & b_{22} & . & a_{22} & a_{42} & c_{32} & g_{12} & h_{32} & e_{32} & k_{22} & h_{12} & i_{32} & . \\ a_{33} & . & b_{53} & c_{43} & b_{43} & f_{53} & d_{23} & d_{33} & f_{13} & f_{23} & . & . & k_{43} \\ . & . & . & . & d_{44} & d_{51} & e_{14} & g_{54} & g_{21} & i_{54} & i_{21} & h_{24} & k_{54} \\ . & a_{55} & c_{55} & d_{15} & e_{45} & b_{15} & c_{15} & e_{35} & i_{15} & h_{55} & k_{35} & k_{15} & . \end{pmatrix}$$

Le premier indice se rapporte à la bande de fréquence dans laquelle se trouvait l'élément et le deuxième celui dans laquelle il se trouve. Comme il y a un changement de fréquence (substitution) et un décalage dans le temps (transposition), on parle de méthode combinée.

Le nouveau brouillage se fait à l'aide d'un détournement des messages intermédiaires. Le premier message intermédiaire est obtenu en ramenant par exemple les cinq bandes de fréquences $x_1 \dots x_5$, dans une même sixième bande désignée par l'indice 6.

$$y_a = \begin{pmatrix} y_{a1} \\ y_{a2} \\ y_{a3} \\ y_{a4} \\ y_{a5} \end{pmatrix} = \begin{pmatrix} a_{16} & b_{16} & c_{16} & d_{16} & e_{16} & f_{16} & g_{16} & h_{16} & i_{16} & k_{16} \\ a_{26} & b_{26} & c_{26} & d_{26} & e_{26} & f_{26} & g_{26} & h_{26} & i_{26} & k_{26} \\ a_{36} & b_{36} & c_{36} & d_{36} & e_{36} & f_{36} & g_{36} & h_{36} & i_{36} & k_{36} \\ a_{46} & b_{46} & c_{46} & d_{46} & e_{46} & f_{46} & g_{46} & h_{46} & i_{46} & k_{46} \\ a_{56} & b_{56} & c_{56} & d_{56} & e_{56} & f_{56} & g_{56} & h_{56} & i_{56} & k_{56} \end{pmatrix}$$

Pour la formation d'un autre message intermédiaire, les éléments sont changés de ligne :

$$y_b = \begin{pmatrix} y_{b1} \\ y_{b2} \\ y_{b3} \\ y_{b4} \\ y_{b5} \end{pmatrix} = \begin{pmatrix} a_{36} & b_{26} & c_{56} & d_{16} & e_{46} & f_{56} & g_{46} & h_{36} & i_{16} & k_{26} \\ a_{56} & b_{46} & c_{16} & d_{56} & e_{26} & f_{36} & g_{56} & h_{46} & i_{56} & k_{36} \\ a_{16} & b_{36} & c_{26} & d_{56} & e_{16} & f_{46} & g_{26} & h_{56} & i_{26} & k_{16} \\ a_{26} & b_{46} & c_{36} & d_{26} & e_{36} & f_{16} & g_{16} & h_{16} & i_{36} & k_{56} \\ a_{46} & b_{16} & c_{16} & d_{36} & e_{56} & f_{26} & g_{36} & h_{26} & i_{46} & k_{46} \end{pmatrix}$$

Par un retard inégal des messages sur les différentes lignes on obtient le message intermédiaire suivant :

$$y_c = \begin{pmatrix} y_{c1} \\ y_{c2} \\ y_{c3} \\ y_{c4} \\ y_{c5} \end{pmatrix} = \begin{pmatrix} a_{36} & b_{26} & c_{56} & d_{16} & e_{46} & f_{56} & g_{46} & h_{36} & i_{16} & k_{26} & \dots \\ \dots & a_{56} & b_{46} & c_{16} & d_{56} & e_{26} & f_{36} & g_{56} & h_{46} & i_{56} & k_{36} & \dots \\ \dots & \dots & a_{16} & b_{36} & c_{26} & d_{56} & e_{16} & f_{46} & g_{26} & h_{56} & i_{26} & k_{16} & \dots \\ \dots & \dots & \dots & a_{26} & b_{46} & c_{36} & d_{26} & e_{36} & f_{16} & g_{16} & h_{16} & i_{36} & k_{56} & \dots \\ \dots & \dots & \dots & \dots & a_{46} & b_{16} & c_{16} & d_{36} & e_{56} & f_{26} & g_{36} & h_{26} & i_{46} & k_{46} \end{pmatrix}$$

Par une nouvelle permutation des termes y_c d'un canal à l'autre on obtient :

$$y_d = \begin{pmatrix} y_{d1} \\ y_{d2} \\ y_{d3} \\ y_{d4} \\ y_{d5} \end{pmatrix} = \begin{pmatrix} \dots & a_{16} & b_{36} & c_{26} & e_{26} & f_{36} & f_{46} & h_{46} & g_{16} & g_{36} & \dots & i_{16} & \dots \\ \dots & b_{26} & a_{56} & a_{16} & c_{56} & g_{46} & h_{36} & e_{56} & k_{26} & h_{16} & i_{36} & \dots & \dots \\ \dots & a_{36} & b_{56} & c_{46} & b_{46} & f_{56} & d_{26} & d_{36} & f_{16} & f_{26} & \dots & \dots & k_{46} \\ \dots & \dots & \dots & d_{46} & d_{56} & e_{16} & g_{56} & g_{26} & i_{56} & i_{26} & h_{26} & k_{56} & \dots \\ \dots & a_{56} & c_{56} & d_{16} & e_{16} & b_{16} & c_{16} & e_{36} & i_{16} & h_{56} & k_{36} & k_{16} & \dots \end{pmatrix}$$

Les messages transmis par chacune des lignes sont maintenant ramenés par un décalage inégal de fréquence dans les cinq bandes de fréquences du spectre de la parole et l'on obtient le message z .

Le déchiffrement du message brouillé par transposition en déplaçant des éléments n'est maintenant plus possible à cause du changement de fréquence. Le déphasage des éléments déforme le diagramme d'énergie utilisé pour le déchiffrement des messages brouillés par substitution or-

динаire. Ce procédé élimine donc les deux possibilités dangereuses de déchiffrement.

La figure 11 montre un dispositif construit pour le brouillage combiné. $M_1 \dots M_5$ sont les modulateurs qui font la différence entre les fréquences du message et les fréquences auxiliaires du générateur G . Le message intermédiaire y_a est formé en filtrant par les filtres BP la bande de fréquence 3100 à 3650 cycles du produit modulé.

En traversant le permutateur Q , les éléments du message sont permutés et l'on obtient à la sortie sur les cinq lignes les signaux $y_{b1} \dots y_{b5}$. Les cinq lignes passent ensuite chacune sur un dispositif retardateur $H_1 \dots H_5$ (ces cinq dispositifs sont entraînés par le même moteur). Les retards produits diffèrent les uns des autres de la durée t_0 d'un élément. Aux têtes de lecture $K_1 \dots K_5$ on obtient les cinq parties $y_{c1} \dots y_{c5}$ du message intermédiaire y_c . On peut, si c'est nécessaire, faire passer chacune des parties par un filtre de bande ZP, qui supprime toutes les fréquences perturbatrices, en dehors de la bande. Une nouvelle permutation des éléments par le permutateur P donne le message intermédiaire y_d puis par une nouvelle modulation avec les fréquences auxiliaires $g_1 \dots g_5$, on obtient le message brouillé z qui couvre toute la bande de fréquence de 200 à 2950 cycles. Les sommes des fréquences produites par la modulation dans $N_1 \dots N_5$ sont supérieures à 6950 cycles et sont supprimées par le filtre passe-bas TP.

L'appareillage est parcouru dans le même sens pour la réception. Il faut seulement inverser les

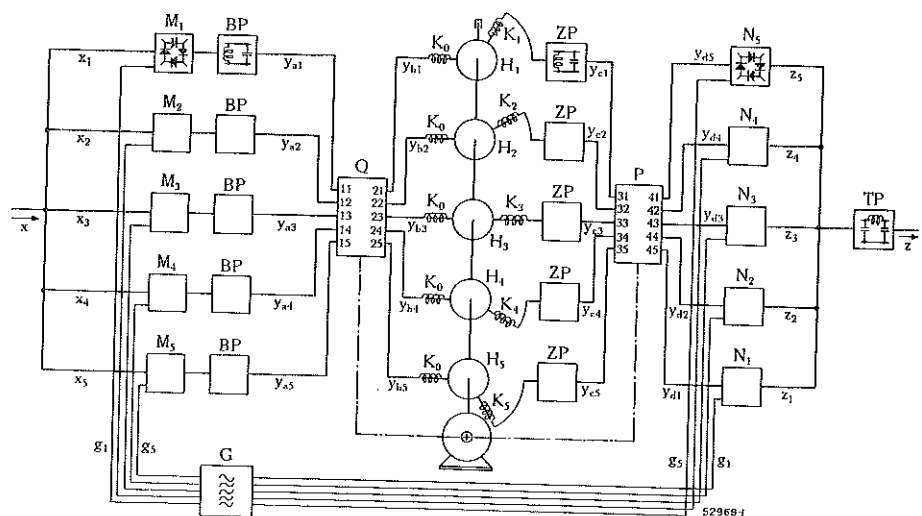


Fig. 11. — Brouillage combiné.

Les bandes de fréquence du message $x_1 \dots x_5$ sont décalées par modulation avec les fréquences auxiliaires $g_1 \dots g_5$ d'une certaine valeur et retardées d'un temps variable par les appareils $H_1 \dots H_5$. Ce double brouillage assure un bon secret.

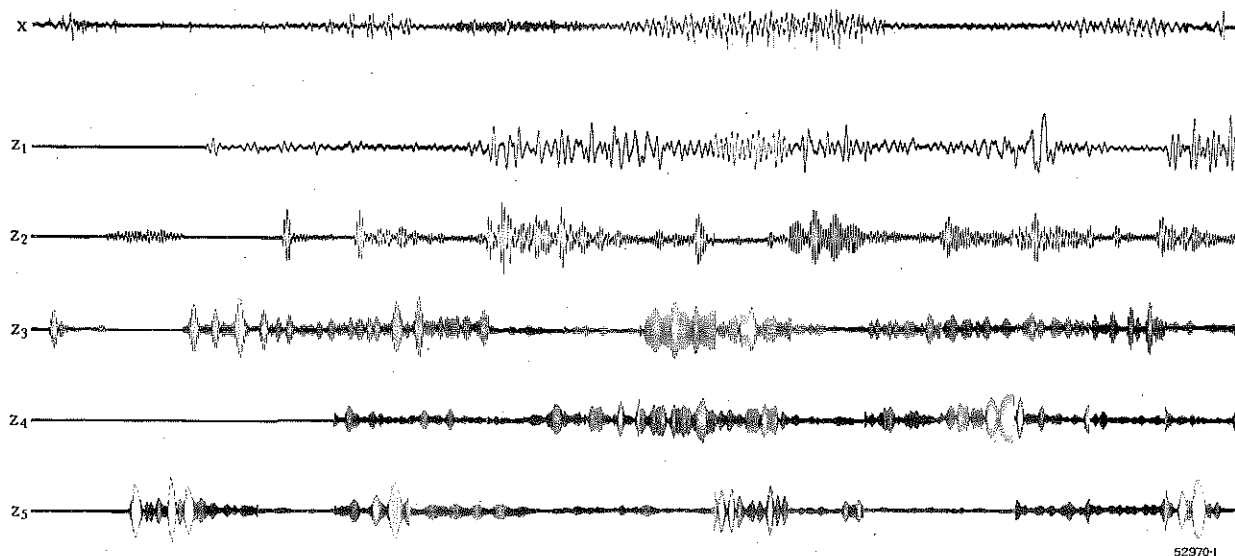


Fig. 12. — Brouillage combiné.

x = Oscillogramme du mot «Gesellschaft» (message non brouillé).

$z_1 \dots z_5$ = Bandes de fréquence du message brouillé après décalage et retardement des bandes de fréquence.

Les caractéristiques essentielles du message non brouillé x ne sont plus reconnaissables dans les bandes de fréquence retardées et décalées $z_1 \dots z_5$ et tout essai de déchiffrement sans connaître le programme du décalage est pratiquement impossible.

commandes des deux permutateurs P et Q. En outre la commande du permutateur P, à la réception, doit être retardée par rapport à celle du permutateur d'émission, d'un temps qui correspond au temps que met chaque élément pour traverser les deux appareils. Ce temps de passage peut être faible, contrairement à ce qu'exige la méthode de transposition, si bien que cela ne rend pas la conversation difficile.

Dans la figure 12, nous avons reproduit sous l'oscillogramme x du spectre d'un mot, l'oscillogramme des émissions sur les différentes bandes $z_1 \dots z_5$ du message z obtenu par brouillage combiné selon le programme donné avec des éléments d'une durée de 0,03 s.

VI⁰ TRANSMISSION DE VALEURS CARACTÉRISTIQUES.

Un autre brouillage efficace est obtenu si, au lieu de modifier la fréquence ou de décaler les fréquences de la parole, on fait les mêmes opérations avec les valeurs caractéristiques du spectre variable des amplitudes. Pour une transmission fidèle, il faudrait un nombre incommensurable de valeurs caractéristiques, correspondant à l'infinité des amplitudes partielles variables avec le déphasage correspondant.

Il est préférable de déterminer par une analyse automatique de la voix les fréquences fondamentales variables des voyelles ainsi que l'amplitude d'un grand nombre

de bandes de fréquences du spectre de la parole. Pour la transmission des valeurs caractéristiques variables correspondantes, un nombre fini de canaux de transmissions suffit, il permet même de couvrir une bande totale de fréquences plus étroite que le spectre de fréquences de la parole. Pour reproduire la parole à la réception, on forme un mélange d'oscillations dont la fréquence fondamentale des voyelles est amenée en correspondance avec celle à l'émission à l'aide de la valeur caractéristique correspondante. Pour les consonnes, le mélange d'oscillation a, comme à l'émission, le caractère d'un bruit.

A l'aide des caractéristiques d'amplitude particulières, les bandes de fréquences de ce mélange d'oscillations sont commandées de façon que la courbe de l'amplitude en fonction de la fréquence corresponde au moins à celle du message initial.

La réalisation de la méthode des valeurs caractéristiques est plus simple, si, déjà à l'émetteur, on filtre la fréquence fondamentale g des voyelles et un mélange d'oscillation, h , des consonnes de faible largeur de bande et si on les transmet au récepteur par un canal spécial. De g ou h on peut produire à la réception par une transmission non linéaire un mélange d'oscillations v_0 ou w_0 qui est composé des harmoniques supérieurs de g ayant une sonorité ou d'un nombre illimité de vibrations partielles ayant le caractère d'un bruit. De ce mélange d'oscillations on forme, par variation de l'amplitude des diverses bandes de fréquences en fonction des valeurs caractéristiques, des syllabes, mots et phrases compréhensibles.

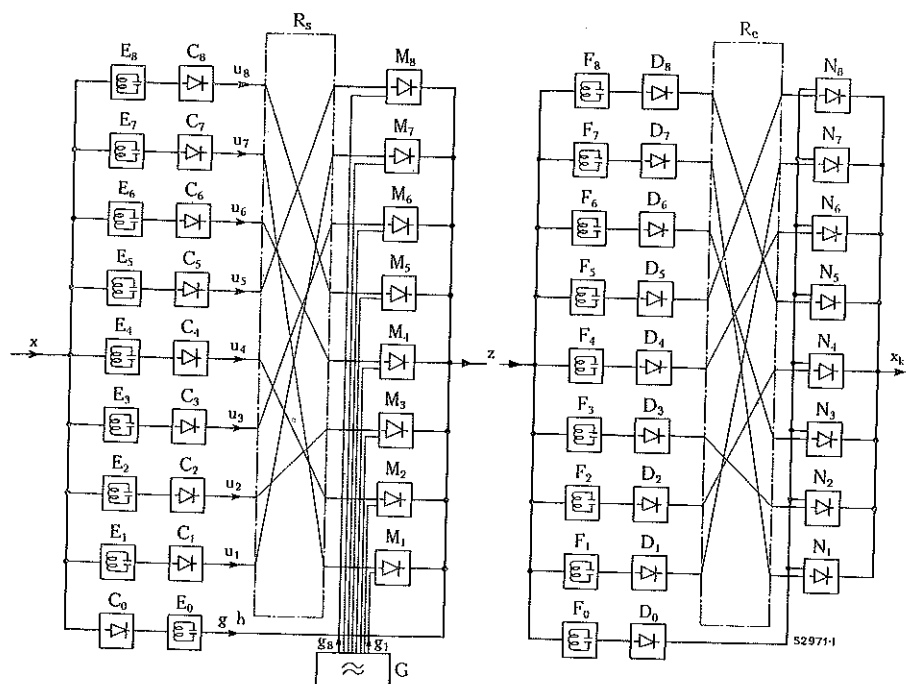


Fig. 13. — Brouillage par valeurs caractéristiques.

Les valeurs caractéristiques $u_1 \dots u_8$ qui caractérisent le spectre d'amplitude du message non brouillé x sont permutées du côté émetteur par le permutateur R_s et rétablies dans le bon ordre du côté récepteur par le permutateur R_c , si bien que le message x_k est de nouveau compréhensible.

Dans le dispositif de secret de la figure 13, on a prévu, à l'émetteur, un élément transducteur non linéaire C_0 et un filtre passe-bas E_0 , pour extraire g ou h . En filtrant les différentes bandes de fréquences $E_1 \dots E_8$ du spectre de la parole et en les démodulant dans $C_1 \dots C_8$, on obtient les valeurs caractéristiques $u_1 \dots u_8$. Pour transmettre ces valeurs caractéristiques par un canal, on les module dans les modulateurs $M_1 \dots M_8$ avec les basses fréquences porteuses $g_1 \dots g_8$ produites par le générateur G .

A la réception, les valeurs caractéristiques sont retrouvées par élimination des ondes porteuses dans les filtres de bandes $F_1 \dots F_8$ et par démodulation dans $D_1 \dots D_8$. A la sortie du filtre passe-bas F_0 apparaissent les oscillations g ou h qui ont été amenées par le transducteur non linéaire D_0 pour former le mélange d'oscillation v_0 ou w_0 . Les bandes de fréquences formant la parole sont maintenant produites par réglage de l'amplitude dans les modulateurs $N_1 \dots N_8$ à l'aide des valeurs caractéristiques $u_1 \dots u_8$ à variation lente.

Un auditeur indiscret pourrait, avec un appareil récepteur convenable, recevoir un message compréhensible. C'est pourquoi on a prévu les permutateurs R_s et R_c qui croisent de façon continue les canaux à l'émission et les rétablissent dans l'ordre convenable à la réception. A la place de cette permutation des canaux on peut aussi employer un brouillage par

transposition des valeurs caractéristiques, avec retard d'un temps variable. Enfin les valeurs caractéristiques peuvent être simultanément permutées et retardées. Ce brouillage des valeurs caractéristiques, soit par la méthode de substitution, de transposition ou combinée est réalisée par les appareils déjà décrits.

La marche synchrone de l'interrupteur de brouillage de l'émetteur et du récepteur est assurée par notre système de synchronisation, qui a fait ses preuves, même pour des commutations très rapides.

Bien que la parole transmise par le procédé des valeurs caractéristiques soit compréhensible, certaines finesses sont perdues car l'analyse et la synthèse sont

faites avec un nombre relativement faible de bandes de fréquences.

VII^o PROGRAMME DE BROUILLAGE ET SYNCHRONISATION.

Dans les procédés de brouillage décrits (substitution, transposition, procédés combinés et par valeurs caractéristiques) il s'agit de la permutation des canaux d'émission ou d'éléments du message selon un chiffre convenu que l'on change périodiquement pour augmenter le secret.

Ce changement se fait selon un certain programme convenu entre l'émetteur et le récepteur, et consigné d'une manière convenable. L'enregistrement du programme sur une bande perforée présente divers avantages par rapport aux nombreuses autres méthodes. La bande de longueur quelconque peut être collée bout-à-bout, c'est-à-dire, que la période des changements de clef peut être choisie à volonté et n'est donc pas connue de l'indiscret, ce qui lui complique le déchiffrement. Cette bande perforée est facile à faire en marquant un jeu de clef à l'aide du perforateur. Elle est aisément interchangeable, légère et rapidement détruite si c'est nécessaire.

Chaque combinaison de trou est, en général, liée à une clef, mais on peut aussi lire simultanément deux ou plusieurs bandes, de façon que chaque combinaison des signes lus simultanément corresponde à une clef déterminée. Si les bandes perforées n'ont

pas le même nombre de trous, on obtient une très grande période du programme de brouillage résultant.

La lecture de la bande et la transposition en un décalage des éléments ou des canaux, se font d'après les méthodes connues. Le changement de clef doit naturellement être simultané dans les deux appareils. C'est pourquoi une synchronisation parfaite du récepteur et de l'émetteur est de toute importance. Des essais approfondis nous ont montré qu'avec des éléments de 0,05 à 0,1 s des erreurs de synchronisme de 5 % de la longueur de l'élément ont déjà un effet perturbateur et que la compréhension baisse beaucoup pour des erreurs de 10 %.

C'est pourquoi nous avons réalisé un système de synchronisation qui garantit une exactitude 1 à 2 % de la longueur de l'élément, soit de 1/1000 s. La synchronisation subsiste lors de perturbations ou d'interruption passagères de la transmission durant de 0,5 à 1 min. C'est-à-dire que le déchiffrement n'est pas troublé par une transmission peu favorable.

VIII^o L'UTILISATION PRATIQUE DES PROCÉDÉS DE BROUILLAGE.

Le choix du système de brouillage est déterminé par les conditions posées dans chaque cas. La méthode de substitution, connue et utilisée déjà depuis des années, a été considérablement améliorée par notre procédé de synchronisation qui a permis de réduire la durée des éléments. Elle convient dans les cas où des essais de déchiffrement qui nécessitent une grande routine et un appareillage compliqué ne sont pas à craindre. Avec la méthode de transposition, le secret peut être augmenté à volonté au dépens de la rapidité de la transmission. Les communications bilatérales sont rendues alors difficiles, ce qui empêche la généralisation de cette méthode pour l'usage commercial.

Les dimensions de l'appareil et son poids sont plus faibles que ceux de l'appareil employé pour la substitution à cause de la suppression des filtres. Le brouillage par transposition convient donc particulièrement aux appareils portatifs pour lesquels certaines concessions sont faites en ce qui concerne la durée de la transmission et le secret. Dans le cas où le maximum de secret est demandé, on préfère le brouillage combiné, car les essais de déchiffrement, qui ont donné de bons résultats avec les deux autres méthodes, sont ici restés sans succès. Comme on peut alors travailler avec un retard faible, cette méthode convient aussi aux besoins pratiques. La transmission de messages sous forme de valeurs caractéristiques brouillées

offre l'avantage d'une bonne utilisation du canal de transmission avec un secret mieux gardé. La tonalité de la parole est légèrement modifiée.

Dans de nombreux cas le brouillage est aussi utilisé pour rendre secret d'autres signaux que la parole, par exemple des télégrammes morse. Le brouillage

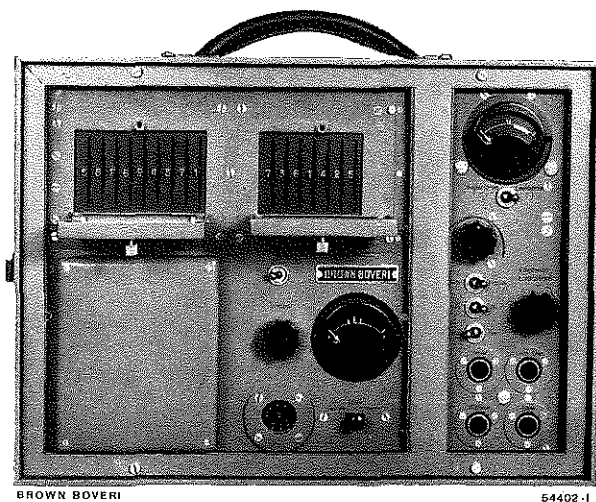


Fig. 14. — Appareil portatif de brouillage.

Des sélecteurs permettent l'ajustage de la clef. L'appareil simple et d'un service facile ne pèse que 23 kg. La sécurité de la clef est par contre limitée.

par transposition ou combiné convient sans modifications. Le brouillage de signes télégraphiques par substitution demande des appareils auxiliaires. Un seul dispositif de brouillage permet le brouillage simultané de plusieurs messages.

La figure 14 montre l'appareillage portatif de brouillage avec son équipement d'alimentation. Ce dispositif sert au brouillage et au déchiffrement. L'inversion du sens de la communication se fait simplement par pression sur un bouton se trouvant sur le combiné téléphonique. La clef est un nombre de plusieurs chiffres. Pour le brouillage de messages en télégraphie morse, on a prévu le raccordement d'un manipulateur. L'appareil travaillant d'après la méthode de transposition a donné de bons résultats lors d'essais dans des conditions de service très sévères. La compréhension du message à la réception est étonnamment bonne.

Le secret assuré par ces appareils relativement simples satisfait presque à toutes les exigences. Il est aussi supérieur à celui de la plupart des dispositifs commerciaux, beaucoup plus encombrants. Nous avons entrepris la fabrication de plus grands appareils, en appliquant les mêmes principes de construction éprouvés et les nouveaux procédés décrits de brouillage assurant le maximum de secret.

(MS 814)

G. Guanella. (J. C.)

INSTALLATIONS MODERNES DE RADIO POUR LA POLICE.

Indice décimal 621.396.99 : 351.74

Nous rappelons, tout d'abord, quelles sont les conditions sévères que pose une organisation de police moderne à ses moyens de liaison. Ensuite, nous montrons que le système de radio pour la police réalisé par Brown Boveri a de grands avantages par rapport aux anciens systèmes; ce système utilise pour la communication bilatérale deux longueurs d'ondes dans des bandes très différentes, la plus courte étant modulée en fréquence. Nous mentionnons aussi un récepteur à ondes ultra-courtes dont la mise en et hors service se fait par l'intermédiaire du réseau de téléphone automatique.

Aujourd'hui, le bon fonctionnement d'une organisation de police ne dépend pas seulement de la formation professionnelle du personnel, mais également d'un grand nombre d'appareils que la technique moderne met à la disposition de la police criminelle. Ainsi, les méthodes modernes de recherches des criminels dépendent beaucoup du fonctionnement parfait du service de renseignements, pour lequel on se servait jusqu'à présent surtout de la téléphonie par fil.

La haute fréquence a déjà plusieurs fois servi à rendre plus efficace et plus rapide la poursuite des criminels, cependant la manipulation de la plupart de ces appareils était compliquée et les perturbations relativement fréquentes. En outre, leur emploi demandait des connaissances techniques spéciales des employés; c'était une raison de plus pour que, dans plusieurs organisations de police, la télégraphie sans fil n'ait pas donné les résultats escomptés.

Brown Boveri a suivi, dans la construction des appareils de radio pour la police, des voies nouvelles à bien des points de vue, afin surtout de simplifier ces appareils et de rendre leurs manipulation et entretien semblables à ceux d'une installation téléphonique normale. Il en résulte un service sûr, même lorsque l'employé pris par ses devoirs professionnels ne serait plus en mesure de faire fonctionner correctement un appareil compliqué. Les appareils, décrits en détail ci-après, sont construits, en admettant que toutes les radiocommunications d'une organisation sont régies et contrôlées d'un poste central. Une installation de radio-police se subdivise comme le montre schématiquement la figure 1 en un *poste fixe émetteur-récepteur* placé avec ses appareils de commande et auxiliaires dans le voisinage immédiat du bureau central, en *postes mobiles émetteurs-récepteurs*, montés dans des automobiles de patrouille, des canots automobiles, des voitures pour le transport de troupe ou dans tout autre véhicule de police, en *émetteurs portatifs* que les agents isolés portent, comme une sacoche, suspendu à l'épaule par une courroie. Lorsque les conditions topographiques ne permettent pas, au poste central, une réception parfaite des émissions des postes mobiles, on installe des *récepteurs* commandés à distance en des points qui permettent la réception des émissions faites dans toute la région contrôlée

par la police. Par une construction appropriée des appareils, on est actuellement à même d'automatiser l'installation de façon que les appareils mobiles comme les portatifs puissent être appelés par groupe du poste central. Le dispositif d'appel des stations mobiles est actionné par fréquence acoustique, tandis que celui des appareils portatifs, par onde porteuse.

L'extension continuelle des centres urbains entraîne des réorganisations et agrandissements fréquents des services criminels de la police. Les installations fixes, ainsi que les mobiles doivent donc être prévues pour une utilisation variée et pour une extension ultérieure, sans grandes modifications. La dernière des installa-

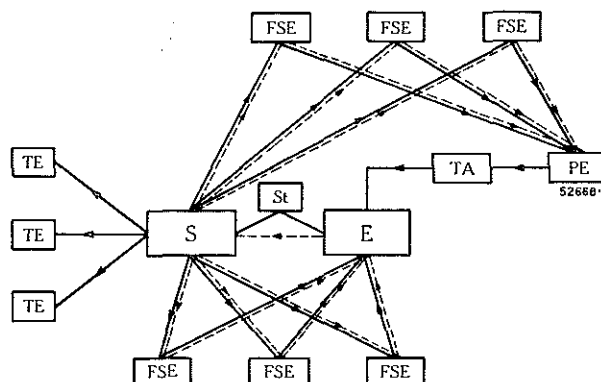


Fig. 1. — Représentation schématique d'une organisation moderne de radio pour la police.

— = Communications entre le poste central et les stations mobiles ou portatifs.

.... = Communications entre les stations mobiles, soit par le poste récepteur de banlieue ou directement par le poste central.

S = Station d'émission

E = Poste récepteur

TE = Postes récepteurs portatifs

PE = Récepteur de banlieue

FSE = Postes émetteurs et récepteurs mobiles

TA = Central téléphonique

St = Commande

Ce diagramme représente toutes les possibilités de liaison que comprend une organisation de police.

tions que nous avons construites a été mise, il y a peu de temps, en service dans une grande ville suisse, dont l'organisation de police est des plus modernes. Elle travaille dans la gamme de 110 à 220 m sur trois longueurs d'ondes fixes, choisies à volonté et sur une longueur d'onde réglable de façon continue. Cette combinaison de fréquences donne le maximum de souplesse au fonctionnement et le maximum de sécurité, vu qu'il est possible, dans certaines conditions, de mettre rapidement en service l'onde variable à la place de l'une des trois ondes fixes. La figure 2 montre l'extérieur d'un tel émetteur pour quatre longueurs d'onde et qui travaille avec trois puissances différentes entre 125 et 500 W.

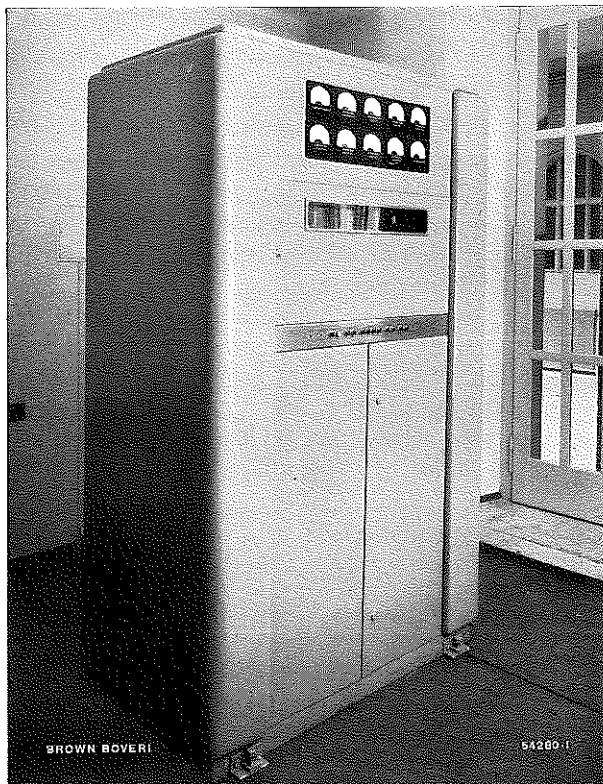


Fig. 2. — Emetteur de police de 500 W, pour téléphonie et télégraphie modulée.

Puissance réglable à 125 et 250 W. Gamme d'ondes: 110 à 220 m; trois ondes fixes et une onde variable.

La construction complètement fermée donne à l'armoire d'émission sa ligne sobre. Cette armoire contient, outre la partie d'émission et de modulation, tous les appareils auxiliaires pour le raccordement direct à un réseau triphasé.

Si l'émetteur sert à différentes utilisations, la répartition des longueurs d'ondes peut être déterminée par des raisons techniques ou d'organisation; par exemple une des ondes fixes peut être accordée sur celle de l'émetteur de radio-diffusion local, de façon que l'émetteur de police puisse remplacer celui-ci. Une deuxième onde sert spécialement au service de police, tandis que la troisième est disponible pour alarmer le service criminel ou les pompiers. La quatrième, l'onde variable, sert de réserve ou à résoudre des problèmes de commande à distance, tels qu'ils se posent actuellement, dans les localités soumises à la défense aérienne passive, pour transmettre les ordres d'obscurcissement. L'appareillage est construit de façon à permettre, sans dispositif spécial d'adaptation, l'émission directe ou d'un message transmis par une longue ligne téléphonique. L'émetteur devant remplacer l'émetteur local de radio-dif-

fusion, tous les éléments sont dimensionnés de façon que la distorsion maximum de la modulation ne dépasse pas 3 %. Pour les besoins particuliers de la police, cette qualité n'est pas nécessaire, aussi monte-t-on dans le circuit de modulation un filtre qui permet d'amortir la transmission des basses fréquences, ce qui augmente sensiblement l'intelligibilité.

La commande et le contrôle d'une installation de radio aussi étendue se fait, comme nous l'avons dit au début, d'un poste central. La figure 3 représente l'intérieur de l'un de ces postes; sur la table devant l'employé se trouvent en plus du microphone les appareils de contrôle et de commande. C'est vers ce bureau que convergent les lignes de commande et de contrôle, si bien que l'employé de service est toujours au courant du fonctionnement de toute l'installation. Il a la possibilité non seulement de régler à distance l'émetteur et les récepteurs montés dans la banlieue, mais encore d'atteindre certains groupes d'appareils et de donner l'alarme en actionnant la clef correspondante. Il établit aussi les liaisons lorsqu'une station mobile veut communiquer directement par le réseau téléphonique avec un abonné ou lorsque deux stations mobiles doivent rester reliées entre elles pendant leurs déplacements.

Ces émetteurs pour la police présentent une nouveauté intéressante; c'est la première fois, en Europe, qu'on utilise les ondes ultra-courtes modulées en fréquence pour la transmission de renseignements. De nombreuses recherches et mesures ont montré que ces ondes conviennent particulièrement bien au

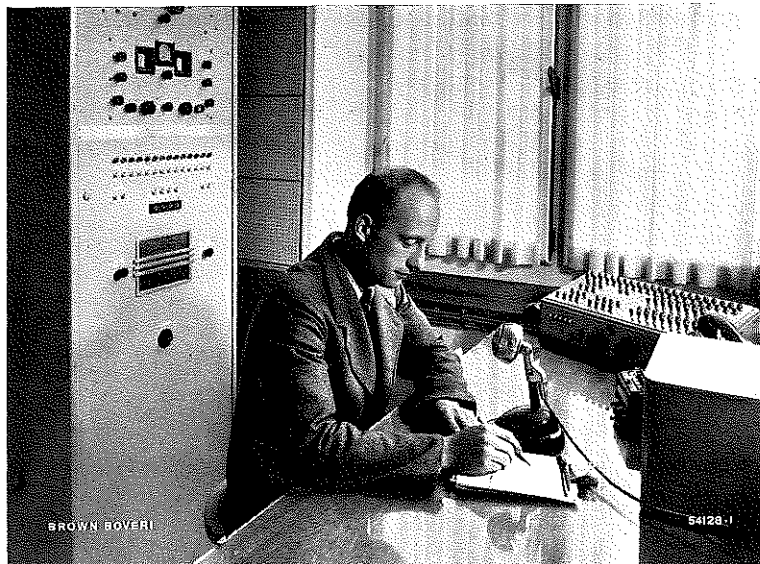


Fig. 3. — Intérieur d'un poste central moderne de radio-police.

A gauche, armoire de réception avec appareils auxiliaires, à droite, microphone, dispositif de commande et de contrôle, sélecteur de lignes téléphoniques et haut-parleur.

La disposition judicieuse de tous les appareils permet un travail des plus rapides de l'employé. Il peut atteindre sans peine de son siège tous les interrupteurs et boutons de commande.

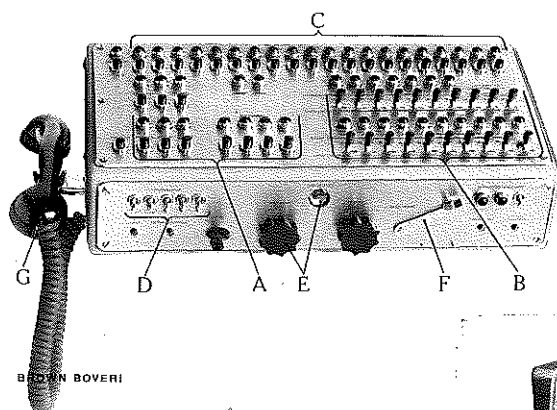


Fig. 4. — Dispositif de commande et de contrôle.

- A = Interrupteurs de commande pour l'émetteur fixe.
 B = Clés de commutation des lignes d'arrivée et des appareils de signalisation sur l'émetteur et celles pour la commande à distance des récepteurs de banlieue.
 C = Lampes de signalisation de l'appel sélectif des stations mobiles.
 D = Interrupteur de sélection préliminaire pour l'appel des groupes.
 E = Contrôle de la modulation.
 F = Sélecteurs à distance pour les récepteurs de banlieue à ondes ultra-courtes.
 G = Combiné téléphonique.

C'est le cerveau de toute l'installation. Dans cet appareil se réunissent environ 300 lignes de commande et de contrôle. D'un seul coup d'œil, l'employé est renseigné sur le fonctionnement de toute l'organisation de radio de la police.

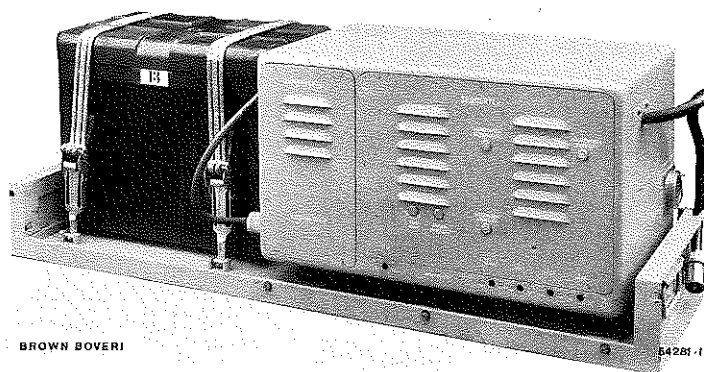


Fig. 5. — Emetteur mobile à ondes ultra-courtes, à modulation de fréquence, d'une puissance de 10 W, avec accumulateur de 105 Ah, monté sur cadre porteur.

La construction robuste et ramassée est particulièrement frappante. Le montage sur un cadre permet un échange très rapide entre tous les véhicules d'une organisation de radio de la police.

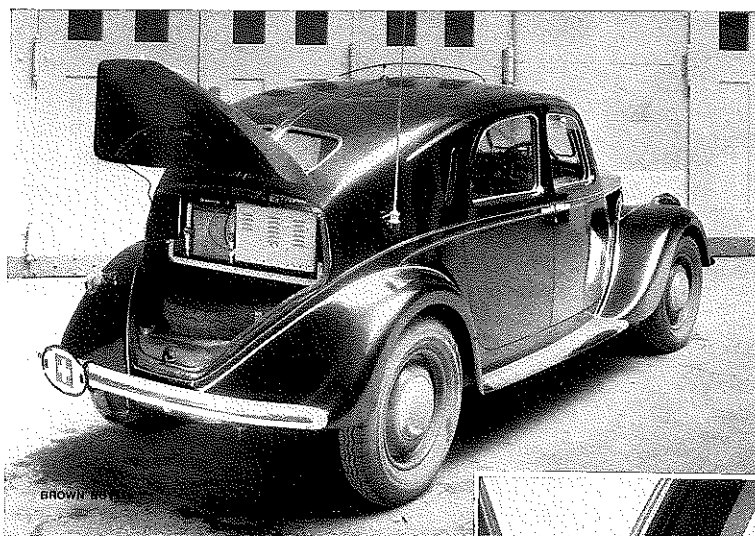
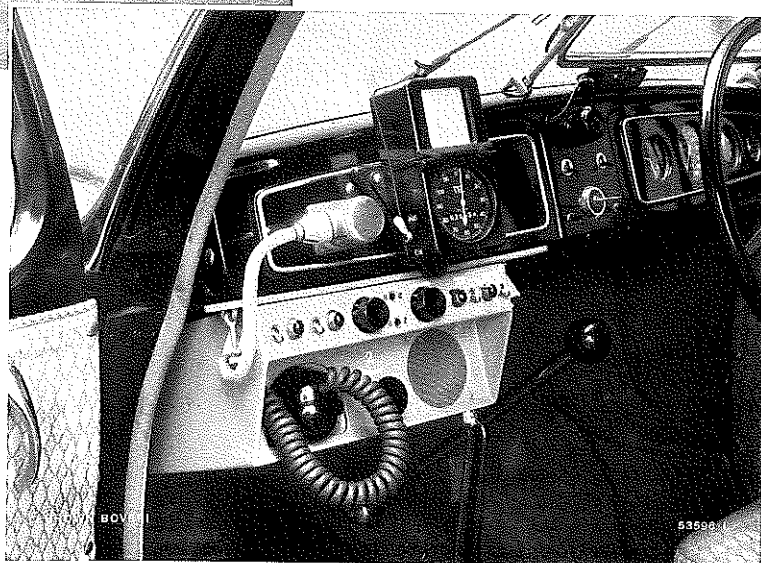


Fig. 6. — Voiture de patrouille avec émetteur à ondes ultra-courtes monté dans le coffre, antenne d'émission et antenne réceptrice montée sur le toit. L'antenne est adaptée le mieux possible à l'extérieur de la voiture et passe ainsi inaperçue. Grâce à la construction compacte, l'émetteur peut être placé même dans les plus petits coffres à bagages.

Fig. 7. — Poste de réception et de commande mobile, monté sous le tableau d'une voiture de patrouille comprenant le récepteur à ondes courtes avec haut-parleur et combiné micro-téléphonique, l'interrupteur de commande pour la commande à distance de l'émetteur, ainsi que le dispositif pour la transmission de l'appel automatique.

Les dimensions de cet appareil sont aussi réduites au minimum, car comme on le sait, sous le tableau d'une automobile de police, il n'y a jamais beaucoup de place. La disposition des signaux et des boutons de commande tient particulièrement compte d'une bonne accessibilité et d'une parfaite visibilité.



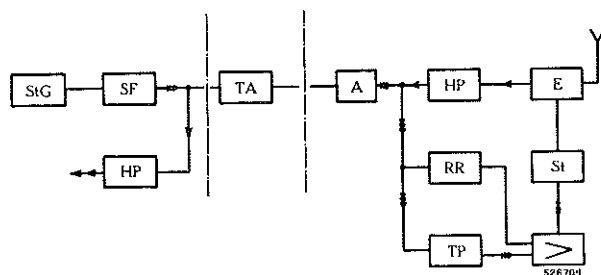


Fig. 8. — Schéma de principe de la commande à distance d'un récepteur de banlieue à ondes ultra-courtes par une ligne téléphonique (ligne choisie par sélecteur).

- | | |
|-----------------------------|------------------------|
| A = Commande. | E = Adaptation. |
| B = Générateur de commande. | F = Filtre passe-haut. |
| C = Filtre de transmission. | G = Filtre passe-bas. |
| D = Central téléphonique. | H = Relais d'appel. |
| > = Amplificateurs. | |

Comme le montre le schéma, le dispositif de commande construit par Brown Boveri travaille avec un minimum d'organes, ce qui assure un bon rendement et une bonne qualité de la transmission.

service dans les villes où les perturbations dues aux moteurs d'automobiles et aux appareils médicaux sont fortes. Cette nouvelle méthode de modulation est à l'abri de tous genres de perturbations si bien qu'elle permet de garantir actuellement dans les centres urbains de configuration normale un fonctionnement pratiquement sans perturbation. Des recherches minutieuses ont, en outre, montré qu'à l'encontre de l'idée admise,



Fig. 9. — Agent de police portant un récepteur à ondes courtes à superréaction, comportant batteries, antenne et dispositif d'alarme. Le récepteur est une réalisation particulièrement intéressante de la technique, car dans un espace très restreint sont placés tous les organes, ainsi que la source d'énergie. La figure montre clairement que le policier n'est aucunement gêné dans son activité par cet appareil.

selon laquelle les ondes ultra-courtes ne se propagent que dans le champ visuel, ces ondes conviennent tout à fait à la transmission dans les villes, malgré les nombreux obstacles. Elles permettent même la liaison entre les automobiles en marche et le poste central. La qualité de la transmission est au moins égale à celle d'une transmission par fil.

Les stations mobiles doivent répondre à des conditions sévères au point de vue mécanique, car les voitures de patrouille circulent souvent à grande vitesse dans des terrains peu praticables. L'appareil émetteur sera le plus souvent placé avec sa batterie d'accumulateurs dans le coffre de la voiture sur un châssis spécial. Le récepteur et la commande sont seuls montés sous le tableau de la voiture, car la place y est réduite. Si une organisation de police dispose de plusieurs véhicules différents, cette construction permet d'échanger très rapidement les appareils d'une voiture à une autre.

Si l'étendue de la ville est telle que les stations mobiles ne peuvent plus communiquer entre elles, ni avec le poste central, on installe dans la banlieue des récepteurs commandés à distance. Nous avons mis au point pour ces stations un nouveau genre de commande qui permet la mise en et hors service, ainsi que le réglage à distance, pendant la réception, par une ligne téléphonique. Les impulsions de commande sont émises par les générateurs du poste central et passent par un système de filtres à une ligne téléphonique choisie préalablement par un sélecteur ordinaire. Ces impulsions actionnent un système de relais dans la station de banlieue, système qui commande le récepteur à ondes ultra-courtes. Dès que la communication est terminée, la boucle téléphonique est déconnectée automatiquement et rendue à d'autres usages. Toute l'installation fonctionne d'une façon parfaitement sûre, car des mesures sont prises pour rendre impossible tout enclenchement, intentionnel ou non, par des personnes non compétentes, soit qu'elles connaissent le numéro ou qu'elles l'aient fait par erreur.

Il faut, dans un service moderne de police, pouvoir atteindre à chaque instant les agents en patrouille. C'est particulièrement important dans le cas d'une alarme générale. Les agents de police sont équipés d'un récepteur portatif pesant environ 2 kg. Le dispositif choisi permet d'alerter à chaque instant les agents, porteurs du récepteur, dans un rayon de 10 km autour de l'émetteur fixe. Il est évident que l'on peut aussi équiper des patrouilles volantes sur motocyclettes de ces appareils pour qu'elles restent en contact permanent avec le poste central.

Une organisation de police établie aussi soigneusement dans ses moindres détails prouvera toutes ses possibilités lorsque des circonstances extraordinaires se présenteront.

(MS 815)

A. Wertli. (J.C.)

ÉQUIPEMENTS PORTATIFS DE RADIO POUR L'ARMÉE.

Indice décimal 621.396.99:623

Les nombreuses qualités exigées des équipements de radio pour l'armée, tant au point de vue mécanique qu'électrique, sont énumérées ci-dessous. Seul un appareillage relativement compliqué répond aux conditions imposées, et afin d'atteindre un poids et un encombrement aussi faibles que possible, il est nécessaire de suivre des directives semblables à celles en vigueur dans la construction aéronautique.

Voici la description d'un appareil portatif, entièrement autonome, réalisé par Brown Boveri.

Il est inutile de signaler l'importance de la télégraphie et de la téléphonie sans fil pour une armée moderne. Les événements ont suffisamment prouvé que des organes de transmission très sûrs sont d'une importance décisive.

Les équipements de radio sont particulièrement intéressants, au point de vue militaire, lorsqu'ils sont de fonctionnement autonome, facilement transportables, très robustes et correspondent aux derniers progrès de la technique. En outre, un maniement simple, un dépannage rapide et aisé, constituent des propriétés très appréciées de ce genre d'appareils.

Un examen plus approfondi de ces conditions, permettra de constater qu'elles se contredisent en grande

partie. Ainsi un récepteur super-hétérodyne, à six lampes, de grande sensibilité, sélectivité et fidélité de reproduction, est certainement plus sujet aux pannes et plus lourd qu'un récepteur à une lampe « Audion » ou à super-réaction. Il en est de même pour un émetteur à sept lampes, conçu pour la télégraphie pure ou modulée, et pour la téléphonie avec microphone à cristal, comparé à un émetteur à une lampe, avec microphone à charbon. Par conséquent, il semble qu'un équipement émetteur-récepteur à une lampe, constituerait pour l'armée l'appareil de radio idéal. Effectivement, un tel appareil serait extrêmement léger, mais ses caractéristiques électriques ne seraient plus compatibles avec l'état actuel de la technique. Il est donc indispensable de construire un émetteur et un récepteur de haute qualité, à montage relativement complexe, et d'en réduire au minimum l'encombrement et le poids. Ceci ne peut être obtenu que par l'utilisation de matières premières de qualité et par l'application des directives usuelles dans la construction aéronautique.

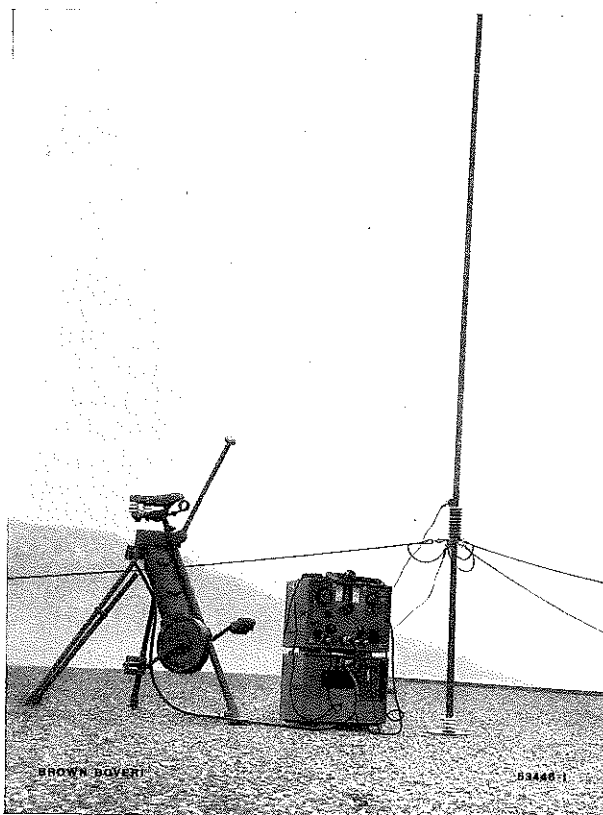


Fig. 1. — Equipement de radio pour l'armée à générateur à pédale, en ordre de marche.

Le poste se compose de l'émetteur et du récepteur, des coffrets à accessoires, de l'antenne et du générateur à pédales. Le montage et l'emploi en sont très simples et son fonctionnement est entièrement autonome.

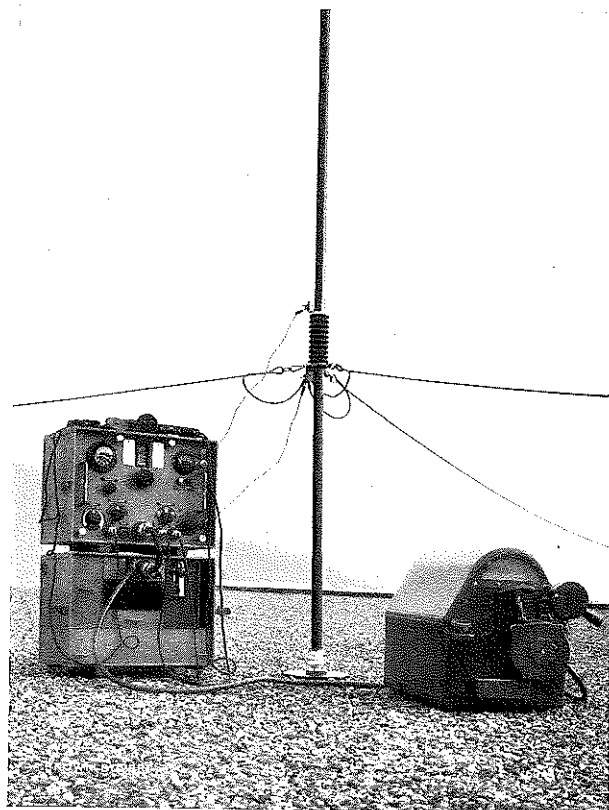


Fig. 2. — Equipement de radio pour l'armée, avec groupe électrogène à benzine, en ordre de marche.

Le groupe électrogène est monté dans un coffret insonore. La puissance est suffisante pour alimenter un dispositif de brouillage.

Brown Boveri a donc mis au point des appareils présentant un sérieux progrès par rapport à ceux utilisés jusqu'ici, en ce qui concerne l'encombrement, le poids et la puissance. Notre poste émetteur-récepteur, type SET 2/10, ne pesant que 23 kg, livre une puissance d'onde porteuse de 10 W, de sorte qu'avec une modulation de 100 %, la puissance fournie en haute fréquence est de 15 W. La possibilité de régler l'énergie de l'antenne représente un progrès évident par rapport aux dispositifs employés jusqu'ici. Cette énergie est directement réglée par un commutateur à plots, de sorte que l'on travaillera toujours avec la plus faible énergie nécessaire. Par l'emploi d'un microphone à cristal, on atteint, avec une bande de 300 à 5000 cycles, une intelligibilité de 95 % et une compréhensibilité de 99 %. Ceci est d'une grande importance dans de nombreux cas, par exemple lors de la transmission d'observations aériennes où il est nécessaire d'éviter à tout prix la répétition des questions, répétition due à une mauvaise compréhension.

L'appareil a une gamme de fréquences, dont les extrêmes sont dans le rapport 1 : 3, la plus grande longueur d'onde étant de 180 m et la plus courte de 20 m. La lecture est faite sur une échelle spirale, de 4 m de longueur, étalonnée directement en kilocycles et permettant une précision de 0,1 à 0,2 %. Cela permet d'éviter une émission de brouillage voisine de

quelques kilocycles et cela également à la fréquence d'émission la plus élevée.

Le récepteur est construit pour la téléphonie et la télégraphie. La sensibilité s'élève en moyenne à 5 μ V pour un mW de puissance de sortie. La sélectivité et la sensibilité à la fréquence d'image répondent aux exigences les plus dures.

De plus, la complète autonomie de notre appareil lui procure un avantage certain sur les appareils utilisés à ce jour. Le récepteur est alimenté par un accumulateur et un convertisseur, logés dans un coffret faisant partie de l'équipement, le convertisseur fournissant la tension anodique obtenue dans d'autres dispositifs à l'aide de piles sèches. Pour l'émission, la source de courant consiste en un générateur à pédales qui assure en même temps la charge de l'accumulateur.

Nous avons mis également au point, pour de plus grandes durées de travail, un groupe électrogène léger à benzine, monté dans un coffret insonore afin d'éviter le bruit gênant du moteur. La puissance de ce groupe est calculée de façon à ce qu'il puisse également actionner un appareillage auxiliaire tel qu'un dispositif de brouillage. Les figures 1 et 2 représentent le poste type SET 2/10 en ordre de marche, avec générateur à pédales et groupe électrogène.

(MS 816)

Karl Lutz.

LA RADIO DANS L'AVIATION.

Indice décimal 621.396.933

Après une brève esquisse de ce qu'est la radio dans l'aviation, nous mentionnons quelques constructions Brown Boveri dans ce domaine : Un émetteur à portée limitée de 150 W, un émetteur à onde courte pour le service météorologique et des buts spéciaux de 2×500 W et une commande à distance entre la place d'aviation et le poste d'émission.

Le maintien de l'horaire dans le trafic aérien avant la guerre et le survol de régions obscurcies pendant les longs raids nocturnes de bombardement de cette guerre, attirent souvent l'attention même des profanes sur les ondes électromagnétiques qui traversent l'espace pour permettre à l'équipage de l'avion de s'orienter et de trouver son chemin. Cela fait comprendre l'énorme importance des appareils qui, inconnus il y a quelques années, doivent actuellement, après une rapide période de perfectionnement, permettre dans des cas innombrables, de diriger des avions.

Le développement de la radio dans l'aviation devait être précédé de celui du vol sans visibilité. En effet, seule la technique du pilotage sans visibilité permet de voler dans le brouillard et les nuages, et de remplacer le sens d'orientation du pilote, s'appuyant surtout sur la vue, par des instruments. Ces instruments, dans une technique plus poussée, dirigent même l'avion à l'aide d'un automate remplaçant complètement le pilote.

La radiogoniométrie utilisée pour déterminer la position de l'avion paraît d'abord d'une importance particulière. Elle repose sur la mesure trigonométrique de la position de l'avion à l'aide de deux postes fixes, ce qui est possible grâce à la directivité de certains systèmes d'antenne. Selon que le repérage est fait de l'avion en relevant la direction de deux postes émetteurs fixes, ou de deux postes fixes d'après les émissions de l'avion un de ces postes transmettant l'indication de sa position au pilote, on parle de repérage direct ou indirect. L'accroissement du trafic aérien a rendu impossible l'emploi du repérage indirect surchargeant trop les stations au sol, si bien que l'avenir appartient au repérage direct. Pour cet usage, on se sert de radiophares commandés par un interrupteur horaire et émettant à intervalles déterminés leur signal d'appel puis un long trait de repérage.

Toute l'organisation de la radiogoniométrie pour le trafic aérien repose sur l'emploi d'ondes moyennes. Dans cette gamme d'ondes, il faut prendre des mesures spéciales pour supprimer l'effet d'ondes polarisées anormales qui se produisent par réflexion, afin d'éviter toute erreur possible de repérage. Ce but a été pleinement atteint. Dans le domaine des ondes

courtes et surtout des ultra-courtes, la réflexion devenait prépondérante aux grandes distances et empêchait l'utilisation de ces ondes pour le repérage. Mais dernièrement ces gammes ont aussi été utilisées avec succès pour le repérage d'avions militaires.

La radiogoniométrie est de toute importance pour le vol par temps bouché; il n'est pas moins important pour de tels vols de permettre l'atterrissage à certains endroits, grâce aux méthodes d'atterrissage sans visibilité. Après cet aperçu sur un domaine touchant de près la radio dans l'aviation, nous traiterons plus à fond cette question. Dans ce cas, une liaison quasi-optique est établie entre l'avion et le champ d'atterrissage au moyen de balises avec émetteur auxiliaire, qui émettent le signal d'approche et le signal principal sur ondes ultra-courtes. Avec les balises, on fixe clairement la ligne d'atterrissage que doit suivre l'avion pour éviter toute collision. La deuxième méthode d'atterrissage par temps bouché connue sous le nom de méthode ZZ, utilise le repérage indirect normal. Comme deux avions ne peuvent pas atterrir à la fois, le poste au sol est à la disposition de l'avion durant toutes ses manœuvres d'atterrissage.

Pour le service en ligne et le service météorologique entre l'avion et le sol et pour le service interstation, on emploie les ondes moyennes. Depuis quelques années, la gamme à la limite entre les ondes moyennes et les ondes courtes, et même les ondes courtes, ont été introduites dans l'aviation. La dernière conférence internationale de la radio du Caire en 1938 a prévu le développement de la radio dans l'aviation et lui a réservé des bandes dans toutes les gammes d'ondes. La plupart des postes au sol sont munis de téléscripteurs, si bien qu'une grande partie du service interstation peut se faire par fil.

Le service aérien, qui déjà avant la guerre était considérable, demandait la création d'un service de sécurité. En Allemagne, par exemple, le pays était divisé en douze régions de sécurité de vol ayant chacune une ou deux stations de repérage. La Suisse, comprise dans la même organisation, devait être partagée en deux ou trois régions. En pénétrant dans une de ces régions, tout avion doit s'annoncer à la station au sol à la fréquence fixée. Dans un rayon d'environ 30 km autour de la place d'aviation, la zone d'approche, l'avion ne peut pénétrer qu'avec l'autorisation de la station. Cette disposition est surtout très importante pour éviter les collisions lors d'atterrissages par temps bouché.

L'équipement des stations au sol est très varié à cause de la diversité des tâches à remplir. Un grand nombre de récepteurs sert aux liaisons avec les avions en vol, un certain nombre de récepteurs avec système d'antenne directive, au repérage. Une installation de téléscripteurs assure une partie des communi-

tions avec les autres stations au sol. En outre, quelques émetteurs sont nécessaires pour la communication avec les avions dans la zone d'approche de la région de sécurité de vol, des radiophares pour le repérage et des émetteurs à ondes courtes pour buts spéciaux, pour le service au sol et pour le service météorologique. La construction par Brown Boveri d'émetteurs pour la sécurité de vol a été bien accueillie. Une des premières commandes que nous avons reçues dans le domaine haute fréquence était celle d'un poste à faible portée pour une place d'aviation suisse. Cet émetteur (fig. 1) était le plus souvent utilisé pour de la télégraphie non modulée, mais il est aussi construit pour la télégraphie modulée et pour la téléphonie. La puissance de la porteuse pour la télégraphie non modulée est de 150 W à l'antenne et de 75 W environ dans les deux autres cas. La gamme d'ondes de 750 à 1200 m est réglable de manière continue. A l'intérieur d'une gamme réduite de 850 à 950 m dans laquelle se trouvent les fréquences le plus souvent utilisées (322, 327, 336 et 340 Kc) quatre fréquences sont fixes, ce qui permet une commutation instantanée sur l'une quelconque de ces fréquences. Tous les boutons de commandes sont verrouillés de façon à éviter toute manœuvre par une personne incompétente. En plus

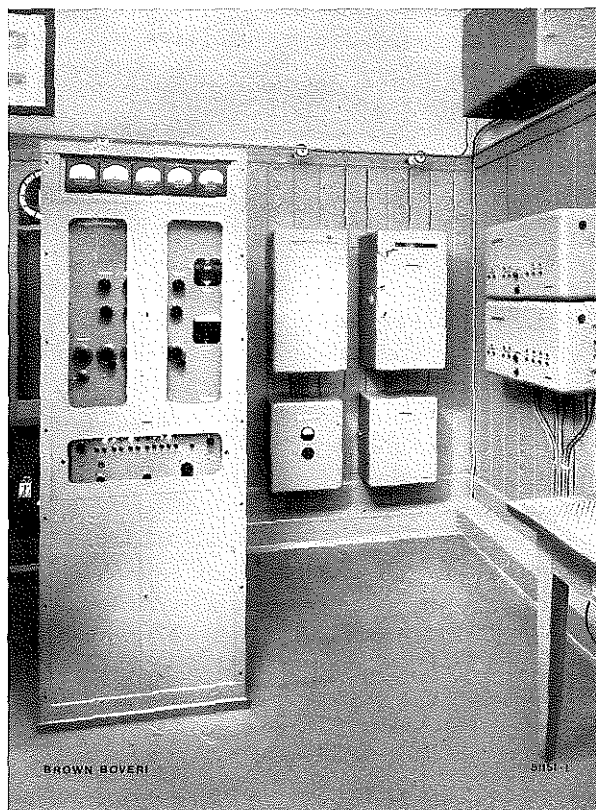


Fig. 1. — Emetteur de 150 W à portée limitée.

A côté de l'émetteur se trouve l'armoire de commande servant à cet émetteur ainsi qu'à deux autres.

du générateur haute fréquence à deux étages, on trouve encore dans l'armoire de l'émetteur, un modulateur de 10 W pour le service en télégraphie modulée et en téléphonie, ainsi que les redresseurs pour la tension d'anode et celle de grille. Le transformateur d'antenne est par contre dans un coffret

sonnel technique du poste, ce qui est important avec l'extension continuelle de ce service. L'installation comprend la commande à distance de trois émetteurs: l'émetteur à faible portée et ceux de la région de sécurité de vol et du service météorologique. Les émetteurs à faible portée et de la région de sécurité de vol peuvent être commandés de deux stations de repérage à choix; celle de l'émetteur météorologique d'un troisième bâtiment.

Etant donné la complication créée par la présence de plusieurs postes de commandement, notre système à impulsion simple (E I-System) convient particulièrement bien. Les signaux morse, les conversations et les ordres sont transmis par la même ligne, les signaux morse par courant alternatif et la commande par courant continu. Un sélecteur pas à pas assure la transmission de 25 signaux doubles comprenant chacun l'ordre et sa confirmation; sa progression synchrone est obtenue par un ordre venant alternativement du poste de commande et de la station, si bien que toute erreur de

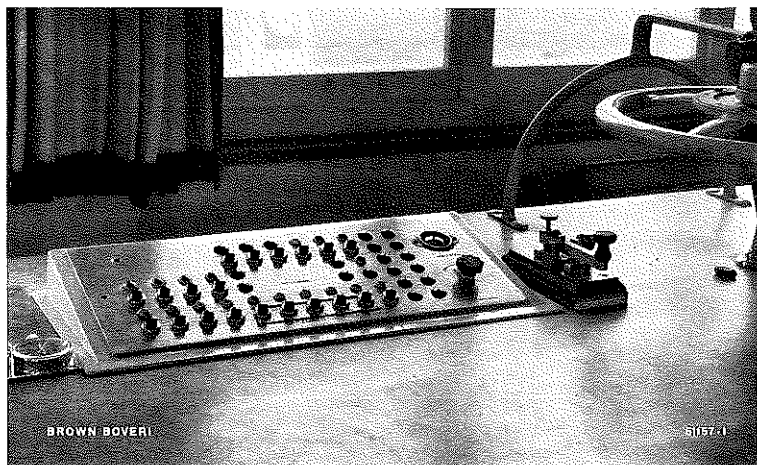


Fig. 2. — Tableau de commande à distance d'un émetteur à faible portée d'une station de balisage.

Le petit pupitre de commande monté dans la table de balisage permet la commande d'un émetteur à portée limitée se trouvant à 10 km.

séparé placé directement près du départ d'antenne; il est commandé automatiquement par relais à chaque changement de la longueur d'onde. Des mesures d'antenne ont montré que la caractéristique d'antenne varie beaucoup avec la formation de givre qui se produit fréquemment. Le transformateur est muni des moyens d'accord nécessaires pour pouvoir s'adapter à ces variations. Cet émetteur a été mis en service en été 1939. Déjà, les premiers essais ont prouvé ses qualités. La constance de la fréquence, la valeur des harmoniques supérieurs de la porteuse, la caractéristique de fréquence et le facteur de distorsion sont inférieurs aux tolérances garanties. La portée mesurée avec une puissance de 150 W atteint même 250 km. La sécurité de service est particulièrement élevée. Pendant un service de plus de deux ans, aucune perturbation ne s'est produite.

Les antennes sont gênantes dans le voisinage immédiat des places d'aviation; elles représentent un obstacle qui peut être dangereux pour les avions atterrissant par temps bouché. Sur l'aérodrome dont nous avons parlé, on a saisi l'occasion de l'installation de cet émetteur à faible portée pour relier la station d'émission à la place d'aviation par notre commande à distance. L'exploitation de l'émetteur, qui était précédemment fait, souvent avec retard, par le personnel du poste d'émission sur ordre téléphonique, est maintenant dirigée directement par le télégraphiste réglant le trafic. Cela a considérablement déchargé le per-

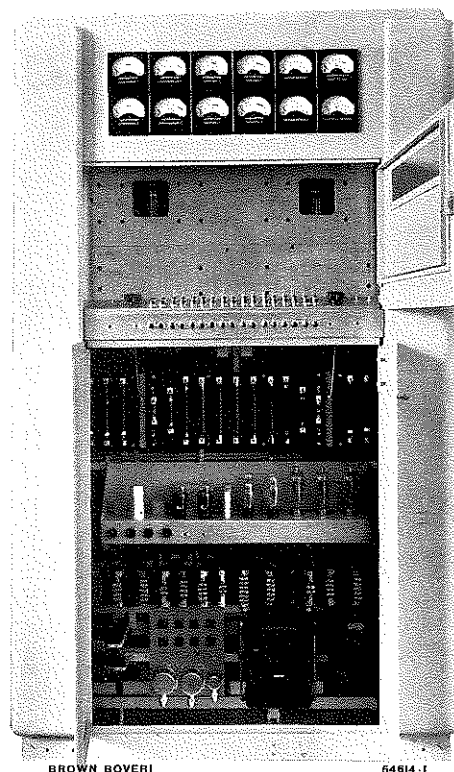


Fig. 3. — Emetteur à ondes courtes de 500 à 1000 W pour le service de sécurité du trafic aérien.

L'émetteur à ondes courtes sert aux communications avec d'autres places d'aviation et pour le service météorologique. Malgré la construction ramassée, tous les éléments sont visibles et facilement accessibles.

connexions est impossible. Un tour complet du sélecteur passant sur les 25 positions doubles dure environ dix secondes. Ce temps est nécessaire pour exécuter l'ordre de la dernière position du sélecteur. Un dispositif sélectionnant les ordres d'après leur importance évite tout retard gênant pour le service. Des parties de la commande à distance de l'émetteur mentionné sont visibles dans la figure 1. La figure 2 montre le tableau de commande d'un poste de repérage; la disposition de ce tableau est clair. Après que les défauts dans les lignes eurent été décelés par des mesures lors de la première mise en service et supprimés, le service se poursuivit pendant deux ans sans accroc et sans perturbation.

Mentionnons encore ici un émetteur à onde courte qui ne sera monté que prochainement sur la même place d'aviation. Cet émetteur est spécialement prévu pour le service météorologique et le service au sol, mais il doit aussi servir à des buts spéciaux. La gamme d'ondes de 20 à 110 m est obtenue par deux générateurs à haute fréquence montés dans la même armoire; l'un pour la gamme de 20 à 45 m, l'autre pour celle de 30 à 110 m. Chacun des émetteurs

fournit une puissance maximum de 500 W; dans la gamme commune, les deux étages finaux sont commandés par l'étage préliminaire de l'un des deux émetteurs, et la puissance maximum totale est de 1 kW. Dans chacune des gammes, deux fréquences sont fixes et peuvent être immédiatement mises en service par commutation. La figure 3 montre l'extérieur de l'émetteur terminé. Cette construction est normalisée et déjà connue par d'autres installations.

La nécessité de pouvoir procéder à une dislocation de l'installation en maintenant la commande à distance qui lors de l'emploi de l'émetteur pour des buts spéciaux doit agir sur plus de 100 km, empêchait de se servir du système EI déjà en service dans l'installation et obligea d'installer le système à impulsions combinées (système IK) qui ne transmet que des impulsions à courant alternatif. Les onze signaux nécessaires pour la commande de l'émetteur peuvent être amplifiés par les amplificateurs normaux d'une ligne téléphonique, ce qui permet d'augmenter à volonté la longueur de la ligne de commande à distance.

(MS 817)

Dr W. Lindecker. (J. C.)

MODULATION EN FRÉQUENCE.

Indice décimal 621.396.619.018.4

Exposé des propriétés de la modulation en fréquence et ses avantages par rapport à la modulation en amplitude.

Nous nous occupons depuis longtemps de ce nouveau procédé de modulation qui est en plein essor en Amérique. Nos essais et les expériences acquises avec une installation que nous avons construite confirment les caractéristiques favorables de la modulation en fréquence.

1^o INTRODUCTION.

La transmission de signaux basse fréquence peut être faite sous forme d'oscillations haute fréquence modulées. Pour de telles oscillations d'amplitude a et de pulsation ω on a en général:

$$e = a \cdot \sin \varphi(t) = a \sin \int \omega dt \quad (1)$$

Lors de la modulation en amplitude (MA) la pulsation ω reste constante tandis que l'amplitude a correspondant à la valeur instantanée du signal basse fréquence s'écarte de la valeur moyenne a_0 , d'une quantité variable a_n . Le signal peut être formé par une tension alternative de pulsation $\omega_N = 2\pi f_N$ (avec une majuscule comme indice). Alors on a dans (1):

$$\begin{aligned} a &= a_0 + a_n = a_0 (1 + m_a \cdot \sin \omega_N t); \\ \omega &= \omega_0 = \text{const.} \end{aligned} \quad (2)$$

Le taux de modulation m_a peut, suivant l'amplitude du signal, atteindre la valeur 1.

Lors de la *modulation en fréquence* (MF) l'amplitude a reste constante, tandis que la pulsation ω correspondant au signal à transmettre s'écarte de la valeur moyenne ω_0 :

$$\begin{aligned} \omega &= \omega_0 + \omega_n = \omega_0 (1 + m_f \cdot \sin \omega_N t) \\ a &= a_0 = \text{const.} \end{aligned} \quad (3)$$

L'écart $\omega_n = 2\pi f_n$ (avec une minuscule comme indice) entre la valeur instantanée $\omega = 2\pi f$ et la valeur moyenne $\omega_0 = 2\pi f_0$ de la pulsation atteint la valeur maximum $\omega_h = 2\pi f_h = m_f \cdot \omega_0 = 2\pi m_f f_0$. La déviation de fréquence, c'est-à-dire l'écart maximum de fréquence $f_h = m_f f_0$ est, en général, beaucoup plus petit que la fréquence moyenne f_0 , donc le taux de modulation m_f reste toujours petit par rapport à 1 en MF.

Une comparaison des deux genres de modulations, malgré l'analogie apparente entre les équations (2) et (3), fait remarquer de très grandes différences dans le caractère des oscillations, dans les perturbations et les déformations qui se produisent, ainsi que dans l'appareillage.

Les oscillations modulées en amplitude se décomposent comme on le sait en une fréquence porteuse f_0 et deux bandes latérales qui s'écartent de f_0 de la fréquence f_N du signal vers le haut et vers le bas. Avec la

modulation en fréquence, la décomposition en série de Fourier donne un spectre de fréquences formé de nombreuses composantes groupées autour de la fréquence moyenne et distantes l'une de l'autre de la fréquence f_N du signal. Déjà cette comparaison montre clairement que la MF demande un canal de transmission plus large que la MA.

II° DÉFORMATIONS ET PERTURBATIONS.

Des éléments non linéaires des circuits (lampes à caractéristiques courbes, etc.) produisent toujours, comme on le sait, des déformations non linéaires qui se traduisent par la formation d'harmoniques supérieurs, de combinaisons de fréquence, etc., dans le signal transmis. Lors de la transmission d'une oscillation sinusoïdale par un circuit à caractéristique non linéaire $a_z = f(a)$ il se forme, outre la fréquence fondamentale a_1 de pulsation ω_N , encore de nombreuses composantes perturbatrices a_μ de pulsation $\mu \omega_N$.

$$a_1 = A_1 \cdot \sin(\omega_N t) \quad (4)$$

$$a_\mu = A_\mu \cdot \sin(\mu \omega_N t) \quad (5)$$

L'importance des composantes perturbatrices est caractérisée par le taux de perturbations c_μ .

$$c_\mu = \frac{A_\mu}{A_1} \quad (6)$$

La suppression des déformations non linéaires présente souvent de grandes difficultés avec la modulation en amplitude.

Avec la modulation en fréquence, seule la pulsation instantanée ω détermine le signal transmis. Les variations d'amplitude dues aux caractéristiques non linéaires peuvent être supprimées par des limiteurs d'amplitude au récepteur. Comme les fréquences ne sont pas modifiées, il ne se produit aucune déformation audible du signal.

Les circuits oscillants servant à la transmission provoquent des *distorsions de phases* de l'oscillation modulée, car la caractéristique de transmission de phase de ces circuits dépend de la fréquence selon la relation $\varphi_z = f(\omega)$. Tandis que ces déformations sont de peu d'importance avec la MA, les variations de fréquence supplémentaires correspondant aux déphasages additionnels peuvent avoir pour conséquence, avec la MF, une déformation non linéaire du signal.

Avec la modulation en fréquence de l'oscillation haute fréquence, des variations périodiques de la pulsation instantanée ω_n se produisent à la cadence du signal basse fréquence $f_N = \frac{\omega_N}{2\pi}$ selon (3). Ces varia-

tions de fréquence provoquent à la sortie du circuit à distorsion de phases, des variations périodiques additionnelles φ_z qui se décomposent de nouveau en de nombreuses composantes:

$$\varphi_1 = \Phi_1 \sin(\omega_N \cdot t) \quad (7)$$

$$\varphi_\mu = \Phi_\mu \sin(\mu \cdot \omega_N \cdot t) \quad (8)$$

Lorsque la caractéristique de la distorsion de phase $\varphi_z = f(\omega)$ a la même allure dans le domaine considéré que la caractéristique de transmission non linéaire $a_z = f(a)$, le rapport des composantes φ_μ entre elles est le même que celui des composantes perturbatrices a_μ d'une des oscillations a_z déformée par la transmission non linéaire:

$$\frac{\Phi_\mu}{\Phi_1} = \frac{A_\mu}{A_1} = c_\mu \quad (9)$$

Les variations supplémentaires de phase entraînent des composantes perturbatrices correspondantes de la pulsation instantanée.

$$\omega_z = \frac{d\varphi_z}{dt} \quad (10)$$

$$\omega_1 = \frac{d\varphi_1}{dt} = \Phi_1 \omega_N \cos(\omega_N \cdot t) \quad (11)$$

$$\omega_\mu = \frac{d\varphi_\mu}{dt} = \Phi_\mu \cdot \mu \cdot \omega_N \cdot \cos(\mu \cdot \omega_N \cdot t) \quad (12)$$

Le signal basse fréquence obtenu par démodulation est proportionnel à l'écart instantané de fréquence $\omega_N + \omega_z$ de l'oscillation haute fréquence. Il contient donc des composantes perturbatrices qui sont dans le même rapport d'amplitude avec la composante utile que les valeurs maxima $\Omega_\mu = \Phi_\mu \cdot \mu \cdot \omega_N$ et $\omega_h = 2\pi f_h$ des variations périodiques de fréquence ω_μ et ω_n correspondantes. Le taux de perturbation k_μ qui caractérise le spectre de distorsion basse fréquence est donné par

$$k_\mu = \frac{\Omega_\mu}{\omega_h} = \frac{\Phi_\mu \cdot \mu \cdot \omega_N}{\omega_h} = c_\mu \cdot \frac{\Phi_1}{\omega_h} \cdot \mu \cdot \omega_N \quad (13)$$

Comme le quotient de la variation maximum de phase Φ_1 par la modification maximum de fréquence ω_h qui entraîne cette variation correspond au temps de transmission T_o , on a:

$$k_\mu = c_\mu \cdot T_o \cdot \mu \cdot \omega_N \quad (14)$$

Le temps de transmission est très faible à cause de la grande largeur de bande nécessaire de tous les circuits accordés, si bien que le produit $T_o \cdot \mu \cdot \omega_N$ reste petit par rapport à 1, même quand le canal de transmission comprend plusieurs circuits filtres. C'est pour cette raison que les déformations non linéaires

du signal peuvent être très facilement réduites avec la MF. Ce genre de modulation se recommande donc dans tous les cas où l'on exige une très grande linéarité de la transmission, par exemple pour les transmissions multiples. La MF permet, contrairement à la MA une transmission fidèle de valeurs à courant continu, comme par exemple les composantes continues des signaux de télévision qui déterminent la luminosité moyenne de l'image. Puisqu'à cause du facteur $\mu \cdot \omega_N$ de k_μ il n'y a pas de composante de distorsion pour les fréquences tendant vers zéro, la MF semble convenir particulièrement bien à la transmission de mesures: C'est pourquoi nous étudions un système de télémètre, dans lequel les valeurs mesurées sont caractérisées par les variations de fréquence d'une onde porteuse.

Les oscillations perturbatrices, bruits de fond des lampes, etc. sont particulièrement importants; ils provoquent des perturbations à la réception qui réduisent considérablement la portée de la MA. Une étude de l'influence de ces perturbations sur la MF peut être faite à l'aide de la figure 1. Le vecteur $\bar{a}_0$, qui fait à chaque instant un angle $\varphi(t)$ avec l'axe d'origine

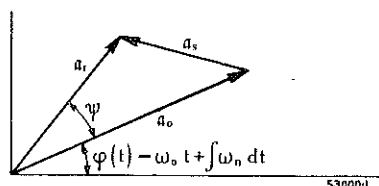


Fig. 1. — Diagramme vectoriel de l'oscillation perturbée à la réception.

$\bar{a}_0$ = Oscillation à la réception non perturbée.
 $\bar{a}_s$ = Oscillation de perturbation.
 $\bar{a}_r$ = Oscillation perturbée à la réception.
 ψ = Erreur de phase résultante.

A la réception l'oscillation perturbée est décalée de l'angle ψ par rapport à l'oscillation non perturbée.

représente l'oscillation utile reçue (1). Le vecteur $\bar{a}_s$ est une oscillation perturbatrice d'amplitude a_s . L'amplitude a_r de l'oscillation reçue peut toujours être ramenée à une valeur constante par une limitation de l'amplitude. En revanche, la phase est décalée de l'angle $\psi(t)$; c'est-à-dire qu'on obtient à la réception l'oscillation déformée:

$$e = a_0 \cdot \sin [\varphi(t) + \psi(t)] \quad (15)$$

La valeur maximum du déphasage ψ peut être déterminée géométriquement:

$$\Psi = \arcsin \left(\frac{a_s}{a_0} \right) \quad (16)$$

Pour de faibles amplitudes perturbatrices Ψ est proportionnel au taux de perturbation de la MA.

$$\Psi \approx \frac{a_s}{a_0} = p \quad (17)$$

Si l'erreur de phase ψ varie sinusoïdalement entre les valeurs extrêmes $\pm \Psi$ en cadence avec ω_s , il se produit une erreur correspondante de fréquence ω_z , variant périodiquement

$$\psi = \Psi \sin (\omega_s t) \quad (18)$$

$$\omega_z = \frac{d\psi}{dt} = \Psi \cdot \omega_s \cdot \cos (\omega_s t) \quad (19)$$

La composante perturbatrice du signal basse fréquence, due à ω_z se rapporte à l'amplitude utile maximum comme les variations maxima de la pulsation correspondante.

Le taux de perturbation est alors:

$$q_s = \frac{\omega_z}{\omega_h} = \frac{\Psi \cdot \omega_s}{\omega_h} = \frac{a_s}{a_0} \cdot \frac{\omega_s}{\omega_h} = p \frac{f_s}{f_h} \quad (20)$$

Les perturbations audibles croissent donc avec la fréquence maximum du signal f_h qui est encore transmise par le filtre basse fréquence.

$$q_H = p \frac{f_H}{f_h} \quad (21)$$

Même dans ce cas défavorable, l'amélioration est encore notable par rapport au taux de perturbation p en MA, et correspond au rapport entre la fréquence limite audible f_H et la déviation de fréquence f_h . Le taux de perturbation f_s est égal à 0 quand la fréquence de la perturbation tend vers 0, la transmission par MF de valeurs continues (par exemple, télémètre), sera donc aussi particulièrement favorable et insensible aux perturbations.

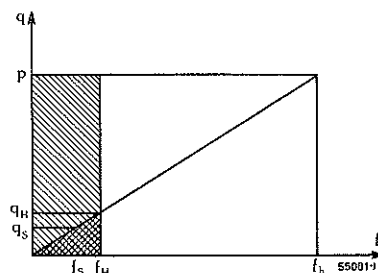


Fig. 2. — Taux de perturbation en MA et en MF.

f_h = Déviation de fréquence.
 f_H = Fréquence maximum du signal.
 f_s = Fréquence de perturbation audible.
 $p = a_s/a_0$ = Taux de perturbation, en MA.
 $q_s = p \cdot f_s/f_h$ = Taux de perturbation, en MF.
 $q_H = p \cdot f_H/f_h$ = Taux de perturbation maximum, en MF.

Pour une grande déviation de fréquence f_h , les taux de perturbation q_s en MF sont petits par rapport au taux de perturbation p en MA. Les taux de perturbations moyens sont dans le même rapport que les surfaces hachées.

Dans ces considérations, on suppose toutefois que l'amplitude a_s de la tension perturbatrice est plus petite que l'amplitude a_o du signal, car ce n'est que dans ce cas que les perturbations à la réception se manifestent seulement par certaines variations de phases ψ relativement faibles à l'intérieur des limites de Ψ selon (15) et (16). Dès que p devient plus grand que 1, les perturbations audibles croissent rapidement et couvrent le signal utile.

La relation (20) est représentée graphiquement par la figure 2. L'affaiblissement moyen Q_i de la perturbation périodique audible obtenu par la MF correspond au rapport des surfaces hachées, c'est-à-dire :

$$Q_i = \frac{P}{q_m} = \frac{2 f_h}{f_H} \quad (22)$$

On obtient un résultat semblable avec les perturbations statistiques dues, par exemple, au bruit de fond des lampes. La puissance de perturbation est ici la somme du carré des amplitudes de toutes les oscillations perturbatrices, c'est-à-dire l'intégrale du carré des ordonnées des deux surfaces hachées. L'affaiblissement Q_r du signal perturbateur correspond donc à la racine du rapport des puissances

$$Q_r = \sqrt{3} \frac{f_h}{f_H} \quad (23)$$

III^o CONDITIONS PRATIQUES.

Les relations (22) et (23) montrent que l'amélioration de la réception croît lorsque la déviation de fréquence f_h augmente. La déviation de fréquence est cependant limitée par la largeur de la bande disponible pour la transmission. En outre, le taux de perturbation p croît avec l'accroissement de la largeur de bande du récepteur, il doit toutefois rester inférieur à 1. Des essais et la pratique ont amené au choix d'une déviation de fréquence, trois à cinq fois plus grande que la fréquence maximum du signal. Pour la transmission de la parole avec une fréquence maximum de 3000 cycles, on prend une déviation de fréquence, de 10 000 à 15 000 cycles tandis que pour la transmission de la musique, une déviation de 50 à 75 kilocycles est adéquate.

Dans ces conditions le facteur Q qui caractérise la diminution du taux de perturbation, due à la MF, est égal à 10 environ, c'est-à-dire que le niveau des perturbations audibles est réduit d'environ 2,5 Neper pour la même amplitude à la réception. Comme le rapport p entre l'amplitude perturbatrice et l'ampli-

tude utile des oscillations reçues peut être augmenté d'autant, on obtient un élargissement de la zone desservie par un émetteur.

Ces résultats favorables ont été confirmés par de nombreux essais pratiques faits dans diverses conditions locales et de réception.

D'après les équations (14) et (20) l'amplitude des perturbations audibles, dues aux distorsions de phases et aux parasites, croît avec la fréquence. En outre, les composantes utiles de fréquence élevée du signal n'ont, en général, qu'une faible amplitude. C'est pour cette raison qu'il est recommandé d'augmenter l'amplitude de fréquence élevée du signal avant la modulation à l'émission et de l'affaiblir dans la même proportion après la démodulation à la réception. La déviation de fréquence n'est ainsi que légèrement augmentée, mais on obtient par contre, par l'affaiblissement mentionné, une réduction considérable des perturbations audibles de fréquences élevées particulièrement gênantes. Un avantage particulier de la MF est qu'à la réception de deux oscillations à haute fréquence modulées, la plus faible est toujours étouffée par la plus forte. Car les oscillations de plus petites amplitudes, qui peuvent être considérées comme oscillations perturbatrices, ne produisent, selon (15) et (17) que des variations relativement faibles $\psi \leq p$ de la phase de l'oscillation la plus forte, donc l'erreur de fréquence correspondante ne provoque que de faibles perturbations audibles. Pour cette raison la zone d'interférence entre deux émetteurs travaillant sur la même fréquence est beaucoup plus petite avec la MF qu'avec la MA.

Ce fait fut aussi confirmé par nos essais faits avec un émetteur à modulation en fréquence et un à modulation en amplitude. Lors d'une émission simultanée des deux émetteurs, la réception était impossible avec le récepteur à MA. Avec un récepteur à MF seule l'émission la plus forte était entendue. Avec une tension à la réception de 12,7 μV pour l'onde à MA et de 14,0 μV , c'est-à-dire 10% plus forte pour l'onde à MF, le signal reçu par le poste à MF était parfaitement compréhensible et n'était troublé que par quelques sifflements. Après avoir réduit le champ de l'onde à MF à la réception à 11,5 μV seule l'émission à MA, légèrement modulée en MF sans qu'on l'ait voulu, est reçue par le récepteur à MF et très compréhensible.

Lors du fonctionnement avec onde commune, c'est-à-dire lors de l'émission du même programme par plusieurs émetteurs, il ne se produit aucune interférence audible si l'on prend soin de n'avoir qu'une légère différence de fréquences entre les différents émetteurs. La synchronisation des oscillations émises,

nécessaire avec la MA, est superflue. Pour la radio-diffusion, les ondes modulées en fréquences permettent d'installer à la place de quelques émetteurs puissants à grande portée, de nombreux émetteurs à faible puissance montés par exemple, dans le voisinage des villes. La dépense et la puissance nécessaires pour la radiodiffusion dans un pays sont ainsi considérablement réduites.

Un autre avantage de la MF est que les variations d'amplitude dues au fading peuvent être complètement compensées par un simple limiteur d'amplitude, si bien que même un fading rapide ne provoque pas de variations notables du signal. Tandis que la surmodulation produit en MA, comme on le sait, des distorsions très désagréables, de brefs excursions au delà de la déviation maximum de fréquence sont le plus souvent sans conséquences, en MF, c'est pourquoi le taux de modulation de l'émetteur ne doit pas être aussi sévèrement contrôlé.

La grande largeur de bande du signal modulé en fréquence, qui correspond environ à 2,5 fois la déviation de fréquence maximum exclut la MF du domaine des ondes moyennes de la radiodiffusion. Pour les très hautes fréquences des ondes ultra-courtes, cette largeur de bande n'est plus un obstacle, c'est pourquoi ce sont, avant tout, les ondes de quelques mètres qui entrent en considération. Naturellement les gammes de fréquences encore disponibles des ondes plus longues conviennent aussi.

Lors de la transmission par ligne, d'oscillations modulées en fréquence, la faible sensibilité aux perturbations se manifeste par une réduction de la diaphonie. Même avec un découplage relativement très mauvais des canaux, on obtient un amortissement suffisant de la diaphonie pour le signal transmis par MF.

IV⁰ APPAREILS D'ÉMISSION ET DE RÉCEPTION.

Avec la MA, l'amplitude d'une onde porteuse à haute fréquence est ordinairement modulée conformément au signal à transmettre. Un procédé analogue ne convient pas, en général, à la MF, car la fréquence d'une porteuse donnée ne se laisse pas moduler avec des moyens simples. C'est pour cette raison que dans les émetteurs à modulation de fréquence, on modifie l'accord du générateur d'oscillation selon le signal à transmettre. L'accord est influencé, par exemple, au moyen de lampes à vide à réaction dévattée, et dont la réactance est fonction de la pente réglable. On peut aussi agir sur l'inductivité des bobines du circuit

oscillant en variant l'aimantation préalable d'un noyau de fer. Enfin l'accord peut être obtenu avec des redresseurs à éléments secs en les soumettant à une tension de polarisation dans la direction de blocage. Ils représentent ainsi des capacités à pertes relativement faibles. Des recherches approfondies nous ont montré que ces capacités peuvent être variées dans un domaine étendu par modification de la tension de polarisation.

Il est très important que la fréquence soit une fonction linéaire de la tension de commande, c'est pourquoi pour la commande par réactance avec des lampes, il faut choisir un type à très forte pente, et dont la variation de pente soit une fonction linéaire de la tension grille dans un domaine aussi étendu que possible. Bien qu'un petit déplacement de la fréquence soit de moindre importance que pour la MA à cause de la grande largeur de bande, il est bon de maintenir la fréquence moyenne à l'émission aussi constante que possible. Par comparaison de fréquence avec une oscillation à haute fréquence constante, on peut obtenir une tension de réglage qui correspond à l'erreur

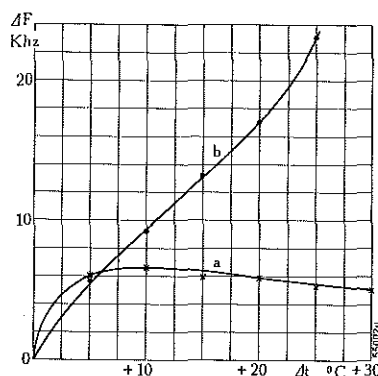


Fig. 3. — Variation de la fréquence en fonction de la température.

a = Variation de fréquence de l'oscillateur compensé.

b = Variation de fréquence d'un oscillateur piloté par quartz (courbe comparative).

La stabilité de fréquence est meilleure avec notre oscillateur compensé qu'avec un oscillateur normal piloté par quartz.

moyenne de fréquence et qui agit sur le générateur d'oscillation de façon à réduire l'erreur d'accord. Avec une compensation soignée des variations de température et de tension dans le générateur d'onde à modulation de fréquence, on obtient même sans ce dispositif de réglage une fréquence qui ne varie que de 0,025 à 0,1 ‰.

Comme une commande directe de l'accord présente certaines difficultés pour les fréquences supérieures à 20 Mc, l'oscillateur commandé fonctionne en général

à plus basse fréquence et la fréquence d'émission est obtenue par multiplication de la fréquence.

Dans une installation que nous avons construite, la fréquence moyenne de l'oscillation est d'environ 10 Mc. En doublant la fréquence à l'étage de sortie et en la triplant dans l'étage préliminaire, on obtient une fréquence moyenne d'émission d'environ 60 Mc. La stabilité mécanique thermique et électrique s'est révélée excellente dans cette installation mobile. Les vibrations pendant les déplacements ne produisent pas de modulations perceptibles. Pour une question de poids, on doit renoncer à un appareil automatique de stabilisation de la fréquence, mais une compensation de la température maintient les variations de fréquences en fonction de la température à moins de 0,1 ‰. Lorsque la température de service a été atteinte après le chauffage des lampes, l'influence de la température est même, d'après la figure 3, plus faible que pour un oscillateur à lampes normal piloté par quartz. Un montage spécial permet d'atteindre une très bonne

tions d'amplitudes correspondantes, puis démodulées. D'après un procédé de démodulation utilisé actuellement, la tension basse fréquence est obtenue par formation du produit de modulation de la tension d'entrée et de celle de sortie d'un circuit dont le déphasage est fonction de la fréquence.

Il est très important, pour une réception sans perturbation, que l'amplitude des oscillations reçues soit limitée; ainsi l'amplitude est indépendante des oscillations perturbatrices qui s'ajoutent à la valeur constante a_0 ; comme les montages connus ne satisfont qu'en partie les conditions posées, nous avons construit un limiteur d'amplitude sur de nouvelles bases; il assure une amplitude pratiquement constante à la sortie même pour de très fortes variations de la tension d'entrée. D'après la figure 4, la caractéristique d'amplitude est pratiquement rectiligne dans un grand domaine, tandis que celle des limiteurs connus est constamment courbe.

L'appareillage pour la MF est remarquable par le rendement élevé de l'étage de puissance qui peut toujours être excité complètement, les déformations non linéaires étant pratiquement sans effet. Par rapport aux émetteurs à modulation en amplitude de même puissance porteuse, une économie de 20 à 40 % peut être réalisée sur le poids, l'encombrement et la consommation d'énergie. Ces économies dépendent naturellement beaucoup des conditions d'exploitation; elles sont encore plus importantes si l'on tient compte que pour le même rayon d'action, de beaucoup plus faibles puissances d'émission suffisent.

Les appareils de réception sont, en revanche, plus compliqués pour la MF que pour la MA, d'une part, parce que la grande largeur de bande et la fréquence élevée rendent nécessaire un grand nombre d'étages amplificateurs, et d'autre part, à cause des circuits limiteurs d'amplitude et démodulateurs.

Il est avantageux pour les transmissions de nouvelles commerciales, comme pour les messages militaires qu'une partie des frais passe de l'émetteur au récepteur. Pour la radiodiffusion où le nombre des récepteurs est très grand par rapport à celui des émetteurs, ce décalage des frais n'est pas favorable. Dans les cas où le minimum de dépenses prime la qualité à la réception, la MF semble ne pas convenir. Si l'on cherche à améliorer la qualité de la réception, ce qui n'est pas possible avec la MA ou seulement avec de très grandes dépenses, la MF aura aussi beaucoup de succès dans ce domaine.

Lors de la réalisation d'une des premières grandes installations d'Europe à modulation en fréquence, nous avons eu l'occasion d'étudier à fond ce procédé de modulation. Les expériences acquises et les bons résultats obtenus nous engagent à appliquer la MF dans d'autres projets.

(MS 819) G. Guanella et J. Schwartz. (J. C.)

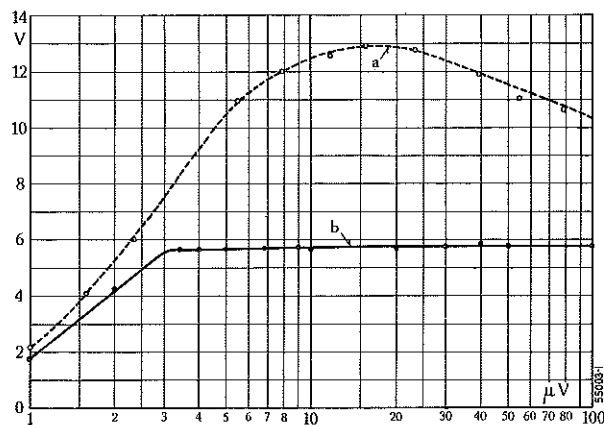


Fig. 4. — Caractéristique d'un limiteur d'amplitude.

(Amplitude de la tension limitée en fonction de la tension à l'entrée du récepteur.)

a = Caractéristique mesurée d'un limiteur d'amplitude normal.

b = Caractéristique mesurée du nouveau limiteur d'amplitude.

Notre limiteur permet une amplitude parfaitement constante de la tension dans de grandes limites, ce qui n'est pas le cas du montage habituel.

stabilité électrique. Une baisse de la tension d'anode de 25 % n'a produit qu'une variation de fréquence de 0,02 ‰.

La déviation de fréquence était fixée dans cette installation à 15 kc. Les variations de fréquences à l'oscillation sont donc de $\pm 2,5$ kc. On a mesuré les taux de distorsion pour toute la transmission, de l'entrée du signal dans l'émetteur à la sortie du démodulateur du récepteur; même avec une surmodulation et une déviation de fréquence de 20 kc, ils restent inférieurs à 1 %. Pour la démodulation, les variations de fréquence des oscillations à amplitude constante sont transformées par des circuits désaccordés en varia-

ÉTUDE THÉORIQUE DU COUPLAGE DE DEUX CIRCUITS OSCILLANTS PAR L'INTERMÉDIAIRE D'UNE LIGNE DE TRANSMISSION.

Indice décimal 621.396.611.3

On utilise fréquemment, entre les différents étages d'un émetteur, le dispositif constitué par deux circuits oscillants reliés par une ligne de transmission. Dans cet article, on étudie les conditions de fonctionnement d'un tel ensemble, pour différentes charges et diverses valeurs des éléments constitutifs.

I^o INTRODUCTION.

L'ensemble de deux circuits accordés, reliés par une ligne, se rencontre très souvent dans un émetteur, soit que la ligne ait pour but d'assurer simplement un couplage entre ces deux circuits, soit que la nécessité de réaliser un transport d'énergie H. F. à quelque distance ait rendu indispensable l'emploi d'une telle ligne intermédiaire. Ce dernier cas se présente notamment presque toujours lorsqu'il s'agit d'alimenter une antenne à partir d'un émetteur. Du côté émetteur, le circuit oscillant est formé par le circuit accordé d'anode de l'étage de sortie, qui est couplé à une extrémité de la ligne de transmission, tandis qu'à l'autre extrémité est couplé le second circuit constitué par l'antenne elle-même. Entre deux étages successifs d'un émetteur, il est souvent nécessaire, surtout en ondes courtes, de prévoir pour l'anode et pour la grille des circuits oscillants séparés qui sont alors couplés par l'intermédiaire d'une ligne, en sorte que là encore on retrouve le même dispositif. Il semble donc bien que celui-ci puisse être considéré comme un type de circuit très fréquemment employé, justifiant une étude spéciale.

Pour simple que puisse paraître ce type de circuit, on ne saurait à première vue découvrir les relations fort complexes existantes entre les grandeurs de ses différents éléments. D'ailleurs, ces relations, alors même qu'elles peuvent se laisser réduire à quelques formules,

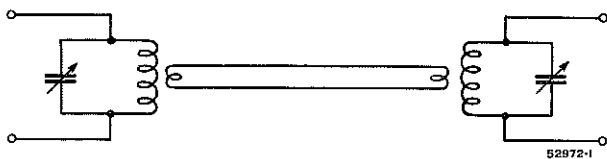


Fig. 1. — Schéma des deux circuits oscillants couplés par l'intermédiaire d'une ligne.

ne sont que de peu d'usage pour l'ingénieur et le technicien d'exploitation, si elles ne sont suivies d'une étude analytique complète des dépendances relatives des diverses grandeurs, en choisissant comme variables, celles d'entre elles qui doivent être normalement ajustées, ou qui peuvent subir des variations en cours d'exploitation. Seule une telle étude peut permettre un choix approprié des grandeurs disponibles, fait en tenant compte des exigences de l'exploitation, ce qui conduit en général, sans aucune indétermination, à une seule et unique solution.

La figure 1 montre le schéma de principe du circuit qui sera étudié dans le sens indiqué par les considérations précédentes :

II^o COTÉ ENTRÉE.

a) Notations.

- L_1 = Inductance du circuit oscillant
- L_2 = Inductance de couplage
- L_2' = Inductance en série avec la bobine de couplage
- M = Inductance mutuelle entre L_1 et L_2
- C_1 = Capacité du circuit oscillant
- C_2 = Capacité du circuit de couplage
- $\mathcal{Z}_1$ = Impédance de la branche couplée du circuit oscillant (valeur vectorielle)
- Z_1 = Impédance de la branche couplée du circuit oscillant (valeur absolue)
- Z_q = Z_1 ramené au secondaire
- R_p = Résistance parallèle de $\mathcal{Z}_1$
- R_s = Résistance série de $\mathcal{Z}_1$
- R_q = R_p ramené au secondaire
- $\mathcal{R}_2$ = Impédance de charge du circuit de couplage
- R_2 = Résistance de charge du circuit de couplage
- R_{20} = Résistance de charge optimum du circuit de couplage
- X_p = Réactance parallèle de $\mathcal{Z}_1$
- X_s = Réactance série de $\mathcal{Z}_1$
- X_q = X_p ramené au secondaire
- $X_{2\lambda}$ = Inductance totale de fuite rapportée au circuit de couplage
- X_{2s} = Réactance additionnelle en série dans le circuit de couplage
- $\mathcal{Y}_1$ = Admittance de la branche couplée du circuit oscillant
- G = Conductance de la branche couplée du circuit oscillant
- B = Susceptance de la branche couplée du circuit oscillant
- k = Coefficient de couplage
- k_s = Coefficient de couplage pour un circuit de charge série
- k_p = Coefficient de couplage pour un circuit de charge parallèle
- k_0 = Valeur minimum du coefficient de couplage
- Q = Facteur de surtension du circuit oscillant
- ω = Pulsation
- U_1 = Valeur vectorielle et valeur efficace de la tension aux bornes du circuit oscillant
- U_2 = Valeur vectorielle et valeur efficace de la tension aux bornes de la bobine de couplage
- U_c = Valeur vectorielle et valeur efficace de la tension aux bornes du condensateur du circuit de couplage
- $\mathcal{I}_1$ = Valeur vectorielle et valeur efficace du courant dans l'inductance du circuit oscillant
- $\mathcal{I}_2$ = Valeur vectorielle et valeur efficace du courant dans la bobine de couplage
- $\mathcal{I}_c$ = Valeur vectorielle et valeur efficace du courant dans le condensateur du circuit de couplage.

b) Charge constituée par une résistance pure.

Pour qu'il ne se produise pas d'ondes réfléchies sur la ligne, il faut rendre celle-ci équivalente à une ligne infiniment longue, c'est-à-dire la terminer sur une résistance pure égale à son impédance caractéristique; dans ces conditions, l'impédance d'entrée de la ligne est égale à cette même résistance. Nous allons donc tout d'abord étudier ce qui se passe lorsque le circuit de couplage est, comme le montre la figure 2, chargé par une résistance pure.

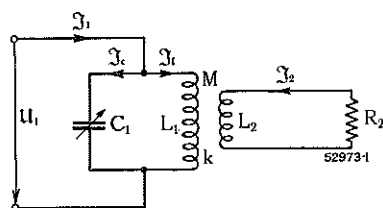


Fig. 2. — Secondaire chargé par une résistance pure.

Il existe entre L_1 et L_2 une inductance mutuelle M , et l'on a :

$$M = k \sqrt{L_1 L_2}$$

k étant le coefficient de couplage. Déterminons l'impédance

$$\mathfrak{Z}_t = \frac{U_1}{I_1}$$

de la branche inductive du circuit primaire, en tenant compte de la charge introduite par le couplage. On a pour cela les deux équations

$$U_1 = \mathfrak{Z}_1 j \omega L_1 + \mathfrak{Z}_2 j \omega M$$

$$U_2 = \mathfrak{Z}_2 j \omega L_2 + \mathfrak{Z}_1 j \omega M + \mathfrak{Z}_2 R_2 = 0$$

qui donnent après élimination de $\mathfrak{Z}_2$

$$\mathfrak{Z}_t = \frac{j \omega L_1 R_2 - \omega^2 L_1 L_2 (1 - k^2)}{R_2 + j \omega L_2} \quad (1)$$

soit en multipliant haut et bas par $R_2 - j \omega L_2$

$$\mathfrak{Z}_t = \frac{\omega_2 L_1 L_2 R_2 k^2 + j \omega L_1 [R_2^2 + \omega^2 L_2^2 (1 - k^2)]}{R_2^2 + (\omega L_2)^2}$$

Cette impédance équivaut à une résistance

$$R_s = R_2 \frac{\omega^2 L_1 L_2 k^2}{R_2^2 + (\omega L_2)^2} \quad (2)$$

en série avec une réactance

$$X_s = \omega L_1 \frac{R_2^2 + (\omega L_2)^2 (1 - k^2)}{R_2^2 + (\omega L_2)^2} \quad (3)$$

et l'accord série du circuit primaire se réalise pour une capacité

$$C_1 = \frac{1}{\omega X_s} = \frac{1}{\omega^2 L_1} \frac{R_2^2 + (\omega L_2)^2}{R_2^2 + (\omega L_2)^2 (1 - k^2)}$$

Lorsque $R_2 \gg \omega L_2$, on a $C_1 = \frac{1}{\omega^2 L_1}$ correspondant

à l'accord propre du primaire. Si R_2 s'annule, il faut pour réaccorder le circuit multiplier cette valeur de

$$C_1 \text{ par } \frac{1}{1 - k^2}.$$

Mais C_1 étant en fait en parallèle avec L_1 , il convient, pour la discussion du circuit de la figure 2, de décomposer l'impédance $\mathfrak{Z}_t$ en une résistance parallèle R_p et une réactance parallèle X_p . On a pour l'admittance :

$$\mathfrak{Y}_t = \frac{\omega L_2 R_2 k^2 - j [R_2^2 + (\omega L_2)^2 (1 - k^2)]}{\omega L_1 [(\omega L_2)^2 (1 - k^2)^2 + R_2^2]} = G + jB$$

$$\text{Avec } R_p = \frac{1}{G} \text{ et } jX_p = \frac{1}{jB}$$

X_p sera à l'accord compensé par C_1 , donc $X_p = \frac{1}{\omega C_1}$ ou $B = \omega C_1$, et :

$$R_p = \frac{L_1}{L_2 R_2 k^2} [(\omega L_2)^2 (1 - k^2)^2 + R_2^2] \quad (4)$$

$$C_1 = \frac{1}{\omega^2 L_1} \frac{R_2^2 + \omega^2 L_2^2 (1 - k^2)}{R_2^2 + \omega^2 L_2^2 (1 - k^2)^2} \quad (5)$$

La figure 3 représente la variation de R_p en fonction de R_2 avec k comme paramètre.

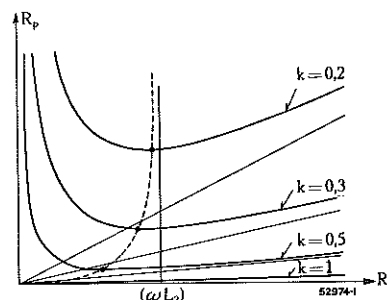


Fig. 3. — Charge parallèle du circuit primaire en fonction de la résistance et du couplage.

A chaque valeur du coefficient de couplage k , correspond une valeur bien définie de la résistance caractéristique de la ligne R_{20} , qui charge au maximum le circuit primaire, c'est-à-dire donne un minimum pour R_p . Cette valeur est :

$$R_{20} = \omega L_2 (1 - k^2) \quad (6)$$

En pratique k est rarement supérieur à 0,5 et est même la plupart du temps nettement plus petit, en sorte que l'on a sensiblement $R_{20} = \omega L_2$. R_{20} peut donc dans ces conditions être considéré comme constant dans la gamme des couplages réalisables.

Pour $R_{20} = \omega L_2$, R_p n'est plus fonction que de L_1 et k et l'on a :

$$R_p = \omega L_1 \frac{1 + (1 - k^2)^2}{k^2}$$

pour $k = 1$, R_p atteint un minimum :

$$R_{p_{\min}} = \omega L_1$$

et pour $k < 1$, R_p prend des valeurs plus grandes. Si maintenant R_2 est constamment adapté au couplage, en sorte que l'on ait $R_{20} = \omega L_2 (1 - k^2)$, R_p est une fonction de L_1 et k donnée par :

$$R_p = 2\omega L_1 \frac{1-k^2}{k^2} \quad (7)$$

et l'on voit qu'il est possible en augmentant le couplage jusqu'à $k=1$ de faire $R_p=0$.

Mis à part l'effet de désaccord (voir plus loin), R_2 est ramené au primaire sous forme d'une résistance parallèle R_p dans un rapport

$$\frac{R_p}{R_2} = \frac{L_1}{L_2 k^2} \left[1 + \left(\frac{\omega L_2}{R_2} \right)^2 (1-k^2)^2 \right] \quad (8)$$

dépendant non seulement de $\frac{L_1}{L_2}$ et k mais aussi de la valeur même de R_2 . (Dans (8), on peut écrire $\left(\frac{L_1}{M}\right)^2$ au lieu de $\frac{L_1}{L_2 k^2}$).

Le rapport $\frac{R_p}{R_2}$ augmente lorsque R_2 diminue, passant d'une valeur idéale $\frac{L_1}{L_2 k^2}$ pour $R_2 = \infty$ à une valeur double pour $R_2 = R_{20}$.

Le diagramme vectoriel de la figure 4 représente, pour une valeur constante de $\left(\frac{L_1}{M}\right)^2$, les relations existant entre R_2 , R_q et X_q , R_q et X_q étant respectivement les valeurs de la résistance et de la réactance parallèle ramenées au secondaire soit:

$$R_p = R_q \left(\frac{L_1}{M} \right)^2 = R_q \frac{L_1}{L_2 k^2} \quad (9)$$

$$X_p = X_q \left(\frac{L_1}{M} \right)^2 \quad (10)$$

La réactance X_q correspond à un désaccord du circuit primaire (relativement à sa fréquence propre); ce désaccord reste faible tant que X_q est grand.

On voit sur ce diagramme que R_q diffère peu de R_2 tant que R_2 lui-même est grand par rapport à la réactance totale de fuite $X_{2\lambda} = \omega L_2 (1-k^2)$. Lorsque R_2 diminue, R_q diminue d'abord, mais de plus en plus lentement, pour atteindre son minimum $R_{q\min} = 2R_2$ pour $R_2 = \omega L_2 (1-k^2)$, puis augmente ensuite. Dans les mêmes conditions, X_q diminue constamment, c'est-à-dire que le désaccord est de plus en plus grand. Pour $R_2 \gg \omega L_2 (1-k^2)$, X_q est très grand, et l'accord du circuit primaire reste pratiquement le même qu'en l'absence de charge. La diminution de R_2 provoque l'apparition d'une réactance parallèle, qui rend nécessaire un réaccord par augmentation de C_1 , la valeur initiale de celui-ci devant être multipliée par $\frac{2-k^2}{2(1-k^2)}$ pour

$R_2 = \omega L_2 (1-k^2)$, et par $\frac{1}{1-k^2}$ pour $R_2 = 0$; dans ce dernier cas, l'inductance de fuite est en parallèle sur le primaire.

Si l'on se fixe R_p , c'est-à-dire R_q et $\left(\frac{L_1}{M}\right)^2$ et que

l'on désire obtenir un désaccord le plus faible possible (accord indépendant du couplage) soit une valeur de X_q la plus forte possible, il faut que $\omega L_2 (1-k^2)$ soit le plus petit possible. Comme $L_2 k^2$ est fixé par $M = k \sqrt{L_1 L_2}$, on est conduit à faire k aussi grand et L_2 aussi petit que possible.

En introduisant le facteur de surtension du circuit primaire $Q = \frac{R_p}{\omega L_1}$ l'équation (7) peut s'écrire:

$$Q = \frac{R_p}{\omega L_1} = 2 \frac{1-k^2}{k^2} \quad (11)$$

Donc, dans les conditions optima, c'est-à-dire pour $\omega L_2 (1-k^2) = R_2$, il faut pour obtenir une charge donnée du primaire une valeur bien déterminée de k . Cette valeur de k est d'ailleurs un minimum, et un couplage plus fort est nécessaire lorsque R_2 diffère de sa valeur optimum.

On a la formule:

$$k_{\min} = \sqrt{\frac{2}{Q+2}} \quad (12)$$

et l'on peut calculer pour divers Q :

$Q =$	5	10	20	50	100
$k_{\min} =$	0,53	0,41	0,30	0,19	0,14

L'examen sur le diagramme de la figure 4 de l'influence de la réactance de fuite $X_{2\lambda} = \omega L_2 (1-k^2)$ amène aux résultats suivants:

$R_2 \gg \omega L_2 (1-k^2)$: une diminution de $X_{2\lambda}$ se traduit par une faible diminution de R_2 et une forte augmentation de X_q (l'accord se rapproche de la fréquence propre du primaire).

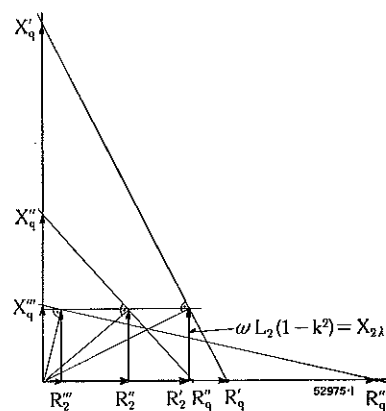


Fig. 4. — Détermination graphique de la résistance et de la réactance parallèle ramenées au secondaire (rapport $\left(\frac{L_1}{M}\right)^2$ maintenu constant), pour une charge constituée par une résistance pure.

$R_2 = \omega L_2 (1 - k^2)$: une diminution de $X_{2\lambda}$ diminue R_q sans agir sur X_q . On a alors la possibilité de régler le rapport $\frac{R_p}{R_2}$ donc R_p , en agissant sur $X_{2\lambda}$, tandis qu'une variation de R_2 agit seulement sur l'accord (déphasage de 90°).

$R_2 \ll \omega L_2 (1 - k^2)$: une diminution de $X_{2\lambda}$ produit une très forte diminution de R_q et de X_q ; l'effet sur l'accord est inverse de celui qui se produisait pour $R_2 \gg \omega L_2 (1 - k^2)$.

c) *Charge constituée par une résistance et une réactance en série.*

L'introduction d'une réactance X_{2s} en série dans le secondaire a le même effet qu'une variation de $X_{2\lambda}$.

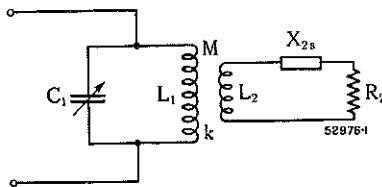


Fig. 5. — Secondaire chargé par une résistance et une réactance en série.

En reprenant alors le calcul précédent avec $\omega L_2 + X_{2s}$ au lieu de ωL_2 on obtient:

$$\frac{R_p}{R_2} = \frac{L_1}{L_2 k^2} \left[1 + \left(\frac{\omega L_2 (1 - k^2) + X_{2s}}{R_2} \right)^2 \right] \quad (13)$$

(X_{2s} ne figure pas au dénominateur).

Comme on pouvait facilement le prévoir, X_{2s} ne modifie pas le facteur $\frac{L_1}{L_2 k^2}$, mais apparaît simplement en série avec la réactance de fuite.

Si la réactance X_{2s} provient d'une inductance série L'_2 , celle-ci peut être simplement ajoutée à L_2 ; le coefficient de couplage rapporté à l'inductance secondaire totale diminue et, M n'ayant pas changé, il prend une nouvelle valeur

$$k = \frac{M}{\sqrt{L_1 (L_2 + L'_2)}}$$

L'introduction de ces expressions dans (8) donne, comme on peut facilement le vérifier, le même résultat que (13).

En prenant une capacité pour X_{2s} , on a le moyen de compenser $X_{2\lambda}$.

$$X_{2s} = \frac{1}{\omega C_2} = X_{2\lambda} = \omega L_2 (1 - k^2) \quad (14)$$

ou
$$\omega C_2 = \frac{1}{\omega L_2 (1 - k^2)}$$

On peut ainsi réduire R_q à la valeur R_2 . La réduction possible est d'autant plus grande que le rapport $\frac{\omega L_2 (1 - k^2)}{R_2}$ est lui-même plus grand. En cas de surcompensation R_q recommence à augmenter.

Les courbes de la figure 6 représentent, en fonction de $\frac{1}{\omega C_2}$ (R_2 étant pris comme unité), et avec $\frac{\omega L_2 (1 - k^2)}{R_2}$ pris comme paramètre, les variations relatives de la charge R_q , les valeurs initiales étant toutes ramenées

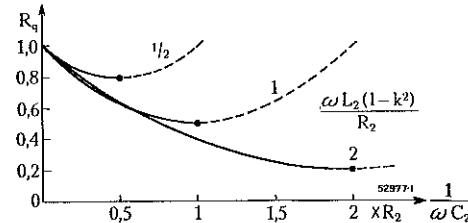


Fig. 6. — Charge relative du circuit primaire pour une capacité C_2 en série, en fonction de $\frac{1}{\omega C_2}$ et pour différentes conditions de couplage.

à 1. Les courbes ainsi tracées donnent une image claire de la réduction de la charge produite par la capacité série.

Pour réaliser cette compensation dans un domaine donné de $\frac{R_p}{R_2}$, c'est-à-dire R_2 , R_{pmin} et R_{pmax} étant fixés, on calculera les éléments comme suit:

Le rapport $\frac{R_{pmin}}{R_2}$ détermine par:

$$\frac{R_{pmin}}{R_2} = \frac{L_1}{L_2 k^2} = \left(\frac{L_1}{M} \right)^2 \quad (15)$$

la valeur de M ou de $L_2 k^2$. D'autre part:

$$\frac{R_{pmax}}{R_{pmin}} = \frac{R_q}{R_2} = 1 + \left(\frac{\omega L_2 (1 - k^2)}{R_2} \right)^2 \quad (16)$$

permet, $L_2 k^2$ étant connu, de calculer L_2 et k . La gamme de variation couverte est d'autant plus grande que k est plus petit et L_2 plus grand (c'est-à-dire que le couplage est plus faible).

d) *Charge constituée par une résistance et une réactance en parallèle.*

La charge $\mathfrak{R}_2$ étant un circuit oscillant, la valeur de la réactance de celui-ci influe aussi sur le rapport $\frac{R_p}{R_2}$. L'accord ne doit agir que sur la composante réactive parallèle, tandis que la composante résistive R_2 reste inchangée. On peut compenser par un réglage capacitif l'effet de la réactance de fuite, diminuant ainsi R_p . Le plus simple pour bien voir ce qui se passe est de transformer le circuit parallèle formé par R_2 et $\frac{1}{\omega C_2}$ (C_2 étant la capacité parallèle résultante) en un circuit série équivalent $R'_2 + jX'_2$ avec

$$R'_2 = \frac{R_2}{1 + (\omega C_2 R_2)^2}$$

$$X'_2 = - \frac{\omega C_2 R_2^2}{1 + (\omega C_2 R_2)^2}$$

Introduites dans les formules précédentes, ces valeurs donnent:

$$R_p = \frac{R_2 L_1}{L_2 k^2} \left(\frac{\omega L_2 (1-k^2)}{R_2} \right) \left[1 + \left(\frac{R_2}{\omega L_2 (1-k^2)} - \omega C_2 R_2 \right)^2 \right] \quad (17)$$

Le facteur entre crochet de cette formule (17) détermine la réduction que subit R_p du fait de la capacité parallèle C_2 . Le facteur correspondant de (13), pour le montage en série, peut s'écrire:

$$\left[1 + \left(\frac{\omega L_2 (1-k^2)}{R_2} - \frac{1}{\omega C_2 R_2} \right)^2 \right]$$

et l'on constate ainsi l'étroite parenté de ces deux expressions: les valeurs littérales y apparaissent respectivement inversées. De ce fait, les variations relatives de R_q , pourront être représentées exactement par les courbes de la figure 6, si l'on utilise pour abscisse, unité d'abscisse et paramètre, les quantités respectivement inverses de celles employées alors. On obtient ainsi la figure 7.

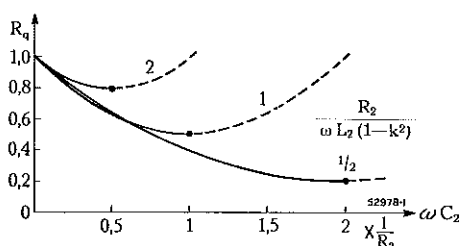


Fig. 7. — Charge relative du circuit primaire pour une capacité C_2 en parallèle, en fonction de ωC_2 et pour différentes conditions de couplage.

Contrairement à ce qui se passe avec une capacité en série, le domaine de variation est ici d'autant plus grand que les fuites sont plus faibles.

Le minimum de R_p s'obtient à nouveau lorsque $\omega C_2 = \frac{1}{\omega L_2 (1-k^2)}$ et a pour valeur

$$R_{p_{\min}} = \frac{\omega^2 L_1 L_2 (1-k^2)^2}{R_2 k^2} \quad (18)$$

Le minimum du rapport $\frac{R_p}{R_2}$ est alors

$$\frac{R_{p_{\min}}}{R_2} = \frac{\omega^2 L_1 L_2 (1-k^2)^2}{R_2^2 k^2} \quad (19)$$

Pour obtenir une meilleure vue d'ensemble des phénomènes et une idée des conditions relatives de charge, traçons le diagramme vectoriel de la figure 8 analogue à celui de la figure 4, mais valable pour un circuit parallèle. Le lieu de l'extrémité du vecteur

représentant l'impédance d'une résistance R_2 en parallèle avec une capacité variable, est un demi-cercle dont le diamètre R_2 est sur l'axe des abscisses. Le lieu de l'extrémité du vecteur représentant la réactance

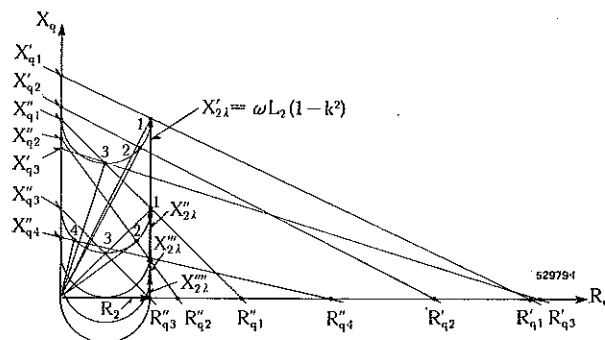


Fig. 8. — Détermination graphique des résistances et réactances parallèles ramenées au secondaire (rapport $\left(\frac{L_1}{M}\right)^2$ maintenu constant), pour une charge constituée par une résistance et un condensateur en parallèle.

de fuite $X_{2\lambda}$ est le même demi-cercle, mais décalé vers le haut de la valeur $X_{2\lambda}$. La construction tracée pour quelques points sur cette figure pour les deux cas

$$\frac{\omega L_2 (1-k^2)}{R_2} = 1 \text{ et } \frac{\omega L_2 (1-k^2)}{R_2} = 2, \text{ permet de}$$

voir que R_q , comme l'indiquaient d'ailleurs les courbes de la figure 7, décroît d'abord pour augmenter ensuite; X_q subit une variation analogue.

e) Comparaison des montages avec capacité-série et capacité-parallèle.

Il faut pour pouvoir comparer les avantages et inconvénients de deux systèmes étudier les différents cas d'emploi entrant en considération.

Voyons d'abord celui où il s'agit d'obtenir par une capacité additionnelle une certaine gamme de variation de R_p (ou R_q). Cette gamme, étant par exemple de 1 : 5 (0,2 : 1), nécessite dans le cas du circuit série un rapport $\frac{\omega L_2 (1-k_s^2)}{R_2} = 2$ (en

général = a), et une variation de $\frac{1}{\omega C_2}$ de 0 à $2 R_2$,

correspondant à une variation de C_2 de l'infini à $\frac{1}{2 \omega R_2}$; dans le cas du circuit parallèle, il faut un

rapport $\frac{\omega L_2 (1-k_p^2)}{R_2} = \frac{1}{2}$ (en général $\frac{1}{a}$) et une

variation de ωC_2 de 0 à $\frac{2}{R_2}$ correspondant à une

variation de capacité de 0 à $\frac{2}{\omega R_2}$.

On remarque tout d'abord que les variations de capacité sont de sens inverse. On peut dire ce qui suit au sujet du dimensionnement du condensateur pour la valeur R_{pmin} : Son volume est déterminé par le nombre de VA; pour faire alors la comparaison, il faut considérer une résistance donnée R_2 , ayant à ses bornes une tension de service donnée U_2 , et traversée

par un courant donné $I_2 = \frac{U_2}{R_2}$ avec dans le pre-

mier cas un condensateur $C_2 = \frac{1}{2 \omega R_2}$ (en général $\frac{1}{a \omega R_2}$) en série, et dans le second cas un conden-

sateur $C_2 = \frac{2}{\omega R_2}$ (en général $\frac{a}{\omega R_2}$) en parallèle,

on obtient alors pour les VA.

$$1^{er} \text{ cas } I_c = I_2 \quad U_c = I_2 \cdot a R_2 \quad U_c I_c = \frac{I_2^2 a R_2}{2}$$

$$2^{ème} \text{ cas } U_c = U_2 \quad I_c = \frac{a U_2}{R_2} \quad U_c I_c = U_2^2 \frac{a}{R_2} = \frac{I_2^2 a R_2}{2}$$

L'encombrement des condensateurs est dans les deux cas exactement le même. Voyons encore ce qui se passe lorsqu'on fait varier C_2 : Dans l'un et l'autre cas, la tension aux bornes du circuit oscillant reste en général constante lorsqu'on augmente R_p , cette augmentation équivaut donc à une diminution de R_2 . Si dans le second cas on agit sur un condensateur rotatif de construction normale, la tension $U_c = U_2$ aux bornes de celui-ci diminue avec la réduction de sa capacité, en sorte qu'il suffit de le dimensionner pour $R_p = R_{pmin}$. Dans le premier cas, le même mode de construction du condensateur conduirait pour pouvoir réaliser l'augmentation de capacité nécessaire au réglage, à lui donner des dimensions notablement peu grandes. Mieux vaut alors faire le réglage par variation de la distance entre plaques; la capacité croît en proportion inverse de la distance, tandis que, pour un courant constant, la tension, qui est inversement proportionnelle à la capacité, décroît proportionnellement à la distance; si l'on tient compte de la réduction du courant lorsque C_2 augmente, on voit que là aussi il suffit que le condensateur soit dimensionné pour $R_p = R_{pmin}$.

Les condensateurs ont donc le même encombrement pour les deux montages, mais différent par leur mode de construction.

Il reste encore à examiner les exigences auxquelles doivent répondre les bobines. R_{pmin} est le même dans les deux cas et L_1 en tant qu'appartenant au circuit oscillant est fixé par d'autres considérations (Q du circuit) et peut donc être considéré comme constant. Établissons les formules nécessaires:

$$1^{er} \text{ cas } R_{pmin} = \frac{R_2 L_1}{L_{2s} k_s^2} \quad \text{voir (15)}$$

$$2^{ème} \text{ cas } R_{pmin} = \frac{\omega^2 L_1 L_{2p} (1 - k_p^2)^2}{R_2 k_p^2} \quad \text{voir (18)}$$

Pour éliminer ωL_1 , introduisons $Q_{min} = \frac{R_{pmin}}{\omega L_1}$:

$$1^{er} \text{ cas } Q_{min} = \frac{R_2}{\omega L_{2s}} \cdot \frac{1}{k_s^2} \quad (20)$$

$$2^{ème} \text{ cas } Q_{min} = \frac{\omega L_{2p}}{R_2} \frac{(1 - k_p^2)^2}{k_p^2} \quad (21)$$

Le facteur $(1 - k_p^2)^2$ mis à part, les deux formules contiennent les facteurs réciproques $\frac{R_2}{\omega L_{2s}}$ et $\frac{\omega L_{2p}}{R_2}$.

Cette réciprocité est nécessaire, car, pour obtenir dans les deux cas une même gamme de variation de R_p , on est conduit aussi à la relation réciproque:

$$\frac{\omega L_{2s} (1 - k_s^2)}{R_2} = a = \frac{R_2}{\omega L_{2p} (1 - k_p^2)}$$

En combinant les trois dernières équations, on obtient la condition très simple:

$$k_s = k_p$$

Le coefficient de couplage doit être le même dans les deux cas; la disposition des bobines doit donc aussi être la même. L'essentielle différence est que L_2 doit être notablement plus grand pour le montage série que pour le montage parallèle; de plus, dans le premier cas, la tension aux bornes de L_2 est plus grande et le courant qui traverse L_2 plus faible. Constructivement, cela signifie qu'il faut bobiner dans le même espace ($k = \text{constant}$) un plus grand nombre de spires et de plus faible section dans le premier cas que dans le second.

Il convient de ne pas perdre de vue qu'après tout réglage de R_p obtenu en variant C_2 , un réaccord du circuit primaire est nécessaire. Comme le montre la figure 4, la position de l'accord aux deux extrémités de la gamme, soit pour R_{pmax} et R_{pmin} , est la même dans les deux cas.

En comparant les deux montages en ce qui concerne la possibilité de réaliser un R_{pmin} (c'est-à-dire un Q_{min}) donné, indépendamment de toute question de variation de la charge, on arrive à un résultat analogue. Ce problème présente de l'intérêt lorsque le Q demandé est trop faible pour pouvoir être obtenu sans faire usage d'une capacité additionnelle. Sans capacité on a:

$$Q_{min} = 2 \frac{1 - k_0^2}{k_0^2} \quad (11)$$

Ce problème se présente encore lorsque, pour quelque raison (distance nécessitée par une tension élevée par exemple), k_0 ne peut atteindre la valeur nécessaire. Pour faire cette comparaison, on peut donc considérer k_0 comme plus ou moins fixé. On a de nouveau les formules :

$$\text{Capacité en série} \quad Q_{\min} = \frac{R_2}{\omega L_{2s}} \frac{1}{k_s^2} \quad (20)$$

$$\text{Capacité en parallèle} \quad Q_{\min} = \frac{\omega L_{2p}}{R_2} \frac{(1-k_p^2)^2}{k_p^2} \quad (21)$$

L'effet essentiel de l'introduction d'une capacité en série ou en parallèle est ainsi de remplacer le facteur 2

de la formule (11) par les facteurs $\frac{R_2}{\omega L_{2s}} \frac{1}{(1-k_0^2)}$

respectivement $\frac{\omega L_{2p}}{R_2}$.

Les considérations du cas précédent sont sans autre applicables et montrent que pour un k constant et le même $Q_{\min}$, les deux montages demandent des condensateurs de même encombrement. *Cette comparaison conduit donc à nouveau à la conclusion que les deux schémas sont équivalents, avec cette seule différence que le montage série demande une bobine L_2 ayant davantage de spires de plus faible section que le montage parallèle.* Le choix de la variante la plus avantageuse dans un cas particulier dépendra de considérations constructives, de la valeur absolue de L_2 ainsi que de la fréquence. On fera usage du schéma série, pour l'augmentation d'inductance qu'il permet, en ondes courtes, ou l'inductance d'une seule spire peut déjà représenter une valeur de L_2 trop grande, plutôt qu'en ondes longues où les conditions sont inverses.

Nous allons considérer en troisième lieu le montage à couplage variable, par rotation ou translation des bobines. Ce montage a une importance particulière par suite des grandes variations de charge qu'il permet de réaliser. Dans un tel montage, le coefficient de couplage peut être amené à zéro, c'est-à-dire que le primaire peut être complètement déchargé ($R_p = \infty$, marche à vide). Pour obtenir une parfaite comparaison des schémas série et parallèle, il faut partir à nouveau dans les deux cas de la même valeur de $R_{p\min}$, en sorte que les résultats puissent sans autre faire suite à ceux précédemment obtenus. La seule différence provenant du fait que les bobines sont mobiles est que le coefficient de couplage est limité à des valeurs plus faibles (environ 0,2 à 0,3 pour des bobines à translation). Restent à étudier les différences des conditions de fonctionnement existantes pour l'un et l'autre des montages. Le montage série se révèle ici nettement plus avantageux, en ce sens que les variations de la

charge réactive (autrement dit les désaccords du circuit) causées par les variations du couplage, sont très faibles. *Le montage série assure une indépendance pratiquement complète de l'accord et du couplage, tandis qu'avec le montage parallèle, on doit toujours, sauf pour quelques cas spéciaux, tenir compte d'une certaine dépendance entre les deux réglages.* Ceci se voit clairement sur les diagrammes figures 4 et 8. Supposons d'abord $k = 0$, le circuit primaire est complètement séparé du circuit de charge, et accordé sur sa fréquence propre. Pour que cet accord reste inchangé lorsque k augmente, il faut que le couplage ne fasse pas apparaître de composante réactive au primaire. Comme on l'a vu précédemment, la réactance de fuite $X_{2\lambda}$ est, pour le montage série, compensée exactement par la réactance capa-

citive $\frac{1}{\omega C_2}$; les vecteurs $X_{2\lambda}$ de la figure 4 sont

ainsi nuls et la construction donne $X_q = \infty$, c'est-à-dire un désaccord nul. Comme la compensation

$X_{2\lambda} = \frac{1}{\omega C_2} = \omega L_2 (1-k^2)$ n'est exacte que pour le

couplage maximum, il reste en fait une légère influence, mais celle-ci est pratiquement négligeable tant que $\omega L_2 (1-k^2)$ n'est pas très grand par rapport à R_2 . On doit en tout cas supposer pour le cas le plus défavorable que l'accord de C_2 a été fait pour le couplage maximum. Pour des couplages plus faibles ($1-k^2$) augmente et la compensation n'existe plus; cependant, la variation est très faible: aussi pour k passant de 0,2 à 0, ($1-k^2$) varie de 0,96 à 1, c'est-à-dire de 4 % seulement. Même lorsque $\omega L_2 (1-k^2) = R_2$, la charge réactive maximum, qui se produit pour k compris entre 0 et 0,2, c'est-à-dire alors que l'erreur de compensation est inférieure à 4 %, reste négligeable. La construction du diagramme vectoriel pour une erreur de compensation de 2 % donne une valeur de X_q 25 fois plus grande que celle de $R_q =$ pratiquement R_2 .

Pour le circuit parallèle, les conditions sont entièrement différentes. On a construit, sur la figure 9, les vecteurs X_q et R_q en faisant pour le couplage maximum $\omega L_2 (1-k^2) = \frac{R_2}{2}$ (ce qui donne le même $R_{p\min}$,

que $\omega L_2 (1-k^2) = 2 R_2$ pour le montage série); on peut voir ainsi que X_q est seulement environ deux fois plus grand que R_q . Cela signifie qu'au couplage maximum l'accord propre du primaire diffère trop de l'accord vrai pour pouvoir être conservé. Même si l'accord primaire est pris pour le couplage maximum, pour $k = 0$ le désaccord capacitif devient tout aussi grand.

On a une possibilité d'obtenir la même indépendance du réglage pour le montage parallèle que pour le montage série, en accordant C_2 non plus pour $R_{p\min}$,

mais de telle sorte que le vecteur résultant Z_q soit sur l'axe des abscisses; c'est ce qui a été dessiné en pointillé sur la figure 9, avec les valeurs C'_2 , Z'_q , R'_q

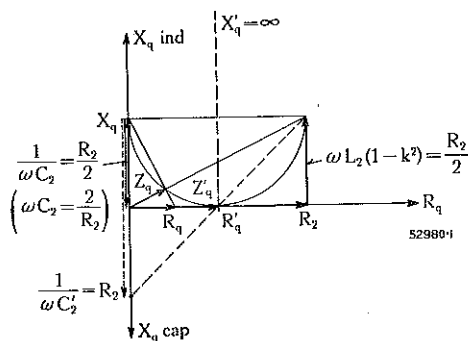


Fig. 9. — Charge parallèle R_q et X_q pour R_{pmin} , ainsi que de R'_q et X'_q pour un désaccord nul, lorsque $\omega L_2 (1 - k^2) = R_2 / 2$.

et X'_q ; C'_2 vaut alors la moitié de C_2 . Ce procédé n'est pourtant possible que lorsque le cercle touche ou coupe l'axe des abscisses, c'est-à-dire lorsque

$$\omega L_2 (1 - k^2) < \frac{R_2}{2}.$$

Mais l'avantage d'avoir des réglages indépendants se paye dans ce cas par l'impossibilité d'obtenir R_{pmin} (dans l'exemple précédent $R'_p = 2 R_p$), et par la nécessité d'un procédé compliqué pour le réglage de C_2 .

La comparaison serait incomplète si l'on ne faisait ressortir encore un autre avantage du montage série sur le montage parallèle: Dans le cas du montage série, une variation de la résistance de charge R_2 est ramenée au primaire avec son amplitude et sa phase exacte. Lorsque la compensation est pratiquement réalisée ($X_{2\lambda} = \frac{1}{\omega C_2}$) on a en effet $R_2 = Z_q$.

Ainsi que le montrent les figures 10 et 11, les conditions sont moins favorables pour le montage parallèle. Si, comme sur la figure 10 le réglage de C_2 est fait, pour obtenir la charge maximum, c'est-à-dire R_{pmin} , la résistance de charge du primaire diminue, lorsque R_2 augmente. (Un déphasage de R_2 modifierait à la fois X_q et R_q .) Si l'on suppose C_2 réglé pour que

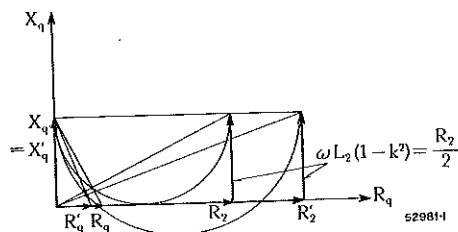


Fig. 10. — Réduction de R_q à la valeur R'_q lorsque R_2 augmentant de 40% devient R'_2 , pour le réglage correspondant à R_{pmin} .

les variations de k soient sans effet sur l'accord, on obtient la figure 11, et, dans ce cas, une variation de R_2 entraîne non seulement une variation de R_q , dans un rapport à vrai dire difficile à voir, mais aussi une variation de $X_q = \infty$ à X'_q , causant l'apparition au primaire d'une charge capacitive.

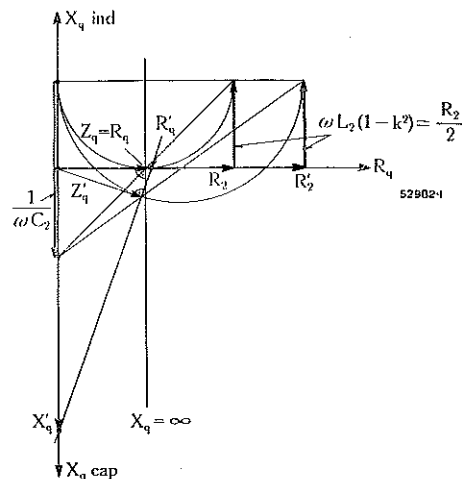


Fig. 11. — Augmentation de R_q à R'_q et production de la charge réactive X'_q , lorsque R_q augmentant de 40% devient R'_q , pour le réglage correspondant à un désaccord nul.

III^e COTÉ SORTIE.

a) Notations pour le côté sortie.

- L_1 = Inductance de couplage
- L_2 = Inductance du circuit oscillant
- M = Inductance mutuelle
- C_2 = Capacité du circuit oscillant
- C_{res} = Capacité de résonance
- ΔC = Variation de capacité correspondant au désaccord
- β_1 = Impédance de la bobine de couplage (valeur vectorielle)
- Z_1 = Impédance de la bobine de couplage (valeur absolue)
- β_2 = Impédance de charge de L_2
- R_p = Résistance parallèle provenant de β_1
- R_2 = Résistance de charge parallèle du circuit oscillant
- k = Coefficient de couplage
- $\dot{u}$ = Rapport de transformation
- ω = Pulsation
- $\dot{U}_1$ = Valeur vectorielle et valeur efficace de la tension aux bornes de la bobine de couplage
- $\dot{U}_2$ = Valeur vectorielle et valeur efficace de la tension aux bornes du circuit oscillant
- $\dot{I}_1$ = Valeur vectorielle et valeur efficace du courant dans la bobine de couplage
- $\dot{I}_2$ = Valeur vectorielle et valeur efficace du courant dans l'inductance du circuit oscillant.

b) Conditions de charge.

Comme pour le côté entrée, nous nous intéressons aux variations de l'impédance primaire et du rapport de transformation en fonction de la charge

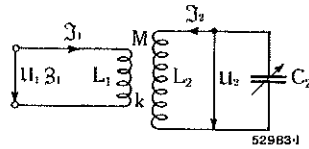


Fig. 12. — Couplage entre la ligne et le circuit oscillant.

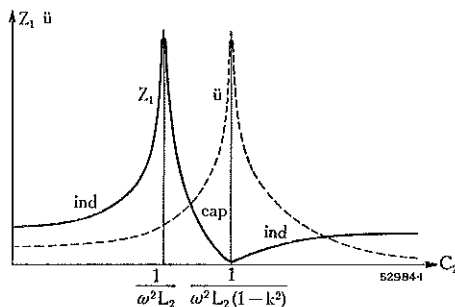


Fig. 13. — Impédance d'entrée Z_1 et rapport de transformation $\bar{u}$, en fonction de la capacité du circuit secondaire.

secondaire, celle-ci étant constituée essentiellement par le condensateur d'accord C_2 . Les grandeurs entrant en ligne de compte sont indiquées sur le schéma de la figure 12.

Pour un transformateur chargé au secondaire par une résistance R_2 , nous avons déjà obtenu l'équation (1).

$$\beta_1 = \frac{j \omega L_1 R_2 - \omega^2 L_1 L_2 (1-k^2)}{R_2 + j \omega L_2} \quad (1)$$

En introduisant maintenant $\frac{1}{j \omega C_2}$ au lieu de R_2 , il vient :

$$\beta_1 = \frac{\frac{\omega L_1}{\omega C_2} - \omega^2 L_1 L_2 (1-k^2)}{j (\omega L_2 - \frac{1}{\omega C_2})} \quad (22)$$

$$\beta_1 = j \omega L_1 \frac{\frac{1}{\omega C_2} - \omega L_2 (1-k^2)}{\frac{1}{\omega C_2} - \omega L_2}$$

Cette relation, représentée graphiquement sur la figure 13, en supposant toutefois en certain amortissement, présente les singularités suivantes :

Pour $C_2 = 0$ $\beta_1 = j \omega L_1$

pour $C_2 = \frac{1}{\omega^2 L_1}$ $\beta_1 = \infty$ résonance

pour $C_2 = \frac{1}{\omega^2 L_2 (1-k^2)}$ $\beta_1 = 0$ maximum du rapport $\bar{u}$ (voir plus loin)

pour $C_2 = \infty$ $\beta_1 = j \omega L_1 (1-k^2)$

entre $C_2 = 0$ et $\frac{1}{\omega^2 L_2}$ β_1 est inductif

entre $C_2 = \frac{1}{\omega^2 L_2}$ et $\frac{1}{\omega^2 L_2 (1-k^2)}$ β_1 est capacitif

entre $C_2 = \frac{1}{\omega^2 L_2 (1-k^2)}$ et ∞ β_1 est inductif.

La variation ΔC de C_2 pour $\beta_1 = \infty$ et 0 vaut :

$$\Delta C = \frac{1}{\omega_2 L_2 (1-k^2)} - \frac{1}{\omega^2 L_2}$$

$$\Delta C = \frac{1}{\omega^2 L_2} \frac{k^2}{1-k^2}$$

$$\Delta C = C_{res} \frac{k^2}{1-k^2} \quad (23)$$

Déterminons le rapport de transformation $\bar{u} = \frac{U_2}{U_1}$

à partir du système d'équations du transformateur :

$$U_1 = \beta_1 j \omega L_1 + \beta_2 j \omega M$$

$$U_2 = \beta_2 j \omega L_2 + \beta_1 j \omega M$$

Avec en plus pour la charge l'équation :

$$U_2 = -\beta_2 \frac{1}{j \omega C_2}$$

on obtient facilement après quelques transformations :

$$\frac{U_2}{U_1} = \frac{M}{L_1 - \omega^2 C_2 (L_1 L_2 - M^2)} \quad \text{ou}$$

$$\frac{U_2}{U_1} = k \sqrt{\frac{L_2}{L_1}} \frac{1}{1 - \omega^2 C_2 L_2 (1-k^2)}$$

La variation de ce rapport en fonction de C_2 est représentée en pointillé sur la figure 13.

Pour

$C_2 = 0$ $\frac{U_2}{U_1} = k \sqrt{\frac{L_1}{L_2}} = \frac{M}{L_1}$

$C_2 = \frac{1}{\omega^2 L_2}$ $\frac{U_2}{U_1} = \frac{1}{k} \sqrt{\frac{L_2}{L_1}} = \frac{L_2}{M}$

$C_2 = \frac{1}{\omega^2 L_2 (1-k^2)}$ $\frac{U_2}{U_1} = \infty$

$C_2 = \infty$ $\frac{U_2}{U_1} = 0.$

Tandis que pour l'entrée nous avons calculé avec le rapport des résistances, dans le cas de la sortie le rapport le mieux adapté aux nécessités pratiques est

le rapport des tensions du côté secondaire au côté primaire. Il faudra donc inverser et élever au carré le rapport obtenu pour l'entrée. Si au lieu de $\frac{1}{j\omega C_2}$ on a en général une impédance $\mathfrak{Z}_2$, il faut remplacer ωC_2 par $\frac{1}{j\mathfrak{Z}_2}$. On obtient ainsi la formule:

$$\frac{U_2}{U_1} = k \sqrt{\frac{L_2}{L_1}} \frac{1}{1 + \frac{j\omega L_2}{\mathfrak{Z}_2} (1 - k^2)} = \frac{M}{L_1 + \frac{j\omega (L_1 L_2 - M^2)}{\mathfrak{Z}_2}}$$

Si l'on remplace maintenant dans cette formule $\mathfrak{Z}_2$ par R_2 , prend le module du résultat et élève son inverse au carré, on obtient la formule (8) donnant le rapport $\frac{R_p}{R_2}$, dans lequel la résistance de charge secondaire R_2 est ramenée au primaire. $R_p = \bar{u}^2 R_2$ avec $\bar{u} = \frac{U_1}{U_2}$.

Une résistance relativement élevée R_2 , connectée en parallèle sur C_2 est ramenée au primaire avec la valeur

$$R_p = \left(\frac{U_1}{U_2} \right)^2 \cdot R_2$$

Pour la résonance par exemple on obtient:

$$C_2 = \frac{1}{\omega^2 L_2} \quad R_p = \left(\frac{M}{L_2} \right)^2 R_2.$$

Les courbes de la figure 14 représentent en fonction de C_2 la variation de X_p telle qu'elle est déterminée par l'équation (22), ainsi que la variation correspondante de R_p . La représentation de ces variations sous forme de lien de l'extrémité d'un vecteur, telle qu'elle est donnée par la figure 15, peut, pour certains buts, être plus utile. Les composantes série R_s et X_s

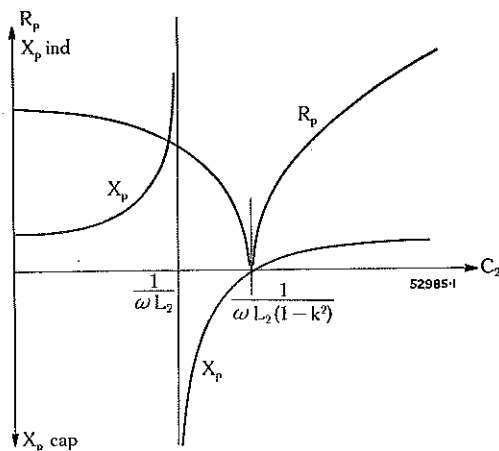


Fig. 14. — Variation des composantes parallèles X_p et R_p de l'impédance du circuit de sortie en fonction de C_2 .

utilisées dans ce dernier diagramme, se déduisent des composantes parallèles R_p et X_p de la figure 14. La transformation peut se faire graphiquement, le vecteur $R_s + jX_s$ étant perpendiculaire à la droite joignant les extrémités des vecteurs R_p et X_p correspondants.

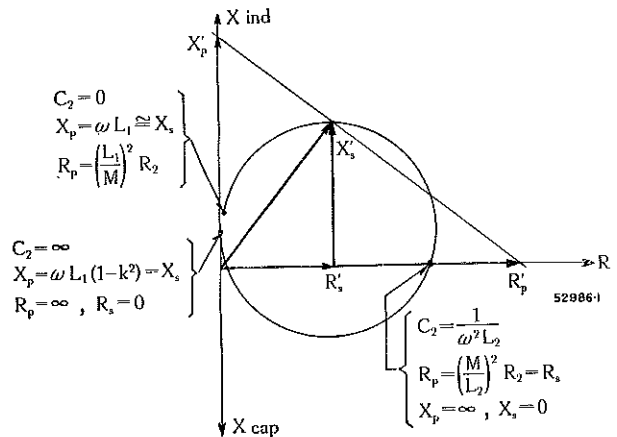


Fig. 15. — Diagramme vectoriel des composantes série X_s , R_s de l'impédance du circuit de sortie.

Pratiquement, R_2 étant connu (charge due à la grille de la lampe suivante), il faut pour obtenir une valeur donnée de la résistance primaire (impédance caractéristique de la ligne) réaliser un rapport déterminé

$$\frac{L_2}{M} = \frac{1}{k} \sqrt{\frac{L_3}{L_1}}. \quad L_2 \text{ étant la plupart du temps fixé par}$$

d'autres exigences (largeur de bande), $k^2 L_1$ doit ainsi prendre une valeur déterminée. *Lors du choix de valeurs relatives de k et L_1 , il est préférable de prendre k aussi voisin de 1 que possible, car ainsi la*

$$\text{différence } \Delta C = C_{res} \frac{k^2}{1 - k^2} \text{ séparant } C_{res} = \frac{1}{\omega L_2}$$

$$\text{de } C_2 = \frac{1}{\omega L_2 (1 - k^2)} \text{ est la plus grande possible.}$$

Dans ces conditions, les variations de C_2 au voisinage de C_{res} , auront peu d'influence sur $\bar{u}$. Il pourrait sinon arriver que l'on trouve un accord ne correspondant pas à une charge purement ohmique, en ce sens qu'avec une valeur de C_2 un peu supérieure à C_{res} , on obtienne bien une diminution de $\mathfrak{Z}_1$, causée par celle de

$$\frac{1}{\omega C_2}, \text{ amenant à des conditions d'adaptation à l'étage}$$

précédent moins favorables, donc à une réduction de tension sur la ligne, mais qu'en même temps le rapport $\bar{u}$ augmente, et vienne ainsi compenser cette réduction.

IV. FONCTIONNEMENT DE L'ENSEMBLE.

Pour un fonctionnement correct de l'ensemble, il faut d'abord que du côté sortie, la charge donnée, vue de la ligne, soit une résistance pure égale à l'impédance caractéristique de celle-ci. *Il ne peut alors se former d'ondes stationnaires sur la ligne; la tension est sensiblement constante sur toute la longueur de celle-ci et la transmission est indépendante de cette longueur.* Du côté entrée, il faut transformer la résistance caractéristique de la ligne en une résistance de charge convenablement adaptée.

Lorsque la ligne est fermée sur une résistance différant de son impédance caractéristique, l'erreur est ramenée à l'entrée avec une valeur dépendant de la longueur de la liaison. Le vecteur représentatif de l'erreur tourne autour de l'extrémité du vecteur impédance caractéristique, faisant par rapport à celui-ci une révolution complète pour une longueur de ligne égale à une demi-longueur d'onde. Par exemple, une résistance de fermeture plus faible que l'impédance caractéristique de la ligne se transforme pour une longueur $\lambda/8$ en une impédance d'entrée égale à l'impédance caractéristique avec un faible déphasage inductif (réactance série sensiblement égale à l'erreur faite sur la résistance de fin). Pour une ligne de $\lambda/4$ on trouve de nouveau à l'entrée une résistance pure présentant par rapport à l'impédance caractéristique un pourcentage d'erreur par excès, égal à celui que présentait par défaut la résistance terminale. Pour $3\lambda/8$ la résistance d'entrée se présente à nouveau comme une impédance égale à l'impédance caractéristique, mais avec un déphasage capacitif, tandis que pour $\lambda/2$ on retrouve à l'entrée la résistance terminale.

Les variations se produisant du côté sortie, par suite d'un réglage ou par toute autre cause, sont donc reportées à l'entrée, en fonction de la longueur de la ligne, mais les relations correspondantes sont compliquées et peu claires. *Par contre, si la ligne est courte, ce qui se présente très souvent lors du couplage entre deux étages successifs d'un émetteur, il est relativement simple de voir ce qui se passe.*

C'est ce que montre le diagramme de la figure 16 pour le montage de la figure 1, lorsque $\omega L_2(1-k^2) = 2R_2$ du côté entrée. Ce diagramme diffère de celui de la figure 8. Le cercle n'est plus tangent à l'extrémité de $\omega L_2(1-k^2)$, comme ce devrait être le cas si le circuit de sortie était directement connecté à la ligne, mais conformément à la figure 15, est un peu décalé par suite de la réactance de fuite due au couplage.

Sous cette influence, le minimum de R_p n'est plus exactement à l'endroit que lui assignait la figure 7; cependant cette influence reste faible. Les valeurs correspondantes du diagramme R'_q, X'_q, β'_1 sont relatives au point d'accord P' . *On voit ainsi que l'impédance de fermeture β'_1 de la ligne n'est pas une résistance pure, mais a un effet capacitif, c'est-à-dire que*

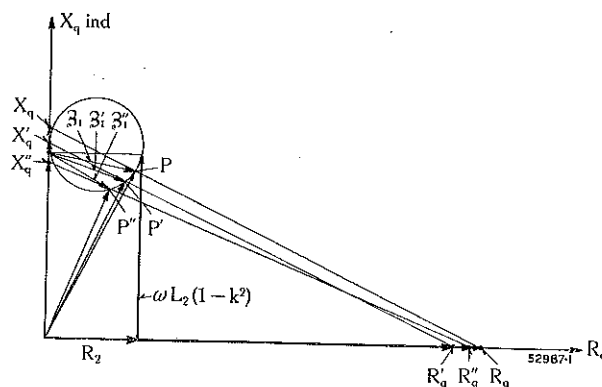


Fig. 16. — Diagramme vectoriel d'ensemble pour une ligne courte, lorsque $\omega L_2(1-k^2) = 2R_2$.

la charge maximum du circuit d'entrée s'obtient par un léger désaccord capacitif du circuit de sortie. Si maintenant le point d'accord P' correspondant à cette position est déplacé soit inductivement vers P , soit capacitivement vers P'' , R_q augmente dans les deux cas et la tension aux bornes du circuit de sortie diminue. Le R_p du circuit d'entrée étant supposé exactement adapté à la lampe précédente, ces variations resteront sans influence sensible sur la puissance. Le réglage sera donc de ce fait relativement peu critique. Cette propriété se manifeste aussi avantageusement lors de l'accord du système: Si, avec le circuit de sortie entièrement déréglé, on accorde le circuit d'entrée, l'accord consécutif du circuit de sortie pourra se faire sans qu'il soit nécessaire ensuite de reprendre l'accord du circuit d'entrée (réaccords successifs). (Il faut pour cela que X'_q coïncide avec la valeur que prend X_q pour un circuit de sortie complètement désaccordé.)

Mentionnons spécialement une propriété que l'on pourrait d'instinct penser être opposée: *La correction sur le circuit d'entrée d'une diminution de la capacité du circuit de sortie s'obtient aussi par une diminution de capacité.* Pour une telle diminution en effet, l'accord devient inductif et le déplacement sur le diagramme figure 16 se fait de P' vers P . La charge réactive du circuit d'entrée passe donc de X'_q à X_q , c'est-à-dire que l'inductance mise en parallèle augmente. Pour rétablir l'accord, il faut donc bien que $\frac{1}{\omega C_1}$ augmente aussi, c'est-à-dire que C_1 diminue.

Les conditions sont toutes autres lorsque $\omega L_2(1-k^2) < R_2$, c'est-à-dire lorsque par exemple $\omega L_2(1-k^2) = \frac{R_2}{2}$, on obtient alors la figure 17.

Supposons qu'un système quelconque de mesure permette de régler l'accord au point P' et que le côté

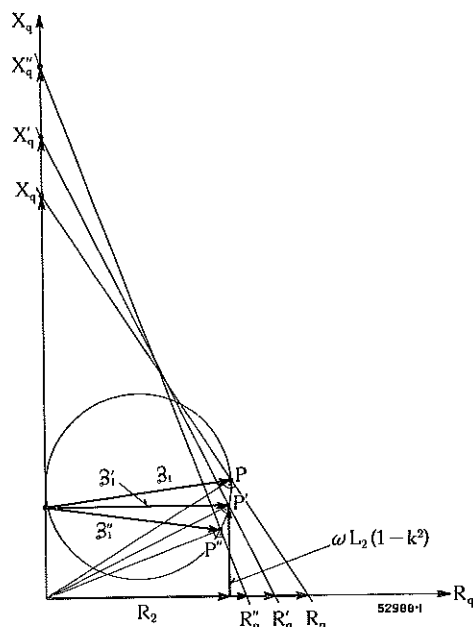


Fig. 17. — Diagramme vectoriel d'ensemble pour une ligne courte, lorsque $\omega L_2(1-k^2) = R_2/2$. Charge résistive de la ligne.

d'entrée soit adapté à la résistance R_q' . On voit alors que R_q' ne correspond pas à un minimum, donc que le maximum de tension obtenu aux bornes du circuit de sortie est dû seulement à la variation de l'adaptation. L'ordre de succession des valeurs X_q , X_q' et X_q'' qui est l'inverse de celui de la figure 16, montre que pour $\omega L_2(1-k^2) < R_2$ les désaccords relatifs des deux circuits varient en sens inverse, c'est-à-dire qu'une diminution de la capacité du circuit de sortie doit maintenant être compensée par une augmentation de la capacité du circuit d'entrée.

En renonçant au réglage spécial du point P' de la figure 17 et accordant le circuit normalement au minimum de R_p , on obtient les conditions de la figure 18. On voit que dans ces conditions, suivant les relations de la page 427, R_q et avec lui R_p , deviennent notablement plus faibles. Par contre, la charge de la ligne de liaison n'est plus une résistance pure; en effet, le réglage du circuit de sortie est tel qu'il existe nécessairement une forte composante capacitive, empêchant

une adaptation exacte à l'impédance caractéristique de la ligne. Ce mode de fonctionnement ne peut donc être employé que si cette dernière est très courte. Il faudra tenir compte, lors d'un calcul éventuel, du déréglage qui réduit notablement la résistance d'entrée du circuit final ($\bar{u}$ de ce circuit voisin du maximum). L'ordre de succession des valeurs X_q , X_q' et X_q'' montre que là encore les circuits doivent être désaccordés dans le même sens.

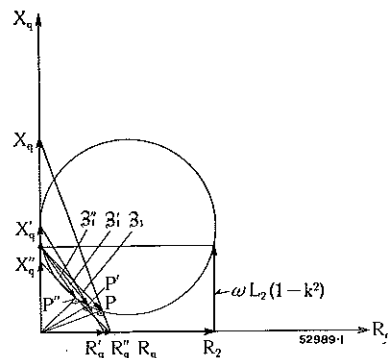


Fig. 18. — Diagramme vectoriel d'ensemble pour une ligne courte, lorsque $\omega L_2(1-k^2) = R_2/2$. Accord pour R_{pmin} .

Si l'on veut, pour une ligne de longueur non négligeable, pouvoir faire le réglage sans pont auxiliaire à l'aide des seuls instruments de mesure normalement prévus sur l'émetteur, par exemple simplement en cherchant le maximum de la charge, seul le mode de fonctionnement avec $\omega L_2(1-k^2) > R_2$ peut donner des résultats acceptables.

Les conditions sont notablement plus favorables lorsque le circuit d'entrée comporte une capacité série. Le diagramme vectoriel prend alors la forme de la figure 19. La grandeur de $\omega L_2(1-k^2)$ n'intervient plus, car elle peut toujours être compensée par la capacité série C_2 . L'accord du circuit de sortie peut être varié dans de larges limites sans que la résistance de charge R_q présente de changement appréciable. Une modification ne se produit que dans la mesure où le centre du cercle se trouve décalé par rapport à $R_3 = \beta_1'$. Si ce centre venait sur R_2 , on aurait constamment, quelque soit le désaccord, $R_q = R_2$ (en supposant toujours la ligne électriquement courte). L'ordre de succession de X_q , X_q' , X_q'' est tel que les deux circuits se réaccordent en sens inverse. L'apparition d'une charge réactive sur la ligne de liaison mise à part, un désaccord du circuit de sortie compensé

par un réaccord du circuit d'entrée, ne signifie rien de plus qu'un échange de capacité entre ces deux circuits.

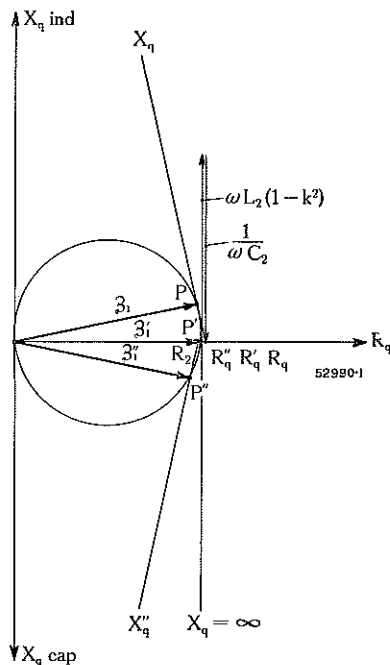


Fig. 19. — Diagramme vectoriel d'ensemble pour une ligne courte avec une capacité série C_2 du côté entrée.

V⁰ RÉSUMÉ.

Le montage constitué par deux circuits oscillants que relie une ligne de transmission, est fréquemment utilisé dans les émetteurs, par exemple pour la liaison à l'antenne, ou entre deux étages d'amplification. On trouve dans la littérature technique de nombreux travaux à ce sujet. Ceux-ci établissent les formules reliant les valeurs des éléments au fonctionnement électrique du système. De tels travaux ne peuvent satisfaire l'ingénieur d'étude ou le technicien d'exploitation que s'ils sont suivis d'une analyse complète du fonctionnement, faite en tenant compte des besoins de la pratique. Celle-ci exige en effet la réalisation de certaines fonctions de réglage, et en même temps demande que les variations correspondantes n'apportent qu'une perturbation négligeable à d'autres relations ne devant pas être modifiées. Les conditions auxquelles on doit alors satisfaire conduisent à choisir et proportionner les éléments dont on dispose selon des principes très différents. Ces principes ne peuvent se déduire des formules mentionnées que lorsque les relations entrant en ligne de compte sont entièrement éclaircies.

Le but de la présente étude est de fournir une solution d'ensemble aux divers problèmes courants que pose dans ce sens le fonctionnement d'un émetteur. Ces considérations sont basées sur un schéma comportant deux circuits oscillants couplés inductivement aux deux extrémités d'une ligne de transmission (link coupling).

Du côté entrée, une analyse approfondie a été faite des composantes résistives et réactives de la charge du primaire (amortissement et désaccord), en fonction soit d'une résistance de charge variable, soit de la combinaison d'une telle résistance et d'un condensateur variable série ou parallèle, et en tenant compte des valeurs des éléments de couplage. Il ressort de cette étude que la charge étant constituée par une résistance pure, la résistance effective de charge rapportée au primaire ne peut être rendue inférieure à un certain minimum, qui pour un circuit oscillant donné, ne dépend que du coefficient de couplage. On obtient ainsi une relation très simple entre le coefficient de couplage et le facteur de surtension du circuit oscillant chargé. Cette relation exige, pour la réalisation d'un facteur de surtension donné, un certain couplage minimum. Si pour une raison constructive ou autre, ce couplage ne peut être obtenu, on a toujours la ressource de connecter un condensateur en série ou en parallèle. Le montage en série permet en même temps d'annuler le désaccord créé par le couplage. La charge du circuit peut être réglée par variation du couplage jusqu'au fonctionnement à vide. On peut obtenir par un réglage exact du condensateur série, une indépendance pratiquement complète entre l'accord du circuit et le réglage de la charge. En y renonçant, on peut, par dérèglage simultané du condensateur d'accord et du condensateur série (ou parallèle), obtenir pour un couplage fixe une certaine gamme de variations de la charge; cette gamme est d'autant plus grande que le coefficient de couplage est petit pour le montage série et grand pour le montage parallèle. Le montage série et le montage parallèle ont fait l'objet de comparaisons approfondies pour les différents modes de fonctionnement entrant pratiquement en considération; alors qu'ils sont sensiblement équivalents en ce qui concerne leur réalisation constructive, le montage série présente de nombreux avantages techniques.

Du côté sortie, on a déterminé, en fonction de l'accord du circuit secondaire, la variation d'impédance aux bornes ainsi que le rapport des tensions

et le rapport des résistances entre ce circuit et la ligne. L'impédance varie et prend deux fois une valeur purement ohmique, l'une étant un maximum et l'autre un minimum. C'est le maximum qui entre normalement en considération comme point de fonctionnement. Le rapport de transformation croît lorsqu'on passe du maximum au minimum de résistance, atteignant son maximum dans ce dernier cas. Lors d'un réglage basé sur une variation d'amplitude, on obtient de ce fait sur la ligne une charge différent d'autant plus d'une résistance pure que les valeurs extrêmes sont voisines. La distance séparant ces valeurs est exprimée sous forme d'une différence de la capacité d'accord qui est une fonction simple du coefficient de couplage. En examinant cette fonction, on voit que, pour éviter le désaccord, il faut utiliser le plus fort couplage possible.

Les phénomènes existant lors du fonctionnement de l'ensemble des deux circuits deviennent extraordinairement compliqués et difficiles à comprendre si la longueur de la ligne n'est pas négligeable par rapport à la longueur d'onde. Même pour une ligne courte, il faut distinguer

plusieurs modes de fonctionnement suivant que la réactance de fuite du circuit d'entrée ramenée au côté ligne est plus grande ou plus petite que l'impédance caractéristique de la ligne. Dans le premier cas, avec l'accord du système basé sur une variation d'amplitude, on obtient une charge capacitive importante de la ligne. Un désaccord du côté sortie peut, si l'on renonce à la possibilité de charger le circuit d'entrée (augmentation de la résistance de charge), être compensé par un désaccord de ce circuit dans le même sens. On peut, par le même procédé, obtenir une charge purement résistive de la ligne, si l'on compense la réactance de fuite du côté de l'entrée par un condensateur en série. On peut, dans ce cas, sans modifier les conditions de charge, compenser entièrement un désaccord du circuit de sortie par un désaccord en sens inverse du circuit d'entrée. La possibilité d'un échange de capacité d'un circuit à l'autre donne beaucoup de liberté à l'accord, qui est de ce fait peu critique.

(MS 820)

Dr Max Dick. (M. F.)

NOUVEAUX RÉSULTATS DES RECHERCHES EN PHYSIQUE ATOMIQUE.

Indice décimal 539.152.1 — 539.17

Parmi les travaux scientifiques de ces dernières années, la première place revient aux recherches sur la transmutation artificielle des atomes. Cette nouvelle branche de la physique prend un rapide développement. Les physiciens peuvent maintenant, dans des expériences sûres, non seulement transmuter entre eux des éléments chimiques connus, mais produire artificiellement un grand nombre de nouvelles sortes d'atomes, encore inconnues. C'est parmi ces nouveaux éléments que l'on trouve les 300 sortes d'atomes radioactifs du plus haut intérêt. Les noyaux d'atome animés de très grandes vitesses, nécessaires à la transmutation des atomes, sont fournis de la façon la plus élégante par le cyclotron. La radioactivité est d'extrême importance pour la chimie et surtout pour les recherches biologiques, et l'énergie libérée lors de la transmutation des atomes, mise en valeur techniquement, réserve des possibilités insoupçonnées.

1° ÉLÉMENTS CONSTITUTIFS DE L'ATOME ET FORCES ATOMIQUES.

Les physiciens sont actuellement d'avis que tous les phénomènes qui se déroulent dans la matière pourraient être ramenés à l'interaction produite par les champs de quelques simples éléments constitutifs de la matière. Nous distinguons parmi ces particules élémentaires que l'on peut considérer d'après la mécanique ondulatoire comme ayant aussi bien une nature ondulatoire que corpusculaire, des particules lourdes et légères.

Les *électrons* sont les principaux représentants du groupe des *particules élémentaires légères*. Ils peuvent aussi bien être chargés négativement, négatons $-e$, que positivement, positons $+e$. La masse d'un électron, exprimée en unités de poids atomique est de 0,000543. (La masse d'un atome d'oxygène est admise

arbitrairement égale à 16,000.) Un électron est donc environ 2000 fois plus léger qu'un atome d'hydrogène. A côté de ces électrons légers, il existe des électrons lourds appelés *mesotons*. Ils portent la même charge que les électrons légers, mais sont environ 160 fois plus lourds. Nous engloberons aussi dans les particules légères les grains de lumière ou photons. Ce sont les particules élémentaires qui permettent le transport d'énergie par tous les genres de rayonnement

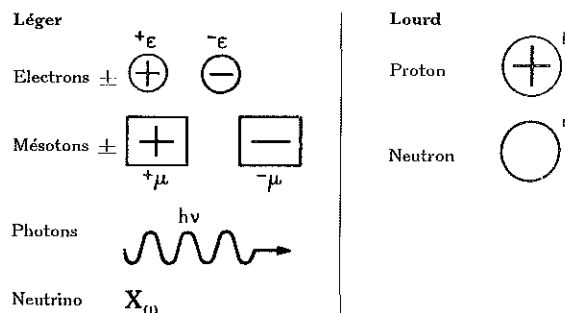


Fig. 1. — Particules élémentaires.

électromagnétique, rayons γ , rayons X, lumière ultraviolette, visible et infra-rouge. Comme dernière particule légère, mentionnons encore l'hypothétique neutrino dont l'existence n'a pu être décelée qu'indirectement par l'expérience et qui n'a pas encore été rendu visible.

Les photons et les neutrini ne sont pas chargés et ont, au repos une masse nulle, c'est-à-dire qu'ils n'existent qu'en mouvement et leur masse est simplement déterminée par leur énergie cinétique et leur impulsion.

Nous connaissons deux particules lourdes, le proton p et le neutron n. Le proton a la même charge positive que l'électron positif et le neutron n'est pas chargé. Les deux particules ont une masse environ égale à 1, exprimée en poids atomique, c'est-à-dire qu'elles ont à peu près le poids d'un atome d'hydrogène.

Toutes ces particules ont un moment cinétique de rotation dont la grandeur est une caractéristique de chaque particule comme sa masse ou sa charge. Ce moment de rotation est un multiple entier d'une unité fondamentale naturelle $\left(\frac{h}{4\pi} = \frac{\text{constante de Planck}}{4\pi}\right)$.

On ne doit pas se représenter ces éléments constitutifs de la matière comme des particules indestructibles et invariables, dont le nombre reste constant dans la nature. Il existe de multiples processus de transformation dans lesquels les particules élémentaires passent d'une catégorie à l'autre. Lors de la «matérialisation de la lumière», un photon γ non chargé,

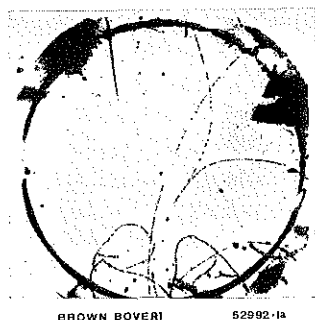


Fig. 2. — Photographie de la matérialisation de la lumière. Dans la chambre de Wilson remplie de gaz une paire d'électrons prend naissance. L'électron positif et l'électron négatif ont, à partir du point de génération, des trajectoires courbées en sens inverse dans le champ magnétique.

qui disparaît dans ce processus, se transforme en une paire d'électrons $+\epsilon$ et $-\epsilon$. Au cours de ces transformations les principes de conservation de la charge électrique et de conservation de «masse + énergie» sont vérifiés, c'est-à-dire qu'il naît toujours simultanément un électron $+$ et un électron $-$ et que l'énergie de la particule de lumière se retrouve complètement dans la masse et dans l'énergie cinétique de la paire d'électrons créée. Nous avons ici une belle preuve expérimentale de l'équivalence de l'énergie et la masse: masse et énergie ne sont que des noms différents pour une même grandeur. C'est pourquoi il existe un facteur déterminé de transformation entre la masse et l'énergie:

$$\text{Masse} = \frac{\text{énergie}}{(\text{vitesse de la lumière})^2}$$

Le phénomène inverse est «l'annihilation de la matière» dans lequel un électron positif et un électron

négatif, se fondent en particules de lumière généralement en deux, et la charge des électrons disparaît sans laisser de trace. Ce phénomène explique pourquoi les électrons positifs n'existent que passagèrement: ils rencontrent toujours dans la matière des électrons négatifs et se transforment en radiation. De la même façon, un mésoton peut se transformer spontanément en un électron ordinaire et un neutrino par une sorte de phénomène radioactif. Ce phénomène est fréquemment observé dans les rayons cosmiques où les mésotons apparaissent à l'état libre.

Il existe aussi un phénomène semblable avec les particules lourdes, le proton et le neutron; c'est ce qu'on appelle le processus β engendrant le rayonnement β des éléments radioactifs. Le neutron peut se transformer spontanément en un proton en émettant un électron négatif et un neutrino. Comme la masse du neutron est de 1,00895, elle est plus grande que les masses du proton et de l'électron créés (1,00759 + 0,00054). Il reste donc encore un excédent de masse qui est transformé en énergie cinétique pour l'électron et le neutrino. En revanche, le proton ne peut se transformer en un neutron et un électron positif que par apport d'énergie.

L'atome, est, comme nous le savons avec certitude depuis les recherches de Rutherford, «un atome nucléaire». Il se compose d'un noyau, petit et lourd, à charge positive, entouré d'une vaste enveloppe ne contenant que des électrons légers à charge négative. Le noyau ne renferme que des particules lourdes, protons et neutrons, et pas d'électrons. Le poids atomique des neutrons et des protons est presque égal à 1, ce qui explique pourquoi le poids atomique de n'importe quelle espèce d'atome est toujours très voisin d'un nombre entier. Le nombre Z de protons, qui fixe la charge du noyau, est égal au nombre atomique de l'élément dans le système périodique. Cette charge positive du noyau est donc déterminante du *comportement chimique de l'atome*, parce qu'elle détermine le nombre et la disposition des électrons négatifs dans l'enveloppe de l'atome et parce que les phénomènes chimiques se produisent uniquement dans l'enveloppe électronique. Le nombre des neutrons dans le noyau est à peu près égal à celui des protons, pour les éléments lourds un peu plus grand. Pour pouvoir lire directement le nombre de protons et de neutrons d'un noyau, nous inscrivons le nombre de protons à gauche, en dessous du symbole chimique, et le poids atomique brut (nombre de protons + nombre de neutrons) à droite, en haut. Par exemple, ${}_8^{16}\text{O}$ est un noyau d'oxygène contenant 8 protons et 8 neutrons et dont le poids atomique est 16.

D'après la loi de Coulomb les protons positifs du noyau doivent se repousser très violemment, étant donné les faibles distances qui les séparent, de l'ordre

de $3 \cdot 10^{-13}$ cm. Cette force de répulsion atteint pour un poids atomique moyen environ 20 kg par proton, c'est une force immense pour un si petit corps. Sous l'action de ces forces répulsives, le noyau de l'atome serait immédiatement dispersé, s'il n'existait pas des forces d'attraction entre les diverses particules qui le composent. Nous savons actuellement qu'il y a une attraction extrêmement forte, mais n'agissant qu'à faible distance, entre les éléments constitutifs du noyau. Cette attraction n'a rien de commun avec l'interaction de Coulomb. La théorie de ces forces nucléaires est actuellement étudiée de façon intensive. Il est très probable que ces forces ont le caractère de force d'échange. De telles interactions, agissant à courte distance, et qui ne peuvent pas être expliquées par la théorie corpusculaire classique, découlent automatiquement d'une application correcte de la mécanique des quanta. L'attraction entre protons et neutrons, par exemple, provient de ce que la charge d'un proton passe très rapidement sur un neutron, le proton devient neutron, et vice-versa. D'une façon semblable, cependant plus compliquée, les protons attirent les protons, et les neutrons, les neutrons. La répulsion électrique de Coulomb des protons entre eux se superpose simplement à cette attraction. Ces forces d'échange agissent aussi dans la couche électronique où elles produisent ces fortes attractions entre atomes conduisant aux combinaisons chimiques: la forte at-

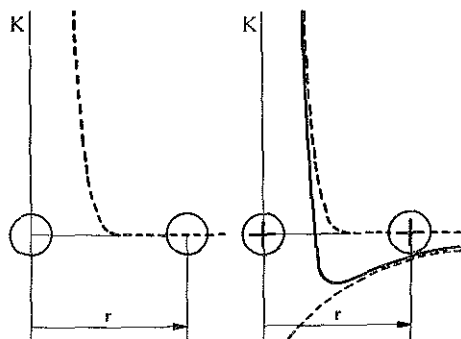


Fig. 3. — Forces nucléaires: Avec deux neutrons, seule existe l'attraction agissant brusquement à une distance r . Lors de l'interaction de protons, une répulsion électrique décroissant plus lentement que $\frac{1}{r^2}$ se superpose à cette force d'attraction.

traction des deux atomes dans une molécule d'hydrogène est expliquée quantitativement par la rapide permutation des deux électrons de l'enveloppe de la molécule.

Lors de la formation d'un noyau atomique en partant de protons et de neutrons, une grande quantité d'énergie est libérée par suite de la forte attraction entre les éléments du noyau. Réciproquement cette énergie devra être fournie pour décomposer un noyau atomique en ses divers éléments.

Les physiciens prennent pour modèle nucléaire simplifié une *gouttelette liquide* chargée d'électricité. De même que dans un liquide ordinaire, les molécules sont maintenues réunies par leur attraction réciproque (Forces de Van der Waals), les protons et les

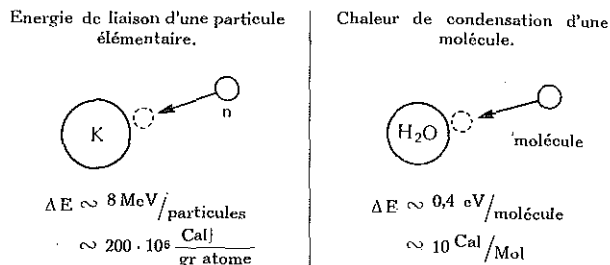


Fig. 4. — Modèle d'un noyau.

neutrons sont tenus ensemble par leur attraction réciproque. Lors de l'adjonction d'une molécule d'eau à une goutte d'eau, il se dégage de l'énergie, sous forme de chaleur de condensation. L'inverse se produit lors de la vaporisation d'une des molécules de la goutte d'eau, il faut fournir cette chaleur de vaporisation. Cette énergie est de 0,4 eV¹⁾, par molécule d'eau. De même, lors de l'adjonction d'un neutron ou d'un proton au noyau d'un atome, une grande énergie est libérée comme «chaleur de condensation» par suite de ces forces attractives à court rayon d'action. Cette énergie est beaucoup plus considérable que celle dégagée par la molécule d'eau: environ 8 millions d'eV par particule. Cette énergie est souvent émise par le noyau sous forme de rayons γ , mais sou-

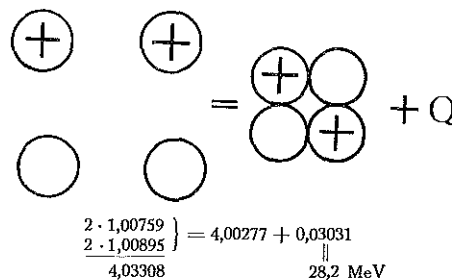


Fig. 5. — Bilan énergétique lors de la construction d'un noyau d'hélium avec un proton et un neutron: le noyau d'hélium se compose de deux protons et de deux neutrons.

La masse des quatre éléments est de 4,03308, mais le noyau d'hélium ne pèse que 4,00277 unités atomiques. Lors de la formation du noyau 0,0303 unité atomique a donc disparu sous forme d'énergie; cette valeur correspond à une énergie de liaison de 28,2 MeV par noyau d'hélium.

vent aussi le noyau subit, après l'adjonction d'un proton ou d'un neutron une transformation radioactive au cours de laquelle l'énergie est libérée.

¹⁾ Dans la physique atomique, on calcule toujours en prenant l'électron-volt eV, comme unité d'énergie. 1 eV est l'énergie qu'un électron reçoit en franchissant une différence de potentiel de 1 V: $1 \text{ eV} = 1,16 \cdot 10^{-19}$ Coulomb. $1 \text{ V} = 1,6 \cdot 10^{-19}$ Joule $= 1,6 \cdot 10^{-12}$ erg. ... $10^6 \text{ eV} = \text{MeV}$.

L'énergie libérée lors de la formation d'un noyau est si grande qu'elle peut même atteindre des valeurs décelables à la balance; elle apparaît directement comme «défaut de masse»: *Le noyau pèse moins que la somme de ses éléments constitutifs*. Une partie de la masse s'est donc perdue sous forme d'énergie, lors de la constitution du noyau. Ainsi, par exemple, la masse des quatre éléments constituant un noyau d'hélium, deux protons et deux neutrons, vaut 4,03308 unités de masse tandis que le noyau d'hélium ne pèse que 4,00277. La différence de masse, environ $30,3 \frac{\text{mg}}{\text{Mo } 1 \text{ He}}$, est donc émise lors de la formation du noyau sous forme d'énergie. La grandeur de cette énergie libérée, d'après le principe de l'équivalence de l'énergie et de la masse, est d'environ 28 millions d'eV par noyau d'hélium ou 600 000 000 grandes calories pour 4 gr He. Si l'on calcule, d'après le modèle de la goutte les énergies de liaison de divers atomes, en ne tenant compte que de forces attractives agissant à courte distance et de la répulsion de Coulomb, on obtient une très bonne concordance entre les défauts de masse calculés et mesurés.

On peut maintenant imaginer des noyaux constitués par la réunion d'un nombre quelconque de protons et de neutrons et se demander si ces noyaux existent dans la nature. En fait, dans la nature, il n'y a entre le plus léger élément, l'hydrogène, et le plus lourd, l'uranium, que 285 sortes d'atomes stables. La plupart des noyaux que nous construisons arbitrairement sont instables: Ils se transforment d'eux-mêmes en noyaux stables de moindre énergie. Le principe d'équilibre est le même qu'en statique: *L'état d'équilibre du noyau est l'état d'énergie minimum*.

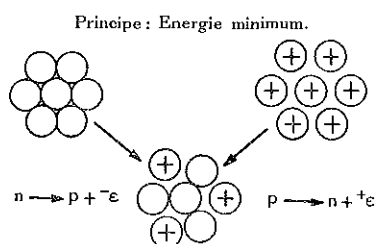


Fig. 6. — Noyaux stables.

Un noyau formé de sept neutrons est instable; il peut, par exemple, atteindre un niveau d'énergie inférieur en transformant quelques-uns de ces neutrons en protons. De même, un noyau composé uniquement de protons est instable. Il se transformerait en un noyau stable (${}^3\text{Li}^7$) si quelques protons se transformaient en neutrons.

Si nous imaginons un noyau formé de sept neutrons, cet assemblage de corpuscules est fortement lié, car les neutrons s'attirent très violemment. Mais il existe un état énergétique inférieur dans lequel ce noyau va se transformer: quelques neutrons se transforment en protons, des électrons négatifs seront ainsi émis et de l'énergie libérée. Notre noyau formé uniquement de neutrons se transformera en un noyau stable de lithium

${}^3\text{Li}^7$ constitué par trois protons et quatre neutrons. Comme avec l'accroissement de la charge du noyau, l'énergie de Coulomb augmente, et qu'ainsi l'énergie du noyau croît de nouveau, tous les neutrons ne se transforment pas en protons, mais seulement trois d'entre eux. D'autre part, un noyau que nous n'aurions construit qu'avec des protons ne serait pas stable non plus. Il pourrait se transformer également en un noyau ${}^3\text{Li}^7$ en fournissant de l'énergie, si quatre protons se transformaient en neutrons en émettant des électrons positifs. La transformation proton neutron nécessite, il est vrai, un apport d'énergie; mais celle-ci peut être tirée de la réserve d'énergie Coulombienne qui diminue lorsque la charge décroît. La cohésion du noyau augmente aussi légèrement par la réduction de la répulsion de Coulomb avec la libération d'énergie.

Outre l'énergie potentielle due aux forces en jeu, les éléments constitutifs du noyau possèdent aussi une *grande énergie cinétique*. De même que les électrons dans l'enveloppe, qui se trouvent animés d'un mouvement rapide déterminé par les lois de la mécanique quantique, les protons et les neutrons du noyau ont aussi une très grande énergie cinétique. On symbolise souvent cette énergie par ce qu'on appelle «la température apparente» du mouvement. Dans une goutte d'eau à la température normale, l'énergie de translation d'une molécule serait d'environ 10^{-14} erg, correspondant à une température absolue de 300° .

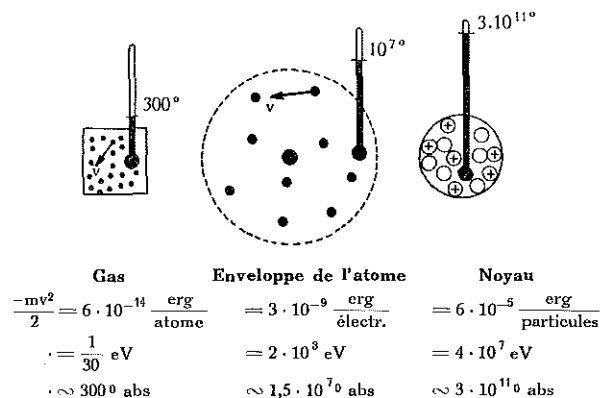


Fig. 7. — Energies atomiques.

Les énergies cinétiques des particules élémentaires peuvent être caractérisées par une «température apparente». Les électrons se déplacent dans l'enveloppe de l'atome avec une énergie aussi grande que s'ils étaient dans un gaz parfait de 10^7°C . Les éléments constitutifs du noyau sont animés d'un mouvement plus rapide encore et leur énergie cinétique correspond à celle d'un gaz à 10^{11}°C .

L'énergie moyenne d'un électron dans l'enveloppe est d'environ 10^{-9} erg, ce qui correspond à une «température apparente» de 10^7 abs. L'énergie de translation d'une particule élémentaire dans le noyau est d'environ 10^{-5} erg, ce qui donne une température apparente de 10^{11} abs. Cette énorme «température apparente» explique pourquoi nous ne pouvons pas influencer les noyaux avec les températures dont nous disposons.

II° TRANSMUTATION DES NOYAUX.

a) En principe, une transmutation nucléaire peut être provoquée par simple *apport d'énergie*. Exactement de même façon que les molécules d'une goutte d'eau peuvent être vaporisées par échauffement, les neutrons ou les protons d'un noyau peuvent être amenés à la «vaporisation» par simple apport d'énergie. L'énergie mise en jeu est naturellement beaucoup plus impor-

Transmutation d'un noyau par rayonnement γ .

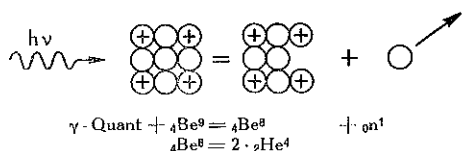
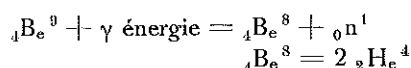


Fig. 8. — Effet de la lumière sur les noyaux.

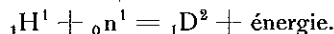
Un grain de lumière ayant suffisamment d'énergie peut provoquer la séparation d'un neutron d'un noyau de béryllium.

tante pour le noyau. Par exemple, un noyau d'atome de béryllium exposé aux rayons γ du radium vaporise un neutron selon la formule



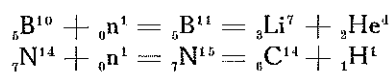
Le noyau de béryllium de poids atomique 9 se transforme en un isotope de poids atomique 8, et un neutron. Le béryllium 8 est instable et se partage en deux noyaux d'hélium. Cet effet des radiations sur les noyaux peut être observé sur beaucoup d'autres noyaux. Malheureusement, les rayons γ naturels sont moins chargés d'énergie (maximum 2,62 MeV); mais on peut actuellement, à l'aide des processus de désintégration atomique, produire des rayons γ d'énergie atteignant 17 MeV qui provoquent mieux l'effet en question.

b) Une autre réaction nucléaire est la simple adjonction d'un neutron à un noyau d'atome stable. On peut, par exemple, partant d'un atome d'hydrogène léger produire un atome d'hydrogène lourd ou deuton, de poids atomique 2, d'après la formule suivante:



Cette adjonction libère de l'énergie sous forme de photons (rayons γ).

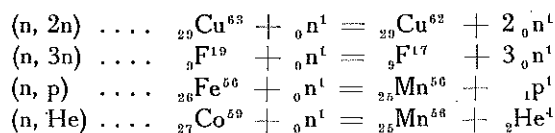
c) L'incorporation d'un neutron dans un noyau existant entraîne souvent une *réaction avec échange de particules*: L'adjonction d'un neutron lent donne naissance à un noyau intermédiaire non stable et qui se partage immédiatement en deux. Donnons des exemples de ce processus:



(Le neutron est échangé, dans le premier cas, contre un noyau d'hélium, dans le second, contre un proton.)

Les noyaux intermédiaires ne sont pas les noyaux de bore ou d'azote ordinaires, qui existent aussi à l'état stable mais des «noyaux excités», qui, lors de l'incorporation du neutron, ont reçu un excédent d'énergie. Un noyau excité peut être comparé à une goutte d'eau surchauffée. Le noyau excité expulse une particule exactement comme la goutte d'eau surchauffée se vaporise brusquement. Il peut aussi revenir à son état primitif en émettant des rayons γ , mais c'est plus rare.

d) Par bombardement avec des neutrons très rapides riches en énergie, il est évidemment possible, de dissocier n'importe quel noyau de diverses façons. L'énergie des neutrons pénétrant dans le noyau «échauffe» tellement celui-ci qu'il se «vaporise» en partie. C'est ainsi qu'il existe des réactions $(n, 2n)$ ou $(n, 3n)$ dans lesquelles un neutron pénétrant dans le noyau permet l'évaporation de deux ou trois neutrons. On connaît également de nombreuses réactions (n, p) , (n, d) et (n, He) , dans lesquelles le neutron pénétrant dans le noyau chasse un proton (p), un deuton (d) ou un atome d'hélium (He). Selon les noyaux, les conditions énergétiques sont très différentes:



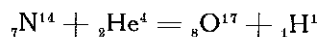
Ces transmutations d'atomes par des neutrons lents ou rapides sont, en principe, les plus simples. Le neutron qui s'approche d'un noyau ne subit aucune force, jusqu'à l'instant où il arrive dans le voisinage immédiat du noyau. Ce n'est qu'à partir de cet instant qu'il est saisi par la force du noyau et attiré par celui-ci.

Les physiciens ne disposent pas directement de neutrons pour la transmutation nucléaire, car ceux-ci n'existent pas à l'état libre dans la nature. Les neutrons libres n'ont dans la matière qu'une courte vie moyenne, car ils sont absorbés avec avidité par les noyaux des atomes.

e) Le physicien doit, pour effectuer la transmutation des atomes, partir de noyaux d'atome stables existant dans la nature. Il doit les amener en contact de façon qu'ils puissent réagir l'un sur l'autre et échanger des éléments. A cause de leur charge positive, les noyaux d'atome se repoussent violemment, et pour vaincre cette répulsion, il faut que les noyaux aient de grandes énergies cinétiques. Nous ne connaissons malheureusement pas encore d'autres moyens pour amener les noyaux en contact que le bombardement des atomes immobiles par d'autres atomes animés d'un mouvement rapide. L'utilisation de ce bombardement ultra-rapide a conduit à la notion de «rupture de l'atome», bien que de telles réactions nucléaires se produiraient

aussi si deux noyaux étaient amenés en contact par de grandes forces, mais sans vitesse. Il est donc plus correct de parler de «réaction nucléaire» que de rupture d'atomes.

La première transmutation d'atomes de ce genre a été réalisée par Rutherford: il utilisa les noyaux ${}^4_2\text{He}$, qui sont émis par le radium C' et qui ont une énergie de 7,83 MeV¹⁾, pour bombarder de l'azote, et il observa la réaction nucléaire suivante:



On réussit avec des rayons α à provoquer des réactions nucléaires dans de nombreux éléments légers.

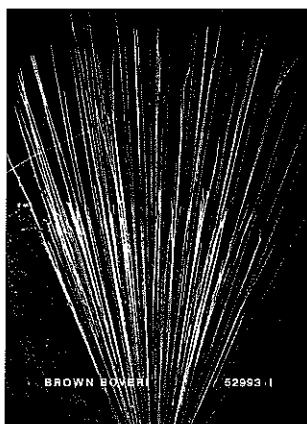


Fig. 9. — Photographie d'une transformation artificielle d'atomes dans une chambre de Wilson.

Une particule α frappe un noyau d'azote avec lequel elle réagit. Il se produit alors un noyau d'oxygène et un d'hydrogène dont les trajectoires sont visibles.

La chambre de Wilson rendit des services inappréciables dans les recherches sur ces phénomènes, car elle permet de rendre directement visibles les trajectoires des noyaux participant à la réaction, et de mesurer les énergies et les impulsions des particules. Dans cet appareil rempli de vapeur d'eau surchauffée, les trajectoires sont rendues visibles par de très fines gouttelettes d'eau qui se condensent le long du parcours de la particule.

Si l'on veut produire artificiellement les noyaux d'atomes à grandes vitesses, nécessaires aux réactions nucléaires, il faut avoir des appareils dans lesquels ces noyaux sont accélérés par de très hautes tensions. Ces dernières années, un grand nombre d'installations ont été construites dans lesquelles des tensions continues de plusieurs millions de volts ont été produites à l'aide de transformateurs et de redresseurs ou d'après le principe électrostatique au moyen de rubans marchant à grand vitesse. Dans les tubes à «rayons ca-

naux» on produit des rayons de particules dont chaque élément a une énergie de plusieurs millions d'eV. La limite actuelle de ces installations est de 4 millions de volts environ.

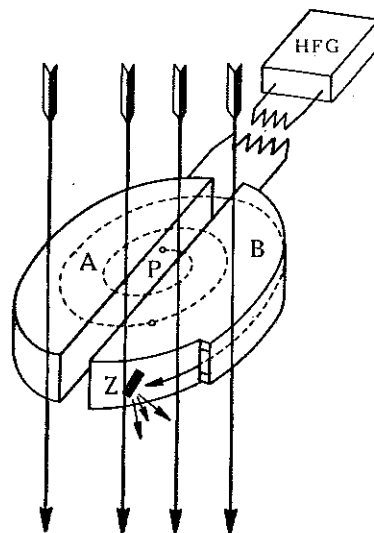


Fig. 10. — Principe du cyclotron.

A et B sont les deux moitiés d'une boîte métallique, P la source de proton; la spirale pointillée représente la trajectoire du proton. Z est le corps visé que frappent les particules et dont les atomes sont rompus. HFG est le générateur haute fréquence.

Le *cyclotron* est le générateur à très grande vitesse, le plus élégant et le plus efficace; dans cet appareil, une tension relativement basse, mais alternative, à haute fréquence, accélère plusieurs fois le même atome. La particule se déplaçant de plus en plus rapidement doit toujours être ramenée en synchronisme dans le champ électrique alternatif; sa déviation est obtenue par un fort champ magnétique.

Entre les deux pôles d'un gros électro-aimant se trouve une chambre cylindrique d'environ 1 m de diamètre et 25 cm de haut dans laquelle on fait un vide très poussé. Cette chambre renferme les deux électrodes d'accélération qui ont la forme des deux moitiés d'une boîte cylindrique coupée par un plan axial. Ces deux électrodes, qui sont isolées l'une de l'autre et des parois de la chambre, sont soumises à une tension haute fréquence d'environ 50 000 V. Au milieu de la chambre se trouve une source de ions qui émet des noyaux d'atomes positifs, protons ou deutons, entre les deux électrodes. Ces noyaux sont saisis par le champ magnétique et attirés d'un côté à l'autre, et commencent à décrire dans la boîte une spirale. Un ion nouvellement émis est, par exemple, tout d'abord attiré par l'électrode de droite si à cet instant elle a un potentiel négatif, et celle de gauche un potentiel positif. A l'intérieur de la boîte où ne règne aucun champ électrique, le ion ne peut pas avoir une trajectoire rectiligne, il est contraint par la force de Lorentz de décrire un

¹⁾ Un de ces noyaux d'He a une vitesse d'environ 1/15 de celle de la lumière, c'est-à-dire 20 000 km/s, soit 72 000 000 km/h.

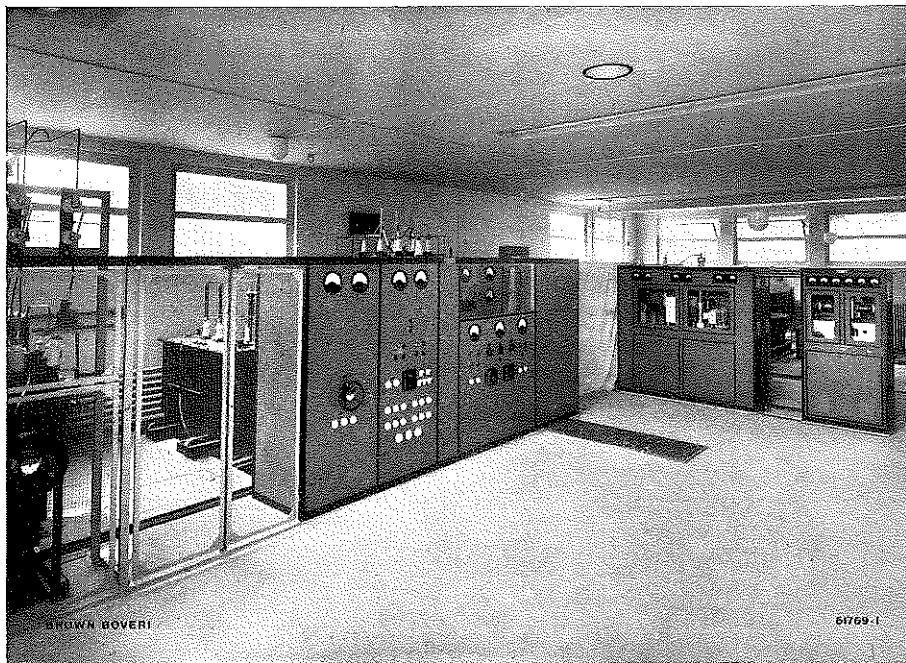


Fig. 11. — Vue d'ensemble de l'installation du cyclotron.

A gauche, installation à courant fort et mutateur de 15 kV pour l'alimentation de l'étage final à haute fréquence. Au milieu, armoire de l'étage final à haute fréquence pour une puissance maximum de 50 kW. A droite, armoire de l'étage pilote haute fréquence. Entre les armoires, on aperçoit une partie de l'aimant.

cerce. Après un certain temps il repasse de nouveau entre les électrodes. Mais, pendant ce même temps, la polarité des électrodes a changé, si bien que le ion est maintenant accéléré vers l'électrode de gauche. Le rayon de la trajectoire croît exactement proportionnellement à la vitesse, si bien que le ion revient toujours entre les deux électrodes après le même temps, celui après lequel la polarité avait changé la première fois (théorème de Larmor). Comme la durée du passage dans l'électrode, déterminée par l'intensité du champ magnétique, est extraordinairement courte, à savoir $1/30\,000\,000$ s, la polarité de la tension alternative doit changer $30\,000\,000$ de fois par seconde, c'est dire qu'il faut une tension à haute fréquence. Et pour obtenir à chaque passage d'une électrode à l'autre une forte accélération il faut une tension élevée. Comme on peut facilement le voir, à une vitesse de rotation $15\,000\,000$ t/s sur une circonférence de 50 cm de rayon correspond une vitesse de $30\,000$ km/s, soit le $1/10$ de la vitesse de la lumière.

L'Ecole Polytechnique Fédérale possède un cyclotron avec lequel les deutions peuvent atteindre une énergie de 14 MeV et les particules α , de 24 MeV.

L'électro-aimant construit par la Fabrique de Machines Oerlikon pèse 40 t. Il a des pôles de 90 cm de diamètre produisant une intensité de champ de $18\,000$ A/cm dans l'entrefer de 15 cm. La puissance aux bornes des bobines de l'électro-aimant est de 200 kW.

L'installation haute fréquence du cyclotron a été construite par Brown Boveri. Il s'agit d'un générateur à ondes courtes de 50 kW avec une longueur d'onde de 20 m. Dans un étage pilote sont produites des oscillations électriques dont la fréquence est maintenue constante par un dispositif adéquat. Dans les

deux étages suivants; cette oscillation excitatrice est amplifiée et double deux fois sa fréquence. On obtient ainsi pour la commande de l'étage final la puissance nécessaire d'environ 1 kW.

L'étage final contient deux lampes démontables Brown Boveri à grande puissance. Il s'agit de triodes refroidies à l'eau et reliées à une pompe moléculaire. Elles travaillent en push-pull, classe C.



Fig. 12. — Etage pilote (à droite) et étage final (à gauche) de l'installation haute fréquence. A l'intérieur de l'armoire de gauche on aperçoit une des lampes Brown Boveri démontables ainsi qu'une partie du circuit oscillant d'anode.

La chambre du cyclotron a été construite par l'Institut de Physique de l'Ecole Polytechnique Fédérale. Elle présente des modifications importantes par

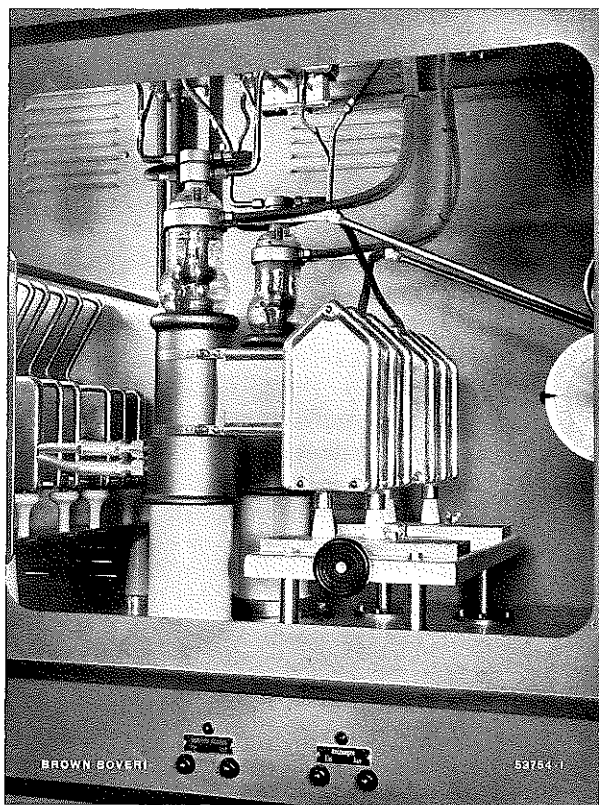


Fig. 13. — Lampes Brown Boveri démontables, refroidies à l'eau avec, à droite, condensateurs de neutralisation et, à gauche, le condensateur du circuit oscillant d'anode.

Les lampes sont continuellement reliées, par le cylindre isolant, à la pompe à vide élevée se trouvant sous la tôle de protection.

rapport aux constructions connues. La partie haute fréquence de la chambre avec les électrodes DD a reçu la forme du système de Lecher raccourci, de longueur effective $\frac{\lambda}{2}$. Cette disposition entraîne toute une série d'avantages:

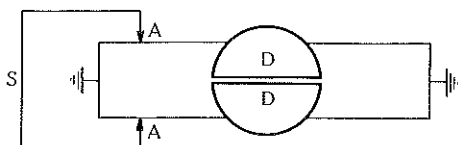


Fig. 14. — Système accélérant les atomes.

Les deux moitiés de la boîte dans lesquelles les noyaux d'atomes sont accélérés sont montées au milieu d'un système de Lecher qui oscille à la demi-longueur d'onde. La tension alternative à haute fréquence est amenée en AA.

1^o Répartition symétrique de la tension dans l'espace entre les électrodes.

2^o Possibilité d'une fixation stable des électrodes en supprimant les isolateurs (les deux extrémités du système se trouvant au potentiel de la terre).

3^o Possibilité d'une adaptation exacte du système à la ligne d'apport d'énergie.

4^o La haute fréquence est amenée au système en un point où la tension est 25 fois plus faible qu'aux électrodes déviateurs. Le système de Lecher agit donc comme un transformateur pour la tension déviatrice. Les ions sont produits dans une petite chambre à gaz par un arc à basse tension. Le vide dans la chambre d'accélération est fait par une pompe à vapeur d'huile, construite à l'Institut et ayant un débit de 1000 l/s.

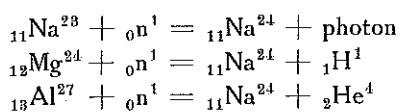
Le plus grand cyclotron construit (à Berkeley, en Californie) a un aimant de 400 t et des pôles de 2,5 m de diamètre. Avec cet appareil on produit des deutons ayant une énergie de 16 MeV et des noyaux He de 32 MeV. Le cyclotron sert surtout comme source très puissante de neutrons et à la production de matières radioactives artificielles. Le plus souvent, on utilise pour la production de neutrons la réaction nucléaire ${}_3\text{Li}^7 + {}_1\text{D}^2 = 2{}_2\text{He}^4 + {}_0\text{n}^1$ qui livre une grande quantité de neutrons. Il existe un projet d'un cyclotron plus grand encore dont l'aimant pèserait 4000 t. Il servirait à des essais biologiques et médicaux; il sera terminé à Berkeley dans trois à quatre ans. Jusqu'à présent, ces appareils ont servi à effectuer de nombreuses réactions nucléaires, à mesurer des bilans d'énergie et à produire une très grande série de nouvelles sortes d'atomes.

III^o MATIÈRES RADIOACTIVES ARTIFICIELLES.

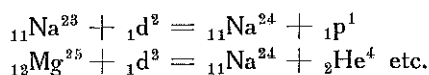
Parmi ces atomes produits artificiellement les plus intéressants sont les *atomes artificiels radioactifs* dont on connaît actuellement plus de 300. Beaucoup d'atomes créés par une réaction nucléaire ne sont pas en équilibre stable; ils se transforment en éléments stables par la transmutation de neutrons en protons ou vice-versa, en émettant des électrons négatifs ou positifs et souvent aussi des rayons γ . Il est très intéressant pour le physicien de pouvoir produire aussi des éléments radioactifs naturels comme le radium E (rayons β) et le polonium (rayon α) en partant de bismuth.

Citons comme exemple de formation d'un élément radioactif artificiel le sodium radioactif, qui a un poids atomique de 24, alors que le sodium stable connu a un poids atomique de 23. On peut produire ce Na^{24}

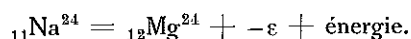
par de nombreux types de transformations, par exemple par adjonction ou échange de neutrons:



ou de deutons:



Le ${}_{11}\text{Na}^{24}$ radioactif se décompose en émettant des rayons β d'après le schéma suivant:

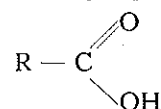


Les électrons émis ont une énergie maximum de 1,4 MeV, on observe en outre la présence de radiations γ . La période est de 14,8 h, c'est-à-dire qu'en environ 15 h la moitié des atomes d'une préparation se sont transformés, après 2 fois 15 heures il ne reste plus qu'un quart des atomes radioactifs, etc. Ces éléments radioactifs artificiels sont appelés à jouer un rôle important dans la chimie et surtout dans la biologie et la médecine. Aucune découverte aussi importante, pour l'étude biologique de l'assimilation et pour la chimie pharmaceutique, que la radioactivité artificielle, n'a été faite depuis celle du microscope.

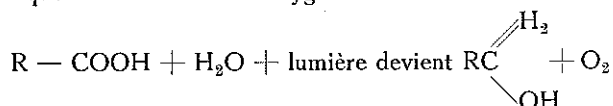
L'utilisation des atomes radioactifs en chimie et en biologie repose sur le fait que les éléments radioactifs ne se différencient *absolument pas* de leur isotope au point de vue chimique et physiologique, tant que l'action des radiations radioactives peut être négligée. Le sodium radioactif se comporte chimiquement tout à fait comme le sodium ordinaire. Cependant, les éléments radioactifs sont *caractérisés par leur radioactivité*, ils sont pour ainsi dire munis d'une étiquette qui permet à chaque instant de les différencier des atomes stables, chimiquement semblables. Il est com-

préhensible que par des mesures radioactives on puisse facilement déterminer la répartition dans les tissus de phosphore radioactif sous forme de NaH_2PO_4 . Le phosphore ordinaire déjà contenu dans le corps avant l'expérience n'est pas gênant lors de ces mesures, car le phosphore introduit se distingue tout de suite de l'autre par sa radioactivité.

Un des exemples les plus intéressants de l'emploi de ces atomes particuliers est celui du carbone radioactif dans la *recherche sur l'assimilation dans les plantes vertes*. A l'aide de la méthode radioactive qui est environ un million de fois plus sensible que la méthode chimique, on a pu montrer que la théorie classique de l'assimilation n'est pas exacte. Cette théorie classique admet que les plantes n'absorbent le CO_2 que si elles sont éclairées et le réduisent en formaldéhyde et oxygène: $\text{CO}_2 + \text{H}_2\text{O} + \text{lumière} = \text{CH}_2\text{O} + \text{O}_2$. Au moyen de dioxyde de carbone dans lequel le carbone est remplacé par du carbone radioactif, on peut facilement montrer que les plantes absorbent le bioxyde de carbone même *dans l'obscurité*; le carbone radioactif provenant du dioxyde se retrouve quantitativement dans un groupe carboxyle:



R est un radical d'un poids moléculaire de 1000 environ. Ce n'est que sous l'effet de la lumière et grâce à l'effet catalyseur de la chlorophylle que le groupe COOH de cet acide se réduit en un groupe alcoolique en libérant de l'oxygène:



A l'aide de mesures radioactives on peut prouver, par exemple, qu'en deux heures les 20 % du dioxyde de carbone absorbé par une plante d'orge sont transformés en sucre. Il est clair que ces recherches n'auraient jamais pu être faites sans le carbone radioactif. On aurait, en effet, été dans l'impossibilité de distinguer les atomes de carbone provenant du dioxyde de ceux se trouvant déjà dans la plante.

Les éléments radioactifs artificiels ont aussi pris une très grande importance pour les recherches sur l'assimilation chez l'homme. On peut maintenant, en faisant absorber des matières contenant des éléments radioactifs, suivre facilement la digestion, le transport et la transformation dans l'organisme de ces matières ou médicaments. Grâce à la grande sensibilité de détection des substances radioactives à l'aide de compteurs d'électrons (compteur Geiger), il suffit d'introduire des quantités très faibles de substances dans l'organisme. Il est même souvent possible de faire ces essais sur des corps intacts, car les rayons γ des éléments radioactifs,

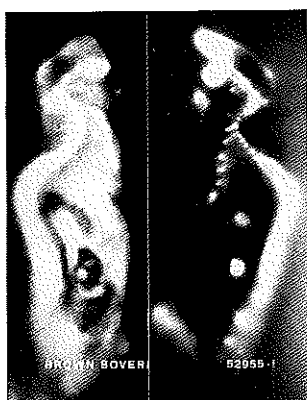


Fig. 15. — Photographie de dépôts de phosphore radioactif (à gauche) et de strontium radioactif (à droite), dans un rat.

Les substances radioactives ont été photographiées par leur propre radiation. On remarque que le phosphore introduit dans l'organisme s'est réparti plus uniformément dans les os et les tissus que le strontium qui se dépose plus fortement dans les os.

par suite de leur grand pouvoir de pénétration, sortent de l'organisme et leur présence peut être constatée par des compteurs d'électrons. Mentionnons comme exemple l'absorption et la transformation de l'iode par la glande thyroïde. Il est connu que la concentration d'iode dans tous les tissus du corps humain est particulièrement faible, sauf dans la glande thyroïde. Dans cet organe où la concentration est 10 000 fois plus grande qu'ailleurs, l'iode se transforme en thyroxine, combinaison organique de l'iode, qui règle la vitesse de l'oxydation dans le corps. Le manque de thyroxine réduit l'absorption d'oxygène par le corps et ralentit toute l'assimilation. Une production exagérée de thyroxine accroît fortement l'absorption d'oxygène et l'oxydation. Si l'on introduit quelques milligrammes de iodure de potassium, contenant de l'iode radioactif, dans le corps, on peut après quelques minutes constater sa présence dans la thyroïde à l'aide d'un compteur d'électrons placé dans la gorge. Des essais sur individus normaux montrent que la saturation d'iode dans la thyroïde est atteinte au bout de deux jours; on retrouve environ 4 % de l'iode dans la thyroïde, tandis que la plus grande partie du reste a été éliminée par le corps. Chez les personnes atteintes de la maladie de Basedow, dont la glande thyroïde travaille trop activement, l'absorption d'iode est considérablement augmentée. Déjà après quelques heures, leur glande thyroïde a absorbé les 12 à 15 % de l'iode radioactif, et tout aussi rapidement l'iode, transformé en thyroxine, est emporté par la circulation du sang, et il accélère les phénomènes d'assimilation. Chez les individus dont la thyroïde est plus faible que normale, l'absorption d'iode est très réduite.

Le dépôt des substances radioactives dans le corps est souvent très sélectif; par exemple, le strontium radioactif se dépose presque exclusivement dans les os. De cette propriété est né l'espoir de pouvoir appliquer ainsi la radiothérapie à certains organes, sans nuire aux autres; cependant, jusqu'à présent, il n'existe que peu de résultats d'essais. On a déjà fait des essais très intéressants sur l'absorption et l'assimilation de phosphore par les plantes, ainsi que sur l'assimilation par étapes du fer radioactif par l'hémoglobine du sang.

Naturellement, les matières radioactives artificielles, telles que le radio-sodium, dont on a pu préparer des compositions très actives, ont déjà une grande importance *en thérapeutique*.

En métallurgie, les métaux radioactifs sont utilisés dans les recherches sur les phénomènes de mélange, d'auto-diffusion et de précipitation.

Malheureusement, on n'est pas arrivé jusqu'à présent à transformer des quantités pesables d'un élément; il s'agit toujours de très petites quantités de substance que l'on ne peut déceler que chimiquement. Cela n'est pas considéré comme un défaut par le physicien,

car la nature a mis plus que suffisamment de matières premières et de matières de remplacement à la disposition du technicien. Cependant, il est très regrettable que l'énorme *énergie* produite lors des réactions nucléaires ne puisse pas être rendue techniquement utilisable, car rien ne remplace l'énergie. Cette énergie de réaction est particulièrement élevée lors de la décomposition d'un noyau d'uranium. L'uranium est le dernier élément dans le système périodique; le noyau d'uranium ${}_{92}\text{U}^{238}$ se compose de 92 protons et de 143 neutrons. On peut, en se basant sur le simple modèle de la goutte d'eau, voir facilement qu'un tel noyau est près de la limite de stabilité et que, en raison de la charge élevée du noyau, la répulsion de Coulomb, est presque supérieure aux forces attractives du noyau. En fait, il suffit d'introduire un neutron lent dans le noyau pour obtenir un *noyau instable d'uranium 235*. Il se dissocie après l'adjonction du neutron en deux noyaux plus petits. On peut se représenter cette dissociation comme un étranglement du noyau, devenu instable par l'apport d'énergie, suivi d'une séparation en deux parties sous l'effet des forces de répulsion. En outre, deux ou trois neutrons «se vaporisent» au cours du processus. Une fois les deux petits noyaux créés, la forte répulsion de Coulomb les sépare, en leur communiquant une énorme énergie cinétique: l'énergie libérée par la dissociation de l'atome d'uranium est d'environ 160 MeV. Si l'on pouvait dissocier un kilogramme d'uranium par ce procédé, on libérerait environ 16 milliards de grandes calories, ce qui correspond à la combustion de deux millions de kilogrammes de charbon.

Le physicien cherche, naturellement, avec assiduité un moyen de rendre cette énergie utilisable, mais il y a d'immenses difficultés à surmonter. D'abord, seul un des trois isotopes dont est constitué l'uranium stable se décompose de la façon indiquée. L'uranium se compose en effet de trois sortes d'atomes de poids atomiques 238, 235 et 234. L'uranium 235, qui se dissocie si facilement, ne forme que le 0,7 % du mélange. On essaie actuellement, par de nouveaux procédés de séparation des isotopes, d'enrichir fortement le mélange en uranium 235. L'uranium 235 une fois obtenu, il faut que la dissociation de l'atome soit une réaction successive, c'est-à-dire que les neutrons libérés par la dissociation d'un atome doivent entraîner tout de suite celle d'un autre atome, de façon que la réaction une fois commencée se poursuive dans toute la masse d'uranium. On sait actuellement que seuls les neutrons lents provoquent la dissociation de l'uranium. C'est pourquoi il est presque certain que la réaction se freinera d'elle-même à cause de la vitesse toujours plus grande des neutrons engendrés par les températures élevées qui sont produites. Cet effet est très désirable, car ainsi il ne se produit

pas d'explosion lors de la dissociation de l'uranium, mais l'énergie se libère progressivement comme dans la combustion du charbon.

Nous savons maintenant avec certitude que la transformation des atomes est la formidable source d'énergie à laquelle *la vie des étoiles* s'alimente. Elle fournit en effet plusieurs millions de fois plus d'énergie que les réactions chimiques. Dans le cas du soleil, qui perd par rayonnement de si énormes quantités d'énergie, qui, si elles étaient produites par la combustion du charbon et de l'oxygène dont le soleil serait constitué, ne pourraient être livrées que pendant 3000 ans, on connaît exactement le phénomène qui se produit. L'hydrogène est transformé en hélium, le carbone servant de catalyseur. Le phénomène se poursuit en cinq étapes et développe environ 150 000 000 de grandes calories par gramme d'hydrogène transformé. (La chaleur de combustion d'un gramme d'hydrogène n'est que de 48 grandes calories.) Malgré la température excessivement élevée de 20 000 000 ° qui doit régner au centre du soleil, cette réaction atomique se produit par bonheur lentement, car seuls quelques noyaux

d'atomes ont la vitesse suffisante pour vaincre la répulsion de Coulomb et déclencher la réaction nucléaire.

L'inflammation brusque de la *Supernovae*, dans laquelle des étoiles émettant un rayonnement 60 millions de fois plus intense que celui du soleil prennent un nouvel équilibre, peut aussi s'expliquer par des considérations de physique nucléaire.

Les recherches sur la transmutation artificielle des atomes progressent rapidement, surtout dans les instituts américains qui disposent de moyens énormes et de véritables états-majors de collaborateurs. Les connaissances de la physique ont été, en peu de temps, accrues et enrichies de façon inespérée par la solution de nombreux problèmes sur les noyaux. En particulier, notre connaissance des éléments constitutifs de la matière et de leurs interactions s'est beaucoup approfondie. Il reste une question passionnante, mais à laquelle on ne peut malheureusement pas donner une réponse certaine: celle de savoir si ce domaine pourra dans un avenir rapproché être rendu accessible à la technique.

(MS 818)

Prof. Dr P. Scherrer. (J. C.)

POSTES ÉMETTEURS DE RADIODIFFUSION.

Indice décimal 621.396.712

Après avoir mené à bien de nombreuses recherches dans le domaine de la haute fréquence, après avoir construit divers émetteurs de faible et moyenne puissance, et surtout les lampes démontables de grande puissance, nous avons entrepris la construction d'émetteurs complets de radiodiffusion.

Brown Boveri est parmi les rares constructeurs qui sont à même de livrer une installation complète, c'est-à-dire la partie courant fort, celles à basse et à haute fréquence, construite dans leurs propres ateliers. Bien avant de nous occuper de haute fréquence, nous avions livré des parties importantes d'émetteurs puissants. Depuis 1929, nous avons mis en service un grand nombre de mutateurs à haute tension atteignant 21 kV pour des puissances jusqu'à 1500 kW, destinés à des postes émetteurs de radiodiffusion, et sur lesquels nous avons introduit les premiers l'extinction des courts-cuits par polarisation des grilles. Nous avons aussi livré en Suisse et à l'étranger des filtres d'harmoniques, des groupes de chauffage des filaments et des transformateurs de modulation à grande puissance.

Pour débiter dans la construction de postes d'émission, nous avons monté dans nos laboratoires un émetteur d'essai d'une puissance porteuse pouvant atteindre 300 kW. Les résultats obtenus furent si favorables que la Direction Générale des PTT à Berne se déclara disposée à utiliser cet émetteur à titre d'essai dans le plus grand poste suisse de radiodiffusion. Au printemps 1941, on put commencer le montage sur place et le réglage de l'émetteur et au milieu de juillet déjà il servait aux émissions journalières. Jusqu'en décembre 1941, cet émetteur d'essai, équipé de nos lampes démontables, type DT 20/150 d'une puissance porteuse de 80 kW, a fonctionné plus de 1500 heures sur l'antenne et n'a donné lieu au total qu'à moins d'une heure de pannes d'émission, ce qui est très peu.

La modulation se faisant provisoirement dans l'anode d'un étage intermédiaire, le facteur maximum de distorsion est de 3,9 à 4,3 % pour une profondeur de modulation de 80 % et un niveau de bruit de 64 db. En outre, la courbe de réponse de l'émetteur est excellente, ainsi que le montre le tableau suivant:

Caractéristique amplitude-fréquence (rapportée à 1000 cycles).

f en c	35	50	100	200	1000	2000	4000	6000	8000	10 000
Ecart en db	+ 1,1	+ 0,4	+ 0,1	0	0	+ 0,1	+ 0,3	+ 6,6	+ 1,2	- 1,3

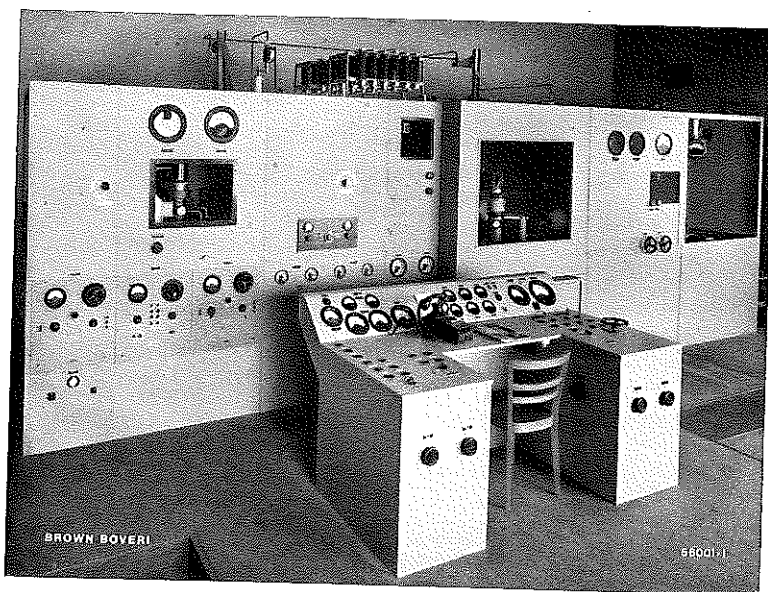


Fig. 1. — Emetteur d'essai Brown Boveri, d'une puissance porteuse de 80 kW.
Au milieu, armoire renfermant les lampes démontables type DT 20/150.

Les conditions de stabilité en haute fréquence imposées par les prescriptions du CCIR sont respectées et les harmoniques supérieures restent dans des limites admissibles.

Les résultats obtenus avec cet émetteur d'essai permettent d'élever la puissance porteuse de 80 à 100 kW avec les mêmes lampes démontables, en améliorant simultanément la qualité de transmission, ce qui est possible grâce au modulateur puissant que nous venons de réaliser. Si les conventions internationales pour la radiodiffusion européenne le permettaient, la puissance de ce poste pourrait être portée à 150 kW.

Les bons résultats mentionnés ont aussi permis le remplacement, à partir de janvier 1942, de la triode bien connue CAT 14 de l'émetteur normal de ce grand poste, par nos lampes démontables du même type que celle de l'émetteur d'essai.

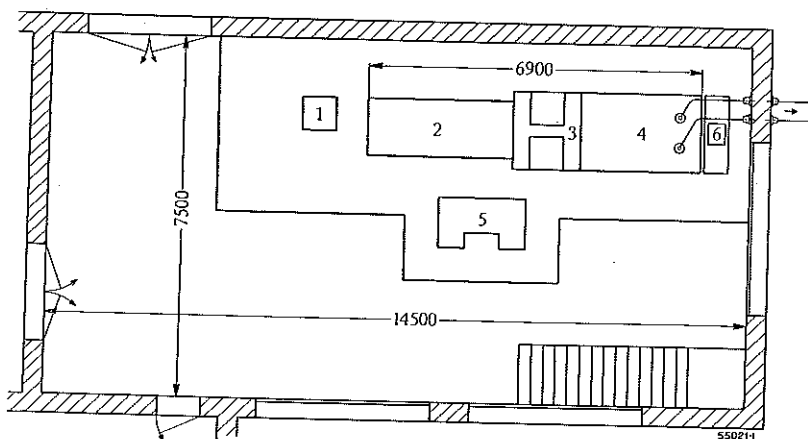


Fig. 2. — Plan de disposition de l'émetteur d'essais.

- | | | |
|---------------------------|-------------------------|----------------------------|
| 1 = Modulateur. | 4 = Etage de puissance. | 6 = Antenne fictive. |
| 2 = Etages préliminaires. | 5 = Pupitre. | 7 = Départ vers l'antenne. |
| 3 = Lampes de puissance. | | |

émetteurs ainsi construits, réglés d'avance, sont facilement transportables et peuvent être rapidement installés et mis en service. Nous pensons en reparler plus en détail dans un article ultérieur.

(MS 821)

F. Schmidlin. E. Aubort.

TABLE ANALYTIQUE DES MATIÈRES DE L'ANNÉE 1941.

Généralités.	Pages	Centrales et sous-stations.	Pages
Progrès constructifs réalisés par la Société Anonyme Brown, Boveri & Cie au cours de l'année 1940	3	Installation de convertisseurs de fréquence de Ski des chemins de fer de l'Etat norvégien	159
L'interdépendance des protections contre les mises à la terre, les courts-circuits et les surtensions dans les réseaux à haute tension	120	Centrales thermiques à mise en marche automatique	217
Aperçu sur l'état actuel de la protection contre les mises à la terre par bobines d'extinction	143	Chaudières.	
Moyens permettant de réduire la contrainte de l'isolation au point neutre des transformateurs et aux bobines d'extinction qui leur sont connectées	150	L'emploi des chaudières Velox dans les installations de chauffage Suralimentation, Velox et turbine à gaz. Leur création et leur développement chez Brown Boveri	89
Protection efficace du personnel des laminoirs à caoutchouc	177	La construction actuelle de la Velox résultant des expériences de plusieurs années	183
Des machines comme vous les désirez	177		221
Alarme! Obscurcissement!	178	Compresseurs et soufflantes.	
Influence du refroidissement à plusieurs étages sur le rendement dans le turbo-compresseur «Isotherme» de Brown Boveri	196	Le compresseur centrifuge «Isotherme»	108
Sur le calcul des échangeurs de chaleur	228	Influence du refroidissement à plusieurs étages sur le rendement dans le turbo-compresseur «Isotherme» de Brown Boveri	196
Détermination des variations d'état dans les turbo-machines à l'aide de l'accroissement de l'entropie	232	Influence de la compressibilité du fluide sur les caractéristiques d'une soufflante radiale	200
Production et réchauffage du vent dans les aciéries	240	Détermination séparée des pertes dans la roue et dans le diffuseur d'une soufflante radiale	203
Mise à la terre du neutre dans les réseaux à courant alternatif à très haute tension (superréseaux)	294	La suralimentation des moteurs à combustion alimentés au gaz de bois, spécialement pour les camions	206
Le problème des télécommunications et transmissions de signaux	333	Plate-forme d'essai pour turbo-compresseur de suralimentation à gaz d'échappement	211
Des débuts de la vapeur à haute pression	353	Considérations théoriques sur la suralimentation des moteurs d'avions pour le vol en altitude par des soufflantes entraînées par des turbines à gaz d'échappement	213
Nouvelles recherches sur les paliers de butée à segments	366		
Contribution à la théorie du graissage des engrenages et de la contrainte des flancs lubrifiés des dents	374	Haute fréquence.	
Appareils.		Alarme! Obscurcissement!	178
Transformateurs de mesure isolés aux gaz comprimés	84	Le problème des télécommunications et transmissions de signaux	333
Projets de protection sélective pour les réseaux d'interconnexion à haute tension	126	Lampes démontables de grande puissance	389
Relais de distance rapides	130	Générateurs d'ondes ultra-courtes	393
Relais primaires indiquant le courant ou la température	134	Procédés de brouillage automatique de la parole	397
Le développement récent du disjoncteur pneumatique ultra-rapide	138	Installations modernes de radio pour la police	409
Aperçu sur l'état actuel de la protection contre les mises à la terre par bobines d'extinction	143	Equipements portatifs de radio pour l'armée	413
La commande automatique des bobines d'extinction réglables en service	147	La radio dans l'aviation	414
Progrès réalisés dans le domaine des parafoudres	153	Modulation en fréquence	417
Résultats d'exploitation de relais de distance rapides Brown Boveri	155	Etude théorique du couplage de deux circuits oscillants par l'intermédiaire d'une ligne de transmission	423
L'interrupteur étoile-triangle avec protection pour moteur, type OS, et ses possibilités d'utilisation	169	Nouveaux résultats des recherches en physique atomique	436
Plate-forme d'essai pour turbo-compresseur de suralimentation à gaz d'échappement	211	Postes émetteurs de radiodiffusion	446
Disjoncteur ultra-rapide à courant alternatif à pouvoir de coupure très élevé, pour très hautes tensions	292	Machines électriques et transformateurs.	
Ejecteur d'air à jet de vapeur	352	Moyens permettant de réduire la contrainte de l'isolation au point neutre des transformateurs et aux bobines d'extinction qui leur sont connectées	150
Application et transmission de l'énergie électrique.		L'exploitation économique des transformateurs et le contrôle de leur exploitation	163
Le séchage électrique de l'herbe, un problème économique actuel pour la Suisse	75	Le transformateur à réglage continu pour basse tension	175
Protection efficace du personnel des laminoirs à caoutchouc	177	Le circuit magnétique des transformateurs de grande puissance pour très haute tension	307
Le labourage à l'électricité	178	Dispositif pour la mesure de la tension aux bornes d'un transformateur	311
Comparaison du coût de transmissions d'énergie à grande distance par courants alternatif et continu	249	Le couplage des transformateurs pour la transmission d'énergie par courant continu à très haute tension	314
Le problème de la stabilité des transmissions d'énergie à grande distance par courant triphasé	264	Mutateurs.	
Considérations techniques et économiques sur les lignes électriques aériennes	279	Courant de vapeur dans les gaines d'anodes des mutateurs et son influence sur le fonctionnement	97
Application de la construction articulée aux lignes à très haute tension	287	Mutateur à vide très poussé pour la transmission d'énergie à courant continu	319
Réalisations modernes de transmission d'énergie électrique par courant alternatif	288	Mutateur de grande puissance pour la transmission d'énergie par courant continu	322
Mise à la terre du neutre dans les réseaux à courant alternatif à très haute tension (superréseaux)	294	Traction.	
Phénomènes d'arcs à la terre dans une ligne de transport d'énergie en courant continu avec point milieu isolé	299	La suralimentation des moteurs à combustion alimentés au gaz de bois, spécialement pour les camions	206
La terre utilisée comme conducteur de retour dans la transmission d'énergie à grande distance	303	La locomotive à turbine à combustion	236
Le couplage des transformateurs pour la transmission d'énergie par courant continu à très haute tension	314	Turbines à vapeur et à gaz avec leurs accessoires.	
Mutateur à vide très poussé pour la transmission d'énergie à courant continu	319	Suralimentation, Velox et turbine à gaz. Leur création et leur développement chez Brown Boveri	183
Disposition d'une station d'extrémité pour la transmission à grande distance d'énergie électrique à haute tension	325	Quelques remarques sur les matériaux pour les turbines	209
Les harmoniques des transmissions à haute tension continue. Données expérimentales sur la première ligne de transport d'énergie Wettingen-Zurich	330	Sur le calcul des échangeurs de chaleur	228
		Histoire de la turbine à vapeur Brown Boveri	343
		Sur les pertes à l'extrémité des aubes de turbine	356
		Régule complexe pour turbines à vapeur	362
		Ejecteur d'air à jet de vapeur	382
		Turbo-alternateurs.	
		Perfectionnement des turbo-alternateurs Brown Boveri	367

LISTE DES ILLUSTRATIONS PUBLIÉES HORS TEXTE.

Généralités.	Pages	Haute fréquence.	Pages
Transport d'un turbo-groupe de 30 000 kW	2	Station d'émission à grande puissance	74
Soupapes de réglage et cylindre d'une turbine d'amont à très haute pression, les bâtis des soupapes sont soudés sur le cylindre et les tubulures sont aussi soudées	341	Lampe d'émission démontable Brown Boveri, de grande puissance	385
Appareils.		Machines électriques et transformateurs.	
Disjoncteur pneumatique de 220 kV en cours de montage dans nos ateliers	70	Atelier de montage des grands transformateurs de la Société Anonyme Brown, Boveri & Cie, à Baden	73
Bobine d'extinction de grande puissance	118	Poste de distribution.	
Application et transmission de l'énergie électrique.		Batterie blindée basse tension	158
Le cultivateur électrique avec moteur Brown Boveri	157	Traction.	
Chaudières.		Installation des machines de la première locomotive électrique du monde à turbine à combustion	1
Deux Velox de 50 t/h et une de 18 t/h peu avant la fin de leur montage dans nos ateliers	195	Turbines à vapeur et à gaz avec leurs accessoires.	
Transport de la chambre de combustion d'une chaudière Velox de 100 t/h	227	Rotor d'une turbine à gaz	181
		Soupapes de réglage et cylindre d'une turbine d'amont à très haute pression, les bâtis des soupapes sont soudés sur le cylindre et les tubulures sont aussi soudées	341
		Turbine à vapeur 50 000 kW, 3000 t/min	342



Editeur: Sté An. Brown, Boveri & Cie, Baden (Suisse). Imprimé par Kreis & Co., à Bâle.
En vente exclusivement aux librairies F. Rouge & Cie, à Lausanne et A. Francke A. G., à Berne ;
pour la France, à la librairie Gauthier-Villars & Cie, 55, quai des Grands-Augustins, à Paris.