

Jon B. Hagen

COMPRENDRE ET UTILISER L'ÉLECTRONIQUE DES HAUTES-FRÉQUENCES

DE LA GALÈNE À LA RADIOASTRONOMIE
PRINCIPES ET APPLICATIONS



PUBLITRONIC/ELEKTOR

JON B. HAGEN

National Astronomy and Ionosphere Center, Cornell University

COMPRENDRE ET UTILISER
L'ÉLECTRONIQUE DES
HAUTES-FRÉQUENCES

DE LA GALENE À LA RADIOASTRONOMIE
PRINCIPES ET APPLICATIONS

TRADUCTION : BRUNO SAVORNIN FIERZ

PUBLITRONIC / ELEKTOR

Sommaire

1 Introduction

Circuits hautes fréquences	2
Faible largeur de bande des signaux HF	3
Analyse des circuits alternatifs – rappel	3
Impédance et admittance	4
Résonance série	4
Résonance parallèle	4
Circuits non linéaires	5
Problèmes	5

2 Adaptation d'impédance 1

Adaptation par transformateur	8
Réseau en L	9
Méthode rapide pour concevoir un réseau en L	10
Les réseaux en Π et en T améliorent le facteur Q	12
Le réseau en double L abaisse le facteur Q	13
Circuits série et parallèle équivalents	13
Éléments réactifs à pertes et rendement des réseaux d'adaptation	14
Résumé sur le facteur Q	14
Problèmes	15

3 Amplificateurs linéaires

Amplificateur à une seule maille	16
Montage émetteur-suiveur	17
Amplificateurs à émetteur commun et à base commune	17
Un transistor, deux alimentations	18
Deux transistors, deux alimentations	19
Amplificateurs de courant alternatif	21
Amplificateurs basses fréquences	21
Amplificateurs hautes fréquences	23
Note sur l'adaptation d'un amplificateur de puissance à sa charge	24
Problèmes	25

4 Filtres 1

Filtres passe-bas normalisés	28
Exemple de filtre passe-bas	29
Évolution vers le filtre passe-bande	32
Bibliographie	33
Annexe 4.1	35
Problèmes	38

5 Convertisseurs de fréquence

Emploi d'un multiplicateur parfait comme changeur de fréquence	41
Changeurs de fréquence à commutation	43
Dispositif changeur de fréquence non linéaire	45
Mélangeur à diode	45
Problèmes	47

6 Récepteur radio

Caractéristiques essentielles	48
Amplification	48
Poste à galène	49
Récepteur à amplification directe	49
Récepteur superhétérodyne	50
Réjection de la fréquence-image	51
Comment résoudre le problème de la fréquence-image ?	51
Récepteur superhétérodyne à double changement de fréquence	52
Commande automatique de gain	53
Réducteur de bruit	53
Traitement numérique du signal dans un récepteur	53
Bibliographie	54
Problèmes	54

7 Amplificateurs en classe C et en classe D

Amplificateurs en classe C	55
Analyse simplifiée du fonctionnement en classe C	56
Analyse générale d'un fonctionnement en classe C avec un tube ou un transistor réel	58
Remarques sur l'attaque	58
Circuits alimentés en série et en parallèle	58
Utilisation d'un amplificateur en classe C comme multiplicateur de tension	60
Amplificateur de puissance en classe C	60
Amplificateur en classe C modifié pour un meilleur rendement	60
Amplificateur en classe D	61
Amplificateur en classe D résonnant en série	61
Amplificateur en classe D résonnant en parallèle	62
Classe C ou classe D ?	63
Bibliographie	63
Problèmes	63

8 Lignes de transmission

Notions fondamentales	65
Détermination de l'impédance caractéristique et de la vitesse de propagation	66
Modification d'une impédance par une ligne de transmission	67
Problèmes	69

9 Adaptation d'impédance 2

Impédances spécifiées par leur coefficient de réflexion	71
Problèmes	77

10 Alimentations

Redresseur à deux alternances	79
Autorégulation d'une alimentation à bobine en tête	80
Ondulation	81
Redresseur à une alternance	81
Alimentation régulée électroniquement	82

Redresseur triphasé	83
Problèmes	84

11 Modulation d'amplitude

Analyse de l'AM dans le domaine temporel	86
Analyse de l'AM dans le domaine fréquentiel	87
Modulation à haut niveau	89
Modulateur en classe A	90
Modulateur en classe B	90
Modulateur en classe S	91
Modulateur numérique/analogique	92
Ce qui se fait actuellement	92
Problèmes	93

12 Modulation à porteuse supprimée

Bande latérale unique	96
Détecteur-produit	96
Autres avantages de la BLU	97
Création d'un signal BLU	97
Méthode de filtrage	97
Méthode de mise en phase	98
Méthode de Weaver	99
BLU avec amplificateurs en classe C ou en classe D	100
Bibliographie	100
Problèmes	100

13 Oscillateurs

Oscillateurs à relaxation	102
Oscillateurs électroniques sinusoïdaux	103
Oscillateur involontaire	106
Oscillateur résonnant en série	107
Oscillateurs à résistance négative	108
Dynamique de l'oscillateur	109
Stabilité	109
Exemple de conception – l'oscillateur Colpitts	110
Exemple numérique	111
Problèmes	113

14 Boucles à phase asservie

Réglage de la phase par la commande de la fréquence	114
Analogie mécanique d'une PLL	115
Dynamique de la boucle	117
Filtre de boucle	118
Analyse linéaire d'une PLL	118
Réponse en fréquence d'une boucle de type I	119
Réponse en fréquence d'une boucle de type II	119
Réponse en régime transitoire	120
Utilisation d'un multiplicateur en détecteur de phase	121
Plage de fonctionnement et stabilité	122
Temps de verrouillage	122
Récepteur à PLL	123
<i>Bibliographie</i>	123
<i>Problèmes</i>	124

15 Synthétiseurs de fréquences

Synthèse directe	125
Synthèse directe par mélange et division	126
Synthèse indirecte	126
Synthèse directe numérique	127
Spectre de bruit du DDS	128
Vitesse de commutation et continuité de phase	130
Bruit de phase dû aux multiplicateurs et aux diviseurs	130
<i>Bibliographie</i>	131
<i>Problèmes</i>	131

16 Convertisseurs à découpage

Analyse d'un convertisseur élémentaire	133
Convertisseur <i>buck</i>	133
Mode continu	133
Mode discontinu	134
Convertisseur <i>buck/boost</i>	135
Mode continu	135
Mode discontinu	136
Convertisseur <i>boost</i>	137

Mode continu	137
Mode discontinu	137
Autres convertisseurs	138
Convertisseur à liaison par transformateur	138
Circuit de sortie de lignes dans les terminaux à tube à rayons cathodiques et dans les récepteurs de télévision	139
<i>Bibliographie</i>	142
<i>Problèmes</i>	142

17 Wattmètres directifs et ondes stationnaires

Wattmètre directif en ligne	144
Pont d'impédance résistif	146
Ondes stationnaires	147
Effets des ondes stationnaires sur la ligne de transmission d'une antenne	147
<i>Problèmes</i>	148

18 Amplificateurs HF de petits signaux

Réseau linéaire à deux accès (quadripôle)	149
Caractéristiques d'un amplificateur – gain, largeur de bande et impédances	150
Stabilité de l'amplificateur	150
Caractéristiques de surcharge	151
Intermodulation	151
Dynamique	153
Amplificateurs à bande étroite	153
Amplificateurs à large bande	153
Schéma équivalent d'un transistor	154
Conception d'un amplificateur	155
Amplificateurs BF simples	155
Amplificateur à base commune	156
<i>Bibliographie</i>	157
<i>Problèmes</i>	157

19 Filtres 2 – Filtres à couplage par résonateur

Inverseurs d'impédance	158
Exemple d'application – filtre passe-bande ayant une largeur de bande fractionnaire de 1%	162

Conséquences liées au fait que le facteur Q est fini	165
Procédures d'accord	166
Autres filtres	166
<i>Bibliographie</i>	166
<i>Problèmes</i>	166

20 Coupleurs hybrides

Couplage directif	168
Hybride à transformateur	169
Applications d'un hybride à transformateur . .	170
Hybrides en quadrature	170
Amplificateur symétrique	172
Combinateur de puissance	174
Autres hybrides	174
Diviseur de puissance (ou combinateur) de Wilkinson	174
Hybride en anneau	175
Hybrides en échelle	175
Hybrides à composants discrets	176
Coupleurs directifs généraux	177
<i>Bibliographie</i>	178
<i>Problèmes</i>	178

21 Bruit de l'amplificateur 1

Bruit thermique	181
Facteur de bruit	182
Amplificateurs montés en cascade	183
Autres paramètres du bruit	184
Mesure du facteur de bruit	185
<i>Problèmes</i>	185

22 Transformateurs et symétriseurs d'antenne

Courants dans le transformateur et transformateur idéal	187
Schéma équivalent en BF d'un transformateur parfaitement couplé et sans perte	188
Fonctionnement d'un transformateur parfaitement couplé et sans perte	189

Analogie mécanique d'un transformateur parfaitement couplé	190
Cas du transformateur imparfaitement couplé . . .	191
Transformateur à accord décalé	192
Transformateurs classiques avec noyaux magnétiques	193
Courants de Foucault et noyaux feuilletés . . .	193
Conception des transformateurs à noyau de fer	194
Température maximale et taille du transformateur	196
Transformateurs à ligne de transmission	197
Symétriseurs d'antenne	198
<i>Bibliographie</i>	202
<i>Problèmes</i>	202

23 Circuits à guides d'ondes

Guides d'ondes	203
Explication simple de la propagation dans un guide d'ondes	203
Propagation du mode fondamental dans un guide d'ondes rectangulaire	204
Longueur d'onde dans le guide d'ondes	205
Forme du champ magnétique	205
Courants dans les parois	206
Comparatif guide d'ondes contre câble coaxial pour une transmission d'énergie à faibles pertes . . .	207
Impédance d'un guide d'ondes	207
Adaptation de circuits à guides d'ondes	208
Jonctions de guides d'ondes à trois accès	209
Jonctions de guides d'ondes à quatre accès	210
Annexe 1 : guide d'ondes à pertes minimales contre ligne coaxiale à pertes minimales	210
Annexe 2 : dimensions d'une ligne coaxiale	213
Pertes minimales	213
Puissance maximale	213
Tension maximale	213
Qualités relatives d'une ligne coaxiale de 50 Ω	213
<i>Bibliographie</i>	214
<i>Problèmes</i>	214

24 Procédés de télévision

Analyse d'une image	215
Système de Nipkow	215
Norme de télévision NTSC	216
Signal vidéo	216
Synchronisation des lignes	217
Synchronisation des trames	218
Modulation	219
Son	220
Autres normes de télévision	221
Télévision en couleur	221
Trois couleurs dans un canal	221
Compatibilité	222
Filtres en peigne	224
Procédé PAL	225
Procédé SECAM	227
Émetteurs de télévision	228
Récepteurs de télévision	228
Récepteurs de télévision en couleur	229
Télévision numérique	234
Procédé ATSC	235
Compression vidéo	235
Couleur, son et paquets	236
<i>Bibliographie</i>	237
<i>Problèmes</i>	237

25 Modulateurs d'impulsions radar

Modulateurs à ligne	239
<i>Bibliographie</i>	243
<i>Problèmes</i>	243

26 Commutation TR

Techniques des radars à auto-duplexage	244
Composants et circuits de commutation TR	245
Commuteurs TR à ligne	246
Duplexeurs équilibrés	246
Commuteurs à diodes	247
Utilisation des diodes en commutation HF	248
<i>Bibliographie</i>	249
<i>Problèmes</i>	249

27 Démodulateurs et détecteurs

Détecteur à diode	251
Analyse du fonctionnement en supposant que le redresseur est parfait	252
Analyse du fonctionnement avec une diode réelle	252
Détecteur à diode à liaison en courant alternatif	254
Détection de la bande latérale unique (BLU) et du code Morse	254
Détecteur de produit pour l'AM	255
Détecteur AM synchrone	255
Démodulateurs FM	256
Démodulateur FM à PLL	256
Détecteur FM tachymétrique	256
Détecteur FM à ligne à retard	257
Démodulateur FM de la composante en quadrature	257
Détecteur de pente	258
Discriminateur de Foster-Seeley	259
Détecteurs de puissance	261
<i>Bibliographie</i>	262
<i>Problèmes</i>	263

28 Modulation de fréquence et modulation de phase

Bases de la modulation angulaire	264
Spectre de fréquences en FM	265
FM ou PM à bande étroite	265
Largeur spectrale de la FM à large bande	266
Multiplication de fréquence d'un signal FM	266
Bruit	266
Comment améliorer le rapport S/B en FM	267
Rapport S/B en sortie d'un signal AM ayant une porteuse de même puissance	268
Étude comparative du bruit FM/AM en présence de signaux forts	268
Préaccentuation et désaccentuation	269
FM, AM et capacité de transmission	269
<i>Bibliographie</i>	271
<i>Problèmes</i>	271

29 Antennes et propagation des ondes radioélectriques

Antennes	272
Ondes électromagnétiques	272
Propagation dans le vide	273
Directivité et gain d'une antenne	273
Aire d'interception équivalente d'une antenne	274
Liaison radio avec un astronef	275
Liaisons radio terrestres	276
Ionosphère	277
Propagation dans l'ionosphère	277
Réflexion des ondes sur l'ionosphère	278
Propagation diurne et nocturne	278
Autres modes de propagation	279
Bibliographie	279
Problèmes	279

30 Bruit de l'amplificateur 2

Adaptation de bruit	281
Schémas équivalents de quadripôles bruités	282
Facteur de bruit d'un schéma équivalent	282
Circuits en parallèle	284
Mesure de bruit	285
Bibliographie	286
Problèmes	286

31 Bruit de l'oscillateur

Spectre de puissance d'un oscillateur linéaire	289
Décroissance du bruit latéral	290
Bruit de phase	291
Effet de la non linéarité	292
Bibliographie	292
Problèmes	292

32 Radioastronomie et radarastronomie

Découverte du bruit cosmique	294
Radiométrie	295
Spectrométrie	296
Interférométrie	296

Interférométrie d'imagerie	297
Radarastronomie	298
Lune	298
Vénus	299
Cartographie à retard Doppler (ou <i>delay-Doppler</i>)	300
Overspreading	301
Bibliographie	301
Problèmes	302

33 Radiospectrométrie

Filtres et batteries de filtres	303
Spectrométrie à autocorrélation	304
Autocorrélateurs matériels	305
Autocorrélation à 1 bit	305
Spectrométrie à transformation de Fourier	307
Spectromètre acousto-optique	308
Spectrométrie à compression	309
Compression d'impulsions radar	310
Bibliographie	311
Problèmes	311

34 Appareils de contrôle de laboratoire

Mesures de la puissance	313
Mesures de la tension	313
Mesures de l'impédance	314
Mesures d'impédance avec balayage de fréquence	316
Problèmes	319

Index 321

1. Introduction

Prenons l'exemple envoûtant de la radio. La portée des postes émetteurs-récepteurs portatifs, voire même portables, peut atteindre plusieurs milliers de kilomètres. De minuscules émetteurs-récepteurs, embarqués à bord de sondes spatiales qui naviguent dans le système solaire, transmettent leurs données dans la bande de fréquences des micro-ondes. Les informations se propagent à la vitesse de la lumière. Rien ne laissait présager jusqu'à la fin du dix-neuvième siècle que la télégraphie sans fil verrait le jour, même en utilisant de puissantes dynamos et encore moins à l'aide d'appareils miniatures de faible coût. Les sujets de Sa Majesté la Reine Victoria pouvaient imaginer, forts des connaissances de l'époque, voler à bord d'un oiseau mécanique mû par des moteurs à vapeur, ou parcourir l'espace dans une fusée extrapolée des fusées de feu d'artifice chinoises. Mais quelle expérience aurait pu laisser entrevoir la possibilité de communiquer *sans fil* ? La découverte de la radio découle de l'étude de la physique théorique. Maxwell a approfondi la connaissance des lois connues sur l'électricité et le magnétisme et a introduit la fameuse notion de courant de déplacement $\delta D / \delta t$. Cela signifie qu'une variation du champ électrique induit une variation du champ magnétique, de même que Faraday avait découvert qu'une variation du champ magnétique créait un champ électrique. Les équations de Maxwell laissaient prévoir que des *ondes électromagnétiques* pouvaient se dégager des courants électriques qui les créaient et se propager indépendamment dans l'espace, les composantes des champs électrique et magnétique se repoussant mutuellement.

Les équations de Maxwell précisent que ces ondes se déplacent à une vitesse égale à $1/\sqrt{\epsilon_0 \mu_0}$; dans cette expression les valeurs des constantes ϵ_0 et μ_0 sont déterminées par la simple mesure des forces statiques qui se développent entre les charges électriques et les conducteurs électriques. Le résultat spectaculaire de ces mesures a naturellement conduit à connaître expérimentalement la vitesse de propagation de la lumière qui est égale à 3×10^8 m/s. C'est ainsi qu'a été mise en évidence la nature électromagnétique de la lumière. Hertz a mené de brillantes expériences, dans les années 1880, dans lesquelles il a produit et détecté des ondes électromagnétiques de longueurs d'ondes bien plus longues que celle de la lumière. L'utilisation des ondes hertziennes (ce sont les ondes radio qui nous sont si familières) pour transmettre des informations s'est développée conjointement avec la naissance de l'électronique.

Quelle est aujourd'hui la situation de la radio ? La radio AM (modulation d'amplitude), ancêtre du service de radiodiffusion, cohabite toujours avec la FM (modulation de fréquence), la télévision et la communication bilatérale. Le terme radio inclut à présent les radars, les satellites de surveillance, de navigation et de radiodiffusion, les téléphones cellulaires, les systèmes de télécommande ainsi que les transmissions de données sans fil. Les applications des hautes fréquences, en dehors de la radio, concernent les fours à micro-ondes, les systèmes d'imagerie médicale et la télévision par câble.

Comme vous pouvez le constater en examinant la figure 1-1, la radio couvre huit décades du spectre électromagnétique (de 10 kHz à 1000 GHz).

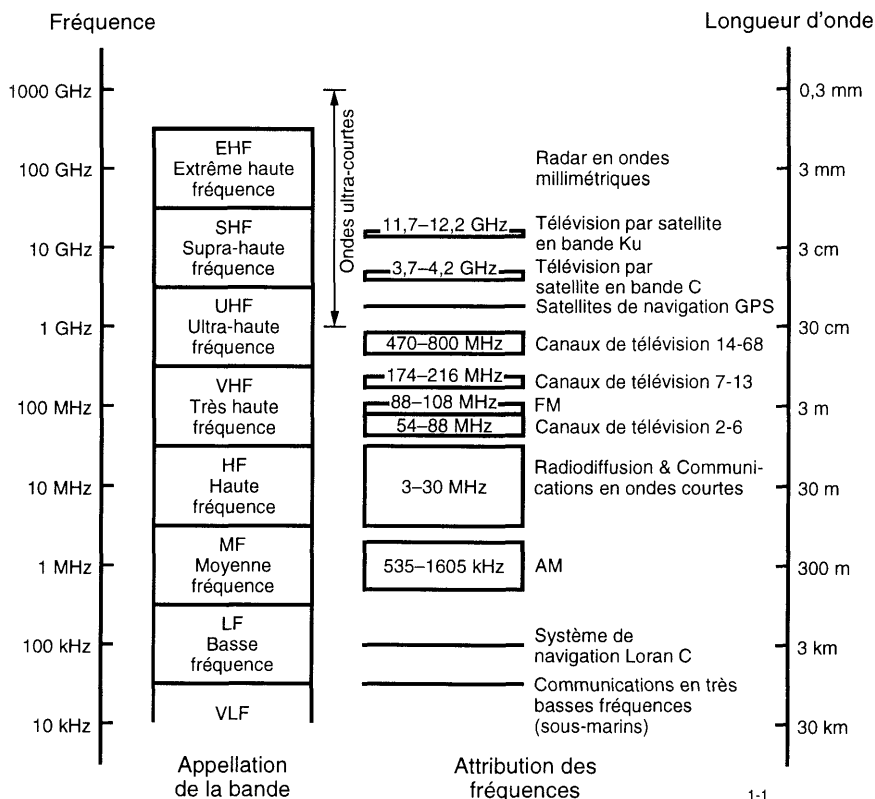


Figure 1-1. Le spectre des fréquences radioélectriques.

Circuits hautes fréquences

Les circuits que nous allons étudier dans cet ouvrage permettent de créer, amplifier, moduler, filtrer, démoduler, détecter et mesurer des tensions et des intensités alternatives dans la bande des hautes fréquences. Ces circuits sont les modules élémentaires qui vont nous permettre de construire des systèmes HF quelles que soient la puissance et la fréquence utilisées. C'est ainsi qu'un filtre passe-bande à six étages, ayant une réponse donnée, peut être d'une taille telle qu'il requiert un refroidissement par eau dans une application particulière, tandis que dans une autre, il sera miniaturisé à l'extrême. Selon la fréquence mise en œuvre, ce filtre sera réalisé en utilisant l'un des jeux de composants suivants : coffrets métalliques et lignes, bobines solénoïdales et condensateurs ou encore résonateurs mécaniques piézoélectriques ; le schéma électrique du filtre, quant à lui, sera toujours le même. Un amplificateur fonctionnant en classe C peut être un sous-ensemble d'un circuit intégré inclus dans un système de transmission de données sans fil ou l'ensemble le plus imposant d'un émetteur de radiodiffusion dont la puissance est de plusieurs mégawatts. Le schéma électrique sera toujours le même.

Faible largeur de bande des signaux HF

Vous remarquerez que la plupart des attributions de fréquences HF concernent de petites plages de largeurs de bandes, ce qui revient à dire que les largeurs de bandes sont faibles vis à vis des valeurs des fréquences centrales. La largeur de bande d'un signal émis représente moins de dix pour-cent (généralement bien moins) de la fréquence d'émission. Cela signifie qu'un signal HF issu d'un système radio est quasiment sinusoïdal. Pour acheminer une information, il faut *moduler* (c'est-à-dire faire varier d'une façon ou d'une autre) une onde « porteuse » HF parfaitement sinusoïdale. Tout type de modulation (qu'il soit audio, vidéo, par impulsion ou codage numérique) fait varier l'amplitude et/ou la phase de la porteuse. La largeur de bande d'une porteuse qui n'est pas modulée est extrêmement étroite ; il s'agit d'une raie spectrale pure. La modulation implique toujours l'élargissement de la raie dans le spectre, mais l'énergie mise en œuvre se concentre autour de la fréquence de la porteuse. Les relevés des tensions HF sur la ligne de transmission ou l'antenne d'un émetteur, effectués à l'aide d'un oscilloscope, montrent qu'elles sont presque sinusoïdales. En présence d'une modulation, l'amplitude et/ou la phase de la sinusoïde ne commencent à changer qu'au bout de plusieurs périodes. C'est ainsi que cette caractéristique de largeur de bande étroite nous permettra d'appliquer à la plupart des circuits HF l'analyse élémentaire des circuits alternatifs sinusoïdaux.

Analyse des circuits alternatifs – rappel

La théorie des circuits alternatifs classiques qui traite des tensions et des courants au sein d'un réseau linéaire, est basée sur la linéarité des éléments du circuit. Si l'on applique à l'entrée d'un circuit une tension ou un courant sinusoïdal, on aura en sortie des tensions et des courants stables parfaitement sinusoïdaux, de même fréquence que celle du générateur. La réponse d'un circuit alternatif répond généralement à une formule mathématique. On remplace un générateur de signaux donné par un générateur hypothétique dont le facteur de dépendance au temps est $e^{j\omega t}$ plutôt que $\cos(\omega t)$ ou $\sin(\omega t)$. Cette fonction de source comporte une partie réelle et une partie imaginaire, puisque $e^{j\omega t} = \cos(\omega t) + j \sin(\omega t)$. Une telle source irréelle (puisque'il s'agit d'une expression complexe) produit une réponse irréelle (complexe). Mais les parties réelle et imaginaire de la réponse, considérées séparément, sont bien les réponses physiques aux parties réelle et imaginaire de la source complexe. L'intérêt de cette façon apparemment indirecte de traiter le problème réside dans le fait que l'utilisation d'une source complexe permet de remplacer un système d'équations *différentielles* par un système d'équations linéaires *algébriques* facile à résoudre. Un circuit de topologie simple, comme c'est souvent le cas, se réduit à une simple boucle composée de branches parallèles et séries. Il existe de nombreux logiciels qui fournissent les valeurs des tensions et des intensités présentes au sein d'un circuit alternatif complexe. La plupart des versions du logiciel SPICE procéderont à cette analyse en régime alternatif stable (bien plus simple à effectuer que l'analyse en régime impulsionnel qui est leur vocation première). Des logiciels d'analyse en régime alternatif linéaire particulièrement adaptés aux HF et aux micro-ondes, comme COMPACT, TOUCH-STONE et MMICAD comportent des modèles de circuit pour les lignes triplaques, les guides d'ondes et autres composants HF. Vous pouvez écrire un programme simple qui vous permettra d'analyser un réseau en échelle (voir le Problème 1.3) et donc la plupart des filtres et des réseaux d'adaptation.

Impédance et admittance

Les coefficients qui apparaissent dans les équations algébriques qui décrivent un circuit dépendent des *impédances* (V/I) ou des *admittances* (I/V) complexes des composants *RLC*. La tension aux bornes d'une inductance est égale à $L di/dt$. Si le courant est $I_0 e^{j\omega t}$, la tension est égale à $j\omega L I_0 e^{j\omega t}$. L'impédance et l'admittance d'une inductance sont respectivement $j\omega L$ et $1/j\omega L$. Le courant dans un condensateur est $C dV/dt$; les impédance et admittance sont ainsi $1/j\omega C$ et $j\omega C$. Les impédance et admittance d'un résisteur sont respectivement égales à R et $1/R$. Des éléments montés en série sont parcourus par un même courant. L'impédance de l'ensemble est donc égale à la somme des impédances individuelles. Des éléments montés en parallèle ont à leur borne une même tension. L'admittance de l'ensemble est donc égale à la somme des admittances individuelles. Les parties réelle et imaginaire d'une impédance sont dénommées la résistance et la réactance tandis que les parties réelle et imaginaire d'une admittance (le réciproque de l'impédance) sont dénommés la *conductance* et la *susceptance*.

Résonance série

L'impédance d'un condensateur et d'une bobine d'inductance (que nous appellerons par la suite bobine) montés en série est égale à $Z_s = j\omega L + 1/j\omega C$. On peut également écrire cette expression sous la forme $Z_s = j(L/\omega)(\omega^2 - 1/LC)$; l'impédance est alors nulle lorsque la pulsation est égale à $1/\sqrt{LC}$. À cette fréquence de résonance, le circuit *série LC* se comporte comme un *court-circuit* parfait. Des tensions égales mais de signe opposé se développent aux bornes de la bobine et du condensateur : la chute de tension résultante est nulle. À la résonance et en régime permanent, il n'y a aucun transfert d'énergie à l'entrée ou à la sortie de ce montage : la tension totale étant toujours nulle, la puissance qui s'exprime par IV , est toujours nulle. Le circuit emmagasine toutefois de l'énergie ; celle-ci circule simplement entre la bobine et le condensateur. Vous noterez que ce circuit n'est ni plus ni moins qu'un filtre passe-bande.

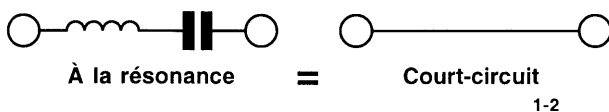


Figure 1-2. Circuit résonant série LC.

Résonance parallèle

L'admittance d'un condensateur et d'une bobine montés en série est égale à $Y_p = j\omega C + 1/j\omega L$; elle est alors nulle lorsque la pulsation est égale à $1/\sqrt{LC}$. À la fréquence de résonance, le circuit *parallèle LC* se comporte comme un *circuit ouvert* parfait – c'est un simple filtre coupe-bande (voir la figure 1-3). De même que dans le cas d'un circuit série LC, le circuit parallèle LC emmagasine une quantité d'énergie déterminée qui dépend de la tension appliquée. Ces deux circuits élémentaires constituent des sous-ensembles fondamentaux dans la conception des circuits hautes fréquences.

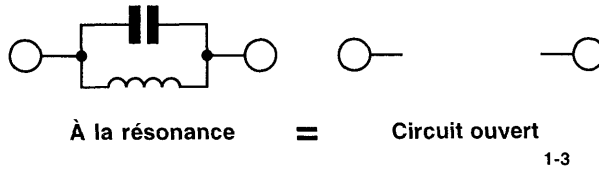


Figure 1-3. Circuit résonant parallèle LC.

Circuits non linéaires

De nombreux circuits fondamentaux hautes fréquences, parmi lesquels nous pouvons citer les changeurs de fréquence, les modulateurs et les détecteurs font appel à des éléments non linéaires tels les diodes et les interrupteurs à transistor saturé. Nous ne pouvons pas dans ce cas appliquer l'analyse linéaire $e^{j\omega x}$ mais nous devons mettre en œuvre l'analyse de signaux dans le domaine temporel. Il sera généralement possible de remplacer les éléments non linéaires par des modèles simples afin d'expliquer le fonctionnement d'un circuit. Une modélisation complète sur ordinateur sera employée lorsqu'il sera indispensable d'obtenir des simulations précises du fonctionnement d'un montage.

Problèmes

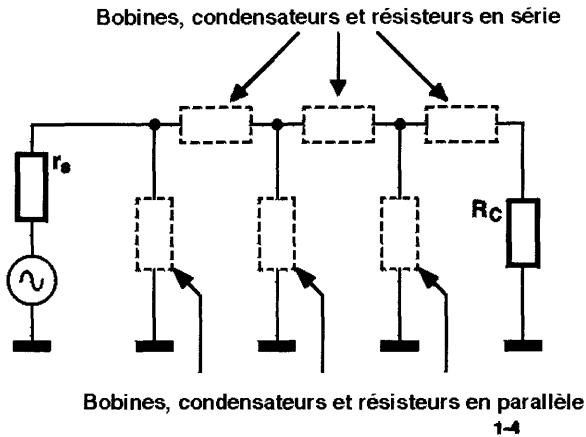
1.1 La résistance de la source d'un générateur est r_s et sa tension efficace à vide est V_0 . Montrer que la puissance maximale que peut fournir le générateur est donnée par la formule suivante : $P_{\max} = U_0^2 / (4r_s)$ et qu'elle sera obtenue lorsque la résistance de charge R_c est égale à la résistance de la source r_s .

1.2 On insère un réseau passif constitué par exemple de résisteurs, de condensateurs et de bobines entre un générateur de résistance de source r_s et une résistance de charge R_c . La *réponse de puissance* du réseau (en fonction de ces résistances) est définie comme étant la fraction de la puissance maximale que peut fournir le générateur qui atteint la charge. La réponse de puissance d'un circuit *sans perte* (c'est-à-dire qui ne renferme aucun résistor ni autre élément dissipatif) dépend de l'impédance $Z_{\text{entrée}}$ vue de l'entrée du circuit auquel a été connectée la charge. Montrer que l'expression ci-dessous exprime la réponse de puissance d'un réseau sans perte :

$$P(\omega) = \frac{4r_b R(\omega)}{[R(\omega) + r_b]^2 + [X(\omega)]^2}$$

où R = partie réelle de $Z_{\text{entrée}}$ et X = partie imaginaire de $Z_{\text{entrée}}$.

1.3 La plupart des filtres et des réseaux d'adaptation prennent la forme d'un *réseau en échelle*. C'est ce qu'illustre la figure ci-dessous. Écrivez un programme qui vienne lire les valeurs des éléments série et parallèle du réseau et qui fournisse la réponse de puissance telle qu'elle est définie dans le Problème 1.2. Ce problème servira à la résolution d'autres problèmes et sera complété à de nombreuses reprises.



Astuces : une façon de procéder consiste à partir de la résistance de charge et de calculer l'impédance d'entrée au fur et à mesure que l'on ajoute un élément. Une fois que tous les éléments ont été pris en compte, la formule du Problème 1.2 donne la réponse de puissance – tant qu'aucun des composants n'est un résisteur. Cette méthode s'applique quelle que soit la fréquence.

Voici un meilleur procédé pas plus complexe que le précédent, qui permet d'insérer des résisteurs : supposons qu'une intensité de $1 + j0$ A circule dans le résisteur de charge. La tension à cet endroit est égale à $R_c + j0$ V. Déplacez-vous d'un élément vers la gauche. S'il s'agit d'un élément série, le courant reste le même mais la tension croît d'une quantité égale à IZ , où Z est l'impédance de l'élément série. S'il s'agit d'un élément parallèle, la tension reste la même mais le courant croît d'une quantité égale à VY où Y est l'admittance de l'élément parallèle. Ajoutez les éléments les uns après les autres et modifiez en conséquence le courant et la tension. Une fois que vous aurez pris en compte tous les éléments, vous obtiendrez les valeurs de la tension et du courant en entrée et vous pourrez alors calculer l'impédance d'entrée totale du réseau terminé par le résisteur de charge. Ce n'est pas tout à fait fini car la dernière étape consiste à tenir compte de la résistance de la source r_s exactement comme d'une autre impédance série. Vous obtenez ainsi la tension fournie par le générateur et vous en déduisez la puissance maximale disponible. Connaissant déjà la puissance fournie à la source, $(I)^2 R_c$, il est très simple de trouver la réponse de puissance. Faites de même pour chacune des fréquences désirées.

Les valeurs des éléments de l'échelle (de même que la fréquence initiale, la fréquence finale et le pas de fréquence) peuvent être traités comme des données, c'est-à-dire qu'ils se situeront dans une même zone du programme ou dans un fichier séparé, ce qui facilitera leur modification éventuelle. Pour le moment, seuls six types d'éléments sont à prendre en considération : les bobines, les condensateurs et les résisteurs montés en série ou en parallèle. Chaque élément du circuit sera donc identifié par « PL », « SL », « PC », « SC », « PR » et « SC » ou encore 1, 2, 3, 4, 5, 6 (ou autre chose) et sa valeur exprimée en henrys, farads ou ohms. Structurez la base de données du fichier du circuit afin qu'elle débute par l'élément qui se trouve à proximité immédiate de R_c et se termine par un identificateur comme par exemple « EOF » (*End Of File* ou *Fin de Fichier*) ou tout autre valeur caractéristique.

Vous découvrez ci-après un exemple d'un tel programme écrit en Microsoft QBasic. Les valeurs des composants (voir la ligne 600) correspondent au schéma qui suit le programme. Aidez-vous de cet exemple pour vérifier votre propre programme. Vous pourrez le compléter en y incluant des en-têtes qui faciliteront sa lecture lorsque vous l'imprimerez. Ajoutez-y, si vous le désirez les instructions qui permettront de tracer les courbes de réponse.

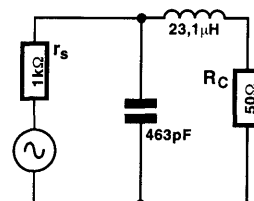
CLS

```

'Programme écrit en Qbasic permettant de calculer
'l'atténuation d'un réseau RCL en échelle
RSOURCE = 1000: RCHARGE = 50
FOR F = 1000000! TO 2000000! STEP 50000!
OMEGA = 2 * 3.14159 * F
IR = 1: II = 0: UR = IR * RCHARGE: UI = II * RCHARGE
'intensité de 1 A dans la charge
READ TYPE$
'"PC" représente un condensateur en parallèle,
'"SL" une bobine en série, etc.
DO UNTIL TYPE$ = "EOF"
'le caractère "EOF" signifie que la fin du fichier
'contenant les paramètres du circuit est atteinte
READ VALEUR
IF TYPE$ = "PC" THEN B = OMEGA * VALEUR: G = 0: GOSUB 400
'mise à jour de I
IF TYPE$ = "SC" THEN X = -1 / (OMEGA * VALEUR): R = 0: GOSUB 500
'mise à jour de U
IF TYPE$ = "PL" THEN B = -1 / (OMEGA * VALEUR): G = 0: GOSUB 400
'mise à jour de I
IF TYPE$ = "SL" THEN X = OMEGA * VALEUR: R = 0: GOSUB 500
'mise à jour de U
IF TYPE$ = "PR" THEN G = 1 / VALEUR: B = 0: GOSUB 400 'mise à jour de I
IF TYPE$ = "SR" THEN R = VALEUR: X = 0: GOSUB 500 'mise à jour de U
READ TYPE$
LOOP
R = RSOURCE: X = 0: GOSUB 500 'tension délivrée par le générateur
'calcul de la fraction de la puissance maximale transférée possible
FRAC = (RCHARGE) / ((UR * UR + UI * UI) / (4 * RSOURCE))
PRINT F; FRAC; 10 / LOG(10) * LOG(FRAC) 'impression de la fréquence, de la
'fraction sous forme décimale et de la fraction exprimée en décibels
RESTORE 600 'retour au début de la liste des paramètres
NEXT F
END
400 IR = IR + (UR * G - UI * B): II = II + (UI * G + UR * B): RETURN
'sous-programme permettant de mettre à jour la partie réelle
'et la partie imaginaire de I
500 UR = UR + (IR * R - II * X): UI = UI + (II * R + IR * X): RETURN
'sous-programme permettant de mettre à jour la partie réelle
'et la partie imaginaire de U
'fichier comportant les paramètres pour un circuit ayant un condensateur
'de 463 pF en parallèle et une bobine de 23,1 H en série
600 DATA SL,23.1E-6,PC,463E-12,EOF

```

Résultats du programme		
Fréquence (MHz)	Fraction	dB
1,00	0,4179	-3,789
1,05	0,4601	-3,372



2. Adaptation d'impédance 1

Le mot adaptation sous-entend normalement l'utilisation d'un réseau sans perte (donc non résistif) inséré entre une source alternative (ici hautes fréquences) et une charge afin d'optimiser la puissance délivrée à cette charge. Par exemple, un dispositif d'accord d'antenne est un circuit qui permet d'adapter une antenne à un émetteur. Ce même circuit s'appellera dispositif d'accord de sortie s'il est inclus dans l'émetteur. Examinez le schéma de la figure 2-1 : dans ce montage à courant continu, la puissance maximale est transférée lorsque la résistance de charge est égale à celle de la source. Vous pouvez aisément vérifier que l'égalité de ces deux résistances entraîne le maximum du produit courant $\times$ tension de la charge.

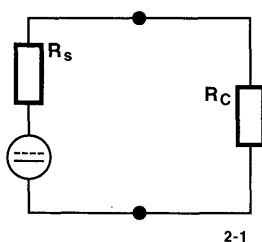


Figure 2-1. La puissance maximale est transférée à la charge R_C lorsque $R_C = R_S$.

Adaptation par transformateur

Lorsqu'on est en présence d'une source de courant alternatif, un transformateur peut adapter la résistance de charge à celle de la source (et réciproquement) : c'est ce qu'illustre la figure 2-2. L'emploi d'une source alternative présente fréquemment une complication : la source et/ou la charge peuvent être réactives, ce qui signifie qu'elles ont une inévitable réactance interne. Une antenne est un exemple de charge réactive ; nombreuses sont les antennes qui ne sont purement résistives qu'à une seule fréquence. Au-dessus de la fréquence de résonance, elles se comportent généralement comme un résisteur en série avec une bobine, tandis qu'en-dessous de la fréquence de résonance, elles sont équivalentes à un résisteur en série avec un condensateur. Une méthode évidente pour résoudre ce problème consiste tout d'abord à annuler la réactance : l'impédance de la charge et/ou de la source devient alors purement résistive. Il suffit alors d'utiliser un transformateur pour adapter les résistances. Dans le schéma de la figure 2-3, une bobine annule la réactance d'une charge capacitive (mais non purement capacitive). Si nous travaillions à 50 Hz, nous pourrions dire que la bobine corrige le facteur de puissance de la charge.

Du point de vue de la charge, le réseau d'adaptation convertit l'impédance de la source, $R_S + j0$, en conjugué complexe de l'impédance de la charge. En règle générale, lorsqu'on insère un réseau d'adaptation entre deux circuits, chaque circuit verra une impédance qui est le conjugué complexe de sa propre impédance. Chaque fois qu'une source et/ou une charge ont une composante réactive, l'adaptation dépendra de la fréquence, ce qui implique qu'elle ne sera parfaite qu'à la fréquence pour laquelle elle a été conçue. De fait, en présence de sources et/ou de charges réactives, *tout* circuit d'adaptation sans perte dépendra, que nous le voulions ou non, de la fréquence – en quelque sorte un filtre.

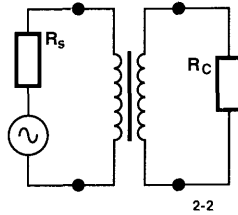


Figure 2-2. Le transformateur convertit R_C en R_S afin que la puissance délivrée soit maximale.

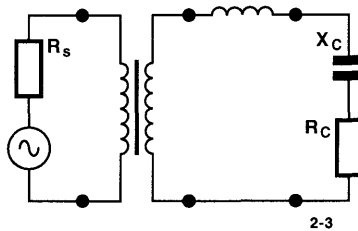


Figure 2-3. La réactance série rend la charge purement résistive.

Réseau en L

La plupart du temps, les circuits d'adaptation n'utilisent pas de transformateur (c'est-à-dire qu'il n'y a aucun couplage par bobine). La figure 2-4 représente un réseau en L à deux éléments (la lettre « L » est inversée) qui adapte la source à une charge dont la résistance est plus faible que celle de la source. L'astuce consiste à placer un composant réactif, X_P , en *parallèle* avec la résistance la plus *élevée*. Prenons un exemple concret : $R_S = 1000 \, \Omega$ et $R_C = 50 \, \Omega$. L'impédance de la partie gauche du circuit est donnée par :

$$\begin{aligned} Z_{\text{gauche}} = R_{\text{gauche}} + jX_{\text{gauche}} &= \frac{1000 jX_P}{1000 + jX_P} = \frac{(1000 jX_P)(1000 - jX_P)}{(1000 + jX_P)(1000 - jX_P)} \\ &= \frac{1000^2 jX_P + 1000 X_P^2}{1000^2 + X_P^2} \end{aligned} \quad (2-1)$$

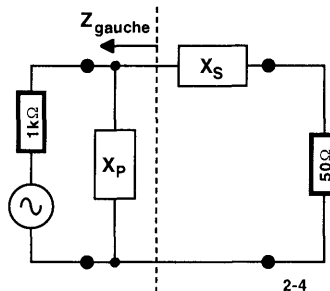


Figure 2-4. Les deux éléments réactifs (X_P , en parallèle ; X_S , en série) du réseau en L adaptent R_C à R_S .

Nous pouvons choisir la valeur de X_P pour que la partie réelle de Z_{gauche} soit égale à $50\ \Omega$, c'est-à-dire égale à la résistance de la charge. Nous trouvons, grâce à l'équation 2-1, que $X_P^2 = 52\,632$, ce qui nous offre le choix entre $X_P = 229$ (une bobine) et $X_P = -229$ (un condensateur). La valeur de la résistance série équivalente de la partie gauche du circuit est correcte (égale à $50\ \Omega$) ; on note également la présence d'une réactance série équivalente, X_{gauche} , donnée par la partie imaginaire de l'équation 2-1. Il nous est facile de la supprimer en insérant en série un élément réactif, X_S , égal à $-X_{\text{gauche}}$. La figure 2-5 représente les circuits d'adaptation obtenus lorsque X_P est une bobine et lorsque X_P est un condensateur.

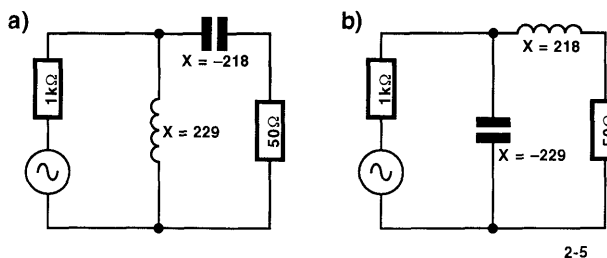


Figure 2-5. Voici deux réalisations possibles du réseau en L de la figure 2-4.

La dernière étape consiste à trouver les valeurs de L et C qui fournissent les réactances spécifiées à la fréquence donnée. Pour le circuit de la figure 2-5b, $\omega L = 218$. Supposons que la fréquence de travail soit égale à $1,5\ \text{MHz}$ ($\omega = 2\pi \times 1,5 \times 10^6$) ; elle se situe à proximité du haut de la bande de radiodiffusion AM. Dans ce cas, $L = 23,1\ \mu\text{H}$ et $C = 462\ \text{pF}$. Vous constaterez que les valeurs de ces deux éléments réactifs ne dépendent que des valeurs des résistances de source et de charge. Ce circuit d'adaptation à deux éléments n'offre aucune liberté sauf celle de choisir quel élément sera la bobine et lequel sera le condensateur. À la fréquence pour laquelle il a été conçu, l'adaptation est parfaite, mais dès que l'on s'en éloigne, nous devons accepter la réponse en fréquence qui en résulte. Les tracés de la figure 2-6 illustrent les réponses en fréquence (c'est-à-dire la fraction de puissance qui atteint la charge en fonction de la fréquence) des deux circuits de la figure 2-5. Vous constatez qu'aux alentours de la fréquence pour laquelle ils ont été conçus (c'est-à-dire aux environs du pic de résonance), les courbes sont quasiment identiques. Autrement, en examinant les schémas des circuits, il est possible de prévoir quelles seront les fréquences de coupure aux très basses fréquences pour le circuit de la figure 2-5a et aux très hautes fréquences pour celui de la figure 2-5b.

Méthode rapide pour concevoir un réseau en L

Si vous avez toujours en mémoire le fait que la réactance parallèle se place sur la résistance la plus élevée, vous allez pouvoir répéter les étapes précédentes et concevoir des réseaux en L. Si vous envisagez d'effectuer plusieurs fois ces opérations, il est préférable de se souvenir, lors de la conception d'un réseau en L, de ce que l'on appelle le « *facteur de qualité* » (également appelé *facteur Q*) :

$$Q_{\text{EL}} = \left(\frac{R_{\text{élevée}}}{R_{\text{basse}}} - 1 \right)^{1/2} \quad (2.2)$$

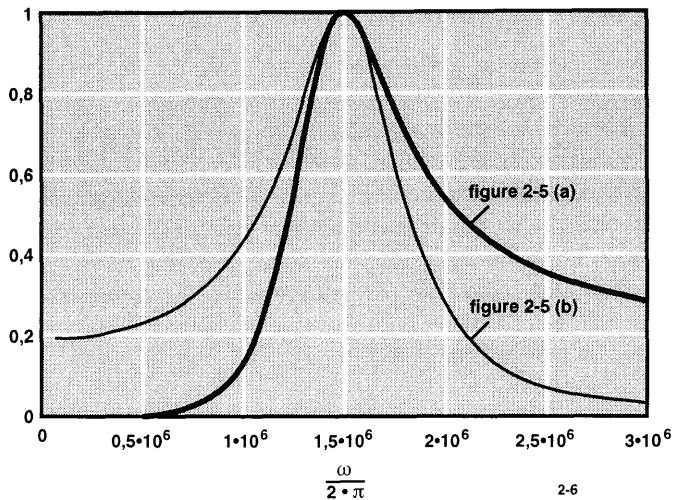


Figure 2-6. Ces courbes de réponse correspondent à celles des réseaux en L de la figure 2-5.

Vous pouvez vérifier (voir le Problème 2.6) que les rapports $R_{\text{élevée}}/X_P$ et X_S/R_{basse} sont bien égaux à ce facteur de qualité Q_{EL} . N'oubliez pas la définition de Q_{EL} et ces rapports vous fournissent immédiatement les valeurs des réactances du réseau en L. Vous pouvez également vous assurer que lorsque la valeur de Q_{EL} est élevée, les réactances des deux éléments du réseau en L sont presque égales et de signe opposé, c'est-à-dire qu'ils résonnent ensemble à la fréquence nominale. La valeur des réactances est dans ce cas égale à la moyenne géométrique de $R_{\text{élevée}}$ et de R_{basse} (ce qui est facile à retenir).

Lorsque le rapport de la résistance de source à la résistance de charge est très différent de l'unité, l'adaptation d'un réseau en L ne se fait que sur une bande de fréquence étroite ; il en résulte que l'adaptation ne sera bonne qu'à proximité immédiate de la fréquence nominale. Réciproquement, lorsque le rapport d'impédance est proche de l'unité, l'adaptation est bonne sur une large bande de fréquence. La largeur de tout phénomène de résonance se caractérise par un facteur, le *facteur Q effectif* (ou plus simplement Q) qui est égal au rapport de la fréquence centrale à la largeur de bande à -3 dB (c'est la différence de fréquences entre les points situés à mi-puissance). On peut dire de même que Q_{eff} est l'expression réciproque de la largeur de bande fractionnaire. Lorsqu'un générateur de tension idéal attaque un simple circuit série RLC , le facteur Q_{eff} est donné par X/R , où X représente soit X_L soit X_C à la fréquence centrale (puisque ces deux termes sont égaux). Le réseau d'adaptation en L est équivalent à un simple circuit série RLC , mais la valeur de Q_{EL} est deux fois celle de Q_{eff} puisque se trouve également en série la résistance de source non nulle. Le circuit d'adaptation rend la résistance de source effective égale à la résistance de charge et la résistance série totale est le double de la résistance de charge. Il en ressort que la largeur de bande fractionnaire est donnée par $1/Q_{\text{eff}} = 2/Q_{\text{EL}}$. Dans de nombreuses applications, la largeur de bande du circuit d'adaptation est un paramètre important : l'adaptation procurée par le réseau en L (qui, rappelons-le ne dépend que des résistances de source et de charge) peut être trop étroite ou trop large. Lorsqu'on adapte par exemple une antenne à un récepteur, on peut désirer obtenir une largeur de bande étroite afin d'éviter que des signaux issus de puissantes stations voisines viennent surcharger le récepteur. Dans un autre cas, il se peut que le signal produit par un émetteur modulé ait une largeur de bande supérieure à celle du réseau en L. Les réseaux que nous allons étudier à présent permettent de résoudre ces problèmes.

Les réseaux en Π et en T améliorent le facteur Q

Il est possible d'améliorer le facteur de qualité Q en plaçant dos à dos deux réseaux en L, chacun d'eux abaisse vers une impédance centrale plus faible la résistance du générateur et celle de la charge. On obtient alors le réseau en Π (il s'agit de la lettre majuscule grecque pi) qu'illustre la figure 2-7. Lorsque nous avons utilisé un simple réseau en L, nous avons obtenu $Q_{EL} = \sqrt{19} = 4,4$. Grâce à ce réseau en Π , la résistance de source de $1000\ \Omega$ et la résistance de charge de $50\ \Omega$ sont adaptées à une impédance centrale égale à $10\ \Omega$ (valeur arbitraire). La largeur de bande obtenue est équivalente à celle d'un réseau en L avec un facteur de qualité $Q_{EL} = 11,95$. Lorsque $R_{\text{élevée}} \gg R_{\text{basse}}$, la largeur de bande d'un réseau en Π est équivalente à celle d'un réseau en L avec $Q_{EL} = \sqrt{R_{\text{élevée}}/R_{\text{centrale}}}$. La largeur de bande fractionnaire est encore donnée par $1/Q_{\text{eff}} = 2/Q_{EL}$. La figure 2-8 représente la courbe de réponse d'un réseau en Π ainsi que celles des deux réseaux en L de la figure 2-5.

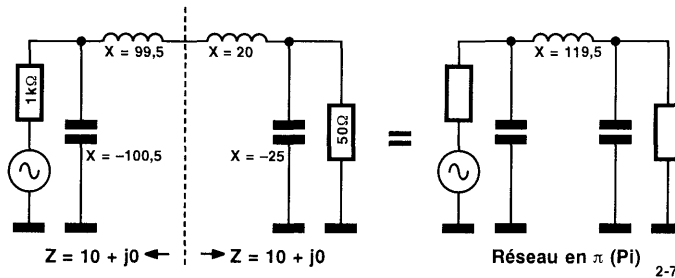


Figure 2-7. Ce réseau en Π (constitué de deux réseaux en L montés dos à dos) offre un meilleur facteur Q .

Vous avez certainement deviné qu'il était aussi possible de placer « face à face » deux réseaux en L, chacun d'eux accroissant vers une impédance centrale plus élevée la résistance du générateur et celle de la charge. On obtient alors le réseau en T de la figure 2-9.

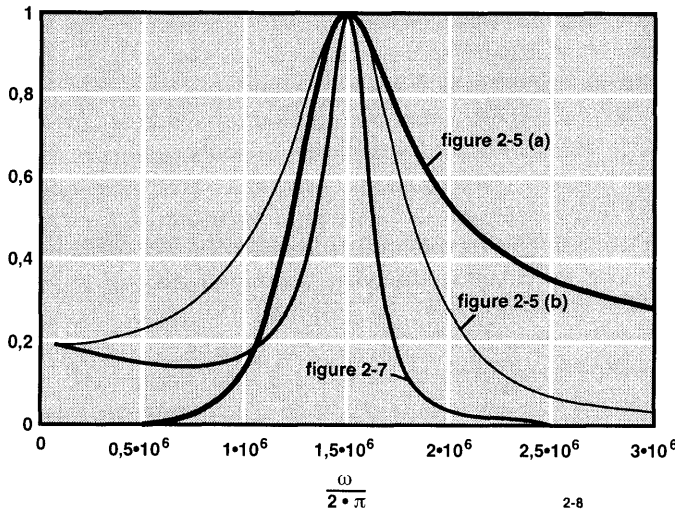


Figure 2-8. Courbe de réponse du réseau en Π de la figure 2-7 (comparée à celles des réseaux en L de la figure 2-5).

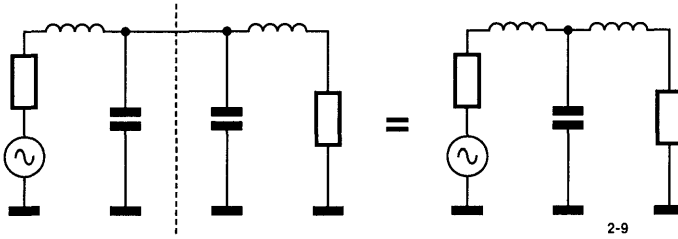


Figure 2-9. De même que dans le cas d'un réseau en Π , le facteur Q d'un réseau en T est meilleur.

Le réseau en double L abaisse le facteur Q

Dans le cas d'un réseau en double L (illustré à la figure 2-10), le premier étage modifie la valeur de l'impédance afin qu'elle soit comprise entre celle de la source et celle de la charge. Le second étage achève l'adaptation. Il est bien entendu possible d'accroître le nombre d'étages. Une longue chaîne de réseaux en L forme une ligne de transmission artificielle dont l'impédance décroît et permet d'obtenir une adaptation indépendante de la fréquence. Les lignes de transmission réelles (à savoir des lignes avec des éléments L et C répartis) sont parfois amincies pour obtenir cette sorte de transformation d'impédance. On appelle quelquefois transformateur une telle ligne, puisqu'elle réalise une adaptation indépendante de la fréquence, de la même façon que le transformateur de la figure 2-2.

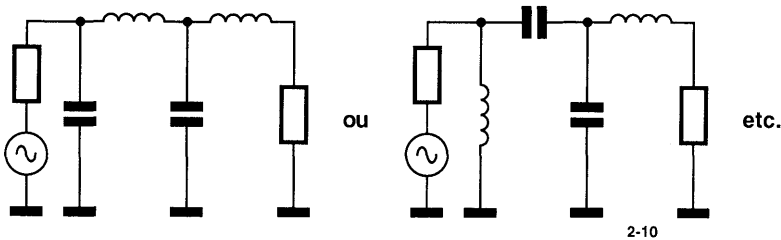


Figure 2-10. Ce réseau en double L permet d'obtenir un facteur Q plus faible (et donc une largeur de bande plus élevée).

Circuits série et parallèle équivalents

Pour concevoir le réseau en L que nous avons étudié, nous sommes partis du fait qu'un circuit parallèle à deux éléments XR , où $1/Z = 1/R_P + 1/(jX_P)$, avait un circuit série équivalent, où $Z = R_S + jX_S$. On se sert tellement fréquemment des conversions entre circuits série et parallèle qu'il est inutile d'en dire plus. Supposons que vous ayez par exemple une antenne ou une boîte noire avec deux bornes et que vous fassiez des mesures à une *fréquence précise* ; vous ne pouvez que déterminer si la boîte noire est « capacitive » (c'est-à-dire équivalente à un circuit R et C) ou « inductive » (c'est-à-dire équivalente à un circuit R et L). Considérons le cas où elle est capacitive. Il est possible de la représenter sous la forme d'un circuit série où $Z = R_S + 1/(j\omega C_S)$ ou celle d'un petit circuit parallèle où $1/Z = 1/R_P + j\omega C_P$. Tant que vous ne travaillez qu'à la fréquence désirée (ou à proximité immédiate), l'une et l'autre représentations sont valables, même si la boîte noire renferme un montage complexe faisant appel à des résisteurs discrets, des condensateurs, des bobines, des lignes de transmission, des assemblages métalliques et résistifs, etc. Si vous mesuriez

l'impédance à d'autres fréquences, vous pourriez vous apercevoir que la boîte noire contient en fait un simple circuit RC parallèle ou série ou que sa variation d'impédance ressemble tout au moins plus à celle d'un simple circuit parallèle qu'à celle d'un simple circuit série.

Éléments réactifs à pertes et rendement des réseaux d'adaptation

Nous avons considéré jusqu'à présent que les éléments constitutifs d'un réseau (bobines et condensateurs) étaient parfaits. Les composants réels, toutefois, présentent des pertes dues à la conductibilité limitée des métaux, aux matériaux diélectriques ou magnétiques et même au rayonnement. La puissance dissipée dans un élément non idéal est perdue puisqu'elle n'atteint pas la charge ; c'est pourquoi, en utilisant un composant à pertes, nous devons prendre en considération le rendement d'un réseau d'adaptation. Comme nous l'avons déjà dit, un composant réactif à pertes peut être modélisé sous la forme d'un composant idéal L ou C auquel on ajoute en série ou en parallèle une résistance. Nous pouvons valablement supposer que les valeurs de L et C et de la résistance série ou parallèle associée sont constantes sur la plage de fréquences considérée. Étudions le rendement d'un réseau en L constitué d'une bobine série et d'un condensateur parallèle. Nous supposons que les pertes dans le condensateur sont négligeables vis à vis de celles de la bobine, ce qui est le cas avec les composants à constante localisée). Nous modéliserons notre bobine à pertes par une bobine parfaite en série avec une résistance r_s . On appelle *facteur de qualité* Q_{NC} , le rapport de la réactance inductive X_C à cette résistance où l'indice « NC » signifie « *facteur Q non chargé* » ou *indice Q* (si l'on fait abstraction du résisteur série, il est certain que l'on obtient un meilleur indice de qualité). Vous remarquerez que cette résistance, de même que la bobine, se trouvent en série avec la résistance de charge et que le même courant I les traverse toutes deux. La puissance délivrée à la charge s'exprime par $I^2 R_C$ et celle qui l'est dans r_s est $I^2 r_s$. Nous allons pouvoir calculer le rendement du réseau d'adaptation en nous servant des relations $X_S = Q_{EL} R_C$ et $Q_{NC} = X_S / R_S$:

$$\eta = \text{rendement} = \frac{\text{puissance de sortie}}{\text{puissance d'entrée}} = \frac{I^2 R_C}{I^2 R_C + I^2 r_s} = \frac{1}{1 + Q_{EL} / Q_{NC}} \quad (2.3)$$

Le rendement est le meilleur lorsque le rapport Q_{EL} / Q_{NC} est le plus petit, c'est-à-dire le rapport du facteur Q chargé au facteur Q non chargé. Si nous modélisons la bobine à pertes sous la forme d'un circuit parallèle LR et si nous définissons le facteur Q non chargé comme étant égal à r_p / X_p , nous obtiendrons pour le rendement la même expression (voir le Problème 2.7). De même, si c'est un condensateur qui présente des pertes, cette expression toujours valable tant que nous définissons le facteur Q non chargé du condensateur comme étant le rapport de la résistance parallèle sur la réactance parallèle ou celui de la réactance série sur la résistance série. Lorsque la résistance de la charge est très différente de celle de la source, le facteur Q effectif d'un réseau en L sera élevé ; aussi, pour obtenir un rendement élevé, faudra-t-il que le facteur Q non chargé des composants soit très élevé. Pour obtenir un meilleur rendement, on pourra mettre en œuvre un réseau en double L qui possède un facteur Q_S chargé plus faible.

Résumé sur le facteur Q

Facteur Q chargé. Le facteur Q associé à des circuits peut être soit élevé soit faible, selon l'application considérée. Celui des filtres à bande étroite est élevé et celui des circuits d'adaptation à large bande est faible. Le facteur Q n'est pas représentatif d'une quelconque qualité.

Facteur Q non chargé. Il fournit une mesure de la qualité puisqu'il tient compte des pertes dans les composants : moins un composant a de pertes, meilleur sera le rendement.

Problèmes

2.1 Un résistor au carbone 1/4 W de valeur nominale $47\ \Omega$ dont les sorties sont constituées de fil de 0,38 mm de diamètre est testé à 100 MHz. L'impédance obtenue est égale à $48 + j39\ \Omega$. Trouver les valeurs des composants qui constituent (a) un circuit équivalent série RL et (b) un circuit équivalent parallèle RL .

2.2a Concevez un réseau en L qui vous permettra d'adapter un générateur dont la résistance de source est de $50\ \Omega$ à une charge de résistance de $100\ \Omega$ pour une fréquence nominale de 1,5 MHz. On suppose que l'élément parallèle est une bobine. Servez-vous du programme d'analyse de circuit (voir le Problème 1.3 du chapitre 1) pour trouver la réponse en fréquence de ce circuit sur une plage de fréquence comprise entre 1 et 2 MHz par incréments de 50 kHz.

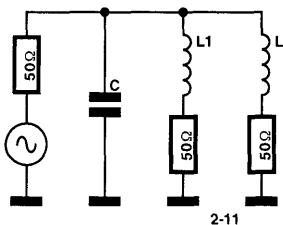
2.2b Faites le même exercice qu'en (a) en supposant que l'élément parallèle est un condensateur.

2.3 Concevez un réseau en double L pour adapter, toujours à la même fréquence, le générateur à la charge du Problème 2.2 (a). Pour obtenir la plus large bande passante possible, prenez comme impédance intermédiaire la moyenne géométrique des impédances de source et de charge (elle est égale à $\sqrt{50 \times 100}$). Servez-vous du programme d'analyse de circuit (voir le problème 3 du chapitre 1) pour trouver la réponse en fréquence de ce circuit (comme pour le Problème 2.2).

2.4 Supposez que les seules bobines que vous ayez à votre disposition pour réaliser les réseaux des Problèmes 2.2 (a) et 2.3 aient un facteur Q_{NC} (facteur Q non chargé) égal à 100 pour une fréquence de 1,5 MHz. Vous supposerez que les condensateurs sont sans perte. Calculez les rendements des réseaux d'adaptation à 1,5 MHz. Vérifiez vos résultats en vous servant du programme d'analyse de circuit (voir le problème 1.3 du chapitre 1).

2.5 Le schéma ci-dessous illustre un réseau qui permet à un générateur de $50\ \Omega$ d'alimenter deux charges (deux antennes par exemple). Le réseau répartit la puissance de telle façon que la charge du haut reçoit une puissance double de celle que reçoit la charge du bas. Le générateur est adapté, c'est-à-dire qu'il voit une charge globale de $50\ \Omega$. Trouvez les valeurs X_{L1} , X_{L2} et X_C .

Astuces : transformez d'abord l'impédance de chaque charge à l'aide d'un réseau en L puis combinez les deux réseaux dans le circuit représenté.



2.6 Vérifiez que dans le cas d'un réseau en L :
 $X_P = \pm R/Q$ et $X_S = \pm rQ$ où $R > r$ et $Q = \sqrt{R/r - 1}$.

2.7 Pour une fréquence donnée, il est possible de modéliser une bobine à pertes sous la forme d'une bobine sans perte en série avec une résistance ou sous la forme d'une bobine sans perte en parallèle avec une résistance. Convertissez la combinaison série r_s, L_s en son équivalent parallèle r_p, L_p et prouvez que Q_{NC} défini comme étant le rapport X_S/r_s est égal à Q_{NC} défini comme étant le rapport r_p/X_p .

3. Amplificateurs linéaires

Nous allons étudier dans ce chapitre les amplificateurs linéaires commandés par une résistance ; leur emploi est varié, amplificateurs à courant continu, amplificateurs basses fréquences, amplificateurs vidéo et amplificateurs hautes fréquences. Nous verrons que si les principes fondamentaux de fonctionnement restent les mêmes, les structures des circuits dépendent de la nature de la charge, de celle du signal et du rendement que l'on désire obtenir. Nous nous concentrerons pour le moment sur la topologie de l'étage de sortie qui est l'« extrémité productrice » de ces amplificateurs. Le mot « linéaire » signifie que le signal de sortie est linéairement proportionnel au signal d'entrée sur une plage continue, ce qui revient à dire que le signal de sortie est la reproduction, à une certaine échelle, du signal d'entrée.

Amplificateur à une seule maille

La figure 3-1 représente le schéma fondamental d'un amplificateur formé d'une alimentation en courant continu, d'un composant actif (transistor ou tube) et d'une charge. La résistance réglable manuellement (rhéostat) que l'on aperçoit dans le schéma de gauche représente un composant actif (c'est le transistor que l'on voit dans le schéma de droite). Il limite ou « laisse passer » le courant, et par là même, commande la chute de tension aux bornes de la résistance de charge R_C . Un composant actif employé comme résistance variable *doit* dissiper de l'énergie (sous forme de chaleur) sauf si la tension de sortie est nulle ou à la tension d'alimentation maximale. On choisit la tension d'alimentation en courant continu égale à la tension maximale nécessaire aux bornes de la charge. Le rendement ne dépend que de la forme d'onde, c'est-à-dire en aucune façon des caractéristiques du transistor. On applique le signal d'entrée ou « signal d'attaque » à la base du transistor pour commander sa résistance. Nous ne mentionnerons pas pour le moment le chemin de retour du signal d'attaque. La puissance que reçoit la charge est supérieure à celle du signal d'attaque (en général très nettement supérieure). Le gain est le rapport entre ces puissances. Lorsqu'on travaille dans le domaine des hautes fréquences, le mot « gain » signifie toujours gain en puissance.

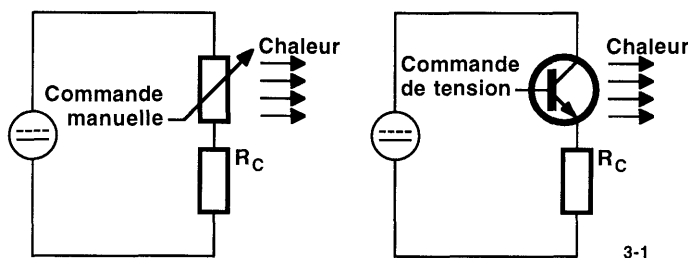


Figure 3-1. Amplificateur fondamental à une seule maille.

Montage émetteur-suiveur

Lorsque le chemin de retour du signal d'attaque se fait par le bas de la résistance de charge, nous obtenons l'amplificateur représenté à la figure 3-2. Nous parlerons dans ce cas d'amplificateur à collecteur commun puisque seul le collecteur se trouve à la masse alternative (le collecteur n'est évidemment pas directement relié à la masse alternative, mais par l'intermédiaire de la source d'alimentation continue). On parle également de montage *émetteur-suiveur* parce que, si la résistance R_c n'est pas trop faible, le transistor régulera son courant pour que la tension instantanée de l'émetteur soit quasiment identique à celle de la base. « Quasiment identique » signifie ici, à un petit décalage près du niveau continu (la tension moyenne de l'émetteur sera inférieure d'environ 0,7 V à celle de la base), et à une légère réduction près de l'amplitude du signal (l'amplitude du signal alternatif présent à l'émetteur sera inférieure d'environ 1% à celle que l'on observe à la base, voir le problème 3.4).

Le montage émetteur-suiveur peut fournir à la charge des tensions positives de façon continue depuis 0 V jusqu'à la tension d'alimentation, mais il est incapable de délivrer des tensions négatives. Le résistor de charge et le transistor constituent un diviseur de tension dans lequel la résistance du transistor commande la chute de tension aux bornes de la charge. Cela ne fonctionne que si la charge est réellement un transistor et non un condensateur, comme par exemple une électrode de déflexion de faisceau au sein d'un tube à rayons cathodiques d'oscilloscope. Malgré ces restrictions, il existe des cas où l'amplificateur à une seule maille convient parfaitement. Prenons l'exemple d'un four régulé électroniquement. La résistance de charge est l'élément chauffant. Le transistor permet de modifier la valeur de la tension aux bornes de la résistance, peu importe la polarité appliquée. De plus, dans cette application, le circuit à émetteur-suiveur sera tout aussi efficace que n'importe quel autre montage de commande de résistance fonctionnant à partir d'une seule alimentation.

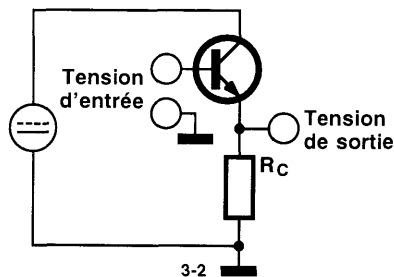


Figure 3-2. Amplificateur à émetteur-suiveur.

Amplificateurs à émetteur commun et à base commune

Avant d'aborder d'autres montages qui permettent de surmonter les limitations de ceux à une seule maille, nous devons mentionner deux variantes courantes. La figure 3-3 représente le schéma d'un amplificateur à émetteur commun. Vous constatez que la charge et le transistor ont changé de place. Puisque l'alimentation en courant continu, le transistor et la charge sont toujours en série, le fonctionnement fondamental et le rendement sont les mêmes que ceux que l'on constate dans un montage à émetteur-suiveur. Mais étant donné que l'émetteur se trouve relié à la masse, la tension d'attaque est directement appliquée à la jonction base-émetteur ; le courant qui traverse le transistor sera alors une fonction non linéaire (exponentielle) de la tension d'attaque. Il serait possible d'adjoindre un circuit externe qui produirait un signal d'attaque inverse (logarithmique) afin de linéariser l'amplificateur. C'est ce qui se passe lorsqu'on utilise une source

de courant à la place d'une source de tension pour attaquer un amplificateur de puissance à collecteur commun. La linéarité de l'ensemble d'amplification dépendra du gain en courant en régime linéaire β qui reste constant sur une grande plage. Le gain en puissance d'un amplificateur à émetteur commun est supérieur à celui d'un amplificateur à émetteur-suiveur.

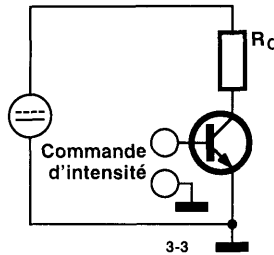


Figure 3-3. Amplificateur à émetteur commun.

La figure 3-4 donne le schéma d'un amplificateur à base commune. Dans ce circuit, le courant d'attaque traverse l'ensemble du circuit. Voici comment s'obtient l'amplification : bien que le circuit d'attaque et la charge soient parcourus par le même courant (il n'y a environ qu'1% de différence entre le courant de base et celui d'émetteur-collecteur), l'excursion de tension au collecteur, limitée par la tension d'alimentation, est bien plus élevée que l'excursion de tension à l'émetteur, limitée par la tension en quasi court-circuit de la jonction émetteur-base.

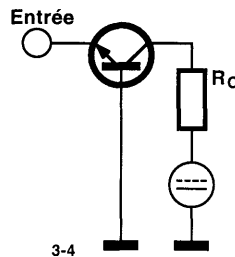


Figure 3-4. Amplificateur à base commune.

Un transistor, deux alimentations

Si l'amplificateur doit pouvoir délivrer des tensions de sortie positives et négatives et amplifier des formes d'onde quelconques, il devra être alimenté par une alimentation double (ou par une alimentation « flottante » comme nous le verrons par la suite). La figure 3-5 représente le schéma d'un circuit à deux mailles qui n'emploie qu'un élément actif. Ici R_E est un résistor d'excursion basse. Elle tire la sortie vers l'alimentation négative lorsque le transistor le permet. Le courant moyen ou *courant de polarisation* qui circule dans R_E dépend de la tension moyenne de sortie. Dans le cas de certains signaux, comme les signaux à basses fréquences, la tension moyenne est nulle. On choisit alors les valeurs de V_{neg} et de R_E de telle sorte que l'excursion de tension négative maximale soit égale à l'excursion de tension positive maximale. Le meilleur rendement est obtenu lorsque $V_{neg} = 2V_{pos}$ et $R_E = R_C$. Vous pouvez vérifier cependant que le rendement maximal est faible : il est seulement égal à 1/12. Les amplificateurs polarisés sont connus sous l'appellation d'amplificateurs fonctionnant en classe A ; on les utilise couramment dans

les applications en petits signaux où la dissipation de puissance est toujours faible et où le rendement n'est pas la première des préoccupations. Cet amplificateur polarisé pourrait attaquer une charge purement capacitive puisque R_E est un chemin de décharge. Le courant maximal négatif qui traverse la charge est toutefois limité par la valeur de R_E . Si l'on réduit la valeur de R_E , on accroît le courant admissible mais aussi celui de polarisation et on affecte le rendement.

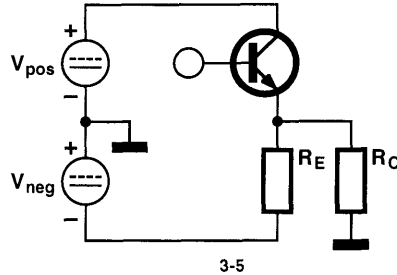


Figure 3-5. Amplificateur à un seul transistor et alimentation double.

Deux transistors, deux alimentations

La configuration symétrique (les anglo-saxons parlent de « *push-pull* ») à deux transistors, illustrée à la figure 3-6, délivre une tension de sortie positive et négative et son rendement est élevé. Ce circuit à deux mailles, et deux transistors, comprend un circuit original à émetteur-suiveur NPN et un circuit complémentaire (PNP) à émetteur-suiveur, au lieu d'une résistance d'excursion basse, pour fournir la tension négative. Le transistor du haut tire la charge vers la tension positive. Celui du bas tire la charge vers la tension négative, ou encore il la repousse de la tension positive. Une éventuelle capacité parasite parallèle présente dans la charge peut être rapidement chargée ou déchargée en surattaquant le transistor opportun jusqu'à obtenir la tension de sortie désirée. Ceci suppose une rétroaction ; l'attaquer doit savoir quand appliquer la surattaque. Même s'il n'y a pas de chaîne de retour, les circuits suiveurs fournissent leur propre rétroaction négative ; si la tension aux bornes de la charge (tension de l'émetteur) ne suit pas immédiatement la tension d'attaque (tension de la base), la tension émetteur-base croît et déclenche la conduction du transistor. Le circuit symétrique est le circuit par excellence pour amplifier des formes d'ondes quelconques ; son rendement est aussi élevé que possible pour un circuit linéaire (il atteint 78% pour un signal sinusoïdal d'amplitude maximale). La vitesse de balayage (*slew rate*) dU/dt lorsqu'on emploie des charges capacitives n'est pas limitée par la valeur des résistances d'excursion basse (ou par la composante résistive de la charge).

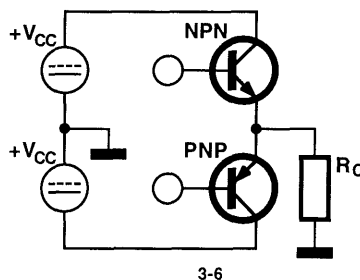


Figure 3-6. Amplificateur symétrique complémentaire (PNP/NPN).

L'amplificateur symétrique fonctionne normalement en classe B, ce qui signifie que lorsque la tension de sortie est nulle, les deux transistors sont bloqués et il n'y a aucune dissipation. On obtient un courant positif lorsque le transistor du haut est actif et un courant négatif lorsque c'est celui du bas qui l'est. Il est essentiel que la transition se fasse en douceur et de façon continue. On utilise parfois les amplificateurs symétriques en classe AB ; dans ce cas, un faible courant de polarisation assure une transition en douceur. Voici un autre avantage du circuit symétrique complémentaire : les signaux d'attaque des deux transistors sont identiques mis à part une petite polarisation en courant continu.

Cet avantage disparaît lorsqu'on emploie deux transistors identiques pour réaliser un amplificateur asymétrique en « *totem pole* » tel qu'on le voit à la figure 3-7. Il faut attaquer dans ce cas les deux transistors de manière déphasée, c'est-à-dire que les signaux d'attaque sur chacune des bases doivent être de polarité opposée. De plus, le transistor du haut est un émetteur-suiveur qui suit sa *tension* d'attaque tandis que le transistor du bas est un émetteur commun qui suit son *courant* de commande. Il est indispensable de disposer, dans ce circuit peu commode, d'attaquiers distincts pour obtenir une réponse linéaire. Nous verrons par la suite que l'utilisation de transformateurs permet de réaliser un amplificateur symétrique avec deux composants identiques (deux transistors NPN, deux transistors PNP ou deux tubes).

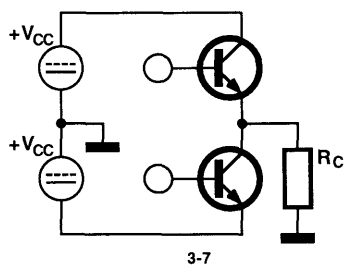


Figure 3-7. Amplificateur symétrique en *totem pole*.

Signalons l'existence d'un amplificateur symétrique qui n'utilise qu'une seule alimentation. La figure 3-8a représente un tel amplificateur dénommé amplificateur en pont. Dans ce circuit à trois mailles, il n'existe aucune connexion directe entre l'alimentation et la charge ; l'alimentation ou la charge doivent être « flottantes », c'est-à-dire n'avoir aucune liaison avec la masse. Deux transistors fonctionnent en interrupteurs. Dans le montage en demi-pont de la figure 3-8b, deux transistors sont remplacés par des résistances ; de même que dans le montage en pont, la charge est alimentée par des courants positif et négatif et il n'est nul besoin qu'un courant de polarisation circule dans la charge. Il faut cependant noter que l'excursion maximale de tension est divisée par deux et que le rendement souffre du fait qu'un courant de polarisation circule dans les résistances du pont.

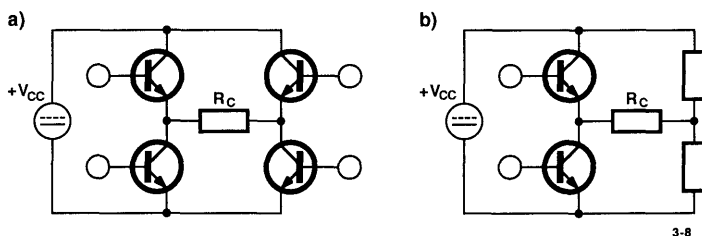


Figure 3-8. Amplificateurs en pont.

Amplificateurs de courant alternatif

S'il est vrai que la configuration symétrique d'un amplificateur est celle des amplificateurs à usages multiples qui a le meilleur rendement, il est possible de concevoir un amplificateur de courant alternatif (signaux à basses fréquences) plus simple, que l'on peut encore simplifier si l'on se limite au domaine des signaux à courant alternatif à bande étroite (c'est le cas des signaux hautes-fréquences).

Amplificateurs basses fréquences

Les signaux à basses fréquences, n'ayant pas de composante continue, sont symétriques par rapport à la masse ; la liaison en courant alternatif permet de se passer d'une alimentation négative. Le schéma situé dans la partie gauche de la figure 3-9 représente une version à une seule alimentation de l'amplificateur en classe A à un transistor que nous avons étudié précédemment. Le rendement n'est alors que de 1/12. Mais si nous remplaçons le résistor d'excursion basse par une bobine d'arrêt (comme on peut le voir dans la partie droite de la figure 3-9), nous éliminons toute dissipation de puissance dans la résistance d'excursion basse. La bobine d'arrêt permet à la sortie de passer de l'état négatif à l'état positif. Nul n'est besoin de placer un condensateur de liaison à courant continu (si la résistance de la bobine d'arrêt en courant continu est faible). Le courant de polarisation qui circule dans la bobine d'arrêt doit avoir une valeur suffisante pour que le transistor soit toujours passant et pour obtenir un fonctionnement linéaire. Vous pouvez calculer que l'efficacité maximale de ce courant atteint 50% ; la bobine d'arrêt améliore le rendement par un facteur égal à six.

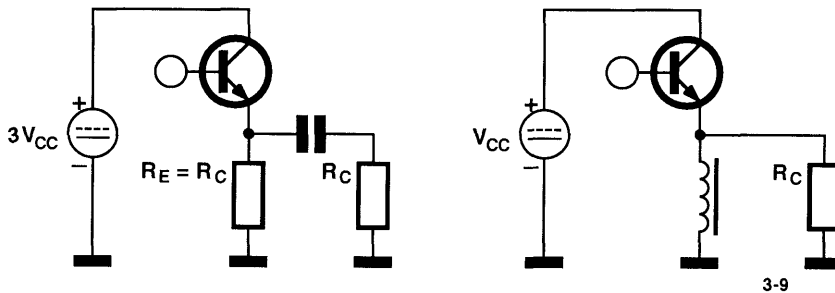


Figure 3-9. Amplificateurs basses fréquences non symétriques (à collecteur commun).

Dans l'application la plus courante de ce circuit (montage à émetteur commun), la charge se trouve au-dessus du transistor comme on peut le constater en examinant la figure 3-10. La variante que vous apercevez à droite montre qu'il est nécessaire d'adjoindre un condensateur de liaison si l'une des extrémités de la charge se trouve à la masse. Ce schéma est celui que l'on rencontre le plus fréquemment pour les amplificateurs de petits signaux ; lorsqu'il n'est pas indispensable d'obtenir un rendement élevé ou de vastes excursions de tension, un résistor remplace souvent la bobine d'arrêt. Comme nous l'avons vu, un courant d'attaque permet d'obtenir une réponse linéaire si l'émetteur est directement relié à la masse. Il est possible de relier l'émetteur à la masse via une résistance, aux dépens d'une perte d'énergie, pour permettre à la tension de l'émetteur de suivre la tension d'attaque (celle de la base). Le courant d'émetteur, le courant de collecteur et la tension de sortie sont alors linéairement proportionnels à la tension d'entrée. On connaît cette méthode de linéarisation d'un amplificateur à émetteur commun sous le nom de « contre-réaction d'émetteur » ou « contre-réaction série ».

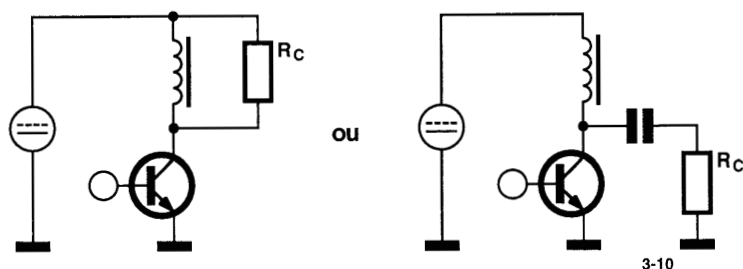


Figure 3-10. Amplificateurs basses fréquences non symétriques (à émetteur commun).

Il arrive fréquemment que l'impédance d'une charge donnée ne soit pas adaptée au transfert de la puissance maximale ; on remplace dans ce cas-là la bobine d'arrêt et le condensateur de liaison par un transformateur. C'est ce que représente la figure 3-11.

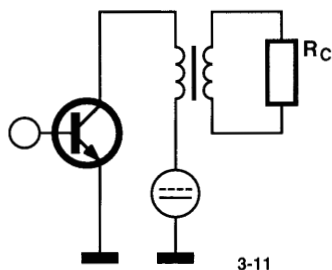


Figure 3-11. Amplificateur basses fréquences non symétriques (à couplage par transformateur).

Dans ces amplificateurs en classe A à liaison par bobine d'arrêt ou par transformateur, l'excursion de tension crête-crête au collecteur est le double de la tension d'alimentation V_{cc} . Notez que l'utilisation du transformateur le plus simple qui soit (le « transformateur idéal ») ne permettrait d'obtenir qu'une excursion de tension de crête à crête égale à seulement V_{cc} puisque R_c et le transformateur seraient remplacés par la seule résistance transformée. S'il est nécessaire de mettre en œuvre un transformateur pour réaliser une transformation d'impédance, on peut ajouter un second composant actif identique et un transformateur à point milieu pour aboutir à l'amplificateur symétrique représenté à la figure 3-12. Il est possible de faire fonctionner en classe B le montage symétrique représenté à la figure 3-12, de même que tous les montages symétriques qui ne font pas appel à un transformateur, pour obtenir un meilleur rendement. Les amplificateurs haute fidélité à tube emploient ce style de montage. Les amplificateurs haute fidélité à transistors ne font pas appel aux transformateurs et préfèrent une combinaison symétrique complémentaire. Ils fonctionnent avec une alimentation unique en utilisant un condensateur de liaison vers le haut-parleur (remplacez les résistances du pont dans le schéma de la figure 3-8b par des condensateurs).

REMARQUE CONCERNANT LE RENDEMENT Il semblerait que le rendement maximal d'un amplificateur en classe B (78%) ne soit que très légèrement supérieur à celui d'un amplificateur en classe A (50%). Ceci n'est vrai que lorsque l'amplificateur délivre un signal sinusoïdal d'amplitude maximale. En ce qui concerne la parole et la musique, la puissance moyenne est bien inférieure à la puissance maximale. Un amplificateur en classe B dissipe peu lorsque l'amplitude du signal est faible, alors qu'un amplificateur en

classe A absorbe une puissance égale au double de la puissance maximale de sortie : cela est dû au courant permanent de polarisation. Un amplificateur en classe A de puissance de sortie nominale 25 W consommera en permanence une puissance de 50 W fournie par l'alimentation, tandis qu'un amplificateur en classe B d'une puissance de sortie nominale équivalente ne consommera, en moyenne, que quelques watts.

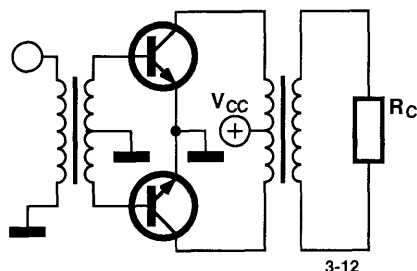


Figure 3-12. Amplificateur symétrique (à couplage par transformateur).

Amplificateurs hautes fréquences

Tout ce que nous avons dit des amplificateurs basses fréquences à courant alternatif s'applique également aux amplificateurs hautes fréquences. Mais du fait que les amplificateurs hautes fréquences sont des amplificateurs à courant alternatif à bande étroite, quelques astuces s'appliquent. Les amplificateurs hautes fréquences de petits signaux, de même que les amplificateurs hautes fréquences de puissance, ressemblent souvent à l'amplificateur basses fréquences en classe A avec une bobine d'arrêt insérée dans le collecteur, tel que le montre la figure 3-13.

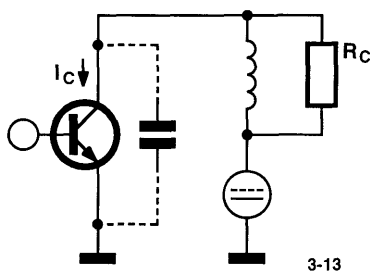


Figure 3-13. Amplificateur hautes fréquences non symétrique.

Aux hautes fréquences la capacité parallèle de sortie du transistor (représentée en traits pointillés) pourrait effectivement court-circuiter la sortie de l'amplificateur. Mais sur une plage de fréquences étroite il est possible d'« annuler » cette capacité parasite à l'aide d'une inductance parallèle accordée. L'alimentation en courant continu se comporte comme un court-circuit en courant alternatif, de sorte que la bobine d'inductance se trouve bien en parallèle avec la capacité de sortie. On peut procéder de la même façon avec la capacité parasite d'entrée. Cet amplificateur hautes fréquences, bien qu'il ait la même structure qu'un amplificateur basses fréquences en classe A non symétrique, fonctionne comme un efficace amplificateur de puissance en classe B grâce à son circuit résonant d'accord. N'oubliez pas que l'architecture des amplificateurs basses fréquences en classe B et des amplificateurs à courant continu en classe B doit être du type symétrique « équilibré ». Ici l'effet de volant¹ du circuit résonant rend possible le fonctionnement

en classe B du circuit non symétrique. Pour en comprendre le fonctionnement, examinez à la figure 3-14 les formes d'onde de la tension et du courant collecteur (niveau maximal du signal). La tension du collecteur décroît au cours de la seconde alternance. Le transistor tire vers le bas la tension (I_c est positif) en absorbant de l'énergie de l'alimentation. Mais la tension croît au cours de la première alternance. Un courant négatif est nécessaire mais il n'y a besoin d'aucun circuit complémentaire pour exécuter la poussée. Au lieu de cela, l'effet de volant du circuit résonant déplace la tension lors des alternances positives. C'est comme lorsqu'on provoque une oscillation : les poussées que vous donnez peuvent être toujours dans le même sens ; l'effet de balancier fait qu'il se produit une oscillation opposée au cours de la seconde alternance. Il faut naturellement que vous exerciez pour cela une poussée double. L'effet de volant doit être assez élevé pour maintenir la forme d'onde sinusoïdale voulue au cours des alternances positives lorsque le transistor est bloqué, c'est-à-dire que le facteur Q du circuit résonant chargé ne doit pas être trop faible. Pour stocker l'énergie nécessaire à l'obtention d'une valeur satisfaisante du facteur Q , on ajoute une capacité en parallèle avec la capacité de sortie du transistor. Il faut bien entendu abaisser en conséquence la valeur de la bobine d'inductance afin que le circuit résonne toujours à la fréquence nominale.

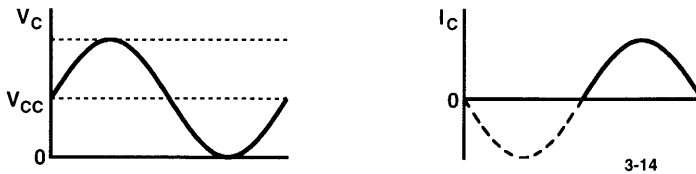


Figure 3-14. Tension et courant de collecteur dans un amplificateur hautes fréquences non symétrique en classe B.

Note sur l'adaptation d'un amplificateur de puissance à sa charge

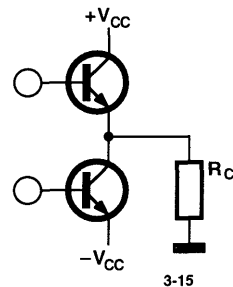
L'examen du circuit d'un amplificateur de puissance montre que la tension de crête de sortie peut atteindre des valeurs aussi élevées que la tension d'alimentation V_{cc} . La puissance maximale P_{max} dissipée dans une charge R est égale à $V_{cc}^2/(2R)$. Pour une puissance maximale P_{max} et une charge R données, la valeur de la tension d'alimentation s'obtient grâce à la formule $V_{cc} = \sqrt{2RP_{max}}$. Si l'amplificateur doit fonctionner en courant continu, cette valeur de V_{cc} permet à l'amplificateur de fournir la puissance maximale désirée. La valeur de V_{cc} peut naturellement être plus élevée, mais le rendement devient plus faible – c'est ce qu'illustre le problème 3.1. Dans le cas d'une amplification en courant alternatif, il est toutefois possible d'adapter l'impédance de charge donnée à la valeur désirée $V_{cc}^2/(2P_{max})$ à l'aide d'un transformateur (pour les basses et hautes fréquences) ou d'un réseau d'adaptation accordé (uniquement pour les hautes fréquences). Vous noterez que le transformateur ou le réseau d'adaptation transforme l'impédance de charge donnée vers la valeur désirée $V_{cc}^2/(2P_{max})$ et non en une impédance quelconque qui soit caractéristique du transistor ou du tube. En ce qui concerne les amplificateurs de puissance, on choisit le tube ou le transistor en fonction de son aptitude à résister aux tensions et aux courants appliqués (il faut également qu'il fonctionne correctement à la fréquence maximale requise). Lorsqu'on travaille sur de petits signaux, les impédances du composant actif n'ont aucune incidence sur la conception du réseau d'adaptation de sortie de l'amplificateur de puissance.

1 Lorsque le courant est sinusoïdal, l'énergie disponible dans un circuit LC parallèle (ou série) se propage de façon cyclique entre la bobine et le condensateur. C'est ce qui se passe également avec l'énergie mécanique emmagasinée dans le mouvement cyclique d'un volant en rotation.

Problèmes

3.1a Analysez l'amplificateur en classe B représenté ci-dessous et montrez que le rendement est égal à $\pi/4$ lorsque l'amplificateur attaque une résistance de charge avec une onde sinusoïdale dont l'amplitude de crête à crête est $2V_{CC}$ (onde sinusoïdale à pleine puissance).

Astuces : souvenez-vous que dans un montage fonctionnant en classe B, le transistor du bas est bloqué lorsque la tension de sortie est positive alors que le transistor du haut est bloqué lorsque la tension de sortie est négative. Le fonctionnement étant symétrique, il suffit de ne considérer qu'une alternance, lorsque le transistor du haut est passant. Trouvez, pour cette alternance, la puissance moyenne dissipée dans la charge (la moyenne de IUC) ainsi que la puissance moyenne fournie par l'alimentation (la moyenne de IV_{CC}).



3.1b Que devient le rendement lorsque l'amplitude de l'onde sinusoïdale n'est qu'une fraction α de l'amplitude maximale ?

3.1c Quel est le rendement lorsque le signal de sortie a une forme d'onde carrée d'amplitude maximale ? (Réponse : 100%)

3.2 L'amplificateur en classe A représenté ci-dessous fonctionne à la puissance maximale et délivre à la résistance de charge une onde sinusoïdale d'amplitude égale à 24 V de crête à crête. Supposons que la résistance en courant continu de la bobine d'arrêt soit nulle, que l'inductance soit assez élevée pour bloquer tout courant alternatif et que la susceptance du condensateur soit assez élevée pour éviter toute chute de tension alternative.

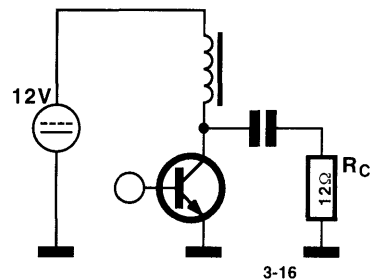
3.2a Représentez la forme d'onde de la tension du collecteur. *Astuces :* n'oubliez pas qu'il n'y a aucune chute de tension continue aux bornes de la bobine d'arrêt.

3.2b Représentez la forme d'onde du courant du collecteur. *Astuces :* n'oubliez pas qu'aucun courant alternatif ne circule dans la bobine d'arrêt.

3.2c Quelle est la puissance fournie par l'alimentation lorsque l'amplitude de l'onde sinusoïdale est maximale ?

3.2d Montrez que dans ces conditions le rendement n'est que de 50%.

3.2e Quelle est la puissance fournie par l'alimentation lorsque l'amplitude de l'onde sinusoïdale est nulle ?



3.3 Un amplificateur symétrique parfait ne connaît pas la limitation du courant d'un amplificateur à émetteur-suiveur ; il peut donc attaquer des charges capacitatives aux hautes fréquences. Mais qu'en est-il des charges inductives ? Supposons qu'une charge ait une inductance série inévitable et qu'il faille produire aux bornes de sa partie résistive des tensions élevées. Comment cela va-t-il se traduire en ce qui concerne la conception de l'amplificateur ?

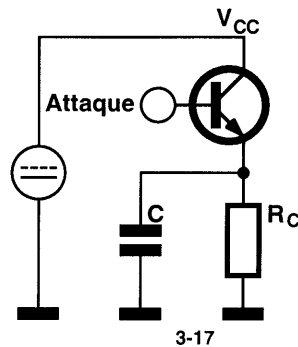
3.4 Justifiez ce qui a été précédemment dit au sujet du gain en tension (de l'ordre de 99%) et de la tension de décalage (environ 0,7 V) de l'amplificateur à émetteur-suiveur. Servez-vous de la relation existant entre le courant d'émetteur et la tension base-émetteur d'un transistor (bipolaire) : $I \approx I_s \exp [(U_b - U_e)/0,026]$. Pour donner une valeur à I_s , supposons que $I = 10$ mA lorsque $U_b - U_e = 0,7$ V. N'oubliez pas que dans un amplificateur à émetteur-suiveur $U_e = I_e R$. Donnez une valeur raisonnable à R , comme par exemple 1000Ω , et trouvez U_e pour différentes valeurs de U_b .

3.5 Trouvez le gain en puissance d'un amplificateur à émetteur-suiveur. Servez-vous du fait que le courant d'entrée (courant de base) est égal au courant d'émetteur multiplié par le facteur $1/(\beta + 1)$ où β est le gain en courant du transistor (une valeur typique est de l'ordre de 100). Souvenez-vous que la tension de sortie est quasiment la même que la tension d'entrée (elle la suit).

3.6 L'amplificateur à émetteur-suiveur représenté ci-dessous attaque une charge qui comprend une capacité parallèle inévitable.

3.6a Quelle est la tension de crête à crête maximale qui peut être fournie à la charge aux basses fréquences (on peut dans ce cas faire abstraction du condensateur) ?

3.6b À quelle fréquence un signal de sortie de forme d'onde sinusoïdale d'amplitude moitié de l'amplitude maximale commencera-t-il à subir une distorsion ? *Astuces : exprimer le courant d'émetteur comme étant la somme du courant qui traverse la résistance et de celui qui traverse le condensateur et notez qu'une distorsion apparaît si ce courant devient négatif (le transistor ne peut fournir qu'un courant positif). (Réponse : $\omega = \sqrt{3}/20(RC)$).*



4. Filtres 1

Dans les montages radio, les filtres passe-bande sont des éléments fondamentaux qui déterminent, par exemple, la sélectivité des récepteurs radio. Nous traiterons ici des filtres à constantes localisées réalisés avec des bobines et des condensateurs. [N.d.T. : un circuit à constantes localisées (les anglo-saxons parlent de *lumped circuit*) est un circuit à courant alternatif composé physiquement de condensateurs et de bobines.] Nous étudierons tout d'abord les filtres passe-bas que nous ferons évoluer vers les filtres passe-bande. Nous commencerons par des filtres passe-bas bien connus – Butterworth, Tchebychev, Bessel et autres. Ces filtres passe-bas sont de simples réseaux en échelle LC constitués de bobines en série et de condensateurs en parallèle (voir la figure 4-1).

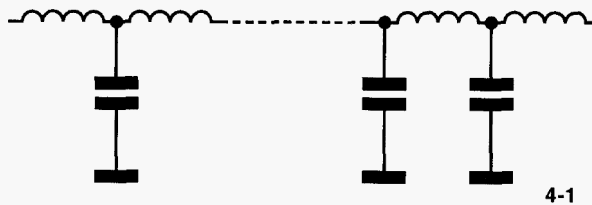


Figure 4-1. Réseau en échelle passe-bas.

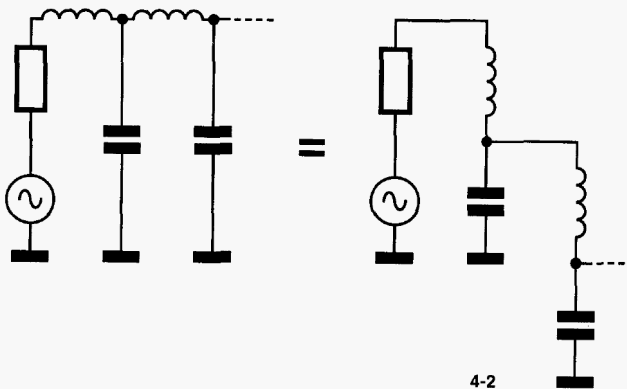


Figure 4-2. Réseau en échelle représenté sous la forme d'une cascade de diviseurs de tension.

Un filtre passe-bas à n sections comporte n ensemble de composants (condensateurs plus bobines). Les composants d'extrémité sont soit des bobines en série, comme le représente la figure 4-1, soit des condensateurs en parallèle. Ces filtres ne comportant (en théorie) aucune résistance, sont des filtres à réflexion ; en dehors de leur bande passante, aucune énergie n'atteindra la charge. On peut redessiner le réseau en échelle sous la forme d'une chaîne de diviseurs de tension comme le montre la figure 4-2. Aux hautes fréquences, le rapport de division croît et la charge est de plus en plus isolée de la source. Pour les fréquences bien inférieures à la fréquence de coupure, chaque cellule du circuit apporte une atténuation (en

puissance) de 6 dB par octave (ou 20 dB par décade). Un filtre passe-bas parfait procure une parfaite adaptation entre charge et source à l'intérieur de la bande passante. Les filtres qui comportent de nombreuses sections se rapprochent du modèle idéal. Il est théoriquement possible, lorsque les impédances de source et de charge ne comprennent aucune réactance (interne ou parasite), d'obtenir une adaptation parfaite sur la totalité de la bande passante.

Filtres passe-bas normalisés

Un filtre Butterworth est extrêmement plat, c'est-à-dire qu'à la fréquence nulle, les dérivées premières d'ordre $2n - 1$ sont nulles en ce qui concerne la fréquence de la fonction de transfert de la puissance. Cette dernière condition (permettant de calculer les valeurs des n éléments) définit la fréquence de coupure f_0 , souvent dénommée fréquence à -3 dB ou fréquence à mi-puissance. La réponse en fréquence d'un filtre Butterworth s'exprime par :

$$\left| \frac{U_{\text{sortie}}}{U_{\text{entrée}}} \right|^2 = \frac{1}{1 + (f/f_0)^{2n}} \quad (4.1)$$

Le filtre Butterworth, s'il offre la réponse la plus plate, n'a pas de flancs aussi raides que le filtre Tchebychev. Cela se traduit par le fait que le filtre Tchebychev présente une certaine ondulation dans la bande passante. Par conception, les ondulations d'un filtre Tchebychev ont toutes la même amplitude. La réponse d'un tel filtre est donnée par la formule suivante :

$$\left| \frac{U_{\text{sortie}}}{U_{\text{entrée}}} \right|^2 = \frac{1}{1 + (U_r^{-2} - 1) \cosh^2 [n \cosh^{-1}(f/f_0)]} \quad (4.2)$$

où U_r exprime le niveau (exprimé en volts) auquel se situe l'ondulation par rapport à la ligne de base.

Vous trouverez dans de nombreux ouvrages des tableaux fournissant les valeurs des composants du filtre. Vous découvrirez dans l'annexe A4-1 située à la fin de ce chapitre de tels tableaux, proposés par Matthaei [1]. Ces tableaux concernent des filtres normalisés, c'est-à-dire que leur fréquence de coupure¹ est de 1 rad/s (1/2π Hz). La valeur du $n^{\text{ième}}$ composant est de g_n F ou g_n H selon que le premier élément du filtre est un condensateur ou une bobine. L'impédance propre de la source est de $1 + j0 \Omega$. C'est également celle de la charge sauf dans la cas des filtres Tchebychev d'ordre pair où elle est égale à $1/g_{n+1} + j0 \Omega$. La figure 4-3 représente le tracé des réponses d'un filtre Butterworth et de plusieurs filtres Tchebychev.

1 La fréquence de coupure d'un filtre Butterworth correspond à la fréquence à demi-puissance (-3 dB). Pour un filtre Tchebychev de n dB, c'est la fréquence la plus élevée pour laquelle l'atténuation est de n dB (voir la figure 4-3).

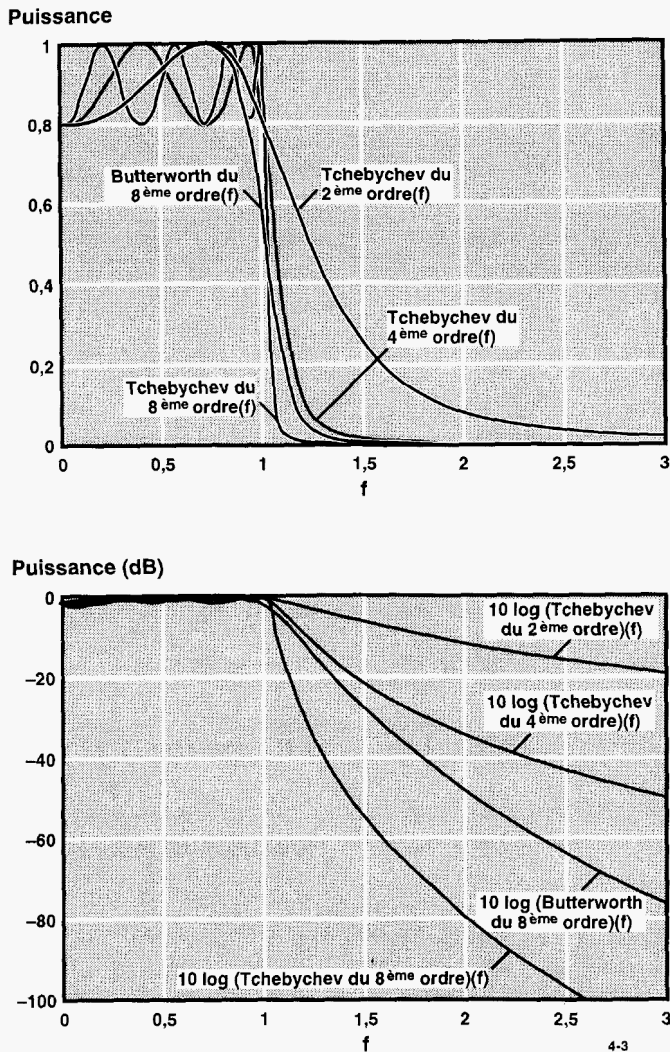


Figure 4-3. Courbes de réponse des filtres Butterworth et Tchebychev.

Exemple de filtre passe-bas

Nous prendrons comme exemple un filtre passe-bas Butterworth à trois éléments. Le tableau A4-1 indique que les valeurs des composants du filtre sont de 1 H, 2 F et 1 H (figure 4-4(a)) ou 1 F, 2 H et 1 F (figure 4-4(b)). Le tableau 4-1 fournit les réponses (identiques) de ces filtres dont le tracé est représenté à la figure 4-5. Vous remarquerez que les résultats sont ceux que nous attendions ; le point à -3 dB se situe à 0,159 Hz.

Supposons que nous ayons besoin d'un filtre Butterworth à trois éléments, de 5 kHz de largeur de bande, inséré entre un générateur d'impédance égale à 50Ω et une charge de 50Ω . Nous allons facilement trouver les valeurs des éléments en mettant à l'échelle le filtre normalisé. Il suffit de multiplier par 50 les

valeurs des bobines (soit 50 fois la valeur de la réactance) puis de diviser le résultat par $2\pi \times 5000$ (puisque la valeur de cette réactance doit être obtenue pour 5 kHz, et non pour 1 rad/s). On divise de même les valeurs des condensateurs par 50 puis par $2\pi \times 5000$. La figure 4-6 représente le schéma résultant de la mise à l'échelle du filtre de la figure 4-4(b). Le tableau 4-2 donne la réponse du filtre mis à l'échelle et la figure 4-7 illustre le tracé correspondant.

Tableau 4-1. Réponse en fréquence des filtres de la figure 4-4.		
Fréquence (Hz)	Puissance	dB
0,00	1,000	-0,0
0,0321	0,000	-0,0
0,0640	0,996	-0,02
0,095	0,955	-0,20
0,1270	0,792	-1,01
0,1590	0,500	-3,01
0,1910	0,251	-6,00
0,2230	0,117	-9,31
0,2540	0,056	-12,5
0,2860	0,029	-15,4

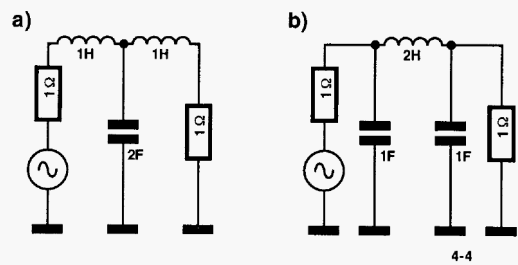


Figure 4-4. Filtres Butterworth (à trois composants) équivalents.

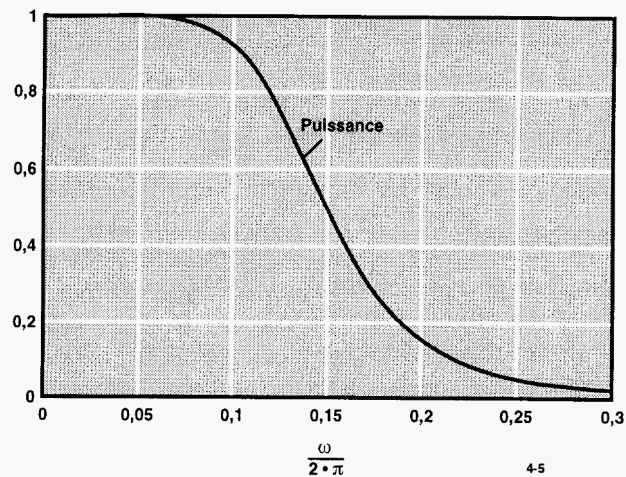


Figure 4-5. Courbe de réponse des filtres de la figure 4-4.

Tableau 4-2. Réponse du filtre passe-bas mis à l'échelle.		
Fréquence (kHz)	Puissance	dB
0,00	1,000	-0,0
1,000	1,000	-0,0
2,000	0,996	-0,02
3,000	0,956	-0,20
4,000	0,793	-1,01
5,000	0,500	-3,01
6,000	0,251	-6,00
7,000	0,117	-9,31
8,000	0,056	-12,5
9,000	0,029	-15,4
10,000	0,015	-18,1

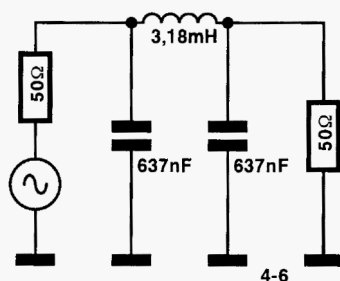


Figure 4-6. Filtre correspondant au schéma de la figure 4-4(b) mis à l'échelle pour des résistances de source et de charge de $50\ \Omega$ et une fréquence de coupure de 5 kHz.

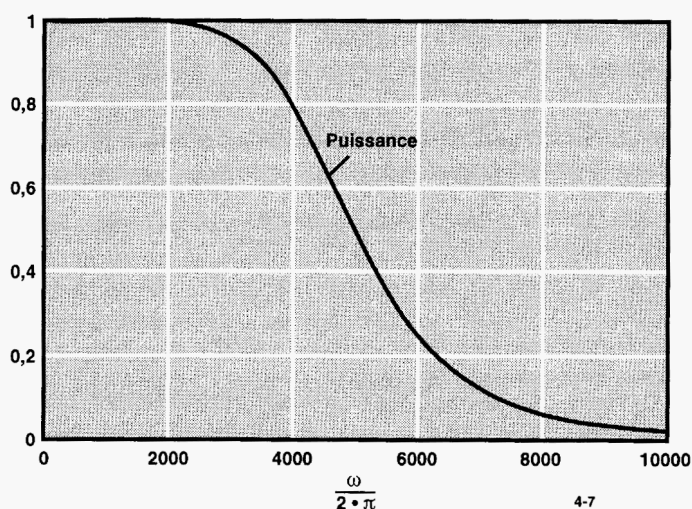


Figure 4-7. Tracé correspondant aux résultats du tableau 4-2.

Évolution vers le filtre passe-bande

Nous allons examiner la façon de convertir un filtre passe-bas en un filtre passe-bande. Souvenez-vous du fonctionnement d'un filtre passe-bas : lorsque la fréquence augmente, la réactance des branches en série (les bobines, qui sont des court-circuits en courant continu) commence à croître. De même, la susceptance des branches en parallèle (les condensateurs, qui sont des circuits ouverts en courant continu) commence également à croître. Comme nous l'avons vu, ces deux effets entravent la transmission d'un signal. Pour passer directement d'un filtre passe-bas à un filtre passe-bande, nous pouvons remplacer les bobines par des combinaisons LC série et les condensateurs par des combinaisons LC parallèle. Les combinaisons en série sont conçues pour résonner (leur impédance est alors nulle) à la fréquence centrale du filtre passe-bande désiré, de même que l'impédance des bobines est nulle en courant continu qui est la « fréquence centrale » d'un filtre passe-bas normalisé.

Il est important de noter que lorsqu'on s'éloigne de la fréquence de résonance, la réactance d'une branche en série LC croît deux fois plus vite que celle d'une simple bobine. C'est facile à voir. La formule suivante donne la réactance d'une branche en série :

$$X_{\text{série}} = \omega L - 1/(\omega C) \quad (4.3)$$

à ω_0 , $X = 0$, donc

$$\omega_0 L = 1/(\omega_0 C) \quad (4.4)$$

et

$$dX/d\omega = L + 1/r(\omega_0^2 C) = 2L \quad (4.5)$$

Lorsque nous nous éloignons de la résonance, la bobine *et* le condensateur participent tous deux à la réactance de l'ensemble. De même la susceptance des circuits en parallèle LC qui remplacent les condensateurs du filtre passe-bas normalisé, croît deux fois plus vite que celle de leurs condensateurs. Ceci étant mémorisé, transformons notre filtre passe-bas à 5 kHz en un filtre passe-bande. Supposons que nous voulions une fréquence centrale de 500 kHz et une bande passante de 10 kHz. Au fur et à mesure que l'on s'éloigne de la fréquence centrale, la réactance des branches en série doit croître au même rythme que celle des bobines du filtre passe-bas normalisé. De même, la susceptance des branches en parallèle doit croître au même rythme que celle des condensateurs du filtre passe-bas normalisé. La courbe de réponse du filtre passe-bande aura la même forme au-dessus de la fréquence centrale que celle du filtre normalisé au-dessus du continu. Si la fréquence du filtre normalisé est de 5 kHz à -3 dB, la fréquence supérieure du filtre passe-bande à -3 dB sera 5 kHz au-dessus de la fréquence centrale. Le filtre passe-bande présentera cependant une réponse miroir lorsqu'on s'éloigne de la fréquence centrale vers le bas. En-dessous de la fréquence centrale, les réactances et les susceptances changent de signe mais la réponse reste la même.

Calculons les valeurs des composants. Lorsque l'on quitte la fréquence centrale, les réactances de L et C des circuits en série varient, comme nous l'avons vu, de la même façon. La valeur de la bobine en série est donc exactement la moitié de ce qu'elle est dans un filtre passe-bas normalisé. *Remarque* : peu importe la valeur de la fréquence centrale, les bobines ont une valeur moitié de celle qu'elles ont dans un filtre passe-bas normalisé. On choisit les valeurs des condensateurs en série pour qu'ils résonnent à la fréquence centrale avec les nouvelles valeurs (moitié) des bobines en série. On détermine de la même façon les valeurs des composants des branches en parallèle ; les condensateurs ont une valeur moitié de celle qu'ils ont dans un filtre passe-bas normalisé. On choisit enfin les valeurs des bobines en parallèle pour qu'elles résonnent à la fréquence centrale avec les nouvelles valeurs (moitié) des condensateurs en parallèle. Ces conversions simples aboutissent au filtre passe-bande de la figure 4-8. Le tableau 4-3 donne la réponse de ce filtre et la figure 4-9 le tracé de celle-ci.

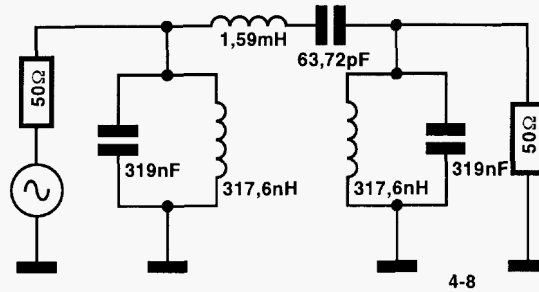


Figure 4-8. Filtre passe-bande.

Bien que ce filtre théorique fonctionne parfaitement bien (puisque les composants sont sans perte), les valeurs des composants sont irréalistes ; des composants réels ayant ces valeurs auraient trop de pertes pour permettre d'obtenir une réponse correcte correspondant au filtre ainsi calculé. Cette conversion directe d'un filtre passe-bas en filtre passe-bande peut être tout à fait satisfaisante si le filtre passe-bande doit avoir une vaste largeur de bande fractionnaire (largeur de bande divisée par la fréquence centrale). C'est lorsque la largeur de bande fractionnaire est faible, comme dans cet exemple, que la conversion directe pose un problème². Nous verrons par la suite qu'il est possible de le résoudre en transformant le filtre passe-bas normalisé en un filtre passe-bande plus complexe connu sous le nom de *filtre à résonateur couplé*. Ces filtres conservent la forme de la réponse désirée (Butterworth, Tchebychev, etc...) ; ils peuvent à leur tour servir de filtres normalisés à des filtres réalisés à partir de résonateurs à quartz ou céramiques et à des filtres utilisant des diaphragmes iris résonnants (minces plaques métalliques qui obstruent partiellement un guide d'onde).

Bibliographie

- 1 G.L. Matthaei, L. Young, E.M.T. Jones (1964), *Microwave Filters, Impedance-Matching Networks and Coupling Structures*. New York: McGraw Hill (réimprimé par Artech House, Boston). (Cet ouvrage comporte de nombreuses réalisations et compare les résultats théoriques aux mesures obtenues dans la pratique ; il traite même des hyperfréquences. La présentation et les rappels théoriques sont excellents.)
- 2 D.G. Fink (1975), *Electronic Engineers' Handbook*. New York: McGraw Hill (chapitre 12). (Ce chapitre, « Filters, coupling networks, and attenuators », de M. Dishal, offre une vaste liste de références.)

2 Le problème qui se pose lors de la simple conversion d'un filtre passe-bas en filtre passe-bande réside dans le fait que les valeurs des bobines en série sont très différentes de celles des bobines en parallèle. C'est d'ailleurs également vrai en ce qui concerne les condensateurs, mais il est facile de se procurer des condensateurs ayant un facteur Q élevé. Dans l'exemple ci-dessus, il existe un facteur de 5000 entre les valeurs des bobines et il est quasiment impossible de trouver des composants ayant un facteur Q élevé dans cette gamme de valeurs. L'emploi de bobines à faible facteur Q se traduit naturellement par la réalisation d'un filtre à pertes et, si l'on n'y prête pas attention, la forme de la courbe de réponse est distordue. Les bobines présentes dans un filtre à résonateur couplé ont pratiquement toutes la même valeur. À condition de se procurer une bobine avec un facteur Q élevé, le filtre à résonateur couplé est conçu pour n'importe quelle valeur d'impédance de la bobine ; on utilise alors des transformateurs ou des sections d'adaptation à chaque extrémité pour obtenir l'impédance voulue.

Tableau 4-3. Réponse du filtre
passe-bande de la figure 4-8.

Fréquence (kHz)	Puissance	dB
490	0,014	-18,1
492	0,053	-12,8
494	0,241	-6,19
496	0,785	-1,05
498	0,996	-0,018
500	1,000	-0,00
502	0,996	-0,018
504	0,801	-0,966
506	0,260	-5,84
508	0,059	-12,9

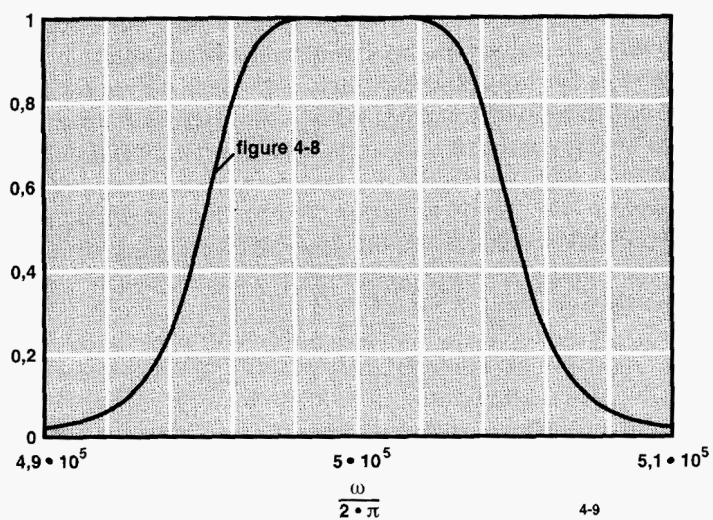


Figure 4-9. Courbe obtenue à partir des données du tableau 4-3.

Annexe 4.1

Valeurs des composants pour un filtre passe-bas normalisé (d'après Matthaei [Bibliographie [1]]).

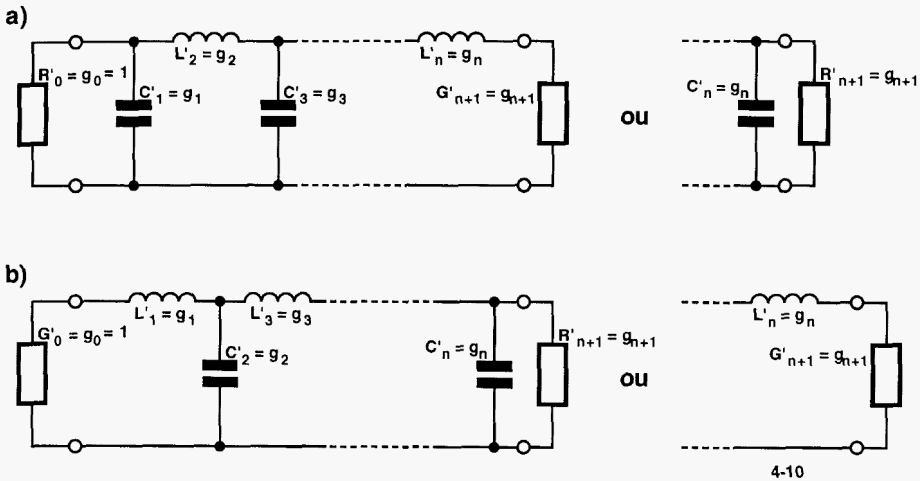


Tableau A4-1. Valeurs des composants pour un filtre passe-bas Butterworth (extrêmement plat).
Le point à -3 dB est obtenu pour $w = 1$ rad/s.

n	g_1	g_2	g_3	g_4	g_5	g_6	g_7	g_8	g_9	g_{10}	g_{11}
1	2,000	1,000									
2	1,414	1,414	1,000								
3	1,000	2,000	1,000	1,000							
4	0,7654	1,848	1,848	0,7654	1,000						
5	0,6180	1,618	2,000	1,618	0,6180	1,000					
6	0,5176	1,414	1,932	1,932	1,414	0,5176	1,000				
7	0,4450	1,247	1,802	2,000	1,802	1,247	0,4450	1,000			
8	0,3902	1,111	1,663	1,962	1,962	1,663	1,111	0,3902	1,000		
9	0,3473	1,000	1,532	1,879	2,000	1,879	1,532	1,000	0,3473	1,000	
10	0,3129	0,9080	1,414	1,782	1,975	1,975	1,782	1,414	0,9080	0,3129	1,000

Tableau A4-2. Valeurs des composants pour un filtre passe-bas Tchebychev (pour un filtre avec une ondulation de n dB, le dernier point à n dB est obtenu pour $\omega = 1$ rd/s).

n	g_1	g_2	g_3	g_4	g_5	g_6	g_7	g_8	g_9	g_{10}	g_{11}
0,01 dB ondulation											
1	0,0960	1,0000									
2	0,4488	0,4077	1,1007								
3	0,6291	0,9702	0,6291	1,0000							
4	0,7128	1,2003	1,3212	0,6476	1,1007						
5	0,7563	1,3049	1,5773	1,3049	0,7563	1,0000					
6	0,7813	1,3600	1,6896	1,5350	1,4970	0,7098	1,1007				
7	0,7969	1,3924	1,7481	1,6331	1,7481	1,3924	0,7969	1,0000			
8	0,8072	1,4130	1,7824	1,6833	1,8529	1,6193	1,5554	0,7333	1,1007		
9	0,8144	1,4270	1,8043	1,7125	1,9057	1,7125	1,8043	1,4270	0,8144	1,0000	
10	0,8196	1,4369	1,8192	1,7311	1,9362	1,7590	1,9055	1,6527	1,5817	0,7446	1,1007
0,1 dB ondulation											
1	0,3052	1,0000									
2	0,8430	0,6220	1,3554								
3	1,0315	1,1474	1,0315	1,0000							
4	1,1088	1,3061	1,7703	0,8180	1,3554						
5	1,1468	1,3712	1,9750	1,3712	1,1468	1,0000					
6	1,1681	1,4039	2,0562	1,5170	1,9029	0,8618	1,3554				
7	1,1811	1,4228	2,0966	1,5733	2,0966	1,4228	1,1811	1,0000			
8	1,1897	1,4346	2,1199	1,6010	2,1699	1,5640	1,9444	0,8778	1,3554		
9	1,1956	1,4425	2,1345	1,6167	2,2053	1,6167	2,1345	1,4425	1,1956	1,0000	
10	1,1999	1,4481	2,1444	1,6265	2,2253	1,6418	2,2046	1,5821	1,9628	0,8853	1,3554
0,2 dB ondulation											
1	0,4342	1,0000									
2	1,0378	0,6745	1,5386								
3	1,2275	1,1525	1,2275	1,0000							
4	1,3028	1,2844	1,9761	0,8468	1,5386						
5	1,3394	1,3370	2,1660	1,3370	1,3394	1,0000					
6	1,3598	1,3632	2,2394	1,4555	2,0974	0,8838	1,5386				
7	1,3722	1,3781	2,2756	1,5001	2,2756	1,3781	1,3722	1,0000			
8	1,3804	1,3875	2,2963	1,5217	2,3413	1,4925	2,1349	0,8972	1,5386		
9	1,3860	1,3938	2,3093	1,5340	2,3728	1,5340	2,3093	1,3938	1,3860	1,0000	
10	1,3901	1,3983	2,3181	1,5417	2,3904	1,5536	2,3720	1,5066	2,1514	0,9034	1,5386

Tableau A4-2 (suite). Valeurs des composants pour un filtre passe-bas Tchebychev (pour un filtre avec une ondulation de n dB, le dernier point à n dB est obtenu pour $\omega = 1$ rd/s).

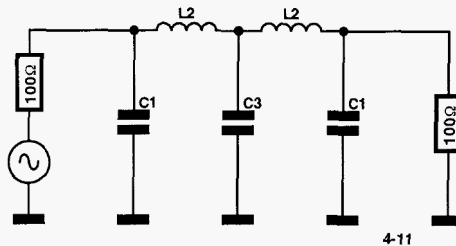
n	g_1	g_2	g_3	g_4	g_5	g_6	g_7	g_8	g_9	g_{10}	g_{11}
0,5 dB ondulation											
1	0,6986	1,0000									
2	1,4029	0,7071	1,9841								
3	1,5963	1,0967	1,5963	1,0000							
4	1,6703	1,1926	2,3661	0,8419	1,9841						
5	1,7058	1,2296	2,5408	1,2296	1,7058	1,0000					
6	1,7254	1,2479	2,6064	1,3137	2,4758	0,8696	1,9841				
7	1,7372	1,2583	2,6381	1,3444	2,6381	1,2583	1,7372	1,0000			
8	1,7451	1,2647	2,6564	1,3590	2,6964	1,3389	2,5093	0,8796	1,9841		
9	1,7504	1,2690	2,6678	1,3673	2,7239	1,3673	2,6678	1,2690	1,7504	1,0000	
10	1,7543	1,2721	2,6754	1,3725	2,7392	1,3806	2,7231	1,3485	2,5239	0,8842	1,9841
1,0 dB ondulation											
1	1,0177	1,0000									
2	1,8219	0,6850	2,6599								
3	2,0236	0,9941	2,0236	1,0000							
4	2,0991	1,0644	2,8311	0,7892	2,6599						
5	2,1349	1,0911	3,0009	1,0911	2,1349	1,0000					
6	2,1546	1,1041	3,0634	1,1518	2,9367	0,8101	2,6599				
7	2,1664	1,1116	3,0934	1,1736	3,0934	1,1116	2,1664	1,0000			
8	2,1744	1,1161	3,1107	1,1839	3,1488	1,1696	2,9685	0,8175	2,6599		
9	2,1797	1,1192	3,1215	1,1897	3,1747	1,1897	3,1215	1,1192	2,1797	1,0000	
10	2,1836	1,1213	3,1286	1,1933	3,1890	1,1990	3,1738	1,1763	2,9824	0,8210	2,6599
2,0 dB ondulation											
1	1,5296	1,0000									
2	2,4881	0,6075	4,0957								
3	2,7107	0,8327	2,7107	1,0000							
4	2,7925	0,8806	3,6063	0,6819	4,0957						
5	2,8310	0,8985	3,7827	0,8985	2,8310	1,0000					
6	2,8521	0,9071	3,8467	0,9393	3,7151	0,6964	4,0957				
7	2,8655	0,9119	3,8780	0,9535	3,8780	0,9119	2,8655	1,0000			
8	2,8733	0,9151	3,8948	0,9605	3,9335	0,9510	3,7477	0,7016	4,0957		
9	2,8790	0,9171	3,9056	0,9643	3,9598	0,9643	3,9056	0,9171	2,8790	1,0000	
10	2,8831	0,9186	3,9128	0,9667	3,9743	0,9704	3,9589	0,9554	3,7619	0,7040	4,0957

Tableau A4-2 (suite). Valeurs des composants pour un filtre passe-bas Tchebychev (pour un filtre avec une ondulation de n dB, le dernier point à n dB est obtenu pour $\omega = 1$ rad/s).

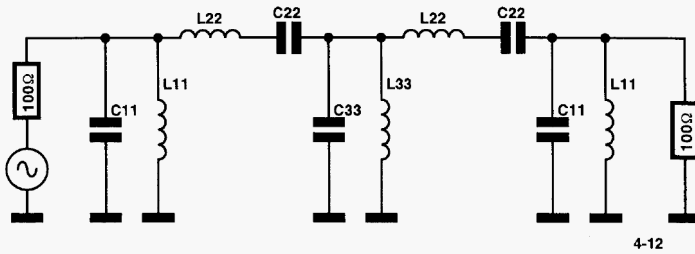
n	g_1	g_2	g_3	g_4	g_5	g_6	g_7	g_8	g_9	g_{10}	g_{11}
3,0 dB ondulation											
1	1,9953	1,0000									
2	3,1013	0,5339	5,8095								
3	3,3487	0,7117	3,3487	1,0000							
4	3,4389	0,7483	4,3471	0,5920	5,8095						
5	3,4817	0,7618	4,5381	0,7618	3,4817	1,0000					
6	3,5045	0,7685	4,6061	0,7929	4,4641	0,6033	5,8095				
7	3,5182	0,7723	4,6386	0,8039	4,6386	0,7723	3,5182	1,0000			
8	3,5277	0,7745	4,6575	0,8089	4,6990	0,8018	4,4990	0,6073	5,8095		
9	3,5340	0,7760	4,6692	0,8118	4,7272	0,8118	4,6692	0,7760	3,5340	1,0000	
10	3,5384	0,7771	4,6768	0,8136	4,7425	0,8164	4,7260	0,8051	4,5142	0,6091	5,8095

Problèmes

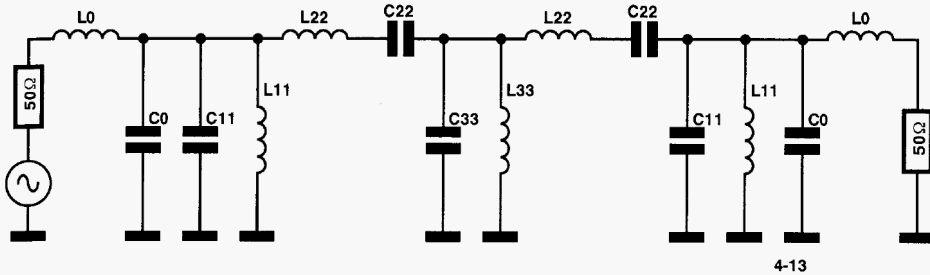
4.1 Concevoir un filtre passe-bas Tchebychev à cinq composants de 0,5 dB d'ondulation. Vous supposez que les impédances d'entrée et de sortie sont de $100\ \Omega$. Utilisez des condensateurs en parallèle aux extrémités. La bande passante (du continu jusqu'au dernier point à 0,5 dB) doit être de 100 kHz (il s'agit de la convention utilisée dans les tableaux Tchebychev de l'annexe 4.1, concernant la bande passante). Servez-vous des tableaux pour trouver les valeurs du filtre normalisé ($1\ \Omega$, 1 rad/s), puis refaites les calculs avec $100\ \Omega$ et 100 kHz.



4.2 Les résultats obtenus précédemment vont vous permettre de concevoir un filtre passe-bande Tchebychev à cinq composants avec une ondulation de 0,5 dB. Vous supposez que les impédances d'entrée et de sortie sont de $100\ \Omega$. La fréquence centrale doit être de 5 MHz et la largeur de bande totale de 200 kHz.

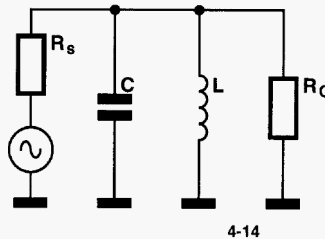


4.3 Modifiez le filtre du problème 4.2 afin que son impédance soit de $50\ \Omega$ vous ajouterez un réseau d'adaptation en L à chaque extrémité.



Testez le filtre du Problème 3.3 en vous servant du programme d'analyse de circuit (voir le Problème 1.3 du chapitre 1) sur une plage de fréquences de 4,5 MHz à 5,5 MHz par pas de 20 kHz.

4.4 Le filtre passe-bande à un composant représenté ci-dessous n'utilise qu'un seul résonateur en parallèle. Dans le filtre passe-bas normalisé, le résonateur est considéré comme étant un simple condensateur en parallèle.



Montrer que la réponse en fréquence de ce filtre est donnée par :

$$\frac{P}{P_{\max}} = \frac{1}{1 + Q^2 [(f/f_0) - (f_0/f)]^2}$$

où f_0 est la fréquence de résonance du circuit LC et Q est défini par $R/(\omega_0 L)$ où R résulte de la mise en parallèle de R_b et R_c .

4.5 Les filtres passe-haut découlent des filtres passe-bas en permutant bobines et condensateurs ; les valeurs des composants des tableaux des filtres passe-bas normalisés sont remplacées par leur réciproque (par exemple, un condensateur de 2 F devient une bobine de 0,5 H). La réponse à ω du filtre passe-haut normalisé sera égale à celle du filtre normalisé passe-bas à $1/\omega$.

Transformez le filtre passe-bas de la figure 4-4(b) en un filtre passe-haut. (*Réponse* : 1 H, 0,5 F et 1 H). Mettez-le ensuite à l'échelle pour que la fréquence de coupure soit de 5 kHz et que son impédance soit de 50 Ω . Transformez enfin ce dernier pour obtenir un filtre coupe-bande centré à 500 kHz avec une largeur de bande de 10 kHz.

4.6 Modifiez le programme d'analyse de circuit (voir le Problème 1.3 du chapitre 1) pour calculer non seulement l'amplitude de la réponse du circuit mais également sa réponse en phase (égale à l'angle de phase de la tension de sortie moins l'angle de phase de la tension d'entrée). Calculez la réponse en phase du filtre Butterworth de la figure 4.4(a). *Note* : les réseaux en échelle appartiennent à une classe de réseaux (les « réseaux à phase minimale ») dans lesquels la réponse en amplitude détermine exclusivement la réponse en phase et vice versa. Nous aborderons ultérieurement les filtres « passe-tout » qui n'en font pas partie ; leur phase varie en effet avec la fréquence tandis que l'amplitude reste constante.

5. Convertisseurs de fréquence

Dans le domaine des hautes fréquences, il est fréquent de procéder à une transposition de fréquence, ce qui signifie que tous les signaux situés dans une bande de fréquences donnée sont décalés dans le spectre vers une bande de fréquences plus élevée ou plus basse. Chaque composante spectrale est décalée de la même quantité. Par exemple, les boîtiers de télédistribution décalent un canal donné du réseau câblé vers un canal VHF situé plus bas en fréquence. La quasi-majorité des récepteurs radio et des téléviseurs sont de type *superhétérodyne*, c'est-à-dire que le canal que l'on souhaite recevoir commence par être décalé vers une bande de fréquence intermédiaire normalisée ou « FI » (les anglo-saxons parlent d'« IF » pour *Intermediate Frequency*). La plus grande partie de l'amplification s'effectue dans la bande FI, ce qui présente l'avantage de n'avoir rien à régler, dans cette partie importante du récepteur, lorsqu'on change de station ou de canal. On applique le même principe dans les émetteurs multi-bandes ; il est souvent plus facile de décaler un signal déjà modulé que de le créer totalement à n'importe quelle fréquence. On nomme également la transposition de fréquence « changement de fréquence » et même « mélange ». Nous allons décrire plusieurs circuits changeurs de fréquence (ou mélangeurs).

Emploi d'un multiplicateur parfait comme changeur de fréquence

Un changeur de fréquence mélange le signal d'entrée, qui doit être décalé, avec un signal de référence dont la fréquence est égale au décalage souhaité. On appelle dans un récepteur radio, oscillateur local (ou « OL »), l'oscillateur, intégré dans le récepteur, qui produit un signal de référence. Pour créer de nouvelles fréquences, les changeurs de fréquence sont forcément de type non linéaire puisqu'un circuit linéaire ne peut que modifier l'amplitude et la phase d'un ensemble d'ondes sinusoïdales. Lorsqu'on travaille en basses fréquences, « mélange » signifie addition (superposition linéaire qui produit de nouvelles fréquences). Cependant, dans le domaine des hautes fréquences, mélangeur est synonyme de multiplication ; les changeurs de fréquence hautes fréquences réalisent le *produit* des signaux d'entrée par le signal de l'OL. Cette opération non linéaire (multiplication) engendre de nouveaux signaux aux fréquences décalées.

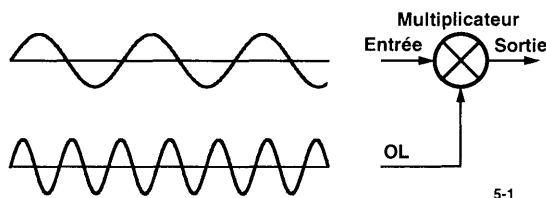


Figure 5-1. Utilisation d'un multiplicateur de tension comme convertisseur de fréquence (changeur de fréquence).

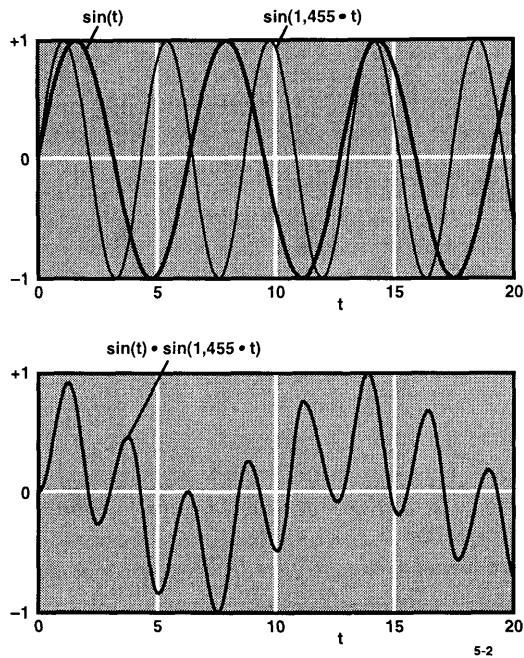


Figure 5-2. Formes d'onde en entrée et en sortie d'un multiplicateur.

Considérons tout d'abord un multiplicateur idéal utilisé comme changeur de fréquence (voir la figure 5-1). La tension de sortie du multiplicateur est égale au produit des deux tensions d'entrée. À la figure 5-2, le signal d'entrée sinusoïdal est multiplié par celui d'un OL dont la fréquence est 1,455 fois plus élevée. Les multiplicandes apparaissent dans le tracé du haut. Le tracé du bas représente leur produit, dans lequel on trouve, comme on peut le constater, des fréquences plus élevées et plus basses que les fréquences originelles. Une formule trigonométrique classique montre qu'en sortie on rencontre deux signaux à des fréquences distinctes : un signal décalé vers le haut à $\omega_L + \omega_R$ et un second décalé vers le bas à $\omega_L - \omega_R$:

$$\sin(\omega_R t) \sin(\omega_L t) = 1/2 [\cos(\omega_R - \omega_L)t - \cos(\omega_R + \omega_L)t] \quad (5.1)$$

Si nous remplaçons le signal hautes fréquences simple par un signal à deux composantes spectrales $U_1(\omega_1 t) + U_2(\omega_2 t)$, la sortie devient :

$$U_1 \cos[(\omega_1 - \omega_L)t] - U_1 \cos[(\omega_1 + \omega_L)t] + U_2 \cos[(\omega_2 - \omega_L)t] - U_2 \cos[(\omega_2 + \omega_L)t] \quad (5.2)$$

Cette combinaison linéaire de deux signaux se retrouve exactement dans une bande décalée vers le haut et dans une bande décalée vers le bas ; de même, toute combinaison linéaire, c'est-à-dire toute distribution spectrale, se retrouvera fidèlement dans ces bandes de sortie décalées. Vu sous cet angle, un changeur de fréquence est un circuit linéaire. On ne souhaite généralement conserver qu'une seule de ces bandes ; un filtre passe-bande adéquat rejettera l'autre. Les multiplicateurs analogiques parfaits ne sont pas courants dans le domaine des hautes fréquences, mais l'un d'entre eux est à présent disponible dans des versions rapides destinées aux changeurs de fréquence ; il s'agit du multiplicateur à cellule de transconductance de Gilbert.

Changeurs de fréquence à commutation

Si le signal de l'OL est un signal carré et non plus un signal sinusoïdal, le signal de sortie du changeur de fréquence contiendra non seulement les fréquences fondamentales décalées vers le haut et vers le bas mais également les fréquences harmoniques impaires du signal de l'OL correspondant à sa décomposition en série de Fourier. Il est généralement facile de filtrer ces nouvelles composantes, c'est pourquoi il est intéressant d'utiliser un oscillateur local qui délivre des signaux carrés. De plus, cette solution présente un avantage : il est possible de remplacer le multiplicateur par un inverseur unipolaire qui relie tour à tour la sortie au signal d'entrée et au moins du signal d'entrée, puisqu'il le multiplie par $+1$ ou par -1 . C'est ce que schématise la figure 5-3.

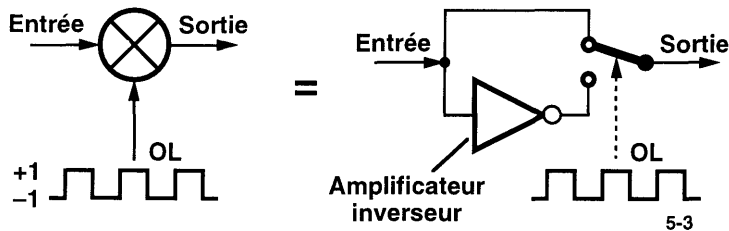


Figure 5-3. Fonctionnement d'un changeur de fréquence à commutation.

On obtient facilement l'inversion de phase nécessaire au fonctionnement du commutateur à l'aide d'un transformateur à point milieu tandis que deux transistors peuvent se charger de la commutation, l'un pour la branche du haut et l'autre pour la branche du bas. Le circuit de la figure 5-4 met en œuvre des FET (transistors à effet de champ). Les changeurs de fréquence dans lesquels se trouvent des transistors sont dits actifs. Un second transformateur procure l'inversion de phase du signal de l'OL de sorte que l'un des FET est passant tandis que l'autre est bloqué.

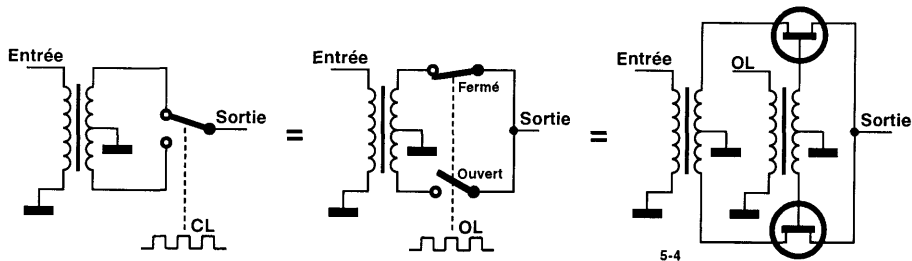


Figure 5-4. Changeur de fréquence à commutation actif.

Nous aurions tout aussi bien pu prendre le signal situé au point milieu et utiliser des FET pour mettre à la masse une extrémité du secondaire puis l'autre. Dans la configuration représentée à la figure 5-5, il est plus simple, puisque les transistors ne sont pas flottants, de leur fournir les signaux d'attaque.

On rencontre plus fréquemment des diodes comme éléments de commutation dans le montage décrit précédemment. Ce type de changeur de fréquence passif à commutation est représenté à la figure 5-6. La tension issue du transformateur de l'OL attaque tour à tour la paire de diodes du haut et la paire de diodes du bas. L'amplitude du signal de l'OL est assez élevée pour que les diodes conductrices aient une

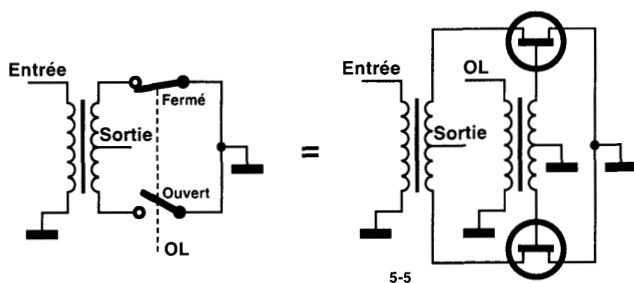


Figure 5-5. Changeur de fréquence actif à commutation alternée.

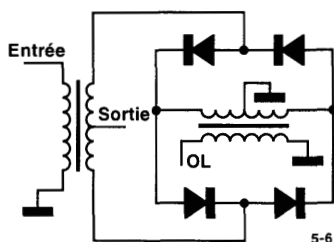


Figure 5-6. Dans ce mélangeur, des diodes remplacent les interrupteurs.

impédance très faible (faible zone de déplétion) et les diodes non conductrices aient une impédance très élevée (large zone de déplétion). L'extrémité du transformateur d'entrée située au point milieu des diodes passantes est en fait reliée à la masse par le biais du secondaire du transformateur de l'OL. Vous remarquerez que le courant passe par les deux côtés du transformateur de l'OL pour retourner à la masse ; par conséquent aucun flux n'est créé dans ce transformateur qui présente, pour ce courant, une impédance nulle. Ce circuit se présente généralement sous l'aspect de celui de la figure 5-7 ; il prend alors le nom de mélangeur symétrique double à diodes.

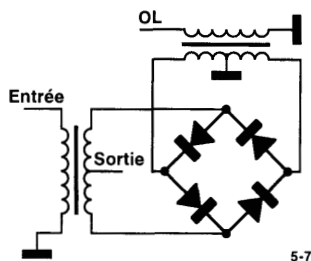


Figure 5-7. Mélangeur symétrique double à diodes en anneau.

Tous les changeurs de fréquence à commutation que nous avons vus jusqu'ici sont du type « symétrique double » puisqu'il existe une isolation entre l'entrée et la sortie. Aucun signal issu de l'OL n'apparaît ni à l'entrée ni à la sortie des hautes fréquences ou de la FI, et aucun signal issu des hautes fréquences, à l'exception des produits de mélange, n'apparaît ni à l'entrée ni à la sortie de la FI. Ainsi, il est préférable d'employer un mélangeur de type symétrique lorsque c'est le premier sous-ensemble d'un récepteur. Un

mélangeur asymétrique permettrait au signal issu de l'OL d'atteindre l'antenne et le rayonnement qui en résulterait pourrait provoquer des interférences dans les autres récepteurs (et révéler également l'emplacement du récepteur). Les mélangeurs simples que nous allons étudier ne sont pas équilibrés. La figure 5-8 représente le schéma d'un mélangeur asymétrique à commutation. Il effectue la multiplication du signal par un signal carré d'amplitude comprise entre +1 et 0 (plutôt qu'entre +1 et -1) : ce n'est ni plus ni moins qu'un signal d'amplitude comprise entre +1/2 et -1/2 avec un décalage de 1/2. Le signal carré produit lui-même le déplacement vers le haut et vers le bas des bandes précédemment, tandis que le décalage permet à l'entrée hautes fréquences d'atteindre la sortie sans subir de déplacement.

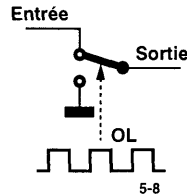


Figure 5-8. Mélangeur asymétrique à commutation.

Dispositif changeur de fréquence non linéaire

Nous venons d'étudier les changeurs de fréquence les plus simples qui soient ; ils utilisent un seul composant non linéaire. Une diode est évidemment un exemple d'un tel composant et nombreux sont les mélangeurs qui n'en utilisent qu'une seule. Les tensions du signal hautes fréquences et du signal issu de l'OL sont simplement additionnées et le signal somme est appliqué à la diode. Le courant, fonction non linéaire de la tension appliquée, contient les produits de mélange aux fréquences $Nf_{HF} \pm Mf_{OL}$, où N et M sont des nombres entiers. De même, un transistor attaqué par la somme des tensions du signal HF et du signal issu de l'OL, produira un courant qui contiendra les produits de mélange, surtout si le signal issu de l'OL produit de vastes excursions non linéaires. On se sert parfois d'un transistor FET à double grille comme d'un mélangeur ; la tension du signal issu de l'OL est appliquée à une grille tandis que celle du signal HF est appliquée à l'autre. Il en résulte une certaine isolation entre l'OL et la HF (qui existe naturellement dans un mélangeur équilibré comme un mélangeur à diodes en anneau).

Mélangeur à diode

La figure 5-9 est le schéma théorique d'un mélangeur à une diode. Le premier amplificateur opérationnel effectue la somme des tensions du signal HF et du signal issu de l'OL. Cette somme est appliquée à la diode. L'entrée du second amplificateur opérationnel est une masse virtuelle, de sorte que la totalité de la tension somme est présente aux bornes de la diode. Le second amplificateur opérationnel est utilisé en convertisseur courant-tension ; il produit une tension proportionnelle au courant qui circule dans la diode. Ce montage à amplificateurs opérationnels est destiné à attirer l'attention sur le fait que la *somme* des tensions du signal HF et du signal issu de l'OL attaque la diode. Les montages que l'on rencontre fréquemment utilisent des composants passifs et le résultat obtenu dans la sommation n'est pas toujours satisfaisant. La réponse d'une diode est en réalité exponentielle ; la formule suivante donne le courant en fonction de la tension appliquée :

$$I = I_s [\exp(U/U_{th}) - 1] \quad (5.3)$$

où $U_{th} = U_{thermique} = kT/e = 26 \text{ mV}$. Dans le cas des petits signaux, c'est-à-dire lorsque $U \ll 26 \text{ mV}$, il est possible d'effectuer le développement de la fonction exponentielle pour obtenir la tension de sortie du mélangeur :

$$U_{uit} = I_s R [U/U_{th} + (U/U_{th})^2/2! + (U/U_{th})^3/3! + \dots] \quad (5.4)$$

Puisque $U = U_{HF} + U_{OL}$, le premier terme représente la tension de traversée (non équilibrée) aux fréquences HF et OL. Le second terme (quadratique) produira les bandes latérales voulues, décalées vers le haut et vers le bas, puisque l'élévation au carré de $U_{OL} + U_{HF}$ renferme le produit $2U_{LO}U_{HF}$. Ce terme crée également des composantes de polarisation et des composantes de fréquence double. Le terme du troisième ordre produira des composantes aux fréquences harmoniques de rang 3 des fréquences HF et OL ainsi qu'aux fréquences $2\omega_{HF} + \omega_{OL}$, $2\omega_{HF} - \omega_{OL}$, $2\omega_{OL} + \omega_{HF}$ et $2\omega_{OL} - \omega_{HF}$. Ces produits sont normalement rejetés loin de la fréquence de sortie désirée et de toute façon filtrés. Il n'est pas nécessaire de pousser le développement de la fonction exponentielle si la tension d'entrée est suffisamment faible. Cependant dans le cas contraire, le terme suivant (du quatrième ordre) produit des composantes indésirables au sein de la bande de fréquence de sortie. Voyons ce qui se passe : considérons un signal d'entrée ayant deux composantes, $A_1 \cos(\omega_1 t)$ et $A_2 \cos(\omega_2 t)$. L'une des composantes de sortie sera, à une constante près :

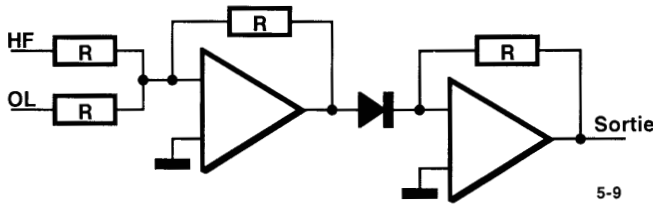


Figure 5-9. Schéma théorique d'un mélangeur à une diode.

$$\cos(\omega_L t) [A_1 \cos(\omega_1 t) + A_2 \cos(\omega_2 t)]^3 \quad (5.5)$$

Cette expression contient un terme

$$\begin{aligned} \cos(\omega_L t) 3A_1^2 \cos^2(\omega_1 t) A_2 \cos(\omega_2 t) = \\ = 3/2 A_1^2 A_2 \cos(\omega_L t) [1 - \cos(2\omega_1 t) \cos(\omega_2 t)] \end{aligned} \quad (5.6)$$

qui contient à son tour

$$3/8 A_1^2 A_2 \cos([\omega_L - (2\omega_1 - \omega_2)]t) \quad (5.7)$$

Lorsque ω_1 et ω_2 sont proches l'une de l'autre, $2\omega_1 - \omega_2$ et $2\omega_1 + \omega_2$ le sont également et peuvent se trouver dans la bande de sortie utile. Dans un récepteur radio, deux signaux très puissants créeront un produit de mélange qui perturbera, à cette fréquence, la réception d'un signal faible.

Nous verrons ultérieurement que la multiplication, base du changement de fréquence, permet également de moduler l'amplitude d'une porteuse et donc de produire une modulation AM. La multiplication, le changement de fréquence et la modulation (AM) découlent de la même opération fondamentale.

Problèmes

5.1 On se sert parfois de deux multiplicateurs, deux déphaseurs et d'un additionneur ou d'un soustracteur pour réaliser un mélangeur qui n'a qu'une bande de sortie (on parle alors de mélangeur à bande latérale unique). La conception d'un mélangeur à bande latérale supérieure découle directement de la formule suivante :

$$\cos(\omega_{HF}t) \cos(\omega_{LO}t) - \sin(\omega_{HF}t) \sin(\omega_{LO}t) = \cos[(\omega_{HF} + \omega_{LO})t]$$

Représentez le schéma d'un tel mélangeur.

5.2 Le mélangeur à commutation à diodes en anneau fonctionne également lorsqu'on permute les entrées HF et OL. Expliquez pourquoi.

5.3 On définit le gain de conversion d'un mélangeur comme étant le rapport de la puissance de sortie (le signal FI) sur la puissance d'entrée (signal HF). Calculez le gain de conversion (en fait c'est une perte, puisqu'il est inférieur à l'unité) d'un mélangeur à commutation parfait.

Astuces : le signal carré issu de l'OL contient non seulement une onde sinusoïdale de fréquence égale à celle du décalage souhaité, mais également des ondes sinusoïdales de fréquences harmoniques impaires. Seule la composante fondamentale permet le transfert de la puissance vers la fréquence FI désirée.

5.4 Considérons le cas de deux signaux de même fréquence mais présentant une différence de phase θ , mélangés indépendamment l'un de l'autre à une nouvelle fréquence. Nous supposons que les deux mélangeurs sont identiques et qu'ils sont attaqués par le même signal OL. Montrez que la différence de phase des signaux mélangés reste égale à θ .

5.5 Dans le domaine des hautes fréquences, on se sert continuellement des formules trigonométriques $\cos(a + b) = \cos a \cos b - \sin a \sin b$ et $\sin(a + b) = \sin a \cos b + \cos a \sin b$. Démontrez ces formules en vous servant de figures géométriques ou de la formule $e^{jx} = \cos x + j \sin x$.

6. Récepteur radio

Caractéristiques essentielles

Tout récepteur radio possède trois caractéristiques fondamentales : l'*amplification*, la *sensibilité* et la *sélectivité*. Est-ce qu'un signal faible, émis par une station que l'on souhaite recevoir et arrivant aux bornes d'une antenne, est capable de produire un signal de sortie (audio, vidéo ou données) suffisamment fort et intelligible et quel est son comportement en présence de signaux puissants émis sur des fréquences proches ? Pour définir la *sensibilité*, on considère souvent le plus petit signal d'entrée détectable, par exemple $1\text{ }\mu\text{V}$ à l'entrée de l'antenne $50\text{ }\Omega$ d'un récepteur, ou comme étant la puissance de bruit ajoutée par le récepteur, appliquée à son entrée. La *sélectivité* est normalement déterminée par les qualités du filtre principal passe-bande et peut s'exprimer sous la forme « -3 dB à 2 kHz de la fréquence centrale et -20 dB à 10 kHz de la fréquence centrale ». Les constructeurs de récepteurs ne mentionnent pas précisément la forme de la courbe de réponse telle qu'on pourrait la décrire par « filtre Tchebychev à huit sections avec une ondulation de 2 dB et des points à -3 dB séparés par 4 kHz ».

Amplification

Essayons de quantifier le niveau d'amplification que doit avoir un récepteur radio classique. Une puissance basses fréquences de 1 mW dans un écouteur miniature standard produit un niveau sonore d'environ 100 dB au-dessus du seuil d'audition. Un signal basses fréquences à peine perceptible sera donc produit par une puissance de -100 dBm (100 dB sous 1 mW , soit 10^{-13} W). Précisons tout de même qu'il est nécessaire de disposer d'un niveau sonore supérieur de 50 dB au niveau minimal pour obtenir un certain confort d'écoute, ce qui nous conduit à une puissance de 10^{-8} W . Vous vous rendez compte qu'avec des circuits efficaces, les piles d'un récepteur radio portatif durent vraiment longtemps ! Les niveaux de puissance acoustique sont étonnamment faibles ; vous ne rayonnez qu'une puissance acoustique d'environ 1 mW lorsque vous criez et approximativement 1 nW lorsque vous chuchotez. Quelle puissance reçoit un récepteur ? Une simple antenne filaire peut réellement recueillir une puissance HF de 10^{-8} W , provenant d'un signal émis par une station radio située à 20000 km avec une puissance d'émission de 10 kW propagée à 1 MHz par une antenne omnidirectionnelle. Considérons donc tout d'abord le cas des récepteurs radio « autonomes ».

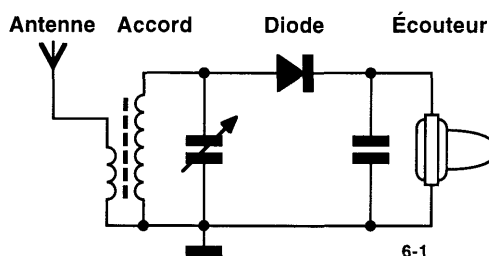


Figure 6-1. Récepteur à galène autonome.

Poste à galène

Les premiers récepteurs radio, les postes à galène, étaient autonomes. Un redresseur à diode à cristal recouvrait l'enveloppe de modulation et convertissait la puissance HF reçue en une puissance acoustique suffisante pour attaquer un écouteur miniature. Un simple circuit accordé *LC* servait de filtre passe-bande pour sélectionner la station recherchée et pouvait également jouer le rôle d'un réseau d'adaptation d'antenne. La figure 6-1 illustre le schéma fondamental d'un récepteur à galène.

Les éléments que nous avons considérés précédemment montrent qu'un récepteur autonome est susceptible de recevoir une plage importante de fréquences. Il est toutefois indispensable de disposer d'une amplification conséquente dès que l'on remplace l'antenne à long fil par un cadre et l'écouteur miniature par un haut-parleur. Nous verrons par la suite que la réponse d'un détecteur à diode, en présence de signaux faibles, est quadratique. Pour que la détection d'enveloppe d'un signal AM se fasse correctement, le signal appliqué au détecteur à diode doit avoir un niveau élevé, de l'ordre de plusieurs milliwatts. L'invention du tube à vide a permis d'obtenir l'amplification requise. Tout récepteur renferme en principe des amplificateurs HF et BF. L'amplification HF fournit assez d'énergie pour que le détecteur fonctionne de manière satisfaisante, tandis que l'amplification BF procure l'énergie nécessaire à l'attaque d'un haut-parleur.

Récepteur à amplification directe

Les premiers récepteurs radios à tubes étaient simplement des postes à galène, semblables à ceux que nous venons de voir, dotés d'une préamplification HF et d'une postamplification BF. Ces postes à amplification directe ou postes TRF (de *Tuned Radio Frequency*) disposaient de réglages d'accord propres à chacun des nombreux étages d'amplification HF. L'utilisateur qui voulait changer de station devait régler individuellement chaque cadran (souvent à l'aide d'un tableau d'accord ou d'un graphique).

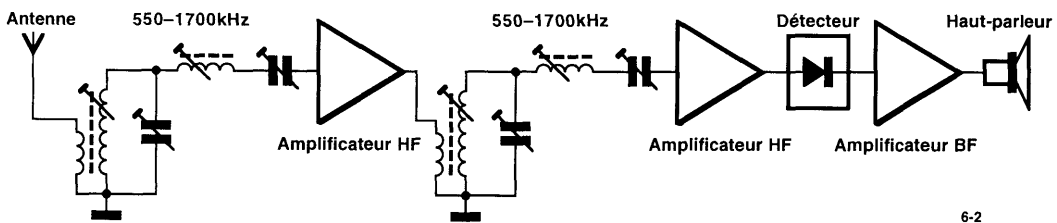


Figure 6-2. Récepteur TRF.

La figure 6-2 représente le schéma d'un récepteur théorique à amplification directe avec ses amplificateurs à plusieurs étages et ses filtres passe-bande. Vous vous apercevez que tous les condensateurs et bobines sont de type variable afin d'accorder la fréquence centrale des filtres passe-bande et conserver également une largeur de bande correcte, environ 10 kHz pour l'AM. Dans la pratique, les filtres passe-bande utilisaient un résonateur couplé plutôt qu'un système à conversion directe passe-bas vers passe-bande tel qu'on le voit ici.

Récepteur superhétérodyne

Les inconvénients du poste TRF résident dans son coût et le désagrément des nombreux réglages d'accord à effectuer. La plupart de ces réglages ont disparu suite à l'invention du montage superhétérodyne que l'on doit en 1917 à Armstrong. Le montage conçu par Armstrong comprend un récepteur TRF à accord fixe, c'est-à-dire accordé sur une seule fréquence, précédé d'un changeur de fréquence (mélangeur et oscillateur local) de sorte que la fréquence de tout signal issu d'une station donnée peut être décalée vers la fréquence du récepteur TRF. Cette fréquence est connue sous le nom de fréquence intermédiaire ou *FI*. Le montage superhétérodyne est utilisé dans la quasi-totalité des récepteurs radios, des téléviseurs et des récepteurs radar d'aujourd'hui. Nous pouvons citer parmi les rares exceptions qui subsistent, les talkies-walkies destinés aux enfants, les récepteurs hyperfréquences utilisés dans les systèmes d'ouverture de portes automatiques et certains radars qui signalent nos excès de vitesse sur route. La figure 6-3 représente le schéma classique d'un récepteur « superhet » destiné à la réception des bandes de radiodiffusion.

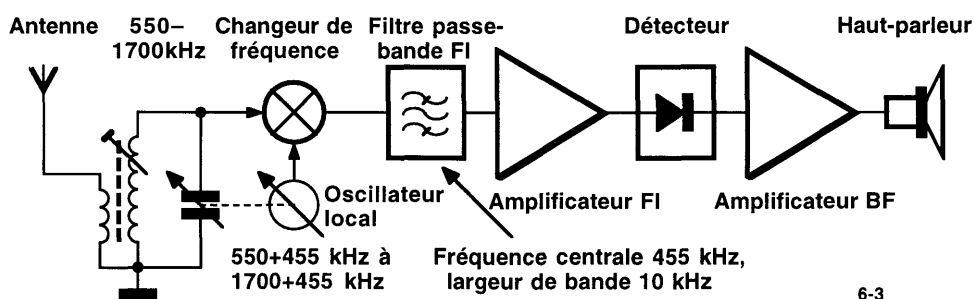


Figure 6-3. Récepteur superhétérodyne classique de radiodiffusion AM.

La sélectivité dépend des filtres passe-bande à accord fixe situés dans la section d'amplification FI. Le détecteur est toujours une diode, à savoir un classique poste à galène. Nous aurons l'occasion de revenir sur le détecteur et d'étudier d'autres circuits de détection. L'amplificateur FI à accord fixe est susceptible de procurer la totalité du gain HF. Nous verrons qu'il est aussi parfois nécessaire d'effectuer une amplification en amont du mélangeur. Le schéma de la figure 6-3 est également valable pour les récepteurs de radiodiffusion FM (à ceci près que la fréquence FI est de 10,7 MHz et que l'on emploie un détecteur FM en lieu et place d'un détecteur d'enveloppe) et les récepteurs de télévision (dans ce cas la fréquence centrale de la FI est de 45 MHz et des circuits de traitement de l'image et du son suivent le schéma de base du récepteur superhétérodyne).

NOTE Il faut mentionner l'existence d'un récepteur *hétérodyne* qui a précédé le récepteur superhétérodyne. Inventé par Reginald Fessenden, l'un des pionniers de la radio, le récepteur hétérodyne convertissait directement le signal HF en signal BF. On rencontre encore parfois de tels récepteurs, connus sous le nom de « récepteurs à conversion directe ». L'assemblage d'un changeur de fréquence FI en entrée suivi d'un récepteur hétérodyne (de préférence à un récepteur TRF) donne un récepteur superhétérodyne avec détecteur de produit (voir le chapitre 27).

Réjection de la fréquence-image

Le récepteur superhétérodyne présente des inconvénients qui lui sont propres. En ce qui concerne les signaux présents à l'entrée du mélangeur, le récepteur détecte aussi bien les signaux de la fréquence recherchée que des signaux présents à une fréquence indésirable connue sous le nom de *fréquence-image*. Pour comprendre ce qui se passe, prenons un exemple concret. Supposons que nous disposions d'un récepteur classique AM ayant une fréquence FI de 455 kHz et que la fréquence de l'OL soit de 1015 kHz pour recevoir une station de radiodiffusion dont la fréquence d'émission est de 1015 kHz – 455 kHz = 560 kHz. Cette fréquence (560 kHz) se situe dans le bas de la bande de radiodiffusion AM. Les mélangeurs que nous avons étudiés produiront également un signal FI de 455 kHz pour un signal d'entrée dont la fréquence est de 1470 kHz, soit 455 kHz au-dessus de la fréquence de l'OL. Si l'on n'effectue aucun filtrage HF avant le mélangeur, tout signal de fréquence égale à 1470 kHz sera détecté en même temps que le signal recherché à 560 kHz. Il est nécessaire de placer en amont du mélangeur un filtre passe-bande pour laisser passer la fréquence désirée et bloquer les signaux de la fréquence-image. Vous remarquerez que dans cet exemple (représentatif de la plupart des récepteurs AM), ce filtre passe-bande « anti-image » doit être accordable et que, pour les récepteurs qui ont une commande d'accord unique, le filtre accordé doit toujours « suivre » exactement à 455 kHz sous la fréquence de l'OL. La poursuite, dans ce cas, n'est pas trop difficile à obtenir ; la fréquence image étant située à plus d'une octave au-dessus de la fréquence recherchée, le simple filtre à un composant de la figure 6-2 est assez large et fournira une bonne réjection de la fréquence-image, peut-être de 20 dB. Malgré tout, une atténuation de 20 dB ne conviendra que si le signal situé à la fréquence-image n'est pas supérieur de plus de 20 dB à celui que l'on veut capter.

Que se passe-t-il dans le cas où un récepteur, ayant la même fréquence FI de 455 kHz, doit également recevoir les bandes d'ondes courtes ? On rencontre la pire des situations à la fréquence la plus élevée, 30 MHz, lorsque la fréquence de l'image n'est que de 3% supérieure à la fréquence recherchée. Un filtre fournissant une atténuation de 20 dB, à seulement 3% de sa fréquence centrale, doit comporter de nombreux étages qu'il faut accorder simultanément à l'aide d'un condensateur variable mécanique à plusieurs cages ou de diodes varicaps. Comme nous l'avons vu, la fréquence centrale du filtre doit suivre, avec un décalage de 455 kHz, la fréquence de l'OL pour que le signal recherché se trouve à l'intérieur de l'étroite bande FI. Le problème de la réjection d'image n'est pas simple à résoudre lorsque la fréquence FI est nettement inférieure à la fréquence du signal d'entrée.

Comment résoudre le problème de la fréquence-image ?

L'utilisation d'une fréquence FI bien plus élevée permet de résoudre le problème. Si le récepteur de radiodiffusion AM que nous avons analysé précédemment disposait d'une FI à 10 MHz plutôt qu'à 455 kHz, on pourrait accorder l'OL sur 10,560 MHz pour recevoir une station à 560 kHz. La fréquence-image se situerait alors à 20,560 MHz. Lorsqu'on accorde le récepteur vers le haut de la bande AM, à 1700 kHz, la fréquence-image prend alors la valeur de 21,700 MHz. Dans ce cas, on peut placer en amont du récepteur un simple filtre passe-bande à accord fixe, assez large pour couvrir la totalité de la bande de radiodiffusion ; il rendra le récepteur insensible aux fréquences-images. La figure 6-4 illustre un tel montage. L'accord du récepteur s'effectue simplement en modifiant la fréquence de l'OL. Le filtre FI à 10 MHz doit évidemment avoir une bande passante étroite de 10 kHz pour que le récepteur conserve sa sélectivité.

Le schéma de la figure 6-4 est parfaitement réalisable bien qu'il soit plus onéreux de construire les indispensables filtres à bande étroite pour des hautes fréquences ; les résonateurs à quartz remplacent alors les filtres LC à constantes localisées. Lorsque la bande de fréquences d'entrée est plus large, c'est le cas

d'un récepteur à ondes courtes qui couvre de 3 à 30 MHz, la fréquence FI doit être bien plus élevée, et il est illusoire de tenter de réaliser des filtres à bande étroite. Même à 10 MHz, une largeur de bande de 10 kHz sous-entend une bande passante fractionnaire de seulement 0,1%. Il est malgré tout possible de résoudre le problème de la fréquence-image en mettant en œuvre une fréquence FI élevée dans un récepteur superhétérodyne à *double changement de fréquence*. C'est ce que nous allons étudier à présent.

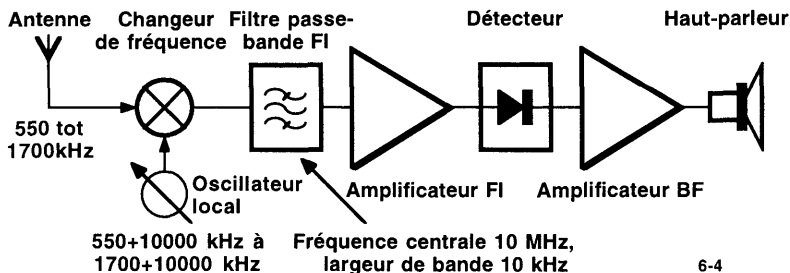


Figure 6-4. L'emploi d'une fréquence FI de 10 MHz permet à ce récepteur de s'affranchir des problèmes de fréquences-images.

Récepteur superhétérodyne à double changement de fréquence

La figure 6-5 montre comment un second convertisseur de fréquence abaisse le premier signal FI (à 10 MHz) à 455 kHz où il est traité par la section FI classique du récepteur de la figure 6-2. La bande passante du premier filtre FI à 10 MHz peut être plus large que celle du dernier étage. Supposons par exemple que sa bande passante soit de 500 kHz. La seconde fréquence de l'OL se situe à 10,455 MHz, de sorte que le second convertisseur de fréquence produit une image d'un signal dont la fréquence est de 10,910 MHz. Mais comme le filtre de notre premier convertisseur de fréquence FI coupe à $10 \text{ MHz} + 0,500 \text{ MHz}/2 = 10,25 \text{ MHz}$, il n'y aura aucun signal à 10,910 MHz.

Ce système présente également certains inconvénients : avec un tel récepteur, il n'est pas possible de recevoir un signal dont la fréquence se situe à proximité immédiate de la première FI, étant donné qu'il est difficile d'éviter la traversée directe du signal dans l'amplificateur FI. Les changements multiples entraînent un accroissement du nombre d'OL, de combinaisons de fréquences sommes et différences dues aux non linéarités, et donc de signaux parasites (également connus sous le nom de « *birdies* »).

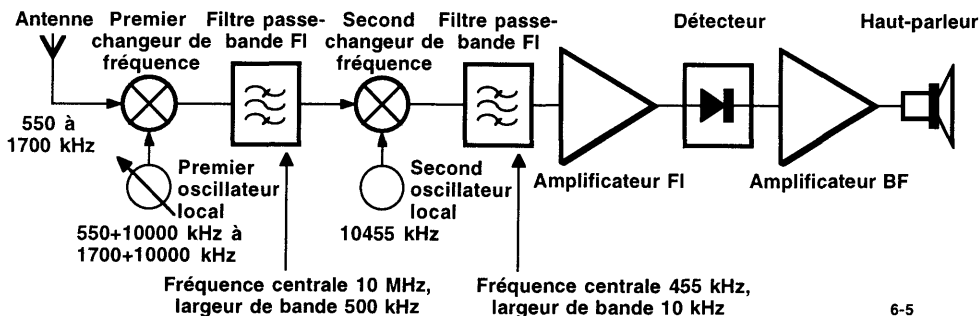


Figure 6-5. Récepteur superhétérodyne à double changement de fréquence.

Les récepteurs de communication modernes emploient un double voire un triple changement de fréquence dans lequel la fréquence de la première FI se situe aux environs de 40 MHz. Le filtre d'entrée destiné à l'élimination de la fréquence-image est généralement un filtre passe-bas dont la fréquence de coupure est de 30 MHz. Les filtres à quartz modernes offrent une bande passante particulièrement étroite, même à 40 MHz ; il est alors possible d'abaisser la fréquence de sortie du premier étage FI vers une seconde FI de fréquence bien plus basse. Il est parfois indispensable d'employer un triple changement de fréquence lorsque la dernière fréquence FI est très basse, par exemple de l'ordre de 50 kHz. Les premières fréquences FI dans la région des VHF (*Very High Frequencies* ou très hautes fréquences) impliquent une excellente stabilité des OL, ce qui ne pose aucun problème grâce à l'utilisation d'oscillateurs à quartz et de synthétiseurs de fréquences. La première génération de récepteurs mettant en œuvre des OL synthétisés souffrait de problèmes liés au bruit de phase des oscillateurs ; les signaux puissants situés à proximité du signal recherché, mais en dehors de sa bande passante, faisaient rentrer le bruit des bandes latérales de l'oscillateur.

Commande automatique de gain

Presque tous les récepteurs disposent, d'une manière ou d'une autre, d'une commande automatique de gain (CAG) qui règle le gain des amplificateurs HF et/ou FI, en fonction de la puissance du signal d'entrée. En l'absence d'un tel dispositif, le récepteur serait surchargé ; la réponse d'un amplificateur suratténué n'est plus linéaire (il y a « écrêtage ») et le signal de sortie est distordu et trop fort. En principe, le niveau sonore de sortie d'un récepteur FM ne varie pas en fonction du niveau du signal, mais un étage d'amplification surchargé produit une distorsion, c'est pourquoi on trouve également un circuit de CAG dans un récepteur FM. Les récepteurs de télévision requièrent une CAG précise pour maintenir un niveau de contraste correct. Tout circuit de CAG est un système asservi. Dans un simple récepteur AM, le détecteur à diode fournit une tension continue commode qui peut commander le courant de polarisation (et donc le gain) des amplificateurs HF.

Réducteur de bruit

De nombreux récepteurs, y compris la plupart des récepteurs de télévision, disposent d'un circuit de réduction du bruit qui élimine les effets des bruits impulsionnels, comme celui des impulsions transitoires produites par les systèmes d'allumage des automobiles. Les impulsions parasites sont tellement fines qu'elles peuvent perturber les étages FI. Le taux d'utilisation du récepteur reste élevé et la pointe de tension est parfaitement audible (ou visible dans le cas d'un récepteur de télévision). Il faut procéder au déblocage avant que la bande passante ne soit très réduite parce que les filtres allongent la durée des impulsions.

Traitement numérique du signal dans un récepteur

Le filtrage et la détection au sein d'un récepteur peuvent en principe être traités de façon numérique, tout au moins après une première amplification et une réduction de la largeur de bande. Le traitement numérique d'un signal FI permet d'obtenir l'amplitude et la réponse en phase recherchées. Il est indispensable que la puissance de traitement soit suffisante pour avoir un bon résultat : il est généralement

fait appel à des processeurs de traitement numérique du signal (puces DSP de *Digital Signal Processing*) ainsi qu'à des convertisseurs analogique-numérique à haute définition rapides. La norme qui est proposée pour le système de télévision de pointe (ATV de *Advanced Television*) nécessite un gros traitement numérique non seulement pour effectuer les tâches purement numériques mais également le traitement du signal, comme l'annulation adaptative du signal par trajets multiples.

Bibliographie

- 1 *Handbook for Radio Amateurs*, 71^e édition, The American Radio Relay League, Newington, Connecticut, 1994. (Il s'agit d'un imposant ouvrage dans lequel vous trouverez des schémas, des explications et des conseils de réalisations.)
- 2 W. Gosling (1986), *Radio Receivers*. New York: Peter Peregrinus. (C'est une bonne analyse des récepteurs modernes.)
- 3 U. Rohde et T. Bucher (1988), *Communications Receivers, Principles & Design*. New York: McGraw-Hill. (Ce livre est à lui-seul un véritable cours.)

Problèmes

6.1 La bande de radiodiffusion FM s'étend de 88 MHz à 108 MHz. Les récepteurs classiques FM ont une fréquence FI de 10,7 MHz. Quelle est la plage d'accord de l'OL ?

6.2 Pourquoi demande-t-on aux passagers d'un avion de ne pas se servir de récepteurs radio durant le vol ?

6.3 Deux signaux sinusoïdaux de fréquences différentes simplement additionnés, apparaissent sous la forme d'un signal ayant une seule fréquence, mais modulé en amplitude. On se sert, par exemple, de ce phénomène de « battement » pour accorder deux cordes de guitare sur une même fréquence. Lorsqu'elles ne se trouvent pas encore rigoureusement à la même fréquence, le son résultant semble battre lentement à une cadence égale à leur différence de fréquence. Montrez que

$$\sin[(\omega_0 - \delta\omega)t] + \sin[(\omega_0 + \delta\omega)t] = A(t) \sin(\omega_0 t) \text{ où } A(t) = 2 \cos(\delta\omega t)$$

L'oreille est sensible à l'intensité du son, c'est-à-dire au carré de l'amplitude, c'est pourquoi elle perçoit une cadence de battement de $2\delta\omega$, soit la fréquence différence. Notez que de nouvelles fréquences sont produites par le processus de détection non linéaire, et non pas par l'opération d'addition qui est linéaire.

6.4 Lorsque vous vous servez de votre récepteur radio dans un environnement envahi par de nombreuses stations, il vous arrive d'entendre parfois une tonalité aigüe à 10 kHz, en même temps que l'émission recherchée. Si vous actionnez d'avant en arrière le bouton d'accord, la hauteur de cette tonalité ne change pas. Quelle en est la cause ?

6.5 Lorsque vous recherchez une station sur un récepteur radio, et plus particulièrement la nuit, vous pouvez entendre des « sifflements d'interférences » (ce sont des tonalités BF sifflantes) qui changent de fréquence lorsque vous manœuvrez lentement le bouton d'accord. Quelle en est la raison ? Le récepteur est-il fautif ? (*Réponse* : oui !)

6.6 Avec des composants modernes et une commande numérique, nous pourrions construire un bon récepteur TRF. Quels sont les avantages que l'on pourrait opposer à un récepteur superhétérodyne ? Quels inconvénients présenterait-il ?

7. Amplificateurs en classe C et en classe D

La non linéarité (signal de sortie en fonction du signal d'entrée) des amplificateurs de puissance en classe C et en classe D est tellement flagrante que l'on ferait mieux de parler de générateurs synchronisés d'ondes sinusoïdales. Ils comprennent une alimentation, au moins un élément de commutation (tube, transistor,...) et un circuit LCR . « R » représente la charge, souvent la résistance de rayonnement d'une antenne, et le circuit résonne à la fréquence de fonctionnement. Le rendement élevé de ces amplificateurs est utilisé dans de nombreuses applications qui vont du plus petit émetteur alimenté par piles jusqu'à des émetteurs de radiodiffusion de plusieurs mégawatts en passant par les fours industriels. L'amplitude de sortie, bien que n'étant pas une fonction linéaire de l'amplitude d'entrée, est linéairement proportionnelle à la tension d'alimentation. On peut donc moduler en amplitude ces amplificateurs en faisant simplement varier leur tension d'alimentation.

Amplificateurs en classe C

La figure 7-1 illustre le schéma d'un amplificateur en classe C parfait. Le composant actif (transistor ou tube) est modélisé sous la forme d'un commutateur dont la résistance à l'état passant, r , est constante et la résistance à l'état bloqué, infinie. Ce modèle convient assez bien à la représentation d'un transistor à effet de champ de puissance (FET de *Field Effect Transistor*).

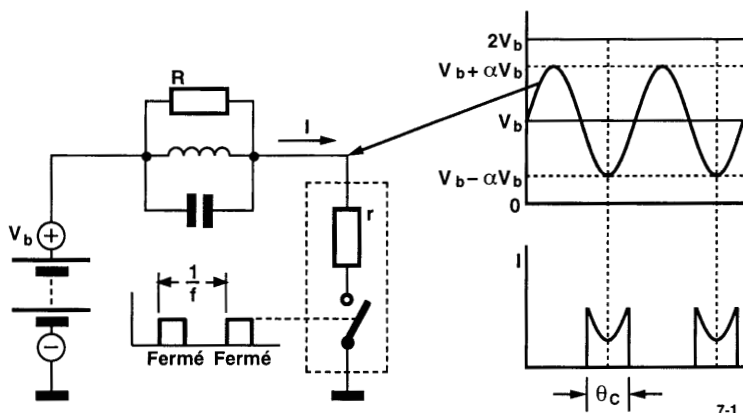


Figure 7-1. Fonctionnement d'un amplificateur en classe C.

Le commutateur est fermé pendant un laps de temps, inférieur à une alternance du signal HF, pendant lequel l'alimentation fournit une impulsion de courant afin de faire pomper le circuit RLC . L'angle de conduction, θ_c , est généralement de l'ordre de 120° (voir la figure 7-1). Le circuit LC a normalement un

facteur Q élevé (environ 5), de sorte que l'effet de volant minimise la distorsion de l'onde sinusoïdale, due à l'action brutale du commutateur et à l'amortissement qui se produit entre les impulsions. Notez qu'il est également possible de n'actionner le commutateur que tous les deux ou trois cycles, et ainsi de suite. Un circuit qui fonctionne de cette façon est dit doubleur ou tripleur en classe C – c'est un multiplicateur de fréquence. On a représenté le signal d'attaque sous la forme d'un signal carré, mais c'est le plus souvent un signal sinusoïdal polarisé de telle façon que la conduction ne se fasse qu'autour de ses extrémités arrondies.

Analyse simplifiée du fonctionnement en classe C

Les amplificateurs en classe C fonctionnent normalement en saturation, cela signifie que le commutateur présente toujours sa plus faible résistance (bien inférieure à la résistance de charge) lorsqu'il est passant. L'amplitude du signal d'attaque doit être suffisante pour placer le commutateur en mode franchement passant ou franchement bloqué. Nous allons procéder à l'analyse du schéma de la figure 7-1 et calculer la tension de sortie (donc la puissance de sortie) et le rendement. Si l'on se réfère à la figure 7-1, θ_c est l'angle de conduction, V_b la tension d'alimentation, r la résistance du commutateur à l'état passant et αV_b la tension de crête de l'onde sinusoïdale. Voici comment trouver α . La puissance d'entrée (puissance fournie par la batterie) doit être égale à la somme de la puissance de sortie (puissance dissipée dans R) plus la puissance dissipée dans la résistance r du commutateur. Ces puissances sont respectivement données par le produit moyenne de la tension de la batterie $\times$ le courant, $(\alpha V_b)^2/(2R)$, et la moyenne de $I^2 r$. L'équation donnant la puissance s'écrit :

$$\frac{1}{2\pi} \int_{-\theta_c/2}^{\theta_c/2} V_b \left(\frac{V_b - \alpha V_b \cos \theta}{r} \right) d\theta = \frac{(\alpha V_b)^2}{2R} + \frac{1}{2\pi} \int_{-\theta_c/2}^{\theta_c/2} \left(\frac{V_b - \alpha V_b \cos \theta}{r} \right)^2 r d\theta \quad (7.1)$$

ou encore :

$$\frac{1}{\pi r/R} \int_0^{\theta_c/2} (1 - \alpha \cos \theta)(\alpha \cos \theta) d\theta - \frac{\alpha^2}{2} = 0 \quad (7.2)$$

Calculons l'intégrale de l'équation 7-2 et résolvons-la pour α , nous obtenons :

$$\alpha(\theta_c, r/R) = \frac{2 \sin(\theta_c/2)}{\theta_c/2 + (\sin \theta_c)/2 + \pi r/R} \quad (7.3)$$

La figure 7-2 représente le tracé de α^2 (proportionnel à la puissance de sortie) pour cinq valeurs de r/R (rapport de la résistance du commutateur en mode passant sur la résistance de charge). La valeur médiane ($r/R = 0,1$) est celle que l'on utilise dans les applications pratiques. On obtient la puissance maximale lorsque l'angle de conduction est de 180° , c'est-à-dire lorsque le commutateur est fermé quand la tension est dans la zone négative. Si l'angle de conduction dépasse 180° , les incursions dans la zone positive prélèvent de l'énergie au circuit accordé, la puissance de sortie est donc réduite.

Vous remarquerez que α^2 , et donc α , peuvent être supérieurs à l'unité, surtout lorsque l'angle de conduction est égal à 180° . Nous verrons cependant ci-dessous que l'on obtient un rendement nettement plus élevé pour des angles de conduction bien inférieurs à 180° . On calcule le rendement en divisant la puissance de sortie par la puissance que fournit la batterie :

$$\eta(\theta_c, r/R) = \left[\frac{(\alpha V_b)^2/2R}{1/2\pi \int_{-\theta_c/2}^{\theta_c/2} V_b [(V_b - \alpha V_b \cos \theta')/r] d\theta'} \right]^{-1} \quad (7.4)$$

Calculons l'intégrale de l'équation 7-4, nous trouvons :

$$\eta(\theta_c, r/R) = \frac{(\pi r/R) \alpha^2}{\theta_c - 2\alpha \sin(\theta_c/2)} \quad (7.5)$$

où la valeur α est fournie par l'équation 7-3. La figure 7-3 représente le tracé du rendement pour les valeurs de r/R déjà choisies. Le rendement est maximal pour 90° lorsque $r/R = 0,1$.

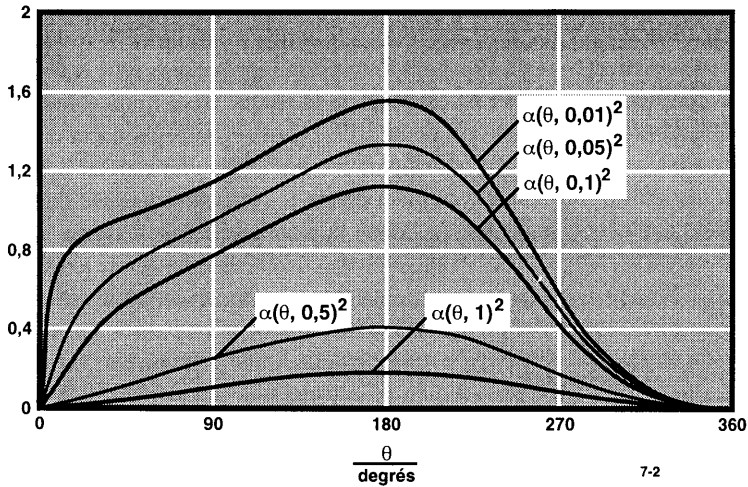


Figure 7-2. Puissance de sortie d'un amplificateur en classe C (α^2) en fonction de l'angle de conduction, pour cinq valeurs de r/R .

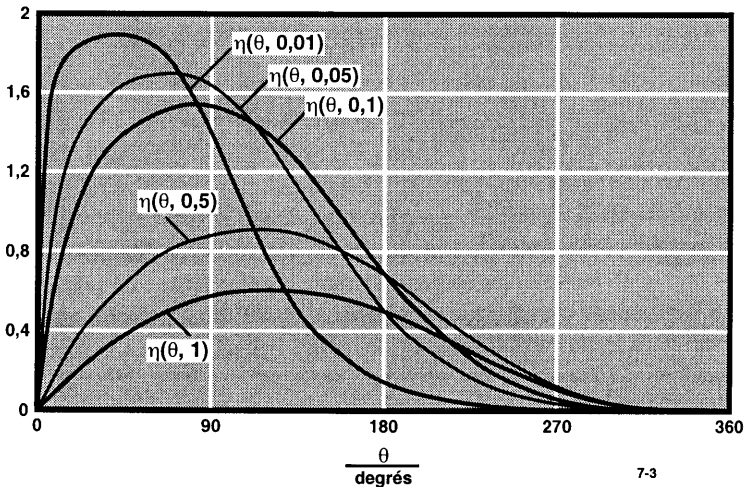


Figure 7-3. Rendement d'un amplificateur en classe C en fonction de l'angle de conduction, pour cinq valeurs de r/R .

Analyse générale d'un fonctionnement en classe C avec un tube ou un transistor réel

L'analyse simplifiée des figures 7-2 et 7-3 ne fournira pas de résultats précis pour un amplificateur en classe C utilisant des tubes ou des transistors bipolaires, parce que ces composants ne possèdent pas la résistance dans l'état passant constante d'un FET. Cependant, leurs caractéristiques non linéaires sont fournies par le constructeur sous forme de courbes, et il est ainsi possible d'analyser précisément le fonctionnement en classe C en s'aidant d'un ordinateur. La façon de procéder est la même que pour l'analyse du circuit simplifié ; on suppose que les circuits résonnants LC en entrée et en sortie ont un facteur Q assez élevé pour forcer les formes d'onde en entrée et en sortie à rester sinusoïdales. Pour effectuer cette analyse, les caractéristiques du composant sont représentées en « courant constant ». Dans le cas d'un tube, les courbes, pour différents courants de la plaque, sont tracées dans un repère dont les axes sont la tension de la plaque et la tension de la grille. On suppose également que les tensions de la plaque et de la grille sont sinusoïdales et que l'intégration numérique du produit courant de la plaque $\times$ tension de la plaque moyennée sur un cycle complet donne la puissance dissipée dans le composant. Le produit tension d'alimentation $\times$ courant *moyen* donne la puissance totale d'entrée. La différence représente la puissance fournie à la charge. Le concepteur choisit un composant et une tension d'alimentation et simule les formes d'onde pour un essai donné (points de polarisation et amplitudes des formes d'onde sinusoïdales). Il faut en général procéder à plusieurs essais, afin de maximaliser, par exemple, la puissance de sortie en fonction du composant choisi ou le rendement pour une puissance de sortie spécifiée.

Remarques sur l'attaque

Les amplificateurs en classe C qui utilisent des tubes à vide¹ rendent presque toujours la grille de commande positive lorsque le tube est passant. La grille absorbe du courant et dissipe de l'énergie. Les feuilles de caractéristiques comprennent les courbes de courant de la grille de sorte que le concepteur peut se servir de la procédure vue précédemment pour vérifier que le cycle de fonctionnement choisi reste dans les limites de dissipation de la plaque et de la grille acceptables. Lorsqu'on emploie des tétrodes, il faut procéder à une troisième analyse qui concerne la dissipation de la grille-écran.

Circuits alimentés en série et en parallèle

La figure 7-4 illustre deux configurations classiques. Le circuit alimenté en série (représenté à la figure 7-4(a)) est équivalent à celui de la figure 7-1, si ce n'est qu'un transistor remplace le commutateur. Vous vous apercevez que le schéma d'un amplificateur en classe C est le même que celui d'un « amplificateur réel » en classe A ou en classe B. Voici les seules différences qui existent : un amplificateur en classe C possède un organe de commande (base, grille de transistor ou de tube) polarisé au-delà du blocage, de sorte que la durée de conduction du composant est inférieure à une alternance et, lorsque le composant est

¹ Un tube à vide triode ressemble à un transistor NPN. La plaque, la grille de commande et la cathode du tube correspondent respectivement au collecteur, à la base et à l'émetteur du transistor. Un tube tétrode possède une grille supplémentaire, la grille-écran, située entre la grille de commande et la plaque. On applique généralement sur la grille-écran une tension de polarisation constante ; cette grille est un écran électrostatique utile entre la grille de commande et la plaque. Voir Bibliographie [3].

passant, il l'est totalement (saturé). Le circuit alimenté en parallèle de la figure 7-4(b) permet de mettre à la masse une extrémité de la charge. Une bobine d'arrêt HF permet l'alimentation du transistor et un condensateur de liaison bloque toute composante continue vers le circuit résonnant. L'inductance de la bobine d'arrêt HF est élevée de sorte que le courant qui la traverse est pratiquement constant. Le commutateur tire des impulsions de charge du condensateur de liaison. Cette charge est à nouveau reconstituée via la bobine.

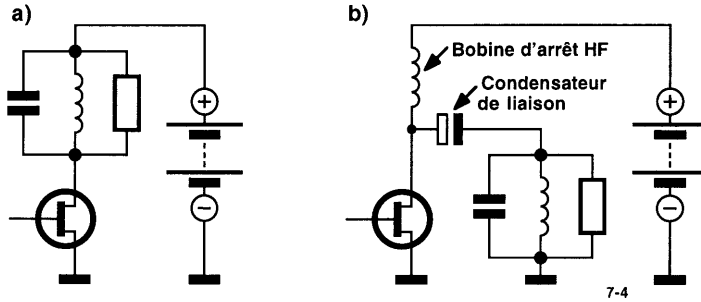


Figure 7-4. Amplificateurs alimentés en série (a) et en parallèle (b).

La figure 7-5 représente, de la gauche vers la droite, les circuits équivalents de circuits alimentés en série et en parallèle. La bobine d'arrêt HF mise en parallèle avec la bobine, à la figure 7-5(b) n'affecte en rien le fonctionnement du circuit. L'inductance de la bobine d'arrêt est si élevée qu'elle modifie peu la valeur de L , mais il est de toute façon possible de s'arranger pour que la mise en parallèle des deux bobines donne une valeur égale à celle de L . À la figure 7-5(c), la bobine d'arrêt laisse passer le courant continu vers le commutateur de sorte que le condensateur de liaison, de valeur très élevée, bloque le courant continu vers le circuit résonnant RLC original. N'oubliez pas que les schémas équivalents des circuits alimentés en série et en parallèle sont également valables pour les amplificateurs en classe A, en classe B et en classe C, qu'ils fonctionnent avec des grands ou des petits signaux.

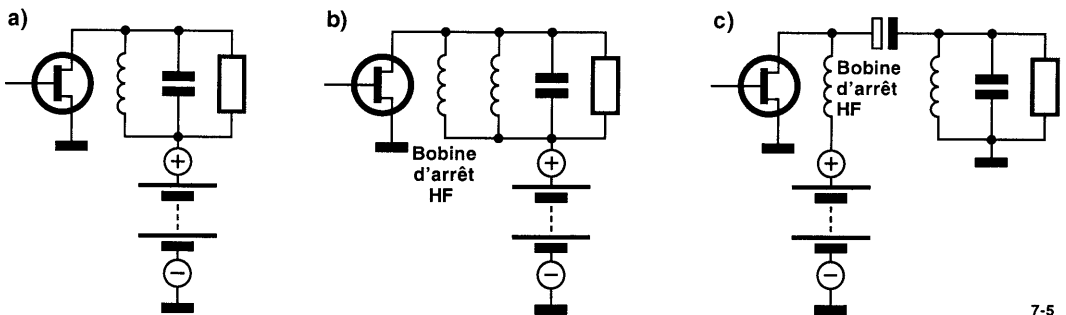


Figure 7-5. Circuits équivalents de circuits alimentés en série et en parallèle.

Utilisation d'un amplificateur en classe C comme multiplicateur de tension

L'une des propriétés intéressantes de l'amplificateur en classe C saturé est que la tension de crête de l'onde sinusoïdale HF de sortie est directement proportionnelle à la tension d'alimentation ($U_{\text{crête}} = \alpha V_b$). Lors du fonctionnement en mode saturé classique, la valeur de la constante de proportionnalité, α , est approximativement égale à 0,9. L'amplificateur en classe C est donc équivalent à un multiplicateur de tension, qui effectue le produit de l'onde sinusoïdale d'entrée (avec un coefficient d'amplitude proche de l'unité) par la tension d'alimentation. La modulation (variation) de la tension d'alimentation d'un amplificateur en classe C est la méthode habituelle utilisée dans les émetteurs-récepteurs AM. Vous en déduisez aisément que cette propriété de l'amplificateur en classe C serait considérée comme étant un défaut pour un amplificateur opérationnel (faible réjection de l'alimentation).

Amplificateur de puissance en classe C

Il existe des tubes dont la dissipation maximale de la plaque peut atteindre, voire dépasser, 1 MW. Le rendement typique d'un amplificateur en classe C est de 75–85%, ainsi un seul tube de puissance peut produire une puissance HF dépassant 6 MW. Si la plus forte puissance utilisée en radiodiffusion est d'environ 2 MW, on rencontre des amplificateurs bien plus puissants (ou des oscillateurs en classe C) dans des applications de fours industriels employés pour le séchage du contre-plaqué et la soudure. Les stations de radiodiffusion AM et FM mettent en œuvre depuis longtemps des amplificateurs en classe C pour leur rendement élevé, et dans le cas de l'AM, pour la caractéristique linéaire de la modulation de la plaque. Les amplificateurs des émetteurs-récepteurs les plus récents sont transistorisés et fonctionnent en classe C ou en classe D (dont nous allons parler). Les émetteurs-récepteurs transistorisés de puissance font appel à de nombreux modules de puissance élémentaires (≈ 1 kW).

Amplificateur en classe C modifié pour un meilleur rendement

Voici un problème fondamental qui se pose avec la fonction de commutation de l'amplificateur en classe C : le commutateur applique directement la tension d'alimentation aux bornes du circuit *LRC*. Le condensateur agit instantanément comme un court-circuit ; puisque les tensions ne sont pas équilibrées, il appelle un courant infini qu'il ne peut pas recevoir puisque celui-ci le courant fourni par l'alimentation est limité par la résistance du commutateur. Donc, pendant l'impulsion, ou tout au moins jusqu'à ce que la différence de tension soit négligeable, il existe un produit tension non nulle $\times$ courant : le commutateur *doit* donc dissiper de l'énergie. On inclut parfois des composants réactifs supplémentaires dans l'étage de sortie pour exciter des courants harmoniques sélectionnés et distordre volontairement la tension de la plaque. En particulier, il est possible d'aplatir le bas de l'onde sinusoïdale afin de réduire le produit tension $\times$ courant dans le commutateur et accroître le rendement. Raab (Bibliographie [2]) parle alors de fonctionnement en *classe F*.

Amplificateur en classe D

Le rendement d'un amplificateur en classe C ne peut en principe s'approcher de 100% que si le rapport cyclique tend vers zéro et si l'impulsion de courant tend vers l'infini. Le rendement d'un amplificateur en classe D peut en principe atteindre 100% avec des rapports cycliques et des courants réalistes. Il est nécessaire d'utiliser au moins deux commutateurs mais aucun des deux ne supporte simultanément la tension et le courant.

Amplificateur en classe D résonnant en série

La figure 7-6(a) illustre le concept de l'amplificateur en classe D résonnant en série. Un inverseur unipolaire produit un signal carré. Un filtre LC permet à l'onde sinusoïdale fondamentale d'atteindre la charge R . Il serait possible de relier l'une des contacts de l'inverseur à une alimentation négative, mais le relier à la masse convient puisque la sortie est à couplage alternatif. La figure 7-6(b) donne le schéma d'un circuit réel ; l'inverseur est composé de deux transistors. Ce montage à transistors est du type push-pull, c'est-à-dire que les deux transistors sont attaqués en opposition de phase : lorsqu'un transistor est bloqué, l'autre est passant.

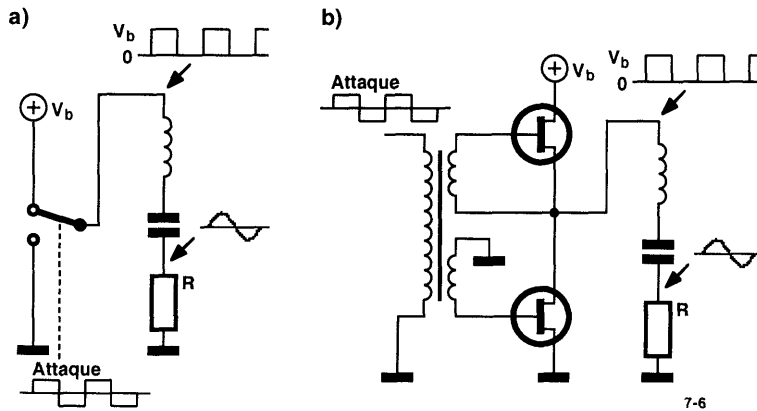


Figure 7-6. Amplificateur en classe D résonnant en série.

Calculons la tension aux bornes de la charge et le rendement. Puisque le condensateur bloque la composante continue, nous pouvons considérer que l'onde carrée est symétrique par rapport à la masse et son excursion en tension est comprise entre $-V_b/2$ et $+V_b/2$. Le circuit série LC extrait la composante sinusoïdale fondamentale du signal carré. Souvenez-vous de la décomposition en série de Fourier ; on peut décomposer notre signal carré sous la forme suivante : $A_1 \sin(\omega t) + A_3 \sin(3\omega t) + \dots$. Intéressons-nous au premier coefficient A_1 que l'on peut trouver en calculant la valeur moyenne du produit de l'onde carrée $\times \sin(\omega t)$. La figure 7-7 détaille ce calcul. Puisque $A_1 = 2V_b/\pi$, la puissance de sortie $A_1^2/(2R)$ est donnée par $2V_b^2/(\pi^2 R)$. Il ne nous arrive pas souvent d'avoir à écrire le terme « $\pi^2 R$ » ! Nous supposons, comme nous l'avons fait lors de l'analyse de l'amplificateur en classe C, que les FET ont une résistance r constante à l'état passant. Le courant qui atteint la charge emprunte l'un ou l'autre des commutateurs de sorte que le rapport de la puissance de sortie sur la puissance dissipée dans le commutateur est $I^2 R / (I^2 r)$ et le rendement $\eta = R/(R + r)$. Vous constatez que le rendement peut s'approcher de 100% tant que $r \ll R$.

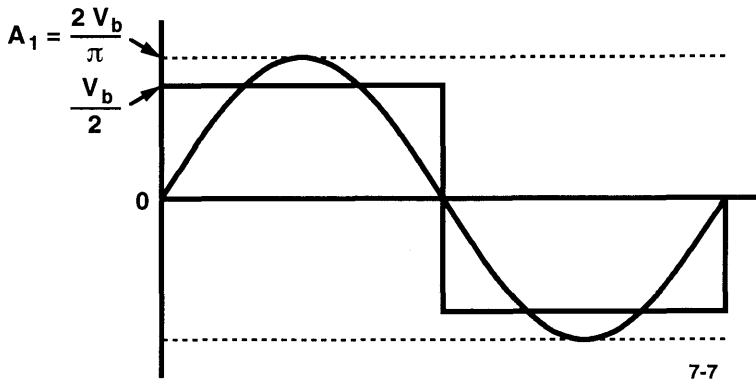


Figure 7-7. Extraction d'une onde sinusoïdale fondamentale à partir d'une onde carrée.

Notez également que l'amplificateur en classe D, tout comme l'amplificateur en classe C, est un multiplieur de tension parce que l'amplitude de l'onde sinusoïdale de sortie est directement proportionnelle à la tension d'alimentation.

Amplificateur en classe D résonnant en parallèle

L'amplificateur en classe D représenté à la figure 7-6 est formé d'une source de signaux carrés qui attaque un circuit *RLC* en série. On peut aussi concevoir qu'une source de courant de signaux carrés attaque un circuit *RLC* en parallèle. Dans le schéma de la figure 7-8, une forte bobine fournit un courant constant. Un interrupteur bipolaire commute la charge. Deux réalisations pratiques d'un tel montage sont dessinées à la figure 7-9. Le circuit de gauche utilise quatre transistors montés en pont, tandis que celui de droite emploie un transformateur et seulement deux transistors.

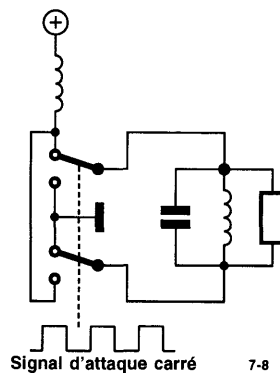


Figure 7-8. Fonctionnement d'un amplificateur en classe D résonnant en parallèle.

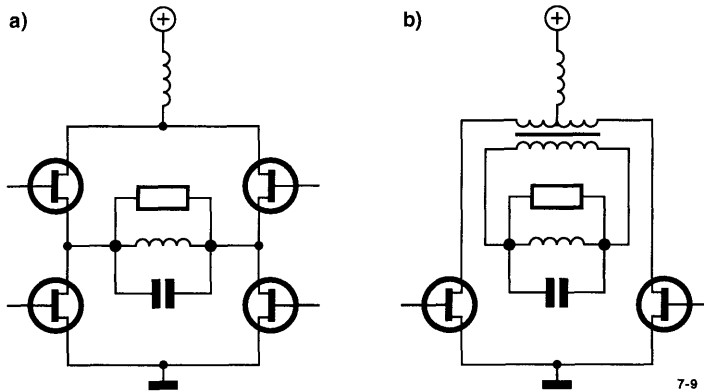


Figure 7-9. Schémas d'amplificateurs en classe D (fonctionnement en parallèle).

Classe C ou classe D ?

L'attrait de l'amplificateur en classe D réside dans le rendement qu'il peut atteindre. Néanmoins le rendement de l'amplificateur en classe C, surtout dans sa version améliorée (en classe F) est loin d'être médiocre. L'amplificateur en classe C ne fait appel qu'à un composant actif et peut fonctionner à la fréquence supérieure limite du tube ou du transistor. Il est parfois difficile de faire fonctionner un amplificateur en classe D. Comme dans tout circuit en totem pole, il faut s'assurer que les composants de chacune des branches ne peuvent pas conduire simultanément, même pendant un bref instant. Un fonctionnement en classe D impose que les composants se comportent comme des commutateurs parfaits, c'est pourquoi ils restreignent le fonctionnement en classe D, à des fréquences relativement basses. Les émetteurs de radiodiffusion AM peuvent travailler en classe D, tandis que les émetteurs FM (VHF) travaillent en classe C. En hyperfréquences, un fonctionnement, même en classe C, est délicat ; les angles de conduction sont élevés et un amplificateur conçu en classe C pourrait fonctionner en réalité en classe A.

Bibliographie

- 1 H.L. Krauss, C.W. Bostian en F.H. Raab (1980), *Solid State Radio Engineering*. New York: John Wiley.
- 2 F.H. Raab (1975), "High efficiency amplification techniques", *IEEE Circuits and Systems* 7: 3-11.
- 3 *Care and Feeding of Power Grid Tubes*. Varian-Eimac Laboratory Staff, San Carlos, California, 1967.

Problèmes

7.1 Imaginons que l'amplificateur en classe D résonnant en série de la figure 7-6(b) attaque une charge de 50Ω à une fréquence de 1 MHz. Les valeurs de la bobine et du condensateur sont respectivement de $9,49 \mu\text{H}$ et de $2,67 \text{ nF}$ (les réactances sont égales et opposées à 1 MHz). Ce circuit résonnant laisse passer

la composante à 1 MHz de l'onde carrée et atténue très fortement les harmoniques. Quel rapport existe-t-il entre la puissance délivrée à la charge à 3 MHz (troisième harmonique) et celle délivrée à la fréquence fondamentale (1 MHz) ?

Astuces : dans un signal carré, l'amplitude de la troisième harmonique est le tiers de celle de la fréquence fondamentale. N'oubliez pas qu'à 1 MHz, $X_L = X_C$ alors qu'à 3 MHz, la valeur de X_L triple tandis que celle de X_C est divisée par trois.

7.2 Un amplificateur en classe C construit autour d'un tube, de rendement égal à 75%, fournit une puissance de sortie en onde entretenue de 500 kW. La tension d'alimentation est de 60 kV et l'angle de conduction de 90° . Quel est le courant moyen fourni par l'alimentation ? (*Réponse* : 11,1 A.) Quelle est la valeur moyenne du courant lorsque le tube est à l'état passant ? (*Réponse* : 44,4 A.)

7.3 Expliquez pourquoi le rendement d'un amplificateur en classe C est faible si les impulsions d'attaque ne saturant pas le tube ou le transistor (en le rendant totalement passant) ?

7.4 Les impulsions de courant représentées à la figure 7-1 correspondent au cas où la tension de crête de l'onde sinusoïdale est inférieure à V_b (pour $\alpha < 1$). Dessinez la forme des impulsions de courant lorsque $\alpha > 1$. Vous supposerez ici que le commutateur laisse passer le courant dans un sens et dans l'autre.

8. Lignes de transmission

Nos schémas comportent des composants discrets reliés entre eux par des traits qui représentent une longueur de câblage nulle. Il est évident que tout câble réel a une certaine longueur et que sa capacité et son inductance interviennent sur le fonctionnement des circuits. Ce phénomène devient sensible dès que les dimensions du montage ne sont pas négligeables vis à vis de la longueur d'onde. D'un autre côté, il est tout à fait possible de prévoir ces effets, et donc en tenir compte lors de la conception d'un montage, si les interconnexions sont réalisées par des lignes de transmission correctes. Nous ne nous occuperons dans ce chapitre que des câbles à deux conducteurs comme les câbles coaxiaux et bifilaires.

Notions fondamentales

Chaque fraction élémentaire d'une ligne de transmission introduit une inductance en série et une capacité en parallèle ; le réseau en échelle représenté à la figure 8-1 devient une ligne de transmission réelle dès que δC et δL tendent vers zéro. *Remarque* : dans certains cas, comme la téléphonie en bande de base et les transmissions de données numériques, le modèle comporte également des résistances en série et en parallèle. En HF, cependant, la réactance en série est généralement bien plus élevée que la résistance en série et la réactance en parallèle est bien plus faible que la résistance en parallèle ; on peut donc négliger l'influence de ces résistances. Un tronçon de ligne de transmission placé en amont d'une impédance Z créera généralement une impédance modifiée Z' : c'est ce que l'on peut voir à la figure 8-2. La seule impédance qui ne soit pas modifiée est la résistance dont la valeur est égale à l'impédance caractéristique Z_0 de la ligne. Cette impédance est égale à $Z_0 = \sqrt{L/C} + j0$ où L et C représentent l'inductance et la capacité par unité de longueur. Il suffit de connaître L ou C puisqu'elles sont liées par la relation $LC = 1/v_{\text{phase}}^2 = \epsilon/c^2$ où v_{phase} est la vitesse de propagation, ϵ la constante diélectrique (par rapport au vide) et c la vitesse de la lumière. La relation qui existe entre L et C est valable pour toute configuration à deux conducteurs possédant une symétrie sur un axe, comme par exemple, une ligne de transmission irréaliste composée d'un conducteur interne de section carrée situé à l'intérieur d'un conducteur externe de section triangulaire.

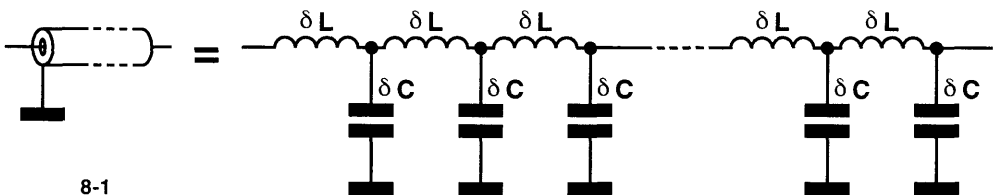


Figure 8-1. Ligne de transmission modélisée sous la forme d'un réseau en échelle constitué de sections LC infinitésimales.

Quelle est la « réalité » de l'impédance caractéristique ? Si nous branchons un ohmmètre pour courant continu classique aux bornes d'un rouleau de câble coaxial d'impédance $50\ \Omega$, allons-nous lire $50\ \Omega$? Réponse : en principe, oui, si le rouleau est assez long pour que la réflexion n'atteigne pas l'ohmmètre pendant que nous effectuons la mesure. Autrement, le fait de court-circuiter l'extrémité la plus éloignée du câble provoque une tension réfléchie négative qui annule la tension appliquée, et nous lirons alors $0\ \Omega$. Si nous laissons l'extrémité ouverte, il se crée une tension réfléchie positive dont le courant négatif (négatif, simplement parce que l'onde réfléchie revient) annule le courant appliqué, et nous lirons alors une impédance infinie. Mais en nous servant d'un générateur d'impulsions et d'un oscilloscope, nous pourrions aisément réaliser un ohmmètre assez rapide pour déterminer la valeur Z_0 .

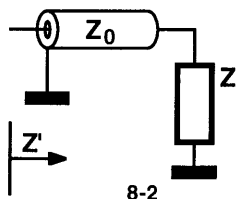


Figure 8-2. Modification d'une impédance située à l'extrémité d'une ligne de transmission.

Détermination de l'impédance caractéristique et de la vitesse de propagation

Pour aboutir à la formule $Z_0 = L/C$, examinons le schéma de la figure 8-3, dans lequel nous avons rajouté un tronçon de câble de longueur infinitésimale (δx mètres) à une terminaison parfaite, à savoir une résistance de $Z_0\ \Omega$. Voyons quelle relation doit exister entre L et C pour que l'impédance reste égale à Z_0 . Le circuit ressemble exactement à une cellule d'adaptation en L qui peut (et même doit) transformer une résistance en une résistance différente. Mais dans notre cas, δL et δC étant infinitésimaux, nous allons montrer que la variation d'impédance δZ_0 est nulle si $Z_0 = L/C$.

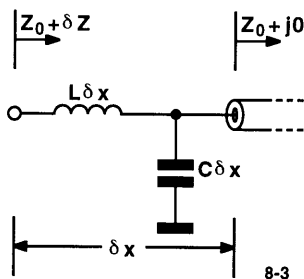


Figure 8-3. L'ajout d'une autre section infinitésimale ne doit pas modifier Z_0 .

Souvenez-vous que L et C sont l'inductance et la capacité par unité de longueur. L'étude du schéma nous conduit à écrire :

$$Z_0 + \delta Z = j\omega L\delta x + \frac{Z_0 [1/(j\omega C\delta x)]}{Z_0 + 1/(j\omega C\delta x)} = j\omega L\delta x + \frac{Z_0}{j\omega C\delta x Z_0 + 1} = \frac{j\omega L\delta x (1 + j\omega C\delta x Z_0) + Z_0}{(1 + j\omega C\delta x Z_0)^2} \quad (8.1)$$

$$= \frac{j\omega L\delta x + Z_0 - \omega^2 L C \delta x^2 Z_0^2}{1 + 2j\omega C\delta x Z_0 + \omega^2 C^2 \delta x^2 Z_0^2}$$

Lorsque δx tend vers 0, nous pouvons simplifier le terme de droite et écrire :

$$Z_0 + \delta Z \rightarrow Z_0 + j\omega\delta x(L - Z_0^2 C)$$

(8.2)

Ainsi, δZ sera nulle si $Z_0 = \sqrt{L/C}$, l'impédance restera constante lorsque nous allongerons la ligne. Vous pouvez vérifier que si l'on permute le condensateur en parallèle et la bobine en série, le résultat reste identique.

Une étude similaire montre la façon dont une onde se propage dans une ligne. Appliquons à son entrée un signal $e^{j\omega t}$ et cherchons l'expression de la chute de tension dans la longueur de ligne élémentaire (voir la figure 8-4). Nous savons déjà que l'impédance d'entrée est Z_0 ; le courant d'entrée est donc égal à U/Z_0 et la chute de tension dans l'inductance, δU est égale à $(U/Z_0)(j\omega L\delta x)$. Mais ceci est une équation différentielle :

$$\frac{dU}{dx} = -j\omega \frac{L}{Z_0} U = -j\omega \sqrt{LC} U \quad (8.3)$$

La solution de cette équation différentielle familière est :

$$U = U_f e^{-j\omega \sqrt{LC} x} = U_f e^{-jkx} \quad (8.4)$$

où $k = \omega \sqrt{LC} = \omega / v_{\text{phase}}$. Nous avons envoyé cette onde dans la direction x positive parce que, si nous faisons intervenir le temps, la tension est égale à :

$$U = U_f e^{j(\omega t - kx)} \quad (8.5)$$

Tout point de phase constante se déplace en respectant $\omega t - kx = \text{constante}$, donc $dx/dt = \omega/k = 1/\sqrt{LC} = v_{\text{phase}}$. Cette vitesse de phase est indépendante de ω ; il n'y a aucune dispersion. Cependant, les lignes de transmission réelles utilisent un matériau isolant dont la constante diélectrique dépend un peu de la fréquence : elles présentent donc une certaine dispersion.

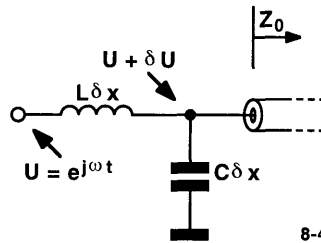


Figure 8-4. Recherche d'une variation de tension δU suite à une variation de longueur δL .

Modification d'une impédance par une ligne de transmission

Poursuivons et cherchons la façon dont un tronçon de ligne modifie une impédance (voir la figure 8-2). Une onde qui se propage selon l'axe positif des x est proportionnelle à $e^{j(\omega t - kx)}$, où $k = 2\pi/\lambda = \omega/v_{\text{phase}}$. De même, une onde qui se propage selon l'axe négatif des x , est proportionnelle à $e^{j(\omega t + kx)}$. La superposition de ces deux ondes représente une solution globale (à une fréquence donnée) pour trouver la tension en un point de la ligne. Considérons un tronçon de ligne de longueur l , dont l'extrémité droite (pour $x = 0$) est reliée à une impédance quelconque Z_C (« C » pour « charge »). Supposons qu'une source produise une onde incidente se propageant vers la droite, e^{-jkx} , et que Z_C crée une onde réfléchie ρe^{jkx} se propageant vers la gauche. La tension en tout point x de la ligne est donnée par $U(x) = e^{-jkx} + \rho e^{jkx}$. Le courant correspondant, comme le montre la figure 8-5, est $I(x) = 1/Z_0 (e^{-jkx} - \rho e^{jkx})$.

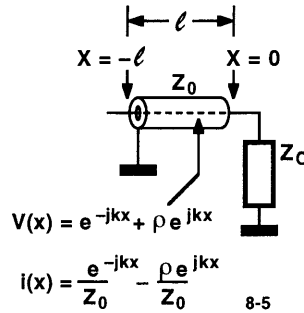


Figure 8-5. Calcul de la tension et du courant à l'entrée d'une ligne de transmission.

Le signe moins provient du fait que l'onde réfléchie se propage selon l'axe négatif des x . Plaçons-nous à droite ($x = 0$) et appliquons la loi d'Ohm : $U(0)/I(0) = Z_C$. Cela nous permet de calculer ρ :

$$\frac{U(0)}{I(0)} = Z_C = \frac{(1 + \rho)}{(1 - \rho)/Z_0}$$

donc

$$\frac{Z_C}{Z_0} = \frac{(1 + \rho)}{(1 - \rho)}$$

et

$$\rho = \frac{(Z_C - Z_0)}{(Z_C + Z_0)} \quad (8.6)$$

Servons-nous à présent de cette expression de ρ pour obtenir ce que nous recherchons, $U(-l)/I(-l)$, c'est-à-dire l'impédance d'entrée au point situé à gauche de la charge, à une distance l :

$$\begin{aligned} Z' &= \frac{e^{-jk(-l)} + \rho e^{jk(-l)}}{e^{-jk(-l)}/Z_0 - \rho e^{jk(-l)}/Z_0} \\ &= Z_0 \frac{(Z_C + Z_0)e^{jkl} + (Z_C - Z_0)e^{-jkl}}{(Z_C + Z_0)e^{jkl} - (Z_C - Z_0)e^{-jkl}} \\ &= Z_0 \frac{Z_C + jZ_0 \tan(kl)}{Z_0 + jZ_C \tan(kl)} \end{aligned} \quad (8.7)$$

Il est facile de mémoriser ce résultat important (la dernière ligne) : il ne comporte aucun signe moins et se présente sous une forme symétrique. Vous n'avez qu'à vous souvenir de :

$$(1 + j\tan)/(1 + j\tan)$$

Ceci étant mémorisé (un moyen mnémotechnique est « tan-tan »), vous n'oublierez pas où mettre les coefficients. Voici quelques cas particuliers :

1. Si $Z_C = Z_0$ alors $Z' = Z_0$, comme prévu.
2. Si $Z_C = 0$ (court-circuit) alors $Z' = jZ_0 \tan(kl)$, c'est une réactance pure, inductive pour $kl < \pi/2$, puis capacitive, et ainsi de suite.
3. Si $Z_C = \infty$ (circuit ouvert) alors $Z' = Z_0/j \tan(kl)$, capacitive pour $kl < \pi/2$, puis inductive, et ainsi de suite.

Notez que « inductive » n'est pas synonyme de bobine discrète parce que $Z_0 \tan(kl) = Z_0 \tan(\omega l / v_{\text{phase}})$ n'est pas proportionnelle à ω (sauf si l est faible, auquel cas un tronçon de ligne en court-circuit est une bonne approximation d'une bobine discrète). De même, un petit tronçon de ligne en circuit ouvert est une bonne approximation d'un condensateur discret.

Grâce à l'équation 8.7, nous avons vu comment une ligne de transmission, de longueur l , transformait une impédance Z en impédance Z' . Notez que l'impédance transformée dépend de la longueur du câble (à moins que $Z = Z_0$). Nous devons parfois effectuer une mesure d'impédance via un câble. Il est, par exemple, malcommode de mesurer l'impédance d'une antenne à son point d'alimentation qui peut se situer assez haut dans les airs ; cette mesure se fait généralement au travers de sa ligne d'alimentation. L'impédance de l'antenne Z_{ant} , fonction de l'impédance mesurée Z' est donnée par la relation :

$$Z_{\text{ant}} = Z_0 [Z' - jZ_0 \tan(kl)] / [Z_0 - jZ' \tan(kl)] \quad (8.8)$$

Cette méthode inverse s'apparente à celle dont on se sert pour déduire les caractéristiques d'un transistor après avoir effectué sur un montage un certain nombre de mesures.

Problèmes

8.1 Un câble coaxial classique de 50Ω , le RG214, possède une capacité parallèle de 101 pF/m . Calculer la valeur de l'inductance en série et la vitesse de propagation.

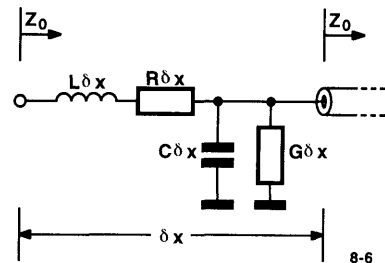
8.2a Servez-vous de la formule « tan-tan » pour démontrer qu'un court tronçon (δx) de ligne de transmission, en circuit ouvert à sa plus lointaine extrémité, se comporte comme un condensateur, c'est-à-dire qu'il possède une susceptance positive directement proportionnelle à la fréquence. Exprimez sa capacité en terme de capacité/unité de longueur de câble. *Astuces* : pour les faibles valeurs angulaires, $\tan \theta \approx \theta$.

8.2b Montrez qu'un court tronçon (δx) de ligne de transmission, en court-circuit à sa plus lointaine extrémité, se comporte comme une bobine, c'est-à-dire qu'il possède une susceptance négative inversement proportionnelle à la fréquence. Exprimez son inductance en terme d'inductance/unité de longueur de câble.

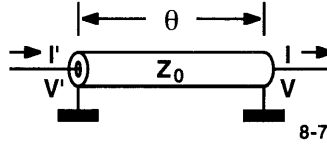
8.3a Trouvez la formule qui donne l'impédance caractéristique d'un câble à pertes dans lequel les pertes peuvent être dues à une résistance en série R par unité de longueur, aussi bien qu'à une conductance en parallèle G par unité de longueur (R représente la perte ohmique des conducteurs métalliques, tandis que G représente la perte diélectrique).

Astuces : vous pouvez généraliser le résultat obtenu pour un câble sans perte en remplaçant L par $L + R/j\omega$ et C par $C + G/j\omega$.

8.3b Trouvez la formule qui donne la constante de propagation k de ce câble à pertes. *Astuces* : appliquez les substitutions mentionnées ci-dessus à la formule $k = \omega \sqrt{LC}$. À quelle distance (exprimée en longueurs d'onde) la puissance d'un signal de fréquence ω_1 est-elle réduite d'un facteur $1/e$ si $R/\omega_1 = 0,01L$? (Vous supposerez que $G = 0$.)



8.4 Si l'on connaît la tension (sinusoïdale) U et le courant I à l'extrémité droite d'une ligne de transmission, trouvez la tension U' et l'intensité I' à son extrémité gauche.

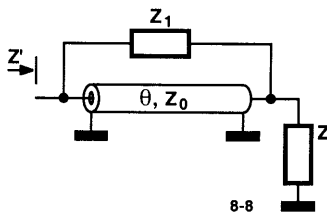


Astuces : supposez que la tension (complexe) de la ligne soit donnée par $U(\varphi) = U_V e^{-j\varphi} + U_R e^{j\varphi}$. Le courant est lié à la tension par la relation $Z_0 I(\varphi) = U_V e^{-j\varphi} - U_R e^{j\varphi}$. Admettons que $\varphi = 0$ à l'extrémité droite de la ligne de transmission. Montrez que $U_V = (U + IZ_0)/2$ et que $U_R = (U - IZ_0)/2$. Montrez ensuite qu'à l'extrémité gauche de la ligne de transmission où $\varphi = -\theta$, que $U' = U \cos \theta + IZ_0 j \sin \theta$ et $I' = I \cos \theta + j \sin \theta U / Z_0$.

8.5 Servez-vous des résultats obtenus au problème 8.4 pour mettre à jour votre programme d'analyse de réseau en échelle (reportez-vous au problème 1.3 du chapitre 1) pour traiter un autre type de circuit, une ligne de transmission en série. Trois paramètres sont nécessaires pour caractériser une ligne de transmission. Ce pourrait être l'impédance caractéristique, la longueur physique de la ligne et la vitesse de propagation. Cependant, pour simplifier les problèmes à venir, ces trois paramètres seront l'impédance caractéristique Z_0 , la longueur électrique θ_0 exprimée en degrés pour une fréquence donnée f_0 . Un câble de 50Ω qui a une longueur électrique de 80° à 10 MHz doit apparaître, dans le fichier des données du circuit, sous la forme « TL, 50, 80, 10E6 ». Pour une fréquence f , la longueur électrique est alors égale à $\theta = \theta_0 f / f_0$.

8.6 On connecte en parallèle une ligne de 50Ω et une ligne de 75Ω , c'est-à-dire qu'à chaque extrémité, les conducteurs internes et externes sont reliés entre eux. Les câbles ont la même vitesse de phase. Prouvez que l'impédance caractéristique de cette ligne de transmission composite est donnée par $(50 \times 75) / (50 + 75)$, c'est-à-dire que les lignes de transmission se comportent comme des résistances mises en parallèle.

8.7 Dans le schéma ci-dessous, l'impédance Z est modifiée par la mise en parallèle d'une ligne de transmission et d'une impédance discrète Z_1 (par exemple, R , C , L ou un réseau). Cherchez l'expression de Z' .



Astuces : commencez par trouver Z' lorsqu'il n'y a aucun élément discret mis en parallèle sur le câble. Supposez que des ondes directe et inverse, d'amplitude 1 et ρ , se propagent dans le câble. La tension en un point du câble est donnée par $U(x) = e^{j\omega x} (e^{-jkx} + \rho e^{jkx})$ et l'intensité par $I(x) = Z_0^{-1} e^{j\omega x} (e^{-jkx} - \rho e^{jkx})$. Le courant dans Z est la somme du courant circulant dans le câble et du courant circulant dans Z_1 .

9. Adaptation d'impédance 2

Impédances spécifiées par leur coefficient de réflexion

Pour savoir comment se modifie une impédance après un tronçon de ligne de transmission d'impédance caractéristique Z_0 , nous avons tout d'abord trouvé le coefficient de réflexion ρ à l'extrémité de la ligne où est connectée cette impédance. Le coefficient de réflexion est le rapport de la tension de l'onde réfléchie sur la tension de l'onde incidente. La loi d'Ohm nous dit que $\rho = (Z - Z_0)/(Z + Z_0)$. Nous pouvons considérer que ρ caractérise une impédance donnée et que l'expression précédente peut aussi s'écrire sous la forme :

$$Z = R + jX = Z_0 (1 + \rho)/(1 - \rho) \quad (9.1)$$

Dans le domaine des antennes ou des hyperfréquences, les connexions se font par des lignes de transmission et il est plus souple de se servir de ρ plutôt que de Z et de travailler dans le plan complexe ρ . Le coefficient de réflexion d'une impédance donnée, située après un tronçon de ligne de transmission de longueur l , est donné par :

$$\rho' = \rho e^{-j2kl} \quad (9.2)$$

C'est facile à vérifier. Lorsqu'on ajoute un tronçon de câble, la phase de l'onde incidente est retardée de kl à l'aller et l'onde réfléchie est également retardée de kl au retour. Le câble a donc pour effet de faire tourner dans le sens des aiguilles d'une montre le nombre complexe ρ d'une quantité égale à $2kl$ pour obtenir ρ' , le coefficient de réflexion modifié. Puisque $e^{j\omega t}$ tourne dans le sens inverse des aiguilles d'une montre, un retard temporel représente un déplacement dans le sens des aiguilles d'une montre. Dans la plupart des cas, on travaille sur ρ sans avoir à revenir à l'expression $Z = R + jX$.

Les équations que nous avons vues précédemment mettent en évidence l'aspect injectif de la relation qui existe entre un point du plan ρ et ses valeurs correspondantes R et X . Notez que le plan ρ est un plan complexe mais qu'il ne s'agit pas du plan $R + jX$. Examinons certains points particuliers du plan ρ :

1. Le centre du plan ($\rho = 0$) correspond à l'absence d'onde réfléchie, ce point représente donc l'impédance $Z_0 + j0$.
2. L'amplitude de ρ doit être inférieure ou égale à 1 lorsque l'impédance est passive (sinon l'énergie de l'onde réfléchie serait supérieure à celle de l'onde incidente).
3. Le point $\rho = -1 + j0$ correspond à $Z = 0$ (court-circuit).
4. Le point $\rho = 1 + j0$ correspond à $Z = \infty$ (circuit ouvert).
5. Les points du cercle $|\rho| = 1$ correspondent à des réactances pures. Tous les points situés à l'intérieur de ce cercle ont des impédances dont la composante R n'est pas nulle.
6. Le point $\rho = 0 + j1$ correspond à l'inductance $Z = 0 + jZ_0$. Tous les points situés dans le demi-plan supérieur sont « inductifs », c'est-à-dire du type $Z = R + j|X|$ ou encore $Y = G - j|B|$.
7. Le point $\rho = 0 - j1$ correspond à la capacité $Z = 0 - jZ_0$. Tous les points situés dans le demi-plan inférieur sont « capacitifs », c'est-à-dire du type $Z = R - j|X|$ ou encore $Y = G + j|B|$.

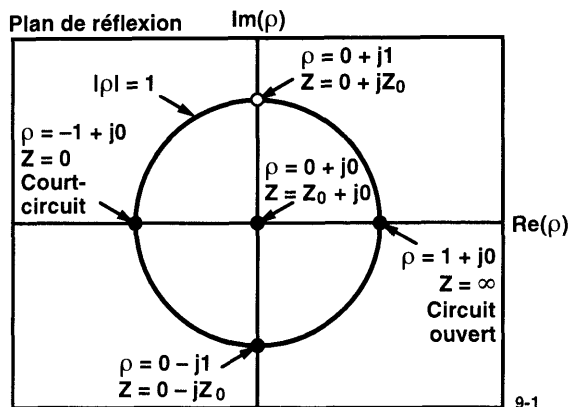


Figure 9-1. Répartition des impédances dans le plan de réflexion.

La figure 9-1 reprend tous ces cas particuliers d'impédances représentées dans le plan p . Si vous reportez dans ce plan les points de la courbe $\rho(R, X)$ où R est constante et X varie, vous obtenez un cercle, dont le centre se situe sur l'axe des réels et tangent à la ligne $\text{Re}(\rho) = 1$. À chaque valeur de R correspond un de ces « cercles de résistance ». Le « cercle de résistance » pour $R = 0$ est le cercle unité du plan p . Le « cercle de résistance » pour $R = \infty$ est un cercle dont le rayon est nul situé au point $\rho = 1 + j0$. De même, si vous reportez les points de la courbe $\rho(R, X)$ où X est constante et R varie, vous obtenez les « cercles de réactance », dont les centres se trouvent sur la ligne $\text{Re}(\rho) = 1$ et tangents à la ligne $\text{Im}(\rho) = 0$. La figure 9-2 représente de tels cercles.

Si vous ajustez à présent ces cercles pour ne conserver que ce qui se trouve à l'intérieur du cercle $|\rho| = 1$ (correspondant aux impédances passives, c'est-à-dire aux impédances dont la partie réelle est positive), vous obtenez la fameuse abaque de Smith représentée à la figure 9-3.

Nous verrons que l'utilisation d'un tracé dans le plan de réflexion (abaque de Smith) permet de suivre les modifications que subit une impédance au cours d'opérations successives, comme l'ajout en série d'un tronçon de ligne de transmission, d'une bobine L ou d'un condensateur C . Commençons tout d'abord par remarquer que nous aurions pu tout aussi bien avoir tracé p en fonction de G et B . Comme vous le devinez, nous aurions obtenu les « cercles G » et les « cercles B » ; c'est ce que représente la figure 9-4. Si nous incluons ces cercles dans le graphique, nous pourrions traiter les éléments montés aussi bien en série qu'en parallèle. La figure 9-5 illustre l'abaque de Smith que l'on obtient : elle est très « épanouie » et très dense.

La conception d'un réseau d'adaptation devient un simple exercice qui consiste à passer dans le plan de réflexion d'une valeur donnée ρ à une valeur souhaitée ρ' . L'emploi de l'outil graphique permet de trouver la stratégie qui permettra de concevoir le circuit d'adaptation le plus simple. On peut lire directement sur l'abaque la valeur des composants (c'est une sorte de calculatrice), mais on préfère aujourd'hui se faire assister par un ordinateur. Cependant, la représentation graphique des impédances permet de mieux visualiser un problème, elle est également plus pédagogique. Parcourez des catalogues de transistors dédiés à la HF ou aux hyperfréquences, vous y découvrirez que les impédances d'entrée et de sortie sont tracées dans le plan de réflexion, c'est-à-dire sur une abaque de Smith.

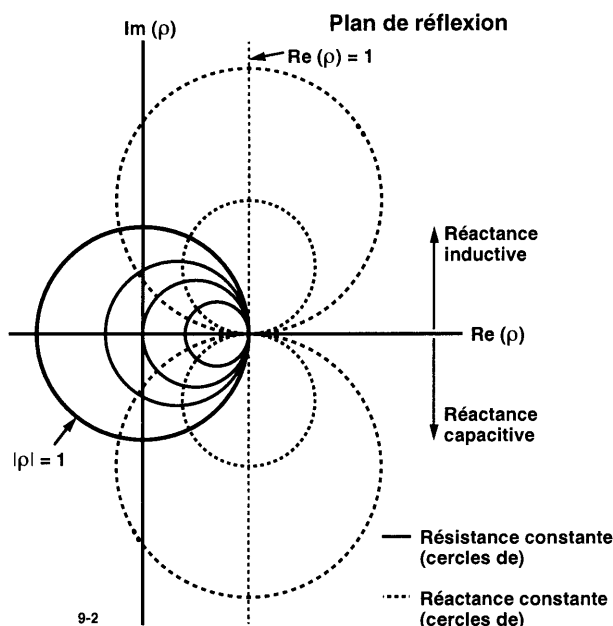


Figure 9-2. Lieux géométriques des points de résistance constante et de réactance constante : ce sont des cercles dans le plan ρ .

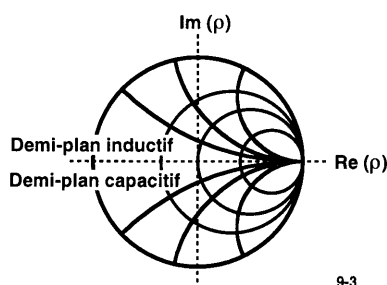


Figure 9-3. Abaque de Smith : elle comprend les cercles de résistance et de réactance situés à l'intérieur du cercle $\rho = 1$.

L'abaque de Smith est parfois mise à l'échelle pour une impédance Z_0 spécifique (comme 50 ou 75 Ω). Les autres tracés sont alors étalonnés ; par exemple, le cercle $R = 1$ devient le cercle 50 Ω si nous utilisons un câble 50 Ω . Souvenez-vous que l'abaque de Smith étant simplement la représentation du plan de réflexion complexe, les impédances à R constante ou X constante se représentent sous la forme de cercles. Ceci est aussi valable pour les lieux des points à B constante et G constante. Nous avons vu qu'il était particulièrement facile de trouver, dans le plan de réflexion, le nouveau coefficient de réflexion ρ' qui résulte de l'ajout d'un tronçon de ligne de transmission en un point dont l'impédance Z correspond à ρ . Voici comment l'abaque de Smith nous permet de passer de Z à ρ , puis à ρ' et enfin à Z' :

1. Connaissant $Z = R + jX$ (ou $Y = G + jB$), servez-vous des cercles de résistance et de réactance (ou de conductance et de susceptance) pour localiser ρ sur l'abaque.

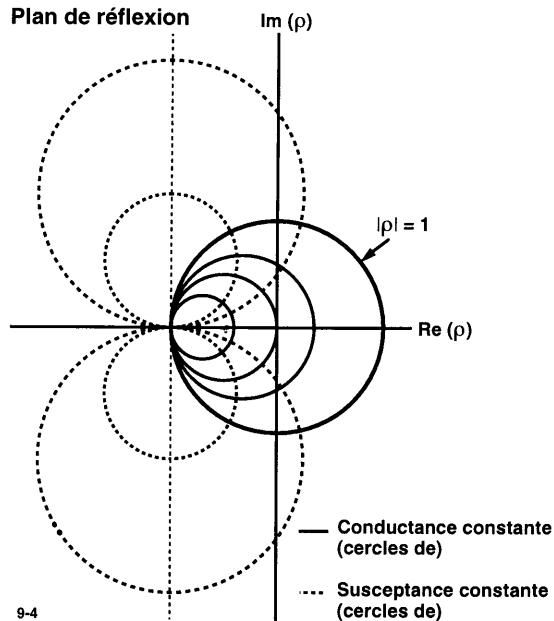


Figure 9-4. Lieux géométriques des points de conductance et de susceptance constantes (les cercles sont plus nombreux).

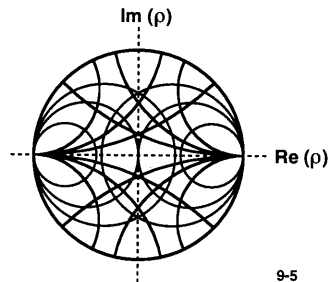


Figure 9-5. Abaque de Smith représentant les cercles R , X , G et B .

2. Faites tourner ce point dans le sens des aiguilles d'une montre autour de l'origine de deux fois la longueur de phase de la ligne pour trouver le point ρ' .
3. Utilisez les cercles pour lire R' et X' (ou G' et B').

Reprenons le problème d'adaptation $1000\ \Omega$ vers $50\ \Omega$ que nous avons vu au chapitre 2 et concevons plusieurs réseaux d'adaptation en nous aidant de l'abaque de Smith.

- I L'impédance que l'on veut transformer ($1000\ \Omega$) et l'impédance que l'on souhaite obtenir ($50\ \Omega$) sont indiquées sur l'abaque de la figure 9-6. On y aperçoit également le cercle $R = 50\ \Omega$. Nous pouvons utiliser une ligne de transmission pour atteindre le cercle de $50\ \Omega$. Cette ligne nous permet de nous déplacer sur le cercle en traits pointillés. Nous avons à présent $R = 50\ \Omega$, mais X est capacitive. Une bobine en série annulera la réactance capacitive, nous ramenant ainsi à $Z = 50 + j0$ (au centre de l'abaque).

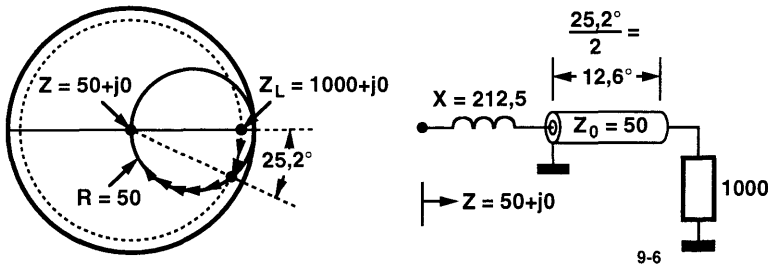


Figure 9-6. Une ligne de transmission et une bobine réalisent l'adaptation de 1000Ω vers 50Ω .

II Une autre solution consiste à se servir d'un tronçon de câble plus long pour parcourir la plus grande partie de l'abaque et couper le cercle de 50Ω dans la moitié supérieure du plan de réflexion (voir la figure 9-7). À cet endroit, nous obtenons $Z = 50 + jX$, où X est positive (donc inductive). Nous pouvons adjoindre un condensateur en série pour annuler X et obtenir à nouveau $Z = 50 + j0$.

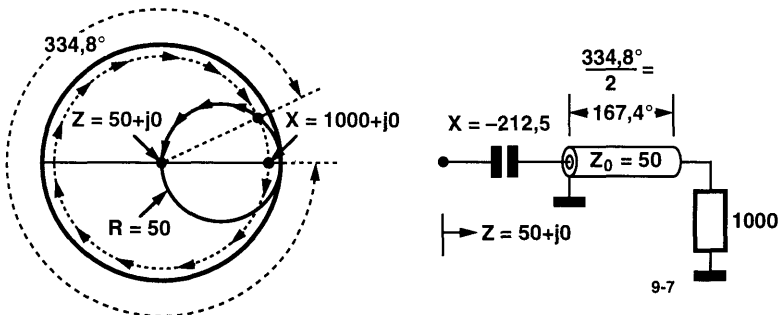


Figure 9-7. Une ligne de transmission et un condensateur réalisent l'adaptation de 1000Ω vers 50Ω .

III Jusqu'à présent, nous n'avons utilisé que des éléments en série. Naviguons à présent sur le cercle $G = 1/50$. Nous pouvons alors monter un élément en parallèle pour atteindre le centre de l'abaque. La première intersection avec le cercle $G = 1/50$ intervient dans le demi-plan inférieur (capacitif) ; nous devons alors mettre une bobine en parallèle. Nous pourrions remplacer la bobine discrète par une ligne de transmission mise en court-circuit, comme le montre la figure 9-8, pour réaliser un circuit d'adaptation ne comportant que des tronçons de ligne de transmission.

IV Voici une solution (voir la figure 9-9) qui ne met en œuvre aucune ligne de transmission. Servons-nous du cercle $G = 1/1000$. Nous nous déplaçons sur ce cercle si nous plaçons une réactance en parallèle. Choisissons une bobine en parallèle qui va nous faire parcourir vers le haut le cercle G en direction du cercle de 50Ω . Nous obtenons à présent $R = 50 \Omega$, mais il subsiste une réactance inductive. Comme dans l'exemple précédent, nous pouvons annuler cette réactance inductive à l'aide d'un condensateur en série. Nous retrouvons alors notre réseau en L originel.

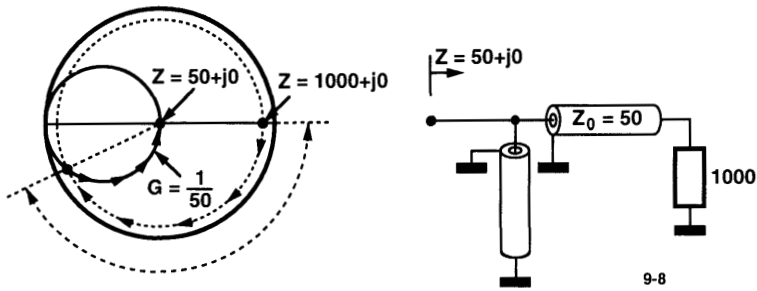


Figure 9-8. Des lignes de transmission mises en série et en parallèle réalisent l'adaptation de 1000 Ω vers 50 Ω .

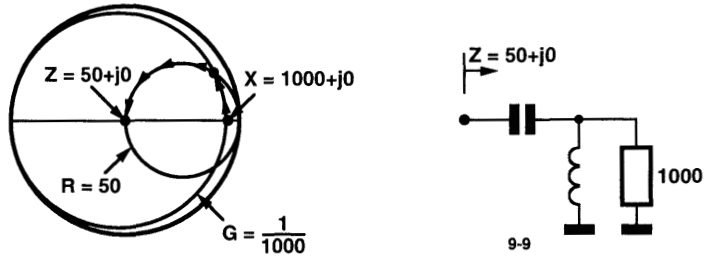


Figure 9-9. Un réseau LC réalise l'adaptation de 1000 Ω vers 50 Ω .

- V Si nous nous étions servi d'un condensateur en parallèle en lieu et place d'une bobine en parallèle, nous nous serions déplacés sur le cercle de 50 Ω vers le bas (voir la figure 9-10). Il est possible d'annuler la capacité en série résiduelle à l'aide d'une bobine. On obtient un réseau en L où les positions de L et C ont été permutes.

Dans les exemples que nous avons étudiés, l'impédance finale se trouvait au centre de l'abaque ($Z = 50 + j0$), mais vous constatez que ces techniques nous permettent de transformer tout point de l'abaque (c'est-à-dire toute impédance) en n'importe quel autre point de l'abaque (donc en n'importe quelle autre impédance).

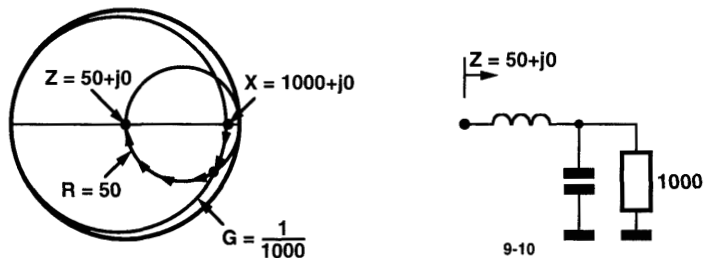


Figure 9-10. Un réseau CL effectue l'adaptation de 1000 Ω vers 50 Ω .

L'abaque de Smith est très utilisée du fait qu'elle s'applique aux réseaux en échelle réalisés aussi bien avec des lignes de transmission qu'avec des composants passifs. Si les lignes de transmission ne nous intéressent pas, les graphiques qui transforment R, X en G, B nous concernent. Par exemple, prenez le plan R, X (en fait le demi-plan, puisque nous excluons les R négatifs). Tracez les courbes pour G et B constants. L'abaque résultante, représentée à la figure 9-9, sert à concevoir des réseaux en échelle à composants discrets L, C et R tels que ceux des figures 9-9 et 9-10.

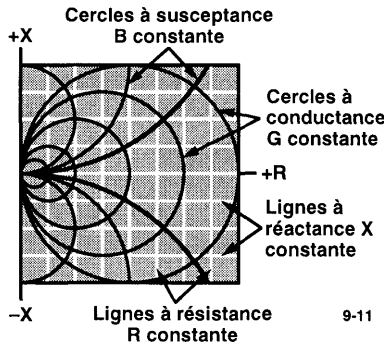


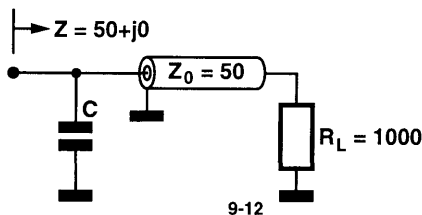
Figure 9-11. Abaque impédance-admittance.

Problèmes

9.1a En se servant d'un analyseur de réseau 50Ω , on trouve qu'un circuit donné, testé à 1 GHz), a un coefficient de réflexion (complexe) de 0,6 pour un angle de -22° (convention des coordonnées polaires : l'axe des x positifs correspond à 0° , et les angles croissent dans le sens inverse des aiguilles d'une montre). Calculer l'impédance $R + jX$.

9.1b Déterminez les valeurs des composants des circuits équivalents (au montage testé) en série $R_s C_s$ et en parallèle $R_p C_p$, à 1 GHz.

9.2a Le circuit représenté ci-dessous adapte, à une fréquence de 10 MHz, une charge de 1000Ω à une source de 50Ω . L'impédance caractéristique du câble est de 50Ω . Représentez sur une abaque de Smith le fonctionnement du montage.



9.2b Trouvez la longueur du tronçon de câble (la plus courte) et la valeur du condensateur. Exprimez la longueur en degrés et la capacité en pF. Calculez ces valeurs plutôt que de les lire sur une abaque de Smith dessinée avec précision.

9.2c Utilisez votre programme d'analyse de réseau en échelle (voir le problème 3 du chapitre 1 et le problème 5 du chapitre 8) pour trouver la transmission de 9 à 11 MHz avec un incrément de 0,1 MHz.

9.3 Déterminez les caractéristiques d'un tronçon de ligne de transmission qui remplacerait le condensateur du problème 9.2.

9.4 Supposez qu'une ligne de transmission possède une petite susceptance en parallèle (capacitive ou inductive) en un point x . Cela se traduit par une légère réflexion. Si l'on place une réactance identique en parallèle, située un quart de longueur d'onde avant la première, sa réflexion compensera la première et le câble présentera une transmission parfaite. Démontrez-le, (a) de manière analytique en vous servant de la formule « tan-tan » pour Z' , et (b) en utilisant l'abaque de Smith (la zone située autour du centre de l'abaque).

10. Alimentations

Lorsqu'on parle d'« alimentation », on pense immédiatement à un ensemble constitué d'un transformateur alimenté par le secteur à 50 Hz (ou 60 Hz aux États-Unis), d'un redresseur et d'une cellule de filtrage. Les alimentations que nous allons étudier sont effectivement de ce type, mais nous verrons par la suite que l'on trouve ces mêmes composants de base dans les alimentations à découpage et les modulateurs AM à découpage et que nous pourrions employer les mêmes méthodes d'analyse.

Redresseur à deux alternances

Il est certain que le circuit le plus simple combine un redresseur à deux alternances et un filtre à bobine en tête. De même que dans les chapitres précédents, une bobine parfaite a une inductance infinie et aucune résistance ohmique. La figure 10-1 en fournit deux exemples. La tension au point « V », pour l'un ou l'autre circuit, est représentée à la figure 10-2.

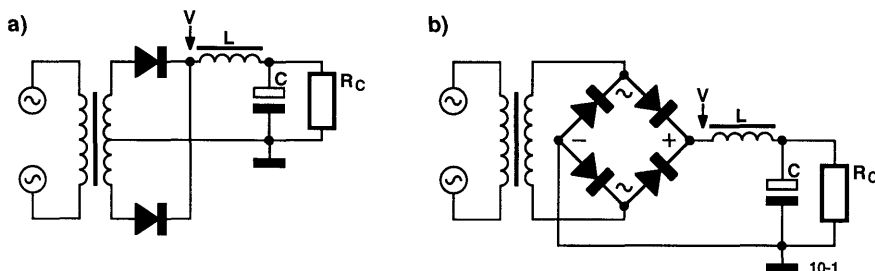


Figure 10-1. Alimentations à bobine en tête :
(a) transformateur avec secondaire à point milieu et, (b) redresseur en pont.

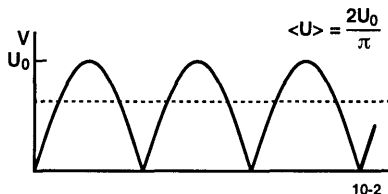


Figure 10-2. Forme d'onde de la tension à l'entrée du filtre.

Le circuit de redressement fournit l'énergie lors des deux alternances de la forme d'onde alternative d'entrée, d'où l'expression à *deux alternances*. Le redresseur en pont (figure 10-1(b)) ne requiert pas de transformateur avec secondaire à point milieu : le bobinage du secondaire est mieux exploité puisqu'il est utilisé pendant les deux alternances du cycle. Par contre, le circuit de redressement à deux diodes ne

présente qu'une chute de tension moitié dans les diodes et c'est pour cela qu'on le préfère pour réaliser des alimentations à basses tensions (par exemple 5 V, voire moins). Le rendement est assez élevé puisque la puissance dissipée dans les diodes est bien plus faible que celle qui est fournie à la charge. Dans ces deux circuits, on peut également placer, la bobine dans le retour à la masse ; c'est en fait ce que l'on fait pour réduire la contrainte galvanique sur la bobine.

Autorégulation d'une alimentation à bobine en tête

Le mot *régulation* est employé ici dans le sens de stabilité de la tension fournie lorsque la charge R_C varie. Nous pourrions également étudier la stabilité en fonction des variations de la tension alternative d'entrée. Il est important de se rendre compte que le circuit à bobine en tête est autorégulé tant que le courant dans la bobine reste positif. Voici pourquoi : si le courant dans la bobine L est toujours positif, l'un des redresseurs est polarisé dans le sens direct et la chute de tension à ses bornes est quasiment nulle. L'extrémité gauche de L est donc toujours reliée au secondaire du transformateur, et la tension de sortie, représentée à la figure 10-2, sera formée d'une succession d'alternances dont la tension varie entre zéro et la tension de crête du secondaire. Dans la mesure où la source de tension est parfaite (d'impédance interne nulle), cette forme d'onde est indépendante du courant dans la charge. Sa valeur moyenne est égale à $2U_0/\pi$ (où U_0 est la tension de crête). Puisqu'il n'y a aucune tension continue aux bornes de la bobine (idéale), la composante continue dans la charge R_C est aussi égale à $2U_0/\pi$ et la régulation est parfaite si l'on fait abstraction des petites pertes ohmiques dans les redresseurs, les enroulements du transformateur et la bobine, ainsi que de l'inductance de fuite du transformateur.

Voyons ce qu'il faut faire pour maintenir en permanence un courant positif dans la bobine. Le filtre LC en L supprime l'ondulation, de sorte que la tension alternative à l'extrémité droite de la bobine est négligeable vis-à-vis de la tension alternative présente à son extrémité gauche. La majeure partie de la tension alternative présente à son extrémité gauche est la composante à 100 Hz de la tension issue du redresseur. Si la tension de crête du redresseur est U_0 , alors la tension de crête de la composante à 100 Hz est égale à $4U_0/(3\pi) = 0,42U_0$. Le courant alternatif de crête dans la bobine est donné par $4U_0/(3\pi\omega L)$. Pour éviter que le courant n'atteigne le niveau le plus bas, il faut que le courant continu soit supérieur à ce courant alternatif de crête. Nous devons donc satisfaire à cette relation $2U_0/(\pi R_C) > 4U_0/(3\pi\omega L)$ ou $L > 2R_C/(3\omega)$. L'inductance minimale est connue sous le nom d'*inductance critique*. Dans certaines applications, la charge peut varier de telle sorte que le courant atteigne parfois une valeur très faible, voire nulle. On peut maintenir un courant minimal en plaçant un résistor (dénommé « résistor stabilisateur ») en parallèle avec la charge. Ce résistor dissipera malheureusement de l'énergie. On trouve fréquemment un résistor stabilisateur dans les alimentations à haute tension pour des raisons de sécurité : il permet au condensateur de filtrage de se décharger lorsqu'on arrête l'alimentation. Pour maintenir le courant constant, une autre méthode consiste à utiliser une bobine « saturable » qui présente une inductance élevée lorsque cela est nécessaire (c'est-à-dire pour un courant faible) et une inductance plus faible, mais de valeur encore suffisante, pour un courant fort (où le noyau se sature).

Ondulation

Les composants L et C forment un simple filtre de lissage en L . Dans l'alimentation monophasée à deux alternances que nous avons évoquée précédemment, la valeur de crête de la composante à 100 Hz, à l'entrée du filtre (extrémité gauche de la bobine), est égale à $0,42U_0$. La réactance du condensateur sera en principe bien inférieure à la résistance de la charge R_C ; nous obtiendrons donc la tension de l'ondulation en considérant que L et C se comportent comme un diviseur de tension :

$$U_{\text{ondulation (crête)}} = \frac{0,42 U_0 X_C}{X_L - X_C}$$

On utilise quelquefois un circuit résonnant LC , comme on le voit à la figure 10-3, pour mieux éliminer la fréquence de l'ondulation résiduelle primaire. On peut filtrer les harmoniques plus élevées en plaçant à la suite un réseau classique LC .

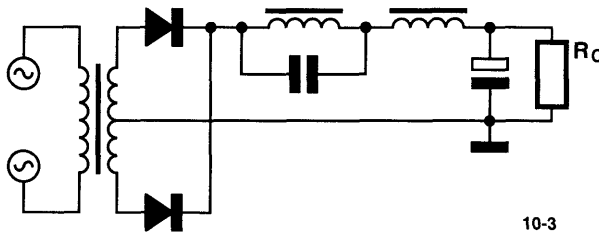


Figure 10-3. Alimentation équipée d'un filtre résonnant.

Redresseur à une alternance

Le circuit de redressement le plus simple ne comporte qu'une seule diode (voir la figure 10-4). Il est évident que ce montage n'exploite pas au mieux le secondaire du transformateur ; de plus la composante alternative est bien plus élevée ici que dans le circuit de redressement à deux alternances. Qui plus est, la fréquence fondamentale de la composante alternative est à 50 Hz au lieu d'être à 100 Hz, ce qui implique un filtrage LC plus musclé pour une ondulation de sortie donnée.

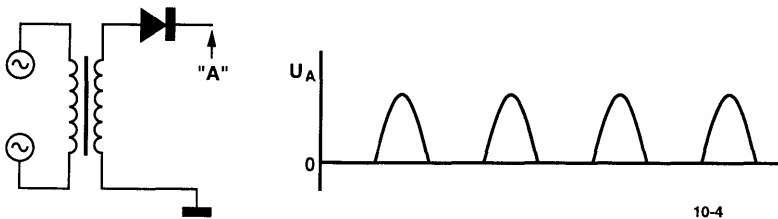


Figure 10-4. Circuit de redressement élémentaire à une alternance.

Une bonne régulation requiert un filtre à bobine en tête, mais aucun courant positif ne circulera en permanence dans une bobine reliée à une seule diode polarisée en inverse la moitié du temps. Une solution consiste à ajouter une diode de « roue libre » (que l'on rencontre dans toute alimentation à découpage). La figure 10-5 illustre ce montage à deux diodes. Lorsque la tension au point haut du secondaire devient

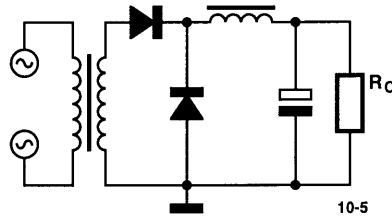


Figure 10-5. Un circuit de redressement à une alternance peut également disposer d'une régulation à bobine en tête si l'on inclut une diode de roue libre.

négative et « déconnecte » la diode du haut, la diode de roue libre continue à fournir du courant à la bobine. La tension présente à l'extrémité gauche de la bobine est donc toujours positive ou nulle et le circuit bénéficie alors de l'autorégulation d'une alimentation à bobine en tête à deux alternances.

Les alimentations les plus simples mettent en œuvre un filtre RC (sans bobine). Le résistor dissipe de l'énergie et la régulation laisse à désirer ; la tension moyenne en sortie du redresseur dépend de la charge. Le résistor R limite le courant pour protéger la diode. Ce résistor est d'ailleurs souvent absent parce que la perte ohmique et l'inductance de fuite du transformateur le remplacent.

Alimentation régulée électroniquement

Si une alimentation de base, comme l'une de celles que nous avons étudiées précédemment, est suivie en série d'un « transistor ballast », son circuit de retour peut alors commander ce transistor ballast pour maintenir une tension de sortie constante. La figure 10-6 représente un tel système. Vous remarquerez que le transistor ballast doit pouvoir dissiper de la puissance et qu'un tel circuit de régulation n'est ni plus ni moins qu'un amplificateur en classe A. On peut remplacer le résistor stabilisateur par un transistor en parallèle pour réaliser une alimentation à régulation en parallèle.

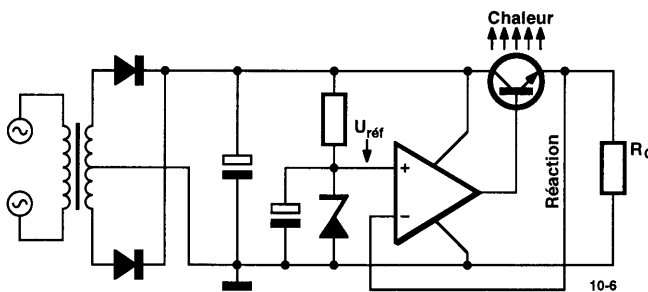


Figure 10-6. Alimentation régulée électroniquement.

Les transformateurs et les bobines sont des composants lourds et onéreux, aussi sont-ils peu employés en électronique grand public. Dans un téléviseur, par exemple, l'alimentation principale, souvent construite autour d'un simple redresseur à une alternance ou à deux alternances alimenté directement par le secteur, délivre plus d'une centaine de volts. Un tel « châssis chaud » impose obligatoirement la présence d'un transformateur d'isolement de rapport 1:1. L'amplificateur horizontal (le sous-ensemble qui consomme le

plus) peut facilement comporter un transistor haute tension. Si le sous-ensemble devait fonctionner à une tension plus basse, il faudrait dissiper de la puissance dans une résistance chutrice ou remplacer le simple redresseur à une alternance par un redresseur à découpage (nous y reviendrons). Les autres sous-ensembles d'un téléviseur fonctionnent sous des tensions d'alimentation inférieures. Plutôt que de les mettre en série et d'alimenter la chaîne ainsi réalisée, on prélève les diverses tensions d'alimentation plus faibles sur la sortie de l'amplificateur horizontal qui, comme nous le verrons par la suite, fonctionne comme un amplificateur et une alimentation à découpage.

Redresseur triphasé

Les grosses alimentations de puissance, de même que les gros moteurs, s'alimentent à partir du réseau triphasé qui permet une meilleure transmission d'énergie et l'emploi de transformateurs plus légers. Examinons quelques circuits usuels.

La figure 10-7 représente une alimentation triphasée à une alternance. Elle offre déjà une amélioration par rapport à l'alimentation monophasée à deux alternances. La fréquence de l'ondulation de la tension à l'entrée de la bobine est plus élevée (150 Hz au lieu de 100 Hz) et une ondulation bien plus faible (mais jamais nulle). On peut utiliser des filtres à bobine en tête sans avoir à rajouter de diode de roue libre.

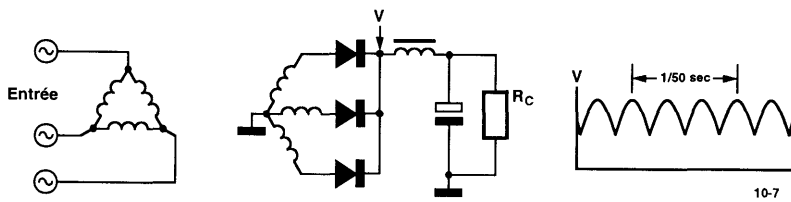


Figure 10-7. Alimentation triphasée à une alternance.

La figure 10-8 représente le schéma d'un redresseur triphasé à deux alternances. Ce circuit emploie des enroulements secondaires à point milieu. L'ondulation de la tension présente à l'entrée de la bobine est à présent très faible et la fréquence fondamentale de l'ondulation est égale à 300 Hz (les composants L et C sont donc plus faibles). Notez qu'il s'agit en fait d'un circuit hexaphasé ; il y a un redresseur par phase.

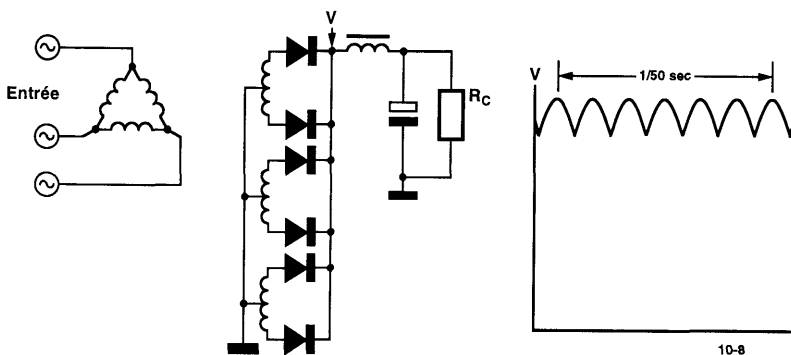


Figure 10-8. Alimentation triphasée à deux alternances.

Un schéma d'alimentation triphasée à deux alternances est représenté à la figure 10-9. Il existe la même relation entre ce redresseur en pont et le circuit précédent qu'entre le redresseur à deux alternances à quatre diodes et le circuit à deux alternances classique à deux diodes. Il y a encore six phases équivalentes. Ce montage ressemble à celui que l'on rencontre dans les alternateurs d'automobiles.

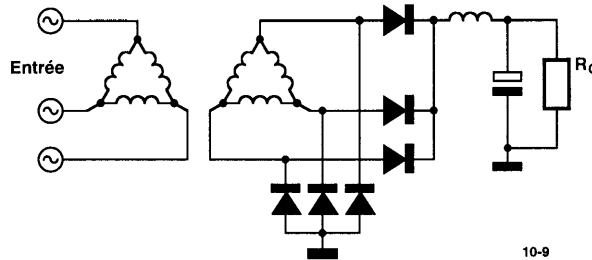
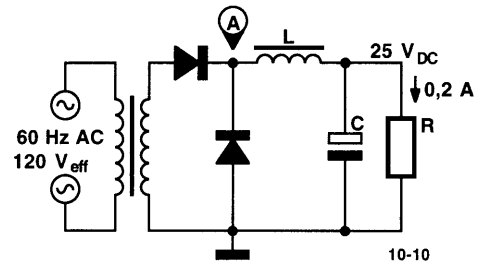


Figure 10-9. Alimentation triphasée à deux alternances.

Problèmes

10.1a L'alimentation à une alternance représentée ci-dessous délivre à une charge (représentée sous la forme d'une résistance) une tension continue de 25 V sous 0,2 A. Par hypothèse, les composants sont parfaits et la capacité du condensateur est suffisamment élevée pour considérer que la tension alternative en sortie (ondulation résiduelle) est très faible. En supposant que l'inductance de la bobine est assez élevée pour maintenir un courant positif, représentez la forme d'onde de la tension au point « A ».



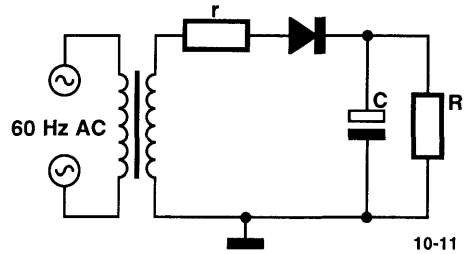
10.1b La valeur de la tension moyenne au point « A » doit être égale à 25 V puisqu'il n'y a aucune chute de tension continue aux bornes de la bobine. Trouvez la tension efficace $U_0/\sqrt{2}$ présente aux bornes du secondaire du transformateur. (Réponse : 55,5 V_{eff}.)

10.1c Calculez la valeur de crête de la composante alternative fondamentale (à 60 Hz) de la tension au point « A », c'est-à-dire trouver son amplitude dans la décomposition en série de Fourier. Si cela vous pose un problème, estimez l'amplitude de cette composante à partir de votre dessin. (Réponse : 39,3 V_{crête}.)

10.1d Servez-vous des résultats obtenus (problème 10-1c) pour trouver la valeur de l'inductance critique de la bobine, c'est-à-dire l'inductance minimale pour que le courant traversant la bobine soit toujours positif. (Réponse : 0,52 H.)

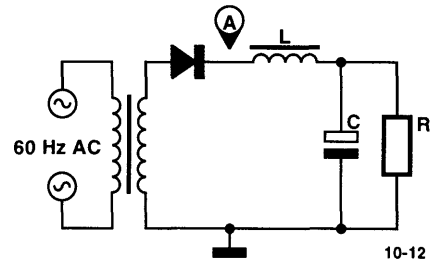
10.1e Quelle doit être la valeur du condensateur pour que la tension de l'ondulation résiduelle de sortie à 60 Hz soit inférieure à 100 mV_{eff} ? (Réponse : 3700 μF.)

10.2 Examinez ci-contre le schéma de l'alimentation classique à une alternance à condensateur en tête. Les résistances r et R représentent respectivement la résistance ohmique de la diode et la résistance de la charge. Par hypothèse, le transformateur est parfait et la capacité du condensateur C est suffisamment élevée pour considérer que la tension alternative en sortie (ondulation résiduelle) est négligeable. Vous supposerez que la tension de crête au secondaire du transformateur est de 10 V et que les résistances r et R sont respectivement égales à 1 et 20 Ω . Calculez la tension de sortie U_{DC} .



Astuces : le courant moyen dans la diode est égal au courant dans la charge : $U_{DC} / 20$. Le courant ne circule dans la diode que lorsque la tension située à l'extrémité haute du secondaire du transformateur est supérieure à U_{DC} ; la valeur instantanée du courant est égale à $(10 \cos \theta - U_{DC}) / 1$. Procédez à l'intégration de ce courant sur un cycle pour trouver sa valeur moyenne qui est égale à $U_{DC} / 20$. Vous aurez besoin d'une calculatrice pour trouver l'angle pour lequel le courant devient nul. (Réponse : 7,51 V.)

10.3a La figure ci-contre donne le schéma d'une alimentation à bobine en tête à une alternance peu commune. Ce circuit, sans diode de roue libre, n'a pas l'autorégulation du montage du problème 10.1. Cependant, si la charge reste constante, la régulation est sans importance et ce montage permet d'abaisser la tension (sans perte de puissance) en augmentant l'inductance de la bobine. Vous supposerez encore que la capacité du condensateur C est suffisamment élevée pour considérer que la tension alternative en sortie (ondulation résiduelle) est négligeable.



Dessinez la forme de la tension au point « A » ainsi que celle du courant dans la bobine.

Astuces : rappelez-vous que $U = L di/dt$. Cela vous aidera peut-être de considérer le cas où R est très élevée et où R est très faible.

10.3b Si la tension (à 60 Hz) au secondaire du transformateur est égale à $10 \sin(\omega t)$ et si la tension de sortie est de 5 V, déterminez le rapport cyclique de la diode, c'est-à-dire l'angle où la diode devient conductrice (polarisée dans le sens direct) et l'angle où la diode se bloque (où le courant dans la bobine devient nul).

10.3c Trouvez la valeur de L lorsque le courant de sortie est de 1 A.

11. Modulation d'amplitude

Moduler signifie ajouter une information à une onde porteuse sinusoïdale pure en faisant varier son amplitude ou sa phase (ou bien les deux). La modulation d'amplitude (ou AM, de *Amplitude Modulation*) la plus simple qui soit est la manipulation par tout ou rien. Cette AM à deux états s'obtient simplement en insérant un interrupteur (le manipulateur télégraphique) en série avec la source d'énergie. Les premières transmissions de la parole se faisaient à l'aide d'un microphone à grenaille de charbon, utilisé comme résistance variable placé en série avec l'antenne. L'AM sert dans les bandes de radiodiffusion¹ à grande ondes, ondes moyennes et ondes courtes.

Analyse de l'AM dans le domaine temporel

En l'absence de modulation (sans musique ni parole), la tension et le courant appliqués à l'antenne sont des ondes sinusoïdales pures à la fréquence de la porteuse. La puissance nominale d'une station est définie comme étant la puissance de sortie de l'émetteur en absence de modulation. La présence d'un signal à basse fréquence (BF) modifie l'amplitude de la porteuse. La figure 11-1 représente un émetteur et un récepteur AM théoriques. Le signal BF (tension du microphone amplifiée) présente des excursions de tension positives et négatives, mais sa valeur moyenne est nulle. Imaginez que la tension BF soit comprise entre $+U_m$ et $-U_m$. On ajoute une tension de polarisation continue de U_m volts à la tension BF. La somme $U_m + U_{\text{audio}}$, toujours positive, est multipliée par l'onde de la porteuse, $\sin(\omega_p t)$. Le produit résultant est le signal AM ; l'amplitude de l'onde sinusoïdale HF est proportionnelle au signal BF polarisé. La simulation représentée à la figure 11-2 montre les différentes formes d'ondes présentes dans l'émetteur et dans le récepteur ; elles correspondent à un segment quelconque d'une forme d'onde BF. On appelle enveloppe de modulation la forme d'onde BF polarisée. Vous constatez que lorsque la modulation est totale, c'est-à-dire pour $U_{\text{audio}} = +U_m$, la porteuse est multipliée par $2U_m$, alors que sans modulation, la porteuse est multipliée par U_m (tension de polarisation). Ce rapport de deux dans l'amplitude signifie qu'un signal totalement modulé (à 100%) a une puissance de crête quatre fois supérieure à celle d'un signal non modulé (porteuse seule). Cela implique que l'antenne prévue pour un émetteur AM de 50 kW doit pouvoir encaisser sans aucun problème des pointes de 200 kW. La puissance moyenne d'un signal modulé dépend du carré de la moyenne de l'enveloppe de modulation. Par exemple, lorsqu'un signal BF pur module à 100% une porteuse, la puissance moyenne du signal modulé est égale à celle de la porteuse multipliée par un facteur $(1 + \cos\theta)^2 = 3/2$.

¹ La bande des grandes ondes (GO ou LW de *Long-Wave*) s'étend de 153 à 179 kHz. La bande des ondes moyennes (OM ou MW de *Middle-Wave*) s'étend de 520 à 1700 kHz. Les bandes des ondes courtes (OC ou SW de *Short-Wave*) sont généralement spécifiées par leurs longueurs d'ondes : 75 m, 60 m, 49 m, 41 m, 31 m, 25 m, 19 m, 16 m, 13 m et 11 m. L'espacement en fréquence entre les stations à grandes ondes est de 9 kHz. L'espacement en fréquence entre les stations à ondes moyennes est de 10 kHz dans les pays occidentaux et de 9 kHz dans le reste du monde. Les attributions des fréquences pour les ondes courtes sont moins bien coordonnées, mais la quasi-totalité des stations à ondes courtes émettent sur des fréquences multiples du kilohertz.

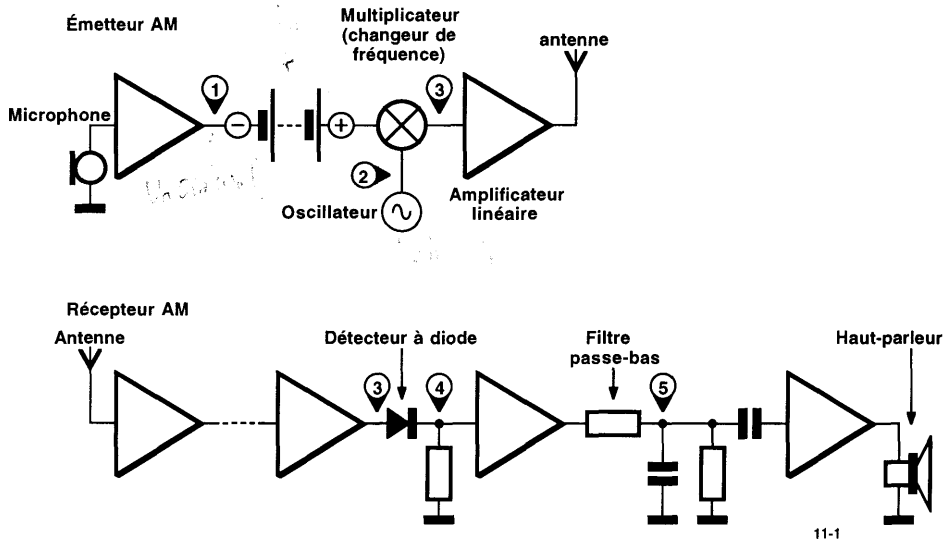


Figure 11-1. Émetteur et récepteur AM théoriques.

La figure 11-1 illustre également la façon dont le récepteur démodule le signal, c'est-à-dire, comment il détecte l'enveloppe de modulation. Le détecteur est une simple diode de redressement qui élimine les portions négatives du signal HF modulé. Un filtre RC passe-bas fournit alors la tension moyenne des portions sinusoïdales positives. Leur tension moyenne est égale à leur tension de crête multipliée par $1/\pi$; elle est donc proportionnelle à celle de l'enveloppe. Enfin une liaison en courant alternatif élimine la polarisation et l'on obtient un signal BF identique au signal issu du microphone.

Analyse de l'AM dans le domaine fréquentiel

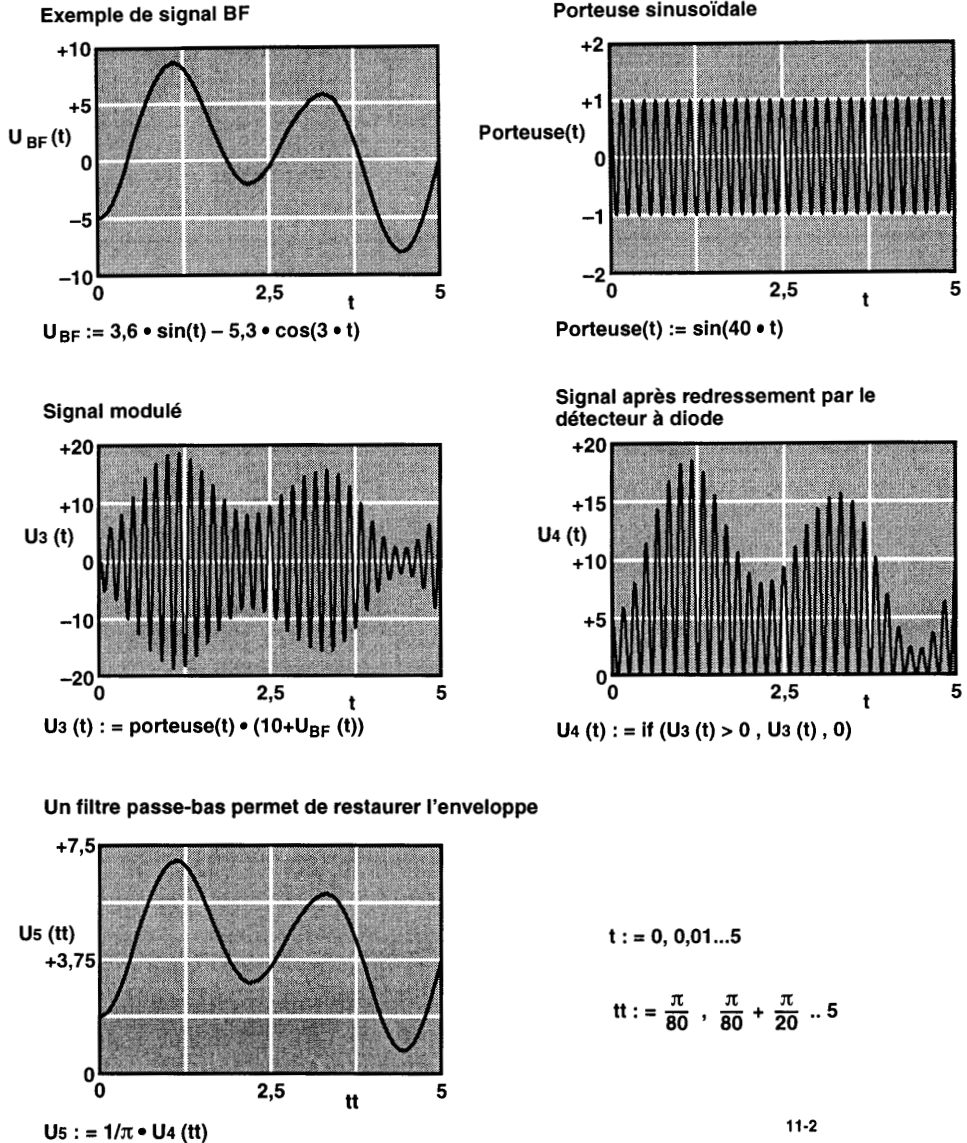
Étudions le spectre d'un signal AM, c'est-à-dire, la façon dont se répartit l'énergie en fonction de la fréquence. C'est simple à effectuer lorsque le signal BF est une onde sinusoïdale pure, de la forme $U_a \sin(\omega_a t)$. La tension présente au point 3 de la figure 11-1 est alors égale à

$$\sin(\omega_p t) [U_m + U_a \sin(\omega_a t)]$$

$$= U_m \sin(\omega_p t) + \frac{1}{2} U_a \cos[(\omega_p - \omega_a)t] - \frac{1}{2} U_a \cos[(\omega_p + \omega_a)t]$$

Ces trois termes représentent respectivement la porteuse de fréquence ω_p , avec une puissance $U_m^2/2$, une bande inférieure de fréquence $\omega_p - \omega_a$ avec une puissance $U_a^2/8$ et une bande supérieure de fréquence $\omega_p + \omega_a$, avec une puissance aussi égale à $U_a^2/8$. La figure 11-3 représente ce spectre.

Vous noterez que la porteuse, de fréquence ω_p est toujours présente et que son amplitude est indépendante du taux de modulation. Pour une modulation à 100% par un signal sinusoïdal, $U_a = U_m$ et la puissance moyenne de chaque bande est le quart de celle de la porteuse. La puissance moyenne maximale est donc une fois et demi celle de la porteuse, c'est ce que nous avons déduit de l'analyse dans le domaine temporel. Les bandes latérales, qui acheminent la totalité de l'information, ne prennent que le tiers de la puissance



11-2

Figure 11-2. Voici les formes d'ondes présentes dans le système AM de la figure 11-1.

totale émise. Lorsque le signal de modulation est du type vocal, la puissance des bandes latérales peut même être plus faible. *Question* : peut-on éliminer la porteuse pour économiser de l'énergie ? Un signal BF typique a des composantes fréquentielles au moins égales à 10 kHz. Chaque composante produit deux bandes latérales situées de part et d'autre de la porteuse. Les stations de radiodiffusion AM limitent leurs signaux BF à 10 kHz, mais la largeur du spectre est de 20 kHz. L'espacement entre les stations de radiodiffusion AM n'est que de 10 kHz. C'est pourquoi, pour éviter tout chevauchement dans une zone de réception donnée, on n'attribue pas à deux stations la même fréquence ni deux fréquences adjacentes.

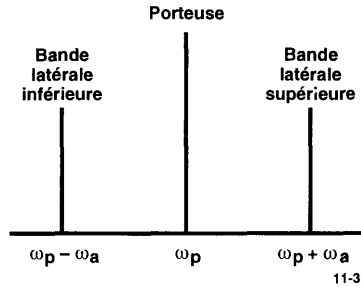


Figure 11-3. Spectre d'un signal AM modulé par un signal BF pur.

Modulation à haut niveau

Le schéma de l'émetteur simple de la figure 11-1 peut parfaitement être mis en pratique. Il faudra évidemment le modifier légèrement pour éliminer la batterie qui fournit la tension de polarisation. L'amplificateur de puissance linéaire HF représente le plus gros sous-ensemble de l'émetteur. Son rendement sera médiocre (voir le problème 11.4) à cause de la présence permanente de la porteuse. Nous avons vu précédemment que les amplificateurs en classe C et en classe D, bien que non linéaires, ont un excellent rendement pour tous les types de signaux ; ils conviennent parfaitement à l'AM parce que leur amplitude de sortie HF est égale à la tension d'alimentation continue appliquée. Ces amplificateurs sont aussi équivalents à des multiplicateurs qui exécutent le produit de la tension d'alimentation par une onde sinusoïdale d'amplitude unitaire à une fréquence HF. En *modulation à haut niveau*, un signal BF polarisé de puissance alimente un amplificateur final en classe C ou en classe D. Examinons plusieurs circuits propres à la modulation à haut niveau.

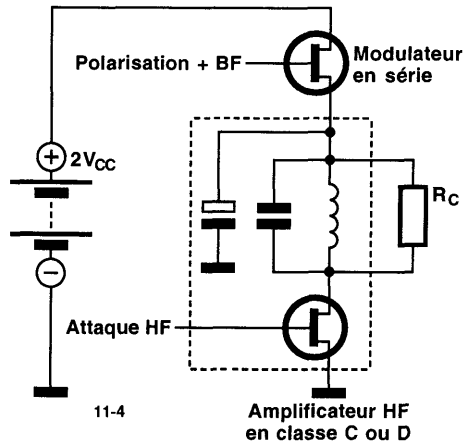


Figure 11-4. Modulateur AM en série de classe A.

Modulateur en classe A

La figure 11-4 donne le schéma d'un modulateur simple à haut niveau² ; il s'agit d'un modulateur en série en classe A. Le « transistor ballast » en série fonctionne en mode suiveur ; le montage est fondamentalement linéaire – ce qui n'est pas négligeable. Qu'en est-il du rendement ? En absence de signal BF (seulement la porteuse), le transistor ballast conduit toujours partiellement et alimente donc l'amplificateur sous V_{cc} . Supposons, pour simplifier, que le rendement de l'amplificateur HF soit de 100%, ce qui est quasiment le cas de tout amplificateur en classe D. En absence de modulation (où il n'y a que la porteuse), la tension d'alimentation totale, $2V_{cc}$, se partage de façon égale entre le transistor en série et l'amplificateur. Les courants étant les mêmes, la puissance dissipée sous forme de chaleur par le transistor ballast est égale à la puissance de sortie HF. Le rendement du modulateur n'est donc que de 50%. Il passe à 75% pour un taux de modulation de 100% (voir le problème 11.5). Ce modulateur est tout simplement trop onéreux pour être mis en œuvre dans une station de radiodiffusion très puissante.

Modulateur en classe B

Un émetteur classique AM utilise un modulateur en classe B dans lequel la tension BF fournie par un amplificateur haute puissance en classe B est ajoutée à la polarisation continue. Dans le schéma de la figure 11-5, la polarisation continue se fait au travers du secondaire du transformateur de modulation. La tension BF produite par le modulateur apparaît au secondaire et s'ajoute à la tension de polarisation.

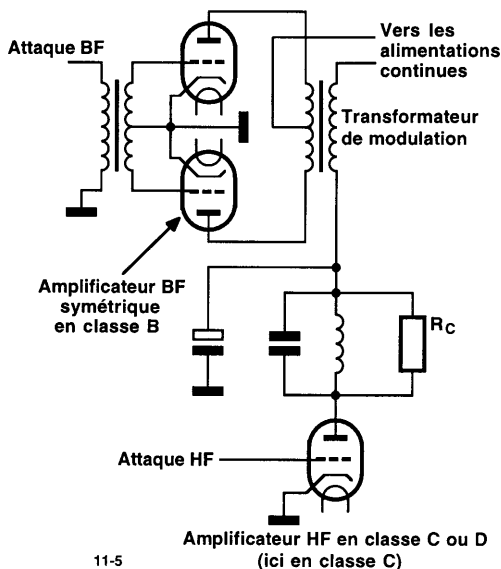


Figure 11-5. Émetteur AM modulé par la plaque en classe B.

- 2 En AM, l'expression *modulateur à haut niveau* fait référence à l'amplificateur BF à haute puissance utilisé pour produire la tension d'alimentation variable destinée à l'amplificateur HF ; mais c'est en réalité l'amplificateur en classe C ou en classe D qui procède à la modulation, c'est-à-dire à la multiplication effective de l'onde sinusoïdale HF par la tension BF (polarisée).

En l'absence de tout signal BF, la consommation de l'amplificateur BF en classe B est négligeable et l'alimentation de la tension de polarisation fournit l'énergie de la porteuse. Pour une modulation à 100% par une onde sinusoïdale pure, le modulateur doit fournir une puissance BF égale à la moitié de celle de l'alimentation de polarisation. Pour obtenir une modulation à haut niveau, un émetteur de 50 kW requiert ainsi un amplificateur BF d'une puissance au moins égale à 25 kW. Ce résultat est valable non seulement pour un signal pur, mais également pour la parole. Le rendement de la porteuse seule est celui d'un amplificateur isolé. Le rendement global d'un émetteur AM utilisant une modulation à haut niveau (avec une modulation de 100%) est donné par :

$$\eta = \frac{P_{\text{sortie}}}{P_{\text{DC}}} = 1,5 \frac{\eta_{\text{HF}}}{(1 + 0,5/\eta_{\text{mod}})} \quad (11.1)$$

où P_{sortie} est la puissance de sortie HF, P_{DC} la puissance fournie par l'alimentation continue et η_{HF} et η_{mod} , respectivement les rendements de l'amplificateur HF et du modulateur. En ce qui concerne l'émetteur de la figure 11-5, nous aurions $\eta_{\text{HF}} = 80\%$ et $\eta_{\text{mod}} = 78\%$, pour un rendement global égal à 73%.

Modulateur en classe S

De nouvelles méthodes ont été essayées pour produire un signal AM avec un rendement encore meilleur. Elles utilisent toutes des techniques de commutation. Le modulateur de largeur d'impulsion, qu'illustre la figure 11-6, est une variante à haut rendement du modulateur en série de la figure 11-4. Il emploie un circuit de régulation à découpage (dénommé parfois amplificateur en classe S) pour fournir la tension d'alimentation V_{cc} à un amplificateur en classe C ou en classe D. Nous étudierons ultérieurement ce genre de circuit avec les alimentations à découpage. Son fonctionnement est identique à celui d'une alimentation à bobine en tête : la tension présente à l'extrémité gauche de la bobine est toujours égale à $2V_{\text{cc}}$ (lorsque le tube est à l'état passant) ou nulle (lorsque le tube est à l'état bloqué). Le taux d'utilisation du tube de commutation détermine la tension moyenne présente à l'entrée de la bobine. La bobine et le condensateur de filtrage éliminent la forme d'onde du découpage dû au tube de commutation. Un circuit à bas niveau produit un train d'impulsions destinées à attaquer le commutateur. Le rapport cyclique du train d'impulsions est proportionnel à la tension BF polarisée.

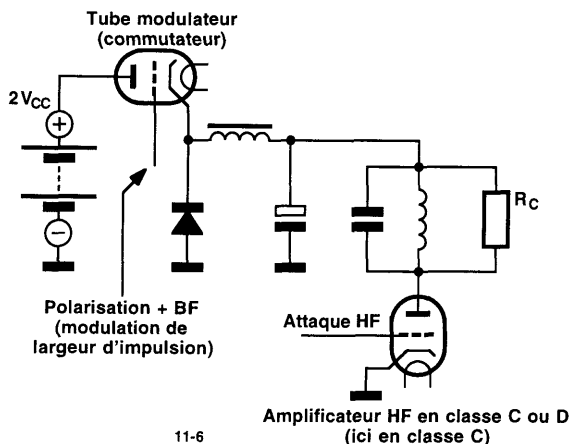


Figure 11-6. Modulateur AM en classe S (à découpage).

Modulateur numérique/analogique

Il est probable que le modulateur à haut niveau qui possède le meilleur rendement est le convertisseur numérique/analogique à grande puissance. Il utilise des interrupteurs à semiconducteurs pour additionner entre elles les tensions des nombreuses alimentations individuelles à basse tension plutôt que des tubes ou des transistors qui feraient chuter la tension d'une seule alimentation à haute tension. Le rendement d'un tel modulateur est supérieur à 95% ; les sociétés Continental Electronics et Brown Boveri en fabriquent. La figure 11-7 fournit le schéma de principe simplifié de l'un de ces modulateurs de type numérique/analogique. On trouve à l'entrée BF un convertisseur analogique/numérique. Ce convertisseur n'est plus nécessaire si le signal BF existe déjà sous une forme numérique. Une autre conception de modulateur consiste à combiner les sorties en signaux carrés de nombreux amplificateurs en classe D. On active un nombre suffisant d'amplificateurs en tout point de la forme d'onde BF pour obtenir l'amplitude de sortie HF désirée. Il y a, dans un tel système, un circuit de commande, mais pas de véritable modulateur.

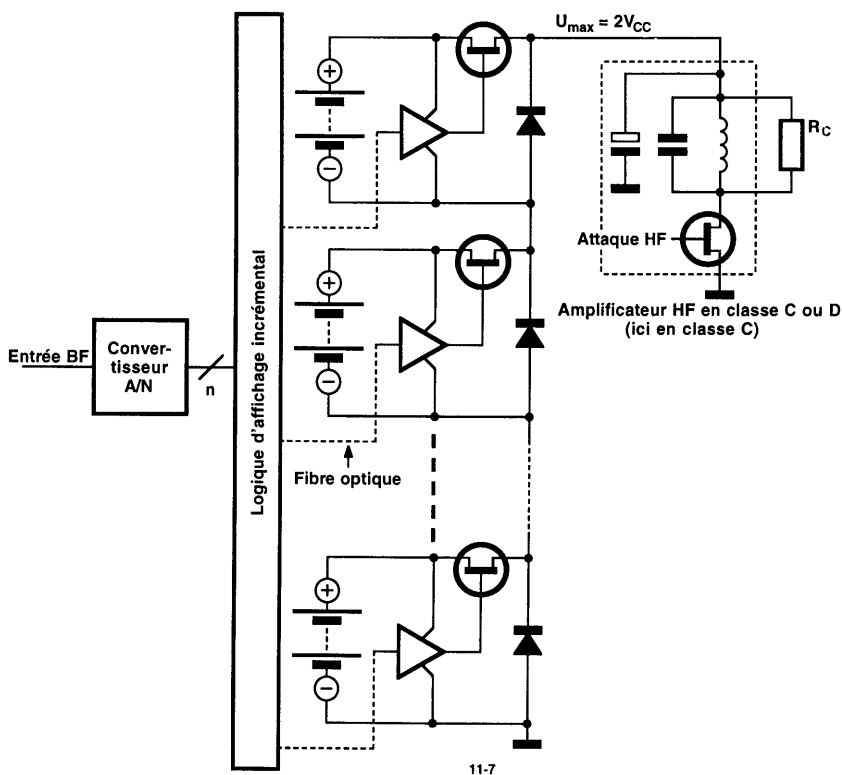


Figure 11-7. Modulateur AM numérique/analogique.

Ce qui se fait actuellement

La puissance des émetteurs de radiodiffusion AM américains est limitée à 50 kW (dans les années 1930, la station WLW, située à Cincinnati, réalisa un émetteur expérimental célèbre de 500 kW). On recense, dans le reste du monde, de nombreuses stations AM de puissance égale de 1 à 2 MW qui émettent en grandes ondes, ondes moyennes et ondes courtes. La totalité de ces stations de très grande puissance emploient des

tubes à vide, montés généralement en amplificateurs en classe C. Les stations AM commencent à utiliser des émetteurs entièrement transistorisés qui allient les puissances fournies par de nombreux amplificateurs d'environ 1 kW. Parmi les avantages offerts par cette solution, nous pouvons citer la mise en œuvre de tensions plus faibles (d'où une meilleure sécurité), la durée de vie quasi éternelle des transistors, alors qu'il faut remplacer tous les ans ou tous les deux ans des tubes chers et enfin une plus grande « disponibilité » ; la défaillance d'un module n'affecte que très légèrement la puissance de sortie et son remplacement peut s'effectuer en cours d'émission.

Problèmes

11.1 Vous avez très certainement observé quelqu'un qui, n'étant pas exactement calé sur la fréquence d'une station AM, entendait un son distordu. On constate ce défaut d'accord lorsque la personne présente une perte d'audition des fréquences élevées et/ou lorsque la bande passante de la radio est insuffisante. Que se passe-t-il ? Pourquoi la bande passante d'une radio serait-elle insuffisante et en conséquence, pourquoi certains auditeurs doivent-ils se décaler légèrement de la fréquence d'émission ?

11.2 Nous avons vu qu'un émetteur AM, modulé à 100% par un signal BF sinusoïdal, possède des bandes latérales dont la puissance totale est égale à la moitié de la puissance de la porteuse. Supposons à présent que ce soit un signal carré BF qui module la porteuse à 100%. Quel rapport existe-t-il entre la puissance d'une bande latérale et celle de la porteuse ?

11.3 Supposons que vous essayiez d'écouter une station AM lointaine et qu'une autre station émette sur la même fréquence avec une puissance quasiment identique. Entendrez-vous distinctement chacune des stations ou constaterez-vous une quelconque interférence entre ces deux stations ?

11.4 Dans l'émetteur de la figure 11-1, la puissance du signal modulé est faible. Un amplificateur linéaire HF délivre la haute puissance nécessaire. Nul n'est besoin de disposer d'un équipement de modulation à haute puissance (commutateurs à haute tension, transformateurs, etc.). Par contre, le rendement est faible, même avec un amplificateur en classe B, puisque l'amplificateur doit être capable de fournir une puissance de crête égale à quatre fois celle de la porteuse. Montrez que lorsque le niveau BF est nul (porteuse seule), le rendement d'un amplificateur en classe B n'est que de 39% ($\pi/8$).

11.5 Montrez que le rendement de l'émetteur de la figure 11-4 atteint 75% à condition qu'un signal BF pur le module à 100%. Vous supposerez que le rendement de l'amplificateur HF est de 100%.

11.6 Lorsque le niveau du signal BF est faible, l'amplitude d'un signal AM varie peu par rapport à celui de la porteuse. L'enveloppe de modulation, qui véhicule l'information, est très proche de la porteuse à haute puissance. On pourrait économiser de l'énergie si l'on pouvait réduire l'amplitude moyenne sans réduire celle de la modulation. Envisagez comment procéder au niveau de l'émetteur, et quelles seraient les conséquences, s'il y en a, au niveau du récepteur.

12. Modulation à porteuse supprimée

Le rendement de l'AM est généralement médiocre parce que la porteuse, qui ne véhicule aucune information, est gourmande en puissance. Pourquoi ne pas en faire l'économie en n'acheminant que les bandes latérales et en rajoutant dans le récepteur une porteuse produite localement ? Il est assez facile de modifier l'émetteur pour supprimer la porteuse. Examinez à nouveau le schéma de l'émetteur AM (voir la figure 11-1) que nous avons analysé précédemment. Remplacez la batterie par une tension nulle (un simple fil), la porteuse disparaît. Le signal qui en résulte s'appelle signal à *deux bandes latérales et porteuse supprimée* (ou DSBSC de *Double-Sideband Suppressed Carrier*). Nous pourrions tester le schéma du récepteur de la figure 12-1 pour restaurer la porteuse manquante. On ajoute simplement une porteuse produite localement (par l'oscillateur de battement, ou BFO, de *Beat Frequency Oscillator*) au signal à la fréquence intermédiaire (FI) immédiatement en amont du détecteur d'enveloppe. Dans un récepteur superhétérodyne, il suffit de régénérer cette porteuse à une seule fréquence, celle de la FI. Comme vous vous en doutez, la fréquence de la porteuse ajoutée doit être exacte. Sa phase doit également être correcte. Supposons par exemple que le signal de modulation soit un signal BF pur de 400 Hz. Si la porteuse reconstituée est déphasée de 45° , le signal en sortie du détecteur d'enveloppe sera très distordu. Si le déphasage atteint 90° , le signal en sortie du détecteur sera un signal de 800 Hz (soit une distorsion de 100%). La figure 12-2 illustre les formes d'onde de ces deux exemples.

Créer une porteuse de substitution avec une phase appropriée est assez simple pour peu que l'émetteur fournisse une porteuse « pilote » d'amplitude faible – mais suffisante pour que le récepteur s'en serve comme d'une référence de phase. Revenons à notre émetteur théorique et remplaçons la batterie de polarisation par une batterie de faible tension. Le récepteur peut alors isoler la porteuse pilote à l'aide d'un filtre passe-bande étroit, l'amplifier et enfin l'ajouter immédiatement en amont du détecteur d'enveloppe, comme nous l'avons vu à la figure 12-1. Nous pourrions réaliser un circuit équivalent pour produire la porteuse locale en utilisant une boucle à phase asservie (ou PLL de *Phase-Locked Loop*), dans laquelle le signal pilote serait le signal de référence.

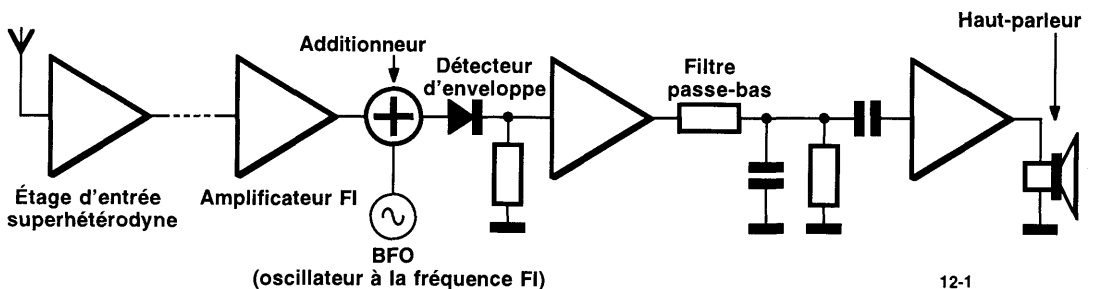


Figure 12-1. Récepteur DSBSC théorique : ajout de la porteuse restaurée et détection d'enveloppe.

La détection d'un signal AM à deux bandes latérales requiert une phase de porteuse correcte (démonstration) :
 $t := 0, 0,1, \text{ à } 150$.

porteuse(t) := sin(t)

BF(t) := sin(0,05 x t)

La BF est un signal sinusoïdal de fréquence égale à 1/20 de celle de la porteuse.

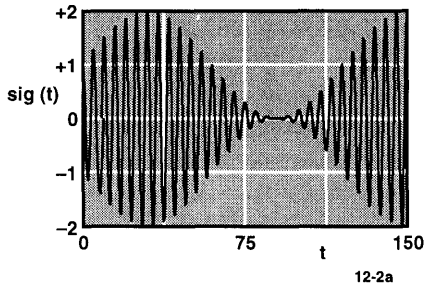
dssc(t) := BF(t) x porteuse(t)

Deux bandes latérales et porteuse supprimée (seules restent les bandes latérales).

dssc

Notez que DSBSC(t) = 1/2 (cos (0,95 t) – cos (1,05 t)).

(1 - 100%)

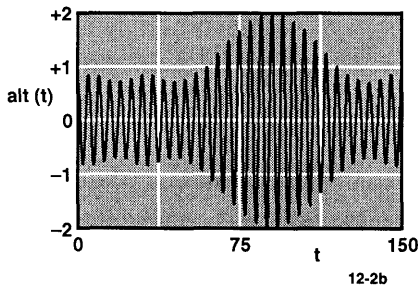


sig(t) := porteuse(t) + dssc(t)

dssc

Porteuse + bandes latérales = (1 + BF) (porteuse)
 = AM classique

Bandes latérales + porteuse correcte
 (notez que l'enveloppe est le signal sinusoïdal BF)

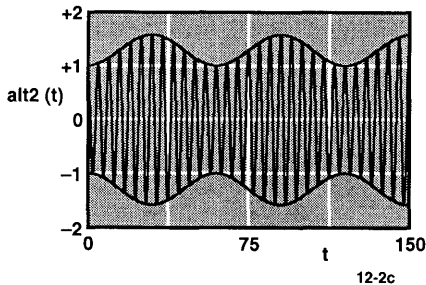


alt(t) := dssc(t) + cos(t + $\pi/4$)

dssc

Donnons à la porteuse réinjectée une erreur de phase de 45°

Bandes latérales avec déphasage de la porteuse de 45°
 (remarquez la distorsion de l'enveloppe)



alt2(t) := dssc(t) + cos(t)

dssc

Ici la porteuse réinjectée présente une erreur de phase de 90°

Bandes latérales avec déphasage de la porteuse de 90° (la fréquence de l'enveloppe est le double de celle de la BF, ce qui dénote une distorsion élevée)

Figure 12-2. Détection d'un signal DSBSC par ajout d'une porteuse locale.

Bande latérale unique

On peut (presque) éliminer le problème de phase du BFO si l'on ne transmet que l'une des bandes latérales. Chacune des bandes latérales est la copie conforme de l'autre : elles véhiculent la même information et nous n'avons pas vraiment besoin de leur présence simultanée. Nous étudierons ultérieurement trois façons d'éliminer l'une des bandes latérales, mais nous pouvons dire que la première méthode, également la plus évidente, consiste à utiliser un filtre passe-bande étroit centré sur la bande latérale désirée. Le signal qui en résulte est connu sous le nom de signal à bande latérale unique à porteuse supprimée (BLU ou SSBSC de *Single-Sideband Suppressed Carrier*) ou plus simplement bande latérale unique (BLU ou SSB de *Single-Sideband*). Reprenons l'exemple du signal BF pur de 400 Hz. L'émetteur BLU produira une fréquence unique, située à 400 Hz au-dessus de la fréquence de la porteuse (supprimée) si nous avons choisi la bande latérale supérieure. Le signal apparaîtra dans le récepteur à une fréquence de 400 Hz au-delà de la fréquence centrale de la FI. Supposons que nous nous servions encore d'un BFO et d'un détecteur d'enveloppe. L'enveloppe de l'ensemble signal plus BFO sera vraiment un signal sinusoïdal pur de 400 Hz (sans distorsion). Jusqu'ici tout va bien. Supposons maintenant que le signal soit composé de deux fréquences, l'une à 400 Hz et l'autre à 600 Hz. Lorsqu'on ajoute la porteuse issue du BFO, l'enveloppe résultante comprendra les fréquences à 400 Hz et à 600 Hz, mais il y aura également une composante à 200 Hz ($600 \text{ Hz} - 400 \text{ Hz}$) qui introduira certainement de la distorsion. Les composantes de la FI battent ensemble et produisent des composantes indésirables ; leur énergie est proportionnelle au produit des deux signaux FI. Par ailleurs, l'énergie de chacune des composantes désirées est proportionnelle au produit de leur amplitude par celle du BFO. Si l'amplitude du signal issu du BFO est nettement supérieure à celle du signal FI, la distorsion peut être réduite.

Détecteur-produit

Il est possible d'éliminer la distorsion mentionnée précédemment si l'on remplace le détecteur d'enveloppe par un détecteur-produit. Dans le récepteur représenté à la figure 12-3, le détecteur-produit est simplement le classique multiplicateur (alias le mélangeur). La sortie de ce détecteur est la somme des produits de chacune des composantes du signal FI par le BFO. Il ne crée pas les produits d'intermodulation des diverses composantes du signal FI. On appelle également le détecteur-produit un « mélangeur en bande de base » parce qu'il décale (vers le bas) les composantes de la bande latérale jusqu'à leurs fréquences BF originales.

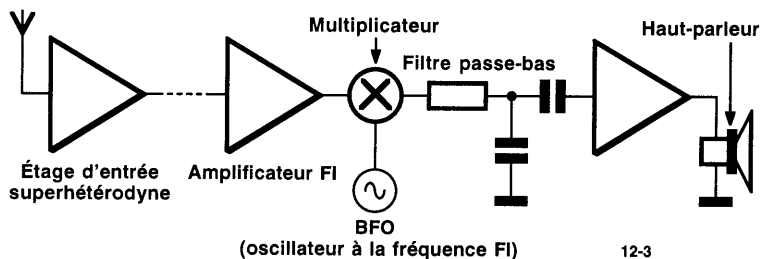


Figure 12-3. Récepteur SSBSC utilisant un détecteur-produit.

Bien que le détecteur-produit ne crée aucun produit d'intermodulation (et donc pas de distorsion d'intermodulation) dû aux composantes du signal FI, la phase de la porteuse injectée doit être exacte. Une phase erronée ne porte pas à conséquence dans le cas où nous n'utilisons qu'un signal à une seule fréquence. Mais lorsque le signal comporte de nombreuses composantes, leurs phases relatives jouent un rôle important. La forme d'onde sera distordue si toutes les composantes spectrales ont un déphasage identique (voir le problème 11.3). Par contre, elle ne sera pas distordue, mais seulement retardée dans le temps, si l'on donne à chaque composante un déphasage proportionnel à sa fréquence. Par conséquent, un émetteur BLU, de même qu'un émetteur à deux bandes latérales, doit vraiment transmettre une porteuse pilote sur laquelle va venir se verrouiller la phase du BFO du récepteur. C'est ce que font certains systèmes BLU. Mais lorsqu'on ne transmet que de la voix, il est fréquent de n'utiliser aucun pilote. La parole reste intelligible et presque naturelle, même quand la phase n'est pas correcte. La fréquence du BFO peut aussi être légèrement décalée. Lorsque la fréquence est trop élevée, le signal vocal démodulé de la bande latérale supérieure (BLS ou USB de *Upper Sideband*) a une hauteur de son plus basse que celle du signal vocal originel. Lorsqu'elle est trop basse, la hauteur de son est plus élevée que celle du signal vocal originel. Que se passe-t-il enfin si l'on essaie d'employer un détecteur-produit avec un BFO non piloté (non verrouillé en phase) pour recevoir un signal DSBSC ? Si le déphasage du BFO est de 90° , il n'y a aucun signal en sortie du détecteur. Si la phase du BFO est bonne, l'amplitude du signal de sortie sera maximale. L'amplitude sera réduite pour toute erreur de phase intermédiaire.

Autres avantages de la BLU

La BLU permet d'économiser la puissance de la porteuse mais ne requiert en plus qu'une bande passante deux fois plus faible : on peut donc placer, dans une bande passante déterminée, deux fois plus de canaux. Lorsqu'on réduit de moitié la bande passante du récepteur, on divise également par deux le bruit de fond ; on obtient donc une amélioration de 3 dB, par rapport à l'AM, du rapport signal/bruit. Lorsqu'un spectre de fréquences comporte de nombreux signaux AM, les porteuses de ces signaux produisent dans le récepteur des signaux de battement audibles gênants (toute porteuse située dans la bande passante du récepteur semble être une bande latérale appartenant au spectre du signal recherché). L'utilisation d'un émetteur BLU permet d'éliminer toute porteuse et donc tout signal de battement. Les radioamateurs, les militaires et les avions transocéaniques effectuent leurs communications vocales à ondes courtes (entre 1,8 et 30 MHz) en BLU.

Création d'un signal BLU

On connaît au moins trois méthodes pour produire un signal BLU.

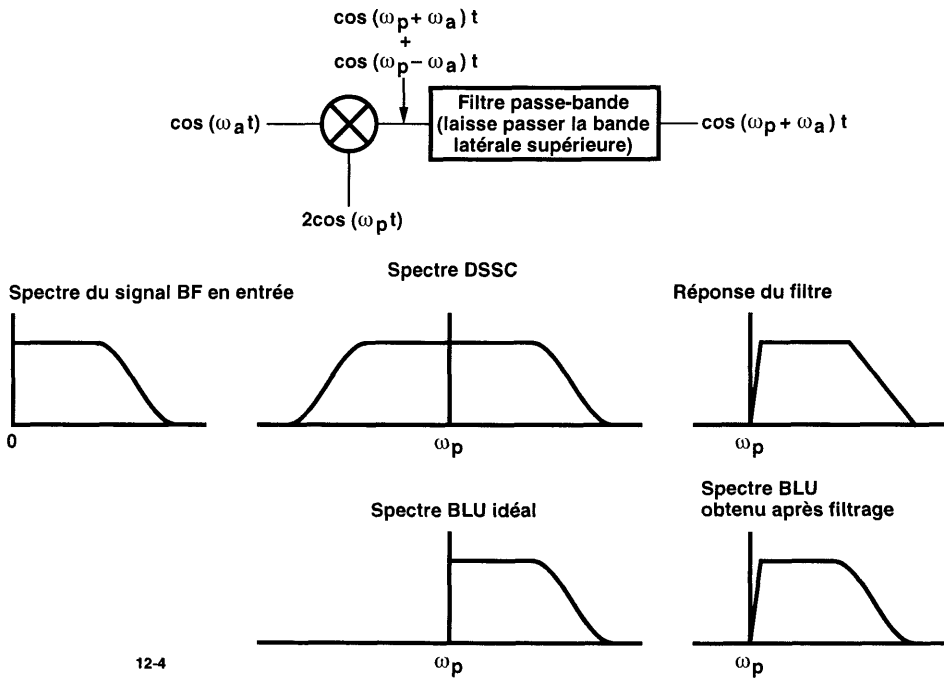
Méthode de filtrage

Dans la méthode de filtrage, illustrée par la figure 12-4, un filtre passe-bande à flancs raides élimine la bande latérale indésirable. Ce filtre est généralement composé de quartz ou autres résonateurs mécaniques à facteur Q élevé, et comprend de nombreux étages. Il n'est pas ajustable et le signal BLU est créé à une seule fréquence (à la fréquence FI d'un émetteur) puis mélangé à la fréquence recherchée. Ce filtre n'a pas des flancs verticaux et coupera certaines fréquences basses du signal BF.

Méthode de mise en phase

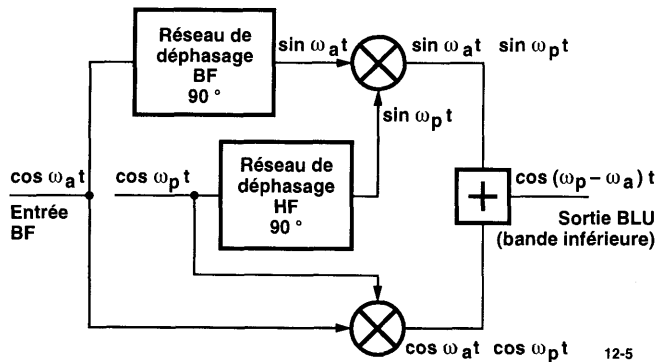
La méthode de mise en phase nécessite deux multiplicateurs et deux réseaux de déphasage afin de satisfaire à l'identité trigonométrique suivante :

$$\cos(\omega_p t) \cos(\omega_a t) \pm \sin(\omega_p t) \sin(\omega_a t) = \cos[(\omega_p \mp \omega_a) t] \quad (12.1)$$



12-4

Figure 12-4. Générateur BLU – méthode de filtrage.



12-5

Figure 12-5. Générateur BLU – méthode de mise en phase.

Les indices a et p se rapportent respectivement à la fréquence BF et à la fréquence de la porteuse (supprimée). Si $\cos(\omega_a t)$ représente le signal BF, nous devons créer le signal $\sin(\omega_a t)$ nécessaire dans le second terme. On peut réaliser, à l'aide de composants passifs, de techniques numériques de traitement du signal, etc., un réseau de déphasage BF de 90° . Il est généralement composé de deux réseaux installés en amont de chacun des mélangeurs. Leur *différence* de phase est proche de 90° sur toute la bande BF.

Le second terme nécessite également que nous créions $\cos(\omega_p t)$ mais, puisque nous connaissons ω_p , tout système qui produit un déphasage de 90° à cette fréquence convient. La figure 12-5 illustre la méthode de mise en phase. Si l'on remplace l'additionneur par un soustracteur (en inversant la polarité d'une entrée), on obtiendra la bande latérale supérieure au lieu d'avoir la bande latérale inférieure.

Méthode de Weaver

La troisième méthode (proposée par Weaver en 1956 [2]) ne requiert ni filtre passe-bande à flancs raides ni réseau de déphasage. La méthode de Weaver, représentée à la figure 12-6, emploie quatre multiplicateurs et deux filtres passe-bande. L'astuce consiste à mélanger le signal BF avec le signal issu d'un premier oscillateur dont la fréquence ω_0 se situe au milieu de la bande BF. Les sorties du premier ensemble de mélangeurs sont bien déphasées de 90° . Le second ensemble de mélangeurs fonctionne de la même façon que les deux mélangeurs utilisés dans la méthode de mise en phase.

Si l'on se reporte à la figure 12-6, on peut écrire :

$$V1 = \text{passe-bas} [\sin(\omega_a t) \sin(\omega_0 t)] = \cos [(\omega_a - \omega_0) t]/2 \quad (12.2)$$

$$V2 = \text{passe-bas} [\sin(\omega_a t) \cos(\omega_0 t)] = \sin [(\omega_a - \omega_0) t]/2 \quad (12.3)$$

$$V3 = \cos[(\omega_p - \omega_0) t] \cos [(\omega_a - \omega_0) t]/2 =$$

$$V3 = \cos[(\omega_a - \omega_p) t]/4 + \cos [(\omega_a + \omega_p - 2\omega_0) t]/4 \quad (12.4)$$

$$V4 = \sin[(\omega_p - \omega_0) t] \sin [(\omega_a - \omega_0) t]/2 =$$

$$V4 = \cos[(\omega_a - \omega_p) t]/4 - \cos [(\omega_a + \omega_p - 2\omega_0) t]/4 \quad (12.5)$$

$$V_{\text{sortie}} = V3 + V4 = \cos [(\omega_p - \omega_a) t]/2 \quad (12.6)$$

L'équation 12.6 nous montre que cette configuration permet de créer la bande latérale inférieure.

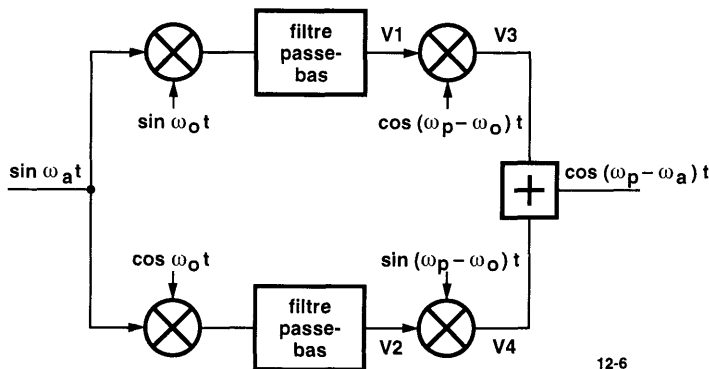


Figure 12-6. Générateur BLU – méthode de Weaver.

BLU avec amplificateurs en classe C ou en classe D

Les trois méthodes que nous venons d'étudier pour produire un signal BLU concernent des signaux de faible niveau. Des amplificateurs linéaires se chargent ensuite de produire la puissance nécessaire pour attaquer l'antenne. Pour créer la BLU, il existe toutefois un moyen d'utiliser un amplificateur en classe C ou en classe D, en modulant simultanément la phase et l'amplitude. Cela vient du fait que tout signal à bande étroite (onde entretenue, AM, FM, modulation de phase, BLU, DSBSC, bruit à bande étroite, etc.) est essentiellement un signal sinusoïdal à la fréquence centrale ω_0 , dont la phase et l'amplitude varient sur un laps de temps bien plus long que $1/\omega_0$. Imaginons que nous ayons produit un signal BLU de faible niveau. Il nous est possible de créer l'enveloppe, l'amplifier et employer l'enveloppe amplifiée pour moduler en AM un amplificateur en classe C ou en classe D. Nous pouvons simultanément moduler en phase l'amplificateur avec la phase du signal BLU de faible niveau. Le signal de faible niveau, simplement limité en amplitude, permet d'attaquer l'amplificateur modulé. Qu'en résulte-t-il ? Nous avons produit un signal BLU de grande puissance avec un amplificateur dont le rendement est voisin de 100%. De plus, il est possible de transformer un émetteur AM existant en émetteur BLU, moyennant certaines modifications mineures.

Bibliographie

- 1 W.E. Sabin et E.O. Schoenike (1987), « *Single-Sideband Systems and Circuits* », New York: McGraw-Hill.
- 2 D.K. Weaver, Jr. (1956), « *A third method of generation and detection of single-sideband signals* », *Proc. IRE* 1703-6.

Problèmes

12.1 Certaines sociétés de radiodiffusion en ondes courtes ont expérimenté la BLU. Quels avantages pourrait présenter une station de radiodiffusion en BLU sur une station de radiodiffusion en AM ? Y a-t-il des inconvénients ? Pensez-vous que dans le domaine de la radiodiffusion, la BLU puisse un jour remplacer l'AM ?

12.2 La BLU offrirait une durée de vie des batteries installées dans les émetteurs-récepteurs portables (téléphones cellulaires, etc...) bien plus grande que dans le cas où l'on utilise l'AM (ou la FM). Y a-t-il des inconvénients à employer la BLU dans ce type d'application ?

12.3 Montrez le type de distorsion de phase produit lorsque la phase du BFO, dans un détecteur-produit, n'est pas identique à celle de la porteuse supprimée. Utilisez la décomposition en série de Fourier d'un signal carré :

$$U(t) = \sin(\omega t) + \frac{1}{3} \sin(3\omega t) + \frac{1}{5} \sin(5\omega t) + \dots$$

Faites tracer par votre ordinateur :

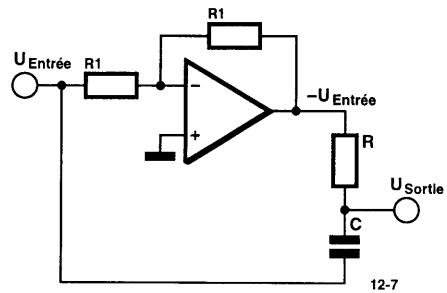
$$U(\theta) = \sin(\theta) + \frac{1}{3} \sin(3\theta) + \frac{1}{5} \sin(5\theta) + \dots + \frac{1}{9} \sin(9\theta)$$

pour θ compris entre 0 et 2π . Tracez ensuite :

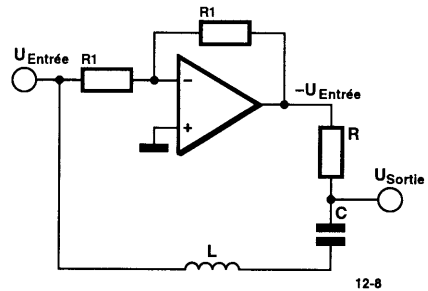
$$U'(\theta) = \sin(\theta + 1) + \frac{1}{3} \sin(3\theta + 1) + \frac{1}{5} \sin(5\theta + 1) + \dots + \frac{1}{9} \sin(9\theta + 1)$$

c'est-à-dire, la même fonction, mais dans laquelle chaque terme de la décomposition en série de Fourier a un déphasage identique (égal ici à 1 radian). Pensez-vous obtenir un résultat similaire ? Examinez également le cas où le déphasage est égal à 180° . La forme d'onde est en opposition de phase. Est-ce un cas particulier ou est-ce une forme d'onde BF inversée, en réalité distordue ?

12.4 Les réseaux de déphasage employés dans la méthode de mise en phase pour créer un signal BLU, présentent une réponse en amplitude plate en fonction de la fréquence – on les appelle des filtres *passé-tout*. Leur nombre de pôles (tous situés dans la moitié gauche du plan ρ) et de zéros (tous situés dans la moitié droite du plan ρ) est le même. Trouvez la réponse en amplitude et en phase du réseau représenté ci-dessous : c'est un filtre passe-tout du premier ordre. Vous noterez que l'amplificateur opérationnel ne sert qu'à inverser le signal d'entrée.



12.5 On obtient un filtre passe-tout universel en mettant en cascade des filtres passe-tout du premier ordre (voir le problème 12.4) et des filtres passe-tout du deuxième ordre. Trouvez la réponse en amplitude et en phase du réseau représenté ci-dessous : c'est un filtre passe-tout du deuxième ordre.



19. Filtres 2 – Filtres à couplage par résonateur

Nous avons vu que la transformation simple d'un filtre passe-bas normalisé en un filtre passe-bande donnait un circuit où alternent circuits résonnants en série et circuits résonnants en parallèle, comme le montre la figure 19-1. Ces filtres passe bande simples fonctionnent très bien – lorsqu'ils sont mis en œuvre dans un programme d'analyse de réseau. Ils font cependant appel à des valeurs de composants qu'il n'est pas facile de se procurer. Les valeurs des bobines des branches en parallèle doivent être plus faibles que celles des bobines des branches en série dans un rapport de l'ordre du carré de la largeur de bande fractionnaire. Si le filtre a une largeur de bande de 5%, le rapport entre les valeurs des bobines est de l'ordre de 1:400. Pour une fréquence centrale donnée, nous aurions déjà beaucoup de chance de pouvoir trouver une bobine à facteur Q élevé de n'importe quelle valeur, sans parler des autres bobines à facteur Q élevé de valeurs tellement différentes. Des bobines à faible facteur Q contribuent à l'obtention d'un filtre à fortes pertes et modifient la forme de la bande passante. Les valeurs des condensateurs en série et en parallèle sont dans le même rapport. En règle générale, le facteur Q n'est pas un problème avec les condensateurs, mais les très faibles valeurs calculées sont peu réalistes lorsqu'elles deviennent comparables aux capacités parasites du câblage.

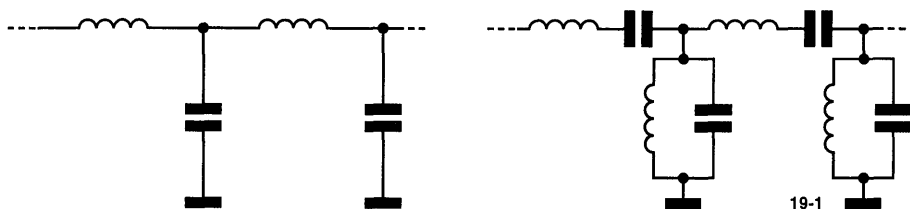


Figure 19-1. Filtre passe-bas normalisé transformé en filtre passe-bande.

Il est possible de résoudre ce problème de valeurs de composants en transformant les filtres passe-bas normalisés en *filtres passe-bande à couplage par résonateur*, que l'on peut réaliser à l'aide de circuits résonnants LC identiques ou presque. La figure 19-2 représente des circuits typiques.

Inverseurs d'impédance

Les filtres à couplage par résonateur sont basés sur les inverseurs d'impédance. La figure 19-3 représente trois exemples d'inverseurs d'impédance : un tronçon de 90° d'une ligne de transmission et deux circuits à composants LC discrets. Dans tous les cas, une impédance Z vue au travers de l'inverseur devient égale à Z_0^2/Z , où Z_0 peut être appelée l'impédance caractéristique de l'inverseur. Pour l'inverseur à ligne de transmission, Z_0 est l'impédance de la ligne. Pour les inverseurs à composants LC , la réactance de la bobine, X_L , et la réactance du condensateur, X_C , doivent être égales à l'impédance Z_0 désirée. Tout comme

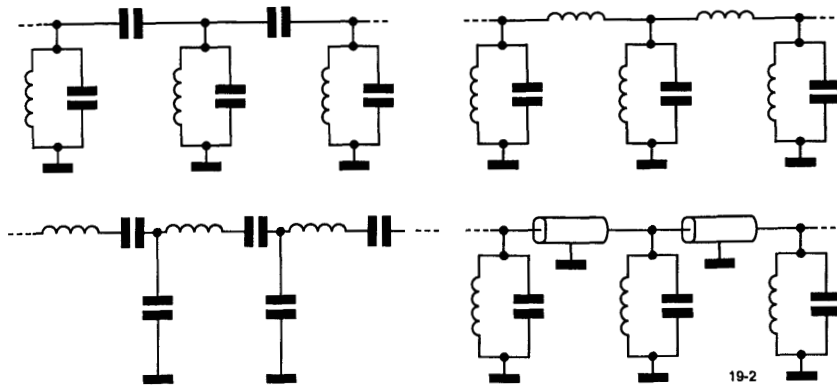


Figure 19-2. Exemples de filtres passe-bande à couplage par résonateur.

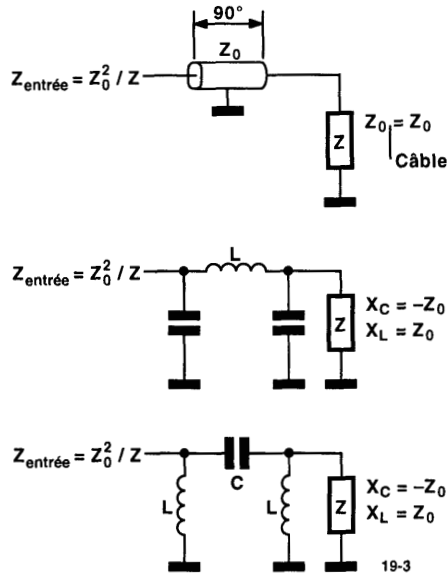


Figure 19-3. Trois circuits inverseurs d'impédance.

le tronçon de câble à 90° , les circuits à composants discrets ne sont en théorie des inverseurs parfaits qu'à une seule fréquence, mais dans la pratique, ils fonctionnent bien sur une large plage de fréquences. Un condensateur inversé devient une bobine, et réciproquement. La figure 19-4 représente un inverseur (dans ce cas, une ligne de transmission à 90°) utilisé pour inverser un circuit en parallèle et le transformer en un circuit en série. Voici l'expression mathématique de cette inversion :

$$Z_{\text{entrée}} = Z_0^2 Y = Z_0^2 \left(\frac{1}{j\omega L_p} + j\omega C_p + \frac{1}{R_p} \right) = \frac{1}{j\omega (L_p / Z_0^2)} + j\omega (Z_0^2 C_p) + \frac{Z_0^2}{R_p} \quad (19.1)$$

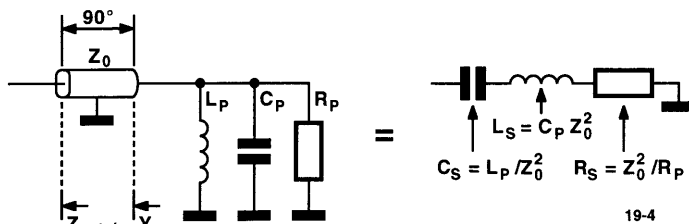


Figure 19-4. Un inverseur d'impédance transforme un circuit en parallèle en un circuit en série.

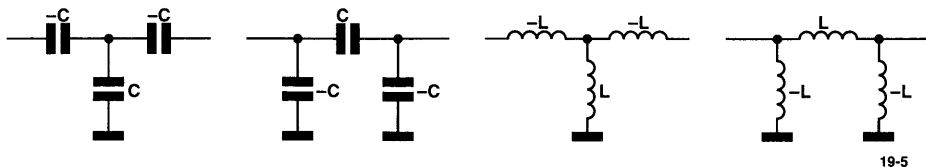
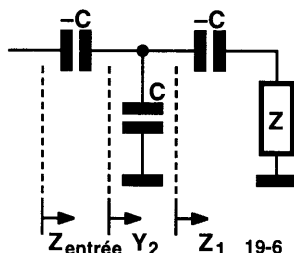


Figure 19-5. Inverseurs d'impédance basés sur des composants de valeur négative.



$$Z_1 = \frac{1}{j\omega(-C)} + Z \quad (19.2)$$

$$Y_2 = j\omega C + \frac{1}{Z_1} = \frac{j\omega C Z}{Z - 1/j\omega C} \quad (19.3)$$

$$Z_{\text{entrée}} = \frac{1}{j\omega(-C)} + \frac{1}{Y_2} = \frac{1}{\omega^2 C^2 Z} = \frac{Z_0^2}{Z} \quad (19.4)$$

Figure 19-6. Fonctionnement d'un inverseur en T à condensateur négatif.

Examinons quatre inverseurs qui comportent des bobines ou des condensateurs de valeurs négatives. La figure 19-5 représente de tels inverseurs pour lesquels $X_C = Z_0$ ou $X_L = Z_0$. Les formules présentées à la figure 19-6 illustrent le fonctionnement d'une section inverseur en T réalisée uniquement à partir de condensateurs. Appliquez le même raisonnement pour vérifier le fonctionnement inverseur des autres circuits.

Étant donné que les quatre circuits inverseurs de la figure 19-5 comportent des capacités ou des inductances négatives, vous pourriez penser qu'il ne s'agit-là que d'une curiosité mathématique. Il n'en est rien ; les éléments négatifs peuvent être « absorbés » par les éléments positifs du circuit contigu, comme le représente la figure 19-7 où un inverseur (section en π) composé de condensateurs est inséré entre deux circuits « bouchon » résonnants en parallèle.

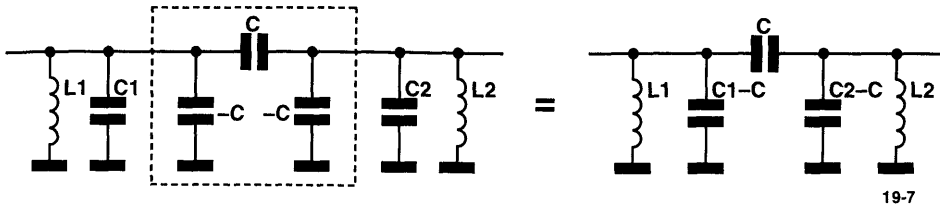


Figure 19-7. Condensateurs négatifs absorbés par des condensateurs contigus.

Les filtres en échelle comportent alternativement des branches en série et en parallèle. Étudions la façon dont sont utilisées des *paires* d'inverseurs dans un filtre en échelle. Supposons que nous placions un condensateur en série entre deux inverseurs en un endroit déterminé du réseau en échelle. La combinaison du condensateur et de la paire d'inverseurs est équivalente à une bobine en parallèle, comme l'illustre la figure 19-8.

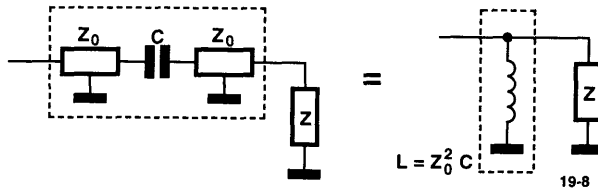


Figure 19-8. Un condensateur en série placé entre deux inverseurs est équivalent à une bobine.

Vous pouvez en déduire aisément que toute impédance en série Z_s , insérée entre deux inverseurs d'impédance caractéristique Z_0 , est équivalente à une admittance en parallèle $Y_p = Z_0^{-2} Z_s$. De même, la combinaison de toute admittance en parallèle Y et d'une paire d'inverseurs est équivalente à une impédance en série $Z = Z_0^2 Y$. C'est ce que montre la figure 19-9 ; on y voit qu'il est possible de remplacer une branche en résonance série située dans un filtre passe-bande classique, par une branche en résonance parallèle située entre deux inverseurs. De la même façon, on peut réaliser une branche en résonance parallèle sous la forme d'une branche en résonance série placée entre deux inverseurs (voir la figure 19-10).

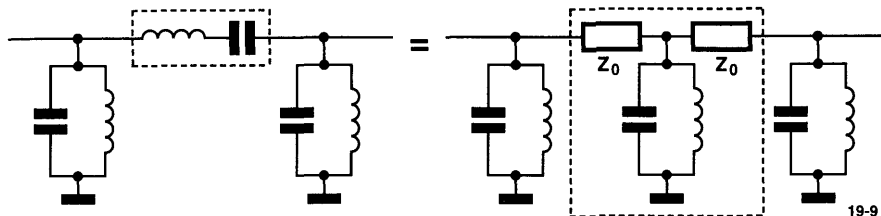


Figure 19-9. Un résonateur en parallèle placé entre deux inverseurs est équivalent à un résonateur en série.

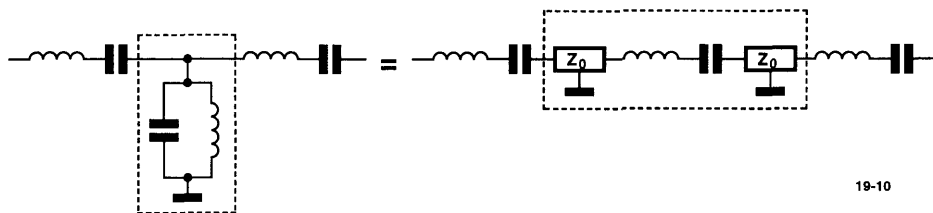


Figure 19-10. Un résonateur en série placé entre deux inverseurs est équivalent à un résonateur en parallèle.

Exemple d'application – filtre passe-bande ayant une largeur de bande fractionnaire de 1%

Considérons le cas d'un filtre de Tchebychev ($50\ \Omega$, 1 dB) ayant une fréquence centrale de 1 MHz et une largeur de bande de 100 kHz entre les points à -1 dB. La figure 19-11 représente le résultat de la transformation simple du filtre passe-bas en filtre passe-bande (reportez-vous au chapitre 4). Le tracé que l'on peut examiner à la figure 19-12 est celui de la réponse en fréquence que fournissent les calculs.

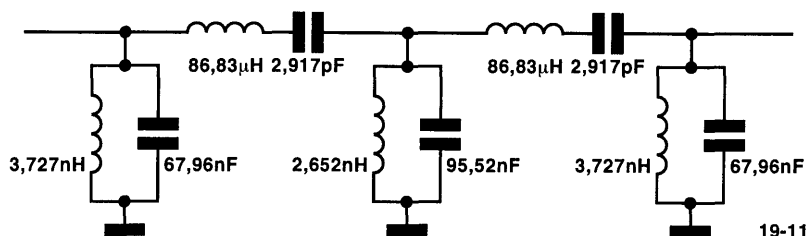


Figure 19-11. Filtre passe-bande direct (mais irréalisable).

Si l'on peut trouver des bobines de $86\ \mu\text{H}$ ayant un facteur Q élevé à 10 MHz, il est difficile de réaliser des bobines de $3,728\ \text{nH}$ et de $2,645\ \text{nH}$ formées d'une minuscule spire qui présenteraient de plus un facteur Q médiocre. Pour ne plus être embarrassés par ces contraintes de valeurs de composants, nous allons transformer ce filtre en un filtre à couplage par résonateur.

Supposons que vous disposions de bobines ajustables de $0,3$ à $0,5\ \mu\text{H}$ dont le facteur Q , à 10 MHz, est élevé. Nous verrons ultérieurement quelle doit être la valeur de ce facteur Q . Commençons par modifier l'impédance de travail du filtre de sorte que les résonateurs en parallèle des extrémités utilisent des bobines

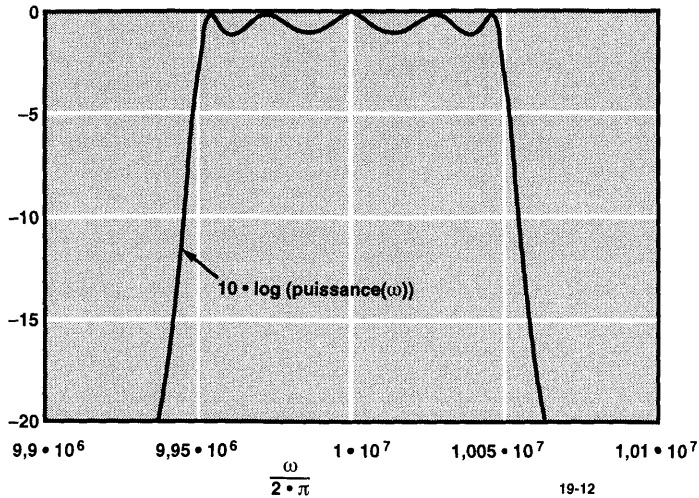
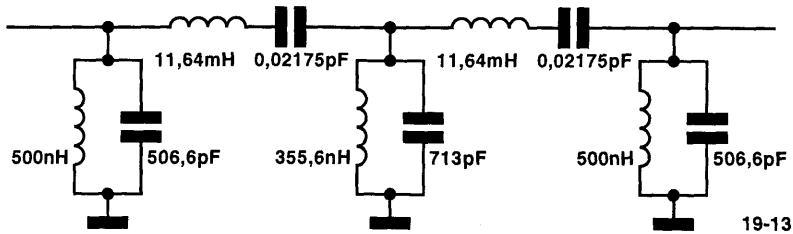


Figure 19-12. Réponse du filtre de la figure 19-11 obtenue par le calcul.

Figure 19-13. Le filtre de la figure 19-11 est mis à l'échelle, de 50 Ω à 6,705 k Ω .

de 0,5 μH soit 134,1 fois la valeur originelle des bobines des extrémités ; l'impédance du filtre sera alors égale à $50 \times 134,1 = 6,705 \text{ k}\Omega$. Nous multiplions la valeur des autres bobines par 134,1 et divisons celle des condensateurs par 134,1 ; nous obtenons alors le circuit de la figure 19-13.

Les résonateurs en parallèle emploient à présent les bobines voulues mais les résonateurs en série font appel à des bobines de 11,6 mH – valeur très forte pour laquelle nous aurons des difficultés à trouver des composants ayant un facteur Q élevé. De plus, il est impossible de trouver les condensateurs en série de 0,02 pF ; cette valeur extrêmement faible est totalement irréaliste. Nous pouvons résoudre ce problème en transformant les résonateurs en série en résonateurs en parallèle à l'aide d'inverseurs d'impédance. Servons-nous d'inverseurs en π composés uniquement de condensateurs et des mêmes résonateurs en parallèle dont nous nous sommes servi pour les extrémités. Nous découvrons à la figure 19-14 comment remplacer un résonateur en série par deux inverseurs et un résonateur en parallèle.

Voici le calcul de l'impédance caractéristique Z_0 de l'inverseur :

$$Z_0 Y = Z ; Z_0 (j\omega C_p + 1/j\omega L_p) = j\omega L_s + 1/j\omega C_s$$

$$Z_0^2 = L_p / C_s = 0,5 \times 10^{-6} / 0,02175 \times 10^{-12} = 4795^2$$

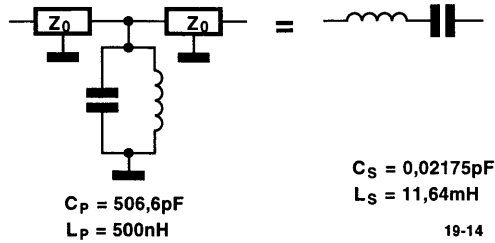


Figure 19-14. Les inverseurs transforment une bobine en parallèle de $0,5\text{ }\mu\text{H}$ en une bobine en série de $11,644\text{ mH}$.

Nous avons vu que dans ce type d'inverseur, $Z_0 = X_c$, donc $C = 3,32\text{ pF}$. Nous avons maintenant réalisé notre filtre à couplage par résonateur, mais étant donné que son impédance est de $6,705\text{ k}\Omega$, nous ajoutons des réseaux d'adaptation en L de façon à obtenir une impédance de $50\text{ }\Omega$. La figure 19-15 représente le schéma du filtre. Tous les résonateurs sont à présent de type parallèle. Dans certains cas, il peut être nécessaire d'employer des inverseurs pour transformer des résonateurs en parallèle en résonateurs en série équivalents pour obtenir un filtre ne faisant appel qu'à des résonateurs en série – voir la figure 19-2.

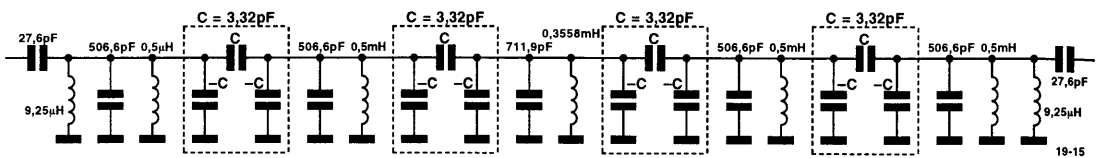
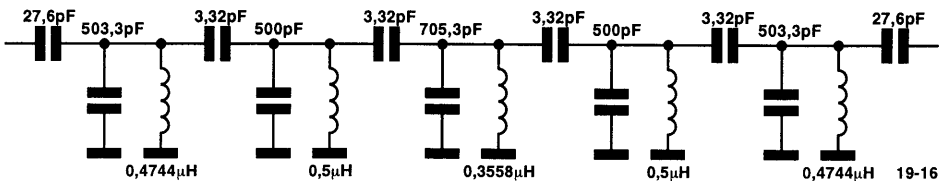


Figure 19-15. Version à couplage par résonateur du filtre passe-bande précédent.

La dernière étape consiste à faire absorber les condensateurs de $-3,32\text{ pF}$ par les condensateurs du résonateur et de combiner les bobines des étages d'adaptation et celles des résonateurs des extrémités, comme le montre la figure 19-16.



Filtre 19-16. Filtre à couplage par résonateur résultant.

La figure 19-17 représente la courbe de réponse du filtre résultant ; elle est presque identique à celle du filtre normalisé de la figure 19-11. La différence constatée, une fraction de décibel, provient du fait que les inverseurs ne fonctionnent parfaitement qu'à la fréquence centrale.

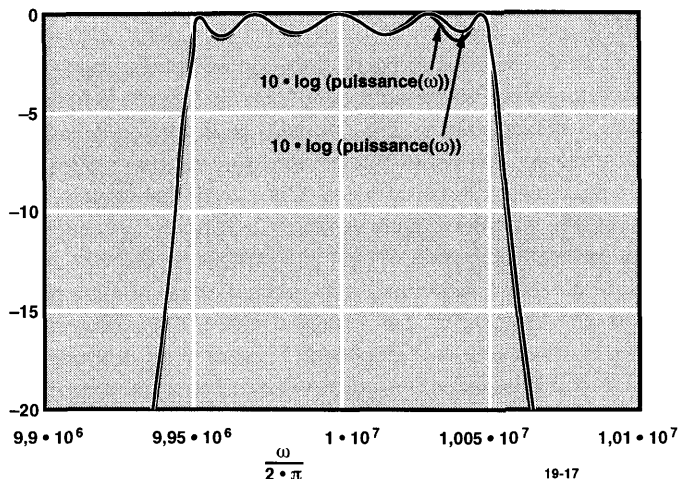


Figure 19-17. Réponse théorique des filtres des figures 19-16 et 19-11.

Conséquences liées au fait que le facteur Q est fini

Les réponses théoriques du filtre supposent que le facteur Q des composants est infini. On peut quantifier les effets d'un facteur Q fini en mettant en parallèle sur les bobines (à pertes) de notre modèle des résisteurs dont les valeurs sont Q fois celles des réactances des bobines à la fréquence centrale. Si le facteur Q est par exemple de 500 (valeur assez élevée pour une bobine), nous devrions mettre en parallèle sur les bobines du filtre de la figure 19-15 des résisteurs d'environ 15 k Ω . Après avoir procédé à une nouvelle analyse de la réponse du circuit, nous découvririons que la perte d'insertion du filtre à mi-bande est de 7 dB et que la réponse plate dans la bande passante (à 1 dB près) est devenue arrondie. L'effet serait moins net si l'on utilisait un filtre à flancs moins raides, comme un filtre de Tchebychev à 0,01 dB ou un filtre de Butterworth. Mais le problème réel se situe toujours au niveau de la faible bande passante fractionnaire. Si l'on veut obtenir un filtre avec une faible bande passante fractionnaire qui ait la courbe de réponse idéale de la figure 19-17, il faut utiliser des résonateurs à quartz, céramiques ou tout autre résonateur dont le facteur Q soit de plusieurs milliers. Une analyse rapide permet de prévoir que la perte à mi-bande *par étage* sera de l'ordre de :

$$\frac{\text{puissance transmise}}{\text{puissance réfléchie}} = \left(1 - \frac{L_0/2}{Q \times \text{largeur de bande fractionnaire}} \right) \quad (19.5)$$

où L_0 représente la valeur de la bobine dans le filtre passe-bas normalisé. Dans le cas de notre filtre à cinq étages, nous pouvons donner à L_0 une valeur d'environ 1,5 H. Si le facteur Q de la bobine est de 500, l'atténuation théorique est égale à $5 \times 10 \log \{ 1 - (1,5/2)/[500 \times (1/100)] \} = -10$ dB, valeur qu'il faut comparer à la valeur pratique de -7 dB.

Procédures d'accord

Les filtres qui présentent de faibles largeurs de bande fractionnaires et des flancs raides sont extrêmement sensibles aux valeurs des composants. Par exemple, dans le cas du filtre de la figure 19-16, il faut accorder très méticuleusement les résonateurs pour ne pas introduire de distorsion dans la courbe de réponse ou accroître l'atténuation. Les valeurs des petits condensateurs de liaison – c'est tout ce qu'il reste des inverseurs d'impédance – ne sont pas aussi critiques. On accorde habituellement chaque résonateur à l'aide d'un condensateur variable ou d'une bobine variable. Tous les réglages interagissent les uns avec les autres et si le filtre est totalement désaccordé, il peut être difficile de détecter un quelconque effet. Une procédure d'accord classique consiste à surveiller l'impédance d'entrée du filtre pendant que l'on accorde, un par un, les résonateurs en commençant par l'étage d'entrée. On court-circuite le résonateur $n+1$ pendant que l'on accorde le résonateur n . On obtient l'accord d'un résonateur lorsque l'impédance d'entrée est maximale tandis que l'on obtient l'accord du suivant pour une impédance d'entrée minimale. Il est parfois nécessaire de personnaliser la procédure pour tenir compte, par exemple, de réseaux d'adaptation placés aux extrémités.

Autres filtres

On emploie la technique du couplage par résonateur dans un domaine qui s'étend des HF aux hyperfréquences. Il faut cependant noter que certains filtres HF n'utilisent pas la technique du couplage par résonateur. La forme de la bande passante de la FI dans un récepteur de télévision est généralement due à l'emploi d'un filtre passe-bande à ondes de surface (ou SAW de *Surface Acoustic Wave*).

Les filtres à ondes de surface sont des filtres *FIR* (de *Finite Impulse Response* ou à réponse impulsionnelle finie) alors que tous les filtres *LC* que nous avons étudiés sont des réseaux *IIR* (de *Infinite Impulse Response* ou à réponse impulsionnelle infinie). On a procédé à cette classification en fonction du comportement de la tension de sortie suite à une impulsion d'excitation delta (impulsion extrêmement fine). On peut concevoir les filtres numériques pour qu'ils se comportent en filtres *FIR* ou en filtres *IIR*.

Bibliographie

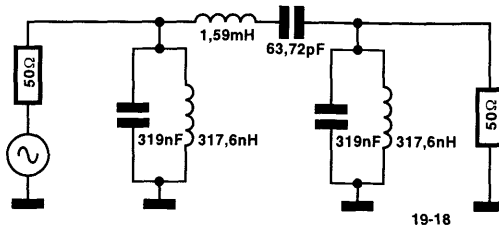
- 1 G.L. Matthaei, L. Young et E.M.T. Jones (1964), *Microwave Filters, Impedance-Matching Networks, and Coupling Structures*. New York: McGraw-Hill (réimprimé par Artech House, Boston, 1980).
- 2 D.G. Fink (1975), *Electronic Engineers' Handbook*. New York: McGraw-Hill.

Problèmes

19.1 Servez-vous de votre programme d'analyse de réseau pour vérifier que le filtre de la figure 19-16 donne bien la courbe de réponse illustrée à la figure 19-17.

19.2 Vérifiez que les deux circuits *LC* de la figure 19-3 sont des inverseurs d'impédance.

19.3 Le filtre représenté ci-dessous, étudié au chapitre 4, est le résultat de la transformation d'un filtre normalisé passe-bas en un filtre passe-bande. La largeur de bande de ce filtre de Butterworth (extrêmement plat) est de 10 kHz et sa fréquence centrale est de 500 kHz.



Supposez que vous ayez à votre disposition des bobines de 30 μH , de facteur Q égal à 100 à 500 kHz. Transformez ce filtre en un filtre à couplage par résonateur qui fasse appel à ces bobines. Utilisez votre programme d'analyse de réseau pour vérifier les caractéristiques du filtre obtenu.

19.4a Nous voulons réaliser un filtre qui ait les caractéristiques suivantes :

- ♦ Fréquence centrale : 10 MHz
- ♦ Forme : filtre de Tchebychev 1 dB à trois étages
- ♦ Largeur de bande : 3 kHz (entre les points à -1 dB)
- ♦ Impédances de la source et de la charge : 50 Ω

Le facteur Q chargé de ce filtre étant très élevé ($10^6/3000 = 333$), il est nécessaire d'utiliser des résonateurs ayant eux aussi un facteur Q très élevé. Supposez que vous disposiez de résonateurs (cavités, quartz ou autres) ayant un tel facteur Q . À 10 MHz, leur résonance est de type parallèle. Lorsque la fréquence varie autour de 10 MHz, la pente de leur susceptance est de 10^{-6} (1 mho/MHz).

Trouvez le circuit équivalent LC de ces résonateurs (au voisinage de 10 MHz).

19.4b Concevez le filtre (ayant les caractéristiques ci-dessus) qui fasse appel à ces résonateurs.

19.4c Utilisez le programme d'analyse de réseau (voir le problème 1.3 au chapitre 1) pour vérifier la réponse en fréquence de votre réalisation.

24. Procédés de télévision

Analyse d'une image

Les procédés classiques de télévision analysent l'image et transmettent séquentiellement les *pixels* (de *Picture Element* ou points d'image) qui la composent [N.d.T. : nous conserverons dans ce chapitre le mot pixel]. L'image est formée de lignes horizontales parcourues séquentiellement pour analyser la luminosité de chacun des éléments qui la forme. L'image est reconstituée au niveau du tube-image ou de tout autre dispositif d'affichage selon un processus inverse ; les pixels sont illuminés les uns après les autres, de la gauche vers la droite et ligne après ligne [N.d.T. : au niveau de la couche sensible de l'écran du tube, on parle de luminophores]. Dans le cas d'un tube à rayons cathodiques (ou CRT de *Cathode Ray Tube*), le signal vidéo commande le courant dans le faisceau et donc la luminosité du point d'impact sur la couche phosphorescente pour chaque position du pixel sur l'écran. Si l'écran est parcouru plus de vingt fois par seconde, le téléspectateur percevra une image stable grâce au phénomène de la « persistance rétinienne de l'œil » (la réponse en fréquence de l'œil est limitée approximativement à 10 Hz). Jusqu'en 1994, toutes les stations de télévision ont employé la modulation analogique, date à laquelle une chaîne commerciale de télédiffusion directe par satellite entra en service et introduisit la modulation numérique. Un boîtier permet d'adapter le système au standard analogique existant. Un standard de télévision numérique, que nous étudierons dans ce chapitre, a été proposé en 1996 à l'*US Federal Communications Commission*.

Système de Nipkow

L'analyse électronique d'une image et sa reconstitution ont été proposées par Nipkow grâce à un système à disques (son brevet a été déposé en 1884) ; il utilisait deux disques tournants synchronisés. Le disque analyseur scrutait l'image tandis que le disque récepteur la reconstituait. L'écran du récepteur, un masque perforé rectangulaire, était en réalité recouvert d'un rideau opaque percé d'un très petit trou illuminé de l'arrière par une lampe à décharge modulée en intensité. La position du trou était semblable à celle d'un point illuminé sur un écran. Ce trou était en fait une série de N trous disposés en spirale sur un disque opaque qui tournait derrière le masque perforé (voir la figure 24-1). L'ouverture dévoilait un seul trou à la fois. Lorsque ce trou actif atteignait l'extrémité droite de l'ouverture, le trou suivant apparaissait du côté gauche mais en se décalant (d'une ligne) vers le bas. Les trous correspondants du disque analyseur permettaient à la lumière issue de l'image originelle d'éclairer une photocellule. Il fallait naturellement que les deux disques soient synchronisés.

Ce procédé primitif a été finalement présenté en 1923 après l'invention de la cellule photoélectrique, de l'amplificateur à tube électronique et de la lampe au néon. Des émissions expérimentales de télédiffusion ont été effectuées aux États-Unis dans la bande des 2-3 MHz avec des canaux d'une largeur de 100 kHz. Les premières émissions de télévision purement électroniques ont été allemandes (1935) et anglaises (1936). Le premier procédé tout électronique américain a été proposé par le *National Television Standards Committee* (NTSC) à la *Radio Manufacturers Association* (RMA). Une fois que cet organisme eût donné son accord à ce qui est encore aujourd'hui une norme, la *Federal Communications Commission* (FCC) donna naissance le 1^{er} juillet 1941 à la télédiffusion commerciale. Les chaînes américaines NBC et CBS

commencèrent à émettre dès ce jour, de New York. Ces stations, bientôt rejointes par d'autres (situées à Philadelphie, Schenectady, Los Angeles et Chicago) émettent en permanence plusieurs heures par semaine durant la Seconde Guerre Mondiale. Le nombre de récepteurs alors en service était de l'ordre de dix à vingt-mille. Malgré cela, la guerre freina le développement commercial de la télévision.

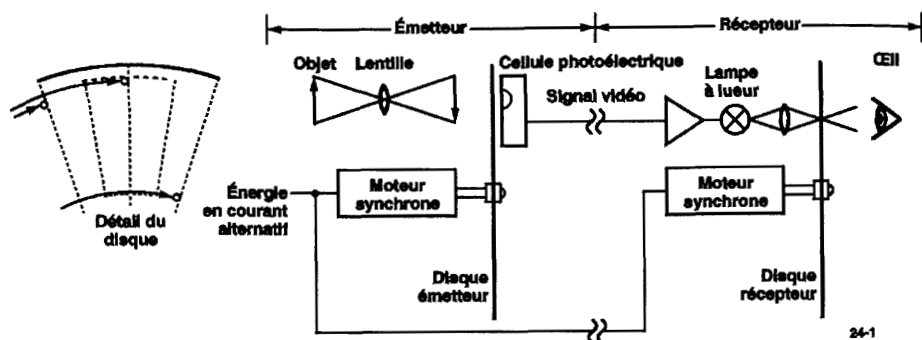


Figure 24-1. Procédé de télévision à disques de Nipkow.

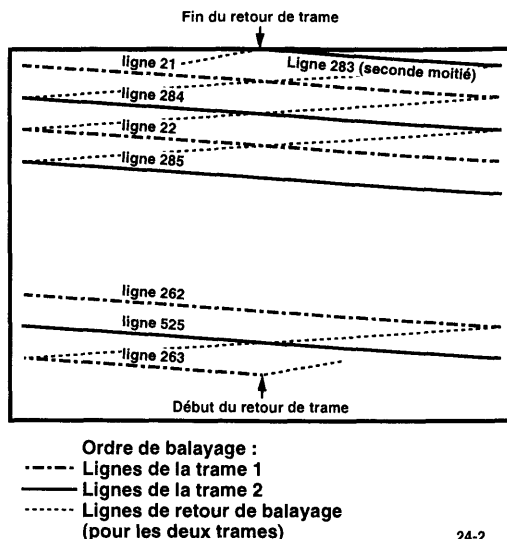
Norme de télévision NTSC

La norme NTSC précise un format de l'image (rapport entre la largeur et la hauteur) de 4:3 pour une trame de balayage avec 525 lignes horizontales, chacune d'elles étant balayée en 62,5 μ s (y compris la période de suppression du faisceau). Une quarantaine de lignes défilent pendant le retour de trame ; l'image comprend donc $525 - 40 = 485$ lignes. Si la définition horizontale était égale à la définition verticale, le nombre de pixels présents sur une ligne serait égal à $4/3 \times 485 = 646$. Le standard NTSC spécifie une définition horizontale moindre, égale à 440 pixels. Le retour de ligne prend approximativement 10 μ s ; la partie active de chaque ligne occupe donc $62,5 - 10 = 52,5$ μ s. La fréquence vidéo maximale est donc donnée par $\frac{1}{2} \times 440/52,5 = 4,2$ MHz. Une onde sinusoïdale vidéo de cette fréquence produirait 220 barres blanches et 220 barres noires.

La fréquence d'image NTSC est de 30 Hz : l'image complète est balayée 30 fois par seconde. La fréquence des lignes est donc égale à $525 \times 30 = 15750$ Hz. Si l'on désire obtenir une meilleure définition verticale, il faut procéder à un balayage double entrelacé. C'est ce que représente la figure 24-2. Les lignes 1 à 262 et la première moitié de la ligne 263 forment la première trame. La seconde moitié de la ligne 263 et les lignes 264 à 525 composent la seconde trame. Les lignes de la seconde trame viennent s'intercaler entre celles de la première. L'entrelacement améliore la définition des images fixes mais peut provoquer des artefacts dus à des objets en mouvement (voir le problème 24-6). Les lignes 1 à 20 se produisent au cours du retour de trame ; durant ce laps de temps, de même que pendant les retours de ligne, elles sont *supprimées* (le faisceau d'électrons est bloqué).

Signal vidéo

Le signal vidéo module en amplitude l'émetteur vidéo. Des impulsions de synchronisation sont insérées entre chaque ligne de données vidéo et la suivante. De même que les disques de Nipkow devaient être synchronisés, l'analyse et la reconstitution électroniques de l'image doivent l'être également ; le point



24-2

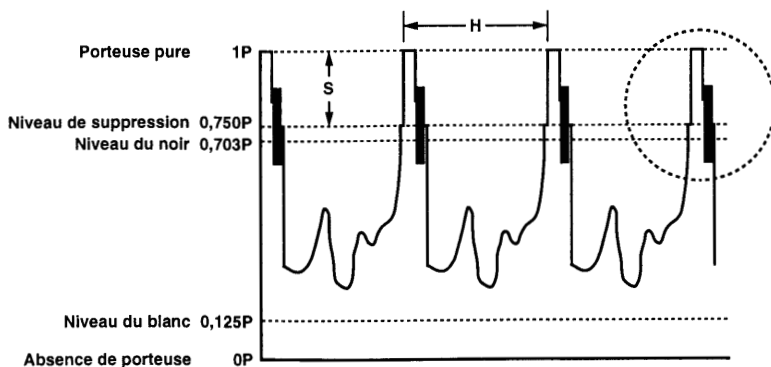
Figure 24-2. Séquence de balayage NTSC.

illuminé sur la trame de balayage du récepteur doit avoir la même position relative que celle du pixel détecté par le tube de prise de vues. Le procédé NTSC utilise une modulation vidéo négative ; ainsi une amplitude décroissante correspond à une luminosité accrue. La principale raison qui a conduit à ce choix est que les parasites impulsionnels produisent plus de points noirs que de points blancs sur l'écran du récepteur et sont donc moins visibles. Les Français ont opté pour une modulation positive afin que les impulsions de bruit ne provoquent aucune erreur de synchronisation ; il faut toutefois noter que les systèmes de synchronisation modernes sont pratiquement insensibles au bruit.

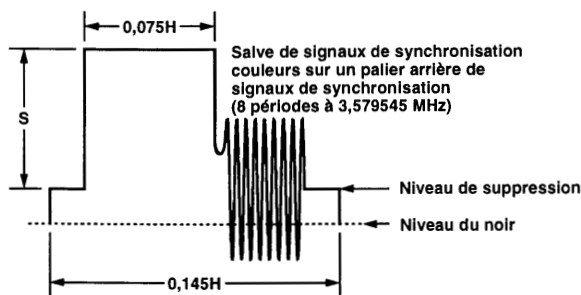
Synchronisation des lignes

Une impulsion de synchronisation de ligne est insérée dans la période de suppression du faisceau entre chaque ligne de balayage ; on la reconnaît à son amplitude plus élevée que la plus élevée des amplitudes des données vidéo. On pourrait dire que l'impulsion de synchronisation est « plus noire » que le niveau du noir qui éteint le faisceau au cours du retour de balayage. La figure 24-3 représente le signal vidéo composite (données vidéo plus impulsions de synchronisation). Cette forme d'onde, telle que l'on pourrait la voir à la sortie d'un détecteur vidéo (après avoir traversé un filtre passe-bas pour éliminer le son à 4,5 MHz) correspond à trois lignes de balayage successives.

Les récepteurs de télévision disposent d'un détecteur à seuil dans le circuit de synchronisation afin de ne détecter que les pointes des impulsions de synchronisation, à savoir la partie de l'impulsion située au-dessus du niveau du noir (par conséquent totalement indépendante des données vidéo). Un différentiateur RC produit une impulsion qui coïncide avec le front (flanc de référence) des pointes de synchronisation. La salve de huit périodes sinusoïdales à 3,579545 MHz du « palier arrière » de chaque impulsion de synchronisation sert de référence au démodulateur de sous-porteuse de chrominance (que nous étudierons ultérieurement).



A. Trois lignes de vidéo



B. Détail de l'impulsion de synchronisation de ligne

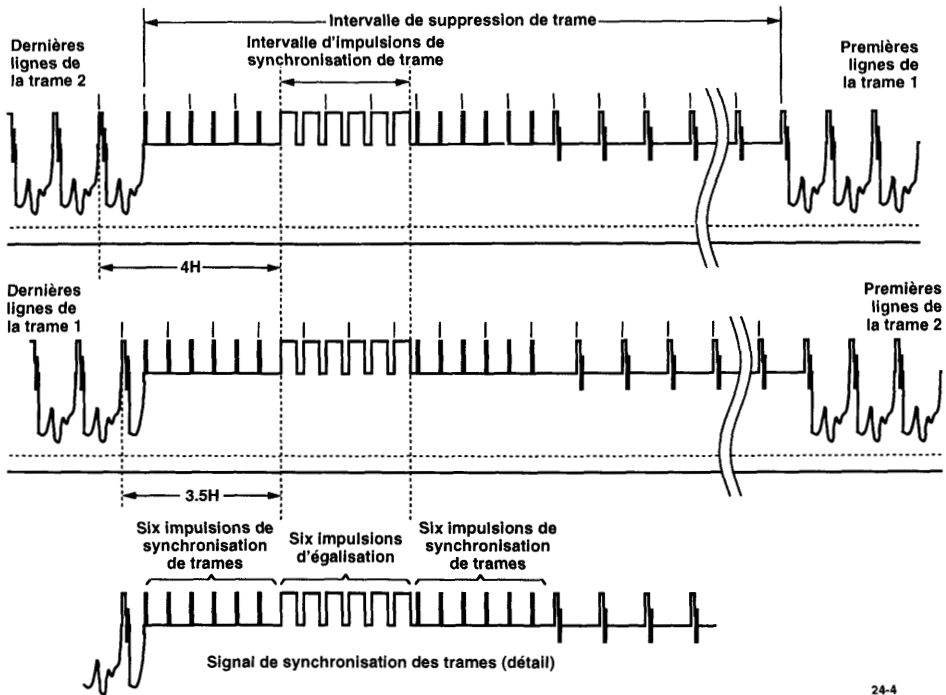
24-3

Figure 24-3. Signal vidéo NTSC.

Synchronisation des trames

La référence de synchronisation des trames est fournie par une série d'impulsions de synchronisation plus larges qui interviennent presque au début de la période de retour de trame (à savoir tous les $1/60$ s à la fin de chaque trame). Un intégrateur RC crée une impulsion (de forme arrondie) chaque fois que ces impulsions larges apparaissent. Un détecteur à seuil, déclenché par cette impulsion, lance le retour de trame. Le processus de synchronisation des trames est légèrement plus complexe à cause de l'entrelacement du balayage. Chaque trame étant composée de 262,5 lignes horizontales, tout autre intervalle de synchronisation des trames est décalé d'une demi-période de balayage horizontal. La figure 24-4 illustre les séquences d'impulsions qui précèdent et suivent l'impulsion de synchronisation des trames dans le cas de trames paires et de trames impaires.

Pour maintenir un entrelacement correct, le signal de sortie de l'intégrateur doit être identique à chaque trame afin que le retour de trame se déclenche précisément tous les $1/60$ s. On insérera une série de six impulsions d'égalisation avant et après les impulsions larges pour que l'action de l'intégrateur ne dépende pas du type de la trame. La constante de temps RC de l'intégrateur est assez courte pour que sa tension de sortie soit toujours fixée par une séquence identique d'impulsions d'égalisation, lorsqu'intervient l'impulsion de synchronisation de trame.



24-4

Figure 24-4. Impulsions de synchronisation NTSC.

Notez que les impulsions d'égalisation et les impulsions larges fournissent une synchronisation permanente des lignes dans et autour de l'intervalle de synchronisation des trames. Dans la trame 1, les références de cadencement des lignes sont les fronts des impulsions d'égalisation et des impulsions larges *paires*, tandis que pour la trame 2, ce sont ceux des impulsions *impaires*. Les petites marques que l'on peut voir à la figure 24-4 signalent les fronts des impulsions de synchronisation des lignes.

Modulation

La transmission radioélectrique de données vidéo (la télévision) se fait en modulant une porteuse par le signal vidéo composite. Le procédé NTSC fait appel à une modulation AM tandis que les liaisons par satellite se font en FM. Puisque le signal vidéo NTSC occupe une largeur de bande de 4,2 MHz, une modulation AM avec les deux bandes latérales nécessiterait une largeur de bande de 8,4 MHz. Pour économiser cette précieuse largeur de bande, on ramène l'intervalle entre les canaux adjacents de 8,4 MHz à 6 MHz en supprimant par filtrage à l'émission, le bas de la bande latérale inférieure. Le signal obtenu conserve la totalité de la bande latérale supérieure, la porteuse et un « moignon » de bande latérale inférieure : c'est ce que l'on voit à la figure 24-5.

Du côté du récepteur, une composante vidéo à basse fréquence (présente dans les deux bandes latérales) produirait une tension double de celle que créerait un signal vidéo à haute fréquence (seulement présent dans la bande latérale supérieure). On résout ce problème à l'aide d'un filtre FI dont la réponse est en pente à son extrémité inférieure, comme le montre la figure 24-6. À la fréquence de la porteuse vidéo, la

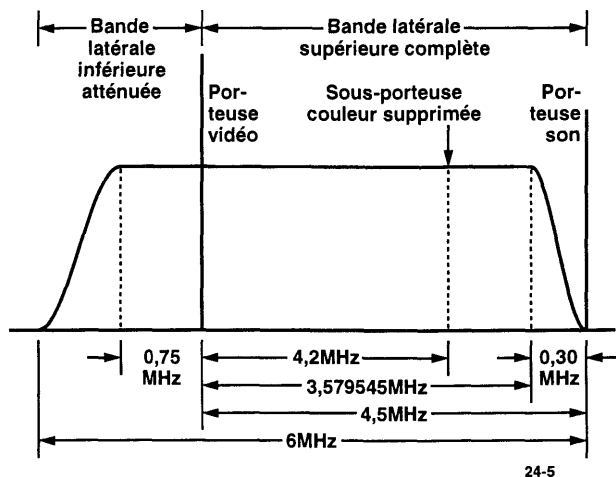


Figure 24-5. Norme FCC d'un canal de télévision de 6 MHz.

réponse en amplitude est de $\frac{1}{2}$. La réponse obtenue est aussi valable et simple à obtenir que celle de la courbe en traits pointillés où la totalité de la zone des bandes latérales est réduite à $\frac{1}{2}$. Pour que cette égalisation FI fonctionne correctement, ce filtre FI doit avoir une réponse linéaire en phase. Aujourd'hui heureusement, les filtres à ondes de surface couramment employés offrent la courbe de réponse désirée et une réponse linéaire en phase (retard constant).

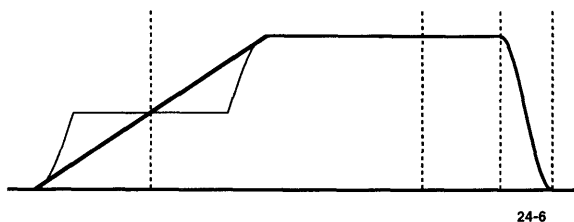


Figure 24-6. Courbe de réponse du filtre FI permettant d'égaliser le « moignon » de bande latérale.

Son

Le signal BF est acheminé sur une porteuse distincte dont la fréquence est supérieure de 4,5 MHz à celle de la porteuse vidéo. Le procédé NTSC utilise la FM pour le son alors que d'autres procédés le font en AM. L'excursion maximale est de 25 kHz ; l'amplitude maximale du signal BF décale donc la fréquence de la porteuse audio de 25 kHz. L'émetteur audio est normalement distinct de l'émetteur vidéo et les deux signaux sont combinés par un duplexeur au niveau de l'antenne. Les premiers postes de télévision comportaient un récepteur FM interne indépendant, destiné à recevoir le son. Cette méthode a été supplantée par le procédé « interporteuse » dans lequel la largeur de bande de l'amplificateur FI est suffisante pour laisser passer le signal BF. Le détecteur vidéo produit alors le signal BF à 4,5 MHz comme s'il s'agissait d'une sous-porteuse du signal vidéo. Le signal traverse un filtre passe-bande à 4,5 MHz, est ensuite amplifié, puis détecté en FM. Un filtre passe-bas correspondant à 4,2 MHz élimine le signal de son du canal vidéo.

Autres normes de télévision

La norme NTSC n'est que l'une des quelques treize autres normes mondiales définies par le CCIR (*Comité Consultatif International des Radiocommunications*). Les différences entre ces diverses normes résident dans le nombre de lignes (525 ou 625), le décalage de la porteuse son, le type de modulation du son (AM ou FM) et la polarité de la vidéo (positive ou négative). Ces systèmes diffèrent aussi légèrement pour assurer une compatibilité avec la télévision en couleur.

Télévision en couleur

La télévision en couleur, de même que la photographie en couleur, est basée sur un système à trois couleurs. Lorsque deux ou plusieurs couleurs (raies spectrales monochromatiques, ensembles de raies monochromatiques ou courbes spectrales continues) impressionnent l'œil, leurs intensités relatives déterminent la teinte perçue. Lorsque l'on superpose trois images identiques, chacune d'elles dans une couleur *primaire* appropriée, l'œil perçoit une image nette en couleur. Le nombre de couleurs primaires mises en œuvre, de même que leur composition spectrale, varient même si trois couleurs primaires sont suffisantes pour restituer la gamme de teintes requises pour transmettre une image de télévision en couleur. Le choix des couleurs particulières rouge, vert et bleu normalisées par le FCC, résulte des luminophores situés sur la dalle du tube à rayons cathodique. Les combinaisons de ces trois couleurs primaires FCC permettent de restituer la quasi totalité des teintes à l'exception des verts très sombres et des verts bleutés. Les ensembles de télévision sur grand écran emploient trois tubes-image à haute intensité, chacun d'eux étant équipé d'un filtre coloré. Les grands récepteurs de télévision à vision directe utilisent un tube unique tricolore muni de trois canons à électrons. L'intérieur de la dalle en verre est couverte de luminophores rouges, verts et bleus (ce sont des grains de matière qui émettent une couleur particulière lorsqu'ils sont frappés par des électrons). Le *masque perforé* (il s'agit d'une plaque métallique perforée de milliers de trous très soigneusement positionnés) se situe immédiatement derrière la dalle. Les électrons émis par le canon « rouge », par exemple, ne peuvent pas atteindre les luminophores bleus ou verts.

Trois couleurs dans un canal

La télévision en couleur nécessite la transmission simultanée de trois images ; il est donc intéressant de savoir comment a été conçu le procédé de télévision en couleur pour les faire tenir toutes les trois dans un canal de 6 MHz attribué à l'origine à une unique image monochrome. Non seulement on y est arrivé, mais de plus la diffusion d'émissions en couleur est compatible avec les récepteurs de télévision en noir et blanc qui subsistent encore. C'est en 1953 qu'a été adoptée la norme NTSC pour une télévision en couleur compatible. Nous étudierons ce système puis nous décrirons deux autres systèmes, à savoir les systèmes PAL et SECAM.

La solution du problème de la largeur de bande se trouve dans le fait que le signal vidéo monochrome n'utilise pas la totalité de la largeur de bande vidéo. Une image est, par nature, fortement redondante. En effet, toute ligne est quasiment identique à sa voisine immédiate (de nombreuses corrélations à 62,5 μ s) ; le signal vidéo ressemble à une forme d'onde répétitive à une fréquence de 15750 Hz (la fréquence de balayage horizontal). Si les lignes étaient toutes identiques, le spectre de fréquences serait un peigne de fonctions delta situées à 15750 Hz et à ses fréquences harmoniques. Étant donné qu'une ligne diffère légèrement de la suivante, ces fonctions delta sont élargies mais l'énergie spectrale se cantonne toujours

autour des harmoniques de la fréquence de balayage horizontal en laissant de la place pour les données couleur. Si l'image est immobile, le spectre est alors un peigne de fonctions delta espacées de 30 Hz et la largeur de bande est très peu utilisée.

Compatibilité

Avant d'insérer le signal de couleur dans les espaces vides, il nous faut préciser comment est assurée la compatibilité avec les récepteurs de télévision en noir et blanc. Au lieu de transmettre de la même manière les signaux rouge, vert et bleu (RVB), on effectue trois combinaisons linéaires. L'une d'elles, le signal de *luminance*, Y , est choisi comme le signal de luminosité qu'aurait produit une caméra monochrome : $Y = 0,299R + 0,587V + 0,114B$. Les deux autres combinaisons linéaires sont $I = 0,74(R - Y) - 0,27(V - Y)$ et $Q = 0,48(R - Y) + 0,41(B - Y)$.

Le signal de luminance module directement la porteuse, exactement comme en télévision noir et blanc, et les récepteurs de télévision en noir et blanc fonctionnent normalement. Chacun des deux autres signaux, I et Q , module une sous-porteuse (de la même façon qu'un signal BF droite moins gauche module une sous-porteuse à 38 kHz dans le cas d'une émission stéréo en FM). Les sous-porteuses couleur sont à la même fréquence de 3,579545 MHz, mais sont déphasées de 90° . La figure 24-7 comment émettre et récupérer deux signaux indépendants dans une seule bande de fréquences d'émission grâce à ce que l'on appelle la modulation en quadrature. Notez qu'avec ce type de multiplexage, il serait possible d'avoir deux stations de radio AM sur chaque fréquence attribuée, mais il faudrait équiper les récepteurs de détecteurs de produits asservis en phase et non de simples détecteurs d'enveloppe.

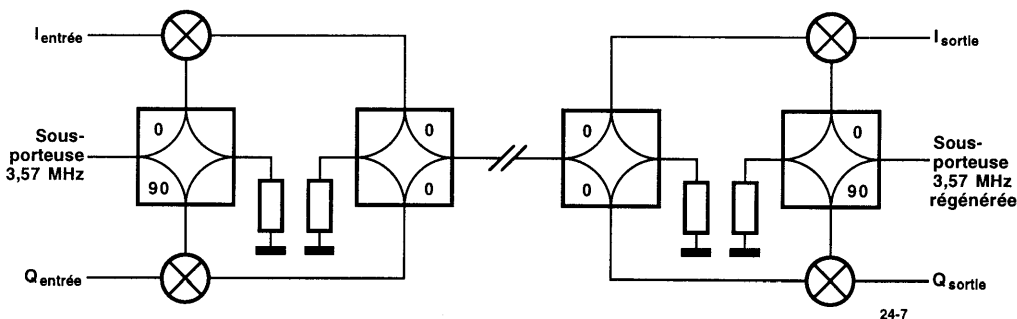


Figure 24-7. Deux signaux partagent une même bande de fréquences grâce à un modulateur/démodulateur en quadrature.

Dans la télévision en couleur, on supprime la sous-porteuse, de sorte que les signaux I et Q (en phase et en quadrature) sont des signaux à deux bandes latérales et porteuse supprimée. Le récepteur doit fournir des porteuses locales (signaux issus de BFO) à la bonne fréquence et avec des phases correctes pour démoduler (détecter) ces signaux. On envoie le signal de référence nécessaire à la restauration de ces porteuses locales sous la forme d'une salve de huit périodes sinusoïdales à 3,579545 MHz sur le palier arrière de chaque impulsion de synchronisation de ligne, comme le représente la figure 24-4. Cette salve couleur sert de référence à une boucle à phase asservie (PLL) du récepteur. La fréquence de la sous-porteuse couleur devant être stable à 10 Hz près, on préfère utiliser un VCXO (de *Voltage Controlled Crystal Oscillator* ou oscillateur à quartz commandé en tension) à la place d'un VCO.

L'information de couleur, comme l'information de luminance, est semblable d'une ligne à l'autre ; son spectre de fréquences est également un peigne dont les composantes ont un espacement égal à la fréquence de balayage horizontal. On choisit la fréquence de la sous-porteuse au milieu de l'un des créneaux laissés libres par le signal de luminance ; le peigne des bandes latérales couleur est imbriqué dans celui des bandes latérales de luminance.

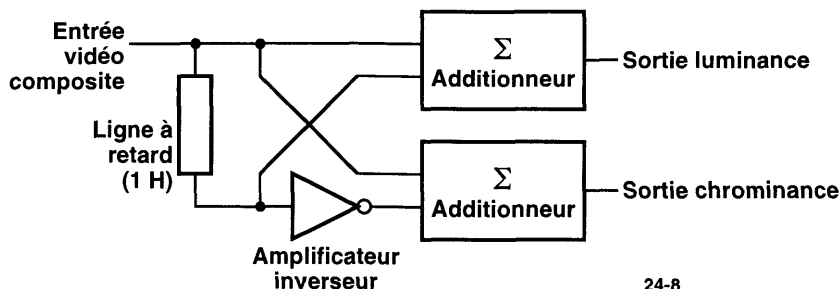
À l'origine, les récepteurs de télévision en couleur ne disposaient d'aucun moyen pour séparer les signaux de luminance et de chrominance ; à l'époque, on ne savait pas réaliser les filtres en peigne nécessaires. La compatibilité résultait du fait que l'entrelacement spectral atténuait fortement la « diaphonie » visible entre les données de chrominance et de luminance. Pour s'en rendre compte, considérons un signal très simple, par exemple une trame de couleur unie comme un écran intégralement jaune. Cette trame a une couleur, ce n'est pas du noir, du blanc ou du gris ; les signaux I et Q ne sont donc pas nuls. Dans cet exemple, l'information couleur est invariante et les signaux I et Q sont sinusoïdaux à la fréquence de la sous-porteuse couleur. Leurs amplitudes relatives fixent la teinte et leurs amplitudes absolues, la saturation. On pourrait s'attendre à ce que cette composante vidéo à 3,58 MHz produise des bandes verticales. Le faisceau, lorsqu'il balaie l'écran, est tour à tour plus clair et plus sombre à une fréquence de 3,58 MHz et tente de visualiser quelques 186 bandes au cours de son balayage de 52,5 μ s. Mais à la ligne suivante, ces bandes se déplacent exactement d'une demi-période. L'écran, au lieu d'afficher 186 bandes verticales montrera un damier très finement quadrillé, motif « à faible visibilité » qu'il sera difficile de distinguer. Des objets en couleur vus sur un récepteur de télévision en noir et blanc peuvent présenter cet effet de damier fin. On peut également le constater sur un vieux récepteur de télévision en couleur ; les récepteurs les plus récents sont équipés d'un filtre en peigne (décrit ci-dessous). La largeur de bande vidéo des récepteurs de télévision en noir et blanc économiques est inférieure à 4,2 MHz ; on ne voit plus le motif en damier mais la qualité de l'image est médiocre.

Il faut savoir qu'il existe quelques rares cas où l'entrelacement spectral ne fonctionne pas. Si l'image représente elle-même un damier ayant le même quadrillage, le signal de luminance se situera dans les créneaux attribués à la chrominance et vice versa. Par exemple, un dessin à chevrons regardé sur un récepteur de télévision en couleur, aura un aspect scintillant flagrant.

Le principe de la faible visibilité s'utilise non seulement pour éviter la diaphonie luminance-chrominance mais également pour réduire l'effet de battement entre la porteuse son à 4,5 MHz et la sous-porteuse couleur. Pour profiter de ce principe, les normes de télévision ont été légèrement modifiées lors de l'introduction de la télévision en couleur. La porteuse son se situe toujours à 4,5 MHz au-dessus de la porteuse vidéo. Le rapport entre la fréquence de balayage horizontal f_h et la fréquence de la sous-porteuse couleur f_{sc} est $f_{sc} = 227,5 f_h$. Ensuite, le battement sous-porteuse son moins sous-porteuse couleur est égal à n (impair) fois la fréquence de balayage horizontal : $4,5 \text{ MHz} - f_{sc} = 58,5 f_h$. Ces deux relations fixent la fréquence de balayage de ligne, 15734,264 Hz et la fréquence de la sous-porteuse couleur, 3,579545 MHz. Le nombre de lignes de balayage étant toujours égal à 525, la fréquence de balayage de trame passe de 60 Hz à $262,5 f_h = 59,940 \text{ Hz}$. Si l'on opte pour ces choix, la fréquence de la porteuse son est 286 fois celle du balayage de ligne. On obtient un motif de grande visibilité, mais la porteuse son se situe au-dessus de la bande de fréquences vidéo nominale et peut donc être facilement filtrée.

Filtres en peigne

Les filtres en peigne destinés à séparer le spectre de luminance du spectre de chrominance peuvent employer une ligne à retard (ultrasonore, à CCD ou numérique). La ligne à retard à CCD, sorte de registre à décalage analogique BBD, a été fréquemment utilisée. Le circuit de la figure 24-8 comporte une ligne à retard de 1 H (ligne à retard d'une période de retour de ligne) qui délivre aux filtres en peigne les signaux de luminance et de chrominance.



24-8

Figure 24-8. Filtres en peigne de luminance et de chrominance.

La figure 24-9 explicite le calcul de la réponse en fréquence de ces filtres (la fréquence est normalisée à la fréquence de balayage horizontal). La réponse en amplitude présente l'entrelacement recherché des peignes, de sorte que le filtre de luminance laisse passer les signaux autour des harmoniques de la fréquence de balayage horizontal tandis que le filtre de chrominance couvre les régions intermédiaires. Le graphique est limité à trois fois la fréquence de balayage lignes mais la réponse périodique du filtre s'étend sur la totalité de la largeur de bande vidéo de 4,2 MHz.

$$F(f) := 1 + e^{j2\pi f} \text{ Fonction de transfert du filtre de luminance}$$

$$G(f) := 1 - e^{j2\pi f} \text{ Fonction de transfert du filtre de chrominance}$$

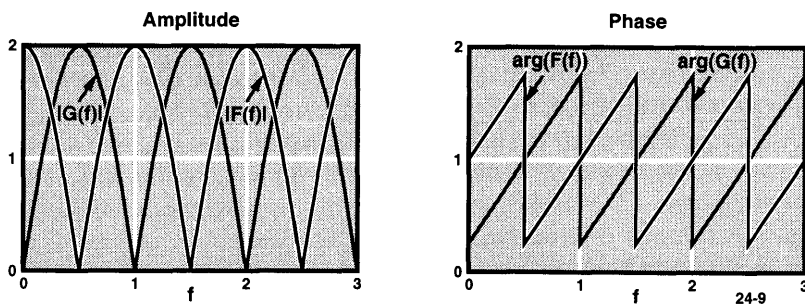


Figure 24-9. Réponses en amplitude et en phase du filtre en peigne.

Notez que le filtre de luminance est un véritable filtre linéaire de phase, il retarde sans distorsion l'information de luminance. Le filtre de chrominance devrait avoir également une réponse en phase linéaire sauf pour un décalage constant de 90°. Mais étant donné que l'information de chrominance est un signal à bande relativement étroite (l'information couleur se cantonne aux alentours de la fréquence de référence couleur), un simple filtre égaliseur passe-tout peut rectifier ce décalage. Le fonctionnement de ce

filtre est très simple lorsqu'on l'étudie dans le domaine temporel ; lorsque deux lignes successives sont superposées, les signaux de chrominance sont de polarité opposée alors que les signaux de luminance ont la même polarité. Tant que deux lignes successives ne sont pas très différentes l'une de l'autre, leur somme ne donne que le signal de luminance et leur différence, le signal de chrominance.

Procédé PAL

Le procédé PAL, très proche du procédé NTSC, est plus souple vis à vis des erreurs de phase des signaux couleur. Nous allons voir que les erreurs de phase introduisent une certaine interférence entre le signal I et le canal Q et réciproquement. On peut corriger une erreur de phase permanente en réglant la commande de la teinte du récepteur de télévision, mais les erreurs de phase qui fluctuent dans le temps distordent l'information couleur. PAL est un acronyme qui vient de *Phase Alternation by Line*. Comme dans le procédé NTSC, le procédé PAL utilise une modulation AM en quadrature pour placer deux signaux couleur sur une sous-porteuse couleur supprimée. Dans le procédé PAL, les deux signaux couleur sont repérés par les lettres U et V . Ces signaux sont des combinaisons linéaires différentes de celles (I et Q) que l'on rencontre dans le procédé NTSC : $U = B - Y$ et $V = \pm(R - Y)$. Vous aurez certainement remarqué le signe $\pm$ qui figure dans l'expression du signal V ; dans le signal émis, la polarité du signal V change d'une ligne à l'autre. Du côté du récepteur, on rétablit cela en combinant les signaux couleur de deux lignes successives pour produire les signaux U et V corrigés. Examinons le fonctionnement de ce procédé : le signal de couleur émis pour la ligne l (une « $+V$ -ligne » plutôt qu'une « $-V$ -ligne ») s'écrit ainsi :

$$T_l = u \cos(\omega_{sc}t) + v \sin(\omega_{sc}t) \quad (24.1)$$

où est ω_{sc} la fréquence de la sous-porteuse couleur. Supposons à présent qu'entre le codeur de l'émetteur et le décodeur du récepteur, le signal de couleur soit entaché d'une erreur de phase $\delta\theta$. Le signal reçu correspondant à la ligne l est :

$$R_l = u \cos(\omega_{sc}t + \delta\theta) + v \sin(\omega_{sc}t + \delta\theta) \quad (24.2)$$

Le signal reçu pour la ligne suivante (n'oublions pas l'inversion de polarité, il s'agit d'une ligne « $-V$ -ligne ») s'écrit :

$$R_{l+1} = u \cos(\omega_{sc}t + \delta\theta) - v \sin(\omega_{sc}t + \delta\theta) \quad (24.3)$$

Le décodeur délivre une sortie U à l'aide d'un détecteur synchrone qui multiplie le signal de couleur reçu par $\cos(\omega_{sc}t)$. Effectuer cette opération sur R_l nous conduit à écrire :

$$U_l = \cos(\omega_{sc}t) R_l = u \cos(\omega_{sc}t) \cos(\omega_{sc}t + \delta\theta) + v \cos(\omega_{sc}t) \sin(\omega_{sc}t + \delta\theta) \quad (24.4)$$

Servons-nous des formules qui donnent le développement de $\cos(a+b)$ et de $\sin(a+b)$ et négligeons les composantes $2\omega_{sc}$; l'équation 24-4 s'écrit alors :

$$U_l = \frac{1}{2} u \cos(\delta\theta) + \frac{1}{2} v \sin(\delta\theta) \quad (24.5)$$

Nous constatons que l'erreur de phase a introduit un signal V dans la sortie U et a atténué le signal U d'un facteur $\cos(\delta\theta)$. Si nous calculons la sortie U pour la ligne suivante, nous obtenons :

$$U_{l+1} = \frac{1}{2} u \cos(\delta\theta) - \frac{1}{2} v \sin(\delta\theta) \quad (24.6)$$

Le canal U dépend de nouveau du signal V , mais ce dernier est égal et opposé à celui de la ligne précédente. Si nous ajoutons les équations 24-5 et 24-6, nous obtenons :

$$U_1 + U_{1+1} = u \cos(\delta\theta) \quad (24.7)$$

La somme est indépendante de V . Le rôle du décodeur PAL est d'effectuer cette addition et de produire un signal corrigé U pour la ligne l qui contient les signaux couleur de cette ligne et de la précédente. Si nous avons toujours présent à l'esprit le fait que le balayage est entrelacé, l'effort dépensé pour la correction de la phase est un processus de lissage vidéo qui prend en réalité trois lignes. Cette réduction de la définition couleur verticale est compatible avec celle de la définition couleur horizontale. Le signal de luminance met en évidence les détails fins, et les signaux de couleur ne requièrent qu'une « mise » en couleur. Remarque : les filtres en peigne évoqués précédemment introduisent une perte de la définition verticale dans les signaux de couleur et de luminance.

Le canal V fonctionne de la même façon si ce n'est que les signaux démodulés sont soustraits, et non plus additionnés, pour obtenir la sortie V corrigée. En procédant de la même façon que pour la sortie U , nous pouvons écrire :

$$V_1 - V_{1+1} = v \cos(\delta\theta) \quad (24.8)$$

Cependant, la prochaine ligne en sortie du soustracteur s'écrit (en tenant compte du changement $\pm V$) :

$$V_{1+1} - V_{1+2} = -v \cos(\delta\theta) \quad (24.9)$$

Le signe moins indique que la polarité de la sortie du soustracteur doit être inversée d'une ligne à l'autre. Ces inversions doivent être parfaitement synchronisées sous peine d'obtenir en sortie $-v \cos(\delta\theta)$ au lieu de $v \cos(\delta\theta)$. On obtient cette synchronisation en déphasant la salve couleur de $\pm 45^\circ$ d'une ligne à l'autre. La boucle à phase asservie ou le filtre à quartz, utilisé pour restaurer la sous-porteuse, filtre ce « bruit » présent sur la salve. La comparaison de la phase de chaque salve à celle de la sous-porteuse stable restaurée fournit le signal nécessaire à la synchronisation des signaux V inversés.

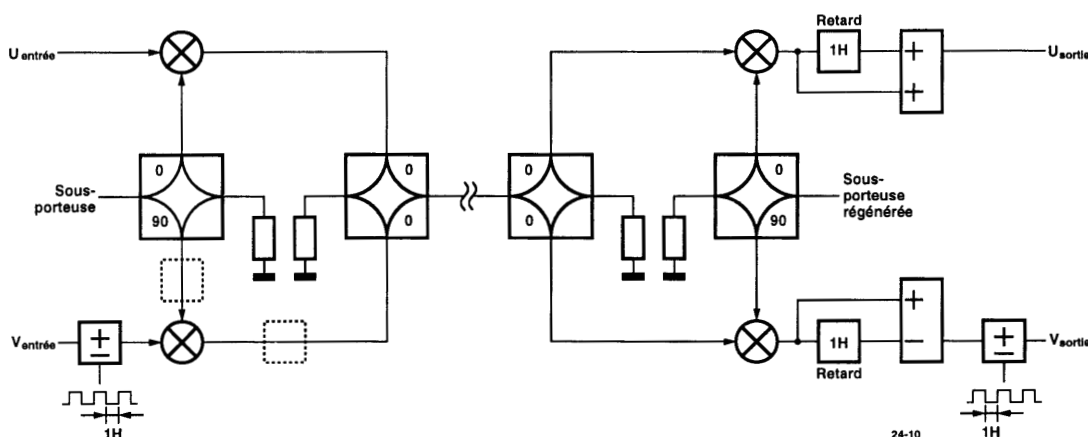


Figure 24-10. Codeur-décodeur PAL. À l'émission, la polarité du signal V est inversée une ligne sur deux. Du côté du récepteur, les signaux U et V sont formés à partir de deux lignes successives. La polarité du signal V décodé doit être restaurée en inversant la polarité une ligne sur deux.

Les signaux U et V issus du décodeur PAL n'ont aucune diaphonie, mais leurs amplitudes sont réduites d'un facteur $\cos(\delta\theta)$. La saturation est légèrement réduite, contrairement aux erreurs de teinte, mais le téléspectateur ne s'en rend pas compte. La figure 24-10 représente le schéma synoptique du système de codage et de décodage que nous venons d'étudier.

Les carrés en traits pointillés indiquent des emplacements possibles pour l'inverseur de polarité au niveau de l'émetteur. L'inverseur de polarité dans le décodeur peut également se situer à un autre endroit. Le signal de l'OL est en général le plus facile à inverser. Dans le cas d'un décodeur PAL très rustique, il n'y a ni ligne à retard ni additionneur. Les couleurs des lignes adjacentes sont « intégrées visuellement » par le téléspectateur. Vous noterez que les fonctions de détection synchrone, retard et sommation sont linéaires et peuvent donc être effectuées dans un ordre quelconque : retard, sommation puis détection. Cette constatation nous permet de n'utiliser qu'une seule ligne à retard comme l'illustre la figure 24-11.

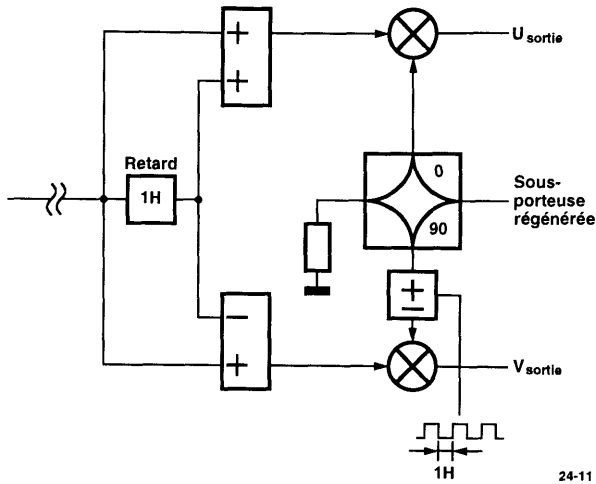


Figure 24-11. Décodeur PAL à une seule ligne à retard.

Dans le procédé PAL, les inversions de phase du signal V une ligne sur deux entraîne une concentration de son spectre sur les multiples impairs de $f_h/2$. Le spectre de U , comme les spectres de I et de Q dans le procédé NTSC, est concentré autour des multiples entiers de f_h . On choisit donc la fréquence de la sous-porteuse PAL égale à un multiple impair de $f_h/4$ afin de séparer les composantes spectrales de la luminance et de la chrominance :

$$f_{sc} = (f_h/4) \cdot 1135 + f_{\text{trame}}/2 = (15625/4) \cdot 1135 + 25 = 4.433.618,75 \text{ Hz}$$

Le petit facteur additionnel de 25 Hz est ajouté pour rendre moins visible l'effet de damier dû au signal de luminance, en le décalant légèrement d'une trame à l'autre sur un cycle de quatre trames.

Procédé SECAM

Le procédé français SECAM (*Séquentiel Couleur à Mémoire*) résoud lui aussi le problème du déphasage. Au lieu de transmettre simultanément les signaux de couleur en quadrature et en AM, les signaux U et V sont envoyés séquentiellement sur des lignes adjacentes. Ce procédé nécessite donc l'utilisation d'une

ligne à retard (la mémoire) dans le décodeur. Comme dans le procédé PAL, l'information couleur d'une ligne donnée est contenue dans l'information couleur de cette ligne plus dans celle que contenait la ligne précédente ; la définition couleur verticale est améliorée. De plus la sous-porteuse couleur SECAM (à 4,4375 MHz) est modulée en FM, et non pas en AM. La sous-porteuse est toujours présente. La figure 24-12 illustre ce système de codage et de décodage.

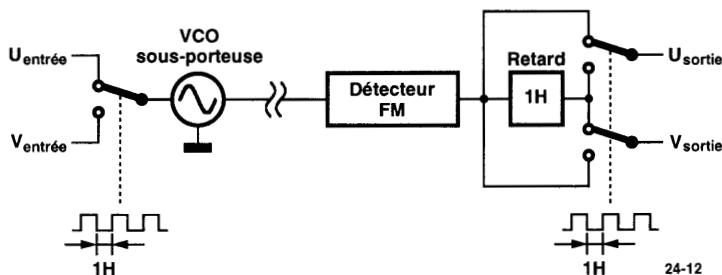


Figure 24-12. Codeur-décodeur SECAM.

Le réseau d'accentuation et de désaccentuation ne figure pas dans ce schéma synoptique ; il est similaire à celui que l'on rencontre en radiodiffusion FM et permet de réduire le bruit HF. Le détecteur FM peut être un simple discriminateur et il n'est pas nécessaire d'envoyer des salves couleur. Un signal de synchronisation est évidemment indispensable pour distinguer les lignes U des lignes V .

Émetteurs de télévision

On emploie des klystrons amplificateurs dans les émetteurs UHF et des tubes dans les amplificateurs VHF. Les émetteurs de télévision travaillent généralement à partir de signaux à bas niveau amplifiés ensuite linéairement. On rencontre parfois un tube amplificateur de puissance modulé par la grille (les hautes tensions nécessaires à une modulation par la plaque sont trop difficiles à utiliser pour produire des fréquences vidéo ; cela impliquerait des milliers de volts et des largeurs de bande de plusieurs mégahertz). L'amplificateur délivre la plus grande partie de sa puissance lors des pointes de synchronisation où la linéarité n'est pas trop préoccupante. La nécessité d'avoir à produire des impulsions de synchronisation de grande puissance se traduit par un rendement faible. Pour l'améliorer, il est possible d'employer des circuits de modulation particuliers qui survolent l'amplificateur au moment de l'envoi des impulsions de synchronisation. La puissance de l'émetteur son est approximativement le quart de celle de l'émetteur vidéo.

Récepteurs de télévision

Les récepteurs de télévision sont du type superhétérodyne à simple conversion avec une bande FI qui place la porteuse vidéo à 45,75 MHz. L'amplificateur d'entrée utilise fréquemment un FET (de *Field Effect Transistor* ou transistor à effet de champ) à l'arséniure de gallium (FET GaAs) et présente un facteur de bruit inférieur à 3 dB. L'OL est généralement un synthétiseur de fréquences verrouillé sur une fréquence de référence ; une commande fine de l'accord n'est plus nécessaire. La courbe de réponse du filtre

passer-bande FI doit être très précise pour égaliser le moignon de la bande latérale, avoir la largeur de bande vidéo totale et éliminer les canaux adjacents. Cela nécessitait la calibration soignée, effectuée en usine, de nombreux circuits accordables, mais à présent cette procédure fastidieuse a été remplacée grâce à la conception de la géométrie d'un simple filtre à ondes de surface. Les récepteurs de télévision modernes emploient des circuits numériques qui fournissent une synchronisation stable asservie en phase et réalisent diverses corrections. Les organismes de télédiffusion envoient des signaux d'étalonnage et d'autres signaux durant l'extinction des lignes au cours du retour de trame. Les signaux d'étalonnage permettent aux récepteurs plus anciens de corriger l'équilibre des couleurs, d'initialiser les systèmes d'annulation des signaux par trajets multiples, d'avoir des informations de télétexte, etc. Les alimentations à découpage font partie des circuits de déviation horizontale à rendement élevé (voir le chapitre 16).

Récepteurs de télévision en couleur

Le schéma synoptique de la figure 24-13 représente l'organisation complète d'un récepteur de télévision en couleur complet. Le premier schéma est simplement celui d'un récepteur de radio AM destiné à recevoir des émissions en VHF, UHF ou via le câble (entre 52 et 400 MHz). La porteuse du canal sélectionné est déplacée à une fréquence FI de 45,75 MHz et la largeur de bande de la FI est d'environ 6 MHz. Le signal FI renferme la vidéo composite (la luminance et la synchronisation) plus le son et l'information couleur. Les signaux de son et de couleur qui sont essentiellement des signaux à bande étroite situés respectivement autour de 4,5 MHz et 3,57 MHz, se situent en haut du signal de luminance. Un filtre passer-bande à 4,5 MHz isole le son dans la section marquée « Récepteur FM à 4,5 MHz ». La sortie de ce récepteur FM attaque directement un amplificateur de puissance BF ou, lorsque le son est stéréophonique, un ensemble stéréophonique dans lequel les canaux gauche et droite sont démodulés.

Son stéréophonique

La chaîne stéréophonique utilisée pour le son en télévision est essentiellement la même que celle que l'on rencontre en radiodiffusion d'émissions en FM. On ajoute, à l'émission, les canaux gauche (G) et droite (D) pour obtenir un signal G+D et on les soustrait pour obtenir un signal G-D. Le signal G-D, le signal stéréophonique différence, est multiplié par une sous-porteuse à 31,4686 MHz. Vous vous souvenez certainement qu'un multiplicateur est un mélangeur symétrique ; le produit est formé d'un ensemble de bandes latérales de modulation inférieure et supérieure centrées autour de la fréquence de la sous-porteuse. Ce signal AM à deux bandes latérales et à porteuse supprimée plus le signal somme G+D et une onde pilote de faible niveau à la fréquence moitié de la sous-porteuse composent le signal de modulation d'un émetteur FM. Au niveau du récepteur, le signal de sortie du démodulateur FM contient le signal G+D déjà utilisable pour écouter une émission monophonique, mais le signal G-D, si l'on en a besoin, doit voir sa fréquence abaissée pour quitter la zone inaudible située autour de la fréquence de la sous-porteuse. Ce changement de fréquence est une simple multiplication. Le signal de l'OL qui doit avoir, rappelons-le, une fréquence et une phase correctes, est issu d'une PLL qui se verrouille sur une onde pilote à 15,75 kHz. Un doubleur de fréquence délivre le signal multipliant à 31 kHz. Vous noterez que la fréquence de cette onde pilote est la même que celle de la fréquence de balayage horizontal ; les impulsions de synchronisation horizontale pourraient très bien délivrer cette référence. Mais en incluant une onde pilote dans le canal son à 4,5 MHz, il est plus simple de démoduler le son stéréophonique de la télévision dans de petits récepteurs pour la réception BF de l'AM, la FM et du son de la télévision. Une matrice additionne ensuite les signaux G+D et G-D pour former le canal gauche et les soustrait pour former le canal droit. Rappelons que l'on appelle une matrice tout réseau qui réalise des combinaisons linéaires. Dans le schéma

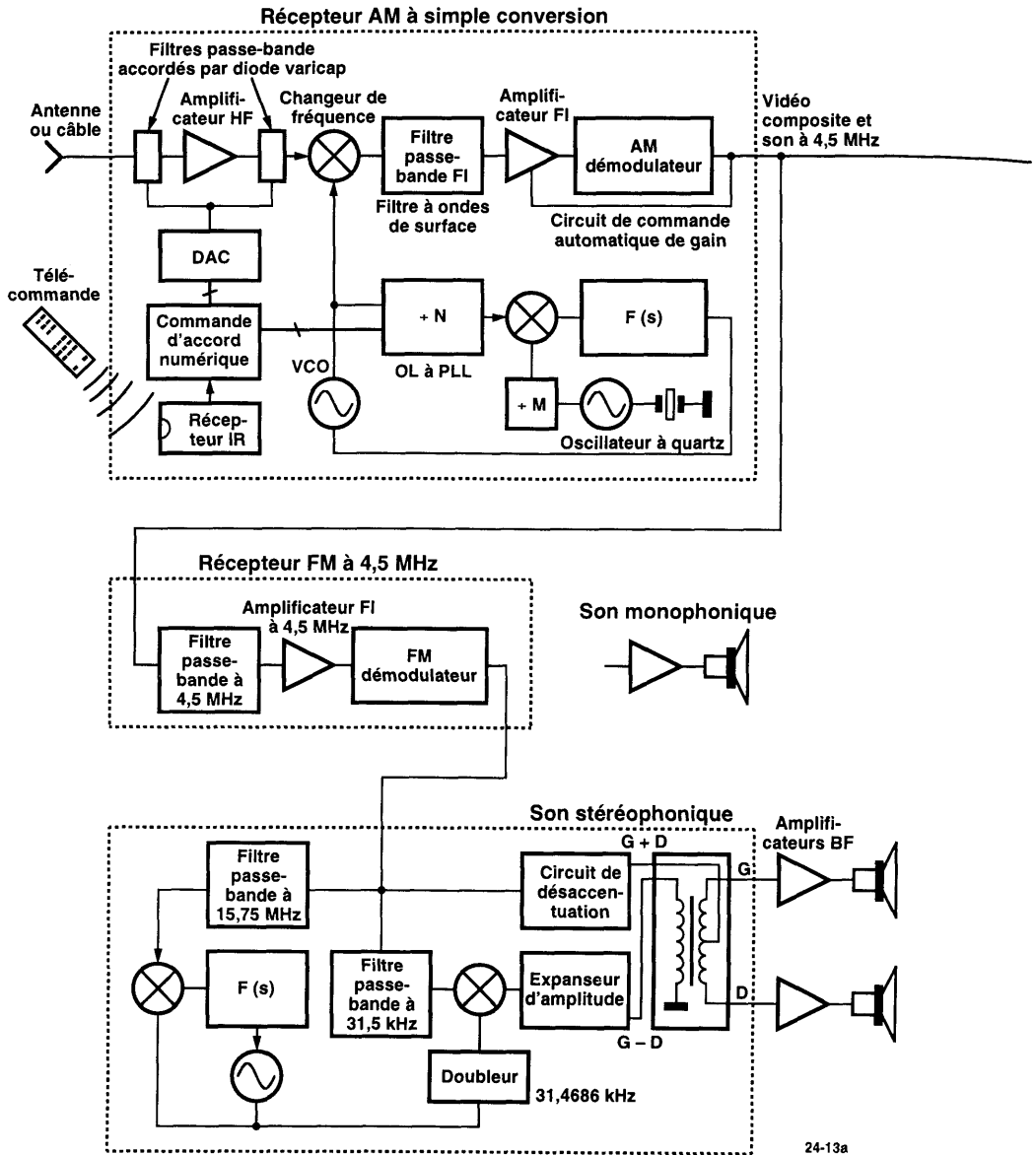


Figure 24-13. Schéma synoptique d'un récepteur de télévision.

synoptique, cette opération de matricage est exécutée par un transformateur. Le plus souvent, un réseau de résistances ou un réseau de résistances et d'amplificateurs opérationnels (comme on peut le voir dans le schéma synoptique de la matrice de codage des signaux de chrominance) s'en charge.

L'expandeur d'amplitude est un processeur de signaux non linéaire qui exécute l'opération inverse de celle qui a consisté à réaliser une compression d'amplitude à l'émission sur le signal $G-D$. Ce système ressemble au système Dolby et permet d'améliorer le rapport S/B du signal $G-D$. Nous étudierons ultérieurement au chapitre 28 le réseau de désaccentuation qui est un « égaliseur » linéaire.

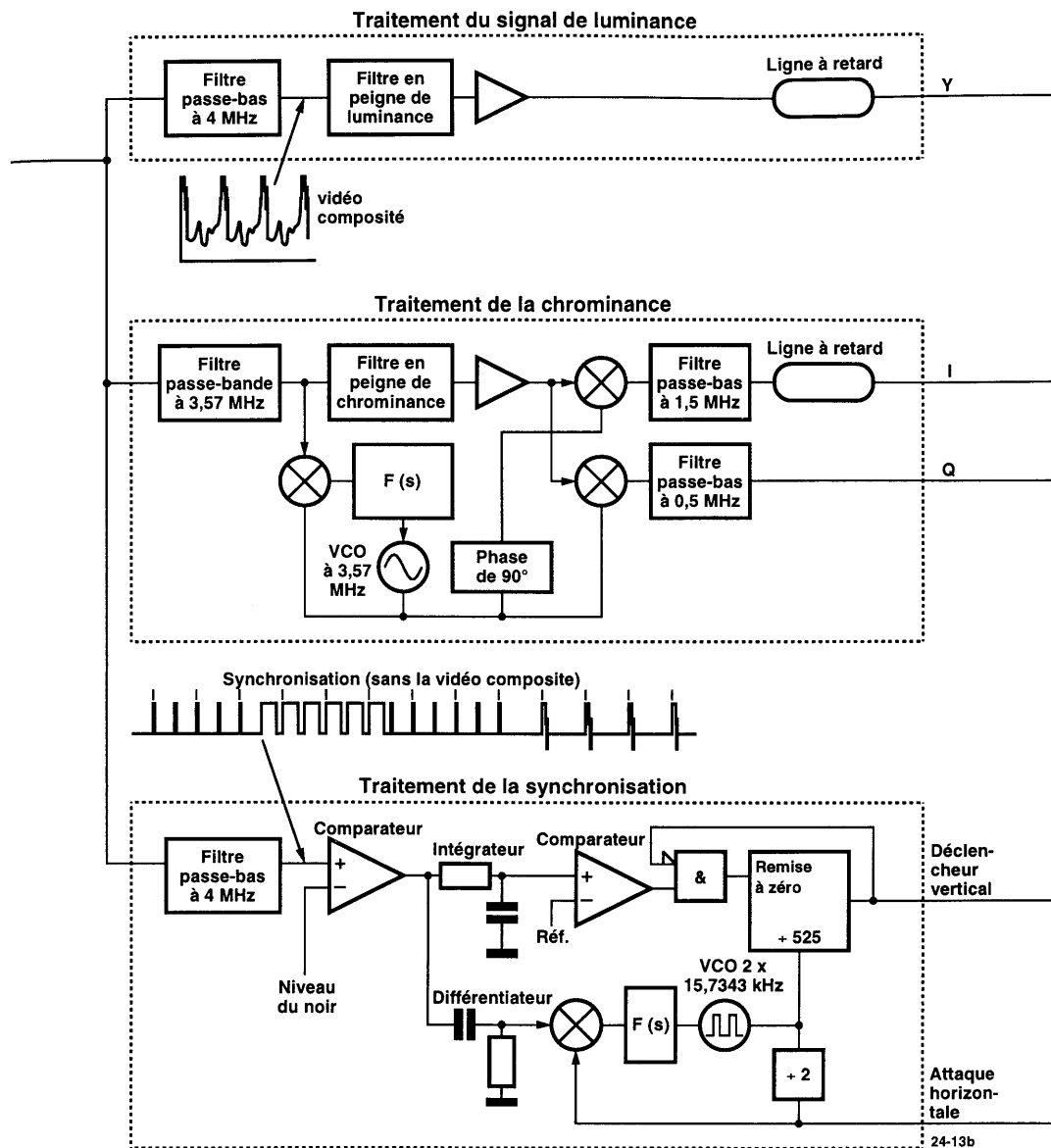


Figure 24-13. Schéma synoptique d'un récepteur de télévision (suite).

La radiodiffusion stéréophonique en FM est identique, si ce n'est que la sous-porteuse est à 38 kHz, l'onde pilote à 19 kHz et qu'il n'y a pas de compression-expansion sur le canal G-D.

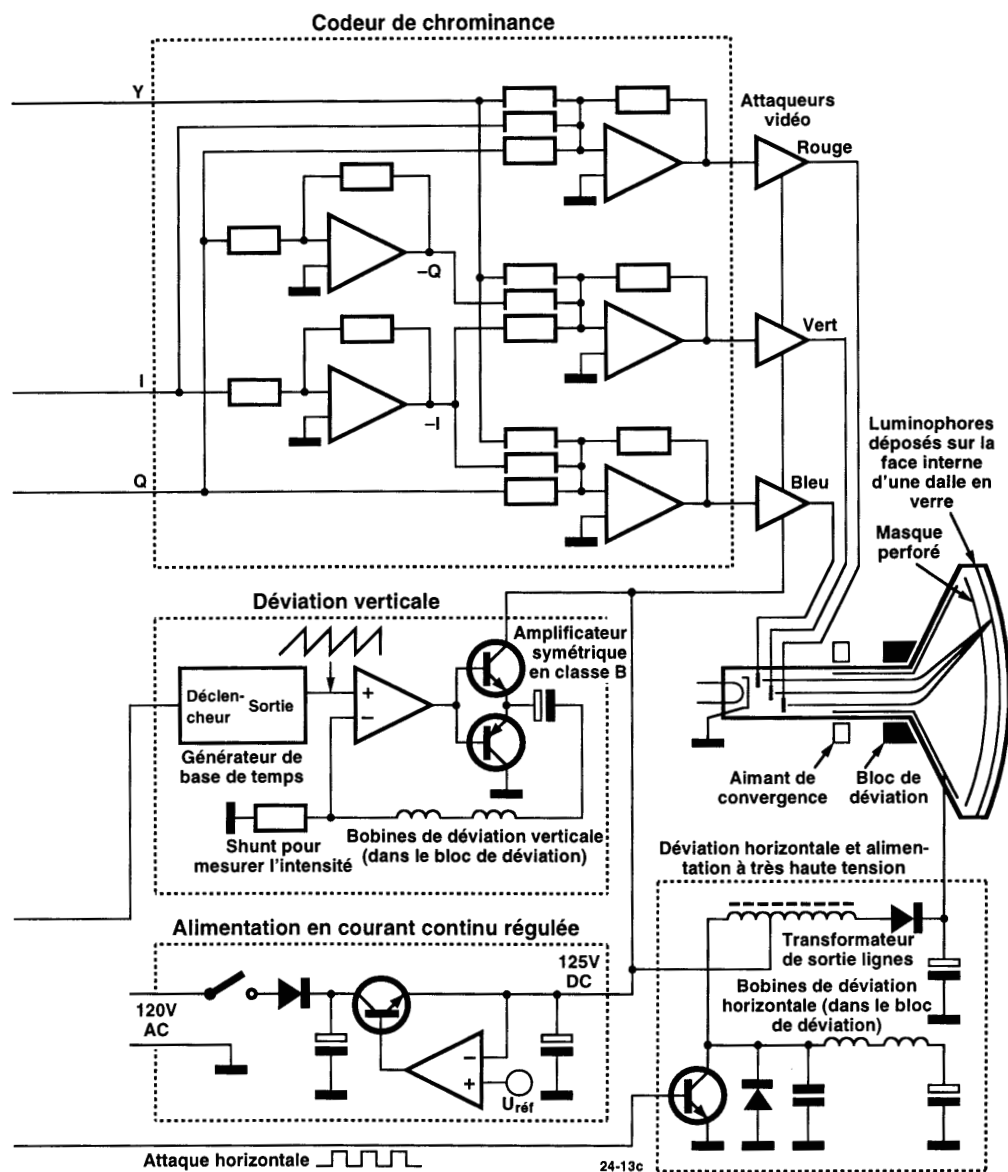


Figure 24-13. Schéma synoptique d'un récepteur de télévision (suite).

Processeur de luminance

Dans ce sous-ensemble, un filtre passe-bas à 4 MHz élimine le signal de son et un filtre en peigne (décrit ci-dessus) le signal de chrominance. Le signal vidéo résultant, Y, ne contient que l'information de luminosité (et les impulsions de synchronisation) et produirait une image correcte si elle était envoyée à un tube-image monochrome (en noir et blanc). Il faut une ligne à retard parce que les filtres passe-bande du

processeur de chrominance retardent forcément les signaux I et Q . Si un retard correspondant n'était pas inclus dans le canal de luminance, la coloration ne serait pas en concordance avec l'image en noir et blanc mais, arrivant plus tard, serait décalée vers la droite.

Processeur de chrominance

Ce sous-ensemble est quasiment semblable à celui du son stéréophonique. Ici, il faut multiplier le signal de chrominance à deux bandes latérales et à porteuse supprimée par une porteuse restaurée localement à la fréquence de la sous-porteuse. L'onde pilote est la salve couleur située sur le palier arrière des impulsions de synchronisation horizontale. Cette salve fournit le signal de référence à une boucle à phase asservie qui délivre elle-même la porteuse locale. On ne voit pas sur le schéma synoptique un interrupteur électronique, la porte de salve, commandé par le circuit de synchronisation pour n'envoyer le signal de référence que pendant la salve. Le rapport S/B de la boucle est ainsi amélioré. La porteuse locale attaque un détecteur de produit (multiplicateur) pour démoduler le signal I . Un réseau déphaseur à 90° délivre la seconde porteuse locale à un second détecteur de produit pour démoduler le signal Q . La largeur de bande du signal Q est inférieure à celle du signal I , c'est pourquoi le filtre passe-bande situé après le démodulateur Q est plus étroit que celui qui se trouve après le démodulateur I . Il est donc indispensable d'insérer une ligne à retard dans le canal I pour égaliser les retards.

Processeur de synchronisation

Un comparateur, dont on a fixé le seuil au niveau du noir, élimine le signal vidéo et délivre un train d'impulsions de synchronisation net. Comme nous l'avons déjà expliqué, un simple différentiateur RC fournit ensuite les impulsions de référence de ligne et un intégrateur RC les impulsions de référence de trame. Un VCO verrouillé sur les impulsions de référence de ligne crée une base de temps de ligne stabilisée par effet de volant. Le VCO fonctionne en réalité à deux fois la fréquence de ligne ; sa fréquence est ensuite divisée par deux pour fournir 15,734 kHz au détecteur de phase à PLL et pour attaquer les circuits de déviation horizontale. On divise par 525 la fréquence du VCO pour disposer d'une base de temps de trame tout aussi stable. Ce diviseur doit fonctionner avec la bonne phase pour que l'alignement vertical de l'image soit correct ; le circuit de division par 525 dispose d'une entrée de remise à zéro qui est activée lorsque la coincidence entre la sortie du compteur et les impulsions de synchronisation de trames a échoué plusieurs fois de suite.

Matrice de décodage et tube-image

Ce sous-ensemble réalise les combinaisons linéaires appropriées des signaux Y , I et Q qui permettent d'obtenir les signaux vidéo R , V et B . Certains coefficients sont négatifs, c'est pourquoi on trouve des inverseurs qui fournissent les signaux $-I$ et $-Q$. Les signaux R , V et B sont amplifiés à des niveaux qui atteignent plusieurs dizaines de volts pour attaquer les canons respectifs d'électrons du tube-image tricolore. Le circuit d'extinction du faisceau d'électrons, qui fonctionne durant les périodes de retour de ligne et de trame n'est pas représenté sur ce schéma synoptique. Si le fonctionnement de l'ensemble est normal, ces signaux d'extinction, produits localement, sont redondants parce que le signal composite vidéo les contient déjà.

Circuits de déviation

Le courant en dents de scie nécessaire aux bobines de déviation verticale est fourni par un amplificateur symétrique linéaire en classe B. Un générateur de tension en dents de scie, déclenché par les circuits de synchronisation, délivre une forme d'onde de référence précise. Un shunt de faible résistance permet d'obtenir une tension proportionnelle au courant de sortie ; cette tension est une contre-réaction qui permet au courant de suivre le signal de référence. Le circuit de déviation horizontale a été étudié dans le chapitre qui traitait des alimentations à découpage.

Alimentation

Un circuit simple de redressement mono-alternance à condensateur en tête délivre approximativement 140 V_{dc}. Un régulateur en série abaisse cette tension à 125 V – ce qui convient aux transistors qui équipent les sous-ensembles de puissance du récepteur (déviation, attaqués vidéo, et amplification BF de puissance). Les circuits du récepteur forment essentiellement une charge constante ; le régulateur n'a donc à corriger que les variations de la tension du secteur. Une alimentation basse tension d'environ 15 V, nécessaire au fonctionnement des circuits de traitement du signal, est souvent prélevée sur un enroulement secondaire du transformateur de sortie lignes. Un autre enroulement secondaire délivre quelques volts pour alimenter le filament du tube à rayons cathodiques. Dans ce genre d'alimentation, l'une des lignes du secteur est mise à la masse, généralement au châssis métallique. Avec un tel « châssis chaud », il est *vital* que le coffret *isole parfaitement* toute personne d'un contact possible avec celui-ci. Il serait évidemment possible d'installer un transformateur de puissance 60 Hz à l'entrée du récepteur pour isoler le châssis, mais son poids et son coût font reculer les constructeurs. On pourrait aussi installer un petit transformateur pour isoler les sous-ensembles auxquels l'utilisateur a accès : les prises de jack audio et vidéo permettant de connecter un magnétoscope, etc.

Télévision numérique

Au moment où cet ouvrage paraît, un successeur à la norme NTSC, agée de plus de cinquante ans, est tout proche. Cette norme naissante, connue sous les noms de *High Definition Television* (HDTV), *Advanced Television* (ATV) ou *Digital Television* (DTV), est basée sur le traitement et l'utilisation de mémoires (mémorisation de la trame) numériques. Aux États-Unis, un consortium dénommé *Grand Alliance*, a proposé un procédé dans lequel le traitement et la mémorisation des données permet d'obtenir une compression des données de l'ordre de soixante fois. Cette norme comporte une technique de compression vidéo (MPEG-2) fondée sur la redondance temporelle et spatiale des informations vidéo. Souvenez-vous que les mêmes redondances sont déjà exploitées, à un degré moindre, dans le procédé NTSC avec l'entrelacement temporel des trames et l'entrelacement spectral des signaux de luminance et de chrominance. Le consortium *Grand Alliance* est devenu l'*Advanced Television Systems Committee* (ATSC). En Europe, le « *European Launching Group* » (ELG), devenu *Digital Video Broadcasting* (DVB), a développé un système semblable, basé également sur la technique de compression MPEG-2. Les deux systèmes ont adopté le même format d'image 16:9 en remplacement du format actuel de 4:3.

Un procédé transitoire, PALplus, a été développé en Europe pour fournir des émissions en 16:9 à des récepteurs de télévision spéciaux PALplus. Ce procédé est compatible avec l'« ancien » PAL. Au cours d'une émission PALplus, un poste de télévision PAL classique visualise une fenêtre de format « boîte aux lettres » de 430 lignes avec des bandes noires en haut et en bas de l'écran. Les 72 lignes contenues dans chaque bande contiennent un signal d'« aide » qui permet aux récepteurs PALplus d'obtenir une image de

576 lignes. Les récepteurs PALplus renferment une mémoire de trame numérique mais le procédé met en œuvre une transmission analogique et non numérique, ce n'est donc pas de la HDTV. Le développement des procédés de télévision numérique s'est effectué tellement rapidement que le procédé PASplus semble être arrivé trop tard. D'un autre côté, le matériel de studio de télévision comporte des ensembles de traitement et de mémorisation numériques et la radiodiffusion directe par satellite a déjà mis en œuvre la compression et l'émission numériques. Ces systèmes de transmission par satellite nécessitent des convertisseurs qui produisent les signaux classiques NTSC, PAL ou SECAM (ou les signaux audio et vidéo en bande de base) à destination des récepteurs de télévision existants et leurs technologies comportent les éléments des prochaines normes de télévision numériques.

Procédé ATSC

Le procédé proposé par ATSC acheminerait deux programmes de résolution classique ou un programme à haute définition sur un seul canal télévision de 6 MHz avec un débit d'informations net de 18 Mbit/s. Dans le cas d'une liaison terrestre, les informations audio et vidéo sont codées en mots de 3 bits. Une variante, réservée à des liaisons par câble, utilisent des mots de 4 bits. Un convertisseur numérique/analogique transforme les mots de 3 bits en un signal analogique à 8 niveaux qui module la porteuse en amplitude. Celle-ci est presque totalement supprimée (on n'en conserve qu'une fraction qui sert d'onde pilote). Au niveau du récepteur, le signal est immédiatement numérisé après une détection synchrone pour restaurer les mots de 3 bits. Comme dans le procédé NTSC, un filtre passe-bande situé dans l'émetteur élimine la quasi totalité de la bande latérale inférieure pour réduire la largeur du spectre – le procédé conserve le « moignon » de bande latérale. Le signal émis est toujours quasiment sinusoïdal, comme dans de nombreux autres systèmes HF à bande étroite. Ici, la distribution d'amplitude se situe autour de chacun des huit niveaux de modulation (il s'agit d'un système numérique « octaire » et non d'un système numérique binaire). Les signaux ATSC et NTSC imposent des récepteurs totalement différents mais les signaux sont au moins compatibles en ce sens qu'ils peuvent coexister sur des canaux adjacents (inutilisés à présent pour éviter toute interférence entre des émissions NTSC adjacentes) sans aucune interférence mutuelle. Il a été prévu que les organismes de télédiffusion fassent cohabiter d'ici à quinze ans, des émissions NTSC et ATSC en éliminant progressivement les émissions NTSC. Vous trouverez ci-dessous une description des éléments du procédé ATSC.

Compression vidéo

Pour réduire le débit d'informations d'un facteur égal à cinquante, voire plus, on se sert des redondances temporelles et spatiales présentes dans la succession d'images fixes qui forment une image animée. L'image est mise à jour par blocs de pixels (8×8 pixels pour les liaisons terrestres ATSC et 16×16 pixels pour les liaisons par câble). D'une trame à l'autre, le codeur vidéo détecte les changements qui se sont produits au sein de chaque bloc et cette information de différences est exploitée au niveau du récepteur pour mettre à jour les blocs, c'est-à-dire modifier les données dans la zone concernée du bloc de la mémoire numérique de trame. À l'exception des changements complets de scènes, les différences d'une trame à l'autre sont dues au déplacement d'éléments au sein de la trame (sujet en mouvement) ou à celui de l'image dans sa totalité (déplacement de la caméra). De trame en trame, un bloc donné va le plus souvent changer simplement de position.

Dans le procédé ATSC, le codeur détermine, pour chaque bloc d'une nouvelle trame, le bloc le plus approprié qu'il convient de déplacer par rapport à la trame précédente (en recherchant le minimum de la somme des valeurs absolues des différences des valeurs de luminosité des pixels). La position du bloc déplacé est spécifiée par un *vecteur de déplacement* : par exemple 1,1 indique qu'il convient de déplacer le bloc approprié d'un pixel vers le haut et d'un pixel vers la droite. Ces vecteurs de déplacement font partie des informations de mise à jour vidéo et permettent au récepteur de construire la nouvelle trame en réutilisant la plus grande partie de la trame précédente mémorisée. Le reste de l'information comprend les différences entre les pixels dans ces blocs compensés en déplacement et ceux des nouveaux blocs. Le résultat de la compensation en mouvement est que les différences sont minimales ; la corrélation temporelle représente la première étape de la compression des données.

L'étape suivante consiste à exploiter la cohérence spatiale de l'image, c'est-à-dire la corrélation généralement élevée qui existe entre des pixels adjacents. Un écran constitué de pixels non corrélés aurait l'aspect d'une neige très fine. Au niveau du codeur, le bloc de 8×8 , qui renferme les différences entre pixels, subit une transformation de Fourier à deux dimensions (en particulier, une transformation cosinusoidale discrète). À cause de la corrélation spatiale, la plupart des 8×8 coefficients résultant de la transformation ont une valeur très faible. Par exemple, si un bloc de pixels donné représente une zone dont la luminosité croît linéairement dans la direction horizontale, seuls deux coefficients prendront des valeurs significatives. Lorsque la valeur des coefficients dans un bloc est négligeable, le flux résultant de données numériques ne comporte pratiquement que des zéros et un codage approprié (indication du nombre de zéros consécutifs) permet d'obtenir un taux de compression élevé. On peut accroître le nombre de zéros consécutifs en procédant à une lecture particulière en zigzag des coefficients résultant de la transformation des différences. De plus, on utilise un codage de longueur variable pour les coefficients qui ne sont pas nuls ; le téléspectateur peut admettre une quantification plus grossière pour les fréquences spatiales élevées.

Il est impossible de prévoir systématiquement une trame à partir de la précédente ; les changements de scènes nécessitent un renouvellement complet des données. On inclut de temps en temps des « trames de rafraîchissement » (*refresh frames*) dans le flux des autres trames.

La norme proposée propose deux résolutions normalisées : 1280×720 pixels (avec un balayage non entrelacé à des vitesses de 60, 30 et 24 trames par seconde) et 1920×1080 pixels (avec un balayage non entrelacé à 24 trames par seconde et un balayage non entrelacé ou entrelacé à 30 trames par seconde). Dans un cas comme dans l'autre, le format de l'image est 16:9.

Couleur, son et paquets

Le procédé ATSC, comme le procédé NTSC, transmet un signal de luminance et deux signaux de chrominance à définition réduite. L'information de chrominance est lissée et échantillonnée à une fréquence spatiale égale à la moitié de celle des signaux de luminance. Le son (de la qualité de celui d'un CD), tout comme la vidéo, subit une compression numérique. Les blocs de données son et vidéo sont multiplexés dans le temps en paquets séparés. Les autres paquets peuvent véhiculer des données destinées à d'autres besoins. Au niveau du récepteur, une mémoire tampon permet de concaténer les paquets audio et vidéo successifs et de délivrer avec les bons débits, les signaux traités et décodés. On appelle *système de transport* la procédure de confection des paquets qui gère le mélange des paquets, leurs longueurs (peut-être variable), les en-têtes, etc. Lors de l'émission, on ajoute des symboles de synchronisation ainsi qu'une information redondante destinée à la détection et à la correction des erreurs.

Bibliographie

- 1 K.B. Benson (réédition) (1986), *Television Engineering Handbook*. New York, McGraw-Hill. (c'est une mise à jour de l'ouvrage *Television Engineering Handbook*, D.G. Fink (réédition), New York, McGraw-Hill, 1957.)
- 2 K. Challapali, X. Lebegue, J.S. Lim, W.H. Paik, R.S. Girons, E. Petajan, V. Sathe, P.A. Snopko et J.D. Zdepksi (1995), « The Grand Alliance System for USA HDTV », *Proc. IEEE* 83 : 139-50.
- 3 K.G. Jackson et G.B. Townsend (réédition) (1991), *TV & Video Engineer's Reference Book*. Oxford, Butterworth Heineman. (cet ouvrage contient les informations les plus récentes à ce jour)
- 4 Engineering Handbook, 8^e édition (1993). National Association of Broadcasters, Washington, DC.
- 5 A. Netravali et A. Lippmann (1995), *Digital Television : A Perspective*, *Proc. IEEE*, 83 : 834-42.
- 6 E. Petajan (1995), « The HDTV Grand Alliance System », *Proc. IEEE*, 83 : 1094-105.

Problèmes

24.1 Un signal capté par l'antenne d'un récepteur de télévision est composé du signal direct (reçu de l'émetteur en trajet direct) et d'un second signal faible (provenant de la réflexion du signal émis sur une tour métallique située en dehors du trajet direct). Si le trajet emprunté par le second signal est plus long de 1 km que celui du signal direct, quelle sera la position de l'image « fantôme » sur l'écran du récepteur ?

24.2 Les images d'un film tourné à 24 images par seconde sont diffusées par la télévision à 60 trames par seconde grâce à une technique dénommée *3:2 pull down*. Pour une image du film, on envoie deux trames de télévision. L'image suivante du film est mise en place et on envoie trois trames de télévision. Montrez que cette procédure revient à réaliser un film à 24 images par seconde.

24.3 Dessinez le schéma synoptique d'un récepteur de télévision NTSC qui dispose d'une fonctionnalité d'incrutation d'une « image dans l'image », c'est-à-dire d'une fenêtre dans laquelle on peut visualiser une autre chaîne. Expliquez le fonctionnement de chacun des sous-ensembles.

24.4 Prenons le cas d'un signal monochrome de télévision à haute définition de caractéristiques 1920×1080 pixels et un balayage non entrelacé de 60 trames par seconde. Si la luminosité de chaque pixel est codée sur 8 bits et si l'on n'utilise aucune compression, quel sera le débit d'informations (réponse : 995 Mbit/s) ? Quel rapport de compression faut-il appliquer, si ce signal doit être émis à 18 Mbit/s sur un canal de largeur classique ?

24.5 Supposons qu'une mire de télévision NTSC soit formée de cinq barres verticales de largeur égale mais de couleurs différentes. Admettons que les barres ont la même luminosité (ou brillance) et qu'elles sont faiblement colorées (aucune saturation). Dessinez la forme d'onde d'une ligne vidéo.

24.6 Le balayage entrelacé procure une meilleure définition dans le cas de scènes fixes mais peut produire des artefacts lorsque les objets se déplacent. Pensez à une situation dans laquelle le balayage entrelacé ferait « disparaître » un objet en mouvement.

Index

A

Abaque de Smith	72	dynamique	153
Adaptation	9	émetteur commun	17, 155
amplificateur	24	en pont	20
bruit	281	gain, largeur de bande, impédances	150
guide d'ondes	208	hautes fréquences	23
impédance	8, 10, 12, 14, 71, 73, 75, 77	HF	149, 150, 152, 154, 156
réseau	72	intermodulation	151
transformateur	8	large bande	153
<i>Adler</i>	285	limiteur	149
Admittance	4	linéaire	16, 149
complexe	4	linéarité	152
Alimentation	79, 80, 82, 84	neutrodynage	151
autorégulation	80	parallèle	58, 62
découpage	138	paramètres Z	149
ondulation	81	petits signaux	23, 149, 150, 152, 154, 156
redresseur à deux alternances	79	puissance	60
redresseur à une alternance	81	puissance, adaptation	24
régulée	82	<i>push-pull</i>	19
télévision	234	rendement	18, 22, 60
triphasée	83, 85	série	58, 61
AM	86	<i>slew-rate</i>	19
capacité de transmission	269, 271	stabilité	150
détecteur à diode	251	surcharge	151
domaine fréquentiel	87	symétrique	19, 22, 172
domaine temporel	86	télévision	228
émetteur-récepteur	86	<i>totem pole</i>	20
formes d'ondes	88	tube	58
Amplificateur		vitesse de balayage	19
asymétrique	20	Analyseur	
bande étroite	153	réseau	316
base commune	17, 156	spectre	317
basses fréquences	21	Antenne	272, 274, 276, 278, 280
BF	155	aire équivalente	274
bruit	181, 182, 184, 186, 281, 282, 284, 286	directivité	273
cascode	154	gain	273
classe A	18, 22	symétriseur	187, 188, 190, 192, 194, 196, 198, 199, 200, 201, 202
classe B	20, 22	transformateur	187, 188, 190, 192, 194, 196, 198, 200, 202
classe C	55, 57, 59	<i>Armstrong</i>	104
classe D	61, 63	ATSC	235, 237
classe F	60	Auto-duplexage	244
collecteur commun	17		
conception	155, 157		
courant alternatif	21, 23, 25		

Autocorrélateur		thermique	181
1 bit	305	transistor	281
matériel	305	<i>Butterworth</i>	27
Autorégulation	80		
B		C	
<i>Balun</i>	187, 198	Câble	
<i>Bessel</i>	27	bifilaire	65
BFO, voir Oscillateur de battement		coaxial	65, 207
BLU	96	Cascode	154
avantages	97	Champ	
classe C	100	électrique	203, 272
classe D	100	électromagnétique	204
création	97, 99, 101	magnétique	204, 205, 272
méthode de filtrage	97	Changeur	
méthode de mise en phase	98	fréquence, à commutation	43
méthode de Weaver	99	fréquence, non linéaire	45
Boucle à phase asservie	114, 116, 118, 120, 122, 124	fréquence, symétrique double	44
analogie mécanique	115	Chrominance, processeur	233
analyse linéaire	118, 119	Circuit, analyse	3
démodulateur	256	<i>Colpitts</i>	105
dynamique	117	Combinateur	
filtre de boucle	118	puissance	174
plage de fonctionnement	122	<i>Wilkinson</i>	174
principe	114	Commutateur TR <i>(Transistor)</i>	244, 246, 248, 250
récepteur	123	diode	247
régime transitoire	120	ligne	246
stabilité	122	Compression	
temps de verrouillage	122	impulsions radar	310, 311
type I	119	vidéo	235
type II	119	Conductance	4
Bruit		bruit	283
adaptation	281	Convertisseur	
amplificateur	181, 182, 184, 186, 281, 282, 284, 286	analyse	133
amplificateurs en cascade	183	autres types	138
circuits en parallèle	284	<i>boost</i>	137
conductance	283	<i>buck</i>	133
cosmique	294	<i>buck/boost</i>	135
facteur	182, 281	découpage	133, 134, 136, 138, 140, 142
FM	266, 267	direct	139
mesure	185, 285, 287	<i>flyback</i>	138
oscillateur	288, 290, 292	<i>forward</i>	139
paramètres	184	fréquence	41, 42, 44, 46
phase	130, 131, 291	indirect	138
résistance	281, 283	liaison par transformateur	138
schémas équivalents	282	Coupleur	
synthétiseur	128	directif	144, 168, 177, 179, 316
température	181	hybride	168, 170, 172, 174, 176, 178, 180
		inverse	177

D

<i>Dahlke</i>	282
Démodulateur	251, 252, 254, 256, 258, 260, 262
FM	256, 257, 259
PLL	256
quadrature	257
<i>Dempingscoefficient</i>	119
Désaccentuation	269
Détecteur	251, 252, 254, 256, 258, 260, 262
diode	251, 253
ligne à retard	257
parfait	252
pente	258
phase	121
produit	255
puissance	261, 263
quadratique	261
réel	252
synchrone	255
tachymétrique	256
Détecteur-produit	96
Détection	123
BLU	254
code Morse	254
enveloppe	251
Déviaton, circuits	234
Discriminateur	257
<i>Foster-Seeley</i>	259
Duplexeur	246

E

Effet Miller	154
Émetteur, télévision	228
Émetteur-récepteur AM	86
Émetteur-suiveur	17
Équations de Maxwell	1, 205, 273

F

Facteur de bruit	182, 281
Facteur de qualité, voir Facteur <i>Q</i>	10
Facteur <i>Q</i>	10
incidence	165
<i>Fessenden</i>	50
Filtre	27
autres types	166, 167
<i>Bessel</i>	27
boucle à phase asservie	118
<i>Butterworth</i>	27

constantes localisées	27
conversion	32, 33
exemple	162, 163
normalisé	28
passe-bas	27, 35
passe-haut	40
passe-tout	40, 101
peigne	224
procédure d'accord	166
radiospectrométrie	303
résonateur	158, 160, 162, 164, 166
<i>Tchebychev</i>	27
FM	264
bande étroite	265
bruit	266, 267
capacité de transmission	269, 271
démodulation	256, 257, 259
désaccentuation	269
excursion	264
indice	265
large bande	266
préaccentuation	269
S/B	267
spectre	265
<i>Foster</i>	259
<i>Foucault</i>	193
<i>Fourier</i>	303
Fréquence	
attribution	2, 86
changeur	41
convertisseur	41, 42, 44, 46
double changement	52
excursion	264
image	51
mesure d'impédance	316
multiplicateur	56
multiplication	266
spectre	265
synthèse directe	125
synthèse directe numérique	127, 129
synthèse indirecte	126
synthétiseur	125, 126, 128, 130, 132

G

<i>Gilbert</i>	42
Guide d'ondes	
adaptation	208
câble coaxial	207
champ	203
courants	206

diaphragme 208
 dimensions 204
 forme 204
 impédance 207
 jonction à quatre accès 210
 jonction à trois accès 209
 longueur d'onde 205
 pertes 210, 211
 propagation 203, 204, 205
 T magique 210
 vitesse de propagation 207

H

Hartley 104
Haus 285
 Hybride 144
 180° 170
 90° 170
 anneau 175
 application 170
 autres types 174, 175
 composants discrets 176
 coupleur 168, 170, 172, 174, 176, 178, 180
 duplexeur 246
 échelle 175
 guide d'ondes 177
 quadrature 170, 171, 173
 T magique 177
 transformateur 169

I

Impédance 4
 adaptation 8, 10, 12, 14, 71, 73, 75, 77
 caractéristique 65
 coefficient de réflexion 71, 73, 75, 77
 complexe 4
 inverseur 158, 159, 161
 mesure 314, 315, 317, 319
 pont 314
 pont résistif 146
 transfert direct 149
 transfert inverse 150
 Inductance
 fuite 191
 magnétisation 189
 mutuelle 189
 Interférométrie 296, 297
 imagerie 297

Ionosphère 277, 279
 propagation 277
 réflexion 278

J

Jansky 294
 Jonction
 guide d'ondes 209, 210
 tourniquet 244

K

Khinchin 304

L

Ligne
 commutateur TR 246
 impédance caractéristique 66
 modification d'impédance 67, 69
 modulateur 239, 241, 243
 notions 65
 pertes 69, 210
 retard 240, 257
 transformateur 197
 transmission 65, 66, 68, 70
 vitesse de propagation 66
 Lignes
 synchronisation 217
 amplificateur 16, 18, 20, 22, 24, 26
 luminance, processeur 232

M

Matrice de décodage 233
Maxwell, équations 1
 Mélangeur
 bande latérale unique 47
 diode 45, 47
 Mesure
 bruit 285, 287
 impédance 314, 315, 317, 319
 pont d'impédance résistif 146
 puissance 144, 313
 tension 313
Miller 154
 Mode
 autres types 204
 fondamental 204, 205

Modulateur		de battement	94
classe A	90	dynamique	109
classe B	90	en série	107
classe S	91	exemple	111, 113
impulsions radar	238, 240, 242	<i>Hartley</i>	104
ligne	239, 241, 243	involontaire	106
numérique/analogique	92	linéaire	289, 291
Modulation	3	non linéarité	292, 293
à haut niveau	89, 91	sinusoïdal	103, 105
à porteuse supprimée	94, 96, 98, 100	stabilité	109
amplitude	86, 88, 90, 92	<i>Overspreading</i>	301
angulaire	264, 265		
fréquence	264, 266, 268, 270	P	
FSK	123	PAL	225
indice	265	Paramètres	
phase	264, 266, 268, 270	Y	157
PSK	123	Z	149, 157
quadrature	222	Filtre	27
télévision	219	Phase, détecteur	121
Monoplexeur	244	PLL	114
<i>Morse</i>	254	PM	264
Multiplicateur	41	bande étroite	265
cellule de Gilbert	42	Préaccentuation	269
détecteur de phase	121	Procédé de télévision	215, 216, 218, 220, 222, 224, 226, 228, 230, 232, 234, 236
fréquence	56, 266	ATSC	235, 237
tension	60	disques de Nipkow	215
		NTSC	216, 217, 219
N		PAL	221, 225
Neutrodynage	151	SECAM	221, 227
<i>Nipkow</i>	215	Propagation	272, 274, 276, 278, 280
NTSC	216, 217, 219	constante	69
		diurne et nocturne	278
O		ionosphère	277
Onde stationnaire	144, 146, 147, 148	vide	273
effets	147	Puissance	
taux	147	combinateur	174
Ondes électromagnétiques	1	densité spectrale	303
bilan d'une liaison	275	détecteur	261, 263
réflexion	278	diviseur	174
Ondes radioélectriques		mesure	144, 313
propagation	272, 274, 276, 278, 280	spectre	289, 291
Ondulation	81		
Oscillateur	102, 104, 106, 108, 110, 112	Q	
à relaxation	102	Q-mètre	315
à résistance négative	108	Quadrature, hybride	170, 171, 173
<i>Armstrong</i>	104	Quadripôle	149
bruit	288, 290, 292	schéma équivalent	282
<i>Colpitts</i>	105, 110		
commandé en tension	103		

R

Radar	
auto-duplexage	244
commutateur TR	244, 246, 248, 250
compression d'impulsions	310, 311
couverture latérale	300
duplexeur équilibré	246
modulateur d'impulsions	238, 240, 242
monoplexeur	244
monostatique	244
ouverture dynamique	300
thyatron	238
Radarastronomie	298, 299, 301
cartographie à retard Doppler	300
Lune	298
<i>overspreading</i>	301
Vénus	299
Radioastronomie	294
Radiométrie	295
Radiospectrométrie	303, 304, 306, 308, 310, 312
<i>Reber</i>	294
Récepteur	
amplification	48
amplification directe	49
caractéristiques	48
commande automatique de gain	53
fréquence-image	51
galène	49
hétérodyne	50
PLL	123
radio	48, 50, 52, 54
réducteur de bruit	53
superhétérodyne	50, 51
télévision	139, 228, 229, 231, 233
traitement numérique	53
Redresseur	
deux alternances	79
triphasé	83, 85
une alternance	81
Réseau	
adaptation	72
analyseur	316
deux accès	149
en double L	13
en échelle	5
en L	9, 11
en p	12
en T	12
ligne à retard	240
phase minimale	40
Résistance, bruit	281, 283

Résonance	
parallèle	4
série	4
Résonateur, filtre	158, 160, 162, 164, 166
<i>Rothe</i>	282

S

SECAM	227
<i>Seeley</i>	259
Signal	
bande latérale unique	96
chrominance	223
deux bandes latérales et porteuse supprimée	94
luminance	222
vidéo	216
<i>Smith</i>	72
Son stéréophonique	229
Spectromètre	
acousto-optique	308
compression	309
transformation de Fourier	307
Spectrométrie	296, 303
autocorrélation	304, 305
Susceptance	4
Symétriseur, antenne	187, 188, 190, 192, 194, 196, 198, 199, 200, 201, 202
Synchronisation	
lignes	217
processeur	233
trames	218
Synthétiseur	
bruit	128
bruit de phase	130, 131
continuité de phase	130
fréquences	125, 126, 128, 130, 132
mélange et division	126
vitesse de commutation	130

T

<i>Tchebychev</i>	27
Télévision	
autres normes	221
circuit de sortie de lignes	139, 141, 143
compatibilité	222
couleur	221, 223, 229
émetteur	228
historique	215
image	215
modulation	219

numérique	234
procédés	215, 216, 218, 220, 222, 224, 226, 228, 230, 232, 234, 236
récepteur	228, 229, 231, 233
son	220
son stéréophonique	229
sous-porteuse	222
Température de bruit	
antenne	181
récepteur	181
Théorème	
réciprocité	275
<i>Wiener-Khinchin</i>	304
Thyratron	238
TOS	147
Trames, synchronisation	218
Transconductance, cellule de Gilbert	42
Transformateur	
accord décalé	192
adaptation	8
analogie mécanique	190
antenne	187, 188, 190, 192, 194, 196, 198, 200, 202
coefficient de couplage	189
conception	194
convertisseur	138
courants	187, 193
hybride	169
idéal	187
inductance mutuelle	189
ligne de transmission	197
noyau feuilleté	193
noyau magnétique	193, 195
paramètres physiques	196
schéma équivalent	188, 189
Transformation de Fourier	303
Transistor	
bruit	281
schéma équivalent	154
Transmission	
capacité	269, 271
ligne	65, 66, 68, 70
longueur électrique	70
Tube-image	233

W

Wattmètre	144, 146, 148
<i>Weaver</i>	99
<i>Wiener</i>	304
<i>Wilkinson</i>	174

V

VCO	103, 114
Voltmètre vectoriel	313
VSWR, voir TOS	

Jon B. Hagen

COMPRENDRE ET UTILISER L'ÉLECTRONIQUE DES HAUTES-FRÉQUENCES

DE LA GALÈNE À LA RADIOASTRONOMIE PRINCIPES ET APPLICATIONS

Cet ouvrage se veut d'abord facile. Ce n'est pas un livre pour spécialistes, mais il est complet.

La première mission que l'auteur s'est assignée consiste à présenter efficacement les fondements et l'essence des circuits pour radiofréquences, ce qu'il fait en passant en revue tous les principes qui régissent la modulation et la démodulation des radiofréquences, aussi bien pour la transmission sans fil de données au moyen de puces semi-conductrices que pour l'émission radiophonique de puissance.

Parmi les sujets abordés on trouve les filtres, les amplificateurs, les oscillateurs, les adaptateurs, les modulateurs, les amplificateurs à faible bruit, les boucles à asservissement de phase, les lignes de transmission et les transformateurs. Pour chacun d'entre eux, la rigueur analytique est mise au profit d'une compréhension en profondeur des propriétés et du fonctionnement. Des applications de systèmes HF sont présentées et décrites dans des domaines aussi divers que les communications, l'émission radio et TV, le radar et la radioastronomie.

Le livre contient certes de nombreux exercices, mais pour tirer profit de cette lecture, il n'est pas nécessaire de disposer d'un gros bagage théorique. Il faut des connaissances élémentaires en électronique, de quoi analyser les circuits de base. Il s'agit donc d'un manuel idéal pour un cours d'électronique, mais aussi d'un excellent ouvrage de référence pour les chercheurs et les ingénieurs déjà engagés sur le terrain.

Publitrone/Elektor-Paris

ISBN 2-86661-110-1 Cat.: 008083



9 782866 611101

PUBLITRONIC/ELEKTOR