

TABLE DES MATIÈRES

	Page
INTRODUCTION	1
CHAPITRE 1 CONVERSION $\Sigma\Delta$	5
1.1 Introduction	5
1.2 Numérisation d'un signal	5
1.3 Rapport signal à bruit (SNR)	7
1.4 dBFS	8
1.5 Plage dynamique	8
1.6 Convertisseur $\Sigma\Delta$	9
1.6.1 Suréchantillonnage d'un signal	10
1.6.2 OSR	10
1.6.3 Filtrage	11
1.6.4 Mise en forme du bruit	11
1.7 Convertisseur d'ordre élevé	14
1.7.1 Filtre de boucle	14
1.7.2 Distribution des zéros	15
1.7.3 Stabilité	16
1.7.4 Quantification	18
1.7.4.1 Quantification binaire	19
1.7.4.2 Quantification multibit	20
1.8 Structure de réalisation	21
1.8.1 CIFB	22
1.8.2 CRFB	23
1.9 Problématique étudiée	24
1.9.1 Niveaux d'abstractions	24
1.9.2 Calcul des coefficients	25
1.10 Conclusion	26
CHAPITRE 2 OPTIMISATION DES PARAMÈTRES DES MODULATEURS $\Sigma\Delta$	27
2.1 Introduction	27
2.2 Paramètres d'optimisation	27
2.2.1 Intégrateur à capacités commutées	27
2.2.2 Sources de bruit d'un intégrateur	29
2.2.2.1 Bruit de scintillation	30
2.2.2.2 Bruit thermique	31
2.2.2.3 Bruit des amplificateurs opérationnels	34
2.2.3 Plage dynamique des intégrateurs	34
2.2.3.1 Réduction de la distorsion harmonique	36
2.2.3.2 Effet des coefficients interétagés	38

2.2.4	Taille du circuit.....	40
2.3	Techniques d'optimisation existantes	41
2.3.1	Delsig Toolbox	41
2.3.2	Algorithmes d'optimisation	44
2.3.3	Analyses des méthodes existantes.....	48
2.4	Conclusion	50
CHAPITRE 3 SOLUTIONS PROPOSÉES		51
3.1	Introduction	51
3.2	Objectif.....	51
3.3	Mise à l'échelle du modulateur vs SNR	52
3.4	Méthode des fenêtres	55
3.4.1	Définition de la fenêtre.....	56
3.4.2	Séquence de minimisation	58
3.4.3	Exemple d'utilisation	60
3.5	Carences de la méthode des fenêtres.....	63
3.6	Minimisation multicritères	64
3.6.1	Variables d'optimisation	64
3.6.2	Contraintes d'optimisation	66
3.6.3	Définition du point de départ	67
3.6.4	Exemple d'utilisation	68
3.6.4.1	Définition des paramètres	69
3.6.4.2	Résultats de simulation.....	70
3.7	Conclusion	73
CHAPITRE 4 VALIDATION PAR SIMULATIONS AU NIVEAU TRANSISTOR		75
4.1	Introduction	75
4.2	Paramètres de réalisation	75
4.3	Conception du circuit.....	77
4.3.1	Amplificateur opérationnel	77
4.3.2	Échantillonneur bloqueur	79
4.3.3	Quantificateur	81
4.3.4	DAC	83
4.3.5	Modulateur	84
4.4	Résultats de simulations	87
4.4.1	SNR.....	87
4.4.2	Amplitude de sortie des intégrateurs	93
4.4.3	Puissance consommée	96
4.5	Conclusion	97
CONCLUSION.....		99
LISTE DE RÉFÉRENCES		103

LISTE DES TABLEAUX

	Page
Tableau 3.1	Résultats obtenus par la méthode des fenêtres 61
Tableau 3.2	Somme des capacités du modulateur et moyenne de l'amplitude de sortie des intégrateurs pour les trois configurations 70
Tableau 3.3	Valeurs finales des variables d'optimisation et amplitude de sortie de chaque intégrateurs..... 70
Tableau 4.1	Paramètres des amplificateurs opérationnels utilisés dans le circuit 79
Tableau 4.2	Performances obtenues (en dB) par les trois configurations 89
Tableau 4.3	Amplitudes de sortie maximales normalisées des intégrateurs pour les trois configurations. Les valeurs théoriques obtenues par la méthode d'optimisation sont indiquées en parenthèse 95

LISTE DES FIGURES

	Page
Figure 1.1	Schéma bloc général d'un ADC..... 5
Figure 1.2	Courbe de transfert d'un quantificateur typique 6
Figure 1.3	Signal d'erreur du quantificateur de la figure 1.2..... 6
Figure 1.4	Modèle mathématique d'un quantificateur 7
Figure 1.5	Graphique typique de la plage dynamique d'un convertisseur 9
Figure 1.6	Impact du suréchantillonnage sur le bruit de quantification 10
Figure 1.7	Schéma bloc d'un convertisseur $\Sigma\Delta$ du premier ordre 12
Figure 1.8	Modèle mathématique d'un convertisseur $\Sigma\Delta$ du premier ordre..... 12
Figure 1.9	Densité spectrale de puissance du bruit de quantification d'un convertisseur $\Sigma\Delta$ du premier ordre 13
Figure 1.10	Modèle mathématique d'un convertisseur $\Sigma\Delta$ d'ordre élevé..... 14
Figure 1.11	Densité spectrale de puissance du bruit de quantification des convertisseurs $\Sigma\Delta$ d'ordre 1 à 4 15
Figure 1.12	Représentation générale des convertisseurs $\Sigma\Delta$ par un filtre de boucle 15
Figure 1.13	Effet de la distribution des zéros à l'intérieur de la bande du signal 16
Figure 1.14	Effet de $\max NTF(e^{j\omega}) $ sur le spectre de la NTF 18
Figure 1.15	Effet de $\max NTF(e^{j\omega}) $ sur la plage dynamique du convertisseur 19
Figure 1.16	Entrée et sortie d'un convertisseur $\Sigma\Delta$ du deuxième ordre ayant un quantificateur binaire 20
Figure 1.17	Modèle d'un convertisseur $\Sigma\Delta$ du premier ordre ayant un quantificateur binaire 20
Figure 1.18	Structure CIFB..... 23
Figure 1.19	Structure CRFB 24

Figure 2.1	Circuit d'un intégrateur à capacités commutées	28
Figure 2.2	Signaux d'horloges sans chevauchement	29
Figure 2.3	Circuit équivalent d'un intégrateur à capacités commutées durant la phase d'échantillonnage	29
Figure 2.4	Circuit équivalent d'un intégrateur à capacités commutées durant la phase d'intégration	30
Figure 2.5	Modèle équivalent du bruit thermique généré par les interrupteurs	32
Figure 2.6	Circuit d'un intégrateur à capacités commutées à trois entrées.....	33
Figure 2.7	Modèle Simulink utilisé pour simuler l'impact de la plage dynamique limitée des intégrateurs	35
Figure 2.8	Résultats de la simulation du modèle de la figure 2.7 avec et sans restriction de la plage dynamique des intégrateurs	36
Figure 2.9	Relation (idéale et réelle) entre l'entrée et la sortie d'un amplificateur opérationnel	37
Figure 2.10	Convertisseur utilisé pour visualiser l'impact des coefficients b_i sur la distribution de la sortie des intégrateurs	38
Figure 2.11	Distribution de la sortie des intégrateurs lorsque le coefficient $b_1 = a_1$ et que les autres coefficients $b_i = 0$	38
Figure 2.12	Distribution de la sortie des intégrateurs lorsque les coefficients $b_i = a_i$ pour $i \leq N$ et $b_{N+1} = 1$	39
Figure 2.13	Application d'un facteur d'échelle sur un système linéaire.....	39
Figure 2.14	SNR maximal obtenu par des simulations empiriques avec un quantificateur de 1 bit.....	42
Figure 2.15	SNR maximal obtenu par des simulations empiriques avec un quantificateur de 2 bits.....	42
Figure 2.16	SNR maximal obtenu par des simulations empiriques avec un quantificateur de 3 bits.....	43
Figure 2.17	Principe d'optimisation par un algorithme génétique	45
Figure 2.18	Structure étudiée afin de déterminer les paramètres optimaux	46

Figure 2.19	Processus de conception pour les convertisseurs $\Sigma\Delta$ passe-bande.....	47
Figure 2.20	Processus d'optimisation par un critère quadratique	48
Figure 2.21	Structure de réalisation générique.....	49
Figure 2.22	Algorithme permettant d'optimiser à la fois la NTF et la structure de réalisation	49
Figure 3.1	Structure CRFB pour un modulateur du quatrième ordre. Les sources équivalentes de bruit thermique à l'entrée de chaque intégrateur sont étiquetées V_{ni}	53
Figure 3.2	Amplitude de la réponse en fréquence des fonctions de transfert des sources de bruit thermique du modulateur de la figure 3.1	54
Figure 3.3	Effet de la mise à l'échelle du modulateur sur le SNR. Tous les intégrateurs sont ajustés à la même amplitude de sortie	55
Figure 3.4	Effet de la mise à l'échelle individuelle de chaque intégrateur du modulateur sur le SNR	56
Figure 3.5	Fenêtre d'opération pour la méthode des fenêtres.....	57
Figure 3.6	Séquence de minimisation de la méthode des fenêtres.....	59
Figure 3.7	Moyenne des amplitudes de sortie en fonction de l'amplitude minimale de sortie. La pénalité en SNR est de 2 dB	62
Figure 3.8	SNR vs. amplitude d'entrée pour les trois configurations présentées dans le tableau 3.3	72
Figure 4.1	Signaux d'horloges sans chevauchement utilisés par le circuit	77
Figure 4.2	Schéma haut-niveau du convertisseur.....	78
Figure 4.3	Échantillonneur bloqueur.....	79
Figure 4.4	Schéma complet du circuit échantillonneur bloqueur	80
Figure 4.5	Simulation temporelle du circuit échantillonneur bloqueur.....	81
Figure 4.6	Spectre de fréquences des deux sortie de l'échantillonneur bloqueur de la figure 4.4.....	82
Figure 4.7	Circuit du quantificateur	82

Figure 4.8	Simulation temporelle du circuit du quantificateur.....	83
Figure 4.9	Circuit du DAC	84
Figure 4.10	Simulation temporelle du circuit du DAC	84
Figure 4.11	Simulation temporelle des deux types d'intégrateurs : sans délai et avec délai. Le même signal d'entrée est appliqué au deux intégrateurs	85
Figure 4.12	Circuit d'un additionneur	86
Figure 4.13	Simulation temporelle du circuit additionneur	86
Figure 4.14	Circuit complet du modulateur	88
Figure 4.15	Circuit du sous-bloc <i>Switch_Int</i> . Contient la configuration d'interrupteurs nécessaire pour réaliser un intégrateur avec délai	89
Figure 4.16	Circuit du sous-bloc <i>Switch_Int_ND</i> . Contient la configuration d'interrupteurs nécessaire pour réaliser un intégrateur sans délai	89
Figure 4.17	Spectre de fréquences de la configuration optimisée avec $\beta = 0,01$	90
Figure 4.18	Spectre de fréquences de la configuration optimisée avec $\beta = 0,05$	91
Figure 4.19	Spectre de fréquences de la configuration de référence	91
Figure 4.20	Courbes SNR vs amplitude d'entrée pour les trois configurations	92
Figure 4.21	Courbes SNR vs amplitude d'entrée pour la configuration optimisée avec $\beta = 0,01$ obtenus de trois façons différentes : calculs théoriques, simulations système avec Simulink et simulations transistors	94
Figure 4.22	Distribution de l'amplitude de sortie du premier intégrateur pour les trois configurations.....	96
Figure 4.23	Distribution de l'amplitude de sortie du quatrième intégrateur pour les trois configurations.....	96

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

$\Sigma\Delta$	Sigma-Delta
ADC	Convertisseur analogique à numérique («Analog to Digital Converter»)
AHDL	<i>Analog Hardware Descriptive Language</i>
CIFB	Cascade d'intégrateurs avec rétroaction («Cascade of Integrators with Feedback»)
CMOS	<i>Complementary Metal Oxide Semiconductor</i>
CRFB	Cascade de résonateurs avec rétroaction («Cascade of Resonators with Feedback»)
DAC	Convertisseur numérique à analogique («Digital to Analog Converter»)
DC	Courant continu («Direct Current»)
ENOB	Nombre effectif de bits («Effective Number Of Bits»)
FS	Pleine échelle («Full Scale»)
NPG	Gain de la puissance de bruit («Noise Power Gain»)
NTF	Fonction de transfert du bruit («Noise Transfer Function»)
OSR	Taux de suréchantillonnage («Oversampling Ratio»)
SNR	Rapport signal à bruit («Signal to Noise Ratio»)
$SNR_{\Sigma\Delta}$	Rapport signal à bruit d'un convertisseur Sigma-Delta
SNR_{max}	Rapport signal à bruit maximal
SNR_{sinus}	Rapport signal à bruit pour un signal sinusoïdal
STF	Fonction de transfert du signal («Signal Transfer Function»)

LISTE DES SYMBOLES ET UNITÉS DE MESURE

β	Paramètre de l'algorithme de minimisation multicritères
Δ	Pas de quantification
ε_q	Erreur de quantification
μm	micromètre
Ω	Ohm
a_i	Coefficients de rétroactions
b_i	Coefficients d'entrées
c_i	Coefficients interétages
C_1	Condensateur d'échantillonnage d'un intégrateur (pF)
C_2	Condensateur d'intégration d'un intégrateur (pF)
C_T	Somme des capacités du modulateur (pF)
dB	Décibel
dBFS	Décibel par rapport à la valeur pleine échelle
f_b	Largeur de bande du signal
f_s	Fréquence d'échantillonnage
g_i	Coefficients de rétroaction locale
$H_{ni}(z)$	Fonction de transfert des sources de bruit
MHz	Megahertz
N	Ordre du modulateur

pF	picoFarad
P_N	Puissance du bruit (W)
P_Q	Puissance du bruit de quantification (W)
$P_{Q\Sigma\Delta}$	Puissance du bruit de quantification d'un convertisseur Sigma-Delta (W)
P_{sinus}	Puissance d'un signal sinusoïdal (W)
$S_f(f)$	Densité spectrale du bruit de scintillation
$S_T(f)$	Densité spectrale du bruit thermique
X_N	Amplitude de sortie des intégrateurs

INTRODUCTION

Plusieurs applications comme le traitement audio et les communications sans fil nécessitent des convertisseurs ayant un rapport signal à bruit (SNR) élevé. Les convertisseurs Sigma-Delta ($\Sigma\Delta$) sont régulièrement utilisés pour ce type d'application car ils combinent la mise en forme du bruit de quantification et le suréchantillonnage du signal afin de diminuer de façon significative le bruit de quantification à l'intérieur de la bande de fréquences du signal [1].

Lors du processus de conception d'un convertisseur $\Sigma\Delta$, de nombreux paramètres doivent être déterminés. Un processus de conception traditionnel débute généralement par la modélisation mathématique du convertisseur. Ce modèle est basé sur la fonction de transfert du bruit (NTF) et la fonction de transfert du signal (STF). Pour définir ces deux fonctions, le taux de suréchantillonnage (OSR), l'ordre du modulateur ainsi que le nombre de bits de quantifications doivent être sélectionnés. Ces paramètres de base sont sélectionnés en fonction de l'objectif en SNR tel que fixé par l'application ciblée.

Après qu'une NTF et une STF satisfaisantes sont trouvées, celles-ci sont réalisées à l'intérieur d'une structure standard. Ces structures sont composées d'un ou plusieurs étages d'amplificateurs configurés en intégrateur. Cette opération nécessite de calculer adéquatement les différents coefficients de la structure afin de réaliser les fonctions de transfert souhaitées. Il n'y a pas de solution unique à ce calcul. Un des critères qui est généralement utilisé pour sélectionner un ensemble de paramètres est l'amplitude de sortie des intégrateurs utilisés dans les circuits du convertisseur. Il est en effet important de s'assurer que l'amplitude de sortie des intégrateurs soit en accord avec la plage d'opération linéaire des amplificateurs.

Finalement, le circuit du convertisseur est réalisé (au niveau transistor). Dans un circuit à temps discret, les différents coefficients de la structure sont réalisés par des ratios de condensateurs. La valeur des condensateurs détermine la puissance du bruit thermique généré par les circuits [2]. Étant donné que le bruit thermique est l'une des principales limitations des convertisseurs $\Sigma\Delta$, les différents condensateurs du circuit doivent être dimensionnés adéquatement.

Pour déterminer tous ces paramètres, il n'existe pas de règle absolue. Devant ce grand nombre de degrés de liberté, il n'est pas aisé de déterminer l'ensemble de paramètres optimaux pour une application donnée. Certains travaux, qui seront présentés au chapitre 2, établissent une méthode pour optimiser ces paramètres. Ce mémoire s'inscrit dans ce domaine de recherche en proposant une nouvelle approche.

Le premier chapitre introduit la théorie générale des convertisseurs $\Sigma\Delta$. Les notions nécessaires à leur conception et leur analyse sont présentées. Le chapitre se termine avec une présentation plus approfondie de la problématique étudiée.

Le second chapitre se divise en deux sections. Tout d'abord, le cadre théorique nécessaire à l'analyse des critères d'optimisation est présenté. Ces critères sont le bruit thermique des intégrateurs, l'amplitude de sortie des intégrateurs et la taille du circuit. Cette section expose également le fonctionnement d'un intégrateur à capacités commutées. Dans la seconde partie du chapitre, une revue de littérature concernant les techniques d'optimisation des convertisseurs $\Sigma\Delta$ est présentée.

Deux nouvelles méthodes pour optimiser les convertisseurs $\Sigma\Delta$ sont présentées au troisième chapitre. La première, appelée méthode des fenêtres, permet de réduire l'amplitude de sortie des intégrateurs en fonction d'une pénalité prédéterminée sur le SNR obtenu, par rapport au SNR théorique maximal pour l'architecture choisie. La seconde est un processus d'optimisation multicritères. En plus de minimiser l'amplitude de sortie des intégrateurs, celle-ci permet de minimiser simultanément la somme des capacités du modulateur. Des simulations au niveau système permettent de valider les deux méthodes.

Au quatrième chapitre, la méthode d'optimisation multicritères est confirmée à l'aide de simulations au niveau transistor. Tout d'abord, la réalisation des différentes parties du circuit à l'aide de la technologie CMOS 0,18 μm est exposée. Ensuite, les résultats des simulations sont analysés. Les performances de la méthode d'optimisation sont comparées aux performances d'un convertisseur issu d'un processus de conception traditionnel.

Finalement, une conclusion fait la synthèse de l'ensemble de cet ouvrage. Certaines suggestions pour des travaux futurs sont également incluses.

Contributions

Ce travail a permis d'apporter les contributions scientifiques suivantes :

- analyse de certaines problématiques liées à la conception de convertisseurs $\Sigma\Delta$;
- revue de littérature des méthodes d'optimisation des convertisseurs $\Sigma\Delta$;
- proposition de deux nouvelles méthodes d'optimisation pour les convertisseurs $\Sigma\Delta$;
- vérification par des simulations au niveau transistor de l'une des méthodes proposées ;
- publication de deux articles lors de conférences [2, 3] :
 - E. Collard-Fréchette and G. Gagnon, "A New Optimization Technique for Coefficient Scaling in Sigma-Delta Modulators," in *Proc. IEEE MWSCAS 2010*, Seattle, WA, Aug. 2010, pp. 312-315.
 - E. Collard-Fréchette, G. Kaddoum, and G. Gagnon, "Dynamic Range Scaling of Sigma-Delta Modulators Based on a Multi-Criteria Optimization Process," in *Proc. IEEE NEWCAS 2011*, Bordeaux, France, June 2011, pp. 193-196.

CHAPITRE 1

CONVERSION $\Sigma\Delta$

1.1 Introduction

Ce mémoire portant sur les convertisseurs analogiques-numériques (ADC) $\Sigma\Delta$, ce chapitre introduit les notions de base nécessaires à la compréhension de ce type de circuit. Lors de la numérisation d'un signal, de la précision est perdue irrémédiablement. Un convertisseur $\Sigma\Delta$ combine deux techniques, la mise en forme du bruit et le suréchantillonnage du signal, afin d'augmenter la résolution du convertisseur et ainsi réduire la perte d'information. De façon générale, la matière présentée dans ce chapitre est fortement adaptée de [1].

1.2 Numérisation d'un signal

Un ADC a pour fonction de discrétiser en temps et en amplitude un signal analogique continu. Ces deux opérations sont nommées échantillonnage et quantification. La figure 1.1 montre le schéma bloc général d'un ADC.

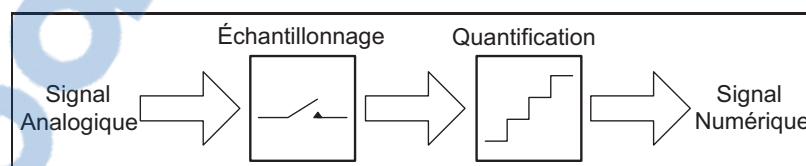


Figure 1.1 Schéma bloc général d'un ADC
Adaptée de [4]

La figure 1.2 montre la courbe de transfert entre l'entrée et la sortie d'un quantificateur typique. L'espacement entre deux paliers, nommé pas de quantification, est généralement constant. Dans cet exemple, il s'agit d'un quantificateur de 2 bits ayant une plage d'entrée de -4 à 4 et un pas de quantification de 2. La quantification d'un signal sur un nombre fini de niveaux ajoute un signal d'erreur au signal entrant [1]. La figure 1.3 trace le signal d'erreur en fonction du signal d'entrée pour le quantificateur de la figure 1.2.

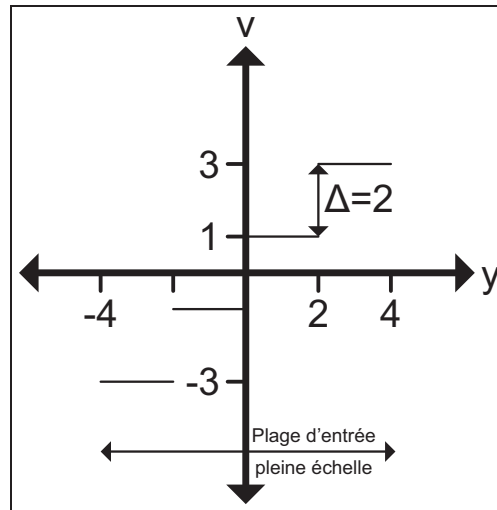


Figure 1.2 Courbe de transfert
d'un quantificateur typique
Adaptée de [1]

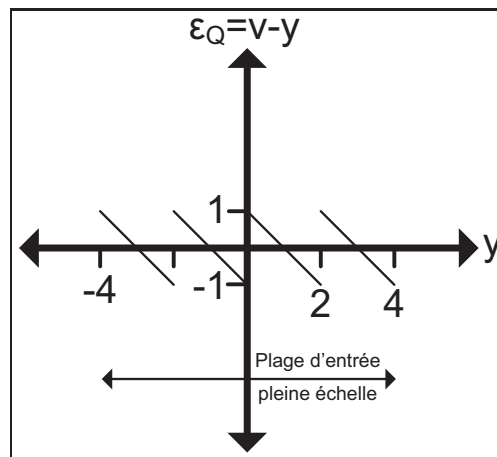


Figure 1.3 Signal d'erreur du
quantificateur de la figure 1.2
Adaptée de [1]

Tel que mentionné ci-haut, un quantificateur peut être modélisé mathématiquement par l'addition d'un signal d'erreur (figure 1.4). L'équation de sortie du quantificateur est alors donnée par :

$$V(n) = U(n) + \epsilon_Q(n) \quad (1.1)$$

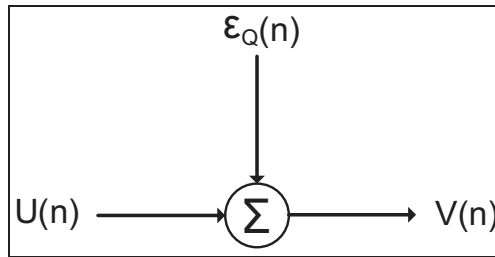


Figure 1.4 Modèle mathématique d'un quantificateur

Le signal d'erreur ε_Q est appelé bruit de quantification. Lorsque les conditions suivantes sont respectées [4] :

- tous les niveaux de quantification ont une probabilité égale d'être utilisés ;
- les pas de quantification sont uniformément distribués ;
- le signal d'entrée franchit plusieurs niveau de quantification entre deux échantillons ;
- le signal d'erreur n'est pas corrélé avec le signal d'entrée.

la puissance du signal d'erreur ε_Q peut alors être calculée par :

$$P_Q = \frac{1}{\Delta} \int_{-\Delta/2}^{\Delta/2} \varepsilon_Q^2 d\varepsilon_Q = \frac{\Delta^2}{12} \quad (1.2)$$

Ces conditions sont généralement respectées lorsqu'un signal d'entrée sinusoïdal est utilisé. Le signal d'erreur ε_Q est alors représenté par un bruit blanc uniformément distribué de 0 à $f_s/2$ ¹, où f_s représente la fréquence d'échantillonnage.

1.3 Rapport signal à bruit (SNR)

Le SNR est défini comme étant le ratio entre la puissance du signal d'entrée et la puissance du bruit dans la bande de fréquences du signal. La puissance d'un signal sinusoïdal dont l'ampli-

1. Tout le long de cet ouvrage, une représentation spectrale avec fréquences strictement positives sera utilisée.

tude occupe la pleine échelle d'un quantificateur sera [4] :

$$P_{\text{sinus}} = \frac{1}{T} \int_0^T \left(\frac{X_{f_s}}{2} \right)^2 \sin^2(2\pi ft) dt = \frac{X_{f_s}^2}{8} = \frac{(\Delta \cdot 2^n)^2}{8} \quad (1.3)$$

où X_{f_s} représente la valeur pleine échelle du quantificateur et n le nombre de bits du quantificateur. En combinant (1.2) et (1.3), on obtient l'équation spécifiant le SNR maximal d'un ADC pour un signal sinusoïdal (en dB) :

$$SNR_{\text{sinus}} = 6,02 \cdot n + 1,78 \quad (1.4)$$

Le SNR d'un convertisseur est parfois représenté par le nombre effectif de bits (ENOB). Le ENOB représente le nombre de bits nécessaire pour obtenir un certain SNR selon l'équation (1.4). Par exemple, un ENOB de 16 bits représente un SNR_{sinus} de 98,1 dB.

1.4 dBFS

L'amplitude du signal d'entrée d'un convertisseur est généralement exprimée en dB relatif à la valeur pleine échelle du quantificateur. Cette mesure est obtenue par :

$$dBFS = 20 \log \left(\frac{A_i}{X_{f_s}} \right) \quad (1.5)$$

où A_i représente l'amplitude d'entrée du convertisseur et X_{f_s} la valeur pleine échelle du quantificateur. Un signal d'entrée pleine échelle a une amplitude de 0 dBFS.

1.5 Plage dynamique

La plage dynamique est définie comme étant le ratio entre les amplitudes minimales et maximales produisant un SNR de 0 dB [4]. La plage dynamique est généralement représentée par un graphique du SNR en fonction de l'amplitude du signal d'entrée en dBFS. La figure 1.5 présente un exemple de ce type de graphique. La courbe en pointillés est celle d'un convertis-

seur 10 bits idéal. La relation entre le SNR et l'amplitude d'entrée est parfaitement linéaire. La plage dynamique est de 60 dB et le SNR_{max} (SNR maximal atteint par le convertisseur) est également de 60 dB. Sur la deuxième courbe, représentant un convertisseur $\Sigma\Delta$ réel, la plage dynamique est toujours 60 dB, mais le SNR_{max} est 53 dB. Il est atteint avant que l'amplitude du signal d'entrée soit maximale (à -5 dBFS). Cela est dû au fait que les convertisseurs $\Sigma\Delta$ ont tendance à saturer ou à devenir instable lorsque l'amplitude du signal d'entrée dépasse un certain seuil. Ce sujet sera approfondi à la section 1.7.3.

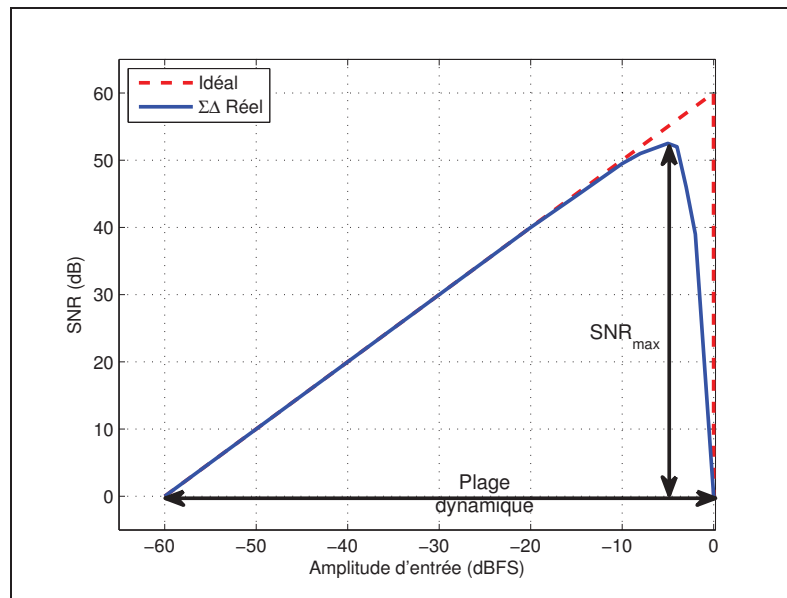


Figure 1.5 Graphique typique de la plage dynamique d'un convertisseur

1.6 Convertisseur $\Sigma\Delta$

Les tolérances des processus de fabrication diminuent la précision des composants analogiques des ADC. Cela a pour effet de limiter l'ENOB des convertisseurs. Bien que l'ENOB maximal obtenu dépend du processus de fabrication, l'ENOB maximal est typiquement limité à environ 12 bits [1]. Il est possible d'obtenir des valeurs supérieures en utilisant des techniques de calibration ou d'ajustage d'appoint (*trimming*) [4]. Cependant, ces techniques sont complexes et dispendieuses.

Plusieurs applications nécessitent un convertisseur ayant une linéarité supérieure à 12 bits [5–7]. Les convertisseurs $\Sigma\Delta$ permettent d’obtenir un ENOB allant jusqu’à 24 bits [8–11]. Pour atteindre ces performances, deux techniques sont utilisées : le suréchantillonnage du signal et la mise en forme du bruit de quantification.

1.6.1 Suréchantillonnage d’un signal

Telle que mentionnée précédemment, la puissance du bruit de quantification dépend uniquement de la valeur du pas de quantification. L’augmentation de la fréquence d’échantillonnage ne modifie pas la puissance totale du bruit de quantification. Cependant, le bruit sera distribué sur une plus grande plage de fréquences. Il est alors possible, par filtrage numérique, d’éliminer le bruit qui se retrouve à l’extérieur de la plage du signal.

La figure 1.6 illustre ce principe lorsque la fréquence d’échantillonnage est doublée. Le bruit de quantification étant distribué de 0 à $f_{s2}/2$ au lieu de 0 à $f_{s1}/2$, la densité spectrale du bruit est réduite de moitié. En appliquant un filtre passe-bas dont la fréquence de coupure est $f_{s1}/2$, la moitié du bruit de quantification est éliminé. Cela conduit à un gain de 3dB pour le SNR.

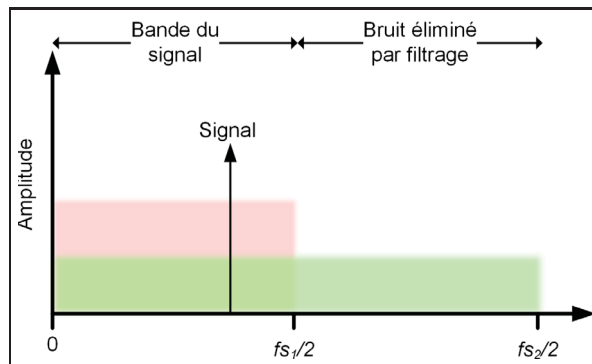


Figure 1.6 Impact du suréchantillonnage sur le bruit de quantification

1.6.2 OSR

L’OSR est défini comme étant le ratio entre la fréquence d’échantillonnage et deux fois la largeur de bande du signal (limite pour respecter le critère de Nyquist). Il peut être calculé par

l'équation suivante :

$$OSR = \frac{f_s}{2 \cdot f_b} \quad (1.6)$$

où f_s représente la fréquence d'échantillonnage et f_b la largeur de bande du signal. La définition du OSR permet de calculer le ENOB obtenu par un quantificateur de n bits [4] :

$$ENOB = n + \frac{\log_2(OSR)}{2} \quad (1.7)$$

1.6.3 Filtrage

À la section 1.6.1, il est mentionné qu'en doublant la fréquence d'échantillonnage, un gain en SNR de 3 dB est obtenu. Cette affirmation suppose qu'un filtre idéal est utilisé à la sortie du quantificateur. La totalité du bruit à l'extérieur de la bande du signal doit être éliminée tout en préservant l'intégralité de la bande du signal. Dans la pratique, ce type de filtre est impossible à réaliser. Cependant, le filtrage se faisant numériquement, des filtres d'ordre très élevés peuvent être utilisés. L'erreur qu'il introduit est donc généralement négligée dans le calcul du SNR. C'est-à-dire que le calcul du SNR se fait directement sur le spectre de sortie du quantificateur en considérant uniquement la bande du signal. Cela permet en outre de comparer les performances de deux convertisseurs $\Sigma\Delta$ différents sans l'ajout d'un biais dû à un filtrage de sortie différent.

1.6.4 Mise en forme du bruit

L'équation (1.7) indique que pour augmenter la résolution d'un quantificateur d'un ENOB, il faut multiplier le OSR par un facteur de quatre. L'efficacité du suréchantillonnage peut cependant être augmentée en effectuant une mise en forme du bruit de quantification à l'extérieur de la bande de fréquences du signal. Ceci est réalisé par une rétroaction qui permet d'intégrer l'erreur de quantification.

La figure 1.7 présente le schéma bloc d'un convertisseur $\Sigma\Delta$ du premier ordre tandis que son modèle mathématique est présenté à la figure 1.8. L'équation de sortie du convertisseur est

donnée par :

$$V(z) = z^{-1}U(z) + (1 - z^{-1}) \varepsilon(z) \quad (1.8)$$

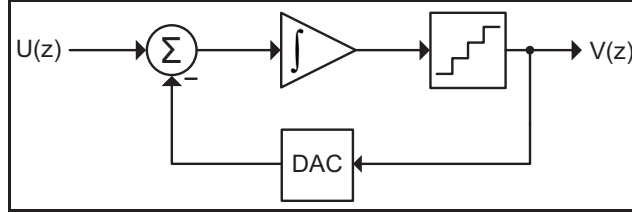


Figure 1.7 Schéma bloc d'un convertisseur $\Sigma\Delta$ du premier ordre
Adaptée de [1]

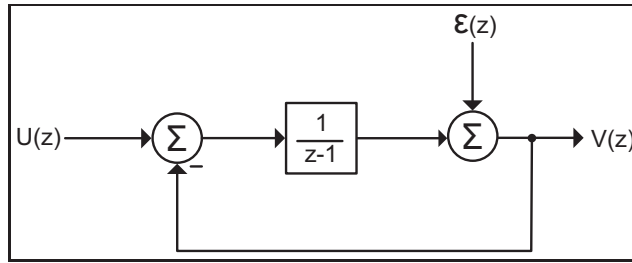


Figure 1.8 Modèle mathématique d'un convertisseur $\Sigma\Delta$ du premier ordre
Adaptée de [1]

En comparant l'équation (1.8) à celle d'un convertisseur de base (équation (1.1)) on remarque l'apparition de deux termes supplémentaires dans l'équation. Le terme qui multiplie $U(z)$ est appelé fonction de transfert du signal (STF) alors que le terme qui multiplie $\varepsilon(z)$ est appelé fonction de transfert du bruit (NTF). Dans le cas du convertisseur du premier ordre de la figure 1.7, la STF est z^{-1} , soit un délai unitaire qui ne change pas la forme du signal. Par contre, la NTF qui est $(1 - z^{-1})$ altère la réponse en fréquence du bruit. Pour connaître son impact sur la puissance du bruit de quantification à la sortie du convertisseur, l'amplitude de sa réponse en fréquence peut être calculée en posant $z = e^{j2\pi f}$:

$$\left| NTF \left(e^{j2\pi f} \right) \right|^2 = \left| 1 - e^{-j2\pi f} \right|^2 = 4 \sin^2 (\pi f) \quad (1.9)$$

Bien que la puissance totale du bruit de quantification soit augmentée par un facteur de 4, le bruit de quantification est mis en forme par la fonction $\sin^2$. Le tracé de $|NTF(e^{j2\pi f})|^2$ est présenté à la figure 1.9. On remarque que l'intégrateur a pour effet d'introduire un zéro à DC dans la NTF. Il en résulte une pente de 20 dB/décade, ce qui est typique d'un système du premier ordre.

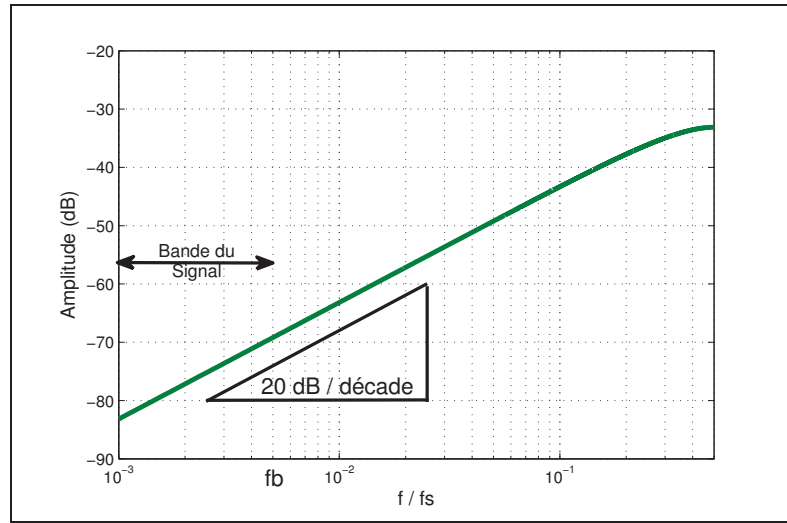


Figure 1.9 Densité spectrale de puissance du bruit de quantification d'un convertisseur $\Sigma\Delta$ du premier ordre

Le bruit de quantification étant repoussé vers les hautes fréquences, l'effet du suréchantillonnage est augmenté. La puissance du bruit dans la bande du signal peut être calculée par :

$$P_{Q\Sigma\Delta} = \int_0^{f_b} \frac{\Delta^2}{12} \cdot 4\sin^2(\pi f) df \approx \frac{\Delta^2}{12} \cdot \frac{\pi^2}{3 \cdot OSR^3} \quad (1.10)$$

Ce résultat est obtenu en utilisant l'approximation $\sin(x) \approx x$, approximation valide lorsque $x \ll \pi/2$.

En combinant les équations (1.3) et (1.10) on obtient l'équation du SNR (en dB) pour un convertisseur $\Sigma\Delta$ du premier ordre pour un signal sinusoïdal pleine échelle :

$$SNR_{\Sigma\Delta} = 6,02 \cdot n + 1,78 + 10\log\left(\frac{3 \cdot OSR^3}{\pi^2}\right) \quad (1.11)$$

1.7 Convertisseur d'ordre élevé

Selon l'équation (1.11), en doublant la fréquence d'échantillonnage le SNR est augmenté de 9 dB (ou 1,5 ENOB). Pour augmenter davantage l'efficacité des convertisseurs $\Sigma\Delta$, des étages d'intégration supplémentaires peuvent être ajoutés tel qu'illustré à la figure 1.10.

L'ajout des étages d'intégration résulte en une NTF d'ordre plus élevé ce qui se traduit par une plus grande atténuation du bruit de quantification dans la bande du signal. La figure 1.11 présente les NTF obtenues par des modulateurs d'ordre 1 à 4. L'équation (1.11) peut être généralisée pour un convertisseur d'ordre N [1] :

$$SNR_{\Sigma\Delta} = 6,02 \cdot n + 1,78 + 10 \log \left(\frac{(2N + 1) \cdot OSR^{2N+1}}{\pi^{2N}} \right) \quad (1.12)$$

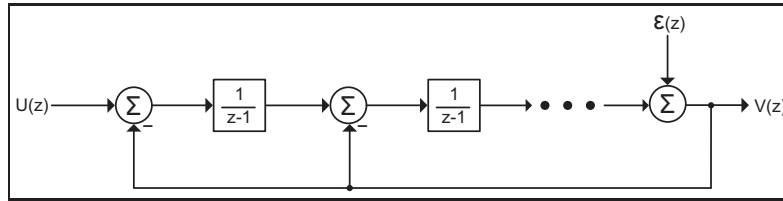


Figure 1.10 Modèle mathématique d'un convertisseur $\Sigma\Delta$ d'ordre élevé

1.7.1 Filtre de boucle

Afin de généraliser l'étude des modulateurs $\Sigma\Delta$ d'ordre élevé, il est avantageux de séparer la représentation mathématique du modulateur de sa structure de réalisation. Ceci est fait en représentant le modulateur par un filtre de boucle [1]. Cela permet d'obtenir la représentation générale des convertisseurs $\Sigma\Delta$ présentée à la figure 1.12. Les performances du modulateur peuvent être étudiées en se basant sur la représentation mathématique fournie par la STF et la NTF. La façon dont ces fonctions sont réalisées à l'intérieur du filtre de boucle n'entre pas en considération à ce stade-ci. L'équation à la sortie du convertisseur est donnée par :

$$V(z) = STF(z) \cdot U(z) + NTF(z) \cdot \varepsilon(z) \quad (1.13)$$

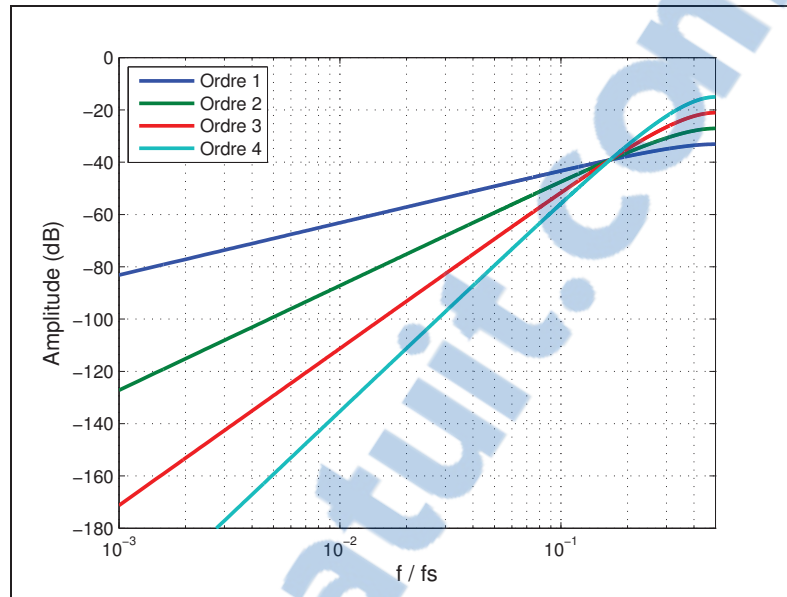


Figure 1.11 Densité spectrale de puissance du bruit de quantification des convertisseurs $\Sigma\Delta$ d'ordre 1 à 4

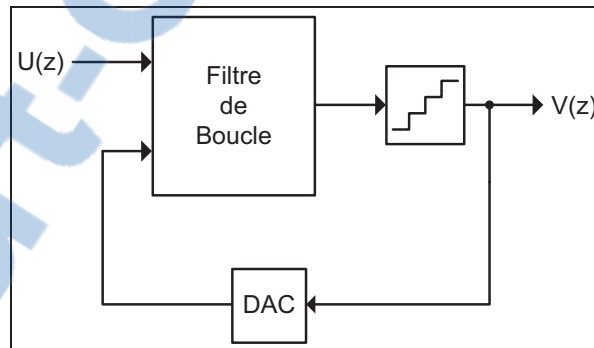


Figure 1.12 Représentation générale des convertisseurs $\Sigma\Delta$ par un filtre de boucle

1.7.2 Distribution des zéros

Jusqu'à présent, les NTF étudiées avaient tous leurs zéros placés à $z = 1$, c'est-à-dire à DC. Il est possible d'optimiser les performances de la NTF en répartissant les zéros à l'intérieur de la bande du signal. Ce principe est illustré à la figure 1.13 pour un modulateur du quatrième ordre optimisé pour un OSR de 32. On constate que les quatre zéros ont été distribués en deux paires de zéros conjugués aux fréquences normalisées 0,0053 et 0,0135. Il en résulte une diminution

du bruit à l'intérieur de la bande du signal procurant un gain de 13dB au niveau du SNR. Pour déterminer les fréquences optimales où placer les paires de zéros conjugués, il faut minimiser l'équation suivante :

$$P_{Q\Sigma\Delta} = \int_0^{f_b} \left| NTF(e^{j2\pi f}) \right|^2 df \quad (1.14)$$

où f_b représente la largeur de la bande de fréquences du signal. Cette équation signifie qu'il faut minimiser l'aire sous la courbe de la réponse en fréquence de la NTF en positionnant adéquatement les zéros. Ce calcul est généralement effectué à l'aide d'outils informatiques comme la librairie *Delsig Toolbox* de Matlab [12].

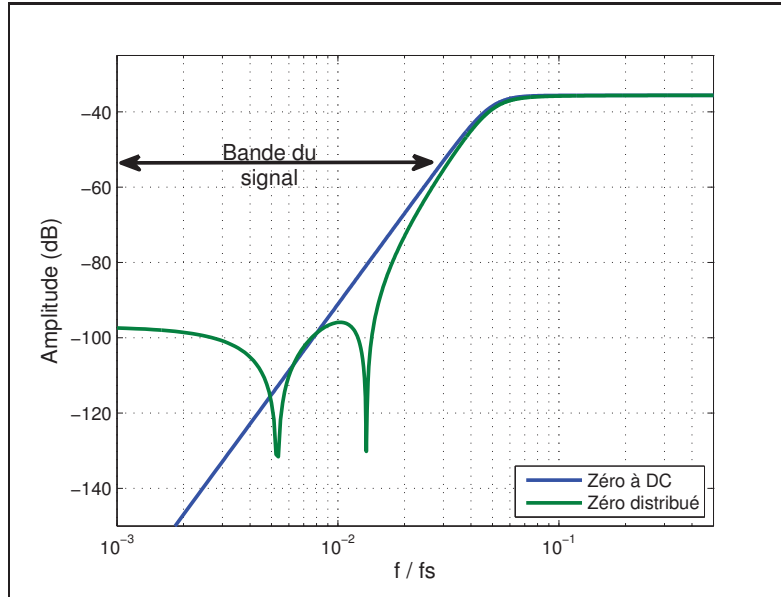


Figure 1.13 Effet de la distribution des zéros à l'intérieur de la bande du signal

1.7.3 Stabilité

La stabilité du modulateur est déterminée par la NTF. Un modulateur $\Sigma\Delta$ est dit stable pour un ensemble d'entrées U et un état initial donné si et seulement si pour toutes entrées $u \in U$ la sortie du modulateur est bornée [13]. Les modulateurs d'ordre 1 et 2 sont naturellement stables tant que l'amplitude du signal d'entrée est comprise à l'intérieur de la plage du quantificateur

[1]. Pour les modulateurs d'ordre supérieur² à 2, la stabilité devient un aspect primordial de la conception du modulateur.

Lors de l'étude d'un système linéaire, il est possible de définir des critères nécessaires et suffisants pour analyser la stabilité d'un système. Dans le cas des systèmes continus, il s'agit du critère de Routh-Hurwitz tandis que pour les systèmes discrets, il s'agit du critère de Jury. Malheureusement, ce type d'analyse ne peut être utilisé pour évaluer la stabilité d'un convertisseur $\Sigma\Delta$ en raison du quantificateur à l'intérieur de la boucle du système. Cet élément non linéaire empêche l'élaboration d'un critère nécessaire et suffisant universel pour évaluer la stabilité d'un convertisseur $\Sigma\Delta$. Il est donc nécessaire de procéder à de nombreuses simulations afin de s'assurer de la stabilité du système avant de le réaliser.

Certaines règles du pouce permettent d'aiguiller les concepteurs lors de la définition de ces paramètres. L'une des plus connues est le critère de Lee [1]. Ce critère qui n'est ni nécessaire, ni suffisant, s'applique uniquement aux convertisseurs ayant un quantificateur binaire (voir section 1.7.4). Il stipule qu'un modulateur $\Sigma\Delta$ binaire est susceptible d'être stable si $\max|NTF(e^{j\omega})| < 1,5$.

Le placement des pôles de la NTF est une opération délicate. Elle doit concilier stabilité et performance. En rapprochant les pôles des zéros, l'amplitude maximale de la réponse en fréquence de la NTF est diminuée ce qui augmente la stabilité du système. Cependant, la mise en forme du bruit devient moins agressive ce qui réduit les performances du système. Pour comparer la stabilité de deux modulateurs, nous pouvons utiliser la plage d'opération stable du convertisseur. Celle-ci représente la plage des amplitudes d'entrées pour laquelle le convertisseur fonctionne correctement.

Les figures 1.14 et 1.15 illustrent la dualité entre stabilité et performance pour le cas d'un convertisseur du quatrième ordre ayant un OSR de 32 et un quantificateur de 2 bits. Tout d'abord, la figure 1.14 présente l'amplitude de la réponse en fréquence de deux NTF ayant les mêmes zéros, mais dont les pôles ont été positionnés de façon à obtenir deux amplitudes

2. Dans certains cas particuliers, les modulateurs d'ordre 2 peuvent également devenir instable [1].

maximales différentes de la réponse en fréquence de la NTF. On constate qu'en éloignant les pôles des zéros, le plancher de bruit est diminué dans la bande du signal. La figure 1.15 montre le SNR obtenue en fonction de l'amplitude du signal d'entrée pour les deux mêmes NTF. On constate qu'en éloignant les pôles des zéros, le SNR est augmenté, mais que le système devient instable pour une amplitude d'entrée plus faible.

En conclusion, la seule façon de s'assurer de la stabilité d'un modulateur $\Sigma\Delta$ est de procéder à de nombreuses simulations représentant les différents points d'opérations. Négliger cette étape peut conduire à de graves problèmes lors de la réalisation.

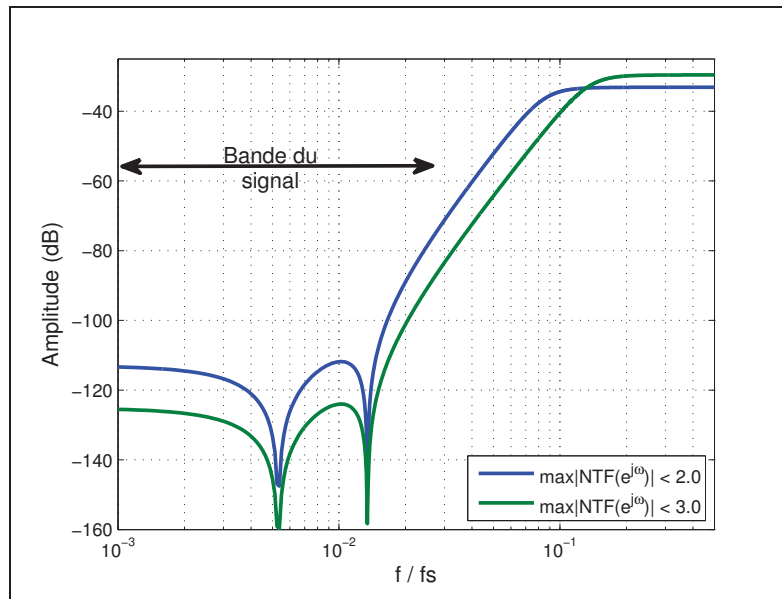


Figure 1.14 Effet de $\max|NTF(e^{j\omega})|$ sur le spectre de la NTF

1.7.4 Quantification

Lors de la conception d'un convertisseur $\Sigma\Delta$, l'un des premiers paramètres à sélectionner est le nombre de bits de quantification. Cette section présente certaines considérations relatives à ce choix.

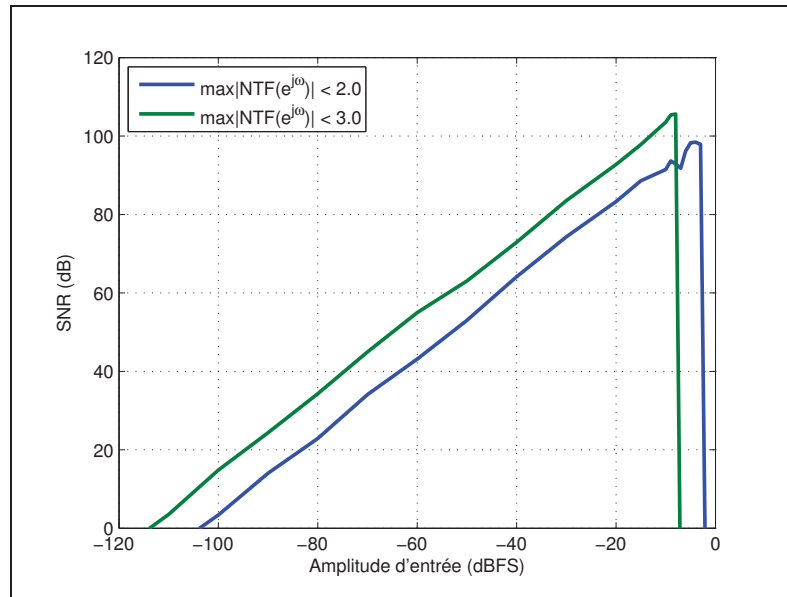


Figure 1.15 Effet de $\max|NTF(e^{j\omega})|$ sur la plage dynamique du convertisseur

1.7.4.1 Quantification binaire

Un convertisseur $\Sigma\Delta$ peut utiliser un quantificateur ayant uniquement deux niveaux de sortie. Il s'agit alors d'une quantification binaire car le convertisseur génère un seul bit à sa sortie. C'est lors du filtrage à décimation que le nombre de bits est augmenté afin de refléter la résolution du convertisseur. La figure 1.16 présente la sortie d'un quantificateur binaire d'un convertisseur du deuxième ordre lorsqu'un signal sinusoïdal est placé à l'entrée du convertisseur. Lorsque l'amplitude du signal d'entrée est maximale, la sortie est majoritairement +1 tandis que lorsque l'amplitude du signal d'entrée est minimale, la sortie est majoritairement -1.

L'utilisation d'un quantificateur binaire permet de simplifier la réalisation du convertisseur. Tel qu'illustré à la figure 1.17, le quantificateur devient un simple comparateur tandis que le convertisseur numérique-analogique (DAC), est réalisé par un simple circuit générant deux niveaux de tension : $-V_{ref}$ et $+V_{ref}$.

Un second avantage des quantificateurs binaires est leur linéarité inhérente. En effet, en n'ayant que deux niveaux, ceux-ci sont nécessairement reliés par une ligne droite. L'importance de cet avantage sera expliquée plus en profondeur à la prochaine section.

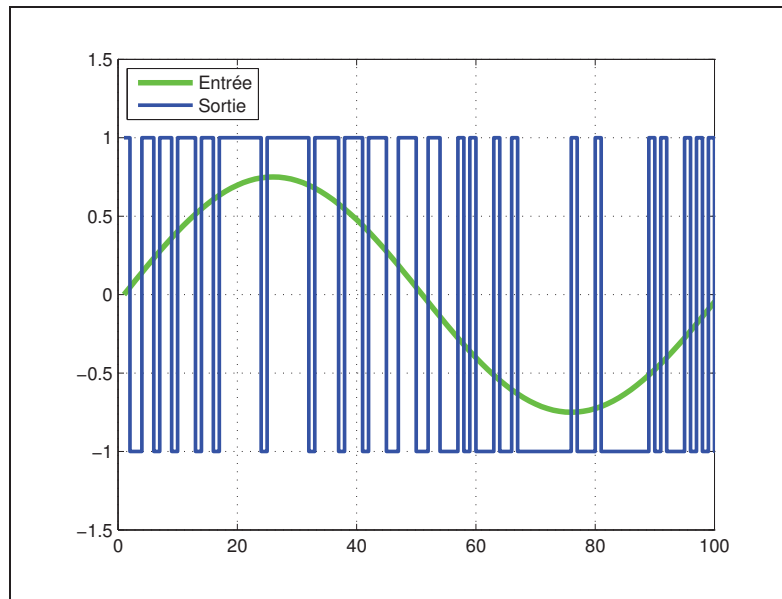


Figure 1.16 Entrée et sortie d'un convertisseur $\Sigma\Delta$ du deuxième ordre ayant un quantificateur binaire

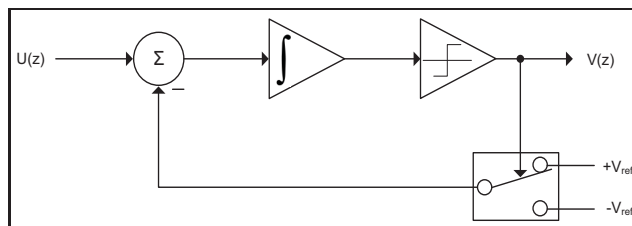


Figure 1.17 Modèle d'un convertisseur $\Sigma\Delta$ du premier ordre ayant un quantificateur binaire

1.7.4.2 Quantification multibit

L'utilisation d'un quantificateur multibit présente de nombreux avantages [1]. Selon l'équation (1.12), chaque bit de quantification supplémentaire augmente le SNR de 6 dB. De plus, le comportement de la boucle de rétroaction est plus linéaire ce qui augmente la stabilité du

modulateur et assouplit les spécifications requises par les composantes du système. Étant naturellement plus stable, la NTF peut être plus agressive afin d'augmenter les performances du convertisseur.

Ces nombreux avantages sont cependant contrebalancés par un désavantage majeur : toute erreur introduite par le DAC de la branche de retour est directement ajoutée au signal d'entrée. Ces erreurs se retrouvent donc dans le signal de sortie sans être mises en forme par le filtre de boucle. Cela conduit à la règle bien connue [1, 4] que pour avoir un convertisseur $\Sigma\Delta$ ayant une linéarité de m -bit, il faut que le DAC de la branche de retour ait une linéarité minimale de m bits.

Dans le cas des quantificateurs binaires, ce problème ne se pose pas, car, comme il a été dit précédemment, ceux-ci ont une linéarité inhérente. Par contre, dès que le nombre de niveaux de quantification est supérieur à deux, la linéarité du DAC est limitée par la précision des processus de fabrication qui, tel que mentionné précédemment, est environ 12 bits. Pour surmonter ces problèmes, différentes techniques ont été développées. Celles-ci incluent des systèmes de calibration [14, 15] et des techniques de moyennage dynamique des éléments [16, 17]. Il est donc possible de réaliser des convertisseurs $\Sigma\Delta$ multibits ayant une linéarité supérieure à la limite imposée par la précision des processus de fabrication [18, 19].

1.8 Structure de réalisation

Dans la section 1.7, la représentation mathématique d'un modulateur $\Sigma\Delta$ a été étudiée. Une fois que la NTF et la STF désirées sont déterminées, celles-ci doivent être réalisées à l'intérieur d'une structure. Pour ce faire, plusieurs architectures sont disponibles. Cette section présente les structures de types cascade d'intégrateurs avec rétroaction (CIFB) et cascade de résonateurs avec rétroaction (CRFB). Ce sont les deux types de structures qui seront utilisées dans cet ouvrage.

1.8.1 CIFB

La première structure présentée, la CIFB, est illustrée à la figure 1.18 pour un modulateur du quatrième ordre. Il s'agit d'une cascade d'intégrateurs avec entrée et rétroaction distribuées. C'est une extension de la structure présentée à la figure 1.10 avec l'ajout des coefficients de rétroactions a_i , d'entrées b_i et d'interétages c_i . En plaçant les coefficients c_i à 1, les équations de la NTF et de la STF sont données par :

$$NTF(z) = \frac{(z-1)^N}{D(z)} \quad (1.15)$$

$$STF(z) = \frac{b_1 + b_2(z-1) + \dots + b_{N+1}(z-1)^N}{D(z)} \quad (1.16)$$

où

$$D(z) = a_1 + a_2(z-1) + \dots + a_N(z-1)^{N-1} + (z-1)^N \quad (1.17)$$

En analysant les équations (1.15) et (1.16), certaines constatations peuvent être faites sur la structure CIFB. Tout d'abord, l'équation (1.15) implique que tous les zéros de la NTF sont situés à DC. La solution pour surmonter cette contrainte sera présentée à la section 1.8.2.

En second lieu, on peut constater que la NTF et la STF se partagent les mêmes pôles. Leurs emplacements sont déterminés par la valeur des coefficients a_i . Il est donc possible de les positionner aux fréquences désirées.

Finalement, les zéros de la STF sont déterminés par les coefficients b_i . Différentes avenues peuvent être envisagées sur la façon d'utiliser la STF. Premièrement, elle peut servir comme un préfiltre du signal. Les coefficients b_i sont alors calculés pour réaliser la fonction de filtrage souhaitée. Une seconde avenue consiste à utiliser $b_i = a_i$ pour $i \leq N$ et $b_{N+1} = 1$. De cette façon, les zéros et les pôles de la STF ont les mêmes valeurs. Il en résulte que la $STF = 1$ [20]. De plus, les intégrateurs se retrouvent à traiter uniquement le bruit de quantification. Le signal est acheminé directement à l'entrée du quantificateur. En ayant uniquement le bruit

de quantification à traiter, l'amplitude à la sortie des intégrateurs est grandement réduite. Ce principe sera approfondi à la section 2.2.3.1.

Une troisième option, intéressante pour des raisons de simplicité, consiste à mettre tous les coefficients b_i excepté b_1 à 0. La STF est alors $b_1/D(z)$. Pour garantir que l'amplitude de la STF est constante dans la bande de fréquences du signal, il faut que $|D(e^{j\omega})|$ soit constant dans la bande du signal, ce qui est généralement le cas..

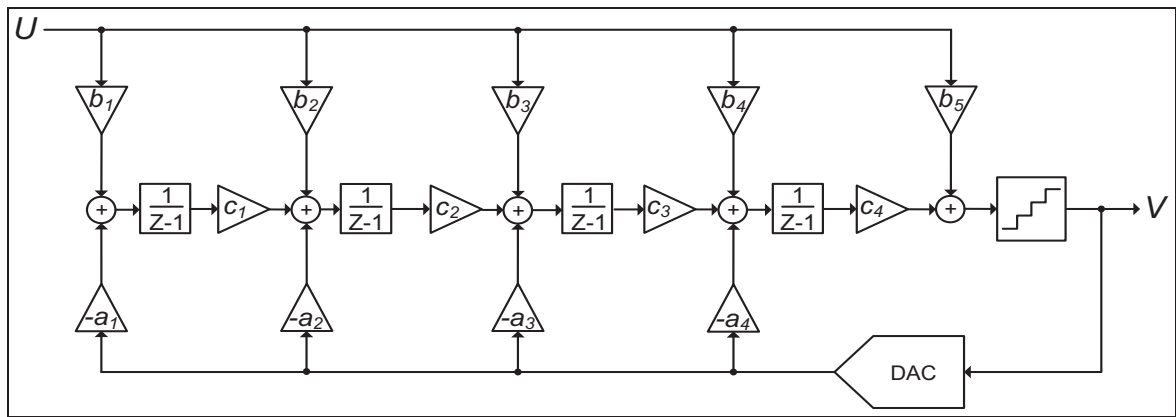


Figure 1.18 Structure CIFB

1.8.2 CRFB

La structure CIFB présentée précédemment force les zéros de la NTF à être à DC. Or, tel que vu à la section 1.7.2, il peut être avantageux de répartir les zéros de la NTF à l'intérieur de la bande du signal. Pour réaliser cette opération, la structure CIFB doit être modifiée afin d'inclure des rétroactions locales appelées résonateur. La structure obtenue, nommée CRFB, est présentée à la figure 1.19.

Chaque résonateur formé par un coefficient de rétroaction g_i permet la réalisation d'une paire de zéros conjugués. Ils seront situés à la fréquence $|\omega_i| = \cos^{-1}(1 - g_i/2)$. Il est à noter que l'un des deux intégrateurs du résonateur doit être sans délai. Si l'ordre du modulateur est impair, un intégrateur supplémentaire est inséré à l'entrée du filtre de boucle. Le zéro supplémentaire qu'il introduit est obligatoirement à DC.

Les résonateurs créés par les coefficients g_i sont intrinsèquement instables. Cependant, étant pris à l'intérieur d'une boucle qui elle est normalement stable, les oscillations locales sont évitées [1].

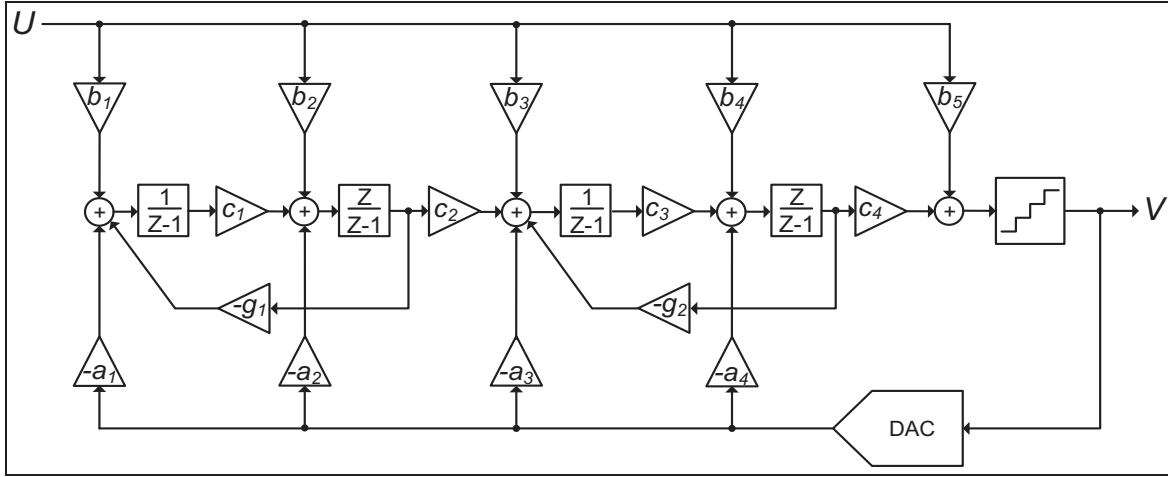


Figure 1.19 Structure CRFB

1.9 Problématique étudiée

Un point qui ressort de la présentation faite jusqu'à présent est qu'un grand nombre de paramètres doivent être déterminés lors de la conception d'un convertisseur $\Sigma\Delta$. Cette problématique a été brièvement énoncée dans l'introduction. Afin d'en cerner adéquatement les différents aspects, cette section présentera tout d'abord les différents niveaux d'abstractions où les paramètres peuvent être optimisés. Ensuite, elle s'attardera sur les difficultés liées au calcul des coefficients du modulateur, difficultés qui sont au coeur de la problématique étudiée.

1.9.1 Niveaux d'abstractions

Le processus de conception d'un convertisseur $\Sigma\Delta$ passe par plusieurs niveaux d'abstractions. À chacun de ces niveaux, plusieurs paramètres doivent être déterminés. Ces choix ont des répercussions sur les niveaux d'abstractions subséquents. Il importe donc de conserver une vue d'ensemble du problème.

Le premier niveau d'abstraction regroupe les paramètres systèmes permettant la synthèse de la NTF et de la STF. C'est à cette étape que sont définis l'ordre du modulateur, l'OSR et le nombre de bits de quantification. Ces choix sont généralement faits en fonction d'un objectif de SNR. Cependant, différentes combinaisons permettent d'atteindre le même objectif. Il faut alors évaluer la taille, la vitesse et la complexité du circuit afin de choisir une combinaison adéquate. Par exemple, l'utilisation d'un modulateur d'ordre élevé permet de diminuer l'OSR du convertisseur. Ce choix a l'avantage de permettre la diminution de la fréquence d'échantillonnage. Par contre, lors de sa réalisation, la taille du circuit sera probablement augmentée. De plus, en augmentant l'ordre du modulateur, il faut surveiller de près la stabilité du modulateur.

Au second niveau, les fonctions de transfert sont réalisées à l'intérieur d'une structure de réalisation. Les différents coefficients du filtre de boucle doivent être calculés adéquatement. Encore une fois, il n'y a pas de solution unique à ce problème. De plus, le choix des coefficients doit tenir compte du circuit qui les réalisera. Par exemple, dans un circuit à capacités commutées, les coefficients sont réalisés par des ratios de condensateurs. Les valeurs des coefficients doivent donc mener à des ratios réalisables compte tenu des tolérances de fabrication.

Finalement, le dernier niveau d'abstraction consiste à réaliser le circuit au niveau transistor. Un convertisseur $\Sigma\Delta$ contient de nombreuses composantes qui, à leur tour, comportent plusieurs paramètres. Le nombre de variables pouvant être optimisées à ce niveau est donc très élevé. Pour éviter de se retrouver avec un problème aux dimensions difficilement gérables, il faut sélectionner judicieusement les paramètres à optimiser. Cet ouvrage se concentre sur la valeur des condensateurs, la taille du circuit ainsi que sur l'amplitude de sortie des intégrateurs. Ces trois aspects seront approfondis au chapitre 2.

1.9.2 Calcul des coefficients

Dans cet ouvrage, un des principaux aspect étudié est l'optimisation des coefficients du modulateur. À la section 1.8, deux structures de réalisation ont été présentées (CIFB et CRFB). Pour chacune de ces structures, les équations permettant le calcul des coefficients a_i , b_i et g_i sont connues [1, 12]. Cependant, dans ces équations, les coefficients interétages c_i sont à 1.

Or, pour chaque ensemble de coefficients interétages c_i , il est possible de calculer un nouvel ensemble de coefficients a_i , b_i et g_i réalisant les mêmes NTF et STF. Tel qu'il le sera démontré à la section 2.2.3.1, les coefficients interétages c_i permettent d'ajuster l'amplitude de sortie des intégrateurs pour une NTF et une STF données.

Alors une question se pose : comment calculer adéquatement les coefficients du modulateur ? Il n'existe pas de solution universelle à cette question. Ce degré de liberté supplémentaire peut être utilisé de différentes façons. Dans cet ouvrage, les relations entre les coefficients interétages c_i , le bruit thermique du circuit, la taille du circuit et l'amplitude de sortie des intégrateurs seront utilisés afin de déterminer un ensemble de paramètres optimaux pour une application donnée. Les notions nécessaires à l'analyse de ces relations seront présentées en détail au chapitre suivant.

1.10 Conclusion

Dans ce chapitre, nous avons présenté la théorie de base des convertisseurs $\Sigma\Delta$. Ce type de convertisseur utilise conjointement les techniques de suréchantillonnage du signal et la mise en forme du bruit de quantification afin d'augmenter le SNR issu de la quantification. Le bruit de quantification est repoussé vers les hautes fréquences pour ensuite être éliminé par filtrage.

La conception du modulateur $\Sigma\Delta$ peut se faire en faisant abstraction du circuit qui le réalise. L'accent est alors mis sur son modèle mathématique. La NTF et la STF idéales pour l'application ciblée sont alors déterminées. Par la suite, ces deux fonctions de transfert doivent être réalisées dans une structure de réalisation. Cependant, il n'y a pas de solution unique à ce problème. Les coefficients interétages c_i offre un degré de liberté qui est généralement utilisé pour ajuster l'amplitude de sortie des intégrateurs.

CHAPITRE 2

OPTIMISATION DES PARAMÈTRES DES MODULATEURS $\Sigma\Delta$

2.1 Introduction

Tel qu'expliqué au chapitre précédent, la réalisation d'un convertisseur $\Sigma\Delta$ nécessite de calculer adéquatement les coefficients du modulateur. Cet ouvrage vise à développer une méthode pour optimiser les coefficients d'un modulateur $\Sigma\Delta$ en fonction du bruit thermique, de la taille du circuit et la consommation de puissance. Cette dernière sera minimisée en réduisant l'amplitude de sortie des intégrateurs. La première partie de ce chapitre présente le cadre théorique nécessaire à l'analyse de ces trois éléments. La seconde partie du chapitre présente une synthèse des principaux travaux publiés dans le domaine de l'optimisation des convertisseurs $\Sigma\Delta$. Cette présentation permettra de situer cet ouvrage à l'intérieur de la littérature existante.

2.2 Paramètres d'optimisation

La première étape de l'élaboration d'un processus d'optimisation est de bien cerner les paramètres. Dans cet ouvrage, il s'agit du bruit thermique, de la taille du circuit ainsi que de l'amplitude de sortie des intégrateurs. La première partie de ce chapitre analyse ces trois éléments afin de bien comprendre leur importance. Le fonctionnement d'un intégrateur à capacités commutées est également exposé.

2.2.1 Intégrateur à capacités commutées

Un modulateur $\Sigma\Delta$ peut être réalisé à l'aide de circuits à temps continu ou à temps discret. Les circuits à temps continu ont l'avantage de posséder un filtrage inhérent du repliement spectral et peuvent généralement fonctionner à des fréquences plus élevées [1, 4]. Cependant, les circuits à temps continu ont généralement une linéarité et une précision inférieure aux circuits à temps discret. De plus, les constantes de temps du circuit ont habituellement besoin d'un système de calibration. Pour ces raisons, cet ouvrage se concentre exclusivement sur les circuits à temps

discret. Par contre, les méthodes qui seront présentées au chapitre 3 pourraient être adaptées aux circuits à temps continu.

Tel qu'expliqué au chapitre 1, un modulateur $\Sigma\Delta$ se compose de plusieurs étages d'intégration. La figure 2.1 présente le circuit d'un intégrateur à capacités commutées. Mathématiquement, l'objectif de ce circuit est de réaliser l'équation suivante :

$$V_{out_{k+1}} = V_{out_k} + V_{in_k} \quad (2.1)$$

Cette équation signifie qu'à chaque période de temps k , la sortie de l'intégrateur doit additionner la valeur de l'entrée à la valeur présente de la sortie. Cette opération s'effectue en deux phases : l'échantillonnage et l'intégration.

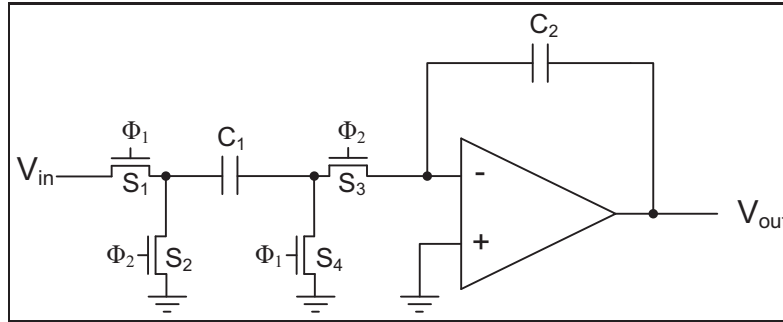


Figure 2.1 Circuit d'un intégrateur à capacités commutées
Adaptée de [1]

Les transistors S_1 à S_4 de la figure 2.1 sont des interrupteurs contrôlés par deux signaux d'horloges sans chevauchement (figure 2.2). La phase Φ_1 est la phase d'échantillonnage tandis que la phase Φ_2 est la phase d'intégration.

Durant la phase Φ_1 , les interrupteurs S_1 et S_4 sont activés tandis que S_2 et S_3 sont bloqués. Le circuit de la figure 2.3 est alors obtenu. Dans cette configuration, le condensateur C_1 se retrouve connecté entre l'entrée et la masse. Il se charge donc à la tension de l'entrée V_{in} .

Durant la phase Φ_2 , ce sont les interrupteurs S_2 et S_3 qui sont activés tandis que S_1 et S_4 sont bloqués. Le circuit de la figure 2.4 est alors obtenu. Cette fois, le condensateur C_1 se

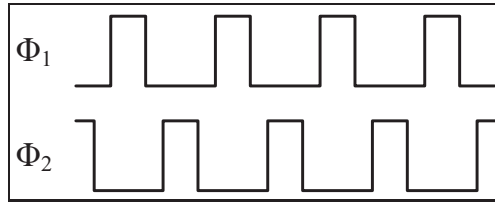


Figure 2.2 Signaux d'horloges sans chevauchement

retrouve connecté entre la masse et la masse virtuelle de l'amplificateur opérationnel. La charge contenue dans C_1 est alors transférée vers C_2 . Ce transfert de charge augmente (ou diminue selon le signe de la tension d'entrée) la tension de sortie d'une valeur proportionnelle à la tension d'entrée. Le circuit effectue donc l'opération suivante [1] :

$$V_{out_{k+1}} = V_{out_k} + \frac{C_1}{C_2} V_{in_k} \quad (2.2)$$

Le ratio entre les valeurs de C_1 et C_2 détermine le coefficient d'intégration.

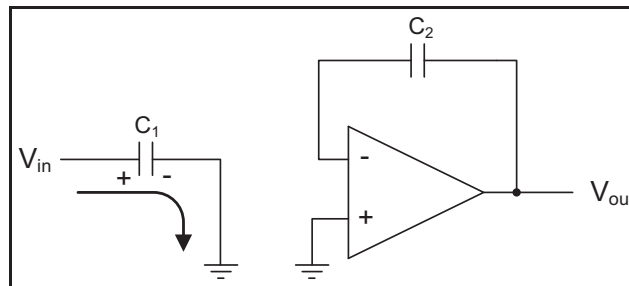


Figure 2.3 Circuit équivalent d'un intégrateur à capacités commutées durant la phase d'échantillonnage

2.2.2 Sources de bruit d'un intégrateur

Tel qu'expliqué au chapitre 1, le bruit de quantification est une limite fondamentale au SNR qu'un ADC peut atteindre. Cependant, ce n'est pas la seule limite imposée par un circuit réel. Les intégrateurs utilisés dans le modulateur génèrent également du bruit qui s'ajoute au signal et limite le SNR obtenu.

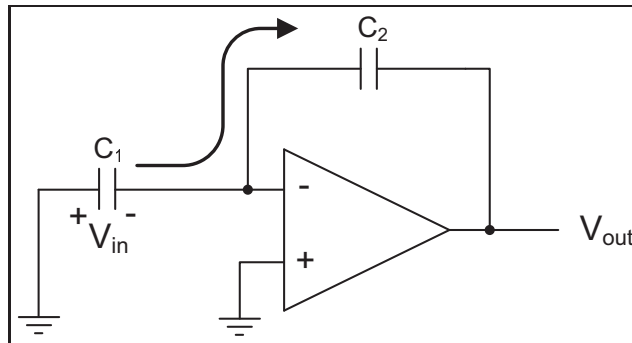


Figure 2.4 Circuit équivalent d'un intégrateur à capacités commutées durant la phase d'intégration

Dans un intégrateur à capacités commutées tel que présenté à la section 2.2.1, les principales sources de bruit sont [21, 22] :

- bruit de scintillation des interrupteurs (*flicker noise*) ;
- bruit thermique des interrupteurs ;
- bruit des amplificateurs opérationnels.

2.2.2.1 Bruit de scintillation

Le bruit de scintillation des interrupteurs, aussi appelé bruit $1/f$, est dû aux porteurs de charges. En se faisant piéger, puis libérer au fur et à mesure qu'ils progressent à l'intérieur du canal du transistor, ceux-ci introduisent un bruit dans le courant du drain [23]. Étant dépendante du processus de fabrication, la puissance moyenne du bruit de scintillation est difficile à évaluer. Cette source de bruit peut être modélisée par une source de tension connectée à la grille du transistor dont la densité spectrale de puissance est donnée approximativement par [21] :

$$S_f(f) = \frac{K}{WLf} \quad (2.3)$$

où W et L sont les dimensions du transistor, f la fréquence et K une constante propre au processus de fabrication. Le spectre de fréquences de cette source de bruit n'est pas blanc. Sa puissance est majoritairement concentrée aux basses fréquences.

Plusieurs techniques existent afin de rendre négligeable l'effet du bruit de scintillation. Parmi celles-ci, notons l'augmentation de la largeur des transistors, le double échantillonnage corrélé ou la stabilisation d'interrupteur périodique [21–23]. Étant donné l'existence de solutions matérielles au bruit de scintillation, celui-ci ne sera pas considéré dans cet ouvrage.

2.2.2.2 Bruit thermique

Le bruit thermique est causé par les mouvements aléatoires des électrons à l'intérieur d'une résistance. Ceci introduit une fluctuation de tension pouvant être modélisée par une source de tension en série avec la résistance. Cette source génère du bruit même si le courant est nul. Pour une résistance, la densité spectrale de cette source de bruit est donnée par [23] :

$$S_t(f) = 4kTR \quad (2.4)$$

où k est la constante de Boltzmann ($1,38 \cdot 10^{-23}$), T est la température en degré Kelvin et R la valeur de la résistance.

Un transistor MOSFET utilisé comme interrupteur peut être modélisé par une résistance. Lorsque le transistor est activé, la résistance entre le drain et la source est appelée R_{on} . Le transistor génère alors du bruit thermique dont la densité spectrale est $4kTR_{on}$.

Dans le circuit de l'intégrateur à capacités commutées de la figure 2.1, les transistors S_1 à S_4 génèrent du bruit thermique qui vient s'ajouter au signal. L'effet de ces sources de bruit peut être représenté par une source de tension placée à l'entrée d'un intégrateur idéal [21] tel qu'illustré à la figure 2.5.

Pour connaître la puissance de la source de bruit équivalente à l'entrée de l'intégrateur, il faut considérer l'impact du condensateur d'échantillonnage C_1 . Lorsque celui-ci échantillonne

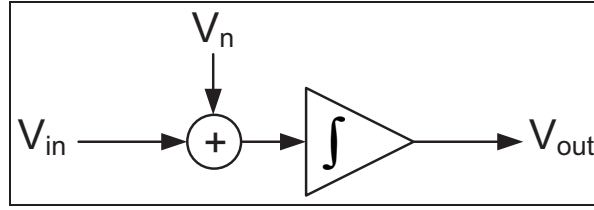


Figure 2.5 Modèle équivalent du bruit thermique généré par les interrupteurs
Adaptée de [21]

l'entrée de l'intégrateur, il échantillonne non seulement la tension d'entrée, mais également le bruit généré par les interrupteurs. Par la suite, ce bruit est intégré à la sortie en même temps que la tension d'entrée. Par définition, la densité spectrale du bruit thermique est blanche. Cependant, la résistance R_{on} des transistors et le condensateur d'échantillonnage C_1 forment un filtre passe-bas. Ce filtre vient limiter la bande passante du bruit thermique. La puissance totale du bruit peut alors être calculée par [22] :

$$P_N = \int_0^{\infty} \frac{4kTR_{on}}{1 + (2\pi f R_{on} C_1)^2} df = \frac{kT}{C_1} \quad (2.5)$$

L'équation 2.5 démontre que la puissance totale du bruit thermique est indépendante de la valeur de la résistance R_{on} . Cela est dû au fait qu'en augmentant la valeur de R_{on} , la densité spectrale du bruit est augmentée, mais sa bande passante est réduite. Ces deux relations qui s'annulent mutuellement font en sorte que la valeur de R_{on} n'a pas d'impact sur la puissance du bruit thermique. Bien que la largeur de bande du bruit ($1/2\pi R_{on} C_1$) est typiquement beaucoup plus grande que la fréquence d'échantillonnage, le repliement spectral lors de l'échantillonnage fait en sorte que le bruit est pratiquement blanc dans la bande 0 à $f_s/2$.

Dans un modulateur $\Sigma\Delta$, des intégrateurs ayant plusieurs entrées sont nécessaires afin de faire la somme des différents signaux. La figure 2.6 montre le circuit d'un intégrateur à capacités commutées ayant trois entrées. Les différents coefficients du modulateur $\Sigma\Delta$ sont réalisés par les ratios des condensateurs C_{1i}/C_2 .

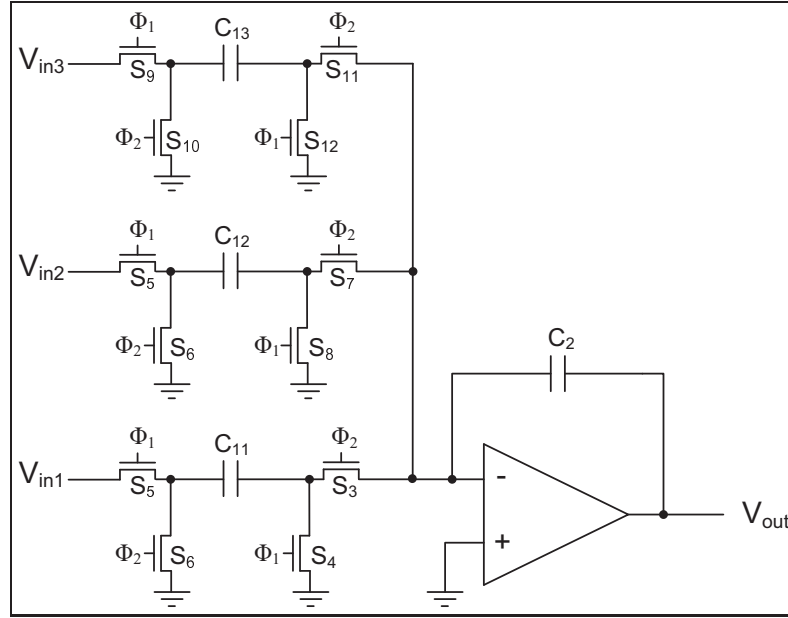


Figure 2.6 Circuit d'un intégrateur à capacités commutées à trois entrées

Les interrupteurs des différentes entrées forment des sources de bruit non corrélées. Leurs puissances peuvent donc être additionnées. La puissance de la source de bruit équivalente à l'entrée de l'intégrateur est alors donnée par :

$$P_N = kT \left(\frac{1}{C_{11}} + \frac{1}{C_{12}} + \frac{1}{C_{13}} \right) \quad (2.6)$$

Pour diminuer la puissance du bruit thermique, il faut augmenter la taille des condensateurs d'échantillonnage. Celle-ci est déterminée par deux facteurs : la taille du condensateur d'intégration C_2 et la valeur des coefficients du modulateur. En augmentant la valeur du condensateur d'intégration C_2 , on augmente *ipso facto* la valeur des condensateurs d'échantillonnage, car les ratios doivent rester constants afin de conserver les mêmes coefficients.

Cependant, en augmentant la taille des condensateurs, la taille du circuit ainsi que la puissance consommée [24, 25] sont également augmentées. Le processus d'optimisation doit donc déterminer la taille des condensateurs permettant d'obtenir le SNR désiré tout en minimisant la taille du circuit et la puissance consommée.

2.2.2.3 Bruit des amplificateurs opérationnels

La troisième et dernière source de bruit d'un intégrateur à capacités commutées est le bruit généré par l'amplificateur opérationnel. La puissance de cette source de bruit dépend grandement du type de circuit utilisé pour réaliser l'amplificateur [23]. Il n'existe donc pas de formules universelles pour calculer sa valeur. Cela complique sa prise en considération dans un modèle de simulation de plus haut niveau. Pour contourner ce problème, une méthode est suggérée dans [22]. Il s'agit d'utiliser une simulation au niveau transistor pour obtenir la puissance de bruit générée par l'amplificateur opérationnel. Par la suite, la valeur obtenue peut être utilisée dans des simulations de plus haut niveau.

Dans cet ouvrage le bruit généré par les amplificateurs opérationnels ne sera pas considéré. Ce choix repose sur deux facteurs. Premièrement, le bruit des amplificateurs opérationnels ne dépend pas directement de paramètres de conception comme la taille des condensateurs d'intégration ou les coefficients du modulateur. Minimiser la puissance du bruit générée par un amplificateur opérationnel est essentiellement un travail de conception analogique. Deuxièmement, il peut être démontré que la source de bruit dominante dans un intégrateur à capacités commutées est le bruit thermique des interrupteurs [21].

2.2.3 Plage dynamique des intégrateurs

Les convertisseurs $\Sigma\Delta$ peuvent être utilisés à l'intérieur d'applications réalisées en technologie CMOS. Dans ces processus de fabrication optimisés pour des circuits numériques, la tension d'alimentation est faible. Par exemple, pour la technologie CMOS $0,18\ \mu\text{m}$ utilisée au chapitre 4, la tension d'alimentation maximale est 1,8V. De plus, la diminution constante de la taille des transistors entraîne inexorablement la tension d'alimentation vers des valeurs de plus en plus faibles.

En réduisant la tension d'alimentation, la plage d'opération linéaire des amplificateurs opérationnels est également réduite. Il est donc primordial lors de la conception d'un convertisseur

$\Sigma\Delta$ de s'assurer que l'amplitude de sortie des intégrateurs soit en accord avec la plage linéaire des amplificateurs.

Lorsque l'amplitude de sortie d'un intégrateur est supérieure à la plage dynamique de l'amplificateur opérationnel, le signal est écrêté au niveau de saturation. Ce fonctionnement non linéaire se traduit au niveau spectral par de la distorsion harmonique [23].

Pour illustrer l'impact de ce problème, le modèle Simulink de la figure 2.7 est utilisé. L'amplitude du signal d'entrée est de -6 dBFS. Le bloc saturation à la sortie de chaque intégrateur modélise la plage dynamique limitée des intégrateurs. Les résultats de la simulation sont présentés à la figure 2.8. La courbe en vert représente la sortie du convertisseur lorsque la plage dynamique des intégrateurs est infinie, c'est-à-dire lorsqu'il n'y a aucune contrainte à la sortie des intégrateurs. La courbe en bleu illustre ce qui se produit lorsque la plage dynamique des intégrateurs est limitée à une valeur égale à la plage d'entrée pleine échelle du quantificateur. On remarque l'apparition de nombreuses raies spectrales causées par la distorsion du signal ainsi qu'une augmentation du plancher de bruit. Même si les raies qui se retrouvent à l'extérieur de la bande de fréquences du signal sont éliminées lors du filtrage, celles qui restent affectent grandement le SNR à la sortie du convertisseur. Dans cet exemple, la distorsion harmonique entraîne une réduction du SNR de 25 dB.

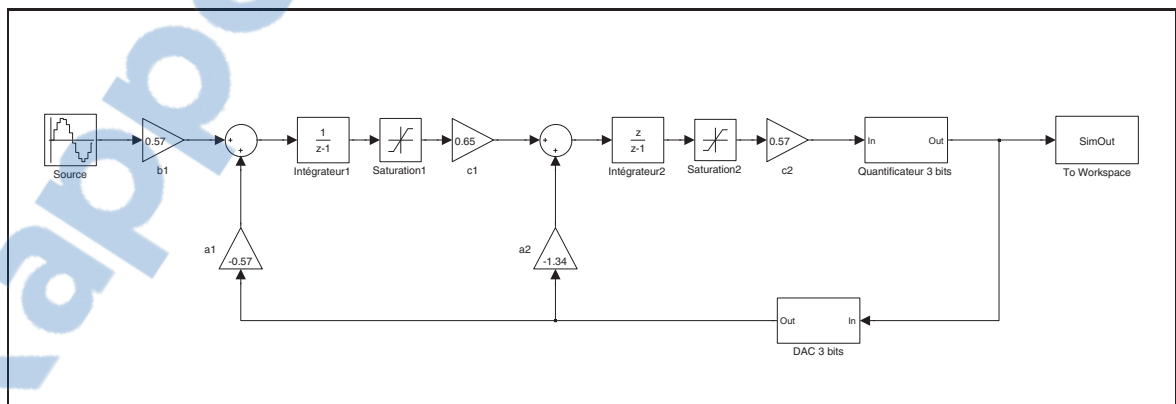


Figure 2.7 Modèle Simulink utilisé pour simuler l'impact de la plage dynamique limitée des intégrateurs

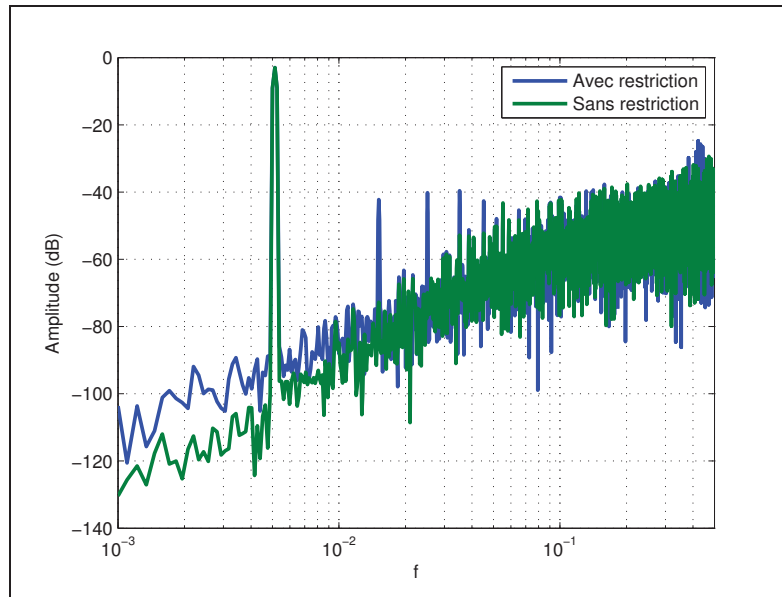


Figure 2.8 Résultats de la simulation du modèle de la figure 2.7 avec et sans restriction de la plage dynamique des intégrateurs

Dans l'exemple précédent, le bloc saturation utilisé écrêtait le signal à une valeur positive et négative précise. Entre ces deux valeurs, le comportement est parfaitement linéaire. Il s'agit d'une approximation du fonctionnement d'un amplificateur opérationnel réel. La figure 2.9 présente la courbe entrée sortie d'un amplificateur opérationnel typique. Plus l'amplitude de sortie est élevée, plus la relation entre l'entrée et la sortie devient non linéaire. En se basant sur ce constat, le concepteur a tout avantage à minimiser l'amplitude de sortie des intégrateurs [23].

2.2.3.1 Réduction de la distorsion harmonique

Pour réduire l'impact de la distorsion harmonique, l'amplitude de sortie des intégrateurs doit être contrôlée. Pour y arriver, diverses techniques sont possibles.

Tout d'abord, au niveau circuit, il est possible de doubler la plage dynamique des intégrateurs en utilisant des amplificateurs entièrement différentiels [23]. Ce type de circuit a également l'avantage d'éliminer le bruit en mode commun et les harmoniques paires. La réalisation de coefficients négatifs est également simplifiée. En effet, il suffit de permuter les deux phases

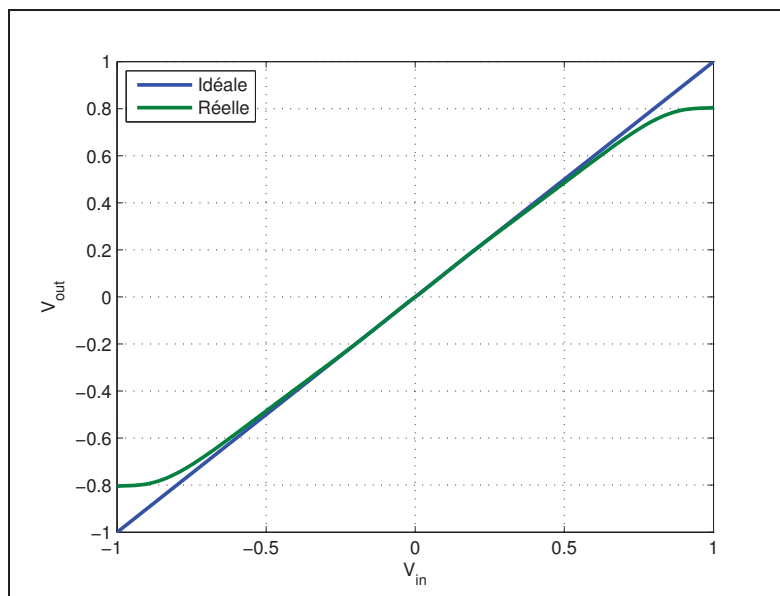


Figure 2.9 Relation (idéale et réelle) entre l'entrée et la sortie d'un amplificateur opérationnel

pour inverser le signal. Les principaux inconvénients liés à l'utilisation de signaux différentiels sont une augmentation de la taille et de la complexité du circuit [23].

Tel que mentionné à la section 1.8.1, les coefficients d'entrées b_i ont un impact sur l'amplitude de sortie des intégrateurs. En choisissant $b_i = a_i$ pour $i \leq N$ et $b_{N+1} = 1$ (où N représente l'ordre du modulateur), les intégrateurs se retrouvent à traiter uniquement le bruit de quantification, ce qui diminue leur amplitude de sortie.

L'impact des coefficients b_i peut être visualisé en traçant la distribution de la sortie des intégrateurs. L'exercice a été fait pour le convertisseur de la figure 2.10. Dans un premier temps, la figure 2.11 montre la distribution de la sortie des intégrateurs lorsque les coefficients b_2 et b_3 sont forcés à 0. Dans cette configuration, la distribution est essentiellement concentrée aux extrémités de la plage dynamique. À l'inverse, la figure 2.12 montre la distribution de la sortie des intégrateurs lorsque les coefficients sont laissés tels qu'indiqués à la figure 2.10. Cette fois, la distribution de la sortie des intégrateurs est concentrée autour de zéro. Cette seconde configuration est donc moins sensible aux non-linéarités des amplificateurs opérationnels [26].



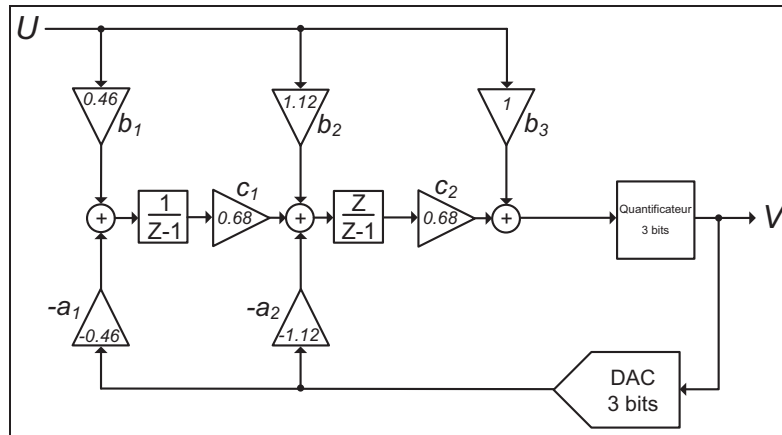


Figure 2.10 Convertisseur utilisé pour visualiser l'impact des coefficients b_i sur la distribution de la sortie des intégrateurs

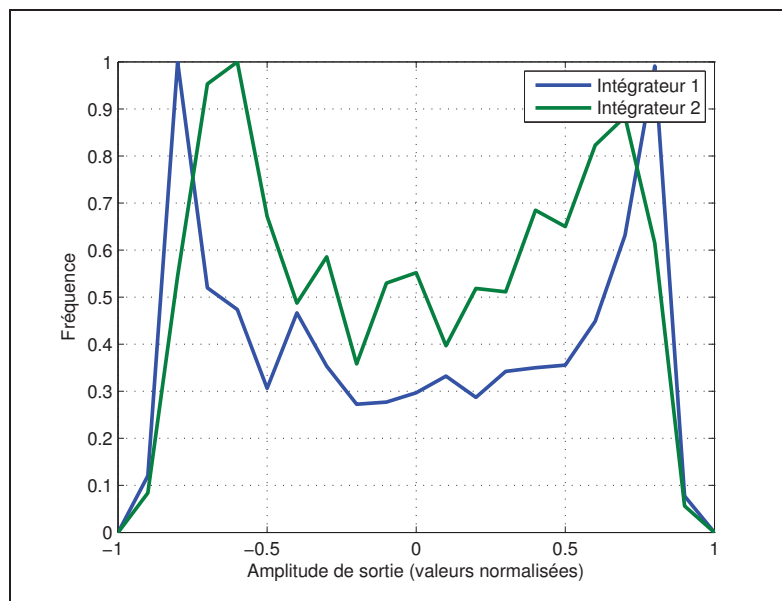


Figure 2.11 Distribution de la sortie des intégrateurs lorsque le coefficient $b_1 = a_1$ et que les autres coefficients $b_i = 0$

2.2.3.2 Effet des coefficients interétagés

Les coefficients interétagés c_i peuvent être utilisés pour contraindre la sortie des intégrateurs à l'intérieur d'une certaine plage. Un modulateur $\Sigma\Delta$ étant un système linéaire, il est possible d'y

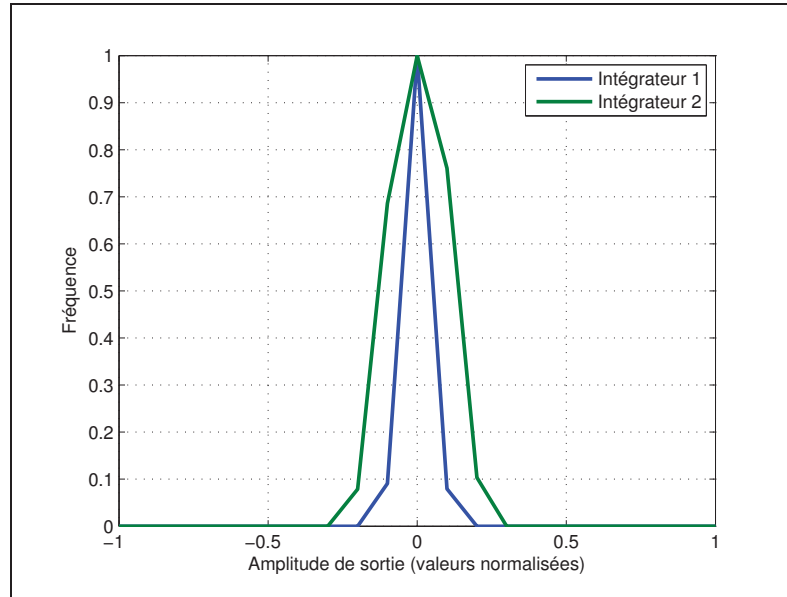


Figure 2.12 Distribution de la sortie des intégrateurs lorsque les coefficients $b_i = a_i$ pour $i \leq N$ et $b_{N+1} = 1$

appliquer un facteur d'échelle k sans en modifier la fonction de transfert. La mise à l'échelle d'un système linéaire se fait en divisant les branches d'entrées par un facteur k et en multipliant la branche de sortie par le même facteur k .

Ce principe est illustré à la figure 2.13 où $H(z)$ représente un système linéaire quelconque. Dans ce système, la valeur du facteur k ne change pas les fonctions de transfert $x_3(z)/v(z)$ et $x_3(z)/w(z)$. Le facteur k peut ainsi être utilisé pour ajuster la plage dynamique des signaux qu'aura à traiter le système linéaire $H(z)$ sans affecter la fonction de transfert du système complet.

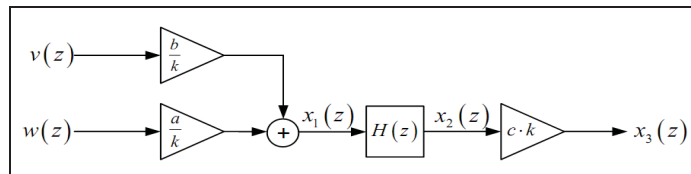


Figure 2.13 Application d'un facteur d'échelle sur un système linéaire

Dans le cas d'un modulateur $\Sigma\Delta$, la mise à l'échelle doit être appliquée à chaque intégrateur. Pour chacun d'eux, la valeur k appropriée est déterminée par :

$$k = \frac{x_{max}}{x_{lim}} \quad (2.7)$$

où x_{max} représente la valeur maximale à la sortie de l'intégrateur lorsque la valeur de k est 1 et x_{lim} représente la valeur maximale désirée. La valeur de x_{max} peut être trouvée par simulation. Avec cette technique, l'amplitude de sortie de chaque intégrateur peut être ajustée de façon à être contenue à l'intérieur d'une certaine plage. De plus, en diminuant l'amplitude de sortie des intégrateurs, la consommation en puissance est diminuée [27, 28]. Toute cette opération de mise à l'échelle peut être réalisée à l'aide de la librairie Matlab *Delsig Toolbox* [12]. Cet outil logiciel sera présenté plus en détail à la section 2.3.1.

2.2.4 Taille du circuit

La taille occupée par le circuit est un critère d'optimisation important. Alors que les circuits numériques bénéficient de la réduction constante de la taille des transistors, cela n'est pas le cas des circuits analogiques et mixtes. Les transistors y occupent généralement une place non significative de la surface totale du circuit. Les composantes les plus volumineuses sont les résistances et les condensateurs. Pour ce type de composantes, la surface occupée est proportionnelle à leur valeur.

Pour optimiser l'espace occupé par un circuit, il faut pouvoir l'estimer en fonction des paramètres de conception. Il n'est pas nécessaire de connaître la valeur exacte de la superficie du circuit. Il faut déterminer les relations entre les paramètres à optimiser et la superficie occupée par le circuit. La taille du circuit peut alors être minimisée sans connaître la superficie exacte que le circuit occupera.

Un convertisseur $\Sigma\Delta$ peut être divisé en trois sections : le modulateur, le quantificateur et le DAC de la branche de retour. Pour ces deux derniers, la superficie occupée est déterminée principalement par le nombre de bits de quantification. En fixant le nombre de bits de quantifi-

cation, la superficie qu'occupent le quantificateur et le DAC devient une constante qui est alors exclue du processus d'optimisation.

Pour le modulateur, la superficie occupée est déterminée principalement par les différents condensateurs des intégrateurs. La superficie d'un condensateur étant proportionnelle à sa valeur, la superficie du modulateur sera proportionnelle à la somme des capacités. C'est cette méthode qui sera utilisée dans cet ouvrage afin d'estimer la taille d'un convertisseur $\Sigma\Delta$. En d'autres mots, pour minimiser la taille du modulateur, il faut minimiser la somme des capacités.

2.3 Techniques d'optimisation existantes

Cette seconde moitié du chapitre présente différentes techniques existantes pour optimiser les convertisseurs $\Sigma\Delta$. Dans un premier temps, un outil logiciel largement utilisé, la librairie Matlab *Delsig Toolbox* [12], est exposé. Dans un deuxième temps, des algorithmes d'optimisation tirés de la littérature sont présentés.

2.3.1 Delsig Toolbox

La librairie *Delsig Toolbox* [12] est un ensemble de fonctions Matlab développées pour faciliter la conception et l'analyse des convertisseurs $\Sigma\Delta$. Cet outil est disponible gratuitement pour téléchargement sur le site de Mathworks. Les différentes fonctions disponibles permettent de définir la NTF et la STF, de calculer les coefficients et de simuler le convertisseur. La librairie Matlab *Delsig Toolbox* étant fréquemment utilisée dans plusieurs publications traitant des convertisseurs $\Sigma\Delta$, elle est présentée en premier. De plus, plusieurs de ses fonctions seront utilisées ultérieurement dans cet ouvrage.

Dans un processus de conception standard, la première étape consiste à déterminer l'ordre du modulateur, l'OSR et le nombre de bits de quantification. Ces spécifications sont généralement déterminées en fonction du SNR requis par l'application. À l'intérieur de [1], ouvrage cosigné par l'auteur de la librairie *Delsig toolbox*, on retrouve trois graphiques fournissant le SNR maximal qu'il est possible d'atteindre pour des modulateurs d'ordre 1 à 8 en fonction de l'OSR

pour des quantificateurs de 1 à 3 bits. À titre de référence, ceux-ci sont reproduits aux figures 2.14 à 2.16. Ces courbes sont le résultat de simulations empiriques. Elles tiennent compte de la réduction d'amplitude nécessaire pour garantir la stabilité du système. De plus amples détails sur la méthodologie utilisée sont disponibles dans [13].

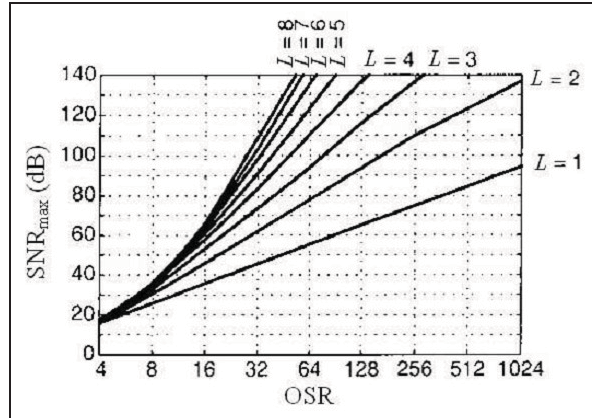


Figure 2.14 SNR maximal obtenu par des simulations empiriques avec un quantificateur de 1 bit
Tirée de [1]

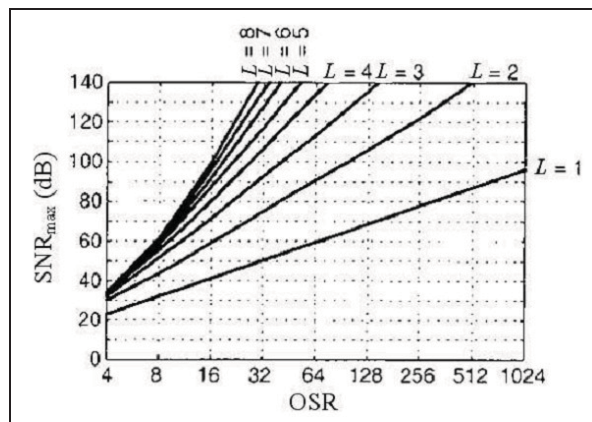


Figure 2.15 SNR maximal obtenu par des simulations empiriques avec un quantificateur de 2 bits
Tirée de [1]

En observant ces courbes, on remarque que pour un même objectif de SNR, plusieurs configurations sont possibles. Par exemple, pour obtenir un SNR de 100 dB, il est possible d'utiliser un

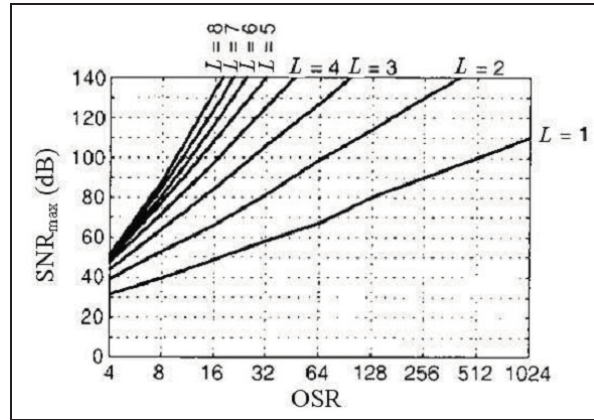


Figure 2.16 SNR maximal obtenu par des simulations empiriques avec un quantificateur de 3 bits
Tirée de [1]

modulateur d'ordre 6 avec un OSR de 32 et un quantificateur de 1 bit ou bien encore, un modulateur d'ordre 3 avec un OSR de 32 et un quantificateur de 3 bits. Il est à noter que ces courbes ne tiennent pas compte des limitations imposées par l'utilisation de composants réelles. Notamment, le bruit thermique n'est pas pris en considération. Il est donc préférable de faire un choix qui offre une certaine marge de design afin de pallier les non idéalités d'un circuit réel.

Une fois la configuration de base est sélectionnée (ordre du modulateur, OSR et nombre de bits de quantification), l'étape suivante consiste à obtenir une NTF correspondant à ces critères. Cette étape est réalisée par la fonction *synthesizeNTF*. Cette fonction reçoit en entrée cinq paramètres : l'ordre du modulateur, l'OSR, le nombre de bits de quantification, la valeur maximale de $|NTF(e^{j\omega})|$ et une variable booléenne spécifiant si oui ou non le placement des zéros doit être optimisé. Tel que mentionné à la section 1.8.2, l'utilisation de zéros optimisés requière la présence de résonateurs dans le circuit tandis que la valeur maximale de $|NTF(e^{j\omega})|$ a un impact majeur sur les performances et la stabilité du système (voir section 1.7.3). En sortie, la fonction retourne un objet fonction de transfert sous la forme des coefficients de puissances descendantes de z . Il s'agit de la NTF générée en fonction des paramètres désirés.

Lorsqu'une NTF et une STF satisfaisantes sont sélectionnées, elles peuvent être réalisées à l'intérieur d'une structure à l'aide de la fonction *realizeNTF*. Cette fonction permet le calcul des coefficients a_i , b_i , g_i et c_i pour différentes structures de réalisation dont celles présentées à la section 1.8. Il est à noter que cette fonction place tous les coefficients interétages c_i à 1.

Tel qu'expliqué à la section 2.2.3.1, les coefficients c_i peuvent être utilisés afin de borner l'amplitude de sortie de chaque étage d'intégration. Cette étape est réalisée par la fonction *scaleABCD*. Cette fonction utilise une représentation d'état du modulateur. L'avantage lié à l'utilisation d'une matrice d'état pour la mise à l'échelle des coefficients est que la fonction devient indépendante de la structure de réalisation. En effet, pour utiliser une structure de réalisation différente, il suffit de créer deux fonctions : l'une pour convertir les coefficients en représentation d'état et l'autre pour convertir la représentation d'état en coefficient.

En conclusion, la librairie *Delsig Toolbox* est un outil logiciel simple d'utilisation qui est largement répandu. En fonction des critères de base, la NTF correspondante est synthétisée. Celle-ci est ensuite réalisée dans la structure désirée. Finalement, les coefficients de la structure peuvent être mis à l'échelle afin de limiter l'amplitude de sortie des intégrateurs à l'intérieur de la plage désirée.

2.3.2 Algorithmes d'optimisation

Dans la littérature scientifique, plusieurs travaux présentent des techniques d'optimisation pour les convertisseurs $\Sigma\Delta$. L'objectif de cette section est de synthétiser les principales publications. Ceci permettra de mettre en contexte les solutions proposées au chapitre suivant.

Un algorithme génétique pour optimiser les modulateurs $\Sigma\Delta$ est présenté en [29]. Un algorithme génétique est une méthode de recherche semi-aléatoire basée sur la sélection naturelle. La méthode consiste à créer une population où chaque individu est représenté par un chromosome contenant les paramètres à optimiser. Une fonction permettant de déterminer les chromosomes offrant les meilleures performances est utilisée afin de sélectionner les parents de la prochaine génération. Les chromosomes enfants sont générés en fonction des chromosomes

parents par une fonction de reproduction. Le processus se répète jusqu'à ce qu'un chromosome offrant les performances désirées soit trouvé. La figure 2.17 illustre le principe de l'algorithme génétique.

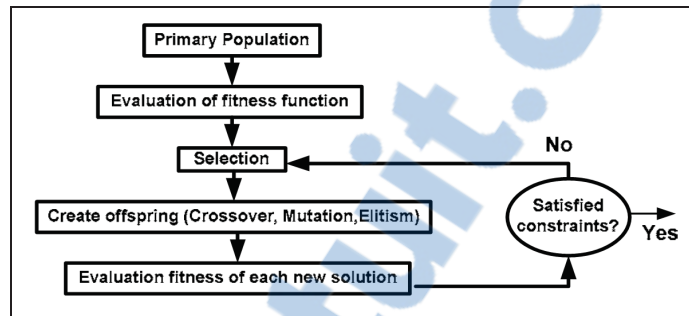


Figure 2.17 Principe d'optimisation par un algorithme génétique
Tirée de [29]

L'algorithme présenté en [29] s'effectue en deux étapes. La première permet de déterminer les paramètres systèmes : l'ordre du modulateur, le nombre de bits de quantification et l'OSR. Ensuite, une deuxième étape permet d'optimiser des paramètres au niveau circuit comme la vitesse de balayage, le gain DC et l'amplitude de sortie des intégrateurs. Cet algorithme utilise plusieurs fonctions de la librairie Matlab *Delsig Toolbox* pour la synthèse des fonctions de transfert et l'évaluation des performances du modulateur.

Une étude visant à trouver les coefficients optimaux est présentée en [30]. La structure étudiée est présentée à la figure 2.18. Dans cette structure, tous les coefficients de rétroaction ont été placés à 1. Les seuls coefficients qui sont donc optimisés sont les coefficients interétages c . Les différentes combinaisons de coefficients sont simulées afin de trouver l'ensemble offrant les meilleures performances. L'objectif n'est donc pas d'optimiser la réalisation d'une NTF dans la structure, mais bien plutôt de trouver directement les coefficients offrant les performances optimales. Trois critères sont utilisés pour évaluer les performances du convertisseur : le SNR pour une amplitude d'entrée normalisée de 0,25, la plage dynamique du convertisseur et la valeur d'entrée maximale assurant la stabilité du système. L'analyse des résultats obtenus per-

met de constater qu'en diminuant les coefficients interétages, la plage dynamique est diminuée tandis qu'en augmentant les coefficients interétages, la stabilité du système est diminuée.

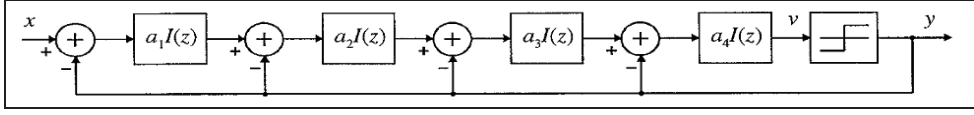


Figure 2.18 Structure étudiée afin de déterminer les paramètres optimaux
Tirée de [30]

Une autre méthode visant à trouver directement les coefficients optimaux est présentée en [31]. Tout comme la méthode présentée en [30], seuls les coefficients interétages sont optimisés. L'algorithme vise à trouver les coefficients offrant le SNR le plus élevé et qui respecte une plage d'entrée stable minimale. Pour imposer cette dernière contrainte, une étude empirique permet d'établir une relation entre la plage d'entrée stable et le gain de la puissance du bruit (NPG) défini par :

$$NPG = \frac{1}{\pi} \int_0^{\pi} |NTF(e^{j\omega})|^2 d\omega \quad (2.8)$$

La méthode de recherche retient donc uniquement les ensembles de coefficients qui résulte en une NTF dont la NPG est susceptible d'offrir la plage d'entrée stable désirée. Parmi les ensembles de coefficients répondant à ce critère, celui offrant le SNR le plus élevé est sélectionné à la fin.

En [32], une méthodologie de conception pour les convertisseurs $\Sigma\Delta$ passe-bande est présentée. La figure 2.19 résume le processus utilisé. Celui-ci se concentre essentiellement à déterminer la NTF optimale pour une application donnée. En fonction des critères de conception, l'ordre du modulateur, l'OSR ainsi que le nombre de bits de quantification idéal est trouvé. Un critère permettant d'estimer la stabilité de la NTF en fonction de la NPG est utilisé pour assurer que le convertisseur reste à l'intérieur des limites de la stabilité. Une fois la NTF idéale trouvée, celle-ci peut être réalisée dans la structure voulue. Les coefficients sont mis à l'échelle pour éviter des problèmes de saturation. Aucun détail n'est cependant fourni sur la façon dont cette mise à l'échelle est effectuée.

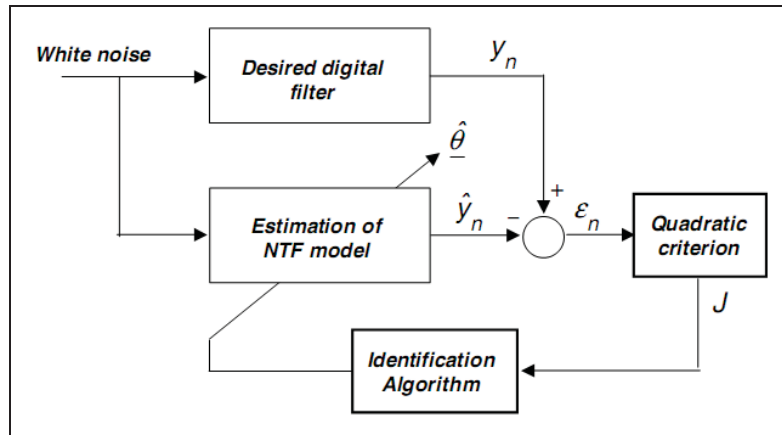


Figure 2.20 Processus d'optimisation par un critère quadratique
Tirée de [33]

convertisseur devant être optimisé : la complexité de la structure, la sensibilité du modulateur aux variations des coefficients et la consommation en puissance. Cette dernière fonction coût est basée sur la taille des capacités du modulateur. Une taille minimale est imposée afin de respecter les contraintes du bruit thermiques. L'ensemble du processus est également soumis à une contrainte sur l'amplitude de sortie des intégrateurs afin de s'assurer que celle-ci respecte la plage linéaire des amplificateurs. L'ensemble du processus d'optimisation est présenté à la figure 2.22. Selon les spécifications initiales (SNR et plage dynamique), les paramètres optimaux pour la NTF (ordre du modulateur, OSR et $\max|NTF(e^{j\omega})|$) sont déterminés par un processus itératif. En conclusion il s'agit d'un algorithme d'optimisation complexe permettant de faire la conception complète d'un modulateur $\Sigma\Delta$ en fonction de plusieurs paramètres. Au chapitre 3, les méthodes présentées seront des algorithmes d'optimisation plus simple permettant de déterminer les coefficients interétages et la taille des capacités optimales pour une structure et une NTF données, en prenant en considération le bruit thermique et l'amplitude de sortie des intégrateurs.

2.3.3 Analyses des méthodes existantes

Suite à la présentation des techniques d'optimisation existantes, certains constats peuvent être faits. La majorité des méthodes se concentrent à déterminer la NTF optimale pour une appli-

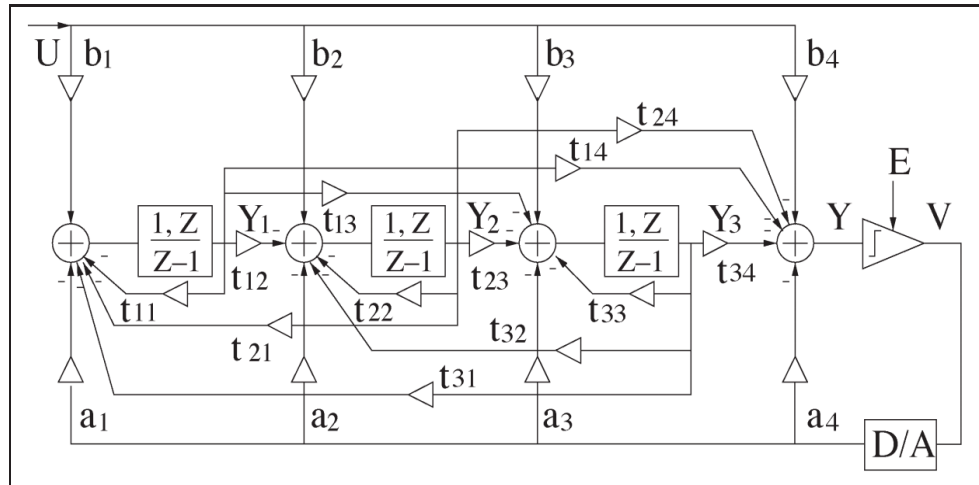


Figure 2.21 Structure de réalisation générique
Tirée de [34]

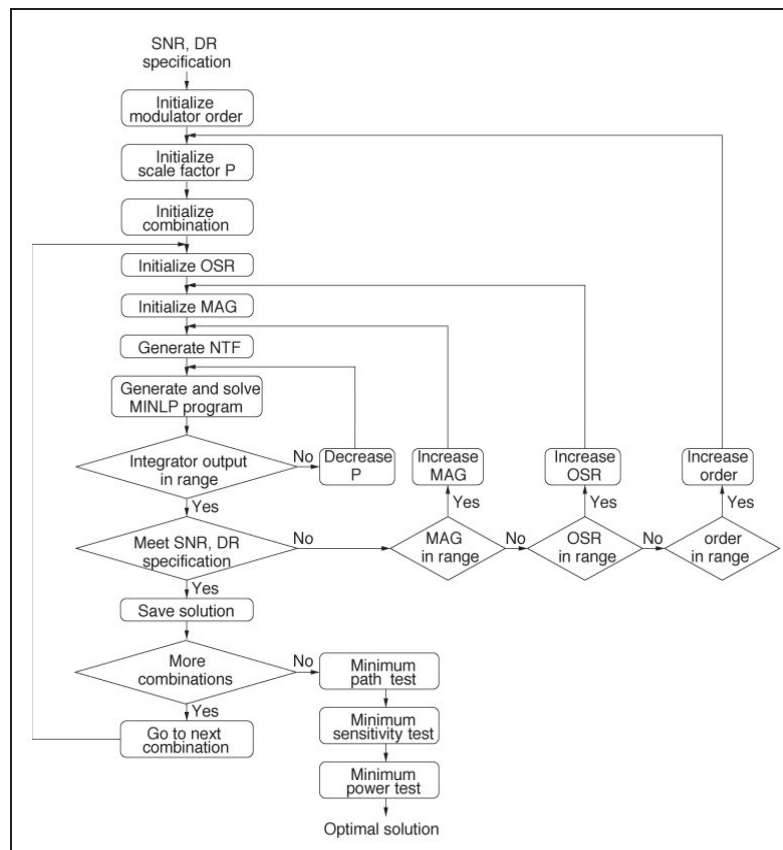


Figure 2.22 Algorithme permettant d'optimiser à la fois
la NTF et la structure de réalisation
Tirée de [34]

cation donnée. Cependant, à notre connaissance, la réalisation optimale d'une NTF donnée, en prenant en considération l'impact du bruit thermique, demeure une problématique de recherche qui a été peu étudiée. La majorité des méthodes présentées ne tiennent pas compte du bruit thermique.

Si certaines méthodes s'assurent de mettre à l'échelle les modulateurs afin de limiter l'amplitude de sortie des intégrateurs, aucune méthode pour déterminer les amplitudes optimales à la sortie de chaque intégrateur n'a été trouvée.

Au chapitre suivant, deux méthodes d'optimisation développées dans le cadre de cet ouvrage seront présentées. Celles-ci s'inscrivent dans la continuité de ce domaine de recherche.

2.4 Conclusion

La première partie de ce chapitre a présenté les différents critères devant être optimisés lors de la conception d'un modulateur $\Sigma\Delta$. Tout d'abord, le bruit thermique généré par le circuit est une des principales limitations du SNR. Pour diminuer la puissance du bruit thermique, il faut augmenter la taille des condensateurs d'échantillonnage. Ceci a cependant un impact majeur sur la puissance consommée et sur la taille du circuit. Celle-ci est généralement proportionnelle à la somme de toutes les capacités du circuit. Finalement, l'amplitude de sortie des intégrateurs doit également être minimisée afin d'éviter des problèmes de distorsion harmonique et pour diminuer la puissance consommée par le circuit.

Dans la deuxième partie de ce chapitre, diverses techniques d'optimisation des paramètres des modulateurs $\Sigma\Delta$ ont été présentées. Les objectifs visés diffèrent d'une technique à l'autre. Certains tentent de déterminer directement les coefficients d'une structure générant le SNR le plus élevé alors que d'autres se concentrent à déterminer la NTF offrant les meilleures performances. Cependant, il a été constaté que de façon générale elles ne prennent pas en considération l'effet du bruit thermique. Il s'agit pourtant de l'un des éléments principaux limitant le SNR des convertisseurs.

CHAPITRE 3

SOLUTIONS PROPOSÉES

3.1 Introduction

Dans le chapitre précédant, les différents aspects pouvant être optimisés durant le processus de conception d'un convertisseur $\Sigma\Delta$ ont été exposés. En se basant sur ces critères, deux méthodes d'optimisation ont été développées dans le cadre de cette maîtrise. [2, 3]. La première, appelée méthode des fenêtres, permet de déterminer l'amplitude de sortie de chaque intégrateur en fonction d'une certaine pénalité en SNR. La seconde est un processus d'optimisation multicritères permettant de déterminer à la fois l'amplitude de sortie des intégrateurs et la taille des capacités du modulateur.

3.2 Objectif

Avant de commencer la présentation des solutions proposées, il est de mise de récapituler les objectifs initiaux qui ont motivé ce travail de recherche. La problématique étudiée est l'optimisation des coefficients interétages des modulateurs $\Sigma\Delta$. Tel que vu au chapitre précédent, les coefficients interétages sont généralement utilisés pour ajuster l'amplitude de sortie des intégrateurs à la plage linéaire des amplificateurs. Il existe cependant peu de travaux rapportés dans la littérature concernant l'optimisation de cette étape. L'absence de solution unique à ce problème confère à cette étape du processus de conception un potentiel d'optimisation. L'objectif de base est donc de déterminer une façon optimale d'utiliser le degré de liberté fourni par les coefficients interétages.

Il est à noter que la synthèse de la NTF ne fait pas partie de l'objectif d'optimisation. Le but est d'optimiser la réalisation d'une NTF donnée à l'intérieur d'une structure. Les deux méthodes d'optimisation qui seront présentées ont donc été développées avec comme prémisse qu'une NTF idéale pour l'application a déjà été trouvée. Ce choix se justifie par la revue de littérature

présentée au chapitre précédent. Tel qu'il y était exposé, il existe déjà de nombreux algorithmes pour déterminer la NTF idéale pour une application donnée.

3.3 Mise à l'échelle du modulateur vs SNR

À la section 2.2.3.1, il a été démontré que les coefficients interétages c_i peuvent être utilisés pour appliquer un facteur d'échelle au modulateur. Cette étape permet d'ajuster l'amplitude de sortie des intégrateurs. Dans cette section, l'étude de cette opération sera approfondie. Plus précisément, nous allons analyser l'impact de cette opération sur le SNR à la sortie du convertisseur.

Lorsqu'un facteur d'échelle est appliqué à un système linéaire, sa fonction de transfert n'est pas modifiée. De prime abord, on peut alors conclure que peu importe le facteur d'échelle appliqué à un modulateur $\Sigma\Delta$, cette opération n'a pas d'impact sur le SNR à la sortie du convertisseur. Cette affirmation serait vraie uniquement si le modulateur était un système idéal. Dans le cas d'un modulateur réel, il faut être prudent dans le choix du facteur d'échelle et analyser adéquatement les répercussions qu'a ce choix sur les performances du convertisseur.

Les deux principales non-idéalités qui influencent le choix du facteur d'échelle sont la plage dynamique des amplificateurs et le bruit thermique des interrupteurs. Dans le premier cas, le facteur d'échelle choisi doit conserver l'amplitude des signaux à l'intérieur de la plage linéaire des amplificateurs. Dans le cas contraire, des problèmes de saturation dégraderont rapidement le SNR du convertisseur (voir section 2.2.3). Pour contrer ce problème, on peut être tenté de minimiser l'amplitude des signaux. Cependant, comme nous allons le démontrer, en raison du bruit thermique des interrupteurs, cela peut également entraîner une dégradation du SNR du convertisseur.

Tel qu'expliqué au chapitre précédent, le bruit thermique des interrupteurs peut être modélisé par une source de bruit ajoutée à l'entrée de chaque intégrateur. Pour quantifier leur impact sur le SNR à la sortie du convertisseur, il faut déterminer leurs fonctions de transfert. Celles-ci sont

de la forme :

$$H_{ni}(z) = \frac{V(z)}{V_{ni}(z)} \quad (3.1)$$

où $V(z)$ est la sortie du convertisseur $\Sigma\Delta$ et $V_{ni}(z)$ sont les sources de bruit thermique tel qu'illustrées à la figure 3.1.

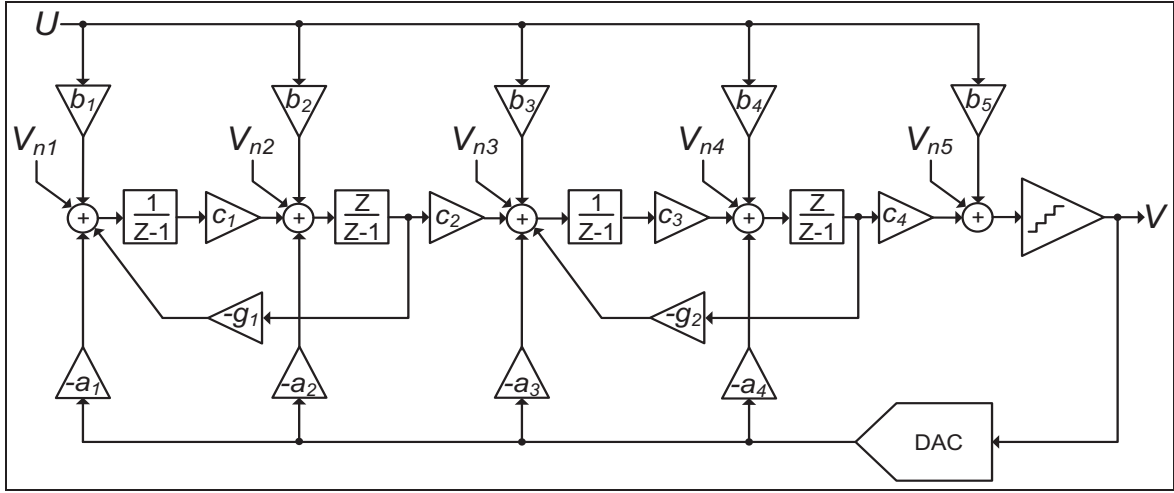


Figure 3.1 Structure CRFB pour un modulateur du quatrième ordre. Les sources équivalentes de bruit thermique à l'entrée de chaque intégrateur sont étiquetées V_{ni}

Les fonctions de transfert du bruit thermique partagent toutes le même dénominateur. Il s'agit du même dénominateur que la NTF qui est indépendante du facteur d'échelle appliqué au modulateur. Donc, seuls les numérateurs des fonctions de transfert du bruit thermique sont affectés par la mise à l'échelle du modulateur. À titre d'exemple, les numérateurs des fonctions de transfert des sources de bruit thermique de la figure 3.1 sont donnés par :

$$N_{n1}(z) = c_1 c_2 c_3 c_4 z \quad (3.2)$$

$$N_{n2}(z) = c_2 c_3 c_4 z (z - 1) \quad (3.3)$$

$$N_{n3}(z) = c_3 c_4 (z^2 + (c_1 g_1 - 2)z + 1) \quad (3.4)$$

$$N_{n4}(z) = c_4 (z^3 + (c_1 g_1 - 3)z^2 + (3 - c_1 g_1)z - 1) \quad (3.5)$$

Ces équations démontrent qu'en augmentant la valeur des coefficients interétages c_i , la puissance du bruit thermique à la sortie du convertisseur est augmentée. Cela démontre également que les intégrateurs précédés par d'autres étages d'intégration bénéficient de la mise en forme de leur bruit de quantification. Cet effet peut être visualisé à la figure 3.2 où l'amplitude de la réponse en fréquence des fonctions de transfert est tracée (les coefficients interétages sont mis à 1). Cette figure confirme le fait bien connu [35,36] que le premier intégrateur de la chaîne est celui qui a l'impact le plus significatif sur les performances du convertisseur.

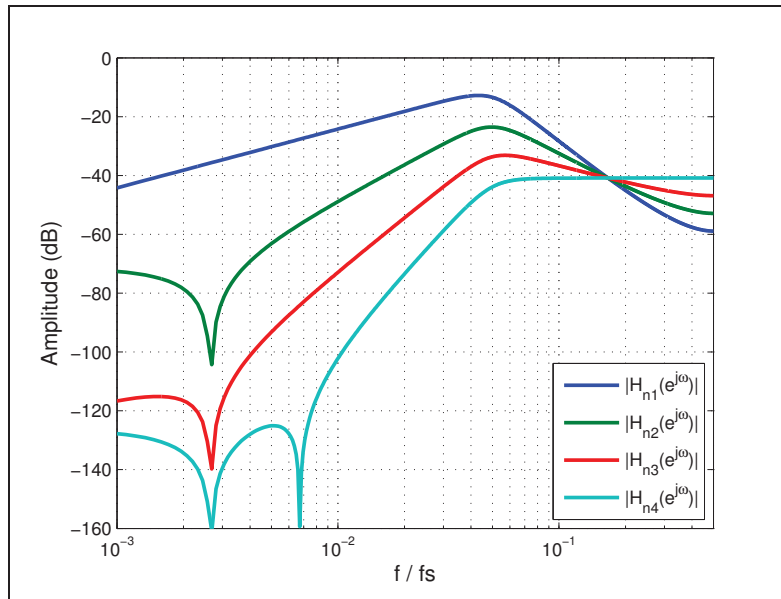


Figure 3.2 Amplitude de la réponse en fréquence des fonctions de transfert des sources de bruit thermique du modulateur de la figure 3.1

Maintenant, observons ce qui se produit lorsqu'un facteur d'échelle est appliqué au modulateur. La figure 3.3 montre le SNR obtenu en fonction de l'amplitude de sortie des intégrateurs (en valeurs normalisées). Les amplitudes de sortie de tous les intégrateurs du modulateur sont ajustées à la même valeur. On remarque que plus l'amplitude de sortie est réduite, plus le SNR est réduit. Ce constat est valide lorsque la réduction d'amplitude est faite de façon globale à tout le modulateur. Pris de façon individuelle, l'effet de chaque intégrateur dépend de sa position dans la chaîne. Ce principe est illustré à la figure 3.4. Dans cette figure, l'amplitude de sortie d'un seul intégrateur à la fois est réduite (l'amplitude de sortie des autres intégrateurs reste

maximale). Cette figure permet de constater que le premier intégrateur est celui ayant l'impact le plus important. Cependant, les deuxième et troisième intégrateurs ont également un effet qui ne peut être entièrement ignoré. Le quatrième intégrateur a un impact non significatif.

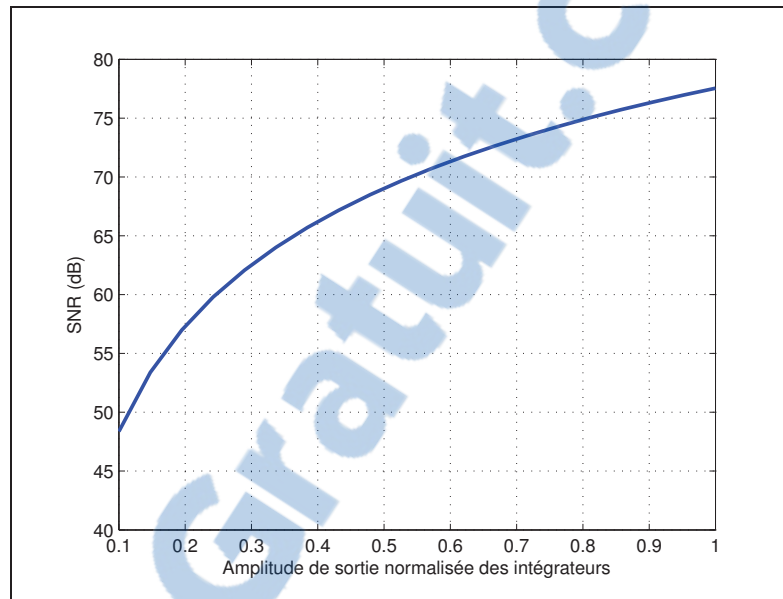


Figure 3.3 Effet de la mise à l'échelle du modulateur sur le SNR. Tous les intégrateurs sont ajustés à la même amplitude de sortie

3.4 Méthode des fenêtres

Cette première méthode est élaborée par l'observation des résultats présentés à la figure 3.4. Selon ce graphique, si l'unique objectif de la conception est de maximiser les performances en SNR, l'amplitude de sortie des intégrateurs doit être maximisée. Il faut cependant mettre deux bémols à ce constat. Premièrement, pour obtenir les résultats de la figure 3.4, les amplificateurs doivent être linéaires sur toute leur plage d'utilisation. Dans le cas inverse, les performances seront détériorées par la distorsion harmonique (voir section 2.2.3). Deuxièmement, pour réduire la consommation en puissance, il est avantageux de diminuer l'amplitude de sortie des intégrateurs [27, 28]. Pour ces deux raisons, accepter une légère pénalité en SNR afin de permettre de diminuer l'amplitude de sortie des intégrateurs peut être bénéfique.

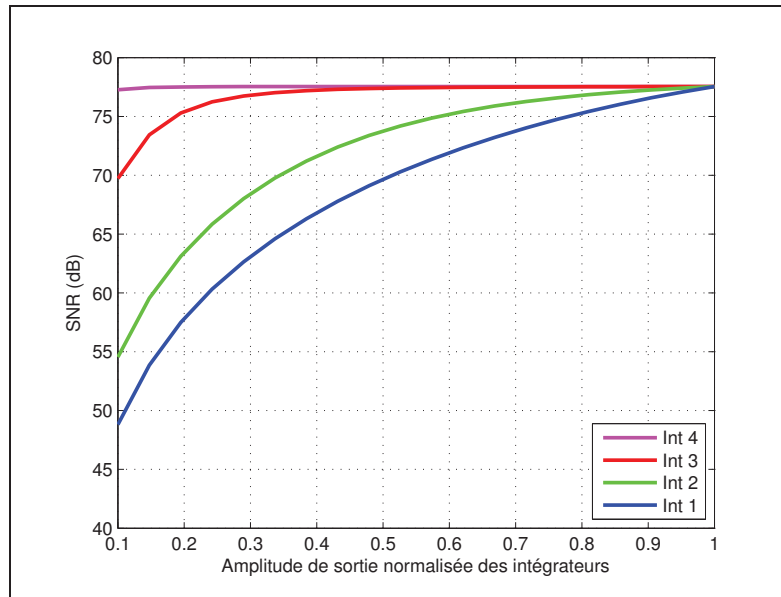


Figure 3.4 Effet de la mise à l'échelle individuelle de chaque intégrateur du modulateur sur le SNR

La méthode des fenêtres a pour objectif de déterminer un compromis entre les performances en SNR et l'amplitude de sortie des intégrateurs. L'objectif est de minimiser l'amplitude de sortie de chaque intégrateur tout en gardant le SNR à l'intérieur d'une plage acceptable pour l'application ciblée. Pour cela, il faut balancer adéquatement la contribution de chaque intégrateur à la puissance de bruit totale du modulateur.

3.4.1 Définition de la fenêtre

La première étape consiste à déterminer une fenêtre d'opération telle qu'illustrée à la figure 3.5. Cette fenêtre délimite la région à l'intérieur de laquelle le processus d'optimisation aura lieu. La fenêtre a deux dimensions, l'amplitude de sortie des intégrateurs et le SNR. Les limites de chacune de ces dimensions doivent être correctement déterminées afin d'assurer l'efficacité de la méthode.

La première dimension à fixer est l'amplitude de sortie des intégrateurs. Cette dimension spécifie la plage à l'intérieur de laquelle l'amplitude de sortie des intégrateurs doit être ajustée. La limite supérieure peut être fixée à la valeur pleine échelle du quantificateur, mais pour évi-

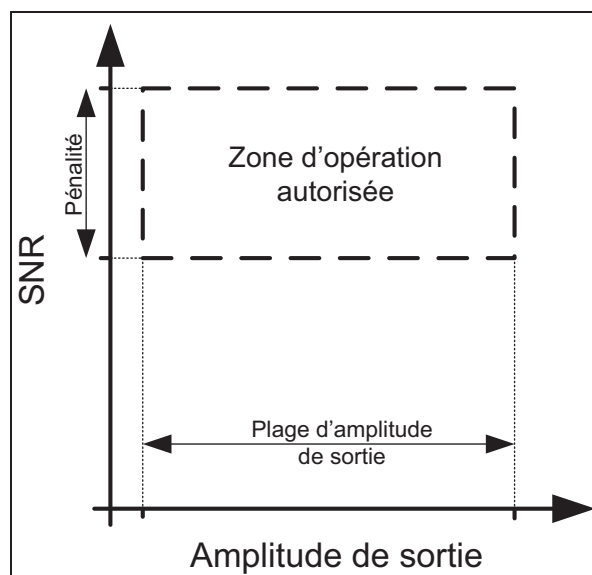


Figure 3.5 Fenêtre d'opération pour la méthode des fenêtres

ter des problèmes de saturation, elle peut également être fixée à un niveau plus bas. La limite inférieure de cette dimension peut aller jusqu'à 0. Cependant, lorsque la limite est trop basse, deux problèmes peuvent survenir. Premièrement, les réductions en amplitude ont tendance à se concentrer sur un nombre restreint d'intégrateurs. Deuxièmement, plus l'amplitude de sortie d'un intégrateur est réduite, plus il faut augmenter son coefficient interétagé. Donc, si les réductions en amplitude sont trop agressives, il peut en résulter des coefficients difficilement réalisables par un ratio de condensateurs en raison des tolérances de fabrication. Pour ces deux raisons, il vaut mieux choisir une valeur normalisée proche de 0,1 comme amplitude de sortie minimale. Le choix de cette valeur est étayé par des résultats de simulations qui seront présentés à la section 3.4.3.

La seconde dimension permet d'établir la pénalité en SNR acceptable. La limite inférieure est le SNR minimal désiré. Cette valeur est déterminée par l'application ciblée. Elle peut être majorée d'une certaine marge de sécurité afin de pallier les non-idéalités d'un circuit réel. La limite supérieure est le SNR obtenu lorsque l'amplitude de sortie de tous les intégrateurs est ajustée à leur valeur maximale (tel que spécifié par la première dimension). Cette valeur est obtenue en simulant le système à l'aide d'un modèle Simulink qui inclut les sources de

bruit thermique des intégrateurs. La différence entre ces deux valeurs de SNR est la pénalité acceptable.

3.4.2 Séquence de minimisation

Lorsqu'une fenêtre d'opération est correctement définie, la recherche de la solution optimale peut commencer. La méthode est un processus itératif s'appliquant individuellement sur chaque intégrateur. L'objectif est de déterminer tour à tour à quelle amplitude de sortie chaque intégrateur doit être ajusté. Il est à noter que tout au long de la méthode exposée ci-dessous, la fonction *scaleABCD* de la librairie Matlab *Delsig Toolbox* est utilisée pour mettre à l'échelle le modulateur afin d'obtenir les amplitudes de sortie désirées pour chaque intégrateur. La séquence de minimisation est résumée à la figure 3.6.

Au départ, tous les intégrateurs sont ajustés à l'amplitude de sortie maximale autorisée par la fenêtre d'opération. Graphiquement, ce point de départ peut être localisé en haut à droite de la fenêtre. Le SNR est alors maximal.

L'étape suivante consiste à diminuer graduellement l'amplitude de sortie d'un intégrateur alors que l'amplitude de sortie des autres intégrateurs reste constante. L'ordre dans lequel les intégrateurs sont traités est l'ordre inverse de l'importance de leur contribution au bruit à la sortie du convertisseur. En fonction de la figure 3.4, le quatrième intégrateur est celui ayant la plus faible influence sur le SNR. Il sera donc le premier à être mis à l'échelle. À l'opposé, le premier intégrateur étant celui ayant le plus grand impact, il sera mis à l'échelle en dernier.

Après chaque baisse d'amplitude, le système est resimulé afin de connaître le SNR obtenu par cette nouvelle configuration. Le processus se répète jusqu'à ce qu'une des limites de la fenêtre soit atteinte. Deux possibilités sont donc possibles, soit l'amplitude de sortie minimale est atteinte, soit le SNR minimal est atteint. Dans ce deuxième cas, l'amplitude est rehaussée à la valeur de l'itération précédente. Cela permet de conserver une marge d'optimisation pour les intégrateurs suivants. L'algorithme passe alors au prochain intégrateur jusqu'à ce que l'ensemble du modulateur ait été traité.

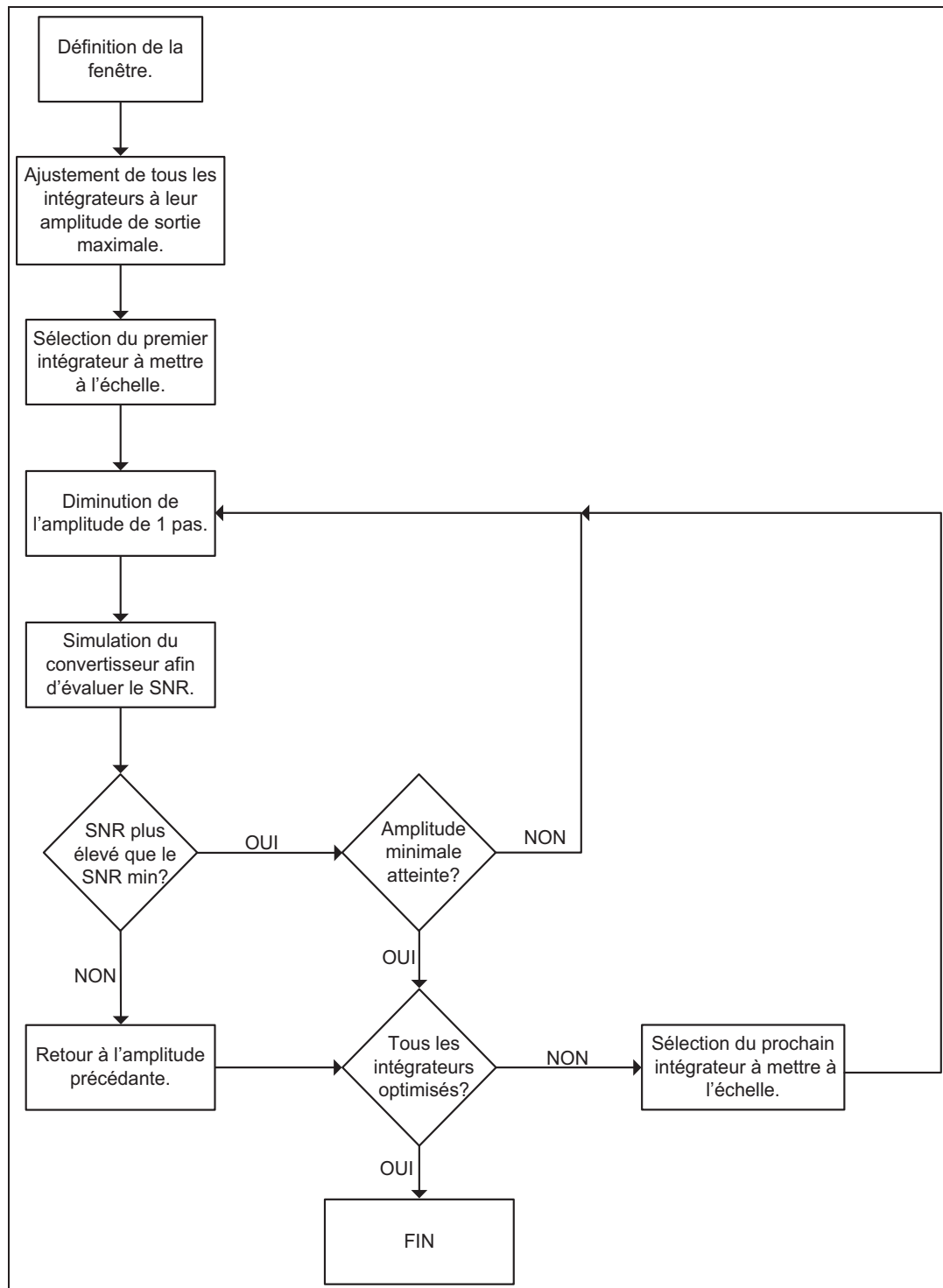


Figure 3.6 Séquence de minimisation de la méthode des fenêtres

Avec cette technique, l'amplitude de sortie de chaque intégrateur est diminuée jusqu'à ce que son bruit thermique soit à la limite de devenir une source de bruit significative. De cette façon, l'amplitude de sortie des intégrateurs bénéficiant de la mise en forme de leur bruit thermique est réduite de façon plus agressive. Cela favorise alors une diminution de la consommation de puissance du modulateur.

3.4.3 Exemple d'utilisation

Afin d'illustrer la méthode des fenêtres, celle-ci est appliquée au convertisseur de la figure 3.1. Le modulateur est conçu pour un OSR de 64 et un quantificateur de 1 bit. La NTF est obtenue à l'aide de la fonction *synthesizeNTF* de la librairie Matlab *Delsig Toolbox*. La valeur maximale de $|NTF(e^{j\omega})|$ est fixée à 1,4.

Avec un OSR de 64, les coefficients g des résonateurs sont 0.0003 et 0.0018. Ces valeurs conduisent à des ratios de condensateurs difficilement réalisables. Pour cette raison, les coefficients des résonateurs g ont été forcés à 0 dans cet exemple. En d'autres mots, tous les zéros de la NTF sont forcés à DC. Les coefficients d'entrées b_i sont calculés de façon à obtenir une $STF = 1$.

L'amplitude du signal d'entrée est -5 dBFS et sa fréquence normalisée est 0,001. Dans cet exemple, toutes les valeurs d'amplitude sont normalisées à la valeur pleine échelle du quantificateur qui est 0,8 V. La puissance du bruit thermique de chaque intégrateur est fixée à $3,16 \cdot 10^{-9} \text{ V}^2$. Cette valeur correspond à la puissance de bruit générée par un intégrateur ayant un condensateur d'échantillonnage de 1,3 pF tel que défini par l'équation 2.5.

Afin de mettre en évidence l'impact des limites de la fenêtre d'opération, la méthode sera réitérée pour différents ensembles de paramètres. Tout d'abord, trois amplitudes de sortie minimales pour les intégrateurs seront utilisées : 0, 0,1 et 0,5. Avec une amplitude minimale de 0, le seul critère d'arrêt utilisé est le SNR, lorsqu'il est inférieur au minimum requis. Ensuite, la méthode est également répétée pour trois pénalités en SNR différentes : 0,5, 1 et 2 dB. Il y a donc un total de 9 ensembles de paramètres qui pourront être comparés.

L'amplitude de sortie maximale sera fixe à 0,9 pour toutes les configurations. Avec les puissances de bruit spécifiées précédemment, le SNR maximal obtenu est 78,5 dB.

L'algorithme est effectué avec des pas de 0,01. Les résultats obtenus sont présentés dans le tableau 3.1. La moyenne de l'amplitude de sortie des quatre intégrateurs est inscrite à l'avant-dernière ligne du tableau tandis que le SNR obtenu avec cette configuration est inscrit à la dernière ligne du tableau.

La première constatation faite en analysant les résultats est que ceux-ci sont en conformité avec la figure 3.3. En effet, plus la pénalité en SNR est élevée, plus la moyenne des amplitudes est réduite. Au premier abord, la méthode fait donc son travail d'établir un compromis entre les performances en SNR et les réductions en amplitude.

Cependant, l'ampleur des réductions n'est pas la même selon l'amplitude de sortie minimale. Pour une amplitude minimale de 0, lorsque la pénalité en SNR passe de 0,5 à 2 dB, la moyenne des amplitudes de sortie passe de 0,52 à 0,48. Par contre, pour une amplitude minimale de 0,1, lorsque la pénalité en SNR passe de 0,5 à 2 dB, la moyenne des amplitudes de sortie passe 0,52 à 0,39.

Tableau 3.1 Résultats obtenus par la méthode des fenêtres

	Amplitude de sortie minimale des intégrateurs								
	0,0			0,1			0,5		
Pénalité en SNR	0,5 dB	1 dB	2 dB	0,5 dB	1 dB	2 dB	0,5 dB	1 dB	2 dB
Intégrateur 1	0,90	0,90	0,90	0,90	0,90	0,90	0,90	0,86	0,72
Intégrateur 2	0,89	0,85	0,89	0,90	0,76	0,48	0,58	0,50	0,50
Intégrateur 3	0,26	0,17	0,11	0,18	0,13	0,10	0,50	0,50	0,50
Intégrateur 4	0,04	0,03	0,02	0,10	0,10	0,10	0,50	0,50	0,50
Moyenne	0,52	0,49	0,48	0,52	0,47	0,39	0,62	0,59	0,55
SNR	78,0	77,5	76,0	78,0	77,5	76,0	78,0	77,5	76,0

Cela est dû au fait que lorsque l'amplitude de sortie minimale est 0, les réductions en amplitude se concentrent dans les intégrateurs 3 et 4. En augmentant l'amplitude de sortie minimale à 0,1, on favorise une plus grande réduction de l'amplitude de sortie du troisième intégrateur. Il en résulte alors une moyenne plus basse.

Lorsque l'amplitude de sortie minimale est fixée à 0,5, les moyennes deviennent plus élevées qu'avec une amplitude minimale de 0. Cela démontre que le choix de l'amplitude minimale influence les performances de la méthode. Afin de déterminer l'amplitude de sortie minimale optimale, la figure 3.7 trace la moyenne des amplitudes de sortie obtenue en fonction de l'amplitude minimale pour une pénalité en SNR de 2 dB. Les moyennes les plus basses sont obtenues lorsque l'amplitude minimale est située entre 0,1 et 0,15.

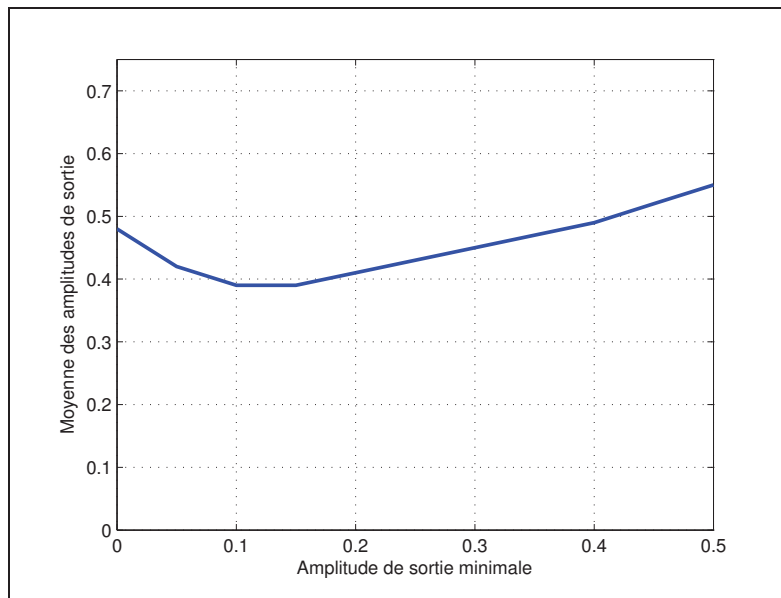


Figure 3.7 Moyenne des amplitudes de sortie en fonction de l'amplitude minimale de sortie. La pénalité en SNR est de 2 dB

Dans un second temps, on constate également que l'importance des réductions en amplitude dépend de la position de l'intégrateur dans le modulateur. En raison de son très grand impact sur le SNR, le premier intégrateur conserve toujours l'amplitude de sortie la plus élevée. À l'inverse, le quatrième intégrateur converge toujours vers l'amplitude de sortie la plus faible.

Ceci est en raison de son trop faible impact sur le SNR. Ces constatations sont en accord avec les courbes de la figure 3.4. Plus un intégrateur est situé loin dans le modulateur, plus son amplitude de sortie peut être réduite avant d’avoir un impact significatif sur le SNR.

Ces résultats de simulation ont permis de quantifier les réductions en amplitude possibles en fonction des paramètres de la fenêtre d’opération. Pour un modulateur du quatrième ordre, il est possible de réduire de 90% l’amplitude de sortie des deux derniers intégrateurs et de 50% l’amplitude de sortie du deuxième en échange d’une pénalité de 2 dB en SNR.

3.5 Carences de la méthode des fenêtres

Bien qu’intéressante pour des raisons de simplicité, la méthode exposée précédemment présente certaines lacunes. Tout d’abord, la méthode repose sur la sélection de plusieurs paramètres de façon empirique afin de définir une fenêtre d’opération. L’absence de règles précises sur la façon optimale de définir cette fenêtre rend la méthode difficilement exploitable par des concepteurs peu expérimentés.

Également, la méthode visant à réduire l’amplitude de sortie des différents intégrateurs en fonction de leur contribution respective au bruit total à la sortie du convertisseur, chaque intégrateur est traité individuellement au lieu de tenter d’optimiser l’ensemble du modulateur simultanément.

Au travers du processus, la puissance du bruit thermique est prise en compte de façon très sommaire. Une valeur est fixée au départ pour chaque étage d’intégration, mais n’est pas réajustée subséquemment durant le processus d’optimisation. Or, lorsque les coefficients de la structure varient, la puissance du bruit varie également (voir section 2.2.2.2). Ces variations devraient être prises en compte afin d’augmenter la précision de la méthode.

Finalement, la simulation d’un système comprenant des signaux aléatoires (les sources de bruit thermique) impose certaines variations dans les résultats à chaque itération. Il y a alors un danger potentiel que le processus d’optimisation converge vers un résultat qui ne soit pas optimal. Une méthode plus rigoureuse devrait être basée sur des calculs théoriques. Non seulement cela

augmenterait la précision des résultats, mais cela augmenterait également la vitesse d'exécution du processus.

3.6 Minimisation multicritères

En se basant sur les lacunes présentées à la section précédente, une nouvelle méthode d'optimisation est développée. Celle-ci ajoute la taille des condensateurs du modulateur dans la fonction objectif. Tel qu'expliqué à la section 2.2.2.2, la puissance du bruit thermique d'un intégrateur est déterminée par la taille des condensateurs d'échantillonnage. Donc, en optimisant la taille des condensateurs du modulateur, la puissance du bruit thermique de chaque étage d'intégration pourra être optimisée en fonction de son impact sur le SNR à la sortie du convertisseur. La taille des condensateurs a également un impact majeur sur la taille du circuit (voir section 2.2.4). De plus, la réduction de l'amplitude des signaux et la réduction de la taille des condensateurs ont toutes deux un impact sur la consommation en puissance [27, 28, 34].

3.6.1 Variables d'optimisation

La première étape pour développer la nouvelle méthode d'optimisation est d'établir la fonction objectif. Pour cela, il faut déterminer les variables d'optimisation et les paramètres devant être minimisés.

La taille des condensateurs d'intégration C_2 et les coefficients interétages c sont les variables devant être optimisées. Étant donné que chaque étage du modulateur a un condensateur d'intégration et un coefficient interétage, un modulateur d'ordre N aura $2N$ variables d'optimisation. Ces variables serviront à minimiser l'amplitude de sortie des intégrateurs ainsi que la somme des capacités du modulateur.

La fonction multivariables suivante est alors obtenue :

$$f(\mathbf{c}, \mathbf{C}_2) = \sum X_n + \beta C_T \quad (3.6)$$

où $\mathbf{c}$ est un vecteur contenant les valeurs des coefficients interétages, $\mathbf{C}_2$ est un vecteur contenant les valeurs des condensateurs d'intégration, X_n est l'amplitude de sortie normalisée des intégrateurs et C_T est la somme de toutes les capacités du modulateur (en pF).

Le paramètre d'optimisation β est fixé par le concepteur afin d'établir la priorité du processus. Une augmentation du paramètre d'optimisation β favorise une réduction accrue des capacités du circuit. À l'inverse, sa diminution favorise une réduction de l'amplitude de sortie des intégrateurs. À la section 3.6.4, un exemple d'utilisation de l'algorithme mettra en relief l'impact du paramètre β .

Une condition nécessaire à la réalisation de l'algorithme est d'être en mesure de calculer la valeur des paramètres à minimiser en fonction des variables d'optimisation.

Pour obtenir la somme des capacités du modulateur, le calcul se fait en trois étapes. Premièrement, la valeur des différents coefficients de la structure est calculée afin de réaliser la NTF et la STF désirées avec les coefficients interétages $\mathbf{c}$ spécifiés. Ensuite, en fonction de la valeur des condensateurs d'intégration $\mathbf{C}_2$, la valeur des condensateurs d'échantillonnage est déterminée. Il ne reste alors plus qu'à faire la somme de toutes ces valeurs.

L'amplitude de sortie maximale des intégrateurs est inversement proportionnelle aux coefficients interétages du modulateur. Elle peut être approximée par :

$$X_k \approx \frac{X'_k}{\prod_{i=k}^N c_i} \quad (3.7)$$

où X_k représente l'amplitude maximale à la sortie du k^e intégrateur, N est l'ordre du modulateur et X'_k est l'amplitude de sortie maximale lorsque les coefficients interétages c_i sont mis à 1. La valeur de X'_k est trouvée à l'aide de simulations extensives [12]. De cette façon, il n'est pas nécessaire de resimuler le système à chaque itération de l'algorithme. Il est simulé une seule fois au départ afin de connaître la dynamique du système.

3.6.2 Contraintes d'optimisation

Afin de s'assurer que le processus de réalisation converge vers une solution viable, certaines contraintes doivent être imposées.

L'algorithme d'optimisation tente de minimiser la taille des condensateurs et l'amplitude de sortie des intégrateurs. Or, la diminution de chacune de ces grandeurs entraîne également une diminution du SNR. Pour cette raison, une contrainte de SNR minimal doit être imposée au processus. La valeur du SNR minimal est dictée par l'application ciblée.

Le SNR est calculé de façon théorique en fonction des variables d'optimisation à chaque itération. Tout d'abord, les coefficients de la structure et la taille des différents condensateurs sont calculés en fonction des variables d'optimisation. Ceci permet ensuite de déterminer la puissance des sources de bruit thermique ainsi que leur fonction de transfert. La puissance de bruit à l'intérieur de la bande de fréquences du signal à la sortie du convertisseur peut alors être calculée (en y ajoutant également le bruit de quantification). Le SNR est finalement obtenu en fonction de la puissance du signal d'entrée.

Les condensateurs d'intégration C_2 ainsi que les coefficients interétages c sont soumis à des contraintes sur leurs valeurs minimales et maximales. Ces valeurs sont choisies en fonction des limites de la technologie cible. Cela permet de s'assurer que les résultats obtenus soient réalisables.

Finalement, l'amplitude de sortie des intégrateurs est soumise à une contrainte sur la valeur maximale. Cela permet de s'assurer que les résultats obtenus respecteront la plage dynamique des amplificateurs. Il est à noter qu'il n'y a pas de contrainte sur l'amplitude de sortie minimale.

L'ensemble du processus d'optimisation est réalisé dans Matlab à l'aide de l'*Optimization Toolbox* [37]. La fonction d'optimisation *fmincon* est utilisée. Cette fonction offre le choix entre quatre algorithmes différents : *trust-region-reflective*, *active-set*, *interior-point* et *sqp*. De ces quatre possibilités, seul l'algorithme *trust-region-reflective* est non recommandé pour ce type de problème. Les trois autres ont été testés pour diverses configurations. Dans tous les cas,

les résultats obtenus étaient identiques peu importe l'algorithme sélectionné. Il a finalement été décidé arbitrairement d'utiliser l'algorithme *interior-point*.

Le modèle d'optimisation complet est donné par :

$$\min f(\mathbf{c}, \mathbf{C}_2) \text{ sujet à } \begin{cases} SNR(\mathbf{c}, \mathbf{C}_2) \geq SNR_{min} \\ \mathbf{C}_{2_{min}} \leq \mathbf{C}_2 \leq \mathbf{C}_{2_{max}} \\ \mathbf{c}_{min} \leq \mathbf{c} \leq \mathbf{c}_{max} \\ X_n \leq X_{n_{max}} \end{cases}$$

3.6.3 Définition du point de départ

Pour initialiser le processus d'optimisation, un point de départ doit être sélectionné. Dans cette méthode, il a été choisi d'utiliser comme point de départ les paramètres obtenus par un processus de conception traditionnel [1].

Les valeurs initiales pour les coefficients interétages $\mathbf{c}$ sont obtenues par la fonction *scaleABCD* de la librairie Matlab Delsig Toolbox. Comme il a été expliquée à la section 2.3.1, cette fonction permet d'ajuster l'amplitude de sortie des intégrateurs afin qu'elle reste à l'intérieur d'une certaine plage. Dans ce cas-ci, la plage sélectionnée est 90% de la valeur pleine échelle du quantificateur.

Les valeurs initiales des condensateurs d'intégration $\mathbf{C}_2$ sont calculées en fonction de la cible en SNR de l'application. Pour que celle-ci soit atteinte, la puissance totale du bruit à l'intérieur de la bande de fréquences du signal à la sortie du convertisseur doit être :

$$P_{NT} \leq \left(\frac{A_{fs}}{2\sqrt{2}} \right)^2 \times 10^{-SNR_t/10} \quad (3.8)$$

où A_{fs} représente la valeur crête à crête d'un signal sinusoïdal pleine échelle et SNR_t représente la cible en SNR de l'application.

Généralement [1], les sources de bruit dominantes sont le bruit de quantification ainsi que le bruit thermique du premier étage. Les autres sources de bruit sont considérées en allouant une marge dans la cible en SNR.

Pour connaître la puissance du bruit thermique allouée au premier intégrateur (P_{N1}), il faut soustraire au résultat de (3.8) la puissance du bruit de quantification (P_Q) :

$$P_{N1} = P_{NT} - P_Q \int_0^{f_b} \left| NTF(e^{j2\pi f}) \right|^2 df \quad (3.9)$$

où f_b représente la largeur de la bande de fréquences du signal. Il est à noter que P_{N1} représente la puissance du bruit thermique du premier intégrateur dans la bande de fréquences du signal à la sortie du convertisseur. La puissance de la source de bruit à l'entrée du premier intégrateur correspondant à P_{N1} est donné par :

$$P'_{N1} = \frac{P_{N1}}{\int_0^{f_b} |H_{n1}(e^{j2\pi f})|^2 df} \quad (3.10)$$

où H_{n1} est la fonction de transfert du bruit thermique du premier intégrateur.

La valeur du condensateur d'intégration du premier étage peut alors être calculée de façon à respecter le budget de bruit alloué :

$$C_2 = \frac{kT}{OSR \cdot P'_{N1}} \left(\frac{1}{a_1} + \frac{1}{b_1} + \frac{1}{g_1} \right) \quad (3.11)$$

où $a_1 = C_{11}/C_2$, $b_1 = C_{12}/C_2$ et $g_1 = C_{13}/C_2$. Les condensateurs d'intégration des étages suivants utilisent la même valeur.

3.6.4 Exemple d'utilisation

Les performances de cette méthode d'optimisation ont été évaluées dans Matlab à l'aide d'un modèle Simulink. Pour faciliter les comparaisons, les mêmes paramètres systèmes que pour

l'exemple d'utilisation de la méthode des fenêtres sont utilisés. Les résultats obtenus par cet exemple seront validés par des simulations au niveau transistor au chapitre suivant.

3.6.4.1 Définition des paramètres

La structure sélectionnée pour tester la méthode est le modulateur passe-bas d'ordre 4 présenté à la figure 3.1. Le système est conçu pour un OSR de 64 et un quantificateur de 1 bit. La cible en SNR est fixée à 78 dB pour un signal d'entrée de -5 dBFS. Le choix de cette configuration s'est fait de façon arbitraire dans l'unique but de démontrer la validité de la méthode d'optimisation. Aucune application précise n'est visée par ces spécifications. Comme pour la méthode des fenêtres, les coefficients des résonateurs g sont forcés à 0.

Tel qu'expliqué à la section 3.6.2, certaines contraintes doivent être définies afin de configurer l'algorithme. Tout d'abord, l'amplitude de sortie maximale des intégrateurs est fixée à 0,9. Cette valeur est normalisée à l'amplitude d'entrée pleine échelle du quantificateur qui est 0,8V. Cette valeur, en dessous de la valeur pleine échelle, a été choisie afin de conserver une marge de manoeuvre à l'intérieur de la plage linéaire des amplificateurs qui est également 0,8V.

Les valeurs minimales et maximales pour les condensateurs d'intégration C_2 sont respectivement 0,1 pF et 10 pF. Les valeurs minimales et maximales pour les coefficients interétages c sont respectivement 0,05 et 20.

Les valeurs pour le point de départ ont été obtenues à l'aide de la méthodologie de conception traditionnelle expliquée à la section 3.6.3. Elles sont rapportées dans le tableau 3.3 dans la colonne références (Réf.). En effet, le point de départ est également utilisé comme valeur de référence pour comparer les résultats obtenus par le processus d'optimisation.

Finalement, afin de mettre en évidence l'impact du paramètre d'optimisation β , la méthode est répétée pour deux valeurs différentes : 0,05 et 0,01.

3.6.4.2 Résultats de simulation

Les valeurs obtenues par l'algorithme de minimisation sont rapportées dans les tableaux 3.2 et 3.3. Le premier tableau présente la somme des capacités du modulateur ainsi que la moyenne de l'amplitude de sortie des intégrateurs tandis que le second tableau présente les valeurs finales des variables d'optimisation (condensateurs d'intégration et coefficients interétages) et l'amplitude de sortie de chaque intégrateurs. La dernière ligne du tableau 3.3 présente la moyenne de chaque colonne. L'amplitude de sortie des intégrateurs est donnée en valeur normalisée.

Tableau 3.2 Somme des capacités du modulateur et moyenne de l'amplitude de sortie des intégrateurs pour les trois configurations

β	0,01	0,05	Réf.
Somme des capacité (pF)	42,80	29,34	51,65
Moyenne de l'amplitude de sortie des intégrateurs	0,62	0,70	0,90

Tableau 3.3 Valeurs finales des variable d'optimisation et amplitude de sortie de chaque intégrateur

	Condensateurs d'intégration (pF)			Amplitude de sortie des intégrateurs			Coefficients interétages		
β	0,01	0,05	Réf.	0,01	0,05	Réf.	0,01	0,05	Réf.
Intégrateur 1	9,66	8,00	4,41	0,90	0,90	0,90	0,077	0,085	0,085
Intégrateur 2	7,46	5,16	4,41	0,80	0,90	0,90	0,082	0,13	0,16
Intégrateur 3	4,64	1,47	4,41	0,43	0,63	0,90	0,19	0,17	0,24
Intégrateur 4	0,39	0,19	4,41	0,33	0,38	0,90	5,79	3,42	2,15
Moyenne	5,53	3,70	4,41	0,62	0,70	0,90	1,53	0,95	0,66

Les résultats du tableau 3.2 permettent de conclure que l'algorithme d'optimisation à l'effet escompté. En effet, la somme des capacités ainsi que la moyenne de l'amplitude de sortie des intégrateurs sont plus faible pour les deux configurations optimisées que pour la configuration de référence. L'effet du paramètre β peut également être constaté. Avec $\beta=0,01$, l'amplitude de sortie des intégrateurs est réduite alors que la taille des condensateurs reste à des valeurs plus élevées. Avec $\beta=0,05$, l'effet est inversé.

L'effet de la réduction de l'amplitude de sortie des intégrateurs sur les coefficients interétages peut être remarqué en analysant les résultats du tableau 3.3. Plus la valeur moyenne des amplitudes est réduite, plus la valeur moyenne des coefficients interétages augmente.

Ces résultats permettent également de quantifier la réduction de la sensibilité aux non-idéalités des intégrateurs en fonction de leur position dans la chaîne. Le processus d'optimisation converge toujours vers une solution où l'amplitude de sortie du premier intégrateur est plus grande. À l'inverse, l'amplitude de sortie du dernier intégrateur est toujours la plus petite, car son impact sur le SNR est limité. Le même principe s'applique pour la taille des condensateurs.

Les trois configurations ont été simulées par un modèle Simulink afin de comparer leurs performances. La figure 3.8 présente les courbes du SNR en fonction de l'amplitude d'entrée. La valeur maximale du SNR ainsi que la plage dynamique sont inchangées (<1 dB) par la méthode d'optimisation. Les trois courbes restent à l'intérieur d'une marge de 2 dB sur toute la plage du signal d'entrée.

Finalement, les résultats obtenus par la méthode de minimisation multicritères peuvent être comparés à ceux obtenus par la méthode des fenêtres. Afin que la comparaison soit valable, il faut comparer deux configurations qui procurent le même SNR maximal. Pour la méthode des fenêtres, la configuration sélectionnée est celle obtenue avec une pénalité de 0,5 dB et une amplitude de sortie minimale de 0,1. Pour la méthode de minimisation multicritères, la configuration sélectionnée est celle obtenue avec $\beta=0,01$. Il s'agit de la configuration qui favorise la minimisation de l'amplitude de sortie des intégrateurs. Ces deux configurations génèrent un SNR maximal de 78 dB.

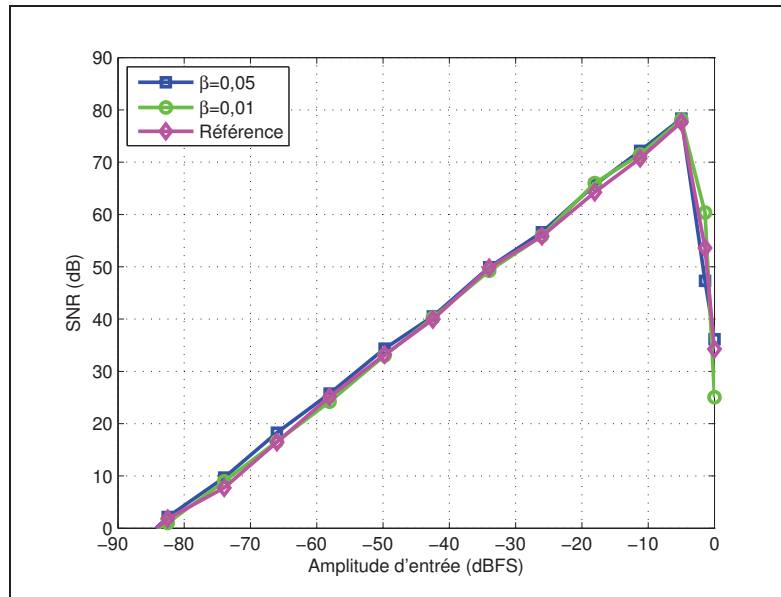


Figure 3.8 SNR vs. amplitude d'entrée pour les trois configurations présentées dans le tableau 3.3

De prime abord, on constate qu'avec la méthode des fenêtres, l'amplitude de sortie des intégrateurs est davantage réduite. La valeur moyenne y est de 0,52 alors que pour la méthode de minimisation multicritères elle est de 0,62.

Cependant, la méthode de minimisation multicritères tente de minimiser à la fois l'amplitude de sortie des intégrateurs et la somme des capacités du modulateur. Pour compléter la comparaison, il faut donc regarder ces deux aspects.

Dans la méthode des fenêtres, une puissance de bruit thermique était allouée pour chaque étage d'intégration. En fonction des coefficients à réaliser dans le modulateur, il est possible de calculer la taille des condensateurs nécessaire afin de respecter cette puissance de bruit allouée. Dans le cas de la configuration qui nous intéresse, la somme des capacités du modulateur est 247 pF. Dans le cas de la méthode de minimisation multicritères, la somme des capacités du modulateur était 42,8 pF. La plus grande réduction de l'amplitude de sortie des intégrateurs par la méthode des fenêtres est donc obtenue au prix d'une augmentation significative de la somme des capacités du modulateur.

3.7 Conclusion

Dans ce chapitre, deux nouvelles méthodes d'optimisation des modulateurs $\Sigma\Delta$ ont été présentées. La première, appelée méthode des fenêtres, vise à établir un compromis entre la performance en SNR et la réduction de l'amplitude de sortie des intégrateurs. Chaque intégrateur est mis à l'échelle individuellement jusqu'à ce que l'amplitude de sortie minimale soit atteinte ou que son impact sur le SNR devienne significatif. Cette technique permet de réduire l'amplitude de sortie des intégrateurs en fonction d'une certaine pénalité en SNR.

Certaines failles ont été relevées sur la première méthode. Celle-ci nécessite de définir une fenêtre d'opération de façon peu intuitive et repose sur un cadre mathématique rudimentaire. Pour pallier à ces lacunes, une seconde méthode est proposée. Celle-ci est basée sur un processus d'optimisation multicritères. L'objectif de l'algorithme est de minimiser simultanément l'amplitude de sortie des intégrateurs et la somme des capacités du modulateur en fonction d'un objectif en SNR. Un paramètre utilisateur permet de prioriser la minimisation de l'une ou l'autre des ces quantités. Des résultats de simulations ont démontré que les performances des convertisseurs optimisés par la méthode sont en accord avec les performances d'un convertisseur issues d'un processus de conception standard.

CHAPITRE 4

VALIDATION PAR SIMULATIONS AU NIVEAU TRANSISTOR

4.1 Introduction

Au chapitre précédent, deux nouvelles méthodes d'optimisation des convertisseurs $\Sigma\Delta$ ont été présentées. Afin de valider la méthode d'optimisation multicritères, des simulations au niveau transistor sont effectuées. Tout d'abord, les paramètres de conception sont définis puis la conception de chaque section du circuit est exposée. Des simulations temporelles permettent de s'assurer du bon fonctionnement individuel de chaque section. Finalement, l'ensemble du circuit est simulé. Les résultats sont analysés et comparés aux résultats obtenus par les simulations au niveau système dans le chapitre précédent.

4.2 Paramètres de réalisation

Avant de débiter la conception à proprement dit du circuit, il est important de définir les différents paramètres de conception. L'objectif est de démontrer la validité de la méthode de minimisation multicritères. À la section 3.6.4, un exemple d'utilisation a été présenté. Pour des fins de comparaisons, les trois mêmes configurations de cet exemple seront utilisées pour les simulations au niveau transistor. Deux de ces configurations sont le résultat de la méthode d'optimisation pour deux valeurs différentes du paramètre d'optimisation β (0,05 et 0,01). La troisième configuration, désignée comme point de référence, est issue d'un processus de conception traditionnel. Certains paramètres ne sont pas déterminés durant le processus d'optimisation. Il s'agit de la technologie cible, la fréquence de l'horloge, le type de circuit et les tensions d'alimentation.

La technologie cible sélectionnée est CMOS 0,18 μm de TSMC. Normalement, en micro-électronique, la façon la plus efficace de démontrer la validité d'une nouvelle technique est la réalisation d'un circuit intégré. Malheureusement, en raison des délais requis pour la conception et la fabrication d'un circuit intégré, cette option était irréalisable dans l'échéancier de ce

travail de recherche. L'exercice n'est toutefois pas vain, car les simulations au niveau transistor sont tout de même plus proches de la réalité que les simulations au niveau système.

Pour la réalisation des différentes sections du circuit, un choix doit être fait entre l'utilisation de signaux entièrement différentiels ou de signaux uniques (*single ended*). Tel qu'énuméré à la section 2.2.3.1, les signaux entièrement différentiels présentent plusieurs avantages. Cependant, le cadre mathématique de la méthode de minimisation multicritères a été développé en considérant des signaux uniques. La méthode pourrait être adaptée à l'utilisation de signaux entièrement différentiels, mais par souci de clarté, il a été choisi d'utiliser des signaux uniques.

Pour la technologie cible, $0,18\ \mu\text{m}$, la tension d'alimentation nominale du coeur (V_{dd}) est fixé à 1,8 V. La plage d'entrée du quantificateur (0,8 V) est placée au centre de la plage d'alimentation du coeur. Les tensions de références (V_{ref_n} et V_{ref_p}) sont donc respectivement 0,5 V et 1,3 V. Un quatrième niveau de tension (V_{com}) est nécessaire pour fixer le point milieu de la plage d'entrée du quantificateur. Sa tension est 0,9 V. Ces quatre tensions d'alimentation sont générées par des sources de tensions idéales.

L'utilisation de la méthode de minimisation multicritères ne nécessite pas que la fréquence d'horloge visée soit déterminée. Cependant, lors de la conception du circuit au niveau transistor, il est important de fixer ce paramètre. Pour éviter qu'entrent en jeu des phénomènes transitoires liés aux circuits à très haute vitesse, il a été choisi de garder la conception simple en utilisant une fréquence d'horloge de 1 MHz.

Tel qu'expliqué à la section 2.2.1, un circuit à capacités commutées nécessite deux signaux d'horloges sans chevauchement nommés $Clk1$ et $Clk2$. Pour chacune de ces deux horloges, le signal inverse est également nécessaire ($Clk1N$ et $Clk2N$). Le circuit a donc un total de quatre signaux d'horloges. La figure 4.1 montre les signaux utilisés par le circuit. Il est à noter que dans cet ouvrage, une horloge est considérée comme étant la phase active lorsque sa version non inversée est au niveau haut. Comme pour les tensions d'alimentation, les quatre signaux d'horloges sont obtenus par des sources idéales. Encore une fois, cela simule qu'ils sont tous générés à l'extérieur du circuit intégré.

Pour chacune des trois configurations, les seuls paramètres qui sont modifiés sont la taille des condensateurs du modulateur réalisant les différents coefficients.

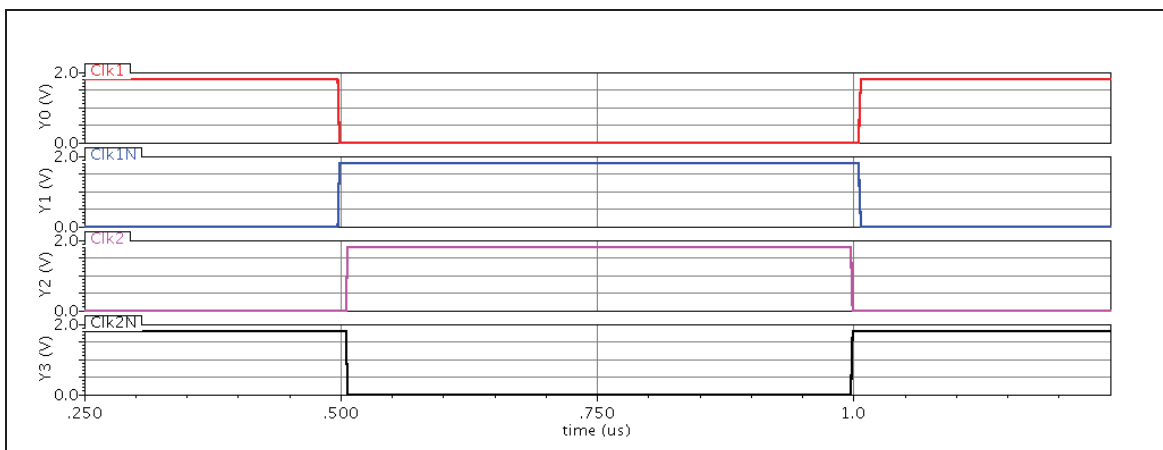


Figure 4.1 Signaux d'horloges sans chevauchement utilisés par le circuit

4.3 Conception du circuit

Une fois que les paramètres de réalisation sont fixés, la conception du circuit peut débuter. La figure 4.2 présente le schéma haut-niveau du convertisseur. Celui-ci peut-être divisé en quatre sous-blocs principaux : l'échantillonneur bloqueur, le modulateur, le quantificateur et le DAC. Cette section décrit le fonctionnement de chacun de ces sous-blocs. Des simulations temporelles sont utilisées afin de démontrer le fonctionnement adéquat de chacune de ces sections.

4.3.1 Amplificateur opérationnel

La conception d'un circuit intégré analogique débute généralement par la conception d'un amplificateur opérationnel répondant aux spécifications requises par le circuit. Celui-ci est par la suite utilisé pour la réalisation des différentes sections du circuit.

Tel que mentionné précédemment, le circuit n'est pas destiné à être fabriqué. Pour cette raison, il a été choisi d'utiliser un modèle AHDL (*Analog Hardware Descriptive Language*) pour réaliser les amplificateurs opérationnels. Un modèle AHDL permet de simuler le comportement d'un amplificateur opérationnel sans avoir à définir son circuit au niveau transistor. Cependant,

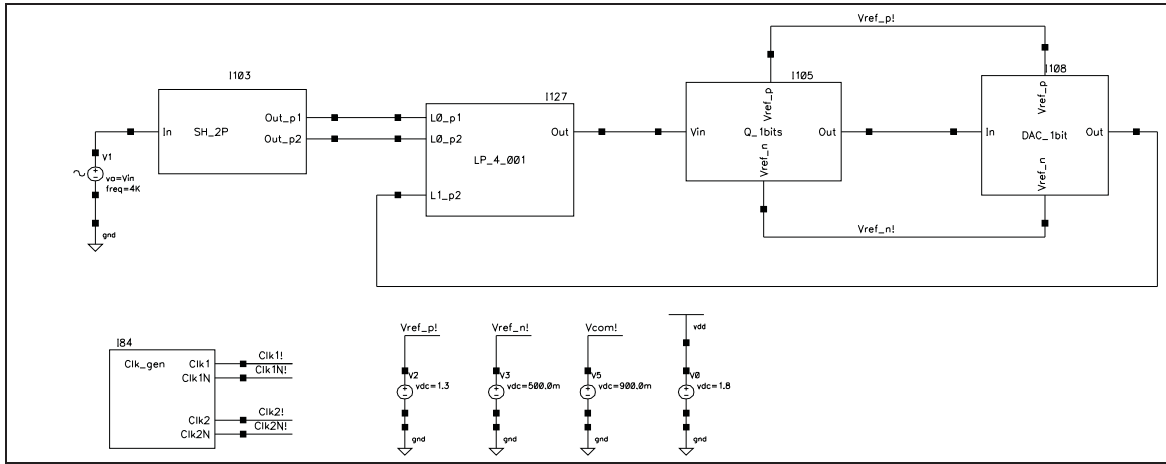


Figure 4.2 Schéma haut-niveau du convertisseur

un modèle AHDL ne génère pas de bruit contrairement à un circuit réel. Mais, tel qu'expliqué à la section 2.2.2.3, il peut être démontré que la source de bruit dominante dans un intégrateur à capacités commutées est le bruit thermique des interrupteurs [21].

Ce choix présente plusieurs avantages. Tout d'abord, le temps nécessaire à la conception du circuit est diminué de façon significative. De plus, les différentes caractéristiques de l'amplificateur opérationnel peuvent être ajustées de façon précise. De cette façon, les performances du convertisseur ne sont pas affectées par une mauvaise conception analogique. L'objectif de ces simulations est de valider la méthode de minimisation multicritères et non de valider l'efficacité d'une conception analogique. Finalement, l'utilisation d'un modèle AHDL diminue le nombre de transistors devant être considérés par l'outil de simulation ce qui accélère sa vitesse d'exécution. Par contre, comme nous le verrons à la section 4.4.3, l'utilisation d'un modèle AHDL empêche de mesurer la puissance consommée par le circuit.

Tous les amplificateurs opérationnels utilisés sont configurés avec les mêmes paramètres. Ceux-ci sont résumés dans le tableau 4.1.

4.3.2 Échantillonneur bloqueur

Tel qu'expliqué à la section 1.2, la première étape d'un ADC est de discrétiser en temps le signal à convertir. Le circuit doit saisir la tension du signal d'entrée analogique à un instant précis et la maintenir le temps de la quantifier. Cette étape est réalisée par le circuit échantillonneur bloqueur présenté à la figure 4.3.

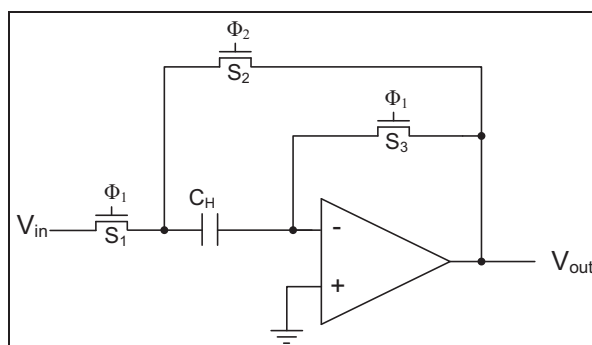


Figure 4.3 Échantillonneur bloqueur

Le circuit fonctionne en deux temps. Lorsque Φ_1 est actif, le condensateur C_H se charge à la tension d'entrée du circuit. Il s'agit de la phase d'échantillonnage. Durant cette phase, la tension de sortie de l'amplificateur opérationnel est maintenue à 0 V. Lorsque Φ_2 est actif, le condensateur C_H se retrouve connecté entre la sortie de l'amplificateur opérationnel et sont

Tableau 4.1 Paramètres des amplificateurs opérationnels utilisés dans le circuit

Caractéristiques	Valeurs
Gain DC (A_0)	1000
Fréquence unitaire (f_u)	100 MHz
Vitesse de balayage	3 V/ μ s
R_{in}	1 M Ω
R_{out}	80 Ω

entrée inversée. Durant cette phase, le circuit agit comme un bloqueur en maintenant la tension de C_H à la sortie du circuit.

Ce type de circuit échantillonneur bloqueur a comme avantage de diminuer les erreurs dues aux injections de charges [23]. Son principal inconvénient est sa vitesse plus lente. Cependant, étant donnée la fréquence du système (1 MHz), cette limitation n'entre pas en ligne de compte.

Pour fonctionner adéquatement, le modulateur a besoin du signal d'entrée sur les deux phases d'horloge. Pour cela, le circuit d'échantillonneur bloqueur est répété deux fois en cascade. L'unique différence entre les deux circuits est au niveau des signaux d'horloges qui sont interchangeables. De cette façon, le second circuit se trouve à échantillonner la sortie du premier. Il est à noter que la sortie du second circuit se retrouve déphasée d'un demi-coup d'horloge. Le circuit complet de l'échantillonneur bloqueur utilisé est présenté à la figure 4.4. La figure 4.5 présente une simulation temporelle du circuit ¹.

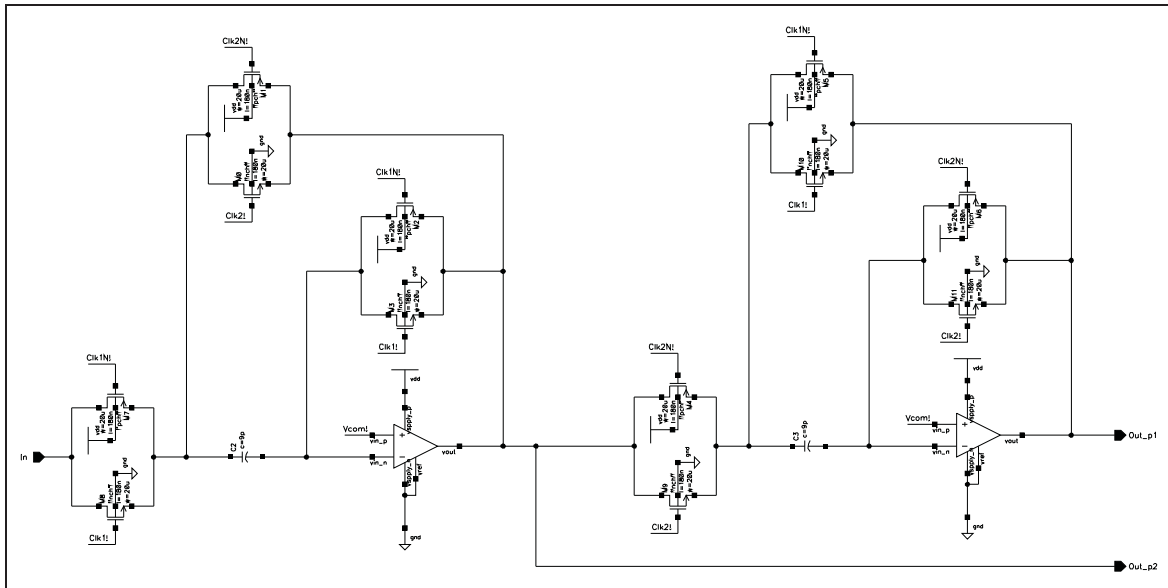


Figure 4.4 Schéma complet du circuit échantillonneur bloqueur

1. Dans ce chapitre, toutes les simulations temporelles sont réalisées avec le bruit thermique activé. La largeur de bande du bruit thermique est fixée à 20MHz, ce qui représente 20 fois la fréquence d'échantillonnage. Le niveau de précision du simulateur est réglé au plus élevé (*conservative*).

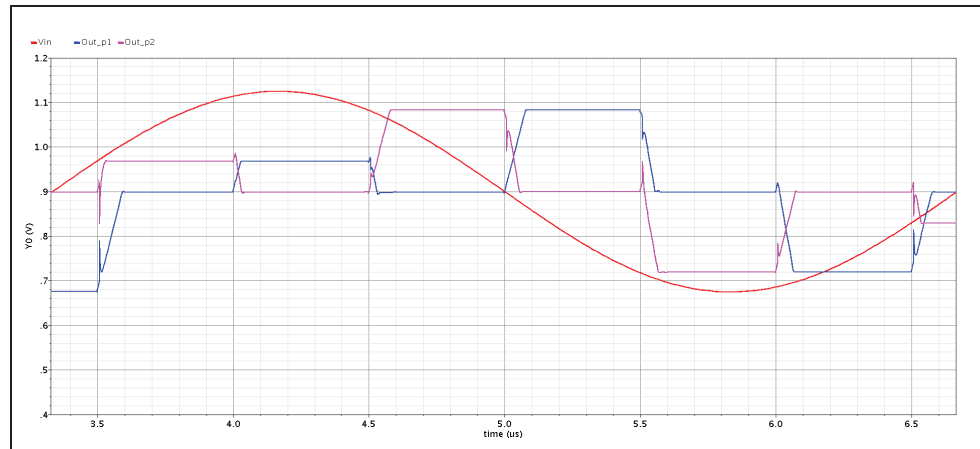


Figure 4.5 Simulation temporelle du circuit échantillonneur bloqueur

Peu importe l'efficacité du modulateur $\Sigma\Delta$, le SNR à la sortie du convertisseur ne peut pas être supérieur au SNR à la sortie de l'échantillonneur bloqueur. Il est donc important de s'assurer que l'échantillonneur bloqueur ne sera pas le noeud d'étranglement des performances du convertisseur. La figure 4.6 montre les spectres de fréquences des deux sorties du circuit de la figure 4.4. Le SNR à l'intérieur de la bande de fréquences du signal est 90,0 dB pour *Out_p1* et 94,5 dB pour *Out_p2*. La différence entre les deux SNR est due au bruit thermique. Comme les intégrateurs à capacités commutés, les échantillonneurs bloqueurs ajoutent du bruit thermique au signal. Étant donné que le signal *Out_p1* passe par deux échantillonneurs bloqueurs, son SNR est plus faible.

Dans les spectres de fréquences de la figure 4.6, la distorsion du signal cause l'apparition de plusieurs harmoniques. Afin d'éviter que celles-ci ne dégradent le SNR, la fréquence d'entrée du signal a été choisie de façon à ce que ces harmoniques se retrouvent à l'extérieur de la bande de fréquences du signal.

4.3.3 Quantificateur

Le circuit du quantificateur est illustré à la figure 4.7. Celui-ci peut être divisé en trois sections. Tout d'abord, un pont diviseur formé de deux résistances permet de déterminer les pas de

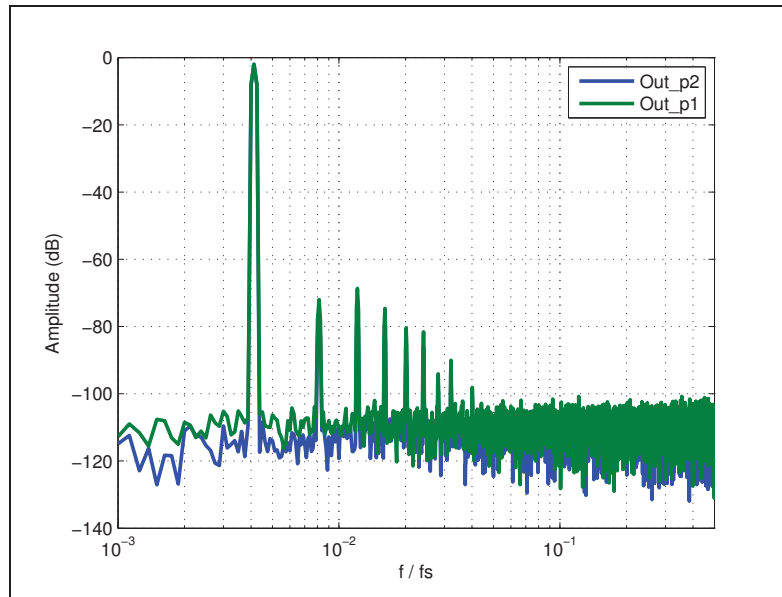


Figure 4.6 Spectre de fréquences des deux sorties de l'échantillonneur bloqueur de la figure 4.4

quantifications. Ce pont diviseur est alimenté par les deux tensions de références définissant la plage d'entrée du quantificateur.

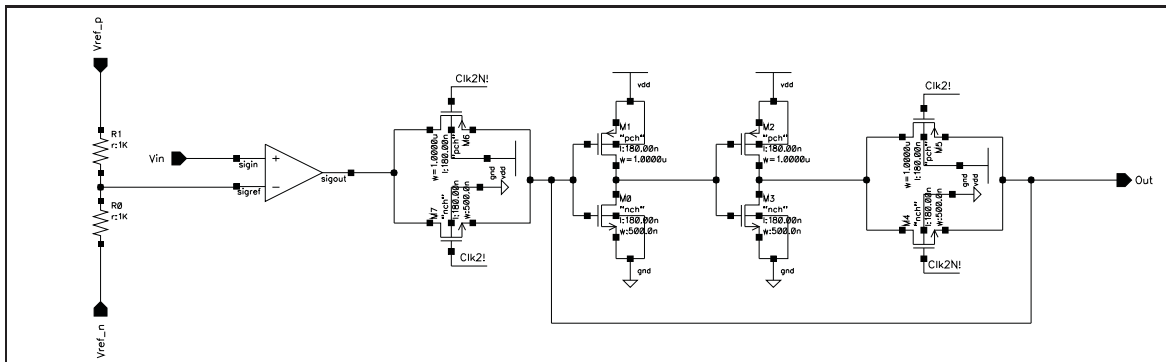


Figure 4.7 Circuit du quantificateur

Un comparateur détermine ensuite la sortie numérique du quantificateur. Comme pour les amplificateurs opérationnels de l'échantillonneur bloqueur et du modulateur, un modèle AHDL est utilisé pour réaliser le comparateur. Les niveaux de sortie du comparateur sont ceux de l'alimentation générale du circuit, c'est-à-dire 0 V et 1,8 V.

Finalement, les huit transistors qui suivent le quantificateur forment une bascule D. L'utilisation d'une bascule D à la sortie du comparateur évite que des soubresauts à la sortie du modulateur dû aux transitions de l'horloge ne se répercutent dans la sortie du convertisseur. La bascule acquiescence le signal lorsque *Clk2* est actif. Comme nous le verrons à la section 4.3.5, la sortie du modulateur est fixe lorsque *Clk2* est actif. La figure 4.8 montre une simulation temporelle du quantificateur.

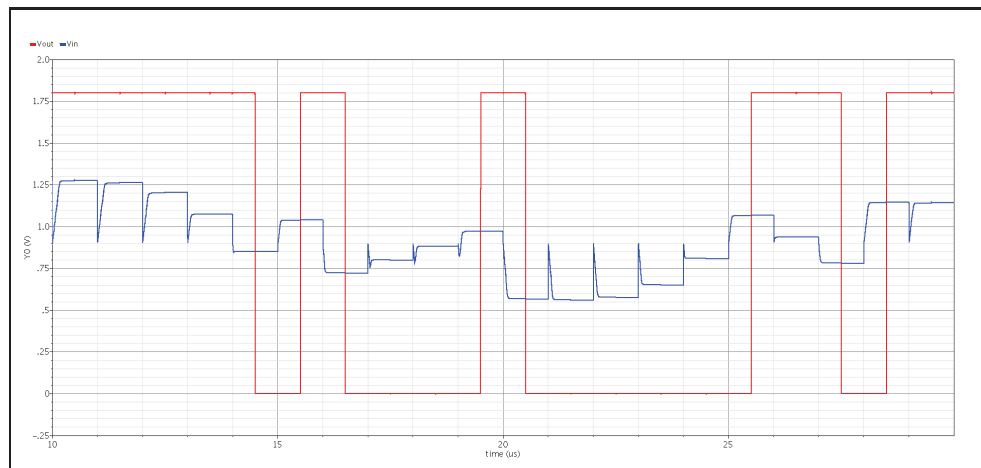


Figure 4.8 Simulation temporelle du circuit du quantificateur

4.3.4 DAC

Le convertisseur nécessite un DAC dans la branche de retour afin transformer la sortie numérique du quantificateur en un signal analogique pour le modulateur. Étant donné qu'une quantification binaire est utilisée, la conception du DAC est grandement simplifiée. De plus, tel qu'expliqué à la section 1.7.4, en utilisant une quantification binaire, les problèmes liés aux non-linéarités des DAC multibits sont évités.

Le circuit utilisé est présenté à la figure 4.9. Le rôle de ce circuit est de transformer les niveaux logiques à la sortie du quantificateur vers les niveaux des tensions de références V_{ref_p} et V_{ref_n} .

Il est à noter que le circuit du DAC a pour effet d'inverser le signal. Lorsque l'entrée est au niveau logique haut, la sortie est V_{ref_n} . À l'inverse, lorsque l'entrée est au niveau logique

BAS, la sortie est V_{ref_p} . Cela facilite la réalisation des coefficients de rétroaction a_i qui sont tous négatifs. La figure 4.10 montre une simulation temporelle du DAC.

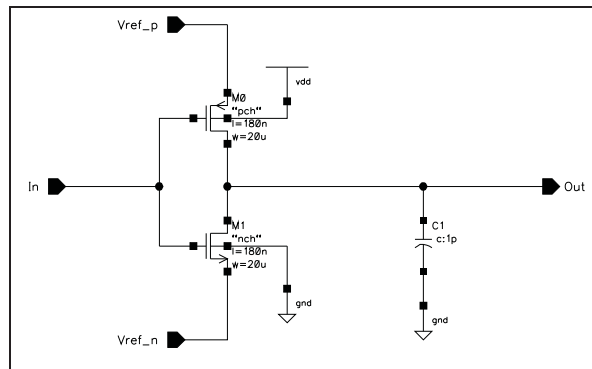


Figure 4.9 Circuit du DAC

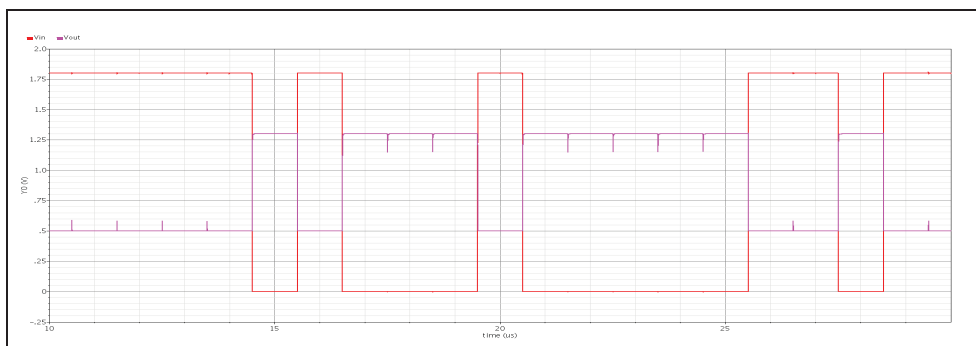


Figure 4.10 Simulation temporelle du circuit du DAC

4.3.5 Modulateur

Le modulateur est le coeur d'un convertisseur $\Sigma\Delta$. Pour sa réalisation, trois types de circuits sont nécessaires : intégrateur avec délai, intégrateur sans délai et additionneur. Chacun de ces trois types de circuits est d'abord conçu individuellement avant de les assembler pour former le modulateur en entier.

À la section 2.2.1, le fonctionnement d'un intégrateur avec délai a été exposé. Dans un modulateur $\Sigma\Delta$, il est généralement nécessaire d'alterner des intégrateurs avec et sans délai pour que celui-ci fonctionne adéquatement [1]. Un intégrateur sans délai utilise le même circuit que

celui présenté à la figure 2.1, mais les signaux d'horloges ne sont pas distribués de la même façon aux différents interrupteurs.

Durant la phase d'échantillonnage, les interrupteurs S_2 et S_4 sont activés. Dans cette configuration, le condensateur d'échantillonnage C_1 est court-circuité et se décharge complètement. Ensuite, durant la phase intégration, les interrupteurs S_1 et S_3 sont activés. Le condensateur d'échantillonnage C_1 se retrouve alors connecté entre l'entrée et l'entrée inverseur de l'amplificateur. Dans cette configuration, la tension d'entrée est soustraite de la tension de sortie. En d'autres mots, le circuit effectue l'opération suivante [38] :

$$V_{out_k} = V_{out_{k-1}} - \frac{C_1}{C_2} V_{in_k} \quad (4.1)$$

Le ratio entre les valeurs de C_1 et C_2 détermine le coefficient d'intégration. Ce circuit est parfois nommé intégrateur inverseur dû au fait que le signal de sortie est inversé par rapport à l'entrée.

La figure 4.11 présente une simulation temporelle pour les deux types d'intégrateurs (dans les deux cas, le coefficient d'intégration est 1). Le même signal d'entrée est appliqué aux deux intégrateurs. Ces résultats confirment que les deux circuits fonctionnent conformément à leurs équations caractéristiques.

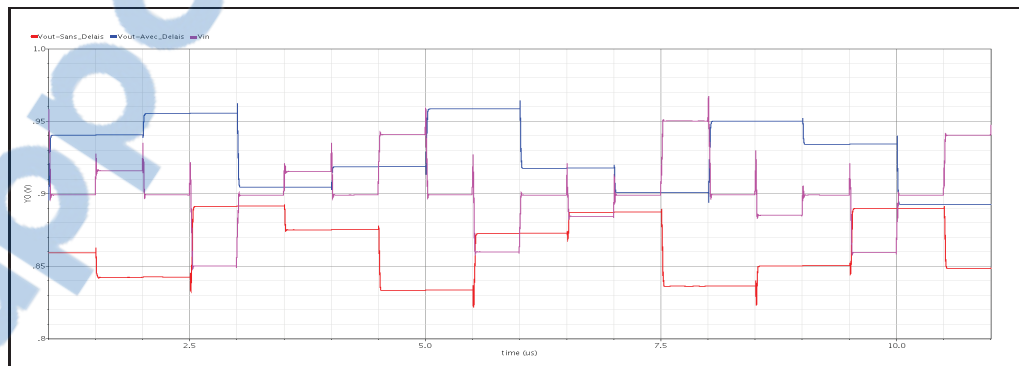


Figure 4.11 Simulation temporelle des deux types d'intégrateurs : sans délai et avec délai. Le même signal d'entrée est appliqué aux deux intégrateurs

La figure 4.12 montre le circuit d'un additionneur. Son fonctionnement est similaire à l'intégrateur sans délai à la différence que durant la phase d'échantillonnage (Φ_1), l'interrupteur S_5

vient décharger complètement le condensateur d'intégration C_2 . Le circuit se trouve donc à effectuer l'opération suivante :

$$V_{out_k} = - \left(\frac{C_{1a}}{C_2} V_a + \frac{C_{1b}}{C_2} V_b \right) \quad (4.2)$$

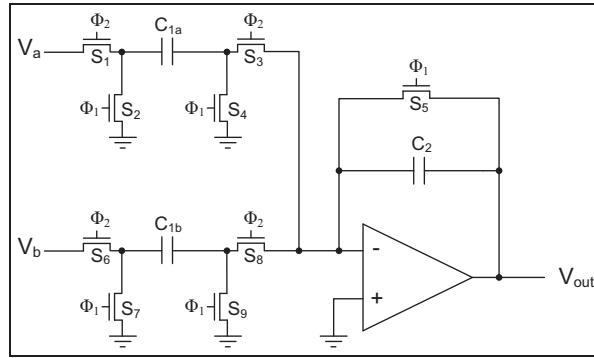


Figure 4.12 Circuit d'un additionneur

Un circuit additionneur est nécessaire pour réaliser l'addition finale des coefficient c_4 et b_5 , tel qu'illustré à la figure 3.1. La figure 4.13 présente une simulation temporelle pour le circuit additionneur. Pour les deux signaux, le coefficient d'addition (C_{1i}/C_2) est 1.

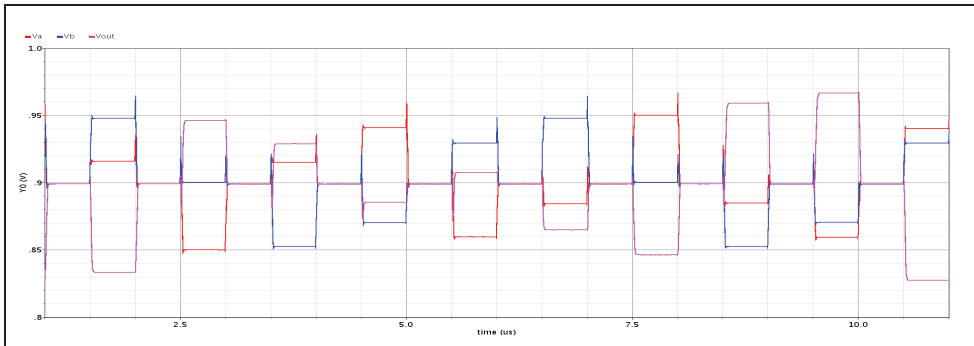


Figure 4.13 Simulation temporelle du circuit additionneur

Le circuit complet du modulateur est présenté à la figure 4.14. La valeur des condensateurs inscrits sur cette figure est celle issue du processus d'optimisation avec $\beta=0,01$. Le sous-bloc *Switch_Int* contient la configuration d'interrupteurs nécessaires pour réaliser un intégrateur avec délai tandis que le sous-bloc *Switch_Int_ND* contient la configuration d'interrupteurs né-

cessaires pour réaliser un intégrateur sans délai. Leurs circuits sont présentés respectivement aux figures 4.15 et 4.16.

Le circuit de la figure 4.14 est conçu pour réaliser le modulateur de la figure 3.1. Cependant, la conception doit également tenir compte du fait que les intégrateurs sans délai inversent le signal. C'est pour cette raison que dans le troisième intégrateur, les coefficients d'entrée b_3 et de rétroaction a_3 utilisent une configuration d'interrupteurs avec délai tandis que le coefficient d'intégration c_2 utilise une configuration d'interrupteurs sans délai. Il peut également être noté que les quatre intégrateurs utilisent le signal de l'échantillonneur bloqueur échantillonné sur $Clk2$ tandis que l'additionneur utilise le signal échantillonné sur $Clk1$. Finalement, à la sortie de l'additionneur, une paire de transistors et un condensateur sont utilisés pour conserver la sortie du modulateur stable pendant que le condensateur de l'additionneur est déchargé.

4.4 Résultats de simulations

Afin de valider le fonctionnement du circuit conçu, de nombreuses simulations ont été faites. L'objectif de ces simulations est de démontrer que la méthode de minimisation multicritères peut être utilisée pour l'optimisation des convertisseurs $\Sigma\Delta$. Tel qu'expliqué précédemment, les trois mêmes configurations qu'à la section 3.6.4 seront simulées. Deux de ces configurations sont issues du processus d'optimisation avec des valeurs de β différentes (0,05 et 0,01) tandis que la troisième configuration est issue d'un processus de conception traditionnel afin de servir de point de référence. Cette section présente et analyse les résultats obtenus.

4.4.1 SNR

Le premier critère de performance vérifié est le SNR. Lors de l'optimisation des paramètres, l'objectif était un SNR de 78 dB pour une amplitude d'entrée de -5 dBFS. La première ligne du tableau 4.2 présente les SNR obtenus pour une amplitude d'entrée de -5 dBFS pour les trois configurations.

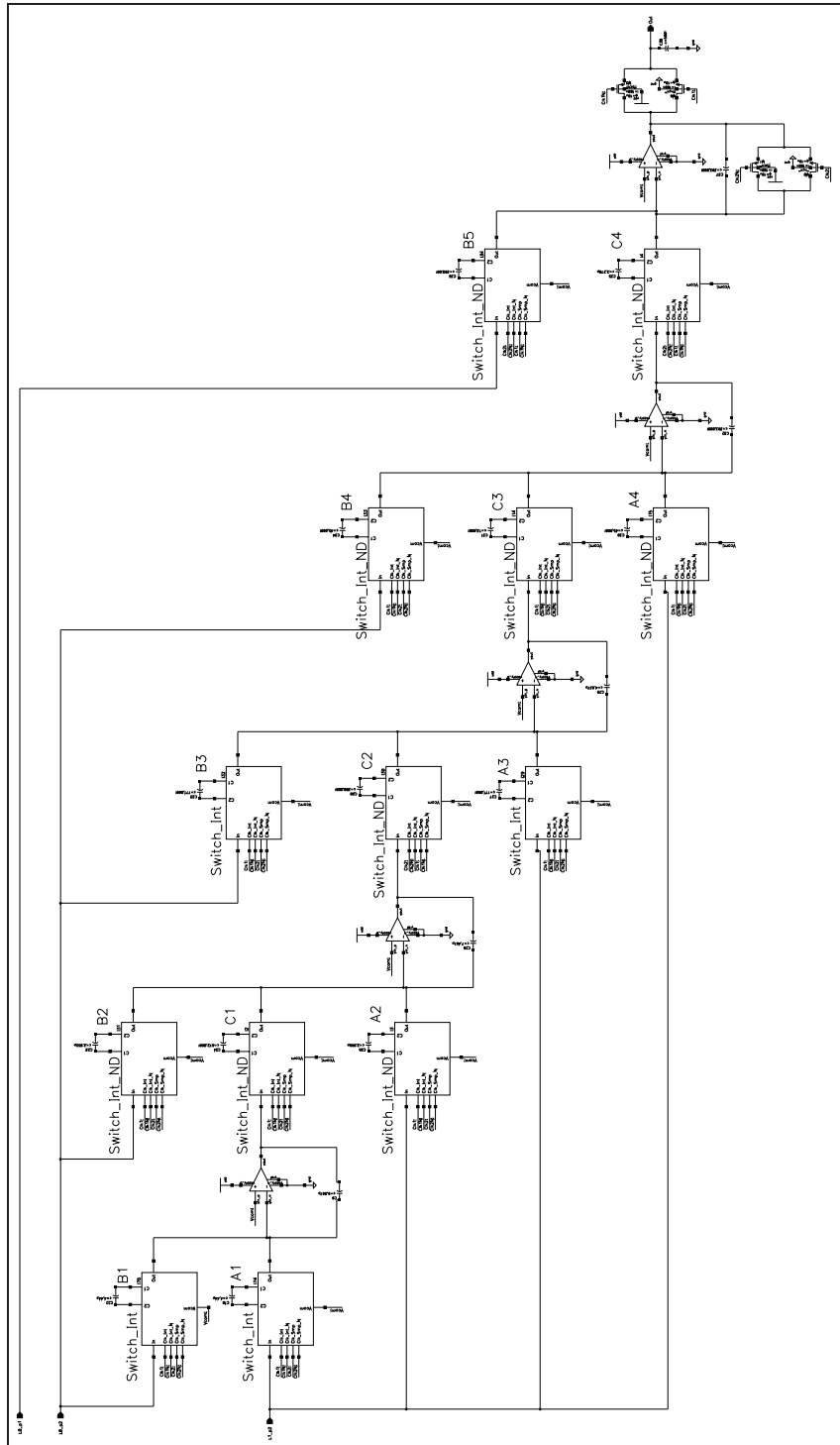


Figure 4.14 Circuit complet du modulateur

La configuration optimisée avec $\beta = 0,01$ obtient les meilleures performances de SNR avec 77,8 dB. Par contre, la configuration optimisée avec $\beta = 0,05$ est 1,6 dB inférieure à l'objec-

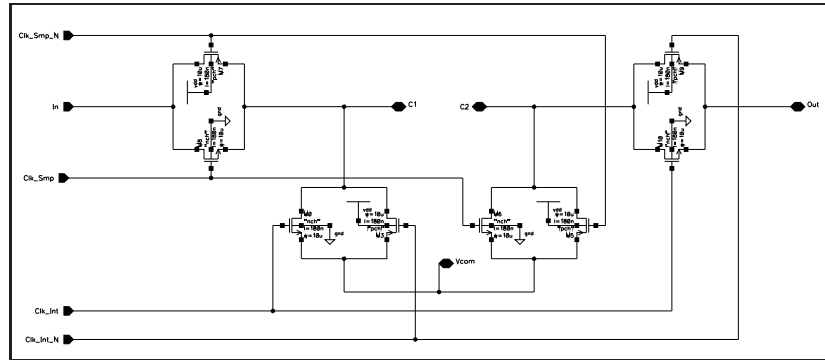


Figure 4.15 Circuit du sous-bloc *Switch_Int*. Contient la configuration d'interrupteurs nécessaire pour réaliser un intégrateur avec délai

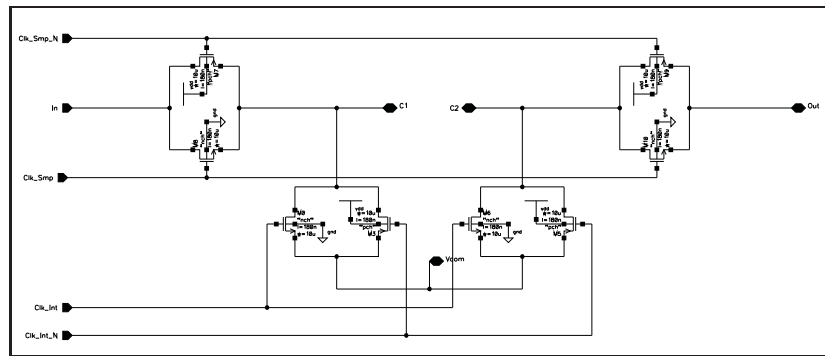


Figure 4.16 Circuit du sous-bloc *Switch_Int_ND*. Contient la configuration d'interrupteurs nécessaire pour réaliser un intégrateur sans délai

Tableau 4.2 Performances obtenues (en dB) par les trois configurations

β	0,01	0,05	Réf.
SNR (-5 dBFS)	77,8	76,4	74,8
SNR maximal (amplitude)	80,0 (à -2,5 dBFS)	79,3 (à -2,5 dBFS)	76,2 (à -2,5 dBFS)
Plage dynamique	83,0	82,5	81,0

tif tandis que la configuration de référence est 3,2 dB inférieure à l'objectif. Ces différences en SNR entre les trois configurations peuvent également être visualisées dans les spectres de fréquences à la sortie des convertisseurs. Les figures 4.17 à 4.19 présentent les spectres de fréquences des trois configurations pour une amplitude d'entrée de -5 dBFS. Dans chacune de ces figures la courbe théorique² de la NTF est ajoutée en pointillés. La ligne verticale en pointillés indique la limite de la bande de fréquences du signal.

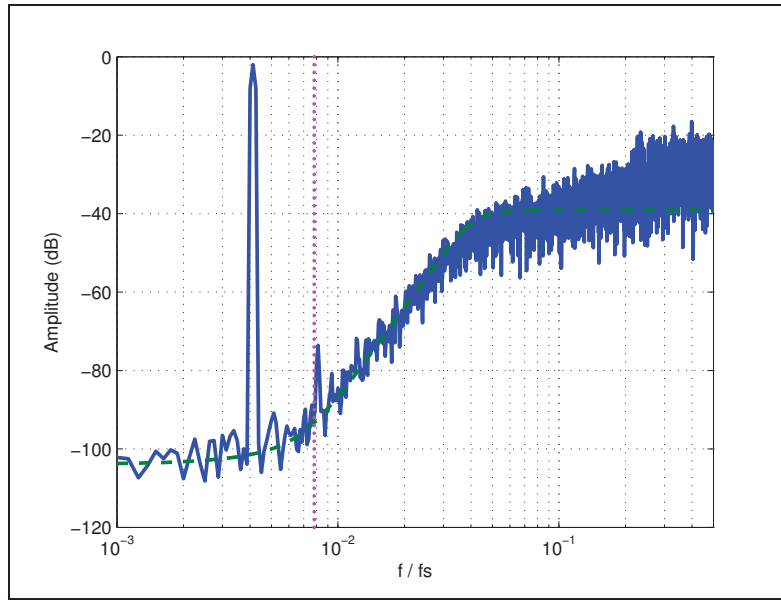


Figure 4.17 Spectre de fréquences de la configuration optimisée avec $\beta = 0,01$

À la figure 4.17 ($\beta = 0,01$), la forme du spectre est similaire à la courbe théorique de la NTF. La seule exception est entre $f=0,1$ et $f=0,5$ où le spectre tend à diverger de la courbe théorique de la NTF. Cependant, étant donné que cela se passe à l'extérieur de la bande de fréquences du signal, cela n'a pas d'impact sur le SNR. À la figure 4.18 ($\beta = 0,05$), le plancher de bruit à l'intérieur de la bande de fréquences du signal est légèrement supérieur à la courbe théorique de la NTF. À la figure 4.18 (configuration de référence), le plancher de bruit à l'intérieur de la bande de fréquences du signal est nettement supérieur à la courbe théorique de la NTF. Dans ces deux derniers cas, cela explique que le SNR n'atteigne pas la valeur visée.

². En plus du bruit de quantification, cette courbe comprend également l'effet des différentes sources de bruit thermique.

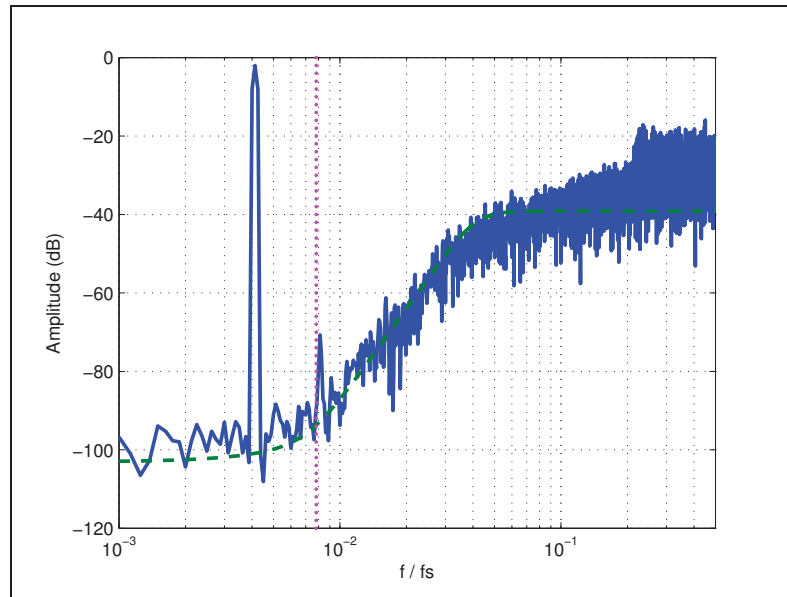


Figure 4.18 Spectre de fréquences de la configuration optimisée avec $\beta = 0,05$

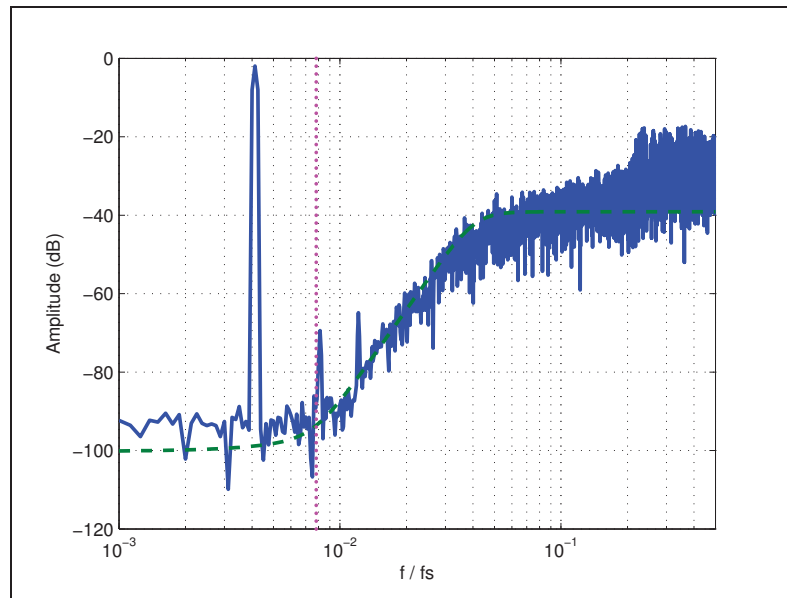


Figure 4.19 Spectre de fréquences de la configuration de référence

La tendance observée lorsque l'amplitude d'entrée est -5 dBFS se répète sur l'ensemble de la plage d'entrée du convertisseur. La figure 4.20 présente les courbes SNR vs amplitude d'entrée

pour les trois configurations. De façon générale, les trois courbes suivent une pente de 1 dB/dB. Sur la majeure partie de la plage d'entrée stable, la courbe de la configuration optimisée avec $\beta = 0,01$ génère le SNR le plus élevé. Suit ensuite la courbe de la configuration optimisée avec $\beta = 0,05$. L'écart entre les deux configurations optimisées est inférieur à 1,5 dB sur l'ensemble de la plage d'entrée stable. L'écart moyen est 1,2 dB. Finalement, la courbe de la configuration de référence est celle offrant le SNR le plus faible sur l'ensemble de la plage d'entrée stable.

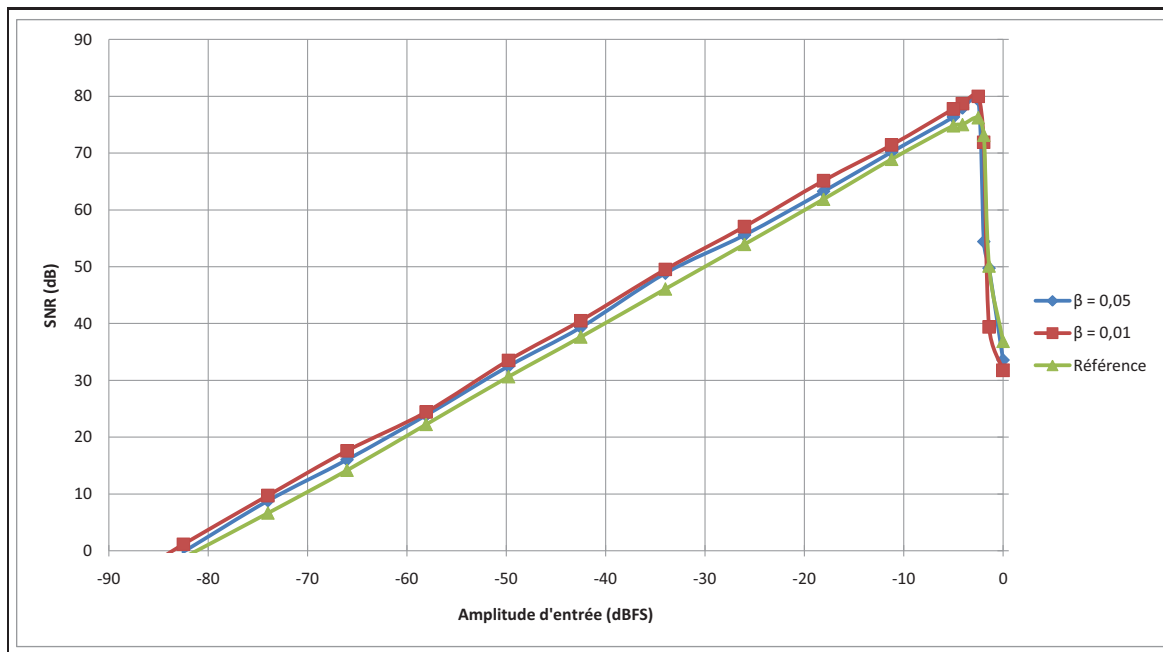


Figure 4.20 Courbes SNR vs amplitude d'entrée pour les trois configurations

Les courbes SNR vs amplitude d'entrée permettent de déterminer le SNR maximal ainsi que la plage dynamique pour chacune des trois configurations. Ces résultats sont présentés dans le tableau 4.2. Les trois configurations atteignent leur SNR maximal à une amplitude d'entrée de -2,5 dBFS. Passé ce point, les modulateurs commencent à devenir instables et le SNR chute abruptement. On remarque également que la configuration optimisée avec $\beta = 0,01$ est celle ayant le SNR maximal et la plage dynamique la plus élevée tandis que la configuration de référence est celle ayant les performances les plus faibles.

L'analyse des résultats permet donc de conclure qu'au niveau du SNR et de la plage dynamique, la configuration offrant les meilleures performances est celle issue du processus d'optimisation

avec $\beta = 0,01$. À l’opposé, la configuration issue d’un processus de conception traditionnel est celle offrant les performances les plus faibles.

Une explication à ces résultats vient de la taille des condensateurs des deux premiers intégrateurs. Avec $\beta = 0,01$, le processus d’optimisation favorise la réduction de l’amplitude de sortie des intégrateurs. Il en résulte que la somme des capacités est plus élevée qu’avec $\beta = 0,05$ mais tout de même moins élevée que pour la configuration de référence. Cependant, les condensateurs des deux premiers intégrateurs ont des valeurs plus élevées. Tel qu’expliqué à la section 3.3, le bruit thermique des deux premiers étages est celui ayant l’impact le plus significatif sur le SNR à la sortie du convertisseur. Il est donc possible d’optimiser le SNR d’un convertisseur en optimisant la répartition des capacités dans le modulateur.

Finalement, la figure 4.21 permet de comparer les courbes SNR vs amplitude d’entrée de la configuration optimisée avec $\beta = 0,01$ obtenu par les simulations au niveau transistor, les simulations au niveau système et les calculs théoriques. L’écart entre les trois courbes est inférieur à 1 dB pour l’ensemble de la plage d’entrée stable tandis que l’écart moyen est 0,5 dB. Il est à noter que la courbe théorique ne tient pas compte de l’instabilité du système.

4.4.2 Amplitude de sortie des intégrateurs

Un des objectifs de la méthode de minimisation multicritères est de diminuer l’amplitude de sortie des intégrateurs. Les amplitudes de sortie maximales normalisées des intégrateurs pour les trois configurations (pour une amplitude d’entrée de -5 dBFS) sont présentés dans le tableau 4.3. Les valeurs théoriques obtenues par la méthode d’optimisation sont également indiquées en parenthèses à titre comparatif. La dernière ligne donne la moyenne des quatre intégrateurs. Il est à noter que toutes ces valeurs sont normalisées à la plage d’entrée pleine échelle du quantificateur.

Afin d’analyser adéquatement ces résultats, il est important de considérer comment les résultats théoriques sont obtenus. À la section 3.6, il a été expliqué qu’avant de débiter le processus d’optimisation, des simulations sont faites afin de connaître la dynamique du système

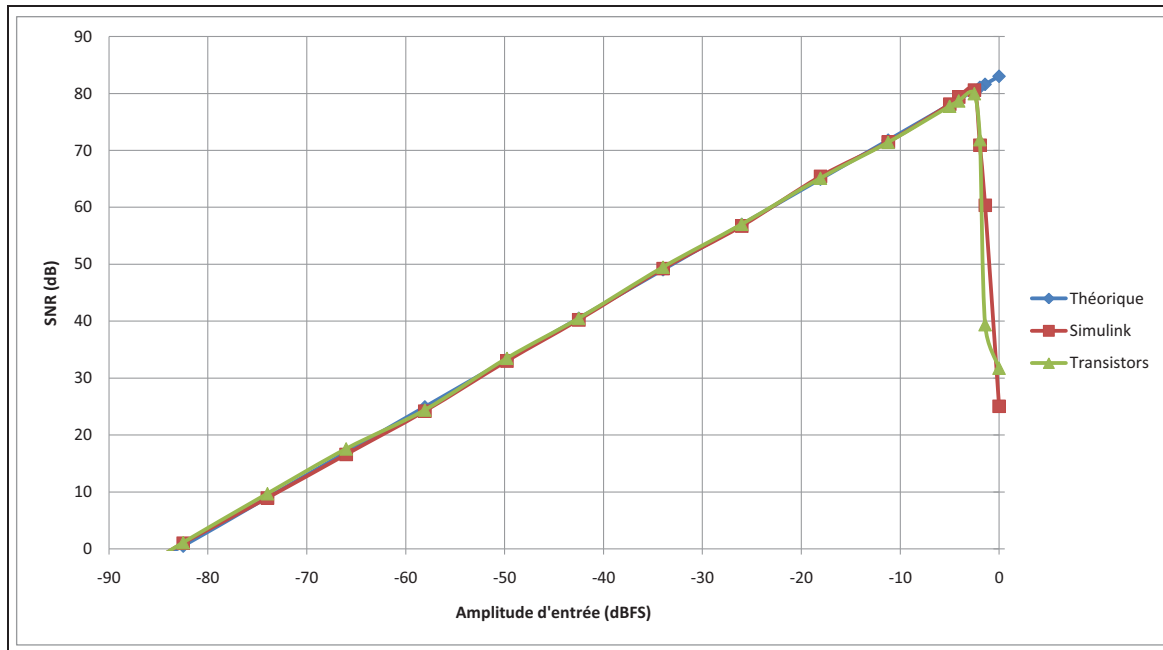


Figure 4.21 Courbes SNR vs amplitude d'entrée pour la configuration optimisée avec $\beta = 0,01$ obtenus de trois façons différentes : calculs théoriques, simulations système avec Simulink et simulations transistors

lorsque les coefficients interétages c_i sont à 1. Par la suite, durant le processus d'optimisation, l'amplitude de sortie des intégrateurs pour un ensemble de coefficients interétages c_i est estimée en fonction de la dynamique du système. Cette façon de faire est inspirée de la fonction *scaleABCD* de la librairie Matlab *Delsig Toolbox*. L'objectif de cette méthode est de s'assurer que l'amplitude de sortie des intégrateurs reste à l'intérieur de la plage désirée. En d'autres mots, les valeurs théoriques représentent les valeurs maximales pouvant être mesurées peu importe le signal d'entrée. En ce sens, les résultats présentés dans le tableau 4.3 sont concluants, car dans la majorité des cas, les résultats obtenus sont inférieurs aux valeurs prévues. Cela signifie que l'amplitude de sortie des intégrateurs est à l'intérieur de la plage désirée. Ce n'est que pour le quatrième intégrateur de la configuration optimisée avec $\beta = 0,05$ que le résultat est supérieur à la valeur théorique.

Finalement, l'effet de la méthode d'optimisation peut être constaté dans les résultats du tableau 4.3. La configuration ayant la moyenne la plus faible est celle obtenue par le processus d'optimisation avec $\beta = 0,01$ tandis que la configuration ayant la moyenne la plus élevée est celle de

Tableau 4.3 Amplitudes de sortie maximales normalisées des intégrateurs pour les trois configurations. Les valeurs théoriques obtenues par la méthode d'optimisation sont indiquées en parenthèse

β	0,01	0,05	Réf.
Intégrateur 1	0,85 (0,90)	0,80 (0,90)	0,83 (0,90)
Intégrateur 2	0,63 (0,80)	0,70 (0,90)	0,70 (0,90)
Intégrateur 3	0,33 (0,43)	0,65 (0,63)	0,60 (0,90)
Intégrateur 4	0,28 (0,33)	0,48 (0,38)	0,73 (0,90)
Moyenne	0,52 (0,61)	0,66 (0,70)	0,71 (0,90)

référence. De plus, pour les deux configurations issues de la méthode d'optimisation, plus un intégrateur est situé loin dans le modulateur, plus son amplitude de sortie diminue. Ces deux constatations démontrent que la méthode d'optimisation a l'effet escompté.

La figure 4.22 présente la distribution de la sortie du premier intégrateur pour les trois configurations. Aucune différence notable ne peut être relevée entre les trois courbes. Par contre, en regardant la distribution de sortie du quatrième intégrateur (présenté à la figure 4.23), l'impact du processus d'optimisation est clairement visible. La courbe de la configuration de référence est celle qui occupe la plus large plage tandis que la courbe optimisée avec $\beta = 0,01$ (configuration favorisant la réduction de l'amplitude de sortie des intégrateurs) est celle ayant la plage la plus étroite. Entre les deux, on retrouve la courbe optimisée avec $\beta = 0,05$.

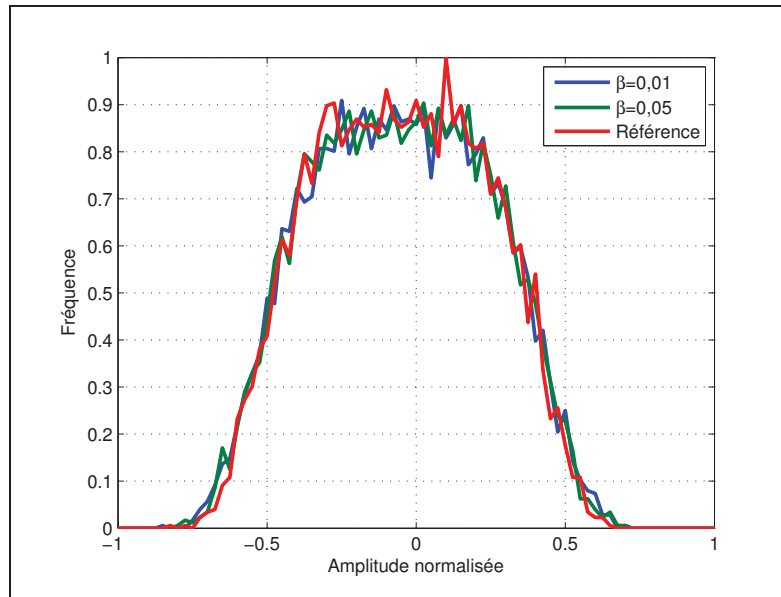


Figure 4.22 Distribution de l'amplitude de sortie du premier intégrateur pour les trois configurations

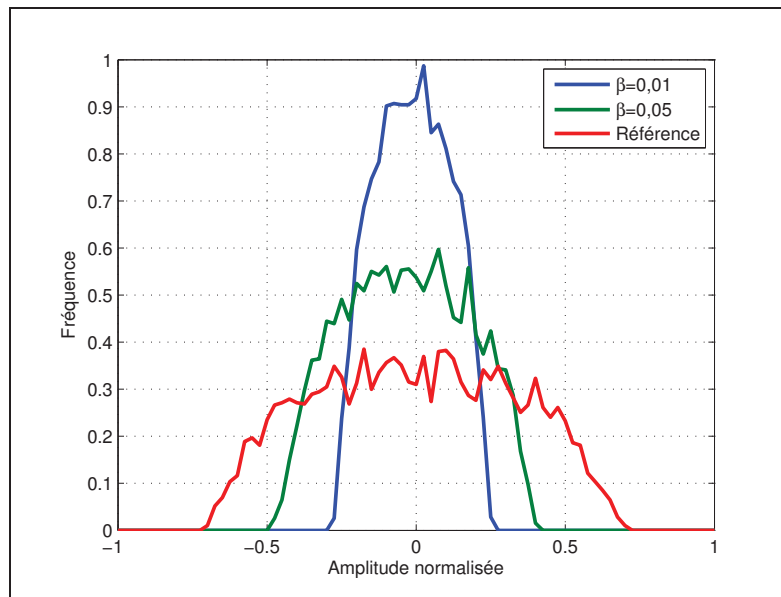


Figure 4.23 Distribution de l'amplitude de sortie du quatrième intégrateur pour les trois configurations

4.4.3 Puissance consommée

La méthode d'optimisation multicritères tente de minimiser deux éléments : l'amplitude de sortie des intégrateurs et la somme des capacités du modulateur. L'une des raisons justifiant ce

choix est l'influence qu'ont ces deux éléments sur la consommation en puissance du circuit. Lors du processus d'optimisation, il n'est pas possible de quantifier ces réductions étant donné que le processus a été conçu pour être indépendant de la technologie cible.

Les simulations au niveau transistor devraient normalement permettre de connaître la consommation en puissance des trois configurations. Malheureusement, l'utilisation de modèles AHDL pour la réalisation des amplificateurs opérationnels et du comparateur rend impossible l'obtention de données probantes. En effet, les modèles AHDL ne consomment aucun courant. En d'autres mots, la puissance consommée par les amplificateurs opérationnels et le comparateur est nulle. Les mesures de la puissance fournie par les alimentations du circuit sont donc non significatives. Dans des travaux futurs, les modèles AHDL des amplificateurs opérationnels devront être remplacés par de vrais circuits afin de permettre une meilleure analyse de la puissance consommée.

4.5 Conclusion

Dans ce chapitre, la méthode de minimisation multicritères présentée au chapitre précédent a été validée à l'aide de simulations au niveau transistor. Ces simulations sont basées sur l'exemple d'utilisation de la méthode de minimisation multicritères présenté au chapitre précédent. Trois configurations sont testées. Les deux premières sont issues du processus d'optimisation avec des objectifs d'optimisation différents. L'une favorise la réduction de l'amplitude de sortie des intégrateurs tandis que l'autre favorise la réduction de la taille des capacités. La troisième configuration, issue d'un processus de conception traditionnel, sert de point de référence.

La première section de chapitre a présenté les étapes de conception du circuit au niveau transistor. Chaque segment a été conçu individuellement. Des simulations temporelles ont permis de vérifier leur fonctionnalité. Finalement, les différents segments ont été réunis afin d'obtenir le convertisseur en entier.

Plusieurs simulations ont été faites afin de comparer les performances des trois configurations. Celle offrant les meilleures performances, tant au niveau du SNR que de la plage dynamique, est celle issue du processus d'optimisation favorisant la réduction de l'amplitude de sortie des intégrateurs. À l'opposé, la configuration offrant les moins bonnes performances est celle obtenue par un processus de conception traditionnel.

L'amplitude de sortie des intégrateurs a également été évaluée. Bien que les résultats diffèrent des valeurs théoriques attendues, l'analyse des résultats permet de conclure que la méthode d'optimisation a l'effet escompté.

CONCLUSION

Un convertisseur $\Sigma\Delta$ utilise un circuit à rétroaction afin de mettre en forme le bruit de quantification à l'extérieur de la bande de fréquences du signal. Combiné au suréchantillonnage du signal, cette technique permet de diminuer l'impact du bruit de quantification sur le SNR.

Le processus de conception d'un convertisseur $\Sigma\Delta$ se divise normalement en trois niveaux d'abstraction différents. Au départ, le convertisseur est modélisé mathématiquement. Ensuite, il est transposé au niveau système dans une structure de réalisation, ce qui nécessite de calculer adéquatement les différents coefficients. Finalement, il est réalisé au niveau transistor. À chacune de ces étapes, de nombreux paramètres doivent être déterminés adéquatement afin d'obtenir les performances désirées. Dans cet ouvrage, deux nouvelles méthodes d'optimisation des convertisseurs $\Sigma\Delta$ ont été développées.

La première section du second chapitre a présenté le cadre théorique nécessaire à l'analyse des différents critères d'optimisation. Ceux-ci sont aux nombres de trois : le bruit thermique, l'amplitude de sortie des intégrateurs et la taille du circuit. Le bruit thermique est l'une des principales limitations des convertisseurs $\Sigma\Delta$. Sa puissance peut être diminuée en augmentant la taille des condensateurs d'échantillonnage. Cependant, cela augmente également la consommation en puissance et la taille du circuit. En effet, pour un nombre de bits de quantification fixe, la taille du circuit sera proportionnelle à la somme des capacités du modulateur. La taille d'un convertisseur $\Sigma\Delta$ peut donc être optimisée en minimisant la somme des capacités du modulateur. Finalement, lors du calcul des coefficients, il est important de s'assurer que l'amplitude de sortie des intégrateurs respecte la plage de sortie linéaire des amplificateurs opérationnels. Dans le cas contraire, des problèmes de saturation et de distorsion harmonique peuvent dégrader les performances du convertisseur. Pour éviter ces problèmes, un facteur d'échelle peut être appliqué au modulateur afin d'ajuster l'amplitude de sortie des intégrateurs aux valeurs désirées.

La seconde moitié du chapitre 2 a présenté une revue de littérature des principaux travaux publiés dans le domaine de l'optimisation des convertisseurs $\Sigma\Delta$. Un outil logiciel largement

utilisé, la *Delsig Toolbox*, est d'abord exposé. Puis, différentes méthodes d'optimisation des convertisseurs $\Sigma\Delta$ ont été résumées. Les objectifs et les types d'algorithmes utilisés diffèrent d'une méthode à l'autre. Cependant, il a été remarqué que de façon générale, l'effet du bruit thermique sur le SNR n'est pas pris en compte par les différents algorithmes.

Une analyse de l'impact de la mise à l'échelle des intégrateurs sur les performances en SNR est présentée au début du chapitre 3. Basée sur les résultats de cette analyse, une méthode d'optimisation visant à minimiser l'amplitude de sortie des intégrateurs est élaborée. Cette méthode utilise une fenêtre d'opération pour établir un compromis entre l'amplitude de sortie des intégrateurs et les performances en SNR. L'amplitude de sortie de chaque intégrateur est réduite jusqu'à ce que l'amplitude de sortie minimale soit atteinte ou que son impact sur le SNR devienne significatif.

Une analyse critique de cette première méthode a permis de déterminer plusieurs lacunes. Celles-ci servent de base au développement d'une seconde méthode. Une fonction de minimisation multicritères est utilisée afin d'optimiser simultanément l'amplitude de sortie des intégrateurs ainsi que la taille des différentes capacités du modulateur. Un paramètre d'optimisation permet à l'utilisateur de prioriser la minimisation de l'une ou l'autre de ces quantités.

Les deux méthodes ont été illustrées par un exemple d'utilisation et les résultats comparés. Avec la méthode des fenêtres, l'amplitude de sortie des intégrateurs est réduite de façon plus agressive. Cependant, la taille des capacités n'est pas prise en compte durant le processus. Il en résulte que les capacités nécessaires pour respecter la puissance de bruit thermique allouée sont beaucoup plus grandes qu'avec la méthode de minimisation multicritères.

Au chapitre 4, la méthode de minimisation multicritères a été confirmée par des simulations au niveau transistor. En premier lieu, les différents paramètres de conception ont été définis. Ensuite, les différents sous-blocs d'un convertisseur $\Sigma\Delta$ ont été conçus individuellement. Le fonctionnement adéquat de ces différents éléments a été validé par des simulations temporelles. Finalement, les performances du convertisseur en entier ont pu être analysées.

Trois configurations différentes ont été testées. Les deux premières étaient issues du processus d'optimisation, l'une favorisait la réduction de l'amplitude de sortie des intégrateurs tandis que l'autre favorisait la réduction de la somme des capacités. La troisième configuration était issue d'un processus de conception traditionnel. Elle était utilisée comme point de référence.

Les résultats des simulations ont permis de démontrer que la configuration optimisée pour favoriser la réduction de l'amplitude de sortie des intégrateurs était celle offrant les meilleures performances en SNR et en plage dynamique. À l'opposé, la configuration offrant les moins bonnes performances était celle issue d'un processus de conception traditionnel.

L'analyse de l'amplitude de sortie des intégrateurs pour les trois configurations a permis de constater l'effet du processus d'optimisation. Pour les deux configurations optimisées, plus un intégrateur était situé loin dans la boucle, plus son amplitude de sortie était diminuée. Finalement, la consommation en puissance n'a pas pu être évaluée en raison des modèles AHDL utilisés dans le circuit.

Pour poursuivre le travail de recherche amorcé dans cet ouvrage, plusieurs axes peuvent être explorés. Tout d'abord, les modèles AHDL pourraient être remplacés par des modèles transistors afin d'évaluer la puissance consommée par les trois configurations. Cela permettrait également de connaître l'impact du bruit thermique généré par les amplificateurs sur les performances du convertisseur. Ultimement, la réalisation d'un circuit intégré permettrait de confirmer totalement le processus d'optimisation.

Pour améliorer le processus d'optimisation, une façon d'intégrer la distorsion harmonique pourrait être étudiée. Cela permettrait probablement d'optimiser davantage l'amplitude de sortie des différents intégrateurs.

Une seconde amélioration serait d'étudier une façon de quantifier la réduction en puissance durant le processus d'optimisation. Étant donné que la puissance consommée varie selon le type de circuits utilisés, le processus d'optimisation pourrait fonctionner en deux temps. Tout d'abord, le processus d'optimisation actuel serait utilisé pour obtenir une configuration initiale.

Celle-ci serait ensuite simulée au niveau transistor afin de connaître la puissance consommée par chaque étage du modulateur. Finalement, un second processus d'optimisation pourrait tenter de minimiser la puissance consommée. Pour cela, les relations entre la taille des capacités, l'amplitude de sortie des intégrateurs et la puissance consommée devraient être déterminées.

Finalement, la méthode pourrait être adaptée à d'autres types d'architecture. Que ce soit des architectures à une seule boucle ou encore des architectures en cascades. Il pourrait également être intéressant de confirmer que l'algorithme fonctionne avec des convertisseurs à temps continu.

BIBLIOGRAPHIE

- [1] R. Schreier et G. C. Temes, *Understanding Delta-Sigma Data Converters*. Hoboken, NJ : Jon Wiley & Sons, 2005.
- [2] E. Collard-Fréchette et G. Gagnon, "A New Optimization Technique for Coefficient Scaling in Sigma-Delta Modulators," in *Proc. IEEE MWSCAS 2010*, Seattle, WA, Août 2010, pp. 312–315.
- [3] E. Collard-Fréchette, G. Kaddoum, et G. Gagnon, "Dynamic Range Scaling of Sigma-Delta Modulators Based on a Multi-Criteria Optimization Process," in *Proc. IEEE NEW-CAS 2011*, Bordeaux, France, Juin 2011, pp. 193–196.
- [4] F. Maloberti, *Data Converters*. Dordrecht, Netherlands : Springer, 2007.
- [5] Y. Ke, J. Craninckx, et G. Gielen, "A Design Approach for Power-Optimized Fully Reconfigurable $\Delta\Sigma$ A/D Converter for 4G Radios," *IEEE Trans. Circuits Syst. II*, vol. 55, pp. 229–233, Mars 2008.
- [6] M. Safi-Harb et G. W. Roberts, "Low Power Delta-Sigma Modulator for ADSL Applications in a Low-Voltage CMOS Technology," *IEEE Trans. Circuits Syst. I*, vol. 52, pp. 2075–2089, Octobre 2005.
- [7] K. B. Weber, W. Skones, C. Talbott, M. Keller, K. Johnson, D. R. Martin, M. Kintis, et I. Robinson, "A Multi-Carrier Basestation Receiver Using a Delta-Sigma Oversampling A/D Converter," in *Proc. IEEE International Conference on Communication*, New York, NY, Août 2002, pp. 877–887.
- [8] D. B. Kasha, W. L. Lee, et A. Thomsen, "A 16-mW, 120dB Linear Switched-Capacitor Delta-Sigma Modulator with Dynamic Biasing," *IEEE J. Solid-State Circuits*, vol. 34, pp. 921–926, Juillet 1999.
- [9] C. B. Wang, "A 20-bit 25-kHz Delta-Sigma A/D Converter Utilizing a Frequency-Shaped Chopper Stabilization Scheme," *IEEE J. Solid-State Circuits*, vol. 36, pp. 566–569, Mars 2001.
- [10] *24-Bit, Pin-Programmable, Ultralow Power Sigma-Delta ADC*, AD7780, Analog Devices, 2009.
- [11] J. Johnston, "A 24-Bit Delta-Sigma ADC with an Ultra-Low Noise Chopper-Stabilized Programmable Gain Instrumentation Amplifier," in *Proc. IEEE Advanced A/D and D/A Conversion Techniques and their Applications*, 1999, pp. 179–182.
- [12] R. Schreier, "The delta sigma toolbox version 7.1," 2004. [Online]. Available : <http://www.mathworks.com/matlabcentral/fileexchange>

- [13] R. Schreier, "An Empirical Study of High-Order Single-Bit Delta-Sigma Modulators," *IEEE Trans. Circuits Syst. II*, vol. 40, pp. 461–466, Août 1993.
- [14] C. Petrie et M. Miller, "A Background Calibration Technique for Multibit Delta-Sigma Modulators," in *Proc. IEEE ISCAS*, Genève, Suisse, Mai 2000, pp. 29–32.
- [15] G. Gagnon et L. MacEachern, "Continuous Compensation of Binary-Weighted DAC Nonlinearities in Bandpass Delta-Sigma Modulators," in *Proc. IEEE ISCAS*, New Orleans, LA, Mai 2007, pp. 253–256.
- [16] D. H. Lee et T. H. Kuo, "Advancing Data Weighted Averaging Technique for Multi-Bit Sigma-Delta Modulators," *IEEE Trans. Circuits Syst. II*, vol. 54, pp. 838–842, Octobre 2007.
- [17] Z. Li et T. S. Fiez, "Dynamic Element Matching in Low Oversampling Delta Sigma ADC," in *Proc. IEEE ISCAS*, Phoenix, AZ, Mai 2002, pp. 683–686.
- [18] O. Nys et R. K. Henderson, "A 19-bit Low-Power Multibit Sigma-Delta ADC Based on Data Weighted Averaging," *IEEE J. Solid-State Circuits*, vol. 32, pp. 933–942, Juillet 1997.
- [19] M. R. Miller et C. S. Petrie, "A Multibit Sigma-Delta ADC for Multimode Receivers," *IEEE J. Solid-State Circuits*, vol. 38, pp. 475–482, Mars 2003.
- [20] P. Benabes, A. Gauthier, et D. Billet, "New Wideband Sigma Delta Convertor," *IEE Electronic Letters*, vol. 29, no. 17, pp. 1575–1577, Août 1993.
- [21] R. Schreier, J. Silva, J. Steensgaard, et G. C. Temes, "Design-Oriented Estimation of Thermal Noise in Switched-Capacitor Circuits," *IEEE Trans. Circuits Syst. I*, vol. 52, pp. 2358–2368, Novembre 2005.
- [22] P. Malcovati *et al.*, "Behavioral Modeling of Switched-Capacitor Sigma-Delta Modulators," *IEEE Trans. Circuits Syst. I*, vol. 50, pp. 352–364, Mars 2003.
- [23] B. Razavi, *Design of Analog CMOS Integrated Circuits*. Singapore : McGraw Hill international edition, 2001.
- [24] O. Oliaei, P. Clément, et P. Gorisse, "A 5-mW Sigma-Delta Modulator With 84-dB Dynamic Range for GSM/EDGE," *IEEE J. Solid-State Circuits*, vol. 37, pp. 2–10, Janvier 2002.
- [25] A. R. Feldman, B. E. Boser, et P. R. Gray, "A 13 Bit, 1.4 MS/s, 3.3V Sigma-Delta Modulator for RF Baseband Channel Applications," in *Proc. IEEE Custom Integrated Circuits Conference*, Santa Clara, CA, Mai 1998, pp. 229–232.
- [26] W. M. Koe et F. Maloberti, "Feed-Forward Path and Gain-Scaling - A Swing and Distortion Reduction Scheme for Second Order Sigma-Delta Modulator," in *Proc. IEEE ISCAS*, Vancouvers, BC, Mai 2004, pp. 413–416.

- [27] F. Maloberti, "Architectures and Design Methodologies for Very Low Power and Power Effective A/D Converters," in *Proc. IEEE NEWCAS 2006*, Gatineau, Qc, Juin 2006, pp. 77–80.
- [28] B. A. Wooley, D. K. Su, S. M. Lee, et K. Y. Nam, "A 1.2V 15-bit 2.5-MS/s Oversampling ADC with Reduced Integrator Swings," in *Proc. IEEE CICC 2004*, Orlando, FL, Octobre 2004, pp. 515–518.
- [29] A. Zahabi, O. Shoaie, Y. Koolivand, et P. Jabehdar-Maralani, "A Two-stage Genetic Algorithm Method for Optimization the $\Sigma\Delta$ Modulators," in *Proc. IEEE ASP-DAC*, Shanghai, Chine, Juillet 2005, pp. 1212–1215.
- [30] A. Marques, V. Peluso, M. S. Steyaert, et W. M. Sansen, "Optimal Parameters for Single Loop $\Delta\Sigma$ Modulators," in *Proc. IEEE ISCAS*, Honk Kong, Juin 1997, pp. 57–60.
- [31] Y. Zhang, E. Hayahara, S. Hirano, et N. Sakakibara, "An Optimal Design Consideration for Higher-order Delta-Sigma AD Converter," in *Proc. IEEE APCCAS*, Tianjin, China, Décembre 2000, pp. 309–312.
- [32] H. W. Wang et T. H. Kuo, "An Automatic Coefficient Design Methodology for High-Order Bandpass Sigma-Delta Modulator With Single-Stage Structure," *IEEE Trans. Circuits Syst. II*, vol. 53, no. 7, pp. 580–584, Juillet 2006.
- [33] D. Gautier, S. Bachir, et C. Duvanaud, "Characterization of BandPass Delta Sigma Modulators in Wireless Transceivers using Parameter Identification," in *Proc. International Conference on Wireless and Mobile Communications*, Cannes, France, Août 2009, pp. 396–400.
- [34] H. Tang et A. Doboli, "High-Level Synthesis of $\Sigma\Delta$ Modulator Topologies Optimized for Complexity, Sensitivity, and Power Consumption," *IEEE Trans. Computer-Aided Design*, vol. 25, no. 3, pp. 597–607, Mars 2006.
- [35] S. Rabii et B. A. Wooley, "A 1.8-V Digital-Audio Sigma-Delta Modulator in 0.8- μm CMOS," *IEEE J. Solid-State Circuits*, vol. 32, pp. 783–796, Juin 1997.
- [36] S. K. Gupta et V. Fong, "A 64-Mhz Clock-Rate $\Sigma\Delta$ ADC With 88-dB SNDR and 105-dB IM3 Distortion at a 1.5-MHz Signal Frequency," *IEEE Trans. Circuits Syst. II*, vol. 37, pp. 1653–1661, Décembre 2002.
- [37] "Optimisation Toolbox Users Guides," The MathWorks Inc., 2010.
- [38] P. E. Allen et D. R. Holberg, *CMOS analog circuit design*. New York, NY : Oxford University Press, 2002.