



Etude et réalisation d'un automate cellulaire opto-électronique parallèle.

Iyad Seyd Darwish

► To cite this version:

Iyad Seyd Darwish. Etude et réalisation d'un automate cellulaire opto-électronique parallèle.. Optique [physics.optics]. Université Paris Sud - Paris XI, 1991. Français. NNT : 1991PA112331 . pastel-00732332

HAL Id: pastel-00732332

<https://pastel.hal.science/pastel-00732332>

Submitted on 14 Sep 2012

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

ORSAY
n° d'ordre :

UNIVERSITE DE PARIS-SUD

CENTRE D'ORSAY

THESE

présentée

pour obtenir

LE TITRE DE DOCTEUR EN SCIENCES

DE L'UNIVERSITE PARIS XI ORSAY

PAR

Iyad SEYD DARWISH

SUJET :

ETUDE ET REALISATION D'UN AUTOMATE CELLULAIRE
OPTOELECTRONIQUE PARALLELE

soutenue le 5 décembre 1991 devant la Commission d'examen

MM.	Serge	LOWENTHAL	<i>Président</i>
	Pierre	CHAVEL	
	Francis	DEVOS	
Mme	Béatrice	St. CRICQ	
	Jean	TABOURY	
	Shamlal	MALLICK	
	J. Paul	POCHOLLE	<i>Rapporteur</i>
	Thierry	MAURIN	<i>Rapporteur</i>

TABLE DES MATIERES

INTRODUCTION GENERALE

Partie 1 : GAZ SUR RESEAU ET AUTOMATE CELLULAIRE **OPTOELECTRONIQUE**

CHAPITRE I : AUTOMATE CELLULAIRE OPTO-ELECTRONIQUE

I-1 INTRODUCTION	9
I-2 PROCESSEUR CELLULAIRE	9
I-3 AUTOMATE CELLULAIRE	11
I-4 EXEMPLE D'ILLUSTRATION	11
I-5 L'APPORT DE L'OPTIQUE DANS LES AUTOMATES CELLULAIRES	14

CHAPITRE II : GAZ SUR RESEAU

II-1 SIMULATION DES ECOULEMENTS DES FLUIDES	19
II-2 PRINCIPE DE L'ALGORITHME	20
II-3 METHODES DE TRAITEMENT	25
III-4 ETAPES DE TRAITEMENT	26
II-5 GAZ SUR RESEAU ET AUTOMATE CELLULAIRE	27

CHAPITRE III : CONCEPTION DE L'AUTOMATE CELLULAIRE

III-1 AUTOMATE CELLULAIRE POUR "GAZ SUR RESEAU"	33
III-2 TRAITEMENT PARALLELE ET SEQUENTIEL	34
III-3 CONCEPTION DE L'AUTOMATE CELLULAIRE OPTO-ELECTRONIQUE	35

Partie 2 : HOLOGRAMME A EFFET TALBOT

CHAPITRE IV : HOLOGRAMME A EFFET TALBOT

IV-1 INTRODUCTION	47
IV-2 ILLUMINER DE TABLEAUX	48
IV-3 EFFET D'AUTO-IMAGERIE DE TALBOT	55
IV-4 TECHNIQUES HOLOGRAPHIQUES	69
IV-5 CONDITIONS D'ENREGISTREMENT	73
IV-6 REALISATION DE L'HOLOGRAMME	77
IV-7 COMPARAISON AVEC LES AUTRES "ILLUMINATEURS DE TABLEAUX"	81

CHAPITRE V : ANALYSE DES ABERRATIONS CHROMATIQUES DE L'HOLOGRAMME

V-1- INTRODUCTION	85
V-2- MODELISATION MATHEMATIQUE	86
V-3 APPROXIMATION DE FRESNEL ET RECONSTRUCTION DE L'IMAGE	95
V-4 TERMES D'ABERRATION	102
V-5 SIMULATION	111
V-6 EFFET DE LA CORRECTION DES ABERRATIONS	117
V-7 AUTRE POSSIBILITE POUR LA CORRECTION DES ABERRATIONS	118
V-8 CONCLUSION	118

Partie 3 : REALISATION MICROELECTRONIQUE ET CARACTERISATION

CHAPITRE VI : CIRCUIT ELECTRONIQUE INTEGRE (VLSI)

VI-1 LES FONCTIONS DU CIRCUIT	125
VI-2 CONCEPTION DU CIRCUIT ELECTRONIQUE	127
VI-3 DESSIN ET REALISATION DU CIRCUIT	133
VI-4 TEST DU CIRCUIT	138
VI-5 PERSPECTIVES	148
VI-5 CONCLUSION	151

CHAPITRE VII : ETUDE OPTO-ELECTRONIQUE

VII-1 CONSTITUTION ET FONCTIONNEMENT DES PHOTODIODES	141
VII-2 MODELISATION DU RENDEMENT QUANTIQUE DES PHOTODIODES	162
VII-3 INFLUENCE DE LA CAPACITE INTERNE DES PHOTODIODES SUR LE TEMPS DE REPONSE	166
VII-4 EVALUATION DE LA SENSIBILITE DES PHOTODIODES EN FONCTION DE λ	167
VII-5 ETUDE DE CAS	168
VII-6 CONCLUSION	173

CHAPITRE VIII : REALISATION EXPERIMENTALE

VIII-1 CIRCUIT DE COMMANDE	177
VIII-2 DIODE LASER D'ECLAIRAGE	181
VIII-3 MODULATEURS OPTO-ELECTRONIQUE	183
VIII-4 MONTAGE EXPERIMENTAL	188

CHAPITRE IX : CARACTERISATION ET ARCHITECTURE DES AUTOMATES CELLULAIRES

IX-1 CARACTERISATION DE L'AUTOMATE EN PROJET	193
IX-2 COMPARAISON ENTRE ARCHITECTURES OPTOELECTRONIQUES ET TOUT ELECTRONIQUES	199
XI-3 CONCLUSION	202
 CONCLUSION GENERALE	 205
ANNEXES	211
BIBLIOGRAPHIE	257

Abstract

The cellular automata, constituted by a large number of elementary processors, permit fast and high performance processing of certain algorithms.

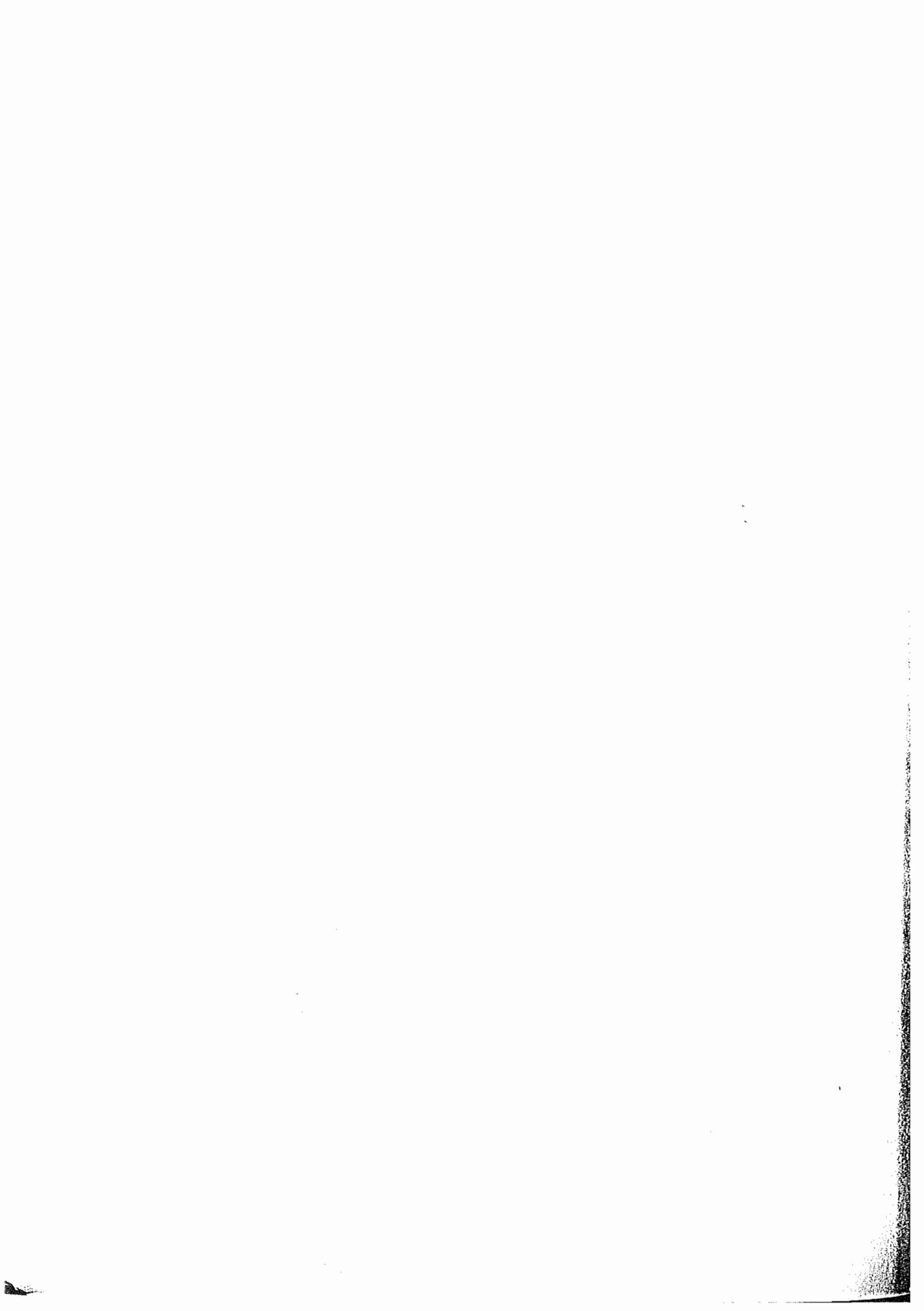
In the present work, we study the possibility of implanting in parallel such an automaton on an integrated circuit using integrated photodiodes for the input optical signals, an array illuminator, and multiple-quantum well modulators for the output signals.

Different components were studied and implemented for this project :

- An array illuminator realized as a hologram using the self-imaging Talbot effect. The chromatic aberrations can be minimized if the recording setup is adequately modified;
- A VLSI electronic circuit containing a single elementary processor, with integrated photodiodes for the input optical data and with special output pads for the modulators;
- Opto-electronic multiple-quantum-well modulators for a parallel output.

An experimental setup was realized by illuminating a photodiode with a diode laser at 999 nm.

An estimation of the performances of the proposed automaton shows its high capacity for calculations and connections.



à Nicole de Beauhoudrey et Line Garnero qui venaient à mon secours aux moments difficiles. Comment oublier Madeleine Calvignac et les services indispensables qu'elle m'a rendus. Je suis extrêmement touché par la gentillesse de Jean Paul Hugonin, il a toujours répondu à mes questions avec modestie. Les discussions avec François Chateaux m'étaient fort agréables.

J'aimerais exprimer ma reconnaissance aux personnes de l'équipe de l'I.E.F qui m'ont témoigné leur amitié et m'ont fait profiter de leurs compétences : Roger Reynaud, Kurosh et Eric.

Je voudrais remercier Jean Claude Rodier et Alain Bellemain de l'équipe électronique pour l'aide qu'ils m'ont accordée et les conseils qu'il m'ont prodigués. Je remercie également les personnes de l'atelier mécanique qui ont réalisé à temps tous mes travaux malgré mon retard.

Un grand merci aux stagiaires, Yan Malet et Philippe Legendre, pour leur précieuse contribution à ce travail.

Je suis extrêmement touché par la gentillesse et l'amitié que m'ont témoignées Dominique Baude, Nicolas Chateau, Dominique Morichère, Jan Svatos et bien d'autres. Je garderai le souvenir de la sympathique ambiance qu'ils ont su créer.

Je ne saurai terminer sans remercier Madame Delmotte qui a assuré la reproduction de ce mémoire, ainsi que toutes les personnes qui m'ont aidé au cours de ce travail et que je n'aurais pas citées.

En guise de prologue, j'aimerais exprimer ma reconnaissance à toutes les personnes qui ont contribué à la réalisation de ce travail.

Je voudrais exprimer toute ma reconnaissance à Monsieur le Professeur Serge Lowenthal de m'avoir accueilli dans son groupe et m'a permis, avec sa bienveillance, de travailler dans de très bonnes conditions matérielles et humaines. Je suis très sensible à l'honneur qu'il me fait de présider mon jury.

Je remercie Monsieur le Professeur Christian Imbert de m'avoir accueilli à l'Institut d'Optique.

Je tiens à exprimer toute ma gratitude à Pierre Chavel qui m'a proposé et encadré cette thèse. Non seulement il m'a fait bénéficier de ses compétences scientifiques, il m'a aussi permis d'apprécier les valeurs humaines comme la gentillesse, la pondération, et l'honnêteté. Jean Taboury a aussi encadré ce travail avec grand intérêt. Qu'il soit vivement remercié pour les constants encouragements, idées et conseils qu'il a su me prodiguer.

Je suis profondément touché par la gentillesse de Shamlal Mallick qui a su porter conseil et soulagement à mes soucis aussi bien sur le plan scientifique que sur le plan humain.

Je tiens à exprimer mes profonds remerciements à Francis Devos qui m'a accueilli dans son équipe à l'I.E.F pendant ma préparation de la partie électronique. Ses compétences et sa bienveillance m'étaient fort utiles.

Mes plus vifs remerciements vont à Thierry Maurin qui a dirigé, avec grande attention, mon travail. à l'I.E.F. Il était toujours disponible pour répondre à mes questions. Je suis honoré qu'il ait accepté la lourde charge de juger mon travail.

Je suis particulièrement reconnaissant à Jean Paul Pocholle qui n'a ménagé ni temps ni compétences pour me fournir le matériel nécessaire à l'expérimentation. Je le remercie vivement pour l'intérêt qu'il a porté sur mon travail et pour avoir accepté d'être rapporteur de ma thèse.

Béatrice St Cricq n'a pas hésité à me fournir le matériel indispensable au montage. Je suis honoré qu'elle soit membre de mon jury. Je souhaite qu'elle trouve ici mes vifs remerciements.

Toute ma reconnaissance va à l'équipe de "Physique des images" de l'Institut d'optique, tout particulièrement à Philippe Lalanne qui était toujours à mon écoute, il a su me soulager de mes soucis avec son humour. Denis Joyeux et François Polack m'ont aidé avec grande patience à résoudre les énigmes de l'ordinateur et ils m'ont fait profiter de leurs compétences. Daniel Phalippou, avec ses qualités humaines et son expérience, a rendu mon travail plus facile et plus agréable. Le grand soin et le sens artistique de Jean Claude Saget m'ont été d'une grande utilité pour la photographie et la réalisation des hologrammes. Grand merci

A ma mère

A mon père

INTRODUCTION GENERALE

INTRODUCTION GENERALE

Les architectures parallèles est un sujet actuel de recherche sur les machines de traitement d'information pour pouvoir exécuter des algorithmes dont le traitement s'avère long sur les machines séquentielles connues jusqu'à maintenant. Comme exemple nous pouvons citer les algorithmes de traitement d'images : reconnaissance de forme, détermination de contour ... etc, ou des phénomènes physiques comme les écoulements des fluides et la détermination des régimes turbulents autour des obstacles qui, vue la complication du traitement, passe au niveau expérimental plutôt que des simulations numériques. Il est certain que les avantages que présentent les machines parallèles sont nombreux mais essentiellement le gain en temps d'exécution des algorithmes.

Plusieurs algorithmes et phénomènes physiques présentent l'aspect d'interaction des entités avec leurs voisins : des pixels pour le traitement d'images, des particules pour les phénomènes des écoulements ou de propagation Les machines tenant compte de cette structure sont de performances supérieures à celles des machines conventionnelles. D'où l'apparition des processeurs et automates cellulaires qui sont des modèles mathématiques simples organisés sous forme de cellules interagissant avec un certain nombre de voisinage. A partir d'un traitement local concernant chaque cellule, ces modèles peuvent simuler des effets divers et compliqués pour les systèmes physiques et biologiques. L'implantation de ces modèles sur des machines spécialisées peut aboutir à des unités de traitement annexes aux ordinateurs capables de réaliser un traitement rapide et efficace. Mais une implantation en architecture parallèle de ces modèles exige un grand nombre de connexions aussi bien entre les cellules ou avec l'extérieur.

De son côté, la recherche en optique est orientée, dans plusieurs laboratoires, vers l'introduction des techniques optiques dans les systèmes informatiques à cause de la performance de ces techniques par rapport à l'électronique surtout en ce qui concerne les connexions entre les unités de traitement. En fait, l'optique a le potentiel de réaliser des connexions de haute densité avec un grand débit d'information bien supérieurs à ceux offerts par les connexions électriques.

Le présent travail est un essai de réalisation d'un modèle d'automate cellulaire alliant les performances de l'électronique pour les unités de traitement (large integration, commutation rapide, fiabilité ...etc) et les performances de l'optique pour les connexions surtout en ce qui concerne les entrées et sorties parallèles.

Comme exemple de démonstration, nous avons choisi l'algorithme "Gaz Sur Réseau" pour la simulation des écoulements des fluides. Cet algorithme est un modèle simple basé sur les interactions entre les particules du gaz se déplaçant sur un réseau régulier. Ces performances ont été montrés et testés sur plusieurs écoulements compliqués et ont donnés des résultats encourageants. Son implantation sur un automate cellulaire a été déjà faite par une machine électronique. Cet algorithme est exécutable sur une machine parallèle.

Ce manuscrit est composé de plusieurs parties :

- La première partie présente le modèle des processeurs et automates cellulaires. Ainsi qu'elle définit l'agorithme "Gaz Sur Réseau" avec ses règles du traitement et démontre la possibilité de son exécution sur les automates cellulaires. Dans le dernier chapitre de cette partie, nous présentons la conception de l'automate cellulaire opto-électronique destiné à exécuter cet algorithme.

- L'hologrammes à effet Talbot est un nouveau précédé pour les illuminateurs de tableaux, un dispositif optique indispensable pour toute réalisation optique ou opto-électronique des machines parallèles. La présentation de ce procédé, ainsi que la réalisation et l'étude de ses aberrations chromatiques est le sujet de la deuxième partie de ce manuscrit.

- Dans la troisième partie, nous présentons la réalisation du circuit électronique intégré (VLSI) qui contient les unités de traitement. La caractérisation des photodiodes intégrées est présentée dans le chapitre VII. Après une description du montage réalisé, nous estimons les caractéristiques générales de l'automate avec les techniques utilisées, et nous menons une comparaison entre les architectures tout-électroniques et opto-électroniques.

PREMIERE PARTIE

GAZ SUR RESEAU

ET AUTOMATE CELLULAIRE

OPTO-ELECTRONIQUE

CHAPITRE I

AUTOMATE CELLULAIRE OPTO-ELECTRONIQUE

CHAPITRE I

AUTOMATE CELLULAIRE OPTO-ELECTRONIQUE

I-1 INTRODUCTION

Dans différentes applications de traitement d'information, les données sont distribuées sur des points, et le traitement pour chaque point est directement lié à l'état des points voisins. L'exemple type de telles applications se trouve dans le traitement d'images comme la détermination des contours ou la reconnaissance des formes. Le traitement séquentiel de telles applications sur les ordinateurs conventionnels s'avère long du point de vue du temps d'exécution. Par contre ces applications peuvent être traitées plus facilement et plus rapidement sur des machines spécialisées qui prennent en compte la structure du problème, d'où l'apparition de la notion du traitement cellulaire. Le traitement cellulaire se fait sur des cellules contenant les informations nécessaires, chaque cellule est liée à un certain nombre de cellules voisines et son évolution dépend de l'état du voisinage. Les machines spécialisées pour un tel traitement s'appellent "processeurs cellulaires". Les "automates cellulaires" sont un cas particulier des processeurs cellulaires.

Dans ce chapitre, nous définissons, dans le cas général, la notion de processeurs cellulaires pour en arriver au cas particulier des automates cellulaires et donnons un exemple simple d'application. A la fin de ce chapitre nous introduisons l'intérêt de l'optique dans la réalisation de tels automates en précisant la notion d'"automate cellulaire opto-électronique".

I-2 PROCESSEUR CELLULAIRE

Le processeur cellulaire est constitué d'un ensemble C d'unité de traitement que nous appelons *cellules*. Nous considérons, pour souci de clarté, que ces cellules sont distribuées sur un réseau bidimensionnel. Nous désignons par c_{ij} l'une de ces cellules. Chaque cellule contient certaines informations qui représentent son état que nous désignons par $[c_{ij}][1-5]$.

Nous distinguons deux catégories de processeurs cellulaires selon la représentation de l'état de chaque cellule :

- processeurs cellulaires analogiques : où l'état de chaque cellule est présenté sous forme analogique; tension électrique, courant, intensité lumineuse.
- processeurs cellulaires binaires : où l'état de chaque cellule est présenté sous forme numérique sur un certain nombre de bits.

L'évolution de chaque cellule dépend de son état et de l'état d'un certain nombre de cellules voisines que nous désignons par l'ensemble V_{ij} qui est, en fait, le sous ensemble des cellules C qui sont connectées à la cellule c_{ij} , la cellule c_{ij} est incluse dans son voisinage. L'évolution est régie par une loi liée dans le cas général à chaque cellule et que nous désignons par la fonction F_{ij} . Ainsi la loi d'évolution de chaque cellule peut être écrite par la formule :

$$[c_{ij}] = F_{ij}(V_{ij})$$

Nous considérons les processeurs synchrones ou à séquençement, c'est à dire que les nouveaux états de toutes les cellules $[c_{ij}]$ sont déterminés simultanément aux temps discrets d'une horloge. Notons les temps d'horloge $t, t+1, t+2 \dots$ etc. Dans ce cas la relation précédente s'écrit sous la forme :

$$[c_{ij}](t+1) = F_{ij}[V_{ij}(t)]$$

Deux fonctions essentielles se présentent pour le fonctionnement de chaque cellule :

- L'interconnexion avec les cellules voisines :

Cette interconnexion apporte l'information de l'état de chaque cellule c_{kl} de l'ensemble V_{ij} à la cellule c_{ij} pour contribuer à la prise de décision concernant l'évolution de son état.

- L'évolution de l'état de la cellule :

Cette évolution est déterminée par la fonction F_{ij} . Cette fonction est une fonction quelconque des cellules de l'ensemble du voisinage V_{ij} . L'équation d'évolution s'écrit sous sa forme la plus générale comme suit :

$$[c_{ij}](t+1) = F_{ij} \left\{ [c_{kl}](t) \right\}_{c_{kl} \in V_{ij}}$$

Nous écrivons cette équation plus explicitement en introduisant les fonctions Φ_{ij}^{kl} qui extraient de l'état de la cellule c_{kl} du voisinage V_{ij} l'information nécessaire à l'évolution de la cellule c_{ij} . De cette façon l'équation d'évolution s'écrit sous la forme :

$$\left[c_{ij} \right](t+1) = F_{ij} \left(\Phi_{kl}^{ij} \left(c_{ij}(t) \right) \right)_{c_{kl} \in V_{ij}}$$

I-3 AUTOMATE CELLULAIRE

La caractéristique essentielle d'un automate cellulaire, par rapport au processeur cellulaire, est l'invariance spatiale ou invariance par translation. Plus précisément, un automate cellulaire est un processeur cellulaire vérifiant les conditions suivantes [3] :

- i - La structure du voisinage est la même pour toutes les cellules. C'est à dire que si pour deux entiers donnés k et l la cellule $c_{i+k,j+l}$ appartient au voisinage de la cellule $c_{i,j}$, ceci est vrai pour tous les i,j . Cette propriété est vraie pour les espaces non limités (infinies ou périodiques); pour les phénomènes physiques limités dans l'espace, un traitement spécial est nécessaire pour les frontières. Les fonctions de contribution ne dépendent que de la position relative des cellules, par conséquent peuvent être écrites sous la forme $\Phi_{i-k,j-l}$. Cette propriété peut être exprimée par l'invariance par translation.
- ii - La fonction de l'évolution est la même pour toutes les cellules. Dans ce cas les fonction F_{ij} seront toutes identiques que nous notons tout simplement F .
- iii - La structure du voisinage et l'équation d'évolution sont indépendantes du temps, ce qui n'est pas nécessaire pour les processeurs cellulaires en général [3].

Dans ces conditions l'équation d'évolution de chaque cellule devient:

$$\left[c_{ij} \right](t+1) = F \left\{ \Phi_{i-k,j-l} \left(\left[c_{kl} \right](t) \right) \right\}_{c_{kl} \in V_{ij}}$$

Un automate cellulaire est dit "élémentaire" si l'état de chaque cellule est codé sur un seul bit; donc sa réponse est "0 ou "1".

I-4 EXEMPLE D'ILLUSTRATION

Pour illustrer la notion d'automate cellulaire nous présentons un exemple simple de traitement d'images.

Il s'agit de déterminer le contour d'une image de 8*8 pixels, l'état de chaque pixel se présentant sous la forme d'un niveau de gris codé sur 8 bits.

Cet exemple peut être implanté sur un automate cellulaire constitué d'une matrice de 8*8 cellules. L'état initial de chaque cellule est codé sur 8 bits comme suit :

$$[c_{i,j}](0) = (e_{ij}^0, e_{ij}^1, \dots, e_{ij}^7)$$

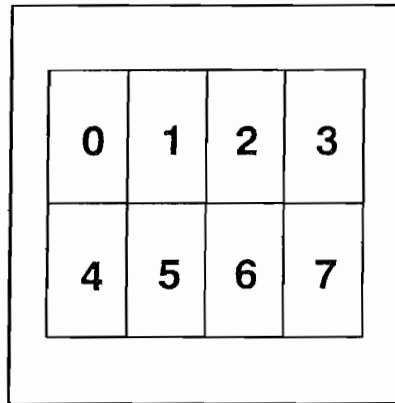


Fig I-1 : Bits d'état d'une cellule.

chaque cellule, à part certaines les cellules de bord, est connectée à quatre cellules voisines comme le montre la figure suivante :

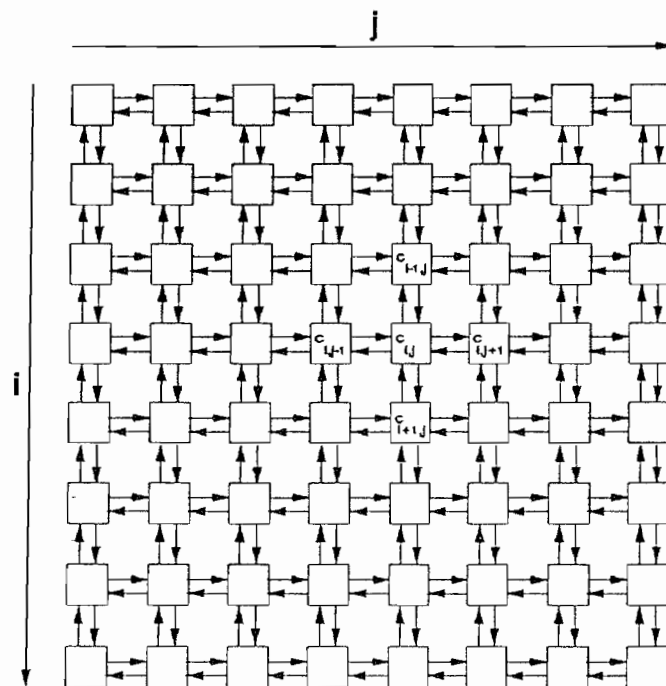


Fig I-2 : Connexion des cellules.

Ainsi le voisinage de la cellule c_{ij} est constitué de l'ensemble suivant :

$$V_{ij} = \{c_{i,j}, c_{i-1,j}, c_{i,j-1}\}$$

La détection du contour se fait en une seule étape en traitement parallèle pour toutes les cellules suivant le passage à "1" le bit n° 6 ou le bit n° 7 de la cellule $c_{i,j}$ par rapport au voisinage, ceci se présente selon la règle suivante :

$$\begin{aligned} \left[c_{i,j} \right](1) &= 1 \quad \text{dans un des cas suivants} & \begin{cases} e_{i,j}^7(0) = 1 \text{ et } e_{i-1,j}^7(0) = 0 \\ \text{ou : } e_{i,j}^7(0) = 1 \text{ et } e_{i,j-1}^7(0) = 0 \\ \text{ou : } e_{i,j}^7(0) = 1 \text{ et } e_{i+1,j}^7(0) = 0 \\ \text{ou : } e_{i,j}^7(0) = 1 \text{ et } e_{i,j+1}^7(0) = 0 \end{cases} \\ \left[c_{i,j} \right](1) &= 0 \quad \text{autrement} \end{aligned}$$

Ainsi nous remarquons que la prise de décision concernant chaque cellule ne dépend que du 8ème bit de chaque cellule dans l'ensemble du voisinage. Les fonctions de contribution se définissent comme suit :

$$\begin{aligned} \Phi_{i-k,j-l} \left(\left[c_{kl} \right](t) \right) &= e_{ij}^7 \quad \text{si } (k,l) = (i,j), (i-1,j), (i+1,j) \\ &\quad (i,j-1), (i,j+1) \\ \Phi_{i-k,j-l} \left(\left[c_{kl} \right](t) \right) &= 0 \quad \text{autrement} \end{aligned}$$

La fonction d'évolution est définie comme suit :

$$F = \left(e_{i,j}^7 \wedge \overline{e_{i-1,j}^7} \right) \vee \left(e_{i,j}^7 \wedge \overline{e_{i+1,j}^7} \right) \vee \left(e_{i,j}^7 \wedge \overline{e_{i,j-1}^7} \right) \vee \left(e_{i,j}^7 \wedge \overline{e_{i,j+1}^7} \right)$$

$\wedge$: "et" logique.

$\vee$: "ou" logique.

La figure suivante montre le résultat d'un tel traitement sur une image simple :

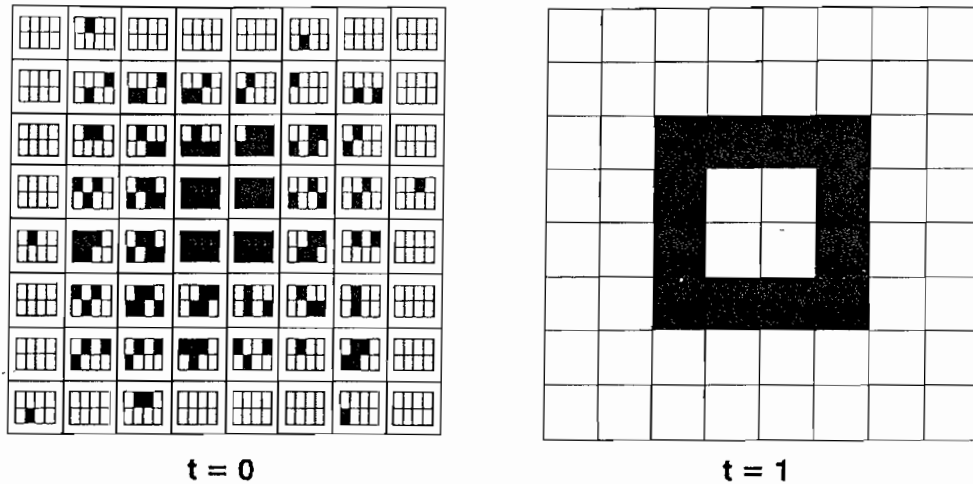


Fig I-3 : Exemple du traitement.

I-5 L'APPORT DE L'OPTIQUE DANS LES AUTOMATES CELLULAIRES

Les processeurs et les automates cellulaires, étant constitués de nombreuses unités de traitement pouvant travailler simultanément, sont traitable sur des machines à architecture parallèle. Ces architectures devraient augmenter considérablement la capacité de traitement des ordinateurs.

Le problème essentiel de ces architectures réside dans le réseau de connexion entre les différentes unités de traitement ou les processeurs élémentaires (P.E.). Pour réaliser des machines parallèles suffisamment puissantes, elles doivent contenir un très grand nombre de P.E. avec une très haute densité de connexions entre eux. Bien que l'électronique offre des composants très performants au niveau du temps et de l'énergie de commutation, elle est très limitée au niveau du transfert d'information. Cette limitation est due essentiellement à la structure planaire des connexions électriques et à l'interaction électrique entre les fils de connexion.

De son côté, l'optique est capable de réaliser des connexions dans la troisième dimension en employant les techniques de l'holographie, de l'acousto-optique ou d'autres technologies. Plusieurs avantages se présentent pour les connexions réalisées par voie optique comme :

- La haute densité de connexion du fait que les faisceaux lumineux peuvent se croiser lors de la propagation en espace libre sans interagir entre eux [6,7].
- Le haut débit d'information dû à l'élimination de l'impédance de propagation existant pour les fils électriques.

Ainsi l'optique peut assurer des connexions nécessaires entre les P.E. d'un automate cellulaire. Reste à signaler les difficultés d'alignement des dispositifs optiques qui demande une précision et une stabilité mécanique non négligeable.

Le type de technologie caractérisant le nom de la machine (optique, électronique...etc) concerne essentiellement le (ou les) support (s) de l'information à l'intérieur de la machine, ou la façon dont d'information est codée et traitée. Ce support concrétise physiquement les fonctions essentielles d'une machine de traitement d'information : la mémorisation, la transformation et le transport d'information à l'intérieur de la machine (Un ordinateur classique reste électronique même s'il est connecté par un câble de fibres optiques sur un réseau car c'est une connexion extérieure). Ainsi, une machine électronique mémorise et transforme les informations par des transistors et les transporte par des fils électrique. De même une machine optique transforme les informations par des dispositifs optiques (ou opto-électroniques) même si l'électronique intervient pour modifier les caractéristiques de ces dispositifs (modulateurs acousto-optiques, SEED'setc).

Dans une machine opto-électronique, le codage de l'information change entre les trois étapes de fonctionnement citées précédemment; par exemple une machine qui utilise l'optique comme support de l'information pour la communication entre les différentes unités, tandis que celles-ci sont électroniques, est une machine opto-électronique.

Ainsi, nous appelons automate cellulaire opto-électronique toute conception et réalisation d'automate cellulaire fondée sur l'emploi des dispositifs des deux technologie (électronique et optique) pour le codage de l'information entre les différentes étapes du traitement.

CHAPITRE II

GAZ SUR RESEAU

CHAPITRE II

GAZ SUR RESEAU

Une réalisation d'un automate cellulaire nécessite le choix d'un algorithme de démonstration. Nous avons choisi l'algorithme du "Gaz sur Réseau" pour la simulation des écoulements des fluides à cause de son intérêt physique, sa simplicité de représentation et le gain intéressant en temps d'exécution que nous pouvons obtenir par son implantation en architecture parallèle.

Dans ce chapitre, nous présentons le "Gaz Sur Réseau" et la simplification qu'il porte pour la simulation des écoulements des fluides.

Après une revue rapide des méthodes de simulation de la mécanique des fluides, nous définissons l'algorithme "Gaz sur réseau" avec son principe, ses règles, et les différentes étapes de traitement utilisant la méthode de substitution symbolique.

Ensuite, nous précisons l'aspect cellulaire de son traitement et évoquons la possibilité d'une implantation sur un automate cellulaire.

II-1 SIMULATION DES ECOULEMENT DES FLUIDES

Du point de vue macroscopique, l'écoulement d'un fluide est décrit par trois équations aux dérivées partielles qui sont obtenues en écrivant en chaque point du fluide :

- la loi de la conservation de la masse : équation de continuité.
- la loi de la conservation de l'impulsion : équation de Navier-Stokes.
- la loi de la conservation de l'énergie : équation de la chaleur.

La résolution de ces équations nécessite l'emploi de programmes compliqués, ainsi que l'utilisation d'ordinateurs puissants avec des temps d'exécution importants. L'approche de la dynamique moléculaire étudie le comportement microscopique des particules du fluide [8]. En se servant de la mécanique statistique nous pouvons retrouver les propriétés macroscopiques en partant de l'approche microscopique. Numériquement, la dynamique moléculaire nécessite la résolution de l'équation de mouvement pour chaque particule considérée. Par conséquent, cette

méthode ne permet de traiter qu'un nombre très limité de particules, ce nombre est trop petit pour simuler un écoulement macroscopique intéressant [9].

Le gaz sur réseau correspond à une simplification extrême de la dynamique moléculaire. Selon la dynamique moléculaire, la position et la vitesse des particules peuvent prendre n'importe quelles valeurs. Par contre, dans le gaz sur réseau, les particules sont contraintes à se déplacer sur un réseau régulier avec une vitesse nulle ou de module constant. Les interactions entre les particules sont schématisées par des collisions se produisant aux noeuds du réseau considéré, et suivant des règles simples. Ses règles sont choisies de telle manière que les lois de conservation du nombre de particules, de la quantité de mouvement et de l'énergie soient respectées. Le premier modèle de gaz sur réseau a été introduit, au début des années 70, pour un réseau carré par Hardy, de Pazzis et Pomeau [10]. En 1985, Frish, Hasslacher et Pomeau ont présenté un modèle similaire sur un réseau hexagonal [11] et ont montré qu'il possédait des propriétés meilleures pour simuler la mécanique des fluides tant que la vitesse du fluide est petite par rapport à la vitesse du son. Ce modèle a été simulé et testé avec succès pour plusieurs écoulements bidimensionnels [12] et tridimensionnels [13]. La figure II-1 montre le résultat de cette simulation pour l'écoulement d'un fluide autour d'un obstacle [14]. Les conditions aux limites de cet exemple sont périodiques dans la direction verticale, tandis que les particules sont injectées sur le bord gauche avec une vitesse constantes et absorbées sur le bord droit. Le nombre de noeuds utilisé est de $1024 * 512$, chaque flèche représente la direction et le module du flux aux noeuds d'un maillage de $32 * 64$ noeuds de traitement. Le nombre de Reynolds est de 90, et ce résultat est obtenu après 7200 itérations.

II-2 PRINCIPE DE L'ALGORITHME

Comme nous venons de le mentionner, les particules du gaz sont contraintes à se déplacer sur un réseau hexagonal avec une vitesse de module constant ou nulle.

A l'instant "t", chaque noeud du réseau voit arriver de ses voisins un certain nombre de particules, 6 au maximum compte-tenu de la symétrie hexagonale du réseau, avec éventuellement une particule immobile en son centre. Ces particules entrent en collision. Les particules résultant de la collision ont une configuration déterminée en fonction des particules entrantes. Ces particules résultant de la collision se propagent vers les noeuds voisins pour constituer les particules incidentes à l'instant "t+1".

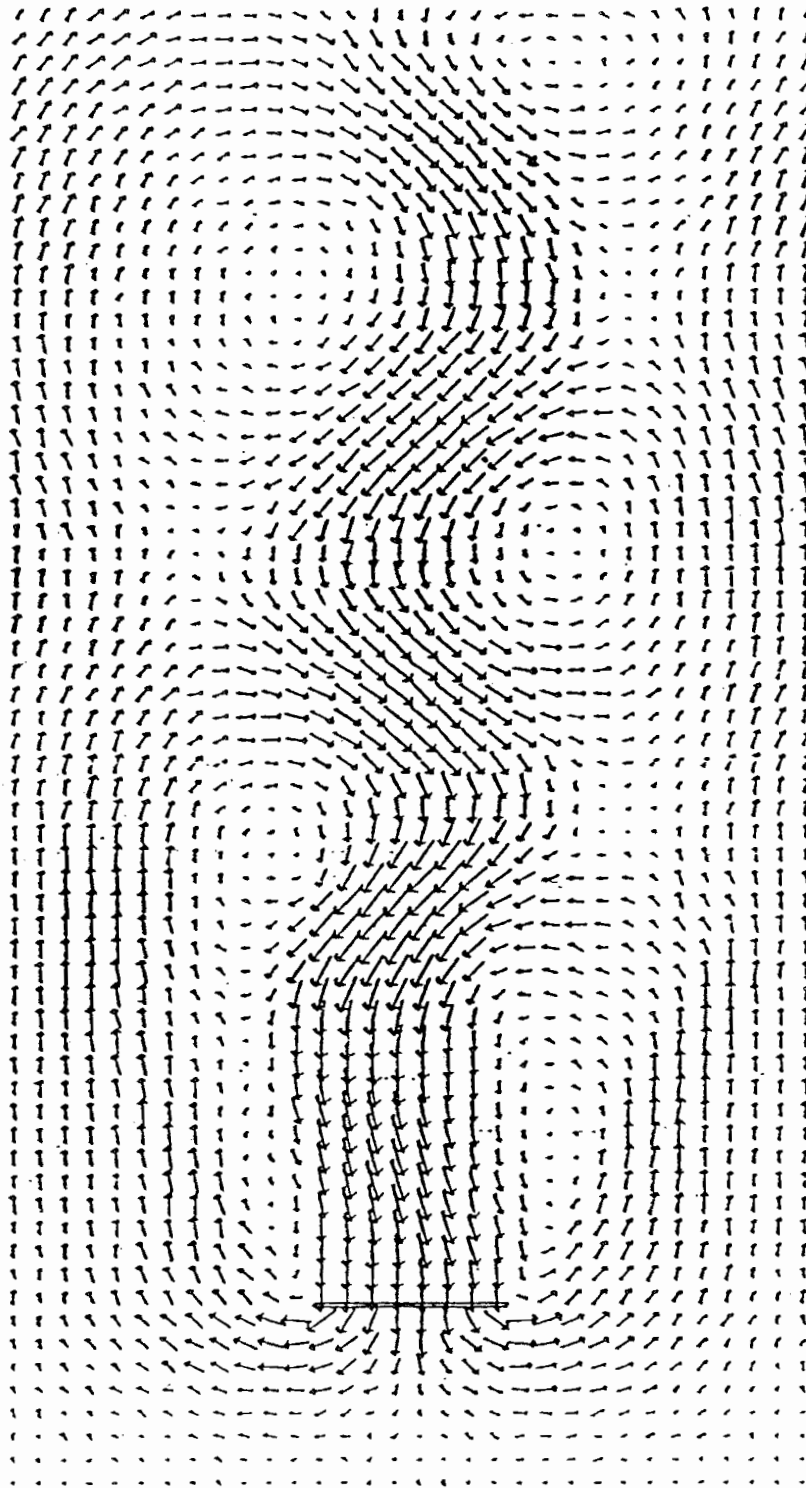


Fig. II-1 : Ecoulement autour d'une plaque mince en mouvement dans le gaz simulé par "Gaz sur Réseau" pour un nombre de Reynolds de l'ordre de 90. Les flèches représentent la direction et le module du flux correspondant aux noeuds d'un maillage 32*64.

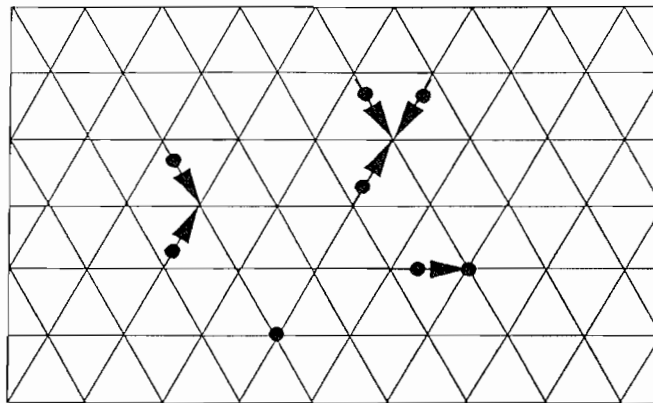


Fig. II-2 : Distribution des particules dans le "Gaz sur réseau"

Nous appelons une itération le traitement consistant à faire la collision des particules incidentes pour tous les noeuds du réseau et la propagation des particules résultant de la collision vers les noeuds voisins. Pour avoir des résultats significatifs de la simulation de l'écoulement, il faut faire, dans les cas simples, plusieurs milliers d'itérations sur plusieurs centaines de milliers de noeuds.

II-2-1 Règles de collision

Les règles de collision sont présentées dans le tableau II-1 qui donne la configuration des particules résultant du choc en fonction des particules incidentes à chaque noeud.

Nous discutons par la suite les différentes lois de conservation à la lumière de ces règles :

i - Conservation de la masse :

Dans le cas le plus simple, les particules du fluide sont identiques et elles ont toutes la même masse. Par conséquent, la loi de conservation de la masse se réduit à la conservation du nombre de particules mises en jeu dans l'ensemble du gaz. Cette conservation est assurée, selon le tableau II-1, par le fait que le nombre des particules incidentes, ou immobiles, à chaque noeud est égal au nombre des particules sortantes, ou restant immobile, du même noeud.

ii - Conservation de la quantité de mouvement :

Compte-tenu que la masse des particules est la même, et la vitesse des particules mobiles est à module constant, la quantité de mouvement est conservée au niveau de chaque noeud et à chaque collision par les règles de collision.

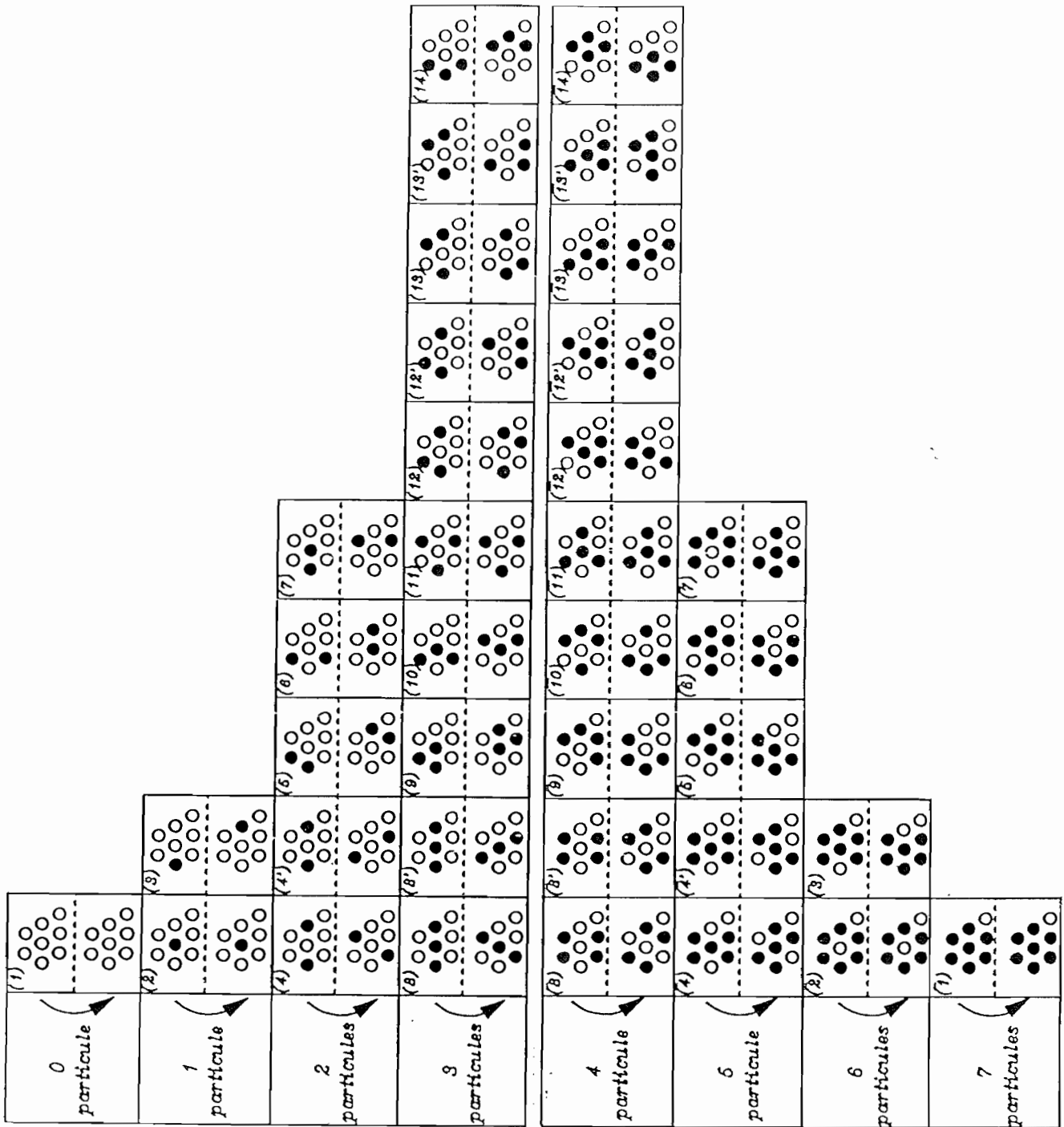


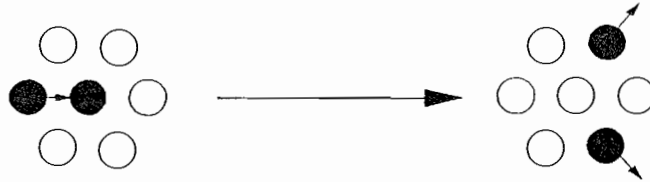
Tableau II-1 : Règles de collision pour le "Gaz sur réseau".

Les configurations des particules sont présentées à une rotation de $\pi/6$ près.

Dans chaque case la configuration de haut représente les particules incidentes, et celle de bas la configuration après choc correspondante.

iii- Conservation de l'énergie

Si nous considérons que l'énergie du gaz est purement cinétique, la conservation de l'énergie n'est pas assurée au niveau de chaque noeud pour certaines collisions qui mettent en jeu des particules immobiles comme la collision présentée suivante :



Dans ce cas la conservation de l'énergie est assurée statistiquement, par des collisions complémentaires. Pour l'exemple précédent la collision complémentaire est donnée comme suit :



II-2-2 Le bit aléatoire

Pour le traitement du gaz sur réseau, chaque noeud est représenté par un mot de 7 bits égal au nombre de particules présentes. Certaines collisions, toutefois, présentent deux possibilités pour la configuration des particules résultant du choc. Pour pouvoir trancher entre ces deux possibilités, il est nécessaire d'ajouter un 8^{ème} bit qui prend des valeurs aléatoires et détermine la configuration à prendre, comme le montre la figure suivante :

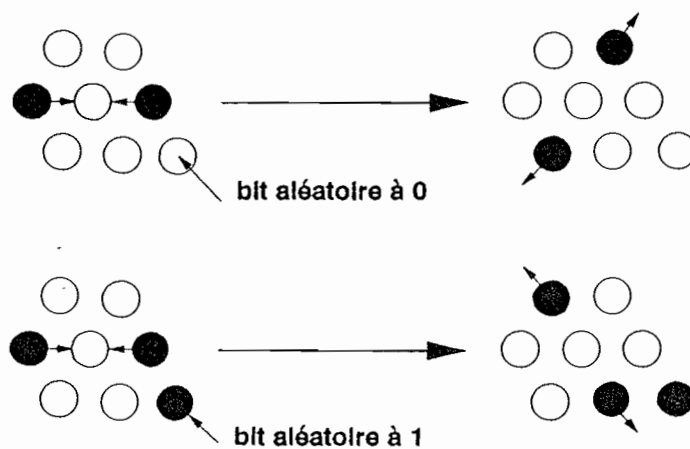


Fig. II-3 : Le bit aléatoire.

L'introduction de ce 8^{ème} bit ne sera pas traitée dans le cadre de cette thèse. En fait, dans le cadre de cette thèse nous testons un motif élémentaire pour une seule cellule de traitement sans aborder le traitement complet de l'algorithme.

II-2-3 Présence de parois ou d'obstacles

Il est évident que les écoulements intéressants sont ceux qui se font à la limite d'une paroi ou en présence d'un certain obstacle. La présence de tels parois est schématisée par des noeuds du réseau interdits à la présence des particules et par des lois particulières de collision sur la frontière de l'ensemble de ces noeuds.

II-3 METHODES DE TRAITEMENT

Le traitement de la collision consiste à remplacer la configuration des particules incidentes à chaque noeud par la configuration des particules résultantes de la collision selon les règles présentées dans le tableau II-1.

Nous utilisons pour effectuer le traitement la méthode de substitution symbolique. La substitution symbolique, introduit par A. Huang [15], est une méthode de calcul utilisant le remplacement de configurations binaires sur une matrice bidimensionnelle. Son traitement consiste à exécuter deux étapes :

- La première est la reconnaissance d'une certaine configuration (particules incidentes dans notre exemple), dans les éléments de la matrice (noeud du réseau).
- La deuxième est la substitution, dans les endroits où la première configuration a été reconnue, d'une autre configuration (particules après choc).

La substitution symbolique est adapté aux architectures de calcul optiques sur des matrices bidimensionnelles comme le cas des automates cellulaires. En fait, la notion de substitution symbolique est équivalente à deux automates cellulaires, une "élémentaire" pour la reconnaissance et l'autre pour la substitution [1].

Pour exécuter les étapes de reconnaissance et de substitution, les différentes configurations sont enregistrées dans un tableau appelé "tableau de correspondance". Deux méthodes existent pour explorer ce tableau :

- La méthode "flash" : le tableau de correspondance est une mémoire organisée de telle manière que la configuration à reconnaître (particules incidentes) sert comme adresse de la

ligne de la mémoire qui contient la configuration substituante (particules après choc) qui lui correspond. De cette façon le résultat de la collision est obtenu pour chaque noeud en un seul accès.

- La méthode "défilement" : la mémoire du "tableau de correspondance" contient les configurations à reconnaître, chacune est suivie de la configuration substituante. Ces différentes configurations sont défilées devant les cellules de traitement, chaque cellule reconnaît configuration correspondante à son état et lui substitue la configuration suivante reçue. L'exécution par cette méthode prend un aspect séquentiel.

Nous optons pour la méthode de défilement à cause de l'économie de place et la possibilité du traitement parallèle (voir chapitre III).

II-4 ETAPES DE TRAITEMENT SELON LA SUBSTITUTION SYMBOLIQUE

La méthode de substitution symbolique appliquée à l'algorithme "Gaz sur réseau" consiste à exécuter trois étapes, la reconnaissance-substitution à chaque noeud du réseau et le déplacement des particules entre les noeuds du réseau :

1- la reconnaissance : Reconnaître la configuration des particules incidentes à l'instant "t". Cette reconnaissance peut se faire par la comparaison de la configuration incidente aux différentes possibilités. Dans le cas d'un réseau hexagonal le nombre des particules incidentes étant de 7 le nombre de possibilité est de $2^7 = 128$.

2- La substitution : Une fois la configuration des particules incidentes est reconnue, il faut lui substituer la configuration des particules résultant de la collision.

3- La propagation : une fois la reconnaissance et la substitution sont traitées pour tous les noeuds du réseau, il faut propager les particules résultant du choc aux noeuds voisins pour préparer la nouvelle itération. Cette opération concerne les plus proches voisins; la particule résultante se déplace au noeud voisin immédiat.

La figure suivante résume les différentes étapes du traitement du gaz sur réseau selon la méthode de substitution symbolique.

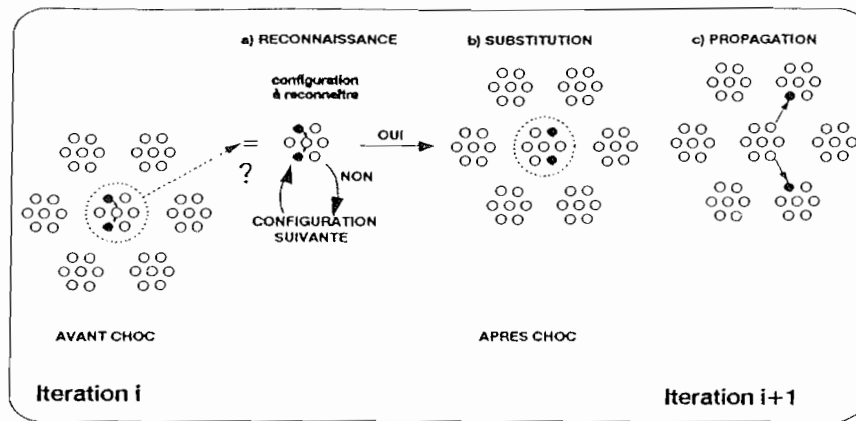


Fig.II-4: "Gaz sur réseau" : Etape du traitement.

II-5 GAZ SUR RESEAU ET AUTOMATE CELLULAIRE

Le traitement au niveau des différents noeuds de l'algorithme "gaz sur réseau" lui donne un aspect de traitement cellulaire où chaque noeud de collision est une cellule de traitement.

Par conséquent, chaque cellule a un état $[c_{ij}]$, un voisinage V_{ij} , et une fonction d'évolution F qui est la même pour toutes les cellules. Les effets de bord ou la présence d'obstacle nécessitent l'introduction de masque pour interdire la présence de particules dans certaine zone du gaz et des lois spécifiques de collision à la frontière de ce masque qui peuvent être traitées dans le cas général tout en marquant les cellules de la frontière pour ne recevoir que les intructions qui les concernent.

Nous nous limitons dans le cadre de cette thèse au cas d'un gaz libre infini ou périodique sans la présence d'obstacle, et nous précisons par la suite les différents éléments : état de cellule, voisinage et fonction d'évolution.

II-5-1 Présentation du réseau

Compte-tenu de la géométrie hexagonale du gaz sur réseau, nous adoptons cette géométrie pour l'automate cellulaire par souci de démonstration. Les différentes cellules dans ce réseau sont indexées comme cela est présenté sur la figure suivante:

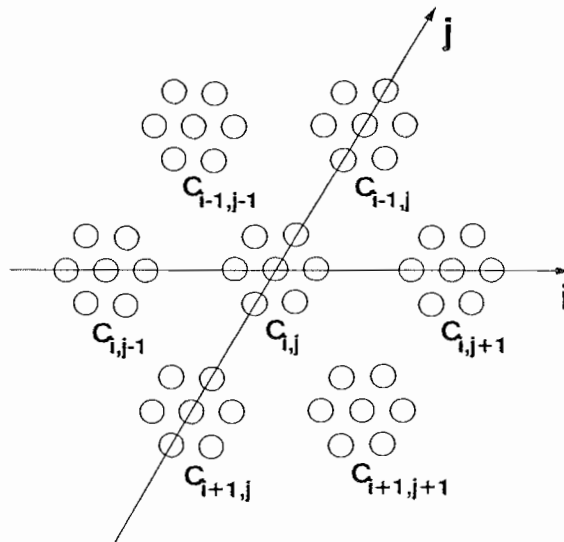


Fig. II-5 : Indexation des cellules.

II-5-2 Etat de chaque cellule $[c_{ij}]$

L'état de la cellule c_{ij} est la configuration des particules résultant de la collision au noeud (i,j) .

Par conséquent, $[c_{ij}]$ est un mot binaire de 7 bits $\{e_{ij}^0, e_{ij}^1, \dots, e_{ij}^6\}$. Le bit e_{ij}^0 correspond à la particule immobile éventuellement présente au centre de ce noeud, les autres bits correspondent aux différentes directions du réseau comme cela est présenté sur la figure suivante :

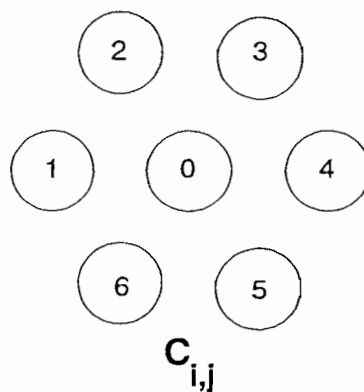


Fig. II-6 : Notation des bits dans chaque cellule.

L'état 1 d'un bit signifie la présence d'une particule résultant du choc et se propageant dans la direction correspondante, et l'état 0 signifie l'absence d'une telle particule.

Par exemple l'état $[c_{ij}] = \{0,0,0,1,0,1,0\}$ représente deux particules sortant de la cellule c_{ij} en se propageant dans les directions indiquées dans la figure suivante :

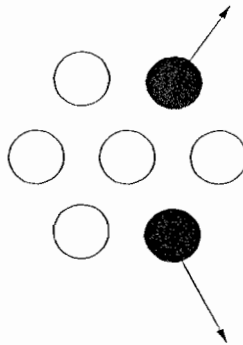


Fig. II-7 : Etat d'une cellule.

II-5-3 Voisinage V_{ij}

Comme cela a été mentionné précédemment, une particule sortant d'une cellule se propage vers la cellule voisine la plus proche. Ainsi, chaque cellule n'est en relation qu'avec ses plus proches voisines au nombre de 6, et avec elle même dans le cas d'une particule résultante immobile en son centre.

Le voisinage V_{ij} de la cellule c_{ij} est constitué de l'ensemble suivant :

$$V_{ij} = \{c_{i,j}; c_{i,j-1}; c_{i,j+1}; c_{i-1,j}; c_{i-1,j-1}; c_{i+1,j}; c_{i+1,j+1}\}$$

II-5-4 Fonctions de contribution $\Phi_{i-k,j-l}$

Chaque noeud du réseau agit sur le noeud voisin par une seule particule qu'il envoie, éventuellement, dans la direction de ce noeud. Par conséquent, la fonction de contribution $\Phi_{i-k,j-l}$ extrait de l'état de la cellule voisine le bit correspondant à la direction de la cellule traitée. Cette fonction prend les valeurs suivantes en fonction de $i-k$ et $j-l$:

$$\begin{aligned}
\Phi_{0,0}(c_{ij}) &= e_{ij}^0 \\
\Phi_{0,-1}(c_{ij-1}) &= e_{ij-1}^4 \\
\Phi_{-1,-1}(c_{i-1,j-1}) &= e_{i-1,j-1}^5 \\
\Phi_{-1,0}(c_{i-1,j}) &= e_{i-1,j}^6 \\
\Phi_{0,+1}(c_{ij+1}) &= e_{ij+1}^1 \\
\Phi_{+1,+1}(c_{i+1,j+1}) &= e_{i+1,j+1}^2 \\
\Phi_{+1,0}(c_{i+1,j}) &= e_{i+1,j}^3
\end{aligned}$$

II-5-5 Fonction d'évolution F

La fonction d'évolution F détermine la configuration des particules résultant du choc en fonction de la configuration des particules incidentes. Dans le cas d'absence de parois, cette fonction est la même pour toutes les cellules de traitement. Les valeurs de cette fonction sont données dans le tableau II-1 de ce chapitre.

CHAPITRE III

CONCEPTION DE L'AUTOMATE CELLULAIRE

CHAPITRE III

CONCEPTION DE L'AUTOMATE CELLULAIRE

Nous avons démontré dans le chapitre précédent que l'algorithme "Gaz sur Réseau" est exécutable sur un automate cellulaire parallèle.

Nous présentons dans ce chapitre d'une façon succincte le premier automate cellulaire destiné à traiter cet algorithme. Ensuite, nous définissons les paramètres nécessaires régissant la conception d'un tel automate, et nous proposons la conception de l'automate en projet en décrivant ses différents éléments.

III-1 AUTOMATE CELLULAIRE POUR "GAZ SUR RESEAU"

L'avantage lié à l'implantation de l'algorithme "Gaz sur réseau" sur un automate cellulaire est représenté par le gain en temps d'exécution. Ce gain est potentiellement dû non seulement à la réalisation d'une machine dont l'architecture est adaptée au problème, mais également au fait qu'il soit possible d'employer un traitement parallèle de toutes les cellules.

Le premier automate cellulaire connu pour le "Gaz sur Réseau" (R.A.P.1) a été réalisé par A. Clouqueur et D. d'Humières de l'E.N.S.[16]. Cette machine est de type électronique et contient 256 lignes chacune composée de 512 cellules chacune. Chaque cellule est définie sur 16 bits. Le traitement se fait séquentiellement au niveau des cellules. La reconnaissance et la substitution se fait en utilisant une table de correspondance à 16 entrées et 16 sorties. La propagation se fait par glissement des plans, chaque plan correspond à une direction de propagation. une itération prend 20 ms ce qui correspond à une vitesse de traitement de $6.5 \cdot 10^6$ sites par seconde. Cette vitesse est à comparer avec celle obtenue sur un ordinateur de type CRAY1: $30 \cdot 10^6$ opérations par seconde. La réalisation de cette machine n'a nécessité que l'utilisation de 10 cartes en circuit imprimé double-face et 'wrapping'. Cette machine démontre bien que le gain en vitesse de traitement est liée à la fabrication de machines spécialisées.

La possibilité d'implanter un automate cellulaire optique pour le "Gaz sur Réseau" à partir des plaques photographiques a été étudiée dans le cadre d'une architecture parallèle [17].

D'une façon générale, un automate destiné à traiter le gaz sur réseau doit contenir un grand nombre de cellules de traitement comparable au volume de R.A.P.1. L'intégration de cet automate sur un nombre limité de "puces" peut diminuer considérablement le temps de calcul, ainsi que l'énergie consommée. Cependant un autre problème peut limiter les performances d'une telle machine électronique : il est représenté par la haute densité des interconnexions qu'il est nécessaire de réaliser au niveau des entrées-sorties et de la distribution des instructions sur les "puces". L'utilité de l'optique pour réaliser un tel automate est représentée par sa capacité à réaliser des interconnexions à haute densité et haut débit comparables aux performances et au temps de commutation des composants électroniques intégrés.

Le problème actuel de la technologie optique est la difficulté de réaliser des bistables optiques avec des performances comparables aux transistors électronique intégrés. De ce fait, nous essayons de profiter de la performance de la technologie électronique pour réaliser les unités de traitement et de la performance de l'optique pour réaliser les connexions. Tout ceci dans un prototype de démonstration destiné au traitement du "Gaz sur Réseau".

Pour fixer les idées, un automate cellulaire efficace doit contenir un nombre comparable au $1.3 \cdot 10^5$ unités de traitement de la machine R.A.P.1, pour un parallélisme complet du traitement. Nous verrons que l'implantation de 10^2 à 10^3 sites de traitement est possible sur une puce électronique de 1 cm^2 de surface. La connexion individuelle de chacune de ces sites avec l'extérieur est presque impossible électroniquement vu le nombre et la densité des fils électriques à utiliser. Une connexion optique permettrait de faire communiquer individuellement ces sites avec l'extérieur en utilisant la troisième dimension. De ce fait une parallélisation complète du traitement sera possible et un temps d'exécution d'une itération d'environ $1 \mu\text{s}$ est envisageable.

III-2 TRAITEMENT PARALLELE ET SEQUENTIEL

L'automate cellulaire est constitué d'un grand nombre de cellules de traitement. Pour chacune de ces cellules, il faut exécuter un certain nombre d'opérations qui dépend de son état et de l'état des cellules voisines. Ce traitement est le même pour toutes les cellules dans le cas de l'automate cellulaire. Par conséquent, son extension à toutes les cellules relève un aspect répétitif que nous pouvons éviter par un traitement parallèle, expression dont le sens peut s'opposer au traitement séquentiel. Donc, il est nécessaire de préciser ces deux notions avant d'aborder la conception de l'automate.

III-2-1 Traitement séquentiel

Le traitement séquentiel, par définition, n'exécute qu'une seule opération à la fois. Cette opération s'effectue sur un mot ayant une longueur déterminée : 8, 16, 32 bits ...etc. Dans le cas de l'automate cellulaire, le traitement séquentiel traite les cellules l'une après l'autre en faisant un balayage sur le réseau. Pour chaque cellule, un certain nombre d'opérations est exécuté. Ce type d'automate est en général constitué d'une unité centrale de traitement, qui peut être aussi compliquée que nécessaire, et d'une mémoire qui représente les différentes cellules. L'unité centrale utilise les données de la mémoire de chaque cellule et y affecte le résultat. C'est le cas, par exemple de la machine R.A.P.1.

Il est évident que la durée de l'exécution d'une itération est une fonction linéaire du nombre de cellules constituant l'automate.

III-2-2 Traitement parallèle

Le traitement parallèle exécute les opérations simultanément sur toutes les cellules de l'automate. Dans ce cas les unités de prise de décision et d'évolution doivent être liées à chaque cellule. L'unité centrale, en cas où elle est présente, se charge de l'initialisation, du recueil des résultats et du séquençement de fonctionnement de l'ensemble des cellules.

Le grand avantage du traitement parallèle est le gain en temps d'exécution qui, dans ce cas, ne dépend pas du volume de l'automate.

Les opérations traitées pour chaque cellule sont d'une nature différente de celles traitées sur une machine séquentielle. En fait l'unité de traitement liée à chaque cellule ne peut pas être de même complexité que l'unité centrale dans la machine séquentielle, sinon la machine parallèle sera extrêmement compliquée. Nous sommes obligés de simplifier ces unités au maximum, pour pouvoir implanter le maximum de cellules. Cette simplification ne peut être effectuée qu'au détriment du nombre d'opérations exécutables par chaque cellule. Le paragraphe suivant montre la différence qui existe entre les opérations traitées sur une cellule selon le traitement séquentiel et parallèle.

En résumé la différence entre le traitement parallèle et séquentiel est le fonctionnement simultané ou non de toutes les cellules.

III-3 CONCEPTION DE L'AUTOMATE CELLULAIRE OPTO-ELECTRONIQUE

La conception de l'automate doit tenir compte des différentes fonctions nécessaires au bon déroulement de l'algorithme, ainsi que des performances des différentes technologies pour la réalisation la plus rapide.

L'automate cellulaire contiendra des cellules de traitement distribuées sur les noeuds d'un réseau hexagonal, Cette option géométrique n'est pas fondamentalement nécessaire mais elle est plus démonstrative. Ces cellules travaillent en parallèle. La réalisation de ces unités de traitement est faite électroniquement sur des circuits intégrés VLSI. Pour implanter le maximum de cellules sur une puce électronique, il est important que l'encombrement de chaque cellule soit minimal, ce qui conduit à réduire les fonctions de chaque cellule aux opérations indispensables de stockage d'information et de comparaison.

III-3-1 Les différentes fonctions de l'automate

Nous définissons par la suite les fonctions nécessaires pour le fonctionnement de l'automate chargé de traiter l'algorithme "gaz sur réseau" :

- i - **La représentation de l'état du gaz** : mémoriser pour chaque cellule du réseau la configuration des particules. Autrement dit, la détermination de l'état de chaque cellule de traitement $[c_{ij}]$.
- ii - **L'initialisation de l'automate** : imposer la distribution initiale des particules sur les noeuds du réseau ($[c_{ij}(0)]$).
- iii - **Le traitement au niveau de chaque cellule** : exécuter les opérations de reconnaissance et de substitution pour déterminer la configuration des particules résultant du choc.
- iv - **La propagation** : propager les particules résultant de la collision à chaque noeud vers les noeuds voisins.
- v - **La lecture de l'état du gaz** : savoir à un instant donné quelle est la distribution des particules résultant de la collision sur les noeuds du réseau ($[c_{ij}(t)]$).

III-3-2 L'implantation des différentes fonctions

Nous proposons dans la suite une solution pour chacune de ces fonctions ce qui nous permettra de définir la structure de l'automate.

a) La représentation de l'état du gaz :

Nous allons voir que l'exécution de chaque itération au niveau de l'automate comportera un certain nombre d'opérations et de séquences. Le long de chaque itération l'état de l'automate ne doit pas changer, autrement dit la distribution des particules sur les différents noeuds doit rester fixe. D'où la nécessité de mémoriser la configuration des particules dans une mémoire liée à chaque noeud. Comme les bistables optiques ne sont pas encore adaptés

pour réaliser des mémoires à haute densité d'intégration, nous avons recours aux mémoires électroniques pour cette mémorisation.

Ainsi, chaque cellule de traitement contiendra deux mémoires à 7 bits, l'une pour les particules incidentes au noeud correspondant à cette cellule (mémoire d'entrée), et l'autre pour les particules résultantes de la collision effectuée à cette cellule (mémoire de sortie). Chaque bit de ces deux mémoires correspond à une particule se propageant dans une direction.

b) L'initialisation de l'automate :

L'initialisation revient à affecter les mémoires électroniques d'entrée des différentes cellules par les valeurs initiales.

Si nous procédons par une initialisation électronique nous avons deux possibilités :

- L'initialisation séquentielle qui demande un nombre réduit d'entrées mais des multiplexeurs pour aiguiller l'information vers les différentes mémoires. Cette solution complique considérablement la conception du circuit électronique.
- L'initialisation parallèle qui demande un grand nombre de plots d'entrées (7 par cellule de traitement). ce qui est obligatoirement encombrant et non réaliste.

Par contre, nous avons préféré une initialisation parallèle optique en associant à chaque cellule de traitement 7 photodiodes directement liées aux 7 bits de la mémoire d'entrée, et par une projection d'une image sur ces photodiodes, nous pouvons tout de suite affecter les mémoires d'entrée des différentes cellules de traitement. Cette initialisation dite "parallèle" ne l'est en fait qu'au niveau du circuit et non au niveau du système. En réalité, il y a report de l'initialisation séquentielle au système d'imagerie (changement du masque si nous voulons changer les données initiales, ou allumage des LED qui se fait séquentiellement). L'avantage que reporte ce système est la simplification du circuit électronique.

c) Le traitement au niveau de chaque cellule :

Le traitement au niveau de chaque cellule consiste à reconnaître le motif des particules incidentes et à lui substituer le motif des particules résultant du choc. Nous avons vu au chapitre II que ce traitement se fait par la substitution symbolique. Les différents motifs incidents avec leurs substitutions sont enregistrés dans une table de correspondance stockée dans une mémoire centrale appelée "mémoire de correspondance". Cette "mémoire de correspondance" communique avec la "mémoire d'entrée" de chaque cellule pour la reconnaissance et affecte le résultat à la "mémoire de sortie".

Nous distinguons deux modes de communication entre la "mémoire de correspondance" et les mémoires attachées à chaque cellule : parallèle et séquentiel. Les notions de parallèle et de séquentiel sont ici propres aux différents motifs de particules, nous les appelons "modes parallèles ou séquentiels" pour les distinguer de celles de "traitements parallèle et séquentiel" traitées dans le paragraphe précédent et qui concernent l'ensemble des cellules. Nous allons traiter ces deux notions et la possibilité de leurs implantations.

Le mode parallèle : Dans ce mode, la reconnaissance et la substitution se font par une seule opération. Le contenu de la "mémoire d'entrée" de chaque cellule sert comme adresse à la "mémoire de correspondance" (voir fig. III_1).

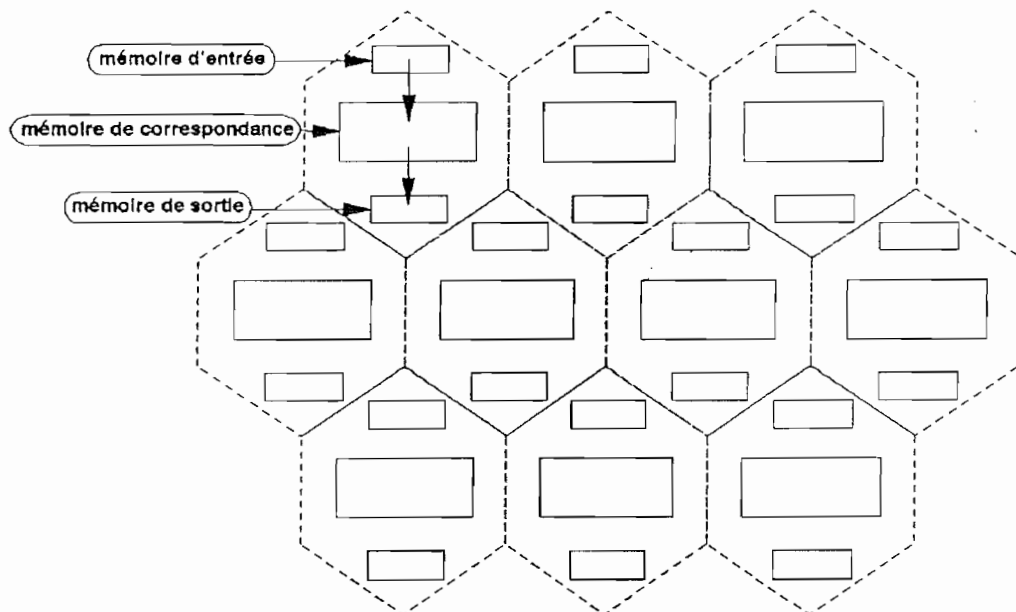


Fig. III-1 : Mode parallèle de traitement.

Si la "mémoire de correspondance" se trouve à l'extérieur des cellules de traitement, elle ne peut être adressée que par une seule cellule à la fois, et nous nous trouvons avec un "traitement séquentiel" au niveau du fonctionnement de l'automate. Pour assurer un "traitement parallèle" au niveau du fonctionnement de l'automate, il faut associer une "mémoire de correspondance" complète de 128 mots à chaque cellule de traitement ce qui augmente l'encombrement de chaque cellule dans un rapport inacceptable, d'autant plus que toutes les "mémoires de correspondance" de toutes les cellules ne sont que la reproduction d'un même motif. Ce mode est le plus parallèle possible au niveau des cellules, parallèle dans chaque cellule au niveau des motifs reconnus, mais il est pénalisé par une densité de cellules par centimètre carré trop faible. Nous ne le considérons pas dans la suite.

Le mode séquentiel : Dans ce cas nous faisons défiler les différents motifs des particules incidentes suivis des motifs substituants. Chaque motif est comparé avec la "mémoire d'entrée" de chaque cellule. En cas de correspondance, le motif substituant est placé dans la "mémoire de sortie" de cette cellule, sinon nous suivons le défilement des motifs. La "mémoire de correspondance" se trouve à l'extérieur du réseau des cellules (d'où l'économie de place), et son contenu est lu et distribué séquentiellement aux différentes cellules. Une distribution simultanée à toutes les cellules est possible, donc un "traitement spatialement parallèle" (et séquentiel dans le temps, est possible (voir fig. III-2).

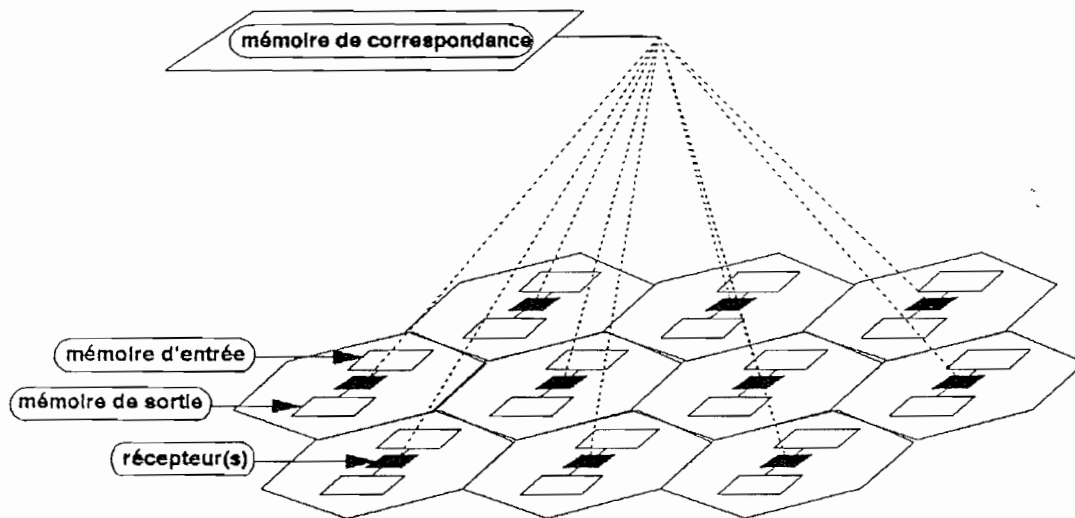


Fig. III-2 : Mode séquentiel de traitement.

Nous optons pour le "mode séquentiels". Avec cette solution la mémoire contenant la table de correspondance se trouve à l'extérieur du circuit électronique contenant les cellules de traitement. Cette "mémoire de correspondance" contient les différents motifs des particules incidentes, chaque motif suivi du motif substituant correspondant. Pour distribuer le contenu de cette mémoire aux différentes cellules de traitement, nous profitons de la faculté que présente l'optique dans l'utilisation de la troisième dimension, et nous projetons par voie optique et séquentiellement dans le temps le contenu de chacun des 128 mots de cette mémoire sur toutes les cellules de traitement. Chaque cellule est équipée de 7 photodiodes (les mêmes pour l'initialisation) pour recevoir les 7 bits de chacun de ces différents motifs. Chaque cellule recevant optiquement la configuration d'entrée, la compare avec le contenu de sa "mémoire d'entrée" et décide la reconnaissance ou non. Si la configuration projetée est reconnue dans la "mémoire d'entrée", la configuration substituant, qui est projetée juste après, est transférée dans la "mémoire de sortie" de cette cellule, sinon la cellule attend les configurations suivantes.

Comme la projection des différentes configurations est possible sur toutes les cellules en même temps, le traitement parallèle des cellules sera assuré. Mais à ce moment nous aurons besoin d'un dispositif optique permettant la projection des différentes configurations simultanément sur l'ensemble des cellules de traitement.

d) La propagation :

La propagation consiste à déplacer les particules résultant du choc dans chaque cellule vers les cellules voisines pour préparer une nouvelle itération. Comme le motif des particules résultantes se trouve physiquement dans la "mémoire de sortie", qui n'est qu'une mémoire tampon classique, il suffit de connecter électriquement chaque bit de cette mémoire au bit de la "mémoire d'entrée" de la cellule voisine. Un signal de séquençement (horloge) déclenche le transfert du contenu des mémoires à la fin de chaque itération.

e) La lecture de l'état du gaz :

Le problème de lecture est similaire à celui de l'initialisation, à savoir le problème de sortie électronique séquentielle qui demande des multiplexeurs ou la sortie électronique parallèle qui demande beaucoup des plots. En fait nous avons besoin de réaliser un contact point à point entre les bits des "mémoires de sortie" de l'ensemble des cellules et un dispositif extérieur pour recevoir les résultats et les analyser. Nous envisageons l'utilisation de moyens opto-électroniques pour extraire optiquement le contenu des mémoires de sortie des différentes cellules en parallèle. Il est possible d'utiliser des modulateurs opto-électronique pour l'extraction de ces résultats. Un dispositif à puits quantiques multiples a été fabriqué au Laboratoire Central de Recherche de Thomson C.S.F. dans ce but.

III-4 PRESENTATION DE L'AUTOMATE CELLULAIRE OPTO-ELECTRONIQUE

A la lumière de la discussion du paragraphe précédent, l'automate cellulaire opto-électronique sera composé des éléments suivants (voir fig. III-3)[18,19] :

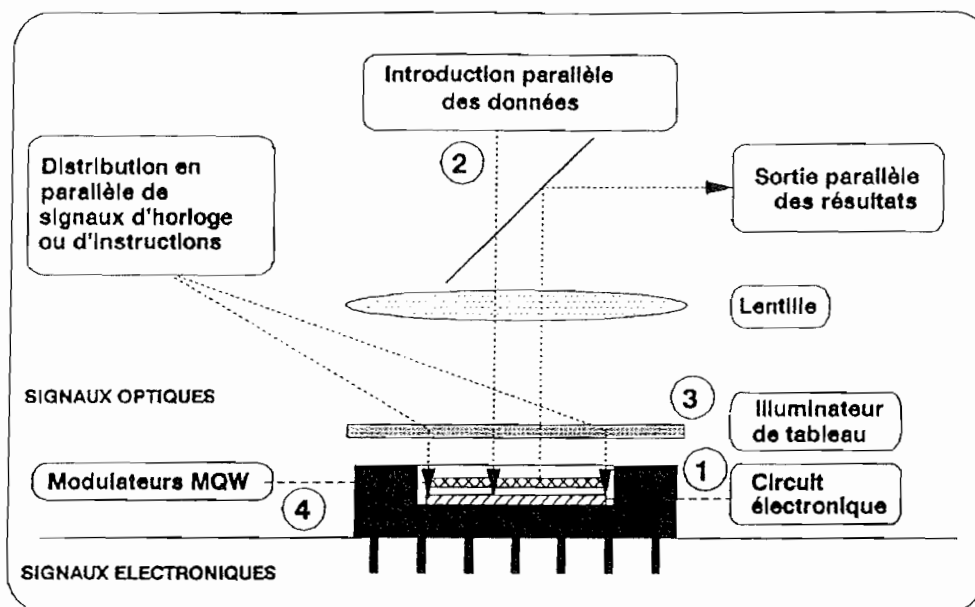


Fig. III-1 : Schéma final de l'automate cellulaire opto-électronique

1- Le circuit électronique

Ce circuit contient des processeurs élémentaires distribués sur un réseau hexagonal. Chaque processeur élémentaire est l'implantation d'un noeud de collision dans le réseau et comporte une "mémoire d'entrée" de 7 bits et une "mémoire de sortie" de 7 bits pour la configuration des particules. Il assure le traitement de reconnaissance et substitution à l'aide des instructions projetées optiquement sur des photodiodes. un réseau de connexion fixe entre les processeurs élémentaires assure la propagation des particules (voir chapitre V).

2- Le système optique d'initialisation

Ce système projette sur les photodiodes de l'ensemble des processeurs élémentaires la distribution initiale des particules.

3- L'illuminateur de tableau

L'illuminateur de tableau duplique l'image d'une matrice de source lumineuses en plusieurs images et les projette en parallèle sur les différents processeurs élémentaires (voir chapitres IV et V).

4- Les modulateurs opto-électroniques pour la lecture des résultats

Ces modulateurs permettent de transformer le contenu des mémoires de sortie de l'ensemble des processeurs élémentaires en une image que nous projetons sur un dispositif extérieur pour l'analyse (voir chapitre VII).

III-5 CONCLUSION

Nous avons présenté dans un chapitre la conception de l'automate cellulaire où le traitement se fait sur un circuit électronique avec des entrées sorties et distribution d'instructions optiques. Le traitement se fait par la substitution symbolique.

DEUXIEME PARTIE

HOLOGRAMME A EFFET TALBOT

CHAPITRE IV

HOLOGRAMME E EFFET TALBOT

CHAPITRE IV

HOLOGRAMME A EFFET TALBOT

IV-1 INTRODUCTION

Nous avons montré, dans le chapitre précédent, l'intérêt qu'il y a de disposer d'un système optique ayant pour fonction de distribuer un signal optique sur un grand nombre de sites correspondant aux unités de traitement dans l'automate cellulaire.

Un tel dispositif optique est en fait nécessaire dans le cadre général des applications optiques dans l'ordinateur et d'interconnexion.

Les nouvelles conceptions des processeurs de traitement d'informations sont orientées vers des architectures massivement parallèles [20]. Ces architectures sont basées sur l'utilisation des matrices pixellisées de composants non-linéaires optiques ou opto-électroniques. Chaque pixel reçoit optiquement certaines informations à traiter. Il est souvent question d'envoyer la même information à tous les pixels de la matrice. Cette information est alors appelée "signal d'horloge" ou "code d'instruction".

Pour éclairer les éléments de ces matrices par le faisceau lumineux porteur de l'information, nous devons concentrer ce faisceau sur les différents pixels pour : d'une part conserver l'énergie lumineuse, et d'autre part éviter l'éclairage des régions entre les pixels. Un éclairage direct en utilisant une onde plane ou sphérique issue d'un laser implique souvent une grande perte de l'énergie lumineuse. Nous prenons par exemple, le problème d'éclairage d'une rétine composée des photodiodes distribuées sur un réseau carré de 300 μm de période, chaque photodiode est de section carré de 50 μm de côté. Cette rétine contient en plus des unités de traitement électroniques situées entre les photodiodes [47]. En utilisant une onde uniforme (fig. 1), seulement $(50 * 50 / 300 * 300) = 3 \%$ de l'énergie lumineuse est utilisée, et le reste (97 %) est perdu et peut perturber le fonctionnement des circuits annexes en éclairant les transistors implantés entre les photodiodes. Le même calcul appliqué au cas de notre circuit destiné au "Gaz sur Réseau" décrit dans le chapitre VI donne pratiquement le même résultat (1.5 %).

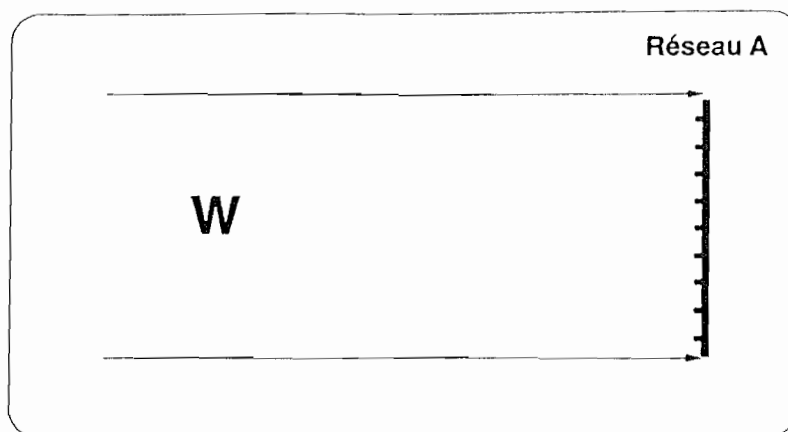


Fig. IV-1 : Eclairage direct des pixels.

Pour cette raison, est apparue une nouvelle génération d'éléments optiques passifs de distribution appelée "Illuminateurs de tableaux". La conception de ces illuminateurs de tableaux a fait appel à plusieurs approches de fabrication qui ont été publiées dans la littérature [22-36]. Nous avons introduit et étudié une nouvelle structure permettant de réaliser un tel illuminateur fondé sur l'emploi de la technique holographique utilisant l'effet d'auto-imagerie de Talbot.

Nous commençons ce chapitre par une définition générale des "Illuminateurs de tableaux", en précisant les principales caractéristiques nécessaires, et nous passons en revue les différents procédés de réalisation existants. Ensuite, nous étudions l'effet Talbot pour les différentes géométries, et nous en déduisons les outils mathématiques nécessaires pour l'étude et la réalisation de l'illuminateur proposé. Après une présentation des techniques holographiques, justifiant le choix des hologrammes épais de phase, nous déterminons les conditions d'enregistrement de l'hologramme à effet Talbot, puis les résultats expérimentaux obtenus. Nous comparons ce procédé avec les autres approches connues pour les illuminateurs de tableaux. Et nous consacrons le chapitre suivant à l'étude des aberrations chromatiques dues au changement de la longueur d'onde entre l'enregistrement et la restitution de l'hologramme.

IV-2 ILLUMINATEUR DE TABLEAUX

IV-2-1 Définition

Un illuminateur de tableaux (AI comme Array Illuminator) est un dispositif optique passif, qui transforme une onde d'éclairage, souvent issue d'un laser, en un nombre d'ondelettes

sphériques, chacune convergente sur un pixel. En d'autres termes, l'illuminateur de tableaux concentre le flux incident sur des petits points séparés.

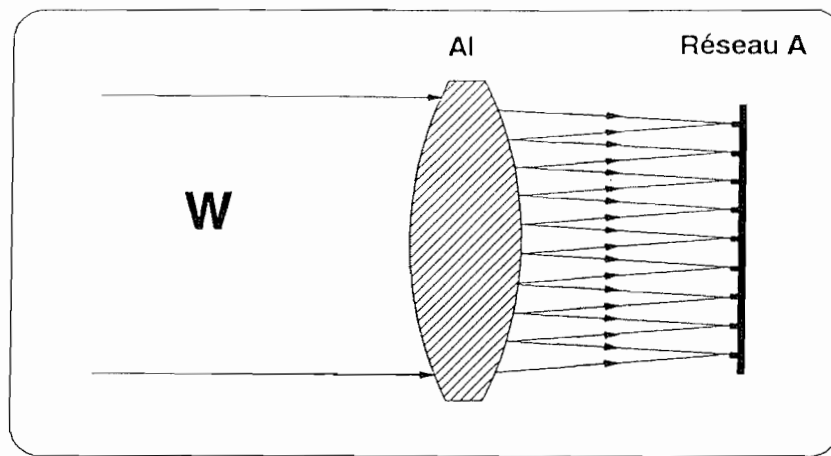


Fig. IV-2 : Illuminateur de tableaux (AI).

Ce même dispositif éclairé par plusieurs faisceaux séparés angulairement éclaire plusieurs réseaux de points séparés spatialement, chaque réseau correspond à un faisceau d'éclairage.

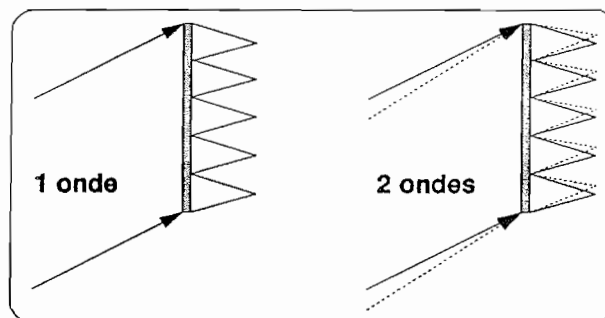


Fig IV-3 : Illuminateur de tableaux avec plusieurs ondes d'éclairage.

IV-2-2 Caractéristiques

L'illuminateur de tableaux doit répondre à plusieurs caractéristiques qui dépendent essentiellement de l'application envisagée. Néanmoins, nous pouvons définir des caractéristiques générales pour tous les illuminateurs de tableaux, et qui sont présentées dans l'article de N. Streibl [21] :

- a) **Le nombre de pixels éclairés** : Le nombre d'éléments opto-électroniques éclairés par un seul illuminateur. Cette caractéristique définit le taux de parallélisation possible et l'encombrement du processeur.
- b) **Taux de compression** : Le rapport entre la surface éclairée et la surface totale de la maille élémentaire du réseau. Ceci donne une indication sur la séparation entre les ondelettes. Un taux de compression de 2 à plusieurs centaines est désirable.
- c) **L'uniformité d'éclairement** : Différence d'intensité concentrée sur les différents pixels. Elle s'exprime en pourcentage de la plus grande différence entre les intensités des pixels par rapport à la moyenne de cette intensité par pixel. Ce paramètre a une importance quand le temps de réponse des composants opto-électroniques occupant les pixels dépend de l'énergie reçue. Elle doit être la plus petite possible pour synchroniser la réponse des différents éléments éclairés.
- d) **Le contrast** : Exprime le rapport entre l'intensité de chaque pixel et la lumière de fond dans les régions non éclairées.
- e) **Le rendement lumineux** : Rapport de l'énergie concentrée sur l'ensemble des pixels à l'énergie totale éclairant l'illuminateur. Ce paramètre est aussi parfois dénommé "efficacité". Il doit évidemment être aussi grand que possible.
- f) **La forme et la superficie des taches lumineuses** : Les dimensions des pixels à éclairer et l'éventuel éclairage en dehors des pixels. Cette caractéristique est importante dans le cas où l'éclairage en dehors des pixels pourrait introduire des perturbations de fonctionnement de la rétine, comme par exemple une rétine électronique contenant des photodiodes et des transistors situés à proximité des photodiodes.
- g) **La cohérence spatiale et temporelle** : Caractérise la lumière utilisée par l'illuminateur : cohérente ou non. Une lumière non cohérente est quelque fois préférable à cause de son rapport signal-bruit élevé.
- h) **La transparence pour d'autres signaux optiques** : La possibilité d'envoyer à la rétine ou d'en extraire d'autres signaux optiques sans qu'ils soient déformés par l'illuminateur. Dans l'automate du "Gaz sur Réseau" nous devons avoir accès à la rétine pour l'initialisation, et nous devons en extraire optiquement les résultats. Ces deux catégories de signaux ne doivent pas être déformés par l'illuminateur, qui est pour sa part consacré à la distribution des "codes d'instructions". Cette caractéristique peut être obtenue soit en

changeant la longueur d'onde, soit en changeant l'angle d'éclairage ou les deux à condition de recourir à des composants dont le principe est sensible à ces paramètres : c'est le cas de certains traitements de surface dits dichroïques dont le taux de réflexion et de transmission varie fortement avec la longueur d'onde ; notre illuminateur en constitue un second exemple.

i) **Le plan focal** : La distance entre le plan des pixels et l'illuminateur. Cette caractéristique est importante pour l'encombrement du système.

j) **La facilité de fabrication** : Ainsi que la sensibilité aux défauts de fabrication.

k) **La sensibilité aux ajustements** : Translation, rotation et vibration.

IV-2-3 Exemples d'illuminateurs

Plusieurs approches pour réaliser des illuminateurs de tableaux ont été proposées. Nous distinguons deux catégories : illuminateurs réfractifs et illuminateurs diffractifs.

a) Illuminateurs de tableaux réfractifs :

Ce sont essentiellement des microlentilles dont les techniques de fabrication sont assez variées [22], et que nous résumons par la suite :

- Technique de l'empreinte mécanique : Agir en forte pression sur un matériau plastique un pinçon constitué d'une matrice munie de microbilles en acier [23]. La fabrication des microlentilles d'un pas de 50 à 400 μm et de focale de 100 μm à 1 mm est possible.

- Technique de l'échange ionique : Modifier localement l'indice de réfraction d'une plaque homogène en faisant diffuser des ions étrangers à l'endroit des microlentilles avec éventuellement un gonflement du matériau [24]. Ainsi des microlentilles de 70 μm de diamètre et de 60 μm de focale ont été réalisées.

- Technique de l'usinage chimique : Cette technique est utilisée pour réaliser des microlentilles de forte ouverture numérique (>0.5). Il s'agit de façonner la surface d'un matériau (InP) par attaque chimique par voie humide [25] ou par voie sèche.

- Dépôt réactif en phase vapeur (CVD) : Après avoir creusé des trous hémisphériques sur un substrat de verre, ces trous sont comblés par le dépôt de SiO_2 et de Si_3N_4 porté à l'état de plasma sous vide [26].

- Technique de formation par photolyse : En irradiant un verre contenant des ions métallique par l'UV suivi d'un traitement thermique, les ions métalliques se réduisent pour former des germes cristalline opaque sauf dans les endroits masqués où se localisent les microlentilles [27].

- Inscription de profils sphérique sur des photorésines : Soit par la formation des surfaces sphériques par la tension superficielle de la résine en phase liquide [28] (diamètre 15 μm , focale 36 μm). Soit par insolation [29].

b) Illuminateurs de tableaux diffractifs :

Les systèmes diffractifs pouvant fonctionner comme illuminateurs de tableaux sont nombreux. Nous essayons de présenter quelques uns de ces systèmes :

- Matrice de microlentilles holographiques : Cette matrice est enregistrée pixel par pixel par interférence d'une onde sphérique avec une onde de référence [30]. Les meilleurs résultats sont obtenus sur hologramme de phase épais avec une efficacité supérieure à 90 % [1]. Le pas des réseaux réalisables est supérieur à 100 μm et le nombre de pixel est inférieur à 100×100 .

Un système de réseau télescopique est proposé en utilisant deux hologrammes de microlentilles [31].

- Réseau de Damman : un réseau de phase binaire optimisé pour créer, par transformée de Fourier, une répartition uniforme de lumière dans certain ordre et minimiser la lumière ou même l'annuler dans d'autres [32]. Des réseaux de 200×200 points ont été réalisés [33]. L'efficacité de diffraction de tels réseaux bidimensionnels est d'environ 40 % [34].

- Contraste de phase : Il s'agit d'inverser le contraste d'un masque formé de régions opaques distribués sur une matrice en introduisant un déphaseur de π à l'ordre 0 dans le plan de Fourier [35]. Un taux de compression de la lumière entre 4 et 9 est réalisable avec un pas supérieur à 10 μm .

- Effet Talbot par objet de phase binaire : L'éclairage d'un réseau carré de phase binaire bidimensionnel forme, par effet Talbot, un réseau de carré lumineux dont le côté est la moitié de la période [36].

IV-2-4 Présentation de notre procédé

Nous avons étudié un procédé holographique pour réaliser un tel illuminateur de tableaux basé sur l'effet Talbot. Nous présentons ce procédé dans un cadre plus général, puis nous introduisons l'intérêt de l'effet Talbot pour simplifier le procédé [37].

Cas général

L'exemple le plus simple et le plus général d'un illuminateur de tableaux est représenté par l'imagerie directe d'un masque (fig. IV-4). Le masque est formé de trous qui ont la même forme et la même distribution que les pixels à éclairer. Les pixels, dans ce cas, peuvent avoir une distribution quelconque et ne sont pas forcément disposés selon un réseau régulier. L'avantage de cette méthode, comparativement à l'emploi de l'éclairage direct, est de ne pas éclairer la région entre les pixels. Au reste, le bilan énergétique reste le même, la lumière perdue est absorbée dans la partie opaque du masque.

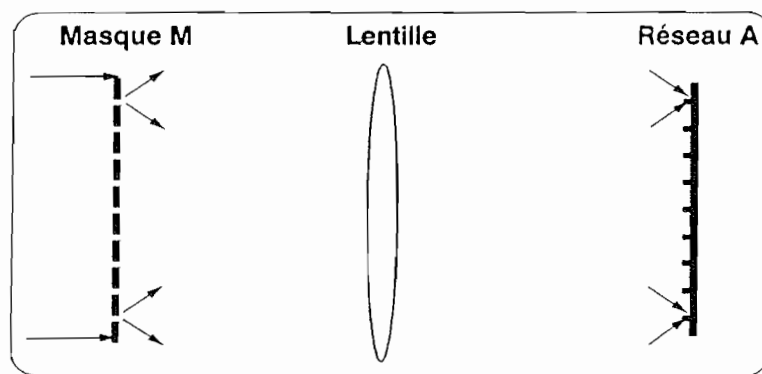


Fig. IV-4: Illuminateur par imagerie.

Notre procédé consiste à enregistrer un hologramme de l'onde diffractée par le masque dans un certain plan entre la lentille et le plan de l'image. Cet enregistrement se fait avec une onde référence plane ou sphérique. A la restitution, nous utilisons la même onde de référence, l'hologramme diffracte une onde similaire à celle résultant du masque et de la lentille, par conséquent cette onde reconstituera l'image du masque. Dans ce cas l'hologramme est directement utilisable comme illuminateur de tableaux sans lentille.

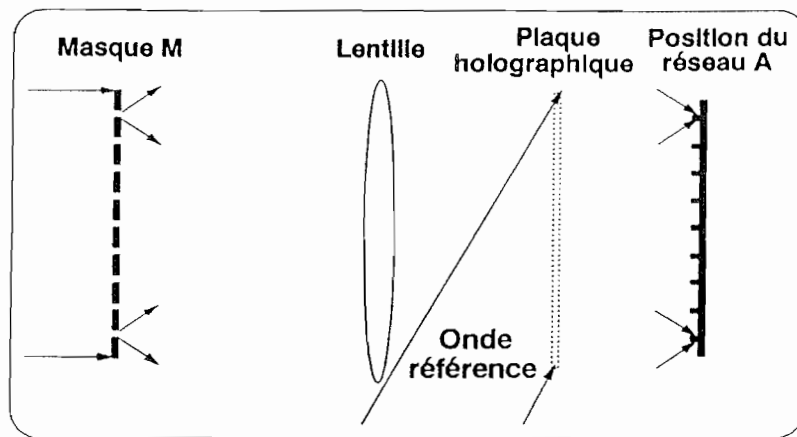


Fig. IV-5: *Illuminateur holographique.*

L'ordre "0" non diffracté par l'hologramme pourrait éclairer la matrice de traitement en dehors des pixels. Pour éviter cet éclairage, l'efficacité de diffraction de l'hologramme doit être suffisamment grande pour que cet ordre soit négligeable. Si ce n'est pas le cas, la géométrie du montage doit être conçue pour que cet ordre se trouve situé en dehors de la matrice de traitement.

Cas d'un réseau périodique

Dans la plupart des cas, les pixels sont distribués sur un réseau régulier carré, rectangulaire ou hexagonal. Le procédé présenté précédemment reste toujours valable à part que nous n'aurons pas besoin de la lentille pour former l'image du masque. Cette image sera directement formée par l'effet d'auto-imagerie de Talbot [38]. Cet effet exige que le masque contienne un grand nombre de périodes.

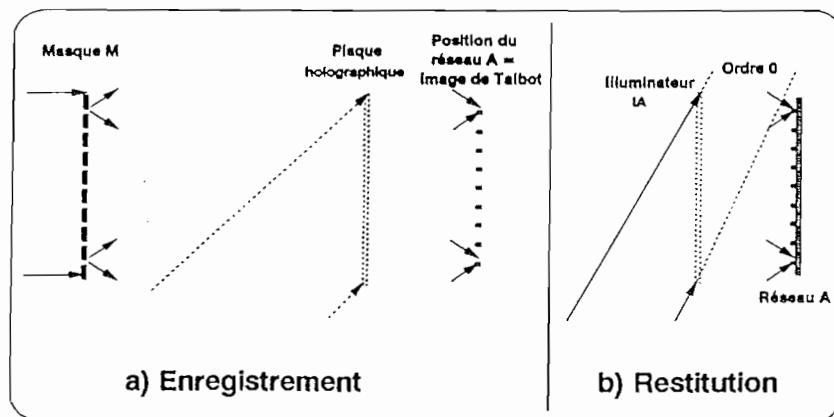


Fig IV-6 : *Hologramme illuminateur par effet Talbot.*

IV-3 EFFET D'AUTO-IMAGERIE DE TALBOT

IV-3-1 Définition de l'effet Talbot

Considérons un objet périodique infini éclairé par une onde plane monochromatique cohérente. Plaçons un écran parallèle à l'objet à une certaine distance z de cet objet. Il est clair que la figure de diffraction sur cet écran sera constituée d'un réseau périodique de franges d'interférences dont la fréquence principale est égale à celle de l'objet.

Si cet objet est un réseau unidimensionnel de période p , la figure d'interférence est une fonction périodique en distance de propagation z , de période :

$$D = 2p^2/\lambda$$

où λ est la longueur d'onde d'éclairage.

Par conséquent, l'image du réseau se reproduit dans des plans séparés par la distance D appelée distance de Talbot en référence à H.Fox Talbot qui a décrit au dix-neuvième siècle ce phénomène d'auto-imagerie des réseaux [38].

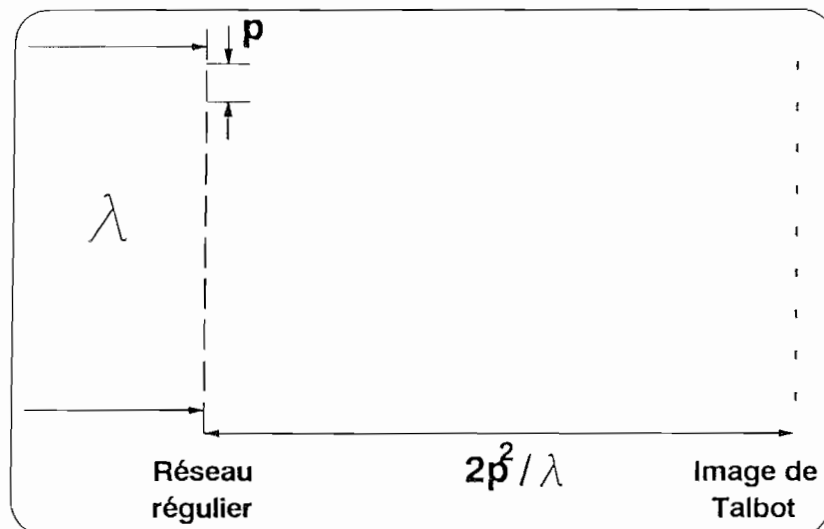


Fig. IV-7 : Effet Talbot.

L'effet Talbot reste valable pour des réseaux bidimensionnels à géométries spécifiques. Nous présentons, par la suite, une étude mathématique pour les réseaux unidimensionnels, bidimensionnels rectangulaires et hexagonaux.

IV-3-2 Analyse mathématique

Nous commençons cette analyse par un rappel mathématique sur la relation entre les fronts d'onde et leurs transformées de Fourier, et la simplification que porte cette transformée pour la propagation des ondes. Après ce rappel, nous analysons l'effet Talbot pour différents réseaux.

a) Rappel

Considérons un front d'onde dont l'amplitude complexe à la distance $z=0$ est notée $U(x,y,0)$. Sa transformée de Fourier est définie par :

$$\tilde{U}(\mu, \nu, 0) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} U(x, y, 0) \cdot \exp(-2i\pi \cdot \vec{\omega} \cdot \vec{r}) \cdot dx \cdot dy$$

tels que :

$$\vec{r} = \begin{cases} x \\ y \end{cases} \quad \text{coordonnées dans l'espace réel,}$$

$$\vec{\omega} = \begin{cases} \mu \\ \nu \end{cases} \quad \text{coordonnées dans l'espace de Fourier.}$$

Cette transformée de Fourier n'est en fait que la décomposition de l'onde en ondes planes. La valeur de la transformée de Fourier en un point (μ_0, ν_0) n'est nulle que si cette fréquence fait partie de la fonction $U(x,y,0)$, c'est à dire que un certain terme :

$$A \cdot \exp(2i\pi \cdot (\mu_0 x + \nu_0 y))$$

est une partie incluse dans $U(x,y,0)$. A ce moment la valeur de la transformée de Fourier en ce point est égale à :

$$\tilde{U}(\mu_0, \nu_0) = A$$

L'expression précédente est celle d'une onde plane se propageant selon la direction donnée par ses cosinus directeurs sur les axes x , y et z :

$$\vec{\Omega}_0 = \begin{cases} \alpha_0 = \lambda \cdot \mu_0 \\ \beta_0 = \lambda \cdot \nu_0 \\ \gamma_0 \end{cases}$$

le troisième cosinus directeur, γ_0 , se déduit de la relation :

$$\alpha_0^2 + \beta_0^2 + \gamma_0^2 = 1$$

Nous en déduisons que chaque point dans le plan de Fourier (μ, ν) représente une onde plane composant de l'onde $U(x, y, 0)$ et se propageant dans la direction $\vec{\Omega}$ dont les deux premiers composants sont donnés par le vecteur :

$$\vec{\Omega}_{xy} = \lambda \cdot \vec{\omega} = \begin{cases} \alpha = \lambda \mu \\ \beta = \lambda \nu \end{cases} \quad ;$$

et le troisième composant γ est donné par la relation précédente.

La densité est donnée par la valeur de la transformée de Fourier en ce point $U(\mu, \nu)$.

Pour la propagation de l'onde, nous partons de la relation de Helmholtz :

$$\Delta U + k^2 U = 0$$

En prenant la transformée de Fourier en x et y de cette relation nous trouvons :

$$\frac{\partial^2 \tilde{U}}{\partial^2 z}(\mu, \nu, z) + \gamma^2 U(\mu, \nu, z) = 0$$

avec
$$\gamma^2 = k^2 - (2\pi)^2 \cdot (\mu^2 + \nu^2) .$$

La solution de cette équation différentielle est donnée par :

$$\tilde{U}(\mu, \nu, z) = A^+(\mu, \nu) \cdot \exp(i\gamma z) + A^-(\mu, \nu) \cdot \exp(-i\gamma z)$$

et comme la propagation se fait dans le sens positif de z nous choisissons seulement le terme en A^+ .

La transformée de Fourier de l'onde diffractée dans le plan z s'écrit :

$$\tilde{U}(\mu, \nu, z) = \tilde{U}(\mu, \nu, 0) \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z \cdot \sqrt{1 - \lambda^2 \cdot (\mu^2 + \nu^2)}\right)$$

et l'onde diffractée dans le plan z est donnée par :

$$U(x, y, z) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \tilde{U}(\mu, \nu, 0) \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z \cdot \sqrt{1 - \lambda^2 \cdot (\mu^2 + \nu^2)}\right) \cdot \exp(2i\pi(\mu x + \nu y)) \cdot d\mu \cdot d\nu$$

Dans le cas de l'approximation de Fresnel nous développons la racine au premier ordre :

$$U(x, y, z) = \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \tilde{U}(\mu, \nu, 0) \cdot \exp(-i\pi \cdot \lambda \cdot (\mu^2 + \nu^2) \cdot z) \cdot \exp(2i\pi(\mu x + \nu y)) \cdot d\mu \cdot d\nu$$

$$U(x, y, z) = \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \cdot U(x, y, 0) * \text{TF}^{-1}\left(\exp(-i\pi \cdot \lambda \cdot (\mu^2 + \nu^2) \cdot z)\right)$$

et nous retrouvons l'intégrale de Fresnel :

$$U(x, y, z) = U(x, y, 0) *_{x,y} \psi(x, y, z)$$

avec :

$$\psi(x, y, z) = -\frac{i}{\lambda z} \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \cdot \exp\left(i\pi \frac{x^2 + y^2}{\lambda \cdot z}\right)$$

Dans le cas où $U(\mu, \nu, 0)$ est constitué des points discrets, la dernière relation s'écrit sous forme d'une série de Fourier.

$$U(x, y, z) = \sum_{ij} \tilde{U}(\mu_i, \nu_j, 0) \cdot \exp\left(\frac{2i\pi}{\lambda} \sqrt{1 - \lambda^2 (\mu_i^2 + \nu_j^2)} \cdot z\right) \cdot \exp(2i\pi(\mu_i x + \nu_j y))$$

ou dans l'approximation de Fresnel :

$$U(x, y, z) = \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{ij} \tilde{U}(\mu_i, \nu_j, 0) \cdot \exp(-i\pi \cdot \lambda \cdot (\mu_i^2 + \nu_j^2) \cdot z) \cdot \exp(2i\pi(\mu_i x + \nu_j y))$$

b) Réseau unidimensionnel infini

Supposons un réseau unidimensionnel infini de période p . Comme l'objet est infini, nous pouvons développer sa transmission en série de Fourier [39] :

$$T(x) = \sum_{j=-\infty}^{+\infty} t_j \cdot \exp\left(2i\pi \cdot j \cdot \frac{x}{p}\right)$$

t_j est le coefficient de Fourier défini par :

$$t_j = \frac{1}{p} \cdot \int_{-p/2}^{p/2} T(x) \cdot \exp\left(-2i\pi \cdot j \cdot \frac{x}{p}\right) \cdot dx$$

En éclairant cet objet par une onde plane définie par :

$$U(x, z=0) = A$$

L'onde transmise après l'objet s'écrit :

$$U(x, 0) = A \sum_{j=-\infty}^{+\infty} t_j \cdot \exp\left(2i\pi \cdot j \cdot \frac{x}{p}\right)$$

Dans le plan de Fourier cette onde se présente sous forme de pics de Dirac aux points situés

aux $\mu_j = j/p$ et de coefficient égal à $U\left(\frac{j}{p}, 0\right) = A \cdot t_j$.

D'après l'étude présentée dans le rappel, l'amplitude de l'onde dans le plan z s'écrit sous la forme :

$$U(x, z) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{j=-\infty}^{+\infty} t_j \cdot \exp\left(2i\pi \cdot x \cdot \frac{j}{p}\right) \cdot \exp\left(-2i\pi \cdot \frac{z}{z_T} \cdot j^2\right)$$

avec $z_T = 2 \cdot p^2 / \lambda$.

Pour les plans $z = k \cdot z_T$, k étant un entier, l'expression précédente s'écrit :

$$U(x, k z_T) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{j=-\infty}^{+\infty} t_j \cdot \exp\left(2i\pi \cdot x \cdot \frac{j}{p}\right)$$

qui est à un déphasage près l'objet initial. Par conséquent, l'image de l'objet se reproduit à des distances périodiques de

$$z_T = 2.p^2/\lambda$$

qui est la distance de Talbot.

Pour les plan $z = (2k+1).z_T/2$, k étant un entier, l'expression de l'onde diffractée devient :

$$U(x, z) = A. \exp\left(\frac{2i\pi}{\lambda} . z\right) \sum_{j=-\infty}^{+\infty} t_j . \exp\left(2i\pi . x . \frac{j}{p}\right) . \exp\left(-2i\pi . \frac{2k+1}{2} . j^2\right)$$

qui peut encore se mettre sous la forme :

$$U(x, z) = A. \exp\left(\frac{2i\pi}{\lambda} . z\right) \sum_{j=-\infty}^{+\infty} t_j . \exp\left(2i\pi . \left(x - \frac{p}{2}\right) . \frac{j}{p}\right) . \exp(-i\pi . (j^2 - j))$$

et comme $j^2 - j$ est toujours un entier pair, le terme $\exp(-i\pi . (j^2 - j))$ est égal à l'unité et l'expression de l'onde devient :

$$U(x, y) = A. \exp\left(\frac{2i\pi}{\lambda} . z\right) \sum_{j=-\infty}^{+\infty} t_j . \exp\left(2i\pi . \left(x - \frac{p}{2}\right) . \frac{j}{p}\right)$$

qui est exactement l'expression de l'objet de Talbot mais décalé d'une demi-période selon l'axe x .

Nous tirons de cette étude que l'image d'un réseau unidimensionnel éclairé par une onde plane se reproduit dans des plans séparés par la distance de Talbot ($z_T = 2.p^2/\lambda$) et aux demi-distances de Talbot avec un décalage d'une demi-période selon l'axe perpendiculaire au réseau.

c) Réseau bidimensionnel rectangulaire infini

Nous pouvons considérer le réseau rectangulaire comme le produit de deux réseaux unidimensionnels disposés selon les deux axes x et y avec deux périodes différentes p_x et p_y .

Comme le réseau résultant est à variables séparables x et y , le développement en série de Fourier de la transmittance d'un tel réseau s'écrit sous la forme :

$$T(x, y) = \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot x \cdot \frac{i}{p_x}\right) \cdot \exp\left(2i\pi \cdot y \cdot \frac{j}{p_y}\right)$$

En éclairant ce réseau par une onde plane, l'onde transmise par le réseau s'écrit sous la forme :

$$U(x, y, z = 0) = A \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot x \cdot \frac{i}{p_x}\right) \cdot \exp\left(2i\pi \cdot y \cdot \frac{j}{p_y}\right)$$

La transformée de Fourier de cette onde est constituée de points discrets :

$$\vec{\omega}_{ij} = \begin{cases} \mu_i = \frac{i}{p_x} \\ \nu_j = \frac{j}{p_y} \end{cases}$$

et les valeurs de cette transformée en ces points sont :

$$U\left(\frac{i}{p_x}, \frac{j}{p_y}\right) = t\left(\frac{i}{p_x}, \frac{j}{p_y}\right) = t_{ij}$$

Et l'onde diffractée dans le plan z aura une amplitude donnée par :

$$U(x, y, z) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \cdot \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot x \cdot \frac{i}{p_x}\right) \cdot \exp\left(2i\pi \cdot y \cdot \frac{j}{p_y}\right) \cdot \exp\left(-i\pi \cdot z \cdot \lambda \cdot \left(\left(\frac{i}{p_x}\right)^2 + \left(\frac{j}{p_y}\right)^2\right)\right)$$

Pour qu'il y ait un plan de reconstruction de l'image du réseau, il faut que l'exponentielle en i^2 et j^2 soit égale à l'unité quelque soit la valeur de i et de j . Par conséquent, il faut que les facteurs de i^2 et j^2 soient des entiers, ceci n'est possible que si $(p_x/p_y)^2$ est un entier rationnel.

$$\text{Soit } \begin{pmatrix} p_x \\ p_y \end{pmatrix} = \frac{n}{m}$$

tels que n et m sont des entiers premiers entre eux, à ce moment l'amplitude de l'onde au plan z s'écrit :

$$U(x, y, z) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot x \cdot \frac{i}{p_x}\right) \cdot \exp\left(2i\pi \cdot y \cdot \frac{j}{p_y}\right) \cdot \exp\left(-2i\pi \cdot z \cdot \frac{\lambda}{2np_y^2} \cdot (mi^2 + nj^2)\right)$$

De même nous pouvons écrire cette formule sous la forme :

$$U(x, y, z) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot x \cdot \frac{i}{p_x}\right) \cdot \exp\left(2i\pi \cdot y \cdot \frac{j}{p_y}\right) \cdot \exp\left(-2i\pi \cdot z \cdot \frac{\lambda}{2mp_x^2} \cdot (mi^2 + nj^2)\right)$$

D'après les formules précédentes nous déduisons l'existence de deux systèmes de plans de Talbot séparés respectivement par les distances:

$$z_T^x = 2 \cdot \frac{p_x^2}{\lambda}$$

et

$$z_T^y = 2 \cdot \frac{p_y^2}{\lambda}$$

La distance de Talbot effective est : $z_T = mz_T^x = nz_T^y$

cas particulier : réseau carré

Dans ce cas nous avons :

$$p_x = p_y = p$$

et

$$m = n = 1$$

et l'amplitude de l'onde au plan z devient :

$$U(x, y, z) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot x \cdot \frac{i}{p_x}\right) \cdot \exp\left(2i\pi \cdot y \cdot \frac{j}{p_y}\right) \cdot \exp\left(-2i\pi \cdot z \cdot \frac{\lambda}{2p^2} \cdot (i^2 + j^2)\right)$$

et la distance de Talbot sera :

$$z_T = 2p^2/\lambda$$

A la demi distance de Talbot $z = z_T/2$, nous pouvons écrire l'expression précédente sous la forme :

$$U(x, y, z) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot \left(x - \frac{p}{2}\right)\right) \cdot \exp\left(2i\pi \cdot \left(y - \frac{p}{2}\right)\right) \cdot \exp\left(-i\pi \cdot (i^2 - i + j^2 - j)\right)$$

et comme $i^2 - i$ et $j^2 - j$ sont toujours pairs, nous retrouvons le développement en série de Fourier du réseau décalé d'une demi-période dans les deux axes x et y .

d) Réseau bidimensionnel hexagonal infini

Nous pouvons décrire le réseau hexagonal par une somme de pics de Dirac sous la forme :

$$\sum_{ij} \delta\left\{x - p \cdot \left[i - \frac{1}{4} \cdot (1 - (-1)^j)\right]\right\} \cdot \delta\left\{y - p \cdot j \cdot \frac{\sqrt{3}}{2}\right\}$$

Nous considérons un objet infini constitué de trous transparents distribués sur un réseau hexagonal. Le développement en série de Fourier d'un tel objet s'écrit :

$$T(x, y) = \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot \frac{x}{p} \cdot \left(i - \frac{1}{4} \cdot (1 - (-1)^j)\right)\right) \cdot \exp\left(2i\pi \cdot \frac{y}{p} \cdot \left(j \cdot \frac{\sqrt{3}}{2}\right)\right)$$

En éclairant cet objet par une onde plane normale à l'objet, et en propageant cette onde jusqu'au plan z , nous obtenons l'amplitude suivante de l'onde :

$$T(x, y) = A \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot \frac{x}{p} \cdot \left(i - \frac{1}{4} \cdot (1 - (-1)^j)\right)\right) \cdot \exp\left(2i\pi \cdot \frac{y}{p} \cdot \left(j \cdot \frac{\sqrt{3}}{2}\right)\right) \\ \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z \cdot \sqrt{1 - \left(i - \frac{1}{4} \cdot (1 - (-1)^j)\right)^2 - \left(j \cdot \frac{\sqrt{3}}{2}\right)^2}\right)$$

que nous pouvons écrire en tenant-compte de l'approximation de Fresnel :

$$T(x, y) = A \cdot \exp\left(\frac{2i\pi}{\lambda} \cdot z\right) \sum_{ij} t_{ij} \cdot \exp\left(2i\pi \cdot \frac{x}{p} \cdot \left(i - \frac{1}{4} \cdot (1 - (-1)^j)\right)\right) \cdot \exp\left(2i\pi \cdot \frac{y}{p} \cdot \left(j \cdot \frac{\sqrt{3}}{2}\right)\right) \\ \cdot \exp\left(-2i\pi \cdot \frac{\lambda \cdot z}{2p^2} \cdot \left(\left(i - \frac{1}{4} \cdot (1 - (-1)^j)\right)^2 - \frac{3}{4} \cdot j^2\right)\right)$$

Nous désignons par : $K = \left(i - \frac{1}{4} \cdot (1 - (-1)^j)\right)^2 - \frac{3}{4} \cdot j^2$

Si j est pair : $j = 2l$ et $K = i^2 + 3 \cdot l^2$ est un entier.

Si j est impair : $j = 2l + 1$ et $K = i^2 - i + 3l^2 + 3l + 1$ est un entier.

Donc nous remarquons que le facteur K est toujours un entier, et l'image du réseau se reproduit à des distances égales à :

$$z_T = 2 \cdot p^2 / \lambda$$

qui est la distance de Talbot.

IV-3-3 Influence des dimensions finies de l'objet

La reconstruction de l'image du réseau par effet Talbot n'est parfaite, en tout rigueur, que lorsque le réseau est infini. En pratique, le réseau a toujours des dimensions finies, ce qui introduit forcément des déformations à l'image reconstruite. Une analyse de ces déformations permet d'estimer les dimensions du réseau nécessaire pour reconstruire correctement une partie de l'image.

Le fait d'utiliser un masque de dimensions finies revient à multiplier la transmission du réseau infini par une fonction de pupille $s(x, y)$.

$$t(x, y) = T(x, y) \cdot s(x, y)$$

La transformée de Fourier de cette transmittance est la convolution de la transmittance du réseau infini avec la transformée de Fourier de la fonction de la pupille :

$$\tilde{t}(\mu, \nu) = \tilde{T}(\mu, \nu) * \tilde{s}(\mu, \nu)$$

Dans ce cas l'amplitude de l'onde diffractée dans le plan à la distance z de l'objet est donnée par :

$$U(x, y, z) = \iint \left\{ \iint \tilde{T}(\mu', \nu') \cdot \tilde{s}(\mu' - \mu, \nu' - \nu) \cdot d\mu' \cdot d\nu' \right\} \\ \cdot \exp(-i\pi\lambda(\mu^2 + \nu^2) \cdot z) \cdot \exp(2i\pi \cdot (x\mu + y\nu)) \cdot d\mu \cdot d\nu$$

Compte-tenu que la transformée de Fourier d'un réseau infini est constituée de somme de pics de Dirac, nous pouvons écrire cette transformée dans le cas d'un réseau carré de période p :

$$\tilde{T}(\mu, \nu) = \sum_{ij} t_{ij} \cdot \delta\left(\mu - \frac{i}{p}\right) \cdot \delta\left(\nu - \frac{j}{p}\right)$$

A ce moment l'amplitude de l'onde s'écrit comme suit :

$$U(x, y, z) = \iint \left\{ \sum_{ij} t_{ij} \cdot \tilde{s}\left(\frac{i}{p} - \mu, \frac{j}{p} - \nu\right) \right\} \cdot \exp(-i\pi\lambda \cdot (\mu^2 + \nu^2) \cdot z) \\ \cdot \exp(2i\pi \cdot (x\mu + y\nu)) \cdot d\mu \cdot d\nu$$

En permutant les signes sommes et intégrales et en faisant un changement de variable nous trouvons :

$$U(x, y, z) = \sum_{ij} t_{ij} \cdot S_{ij} \cdot \exp\left(\frac{2i\pi}{p} \cdot (ix + jy)\right) \cdot \exp\left(-\frac{2i\pi\lambda}{p^2} \cdot (i^2 + j^2) \cdot z\right)$$

tel que :

$$S_{ij}(x, y) = \iint \tilde{s}(a, b) \cdot \exp(-i\pi\lambda \cdot (a^2 + b^2) \cdot z) \\ \cdot \exp\left(-2i\pi \left(a \left(x - \frac{i\lambda}{p} z\right) + b \left(y - \frac{j\lambda}{p} z\right)\right)\right) \cdot da \cdot db$$

D'après cette analyse, l'image de reconstruction est affectée localement par les fonctions $S_{ij}(x,y)$.

Nous allons analyser cette déformation dans le cas particulier d'une pupille rectangulaire de longueur L_x et de largeur L_y .

La fonction caractéristique de cette pupille est donnée par :

$$s(x,y) = \prod \left(\frac{x}{L_x} \right) \cdot \prod \left(\frac{y}{L_y} \right)$$

et sa transformée de Fourier :

$$\tilde{s}(\mu, \nu) = L_x \cdot L_y \cdot \text{sinc}(\mu L_x) \cdot \text{sinc}(\nu L_y)$$

Les fonctions S_{ij} sont données par la relation :

$$S_{ij} = S_{i,L_x}(x,z) \cdot S_{j,L_y}(y,z)$$

telle que :

$$S_{i,L}(x,z) = \int \text{sinc}(\alpha) \cdot \exp\left(-i\pi\lambda \left(\frac{\alpha}{L}\right)^2 z\right) \cdot \exp\left(-2i\pi \cdot \frac{\alpha}{L} \left(x - \frac{\lambda_i}{p} z\right)\right) \cdot d\alpha$$

Bien que l'intégrale se fasse sur α jusqu'à l'infini, le sinus cardinal et les oscillations des exponentielles limitent vite les valeurs utiles pour l'intégrale autour de $\alpha = 0$. Par conséquent, le terme exponentiel en α^2 est proche de l'unité dans la condition :

$\lambda z/L^2 \ll 1$. Dans cette condition, nous trouverons :

$$S_{i,L}(x,z) \approx \prod \left(\frac{x - \frac{\lambda_i}{p} z}{L} \right)$$

qui est la même fonction de la pupille décalée selon l'axe x . Sachant que $\lambda_i/p = \sin \alpha_i$, tel que α_i est l'angle de propagation de l'ordre de diffraction i qui est une onde plane, nous

pouvons interpréter ce résultat par le fait que l'onde plane de chaque ordre de diffraction est limitée spatialement par la pupille le long de sa propagation.

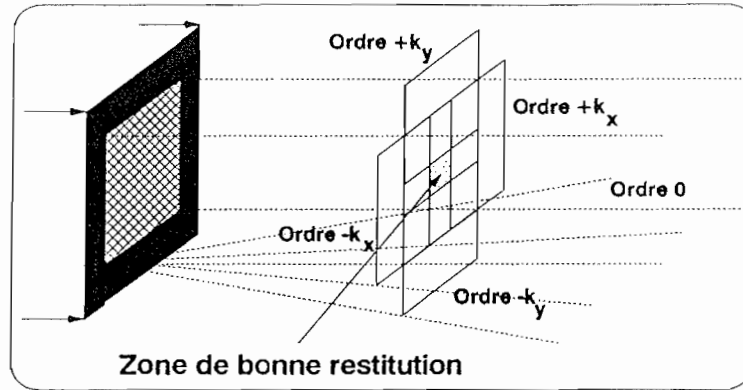


Fig. IV-8 : Effet de la pupille.

La reconstruction du motif principal du réseau est l'interférence de tous les ordres de diffraction. Mais un certain nombre de ces ordres est nécessaire à une reconstruction correcte, tandis que les autres ordres peuvent être négligés. Supposons qu'une reconstruction correcte de ce motif demande $2k+1$ ordres ($-k \leq i, j \leq +k$). Si nous voulons une reconstruction correcte de ce motif à l'abscisse x , il faut que cet endroit soit à l'intersection des deux pupilles suivant les ondes d'ordres $-k$ et $+k$, ou bien :

$$-\frac{L}{2} + \frac{\lambda kz}{p} < x - \frac{p}{2} \quad \text{et} \quad x + \frac{p}{2} < +\frac{L}{2} - \frac{\lambda kz}{p}$$

IV-3-4 Conditions d'une reconstruction correcte de l'image

Nous prenons l'exemple d'un réseau carré de trous de période p , chaque trou est un carré de côté a . Ce réseau est limité par une pupille carrée de largeur L . Dans ces conditions, les coefficients de Fourier sont donnés par la relation :

$$t_{ij} = \varepsilon^2 \text{sinc}(\varepsilon i) \cdot \text{sinc}(\varepsilon j)$$

$\varepsilon = a/p$ est le rapport cyclique du réseau.

Nous allons déterminer la largeur minimale de la pupille pour reconstruire correctement une matrice de $N * N$ points.

Détermination des ordres de diffraction nécessaires

Nous avons pris pour valeurs numériques :

$$p = 300 \quad \mu\text{m}$$

$$a = 50 \mu\text{m}$$

Les simulations montrent qu'il faut 57 ($k=28$) ordres de diffraction pour une reconstruction correcte du trou carré :

$$-28 \leq i, j \leq +28$$

Détermination de la distance de Talbot

En éclairant ce réseau avec une onde plane de $\lambda = 488 \text{ nm}$, nous trouvons que la demi-distance de Talbot :

$$p^2/\lambda = 18.443 \text{ cm}$$

Influence de la pupille

Pour pouvoir approcher la formule (1) du paragraphe précédent par la fonction de la pupille, il faut que le terme exponentiel soit proche de l'unité pour les valeurs de α où le $\text{sinc}(\alpha)$ est important. Nous pouvons considérer que cette approximation est valable quand le premier zéro de la partie réelle de l'exponentielle est situé très loin par rapport au premier zéro de $\text{sinc}(\alpha)$. Ceci conduit à la condition suivante :

$$2\lambda z/p^2 \ll \pi^2$$

Pour les valeurs citées précédemment de λ et z , la largeur de la pupille doit vérifier :

$$L \gg 1.35 \cdot 10^{-2} \text{ cm}$$

Donc nous pouvons considérer que pour L de l'ordre de cm, cette approximation est valable.

Pour reconstruire correctement une matrice de $10 * 10$ points, il faut que les conditions citées précédemment soient vérifiées, ce qui revient à satisfaire la condition :

$$L > 2 \cdot \left(5p + \frac{\lambda k z}{p} \right)$$

ce qui donne une largeur minimale de la pupille de : $L = 9 \text{ mm}$ correspondant à 30 périodes. Donc le masque à utiliser doit être une matrice de $30 * 30$ points minimum.

IV-4 TECHNIQUES HOLOGRAPHIQUES

Le choix du type de l'hologramme est déterminé à partir de son efficacité de diffraction, ce qui précise son rendement énergétique lors de son utilisation. Nous avons vu au paragraphe IV-2-2-a que l'efficacité de diffraction est le rapport entre le flux lumineux diffracté par l'hologramme et formant d'image désirée et le flux lumineux total éclairant l'hologramme lors de la restitution. Toute lumière diffractée dans d'autres ordres est perdue.

L'importance de l'efficacité de diffraction est évidente pour l'énergie nécessaire pour éclairer les modulateurs optiques ou opto-électroniques. Ces modulateurs nécessitent un minimum d'énergie lumineuse pour le basculement en un temps déterminé. Alors l'efficacité de diffraction de l'hologramme est un facteur essentiel pour le choix du laser illuminateur et pour sa puissance.

Nous passons en revue les différents types d'hologrammes avec leurs efficacités de diffraction théoriques. Ensuite nous analysons les propriétés des hologrammes de phase épais.

IV-4-1 Présentation générale

La classification des hologrammes se fait selon plusieurs critères :

-Méthode de fabrication :

- * hologrammes naturels enregistrés par interférences d'ondes objets et références.
- * hologrammes synthétiques où le motif diffractant est calculé par ordinateur.

-Type de fonctionnement :

- * en réflexion.
- * en transmission.

-Type de modulation de l'onde diffractée :

- * modulation d'amplitude.
- * modulation de phase.

-Type d'enregistrement sur la plaque :

- * hologrammes minces où la modulation est la même dans l'épaisseur.
- * hologrammes épais où le système de franges varie en fonction de la profondeur.

Nous nous intéressons ici aux hologrammes naturels : hologrammes enregistrés par l'interférence de deux ondes, différenciés des hologrammes synthétiques calculés par l'ordinateur.

Les hologrammes d'amplitude sont constitués de franges sombres et transparentes, enregistrées en général sur des plaques argentées comme dans la technique photographiques. L'efficacité de ces hologrammes est très faible compte-tenu de l'absorption dans les parties sombres de l'hologramme.

Pour les hologrammes de phase minces, les franges sont enregistrées sous forme de modulation de l'indice de refraction ou modulation d'épaisseur sur une plaque transparente. L'onde diffractée est modulée en phase et non en amplitude. L'efficacité de diffraction peut atteindre 100% pour certain réseau comme les réseaux dits blasés dont le contrôle de l'exposition des franges à profil dissymétrique nécessaire pour "blaser" un réseau est particulièrement critique. Mais les hologrammes de phase mince obtenus par blanchiment des hologrammes d'amplitude présente une efficacité plus grande que celle des hologrammes d'amplitude mais n'atteint pas 100% à cause des ondes non diffractées.

Pour les hologrammes de phase épais, la modulation de l'indice de refraction est faite dans l'épaisseur de la plaque ce qui fait que les ondes non diffractées sur la surface de la plaque sont diffractées dans l'épaisseur. Les franges d'égale modulation jouent le rôle d'un miroir, comme nous allons voir par la suite. Par conséquent, l'efficacité de diffraction peut atteindre 100%.

Nous présentons dans le tableau suivant les différents types d'hologrammes dont la modulation (de phase ou d'amplitude) est sinusoïdale obtenus par enregistrement holographique simple, avec leurs efficacités de diffraction théoriques [40]:

Type	Transmission mince		Transmission épais		Réflexion épais	
	Amplitude	Phase	Amplitude	Phase	Amplitude	Phase
Efficacité	6.25 %	33.9 %	3.7 %	100 %	7.20 %	100 %

D'après le tableau précédent, il est évident que les hologrammes de phase épais sont les plus intéressants.

IV-4-2 Hologramme de phase épais

Nous traitons dans ce paragraphe les caractéristiques d'un hologramme de phase épais constitué de l'interférence de deux ondes planes.

A l'enregistrement de l'hologramme, l'interférence des deux ondes planes, d'une longueur d'onde λ_1 , introduit une modulation de l'indice de réfraction dans l'épaisseur du matériel.

Compte tenu de l'utilisation de deux ondes planes, cette modulation prend la forme des strates inclinées, selon les angles de ces deux ondes, avec un profil sinusoïdal dans la direction perpendiculaire à ces strates $K = 2\pi/\Lambda$. La modulation de l'indice de réfraction s'exprime sous la forme :

$$n = n_0 + n_1 \cos(K \cdot X)$$

où n_0 est l'indice de refraction moyen et n_1 l'amplitude de modulation spatiale de l'indice.

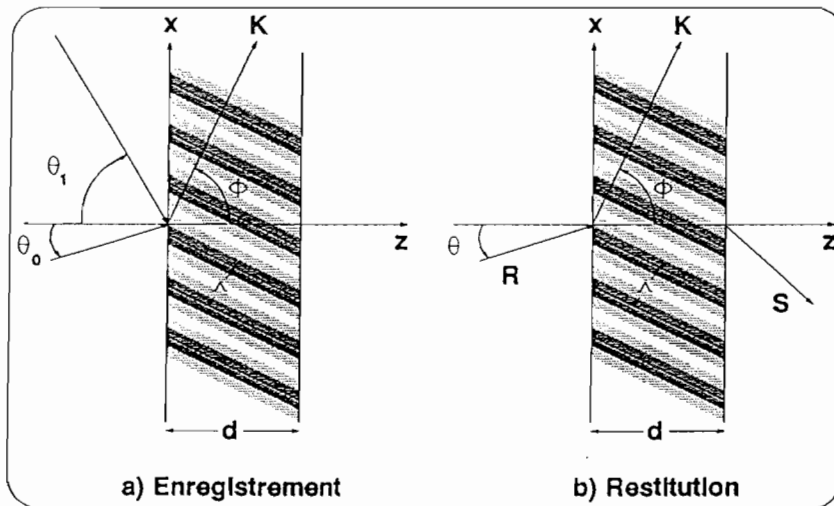


Fig. VI-9 : Hologramme de phase épais.

Le pas de ces strates est donnée par la relation :

$$\Lambda = \frac{\lambda_1}{2 \sin\left(\frac{\theta_0 - \theta_1}{2}\right)}$$

et leur inclinaison :

$$\sin \Phi = -\cos\left(\frac{\theta_0 + \theta_1}{2}\right)$$

En considérant que les strates sont suffisamment profondes dans l'épaisseur de l'hologramme, la diffraction d'une onde plane $\mathbf{R}$ sur cet hologramme est strictement conditionnée par l'effet Bragg. L'effet Bragg détermine la condition pour qu'une onde soit diffractée par le réseau selon sa longueur λ d'onde et son angle d'incidence θ (dans le milieu). Cette condition est donnée par la relation :

$$\cos(\Phi - \theta) = \frac{\lambda}{2\Lambda}$$

Toute onde violant cette condition ne sera pas diffractée par l'hologramme. Par contre, une onde s'approchant cette condition sera diffractée en une onde $\mathbf{S}$.

Kogelnik étudie ce phénomène de diffraction par la théorie d'onde couplée, et en déduit l'efficacité de diffraction η [41]. Cette efficacité est donnée dans le cas d'un hologramme de phase épais sans absorption et travaillant en transmission par la relation :

$$\eta = \sin^2 \left(v^2 + \xi^2 \right)^{1/2} / \left(1 + \xi^2 / v^2 \right)$$

où :

$$v = \pi n_1 d / \lambda (c_R c_S)^{1/2}$$

$$\xi = \Delta\theta . K d \sin(\Phi - \theta_0) / 2c_S = -\Delta\lambda . K^2 d / 8\pi n c_S$$

$$c_R = \cos \theta$$

$$c_S = \cos \theta - \frac{\Lambda}{\lambda} \cos \Phi$$

$\Delta\theta$: l'écart de l'inclinaison de l'onde par rapport à l'angle de Bragg.

$\Delta\lambda$: l'écart de la longueur d'onde de l'onde de Bragg.

Si l'onde répond parfaitement à la condition de Bragg $\xi = 0$, l'efficacité de diffraction devient :

$$\eta = \sin^2 v$$

et pour $v = \pi/2$, l'efficacité de diffraction est de 100%. Compte tenu que v est directement lié à l'amplitude de modulation de l'indice de refraction n_1 , le maximum de l'efficacité de diffraction est obtenu pour une valeur déterminée de cette amplitude. Cette amplitude de

modulation n_1 est déterminée en fonction de l'épaisseur du milieu, de la longueur d'onde et des angles d'inclinaison.

IV-5 CONDITIONS D'ENREGISTREMENT

IV-5-1 Gélatine dichromatée

La gélatine dichromatée est une des matières les plus utilisées pour enregistrer des hologrammes de phase épais.

Les caractéristiques essentielles de la gélatine dichromatée sont les suivantes [42] :

- Modulation d'indice** : de fortes modulations d'indices d'environ 0.08 ont été réalisées.
- Résolution** : environ 5000 traits/mm.
- Sensibilité** : de l'ultra-violet au bleu vert. Cette sensibilité dépend du sel de dichromate utilisé comme sensibilisant.
- Absorption Diffusion** : absorption extrêmement faible et excellent rapport signal sur bruit.
- Stabilité** : assez sensible à l'humidité qui peut causer une baisse importante de l'efficacité de diffraction.
- Variation d'épaisseur** : peut être contrôlée par la préparation, l'exposition ou le développement.
- Retraitement** : la modification de la modulation d'indice est possible par un traitement ultérieur.

Le mécanisme d'enregistrement est basé sur le changement des liaisons de molécules causé par les ions . Ces ions sont dus à la réduction des ions Cr^{6+} sous l'influence de l'éclairage lumineux.

La variation de l'indice de réfraction est directement liée à l'énergie lumineuse reçue. D'après les études faites sur les hologrammes simples d'ondes planes, le maximum d'efficacité est obtenu pour une fluance d'éclairement entre 60 et 100 mJ/cm^2 [43]. Le traitement de la gélatine bichromatée est donné dans l'annexe A.

IV-5-2 Recherche du plan de moindre dynamique

Comme nous l'avons vu, l'efficacité de diffraction d'un hologramme de phase épais est étroitement liée à la modulation de l'indice de réfraction. Pour avoir une efficacité de

diffraction proche de 100% il faut que la modulation de l'indice de réfraction soit uniforme le long de la plaque holographique. Comme la modulation de l'indice de réfraction dépend de l'intensité lumineuse des faisceaux objet et référence lors de l'enregistrement, les variations de l'intensité de l'onde enregistrée doivent être minimales sur la plaque holographique. Cette condition est facilement vérifiée pour les hologrammes simples d'ondes planes ou d'ondes sphériques issues d'un trou d'une petite ouverture. Par contre, l'onde diffractée par l'objet "Talbot" est loin d'être une onde plane ou sphérique, elle s'approche plutôt de la somme d'ondes sphériques convergeant sur les pixels de l'image de Talbot et interférant entre elles au niveau de la plaque holographique.

Par conséquent, il faut s'assurer que la figure d'interférence précédente est la plus proche possible d'une figure à intensité constante avec modulation de phase. L'écart résiduel de l'intensité constante explique que l'efficacité ne peut atteindre 100%.

Deux sortes de paramètres physiques sont susceptibles de rendre la figure d'interférence précédente la plus uniforme possible en intensité :

- la position de l'hologramme.
- la phase relative des différents pixels de l'objet de Talbot.

La deuxième solution, qui exige le contrôle d'une phase variable, est assez difficile à réaliser du point de vue pratique.

C'est pourquoi nous avons choisi de positionner l'hologramme à un endroit où la figure diffractée par l'objet présente la moindre variation de l'intensité, ou "plan de moindre dynamique". Pour déterminer ce plan, nous avons simulé la figure diffractée sur des plans à distances variables de l'objet. Cet objet est défini comme suit :

Un réseau carré de pixels de pas $p=300\ \mu\text{m}$ dans les deux directions. Chaque pixel est un carré de $a=50\ \mu\text{m}$ de côté. La longueur d'onde d'éclairage est de 488 nm (une longueur d'onde convenable pour la sensibilité de la gélatine dichromatée). La demi-distance de Talbot est $p^2/\lambda = 184\ \text{mm}$.

L'étude consiste à simuler numériquement la figure diffractée par un tel objet en fonction de la distance z de l'objet (Fig. IV-10). Cette simulation est faite à partir de la décomposition d'un objet unidimensionnel en série de Fourier, et la propagation de chaque ordre jusqu'au plan z , ensuite l'interférence entre ces différents ordres au plan z (formule du paragraphe IV-3-2-b). L'objet bidimensionnel, ainsi que la figure de diffraction, sont obtenus par la multiplication selon la deuxième dimension.

Cette étude cherche à minimiser la variance de l'intensité de l'onde diffractée par l'objet, ou minimiser le facteur :

$$\sigma^2 = \int_0^p (I - I_m)^2 dx$$

où I_m est la moyenne de l'intensité sur une période :

$$I_m = \frac{1}{p} \int_0^p I(x) dx$$

où $I(x)$ est l'intensité à l'abscisse x .

Le plan d'uniformité maximale en intensité présente le profil montré dans la figure IV-11 b. Il est situé à 31 mm de l'image de Talbot.

Nous avons présenté sur la même figure trois autres profils d'intensité de l'image diffractée en trois plans différents positionnés selon la figure IV-12.

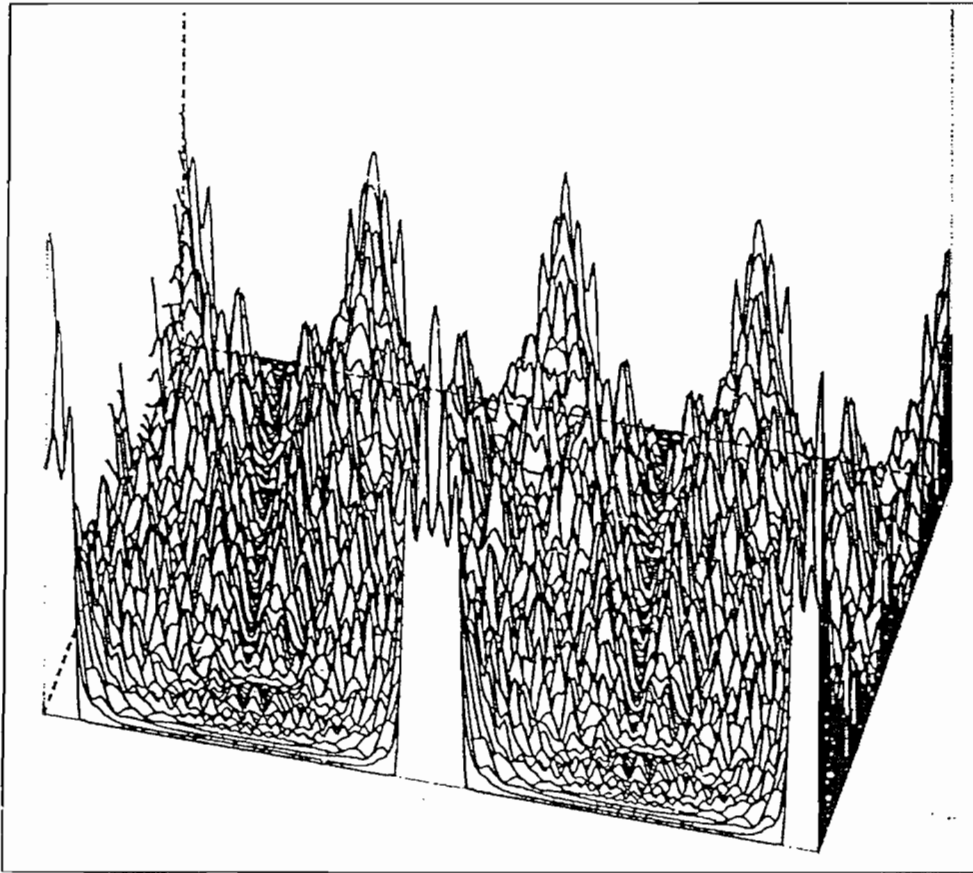
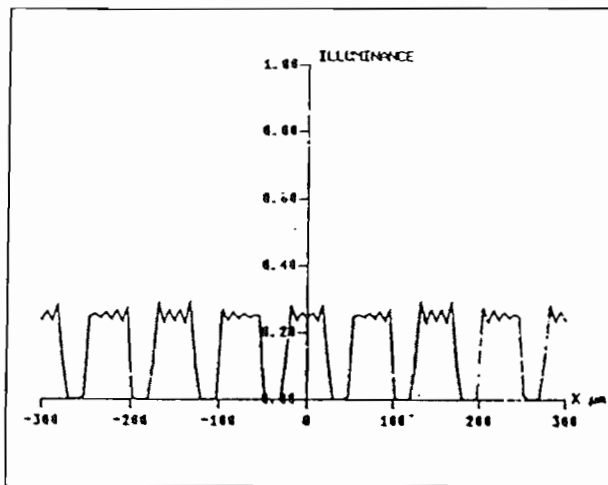
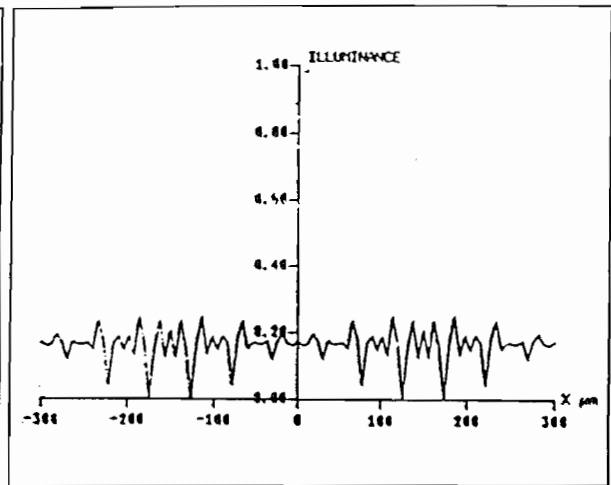


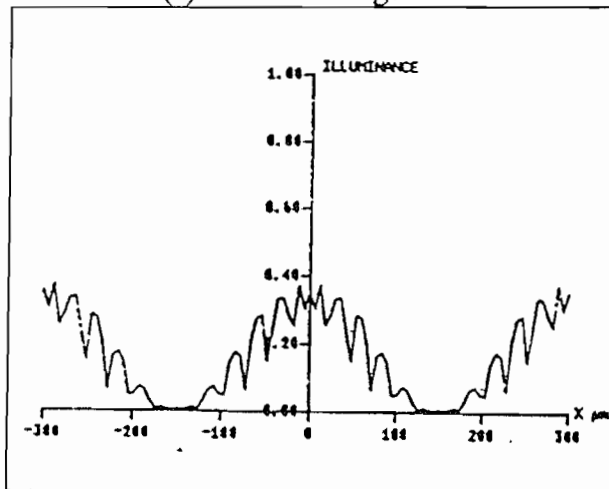
Fig IV-10 : Simulation de l'onde diffractée.



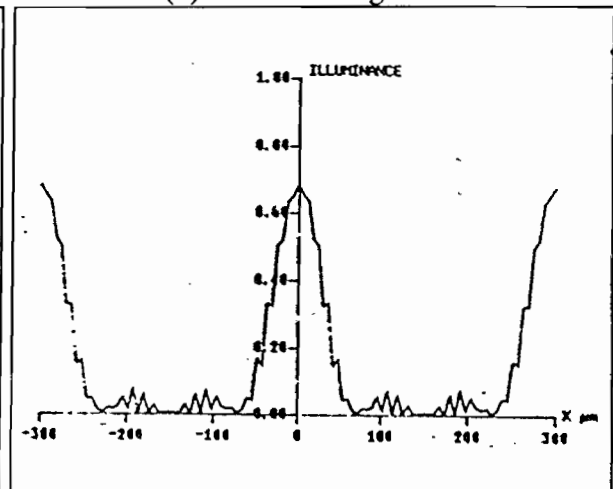
(a) Plan 1 de Fig. IV-12



(b) Plan 2 de Fig. IV-12



(c) Plan 3 de Fig. IV-12



(d) Plan 4 de Fig. IV-12

Fig. IV-11 : Profils d'intensité.

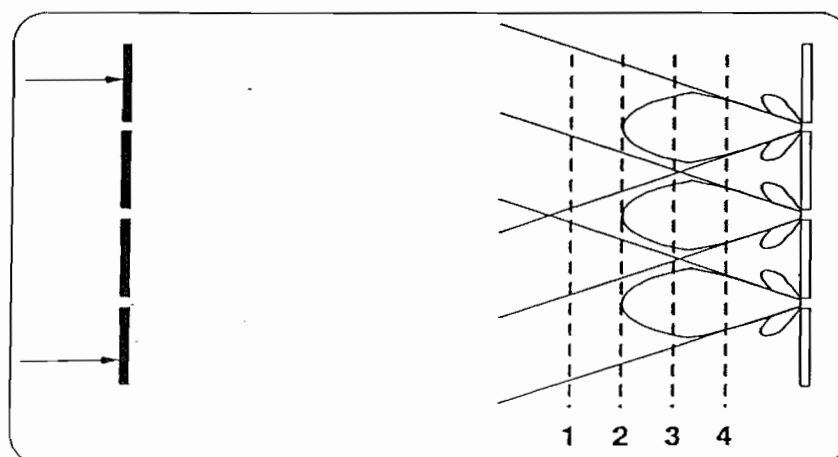


Fig. IV-12: Plans de diffraction.

Une explication physique simple de la position du plan de moindre dynamique en intensité peut être donnée en considérant des ondes sphériques convergentes :

La lumière diffractée par chaque trou isolé prend la forme des lobes dont la section suivant une direction est un "sinc". Au niveau de l'image de Talbot, nous retrouvons ces lobes de diffraction convergeant aux pixels de l'image. Les lobes principaux de diffraction se rejoignent au niveau du plan 3. Dans un plan plus proche de l'image (plan 4) les lobes ne se recouvrent pas et nous pouvons distinguer les différents lobes. Dans un plan éloigné de l'image (plan 1) l'interférence entre les lobes est très forte et nous pouvons distinguer les franges d'interférence. L'optimum est le plan 2 où les lobes principaux interfèrent légèrement mais le contraste des franges est assez faibles.

Du fait que dans le cas de notre illuminateur les lobes de diffraction interfèrent au niveau de l'hologramme, nous ne pouvons pas le considérer comme un hologramme d'ondelettes indépendantes comme c'est le cas dans les microlentilles holographiques.

IV-6 REALISATION DE L'HOLOGRAMME

IV-6-1 Objet de Talbot

Nous avons utilisé un masque carré constitué du réseau précédent contenant 90×90 trous. La réalisation de ce masque est faite par la photoréduction sur plaque photographique d'une image de ce masque calculée par ordinateur et sortie sur table traçante.

IV-6-2 Montage holographique

La figure IV-13 montre un schéma du montage holographique utilisé pour l'enregistrement des hologrammes. Le laser utilisé est un laser continu à Argon à la longueur d'onde 488 nm avec une puissance de 5 mW. L'angle entre les deux faisceaux est de 47° . Le masque est placé perpendiculairement à un faisceau. La plaque holographique est placée parallèlement au masque à une distance correspondant à celle du plan de moindre dynamique. L'ajustement de ce plan est fait par un viseur dont le plan frontal coïncide avec le côté gélatine de la plaque. La puissance moyenne, due aux deux faisceaux, reçue par la plaque holographique était d'environ $400 \mu\text{W}/\text{cm}^2$. L'enregistrement est fait pour plusieurs temps de pause, l'optimum de l'efficacité a été obtenu pour un temps de pause de 180 s, ce qui correspond à une fluance totale de $72 \text{ mJ}/\text{cm}^2$.

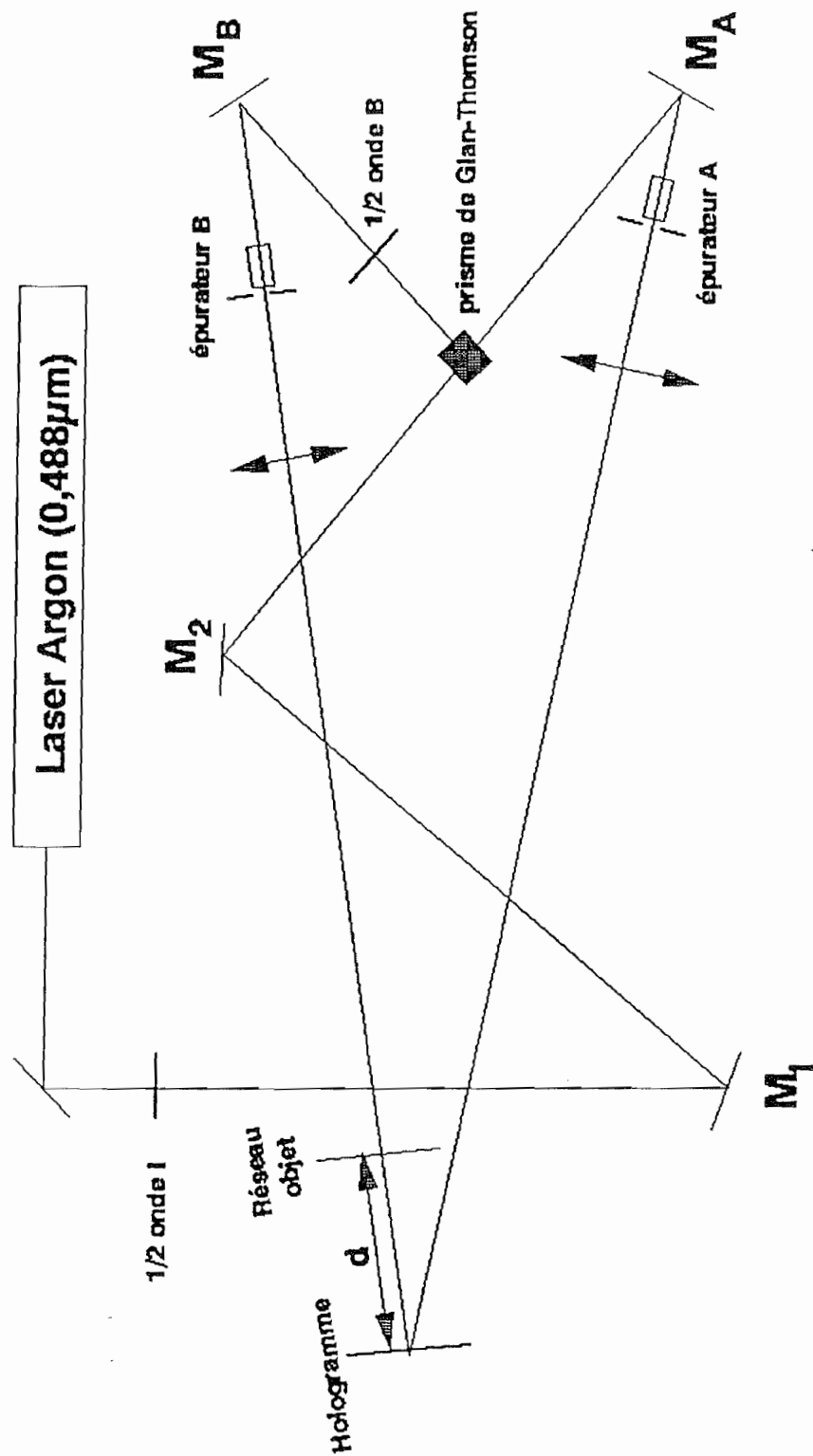


Fig. IV-13 : Montage holographique.

IV-6-3 Résultats expérimentaux

La figure IV-14-a montre le masque, la figure IV-14-b est son image de Talbot, et la figure IV-14-c l'image de Talbot restituée par l'hologramme. La figure IV-14-b correspond à seulement 3% de l'énergie d'éclairage tandis que la figure IV-14-b correspond à 70%.

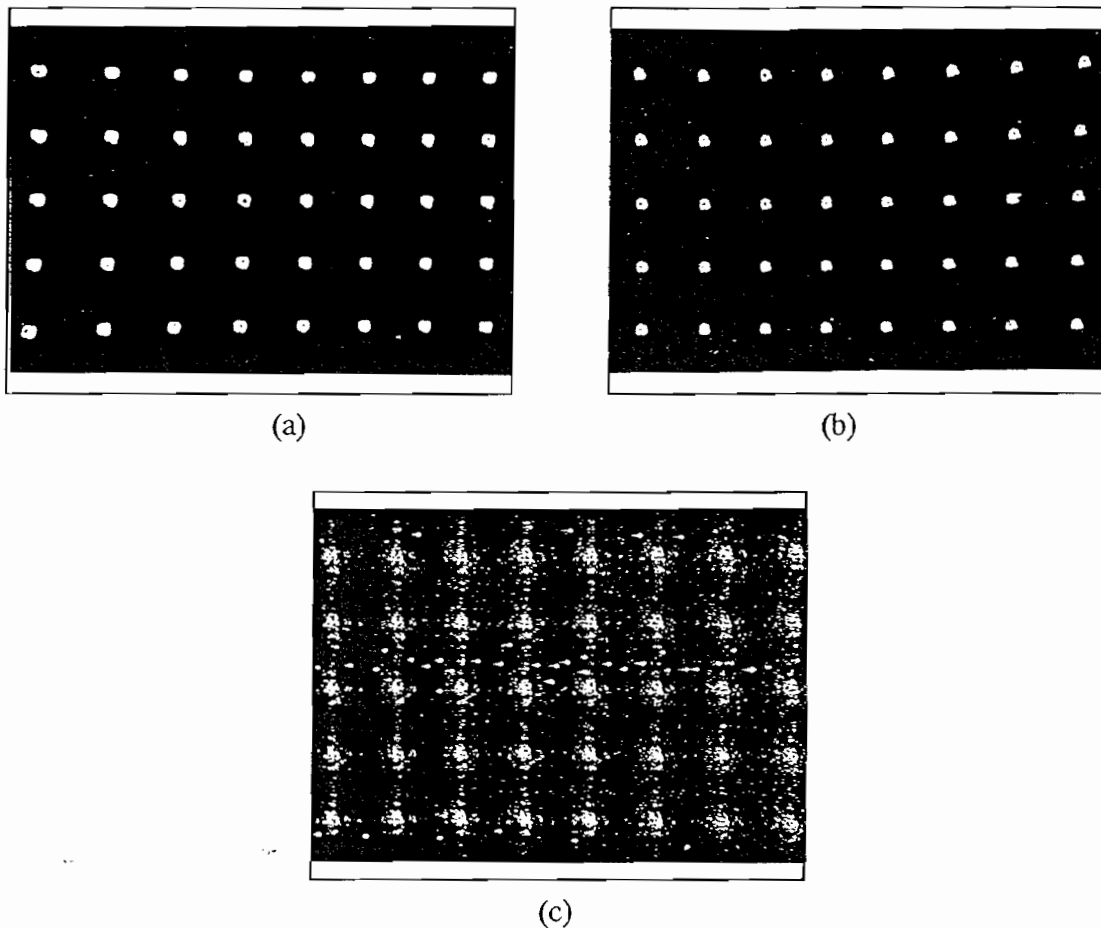


Fig IV-13 : Résultats expérimentaux.

Nous remarquons un bruit lumineux entre les pixels sur l'image de restitution. Plusieurs hypothèses sont possibles pour expliquer la présence de ce bruit :

- La diffusion de la lumière dans la gélatine : cette diffusion génère en général un fond lumineux sous forme de speckle (phénomène aléatoire), ce qui n'est pas tout à fait le cas de l'image restituée où le bruit a plutôt une forme régulière.

- L'intermodulation de l'onde objet : l'onde objet diffractée par l'objet de Talbot est constituée de plusieurs ondes se propageant dans des directions différentes. L'interférence entre ces ondes peut générer un système de franges parasites qui diffracterait la lumière à la restitution pour former la structure trouvée entre les pixels.

- La non linéarité de la réponse de la gélatine : les différents ordres diffractés ont des intensités non proportionnelles à celles de l'image enregistré, par conséquent l'interférence laisse des ordres résiduelles.

Une bonne compréhension de cette structure demande une étude détaillée de la structure des franges enregistrés dans le volume de l'hologramme.

IV-6-4 Structure de l'image enregistrée sur l'hologramme de Talbot

Nous avons vu que la structure de l'image enregistrée sur la plaque holographique correspond à une structure de lobes de diffraction convergentes dans les pixels de l'image. Chaque pixel de l'image reçoit la lumière diffractée de tout l'hologramme, mais essentiellement de la partie correspondant à son lobe de diffraction principal. La structure de l'hologramme s'approche donc de la structure d'une matrice de microlentilles. Nous avons simulé la phase de l'onde dans le plan de moindre dynamique (le plan de l'hologramme) et nous trouvons un profil similaire à celui d'ondes sphérique (Fig. VI-15). Cette faculté nous permet d'utiliser une partie de l'hologramme pour éclairer un nombre réduit de pixels.

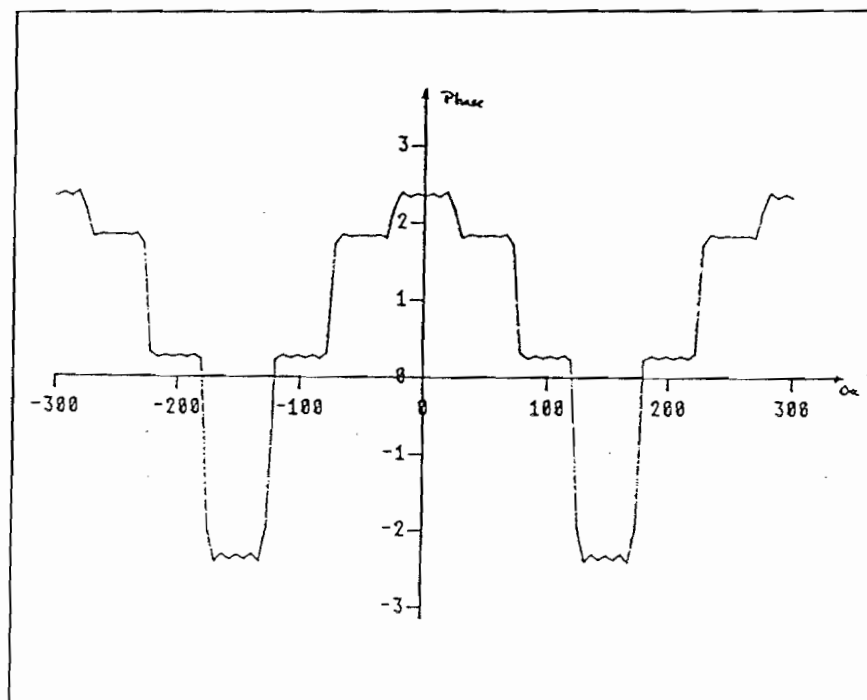


Fig IV-15 : Structure de phase au plan de moindre dynamique.

IV-7 COMPARAISON AVEC LES AUTRES "ILLUMINATEURS DE TABLEAUX"

Notre hologramme illuminateur de tableaux est bien convenable pour l'enregistrement simultané d'un front d'onde reconstituant l'image d'un grand nombre de pixels. Ainsi 100 * 100 pixels ou plus peuvent être enregistrés en un seul étape.

Les réseaux de Dammann et nos hologrammes à effet Talbot, bien qu'utilisant des techniques de fabrications très différentes, sont directement concurrents. L'un et l'autre utilisent la diffraction et sont utilisables pour réaliser des réseaux dont les tailles varient de quelques pixels à plus de 100*100. Notons toutefois trois différences :

- Pour les tailles supérieures à quelques unités, la finesse du contrôle des traits du masque de Dammann devient vite un obstacle sérieux (nettement moins qu'un micromètre pour la position et la largeur des traits);
- Les réseaux de Dammann, placés dans un plan pupillaire, compensent automatiquement le profil du faisceau d'éclairage, en général gaussien, ce que les hologrammes à effet Talbot ne font que très approximativement.
- Enfin, l'efficacité de diffraction de 70% que nous avons obtenue pour notre réseau de 90*90 points excède sensiblement celle des réseaux de Dammann (en général 40%).

Par contre, l'avantage que présente l'emploi des micro-lentilles holographiques par rapport à notre méthode est l'efficacité de diffraction. Dans le cas de microlentilles chaque pixel est enregistré individuellement à partir d'une onde sphérique convergente. Compte-tenu de la simplicité de l'onde sphérique, l'efficacité de diffraction est nettement plus grande que notre hologramme où le front d'onde enregistré a une structure compliquée. L'avantage de notre hologramme réside dans deux points essentiels :

- enregistrement collectif.
- limitations dues à la diffraction plus tolérantes, comme nous allons expliciter ci-dessous :

Pour des raisons de compacité, d'encombrement de montage et de stabilité mécanique, il est souvent désirable que la distance d entre l'illuminateur de tableaux et la matrice de pixels à éclairer soit la plus petite possible. Par conséquent l'illuminateur de tableaux doit avoir une distance focale d la plus petite possible. A ce moment la diffraction joue un rôle limitatif. Considérons par exemple un réseau carré de période p , éclairé par des micro-lentilles de telle sorte que chaque micro-lentille éclaire un pixel. Chaque micro-lentille occupe, par conséquent, une pupille de diamètre p . La demi-largeur de la tache éclairée, à ce moment, est typiquement égale à $\lambda d/p$. Cette tache doit être inférieure à p d'un facteur α de l'ordre de 2 à 10. Il en résulte que la période p doit être au minimum égale à $\sqrt{\alpha \lambda d}$. Par exemple, si $d=10$ cm, $\alpha=5$ et $\lambda=1$ μ m, la séparation entre les micro-lentilles doit être de 700 μ m,

qui est trop grande pour les grands réseaux. Dans le cas de l'hologramme de Talbot, la lumière construisant chaque pixel ne vient pas d'une petite portion de l'hologramme mais peut être diffractée par tout l'hologramme. D'où la possibilité d'éclairer de très petites taches pratiquement non limitées par la diffraction. D'autres limitations peuvent être mises en cause en ce qui concerne l'effet Talbot ou la dimension de l'hologramme : cette analyse fait l'objet du chapitre suivant.

CHAPITRE V

ANALYSE DES ABERRATIONS CHROMATIQUES DE L'HOLOGRAMME A EFFET TALBOT

CHAPITRE V

ANALYSE DES ABERRATIONS CHROMATIQUES DE L'HOLOGRAMME A EFFET TALBOT

V-1- INTRODUCTION

L'hologramme à effet Talbot est enregistré sur gélatine bichromatée par l'intermédiaire d'un laser Argon ($\lambda = 0.488 \mu\text{m}$) dont la longueur d'onde d'émission correspond au domaine spectral de sensibilité du bichromate. L'hologramme peut être restitué avec une longueur d'onde différente, par exemple dans le proche infra-rouge à partir de diode laser utilisée dans le montage final de l'automate cellulaire.

Le changement de la longueur d'onde entre l'enregistrement et la restitution introduit des aberrations chromatiques dont l'influence est gênante si la tache restituée est plus large que la photodiode à éclairer. En effet, il est important que l'éclairage du circuit électronique en dehors des photodiodes soit minimisé, car cet éclairage peut perturber le bon fonctionnement du circuit par excitation des transistors autour des photodiodes. Par ailleurs, cet éclairage représente une perte d'énergie.

Nous présentons dans ce chapitre une analyse détaillée des aberrations introduites par le changement de la longueur d'onde à la restitution.

Dans le § V-3 nous étudions l'onde restituée dans l'approximation de Fresnel pour déterminer les paramètres principaux de la reconstruction de l'image de Talbot, à savoir la distance des différents plans de reconstruction et le décalage de l'image par rapport au masque. Les termes aberrants sont étudiés dans le § V-4 en développant le calcul prépondérant d'aberration et les moyens de minimiser son effet en changeant les conditions d'enregistrement en ce qui concerne l'inclinaison de l'onde éclairant le masque.

Pour estimer l'effet de ces corrections, nous présentons dans le § V-5 les résultats de simulation de la tache restituée par l'hologramme : tout d'abord sans changement de longueur d'onde (voir figure V-6), ensuite l'effet de changement de la longueur d'onde (figure V-7) et enfin l'effet de la correction (figures V-8 à V-10).

Nous récapitulons les résultats de ces études dans le § V-6 pour juger l'efficacité de la correction des aberrations. A la fin du chapitre nous suggérons une autre méthode pour corriger ces aberrations.

V-2- MODELISATION MATHEMATIQUE

L'objectif de cette modélisation est de déterminer l'onde reconstituant l'image de Talbot dans l'approximation scalaire, et sans avoir recours à l'approximation de Fresnel.

A partir de la définition du masque, nous suivons la propagation de l'onde diffractée jusqu'au plan de l'hologramme. L'enregistrement de l'hologramme est présenté sous forme d'interférence entre l'onde diffractée et l'onde référence. Ensuite, nous simulons la restitution de l'hologramme par une autre longueur d'onde dont l'inclinaison est déterminée selon la condition de Bragg pour l'onde diffractée dans l'ordre 1. L'amplitude de l'onde image dans un plan quelconque est déduite par la propagation de l'onde restituée.

Nous avons présenté dans le paragraphe 3-2-1 du chapitre IV un moyen commode pour la propagation des ondes à partir de leur transformée de Fourier. Nous rappelons ici que la transformée de Fourier de l'amplitude d'une onde dans un plan (x,y) représente sa décomposition en ondes planes. La relation entre les coordonnées dans l'espace de Fourier $\vec{\omega}$ et le vecteur des cosinus directeurs $\vec{\Omega}_{xy}$ de l'onde plane suivant x et y est donnée par :

$$\vec{\omega} = \lambda \cdot \vec{\Omega}_{xy}$$

V-2-1 Description de l'objet

Nous considérons, pour la simplicité de l'analyse, un objet périodique bidimensionnel à ouverture de pupille infinie. Cet objet est constitué de trous distribués sur un réseau carré de période p.

La transmission de cet objet est mathématiquement présentée par la formule suivante :

$$t(\vec{r}) = m(\vec{r}) * \left(\frac{1}{p} \Psi\left(\frac{x}{p}\right) \times \frac{1}{p} \Psi\left(\frac{y}{p}\right) \right)$$

$m(\vec{r})$ représente la forme de chaque trou, et Ψ est le peigne de Dirac.

La transformée de Fourier de cette transmittance est donnée par :

$$\tilde{t}(\vec{\omega}) = \tilde{m}(\vec{\omega}) \cdot \Psi(p\mu) \cdot \Psi(p\nu)$$

La transformée de Fourier de l'objet comporte une série périodique de pics de Dirac sur une maille carrée. Chaque pic de Dirac est pondéré par la valeur de la transformée de Fourier du trou $\tilde{m}(\vec{\omega})$ en ce point. Ce qui revient à discrétiser la fonction $\tilde{m}(\vec{\omega})$ par les pics de Dirac. Cette interprétation est équivalente au développement en série de Fourier de la transmittance de l'objet.

V-2-2 Onde objet

Nous considérons que l'onde d'éclairage, de longueur d'onde λ_1 est une onde plane inclinée dont le vecteur d'onde est situé dans le plan (x,z) faisant un angle θ_0 avec l'axe z.

Le vecteur des cosinus directeurs de cette onde est défini par :

$$\vec{\Omega}_0 = \begin{cases} \alpha_0 = \sin \theta_0 \\ \beta_0 = 0 \\ \gamma_0 = \cos \theta_0 \end{cases}$$

Le vecteur d'onde est donné par :

$$\vec{k}_0 = \frac{2\pi}{\lambda_1} \cdot \vec{\Omega}_0 = \begin{cases} \frac{2\pi}{\lambda_1} \cdot \alpha_0 \\ \frac{2\pi}{\lambda_1} \cdot \beta_0 \\ \frac{2\pi}{\lambda_1} \cdot \gamma_0 \end{cases}$$

L'amplitude de cette onde dans le plan du masque ($z=0$) est décrite par :

$$U_{\text{incident}}(\vec{r}, z=0) = \exp(i \vec{k}_0 \cdot \vec{r}) = \exp\left(2i\pi \cdot \frac{\vec{\Omega}_0}{\lambda_1} \cdot \vec{r}\right)$$

Après le masque, l'onde transmise, appelée onde objet, est :

$$U_{\text{incident}}(\vec{r}, 0) \cdot t(\vec{r})$$

La transformée de Fourier de cette onde s'écrit :

$$\begin{aligned}
 U_{\text{transmise}}(\vec{\omega}, 0) &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \exp(2i\pi\vec{\omega}_0 \cdot \vec{r}) \cdot t(\vec{r}) \cdot \exp(-2i\pi\vec{\omega} \cdot \vec{r}) \cdot d\vec{r} \\
 &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} t(\vec{r}) \cdot \exp(2i\pi(\vec{\omega}_0 - \vec{\omega}) \cdot \vec{r}) \cdot d\vec{r} \\
 &= \tilde{t}(\vec{\omega}) * \delta(\vec{\omega} - \vec{\omega}_0) \\
 &= \{\tilde{m}(\vec{\omega}) \cdot \Psi(p\mu) \cdot \Psi(p\nu)\} * \delta(\vec{\omega} - \vec{\omega}_0)
 \end{aligned}$$

Cette expression indique simplement que la transformée de Fourier de l'onde transmise dans le plan de l'objet est la même que celle de l'objet mais décalée du vecteur $\vec{\omega}_0$ avec :

$$\vec{\omega}_0 = \begin{cases} \mu_0 \\ \nu_0 \end{cases} = \begin{cases} \frac{\alpha_0}{\lambda_1} = \frac{i_0}{p} \\ \frac{\beta_0}{\lambda_1} = 0 \end{cases}$$

Pour la commodité de notation, nous définissons i_0 comme suit :

$$\frac{\sin \theta_0}{\lambda_1} = \frac{i_0}{p}$$

Nous interprétons ce décalage par l'inclinaison de l'ordre de diffraction $(i,j) = (0,0)$, ou ordre principal, selon la direction $\vec{\Omega}_0 = \lambda_1 \cdot \vec{\omega}_0$. Par conséquent, ce décalage est nul quand l'onde d'éclairage est normale au masque.

Comme le décalage dans l'espace de Fourier induit seulement un déphasage dans le plan réel, l'onde transmise décrit parfaitement l'objet. Ce résultat évident nous sera utile ultérieurement.

L'espace de Fourier étant composé de points donnés par les vecteurs :

$$\vec{\omega}_{ij} = \begin{cases} \mu_i = \frac{i + i_0}{p} \\ \nu_j = \frac{j}{p} \end{cases}$$

l'onde transmise s'écrit sous forme d'une série de Fourier :

$$U(\vec{r}, 0) = \sum_{ij} \tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(2i\pi\left(\frac{i+i_0}{p}\right)x\right) \cdot \exp\left(2i\pi\left(\frac{j}{p}\right)y\right)$$

Cette formule indique clairement que l'onde transmise par le masque se décompose en somme d'ondes planes. Chaque onde plane a une amplitude complexe :

$$\tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(2i\pi\left(\frac{i+i_0}{p}\right)x\right) \cdot \exp\left(2i\pi\left(\frac{j}{p}\right)y\right)$$

Par conséquent sa direction de propagation est donnée par le vecteur :

$$\vec{\Omega}_{ij} = \begin{cases} \mu_i = \frac{i+i_0}{p} \cdot \lambda_1 \\ \nu_j = \frac{j}{p} \cdot \lambda_1 \end{cases}$$

et sa densité est donnée par $\tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right)$

Remarque : Nous avons supposé dans cette démonstration que $t(\vec{r})$ est à ouverture infinie, c'est pourquoi l'espace de Fourier est composé de points. En utilisant un masque de dimensions finies, il faut convoluer l'expression précédente avec la transformée de Fourier de la pupille du masque, ce qui transforme chaque pic de Dirac en une tache dont les dimensions sont liées à l'ouverture de cette pupille. De ce fait la décomposition en somme d'ondes planes n'est plus valable, et l'étude devient plus compliquée. Une étude détaillée sur l'influence de la pupille a été faite dans le chapitre IV. Pour simplifier le calcul, nous nous limitons dans ce chapitre à l'hypothèse d'une ouverture infinie de la pupille.

V-2-3 Amplitude de l'onde dans le plan de l'hologramme

Suivant la décomposition de l'onde transmise par le masque, et en se servant des résultats du chapitre IV § 3-2-1, la propagation de l'onde jusqu'au plan de l'hologramme à la distance z_h donne l'amplitude suivante :

$$U(\vec{r}, z_h) = \sum_{ij} \tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(2i\pi\left(\frac{i+i_0}{p}\right)x\right) \cdot \exp\left(2i\pi\left(\frac{j}{p}\right)y\right) \cdot \exp\left(\frac{2i\pi}{\lambda_1} \sqrt{1-\alpha_i^2-\beta_j^2} \cdot z_h\right)$$

que l'on peut encore écrire sous la forme :

$$U(\vec{r}, z_h) = \sum_{ij} \tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(2i\pi\left(\frac{i+i_0}{p}\right)x\right) \cdot \exp\left(2i\pi\left(\frac{j}{p}\right)y\right) \cdot \exp\left(\frac{2i\pi}{\lambda_1} \sqrt{1-\lambda_1^2\left[\left(\frac{i+i_0}{p}\right)^2 - \left(\frac{j}{p}\right)^2\right]} \cdot z_h\right)$$

Cette onde détermine l'onde objet $\theta = U(\vec{r}, z_h)$ au niveau de la plaque holographique.

V-2-4 Enregistrement de l'hologramme

Nous définissons l'onde de référence, de longueur d'onde λ_1 comme une onde plane inclinée dont le vecteur d'onde est dans le plan (x,z).

Sa direction de propagation est donnée par :

$$\vec{\Omega}_1 = \begin{cases} \alpha_1 = \sin \theta_1 \\ \beta_1 = 0 \\ \gamma_1 = \cos \theta_1 \end{cases}$$

Comme l'onde de référence se trouve dans le plan (x,z), le cosinus directeur β_1 est nul.

L'amplitude complexe de cette onde est donnée par :

$$\mathcal{R}(\vec{r}, z) = \exp\left(2i\pi \frac{\vec{\Omega}_1}{\lambda_1} \cdot \vec{r}\right)$$

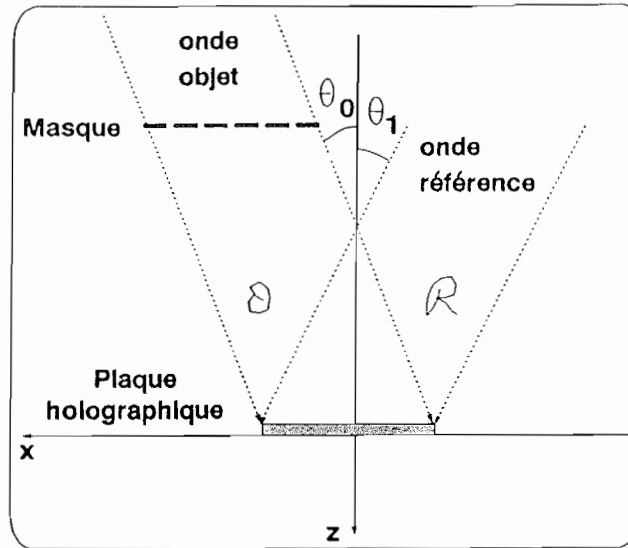


Fig. V-1 : Enregistrement de l'hologramme.

Cette onde interfère avec l'onde objet. L'éclairement du champ d'interférence enregistré sur l'hologramme est donné par :

$$\mathcal{R} + \mathcal{O} \quad \mathcal{R} \mathcal{O} + \mathcal{O} \mathcal{R} + \mathcal{O} \mathcal{O}$$

A la restitution nous nous intéressons uniquement au terme $\mathcal{R}^* \mathcal{O}$ qui est la partie de l'onde objet qui restitue l'image, dont l'expression se présente sous la forme :

$$\begin{aligned} \mathcal{R}^* \mathcal{O} = \sum_{ij} \tilde{m} \left(\frac{i}{p}, \frac{j}{p} \right) \cdot \exp \left(2i\pi \left(\frac{i + i_0}{p} - \frac{\alpha_1}{\lambda_1} \right) x \right) \cdot \exp \left(2i\pi \frac{j}{p} y \right) \\ \cdot \exp \left(\frac{2i\pi}{\lambda_1} \sqrt{1 - \lambda_1^2 \left[\left(\frac{i + i_0}{p} \right)^2 - \left(\frac{j}{p} \right)^2 \right]} \cdot z_h \right) \end{aligned}$$

V-2-5 Restitution de l'hologramme

L'onde de restitution est plane à la longueur d'onde λ_2 qui est, en général, différente de la longueur d'onde de l'enregistrement. L'inclinaison de cette onde dépend de la position voulue de la tache restituée. Pour la tache centrale, le vecteur d'onde de restitution se trouve dans le plan (x,z) faisant un angle θ_2 avec l'axe z, cette inclinaison doit respecter la condition de Bragg donnée dans le chapitre IV, et qui se résume par la formule :

$$\sin\left(\theta_2 - \frac{\theta_0 + \theta_1}{2}\right) = \frac{\lambda_2}{\lambda_1} \cdot \sin\left(\frac{\theta_1 - \theta_0}{2}\right)$$

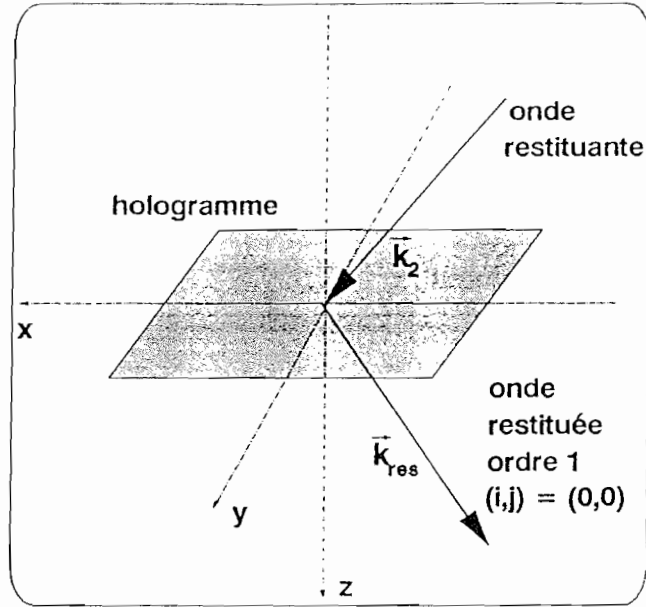


Fig. V-5 : Restitution de l'hologramme.

Le vecteur des cosinus directeurs est défini comme suit :

$$\bar{\Omega}_2 = \begin{cases} \alpha_2 \\ \beta_2 \\ \gamma_2 \end{cases}$$

et le vecteur d'onde est donné par la relation :

$$\bar{k}_2 = \frac{2\pi}{\lambda_2} \cdot \bar{\Omega}_2$$

et l'amplitude complexe de cette onde fait intervenir le facteur :

$$\mathcal{O}'(\bar{r}, z_h) = \exp\left(2i\pi \frac{\bar{\Omega}_2}{\lambda_2} \cdot \bar{r}\right)$$

A la restitution nous nous intéressons au terme qui correspond à l'amplitude de l'onde restituée par l'hologramme dans l'ordre utile pour la reconstruction de l'image. L'expression de cette onde est donnée par :

$$U(\vec{r}, z_h) = \sum_{ij} \tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(2i\pi\left(\frac{i+i_0}{p} - \frac{\alpha_1}{\lambda_1} + \frac{\alpha_2}{\lambda_2}\right)x\right) \cdot \exp\left(2i\pi\left(\frac{j}{p} + \frac{\beta_2}{\lambda_2}\right)y\right) \\ \cdot \exp\left(\frac{2i\pi}{\lambda_1} \sqrt{1 - \lambda_1^2 \left[\left(\frac{i+i_0}{p}\right)^2 - \left(\frac{j}{p}\right)^2\right]} \cdot z_h\right)$$

La transformée de Fourier de cette onde s'écrit sous la forme :

$$U(\vec{\omega}, z_h) = \left\{ \left[\tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(\frac{2i\pi}{\lambda_1} \sqrt{1 - \lambda_1^2 \left[\left(\frac{i+i_0}{p}\right)^2 - \left(\frac{j}{p}\right)^2\right]} \cdot z_h\right) \right] \cdot \Psi(p\mu) \cdot \Psi(p\nu) \right\} \\ * \delta(\vec{\omega} - \vec{\omega}_{res})$$

tel que :

$$\vec{\omega}_{res} = \begin{cases} \mu_i = \left(\frac{i_0}{p} - \frac{\alpha_1}{\lambda_1} + \frac{\alpha_2}{\lambda_2}\right) \\ \nu_j = \left(\frac{j}{p} + \frac{\beta_2}{\lambda_2}\right) \end{cases}$$

L'espace de Fourier est composé de points discrets dont les coordonnées sont :

$$\vec{\omega}_{ij} = \begin{cases} \mu_i = \left(\frac{i+i_0}{p} - \frac{\alpha_1}{\lambda_1} + \frac{\alpha_2}{\lambda_2}\right) \\ \nu_j = \left(\frac{j}{p} + \frac{\beta_2}{\lambda_2}\right) \end{cases}$$

La valeur de cette transformée pour ces points est donnée par :

$$U(\vec{\omega}_{ij}, z_h) = \tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(\frac{2i\pi}{\lambda_1} \sqrt{1 - \lambda_1^2 \left[\left(\frac{i+i_0}{p}\right)^2 - \left(\frac{j}{p}\right)^2\right]} \cdot z_h\right)$$

Nous voyons que la distribution des pics est décalée par rapport au masque de la valeur $\vec{\omega}_{res}$, ceci se traduit par une inclinaison de l'ordre principal diffracté par l'hologramme. Cette inclinaison est donnée par le vecteur :

$$\vec{\Omega}_{\text{res}} = \begin{cases} \alpha_{\text{res}} = \lambda_2 \cdot \left(\frac{i_0}{p} - \frac{\alpha_1}{\lambda_1} + \frac{\alpha_2}{\lambda_2} \right) \\ \beta_{\text{res}} = \beta_{\text{res}} \\ \gamma_{\text{res}} = \sqrt{1 - \alpha_{\text{res}}^2 - \beta_{\text{res}}^2} \end{cases}$$

dont le vecteur d'onde s'écrit comme suit :

$$\vec{k}_{\text{res}} = \frac{2\pi}{\lambda_2} \cdot \vec{\Omega}_{\text{res}}$$

L'ordre 1 diffracté par l'hologramme est celui qui détermine la position de l'image.

V-2-6 Propagation jusqu'au plan de cote z

Selon l'expression précédente de l'onde restituée et sa transformée de Fourier, son amplitude dans un plan à la distance z est donnée par (voir chapitre IV § 3-2-1) :

$$U(\vec{r}, z_h) = \sum_{ij} \tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(2i\pi\left(\frac{i}{p} + \frac{\alpha_{\text{res}}}{\lambda_2}\right)x\right) \cdot \exp\left(2i\pi\left(\frac{j}{p} + \frac{\beta_{\text{res}}}{\lambda_2}\right)y\right) \\ \cdot \exp\left(\frac{2i\pi}{\lambda_1} \sqrt{1 - \lambda_1^2 \left[\left(\frac{i + i_0}{p}\right)^2 + \left(\frac{j}{p}\right)^2\right]} \cdot z_h\right) \\ \cdot \exp\left(\frac{2i\pi}{\lambda_2} \sqrt{1 - \lambda_2^2 \left[\left(\frac{i}{p} + \frac{\alpha_{\text{res}}}{\lambda_2}\right)^2 + \left(\frac{j}{p} + \frac{\beta_{\text{res}}}{\lambda_2}\right)^2\right]} \cdot (z - z_h)\right)$$

Cette formule donne l'amplitude complexe de l'onde reconstruisant l'image dans le plan z suivant l'approximation scalaire, sans approximation de Fresnel.

Cette formule peut être arrangée sous la forme :

$$U(\vec{r}, z_h) = \sum_{ij} \tilde{m}''\left(\frac{i}{p}, \frac{j}{p}\right) \cdot \exp\left(2i\pi\left(\frac{i}{p} + \frac{\alpha_{\text{res}}}{\lambda_2}\right)x\right) \cdot \exp\left(2i\pi\left(\frac{j}{p} + \frac{\beta_{\text{res}}}{\lambda_2}\right)y\right)$$

avec :

$$\tilde{m}''(\mu, \nu) = \tilde{m}(\mu, \nu) \cdot \tilde{A}(\mu, \nu)$$

et

$$\begin{aligned} \tilde{A}\left(\frac{i}{p}, \frac{j}{p}\right) = & \exp\left(\frac{2i\pi}{\lambda_1} \sqrt{1 - \lambda_1^2 \left[\left(\frac{i + i_0}{p}\right)^2 + \left(\frac{j}{p}\right)^2\right]} \cdot z_h\right) \\ & \cdot \exp\left(\frac{2i\pi}{\lambda_2} \sqrt{1 - \lambda_2^2 \left[\left(\frac{i}{p} + \frac{\alpha_{res}}{\lambda_2}\right)^2 + \left(\frac{j}{p} + \frac{\beta_{res}}{\lambda_2}\right)^2\right]} \cdot (z - z_h)\right) \end{aligned}$$

La transformée de Fourier de cette onde s'écrit :

$$U(\vec{\omega}, z_t) = [\tilde{m}''(\vec{\omega}) \cdot \Psi(p\mu) \cdot \Psi(p\nu)] * \delta(\vec{\omega} - \vec{\omega}_{res})$$

Cette transformée de Fourier se présente toujours sous forme d'une distribution des peignes de Dirac orthogonaux mais décalée par rapport à l'origine de la quantité $\vec{\omega}_{res}$.

Et nous savons que ce décalage dans l'espace de Fourier n'agit pas sur la qualité de l'image. Donc nous obtenons toujours un réseau carré de motifs. Par contre, la transformée de Fourier de ce motif répété a changé de $\tilde{m}(\mu, \nu)$ à $\tilde{m}''(\mu, \nu)$

La forme de la tache est en fait la convolution entre la forme du trou initiale et la fonction $A(x, y)$ la transformée de Fourier inverse de $\tilde{A}$. Toute l'étude qui suit portera sur la fonction $\tilde{A}$, qui est le terme responsable de la déformation de la tache restituée.

V-3 APPROXIMATION DE FRESNEL ET RECONSTRUCTION DE L'IMAGE

V-3-1 Etude approximative du terme $\tilde{A}$

La formule précédemment déduite est valable pour toutes les inclinaisons des angles dans l'approximation scalaire des ondes. En fait, les termes exponentiels sont mêmes valables aux grands angles, en dehors de l'approximation de Fresnel. Toutefois, nous nous intéressons essentiellement à des objets dont la période p est relativement grande par rapport à la longueur d'onde. Ainsi, sans faire d'hypothèse sur les inclinaisons des trois ondes

d'éclairage, de référence et de restitution, nous pouvons procéder à un développement limité des racines par rapport aux ordres i, j de la série de Fourier de l'objet périodique. Ce développement est justifié dans les deux conditions :

- $\lambda i/p$ et $\lambda j/p$ sont, pour les ordres les plus importants, petits devant l'unité.
- Pour i et j élevés, $\tilde{m}\left(\frac{i}{p}, \frac{j}{p}\right)$ a des valeurs négligeables. Cette condition est vérifiée dans le cas où l'ouverture de chaque trou est suffisamment grande pour que les lobes de diffraction soient limités dans l'espace.

Dans ce paragraphe, nous travaillons dans ces hypothèses.

Pour clarifier le développement, nous écrivons le terme $\tilde{d}$ sous la forme :

$$\tilde{d}\left(\frac{i}{p}, \frac{j}{p}\right) = \exp\left(\frac{2i\pi}{\lambda_1} \cdot R_1 \cdot z_h\right) \cdot \exp\left(\frac{2i\pi}{\lambda_2} \cdot R_2 \cdot (z - z_h)\right)$$

avec :

$$R_1 = \sqrt{1 - \lambda_1^2 \left[\left(\frac{i + i_0}{p} \right)^2 + \left(\frac{j}{p} \right)^2 \right]}$$

$$R_2 = \sqrt{1 - \lambda_2^2 \left[\left(\frac{i}{p} + \frac{\alpha_{res}}{\lambda_2} \right)^2 + \left(\frac{j}{p} + \frac{\beta_{res}}{\lambda_2} \right)^2 \right]}$$

Comme les termes infinitésimaux pour le développement sont $\lambda i/p$ et $\lambda j/p$, il faut réécrire les deux termes R_1 et R_2 de telle sorte que nous puissions isoler ces termes infinitésimaux.

$$R_1 = \cos \theta_0 \sqrt{1 - \frac{\lambda_1^2 \cdot \left[\left(\frac{i}{p} \right)^2 + \left(\frac{j}{p} \right)^2 + 2 \cdot \frac{i_0 \cdot i}{p^2} \right]}{\cos^2 \theta_0}}$$

$$R_2 = \gamma_{\text{res}} \sqrt{1 - \frac{\lambda_2^2 \cdot \left[\left(\frac{i}{p} \right)^2 + \left(\frac{j}{p} \right)^2 + 2 \cdot \frac{i}{p} \frac{\alpha_{\text{res}}}{\lambda_2} + 2 \cdot \frac{j}{p} \frac{\beta_{\text{res}}}{\lambda_2} \right]}{\gamma_{\text{res}}^2}}$$

avec :

- $\cos \theta_0 = \sqrt{1 - \left(\frac{\lambda_1 i_0}{p} \right)^2}$ où θ_0 est l'angle d'inclinaison du faisceau objet par rapport à la normale au masque.

- $\alpha_{\text{res}} = \lambda_2 \left(\frac{i_0}{p} - \frac{\alpha_1}{\lambda_1} + \frac{\alpha_2}{\lambda_2} \right)$ est le cosinus directeur selon l'axe x de l'ordre $(i, j = 0, 0)$, qui est l'ordre 1 diffracté par l'hologramme.

- $\beta_{\text{res}} = \beta_2$ est le cosinus directeur selon l'axe y de l'ordre 1 diffracté par l'hologramme.

- $\gamma_{\text{res}} = \sqrt{1 - \alpha_{\text{res}}^2 - \beta_{\text{res}}^2}$ est le cosinus directeur selon l'axe z de l'ordre 1 diffracté par l'hologramme.

Dans le cas où le faisceau de restitution se trouve dans le même plan que les faisceaux objet et référence (plan (x, z)), et en notant θ_2 l'angle que fait ce faisceau avec l'axe z (voir Fig. V-3), l'ordre 1 diffracté par l'hologramme se trouve dans le même plan et fait l'angle θ_{res} avec l'axe z.

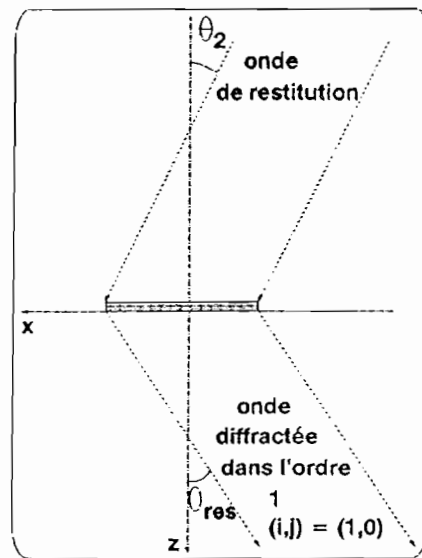


Fig V-3 : Restitution dans le plan (x, z) .

Dans ces condition nous aurons :

$$\begin{aligned}\alpha_2 &= \cos \theta_2 & \alpha_{\text{res}} &= \sin \theta_{\text{res}} \\ \beta_2 &= 0 & \beta_{\text{res}} &= 0 \\ \gamma_2 &= \cos \theta_2 & \gamma_{\text{res}} &= \cos \theta_{\text{res}}\end{aligned}$$

V-3-2 Approximation de Fresnel

Dans l'approximation de Fresnel, nous supposons que les angles des ordres diffractés sont petits. Cette condition est vérifiée quand l'ouverture des trous est suffisamment grande. A ce moment, nous pouvons limiter le développement de R_1 et R_2 au premier ordre.

Après le développement, les deux termes s'écrivent sous la forme:

$$\begin{aligned}R_1 &= \cos \theta_0 \left\{ 1 - \frac{\lambda_1^2}{2 \cos^2 \theta_0} \left[\left(\frac{i}{p} \right)^2 + \left(\frac{j}{p} \right)^2 + 2 \frac{i_0 i}{p^2} \right] \right\} \\ R_2 &= \gamma_{\text{res}} \left\{ 1 - \frac{\lambda_2^2}{2 \gamma_{\text{res}}^2} \left[\left(\frac{i}{p} \right)^2 + \left(\frac{j}{p} \right)^2 + 2 \frac{i}{p} \frac{\alpha_{\text{res}}}{\lambda_2} + 2 \frac{j}{p} \frac{\beta_{\text{res}}}{\lambda_2} \right] \right\}\end{aligned}$$

L'onde reconstituant l'image devient :

$$\begin{aligned}U(\vec{r}, z_T) &= \sum_{ij} \tilde{m} \left(\frac{i}{p}, \frac{j}{p} \right) \cdot \exp \left\{ 2i\pi \left[\frac{i}{p} + \frac{\alpha_{\text{res}}}{\lambda_2} \right] \left[x - \left(\text{tg } \theta_0 \cdot z_h + \frac{\alpha_{\text{res}}}{\gamma_{\text{res}}} (z - z_h) \right) \right] \right\} \\ &\quad \cdot \exp \left\{ 2i\pi \left[\frac{j}{p} + \frac{\beta_{\text{res}}}{\lambda_2} \right] \left[y - \frac{\beta_{\text{res}}}{\gamma_{\text{res}}} (z - z_h) \right] \right\} \\ &\quad \cdot \exp \left\{ -i\pi \frac{1}{p^2} \left[\frac{\lambda_1}{\cos \theta_0} z_h + \frac{\lambda_2}{\gamma_{\text{res}}} (z - z_h) \right] \cdot (i^2 + j^2) \right\}\end{aligned}$$

Cette amplitude de l'onde reconstruisant l'image de restitution peut être écrite sous la forme suivante :

$$U(\vec{r}, z_h) = \sum_{ij} \tilde{m}'' \left(\frac{i}{p}, \frac{j}{p} \right) \cdot \exp \left\{ i\pi \cdot \left[\left(\frac{i}{p} + \frac{\alpha_{\text{res}}}{\lambda_2} \right) \cdot X + \left(\frac{j}{p} + \frac{\beta_{\text{res}}}{\lambda_2} \right) \cdot Y \right] \right\}$$

Ceci n'est en fait que le développement en série de Fourier d'un réseau carré avec un motif dont la transformée de Fourier est $\tilde{m}''(\mu, \nu)$ tel que :

$$\tilde{m}''(\mu, \nu) = \tilde{m}(\mu, \nu) \cdot \tilde{A}(\mu, \nu)$$

avec :

$$\tilde{A}\left(\frac{i}{p}, \frac{j}{p}\right) = \exp\left\{-i\pi \frac{1}{p^2} \left(\frac{\lambda_1 \cdot z_h}{\cos \theta_0} + \frac{\lambda_2 \cdot (z - z_h)}{\gamma_{\text{res}}} \right) \cdot (i^2 + j^2) \right\}$$

Ce réseau est dans système d'axe (X,Y) décalé par rapport au système d'axe d'origine (x,y) par le vecteur :

$$\vec{r} = \begin{cases} \text{tg } \theta_0 \cdot z_h + \frac{\alpha_{\text{res}}}{\gamma_{\text{res}}} \cdot (z - z_h) \\ \frac{\beta_{\text{res}}}{\gamma_{\text{res}}} \cdot (z - z_h) \end{cases}$$

Ce décalage est nul dans les conditions suivantes :

$$\theta_0 = 0, \alpha_1 = \alpha_2, \lambda_1 = \lambda_2$$

Ce qui revient au cas d'un hologramme de Talbot avec éclairage normal, enregistré et restitué avec la même longueur d'onde.

V-3-3 Image de Talbot d'ordre 1

Nous remarquons que le terme $\tilde{A}$ est égal à l'unité, quelque soit i et j, si le plan de restitution à la distance z_T^1 vérifie la relation :

$$\frac{1}{2p^2} \left[\frac{\lambda_1 \cdot z_h}{\cos \theta_0} + \frac{\lambda_2 \cdot (z_T^1 - z_h)}{\gamma_{\text{res}}} \right] = 1$$

et par conséquent le motif d'origine est parfaitement restitué. Et le plan de restitution de l'image de Talbot se situe par rapport au plan de l'hologramme à une distance :

$$z_T^1 - z_h = \left(2p^2 - \frac{\lambda_1 \cdot z_h}{\cos \theta_0} \right) \cdot \frac{\gamma_{res}}{\lambda_2}$$

V-3-4 Image de Talbot d'ordre 1/2

Si nous écrivons l'amplitude de l'onde restituant l'image sous la forme :

$$U(\vec{r}, z) = \exp \left(2i\pi \cdot (\gamma_{res} + \beta_{res}) \cdot \frac{p}{2 \cdot \lambda_2} \right) \cdot \sum_{ij} \tilde{m}'' \left(\frac{i}{p}, \frac{j}{p} \right) \\ \cdot \exp \left\{ 2i\pi \left[\left(\frac{i}{p} + \frac{\alpha_{res}}{\lambda_2} \right) \cdot \left(X - \frac{p}{2} \right) + \left(\frac{j}{p} + \frac{\beta_{res}}{\lambda_2} \right) \cdot \left(Y - \frac{p}{2} \right) \right] \right\}$$

La fonction $\tilde{A}$ devient :

$$\tilde{A} \left(\frac{i}{p}, \frac{j}{p} \right) = \exp \left(-i\pi \cdot \frac{1}{p^2} \left(\frac{\lambda_1 z_h}{\cos \theta_0} + \frac{\lambda_2 \cdot (z - z_h)}{\gamma_{res}} \right) \cdot (i^2 + j^2) \right) \cdot \exp(i\pi (i + j))$$

et en mettant la condition pour le plan de restitution :

$$\frac{1}{2p^2} \left\{ \frac{\lambda_1 \cdot z_h}{\cos \theta_0} + \frac{\lambda_2 \cdot (z_T^{1/2} - z_h)}{\gamma_{res}} \right\} = \frac{1}{2}$$

nous retrouvons :

$$\tilde{A} \left(\frac{i}{p}, \frac{j}{p} \right) = \exp \left\{ -i\pi (i^2 - i + j^2 - j) \right\}$$

et comme $i^2 - i$ et $j^2 - j$ sont toujours pairs, la fonction $\tilde{A}$ est encore égale à l'unité. Dans cette condition, nous avons une restitution parfaite de l'image décalée d'une demi-période dans les deux directions. Le plan de restitution se situe à une distance $z_T^{1/2}$ de l'hologramme plus proche que la distance de Talbot. Cette distance est donnée par la relation :

$$z_T^{1/2} - z_h = \left(p^2 - \frac{\lambda_1 \cdot z_h}{\cos \theta_0} \right) \cdot \frac{\gamma_{res}}{\lambda_2}$$

que nous appelons "distance de Talbot d'ordre 1/2".

V-3-5 Image de Talbot d'ordre 0

Un autre moyen pour annuler l'effet du terme aberrant $\tilde{A}$ est d'imposer la condition :

$$\left(\frac{\lambda_1 \cdot z_h}{\cos \theta_0} + \frac{\lambda_2 \cdot (z_T^0 - z_h)}{\gamma_{res}} \right) = 0$$

Ce qui définit un plan à la distance z_T^0 donnée par :

$$z_T^0 - z_h = - \frac{\lambda_1}{\lambda_2} \cdot \frac{\gamma_{res} \cdot z_h}{\cos \theta_0}$$

L'image restituée est située avant l'hologramme. En fait, cette image est l'image conjuguée. Dans le cas où nous restituons avec la même longueur d'onde :

$$\lambda_1 = \lambda_2$$

$$\gamma_{res} = \cos \theta_0$$

$$z_T^0 = 0$$

et nous obtenons le masque lui même.

Le résultat de cette étude montre que dans l'approximation de Fresnel (angles de propagation faibles donc le terme du développement au second ordre est négligeable), l'hologramme restitue parfaitement l'image de Talbot de l'objet initial sans aberration dans plusieurs plans.

Remarque : Le fait que nous pouvons trouver un plan de reconstruction parfait du motif est dû à ce que $\tilde{A}$ est la transformée de Fourier d'une onde sphérique (le facteur de i^2 est le même que celui de j^2). Cette onde sphérique converge en un seul point pour donner un pic de

Dirac. Le motif restitué est la convolution du motif original avec la transformée de Fourier inverse de $\tilde{A}$ qui dans le plan z_T est un Dirac, donc égal au motif original.

V-4 TERMES D'ABERRATION

Au début du paragraphe précédent, nous avons précisé les conditions de l'approximation de Fresnel. Ces conditions ne sont plus vérifiées dès que l'ouverture des trous est petite ce qui rend leur tache de diffraction très étendue, il en est de même pour l'étendue de leur transformée de Fourier. Nous vérifions ces conditions dans l'exemple suivant :

Dans cet exemple, nous considérons le masque qui a servi à la réalisation de l'hologramme, et qui est constitué de trous carrés distribués selon un réseau carré de $300 \mu\text{m}$ de période. Nous avons cité dans le chapitre V que les ordres de diffractions nécessaire pour la reconstitution de l'objet sont compris entre $-28 < i, j < +28$, ce qui représente une étendue non négligeable dans l'espace de Fourier. L'enregistrement et la restitution de l'hologramme se font avec les paramètres suivants :

$$\lambda_1 = 0.488 \mu\text{m}, \lambda_2 = 1 \mu\text{m}, \theta_0 = 0, \theta_1 = 30^\circ$$

Il en résulte suivant la condition de Bragg que $\theta_2 = 47^\circ$. Dans ces conditions la valeur du deuxième terme du développement de R_2 pour $i = 28, j=0$ et pour une distance de reconstruction de l'image égale à 1 cm de l'hologramme, est égale à 5 fois la longueur d'onde de restitution. Cette valeur est trop grande pour que nous puissions négliger les termes du développement au deuxième ordre. Et nous nous attendons à obtenir de fortes aberrations.

V-4-1 Détermination du type prépondérant d'aberration

Nous faisons le développement au deuxième ordre dans le cas où $\beta_2 = 0^\circ$, pour la simplification du calcul, c'est à dire que le faisceau de restitution se trouve dans le plan (x,z) . Dans ce cas, nous avons vu qu'il est possible d'introduire l'angle θ_2 que fait le vecteur d'onde de restitution avec l'axe z (fig. V-3). Dans ces conditions les cosinus directeurs de l'onde de restitution ont les valeurs suivantes :

$$\alpha_2 = \sin \theta_2$$

$$\beta_2 = 0$$

$$\gamma_2 = \cos \theta_2$$

de même ceux de l'onde restitué dans l'ordre (1,0) ont les valeurs suivantes :

$$\alpha_{res} = \sin \theta_{res}$$

$$\beta_{res} = 0$$

$$\gamma_{res} = \cos \theta_{res}$$

Nous ne gardons de ce développement que les termes quadratiques en i et j en considérant que les termes d'ordre supérieur sont négligeables. Nous remarquons que la fonction $\tilde{A}$ devient :

$$\tilde{A}(\mu, \nu) = \exp \left\{ -2i\pi \frac{1}{2p^2} \left(\frac{\lambda_1 \cdot z_h}{\cos^3 \theta_0} + \frac{\lambda_2 (1 + \tan^2 \theta_{res})}{\cos \theta_{res}} (z - z_h) \right) \cdot i^2 \right\} \\ \cdot \exp \left\{ -2i\pi \frac{1}{2p^2} \left(\frac{\lambda_1 \cdot z_h}{\cos^3 \theta_0} + \frac{\lambda_2 (z_T - z_h)}{\cos \theta_{res}} \right) \cdot j^2 \right\} \quad .$$

D'après cette expression, il est évident que le plan de reconstitution selon x est différent de celui selon y . En fait, le plan z_T^x qui neutralise l'influence du terme quadratique en i est déterminé par la relation :

$$\frac{1}{2p^2} \left(\frac{\lambda_1 \cdot z_h}{\cos^3 \theta_0} + \frac{\lambda_2 (1 + \tan^2 \theta_{res})}{\cos \theta_{res}} (z_T^x - z_h) \right) = 1$$

Il est différent du plan z_T^y qui neutralise l'influence du terme quadratique selon j déterminé par la relation :

$$\frac{1}{2p^2} \left(\frac{\lambda_1 \cdot z_h}{\cos^3 \theta_0} + \frac{\lambda_2}{\cos \theta_{res}} (z_T^y - z_h) \right) = 1$$

Cette dernière relation s'identifie à celle qui donne le plan de Talbot précédemment déterminé, ceci est dû au fait que β_2 est nul.

Par exemple dans le plan z_T^y nous nous attendons à une tache bien limitée selon la direction y et étendue selon la direction x , de même dans le plan z nous nous attendons à une tache bien limitée selon la direction x et étendue selon la direction y . Cet astigmatisme entraîne

l'existence d'un plan intermédiaire entre z_T^x et z_T^y où la tache s'approche de la forme circulaire avec des dimensions qui dépassent les dimensions de la tache dans l'objet de Talbot.

V-4-2 Extraction du terme aberrant

D'après l'étude précédente, nous avons vu que le terme responsable des aberrations est la fonction $\bar{A}$ qui est exponentielle quadratique en i et j . Nous représentons cette fonction sous la forme simplifiée suivante :

$$\bar{A}(i, j) = \exp \left\{ -2i\pi \cdot a \cdot \left[\left(1 + \frac{b}{a} \right) \cdot i^2 + \left(1 + \frac{c}{a} \right) \cdot j^2 \right] \right\}$$

tel que :

$$a = \frac{1}{2p^2} \cdot \left(\frac{\lambda_1 \cdot z_h}{\cos \theta_0} + \frac{\lambda_2 \cdot (z - z_h)}{\gamma_{res}} \right)$$

est le terme résultant du développement des racines au premier ordre, c'est le terme qui détermine le plan de reconstruction de l'image selon une des conditions :

$$a = a_0 = 1, 1/2 \text{ ou } 0$$

$$b = \frac{1}{2p^2} \cdot \left(\frac{\lambda_1 \cdot \tan^2 \theta_0 \cdot z_h}{\cos \theta_0} + \frac{\lambda_2 \cdot \alpha_{res}^2 \cdot (z - z_h)}{\gamma_{res}^3} \right)$$

et

$$c = \frac{1}{2p^2} \cdot \left(\frac{\lambda_2 \cdot \beta_{res}^2 \cdot (z - z_h)}{\gamma_{res}^3} \right)$$

qui sont les termes supplémentaires résultant du développement des racines au deuxième ordre.

Comme l'onde de restitution est dans le plan (x, z) , $\beta_{res} = \beta_2$

et par conséquent $c = 0$. Nous nous plaçons dans un plan de reconstruction, sa distance z sera déterminée par l'une des conditions $a = a_0 = 1, 1/2$ ou 0 . La fonction $\bar{A}$ devient :

$$\tilde{A}(i, j) = \exp \left\{ -2i\pi \cdot \left[(1 + b(a_0)) \cdot i^2 + j^2 \right] \right\} \Bigg|_{z \text{ à un plan de reconstruction}}$$

Nous remarquons que si $b(a_0) = 0$ modulo 1, il n'y aura pas d'aberrations dues aux termes quadratiques (en i^2 ou j^2).

La fonction $b(a_0)$ est une fonction de θ_0 (l'angle de l'onde objet par rapport à la normale au masque), et de θ_1 (inclinaison du faisceau référence). θ_1 étant souvent imposé par la géométrie du montage holographique (angle entre les deux bras du montage), le problème se réduit à minimiser $b(a_0)$ modulo 1 en fonction de θ_0 .

Dans les conditions précisées précédemment $b(a_0)$ est donné par la formule :

$$b(a_0) = \frac{1}{2p^2} \left[\lambda_1 \frac{\tan^2 \theta_0}{\cos \theta_0} z_h + \lambda_2 \frac{\tan^2 \theta_{\text{res}}}{\cos \theta_{\text{res}}} (z - z_h) \right] \Bigg|_{z \text{ à un plan de reconstruction}}$$

V-4-3 Etude du terme aberrant

Nous avons étudié la fonction b déterminée précédemment en fonction de l'angle d'enregistrement θ_0 . Cette étude est faite pour plusieurs longueurs d'onde et plusieurs valeurs de l'angle entre le faisceau objet et le faisceau référence.

Nous avons pris pour angle du faisceau référence à l'enregistrement les valeurs suivantes :

$$\theta_1 = 30^\circ$$

$$\theta_1 = 45^\circ$$

$$\theta_1 = 30^\circ + \theta_0$$

$$\theta_1 = 45^\circ + \theta_0$$

Les deux dernières valeurs de θ_1 sont particulièrement intéressantes car l'angle entre les deux bras du montage holographique est fixé à 30° et 45° .

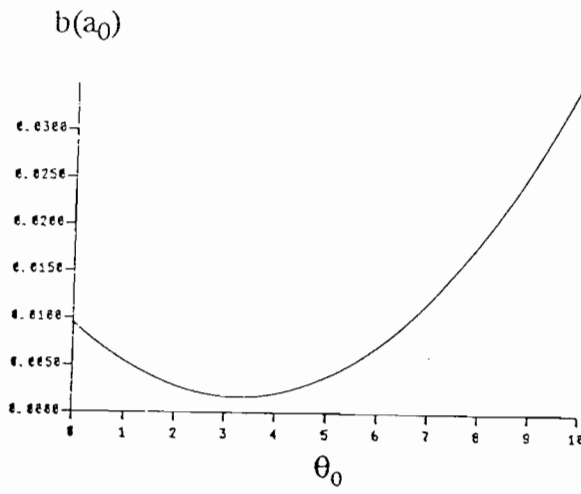
Nous avons encore étudié cette fonction dans les plans de reconstruction de l'image de Talbot d'ordre 1 : z_T^1 et d'ordre 1/2 : $z_T^{1/2}$.

L'étude de cette fonction a porté sur la recherche de son minimum en fonction de θ_0 .

Nous présentons par la suite les résultats de ces études pour certaines longueurs d'onde de restitution. Ces résultats sont présentés sous forme de courbes de variation du terme $b(a_0)$ en fonction de θ_0 .

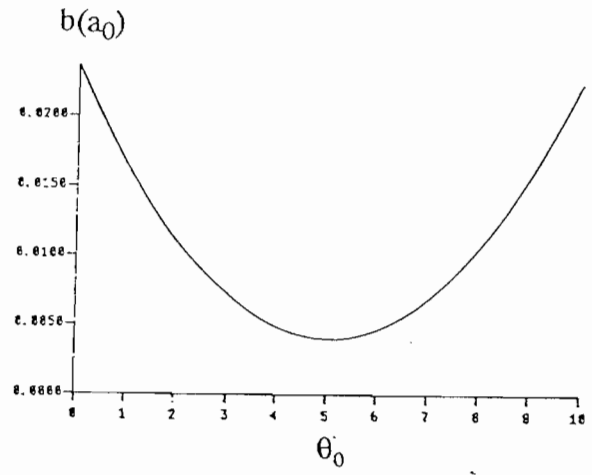
L'étude a été faite pour deux longueurs d'onde :

- $\lambda_2 = 0.6328 \mu\text{m}$ La longueur d'onde du laser He-Ne, qui est un laser souvent utilisé pour les montage de démonstration et d'expérimentation à cause de sa longueur d'onde visible et de sa disponibilité.
- $\lambda_2 = 0.904 \mu\text{m}$ Une longueur d'onde située entre le gap d'énergie de l'AsGa ($0.870 \mu\text{m}$) et celui du Silicium ($1.1 \mu\text{m}$). Cette longueur d'onde est importante dans notre montage d'automate cellulaire car elle permet d'éclairer le circuit électronique du Silicium à travers les modulateur opto-électronique fabriqués sur l'AsGa et qui seront posés sur le circuit Silicium. Cette longueur d'onde correspond à celle des diodes laser implusionnelles de grande puissance disponibles dans le commerce.

Plan de Talbot d'ordre 1

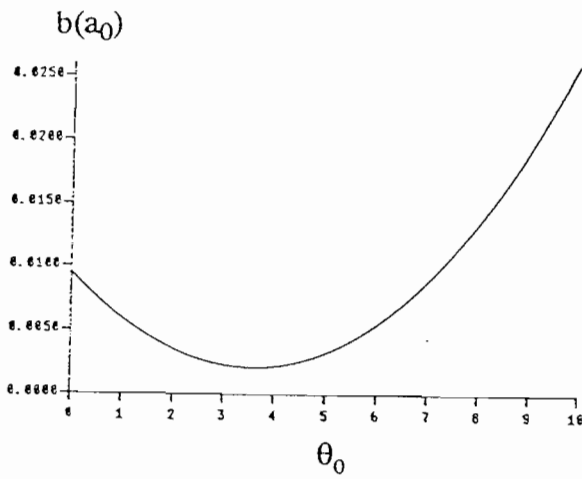
$$\theta_1 = 30^\circ$$

$$\theta_0 \text{ optimal} = 3.3^\circ$$



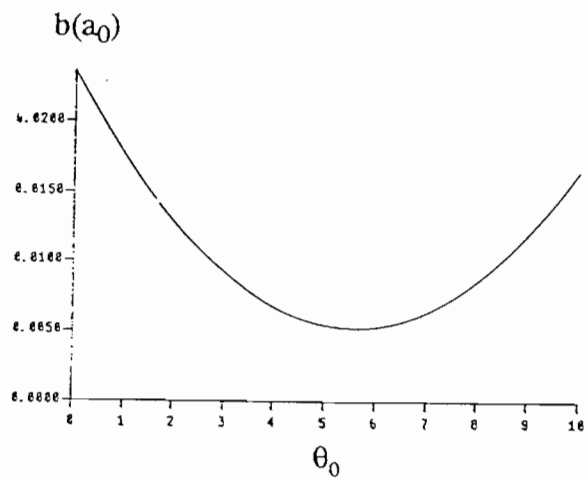
$$\theta_1 = 45^\circ$$

$$\theta_0 \text{ optimal} = 5.1^\circ$$



$$\theta_1 = 30^\circ + \theta_0$$

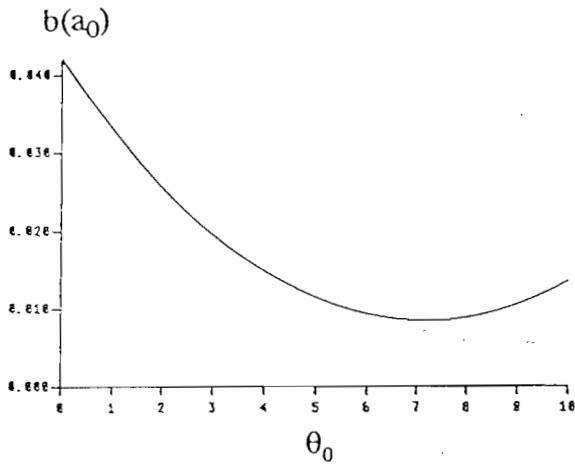
$$\theta_0 \text{ optimal} = 3.6^\circ$$



$$\theta_1 = 45^\circ + \theta_0$$

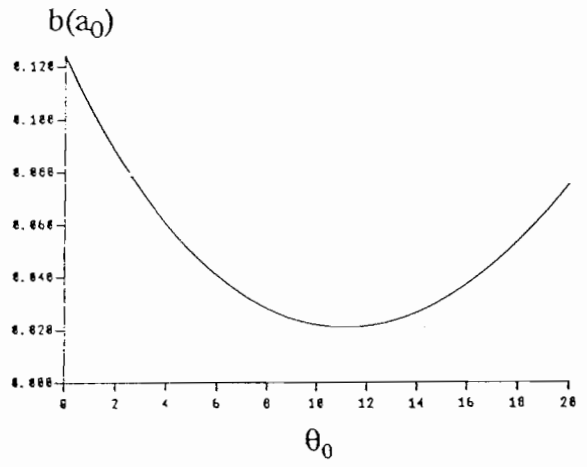
$$\theta_0 \text{ optimal} = 5.6^\circ$$

Fig. V-4-a : $b(a_0)$ pour $\lambda_2 = 0.6328 \mu m$



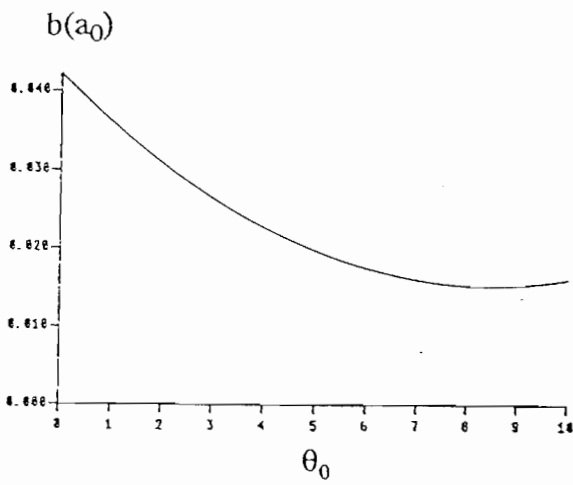
$$\theta_1 = 30^\circ$$

$$\theta_0 \text{ optimal} = 7.2^\circ$$



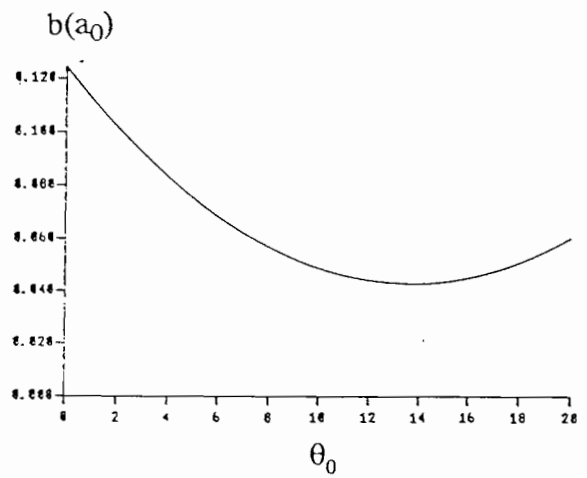
$$\theta_1 = 45^\circ$$

$$\theta_0 \text{ optimal} = 11.1^\circ$$



$$\theta_1 = 30^\circ + \theta_0$$

$$\theta_0 \text{ optimal} = 8.5^\circ$$

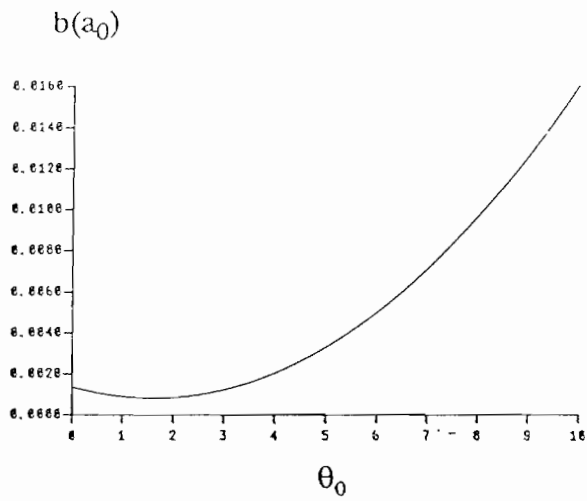


$$\theta_1 = 45^\circ + \theta_0$$

$$\theta_0 \text{ optimal} = 13.7^\circ$$

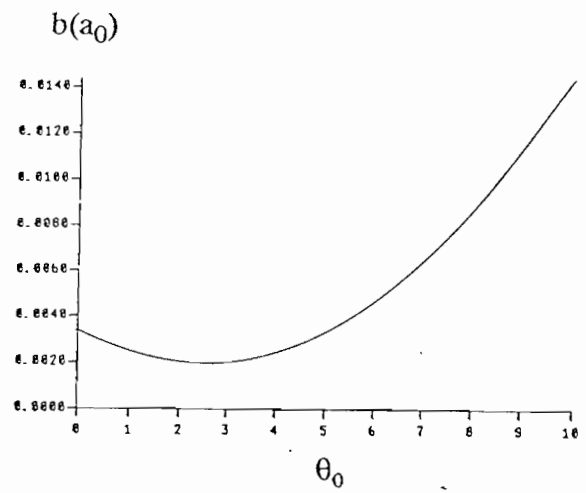
Fig. V-4-b : $b(a_0)$ pour $\lambda_2 = 0.904 \mu m$

Plan de Talbot d'ordre 1/2



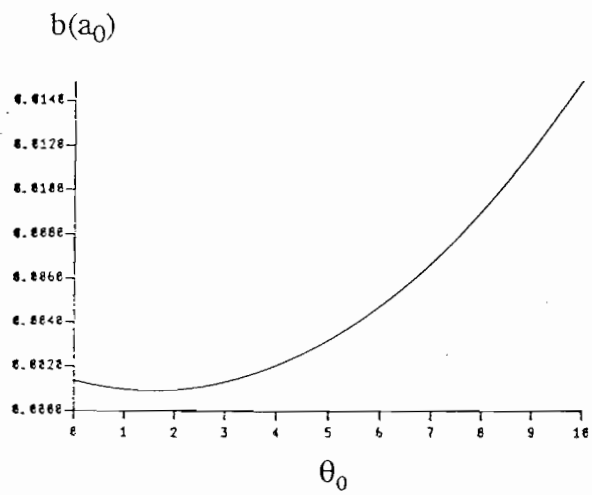
$$\theta_1 = 30^\circ$$

$$\theta_0 \text{ optimal} = 1.6^\circ$$



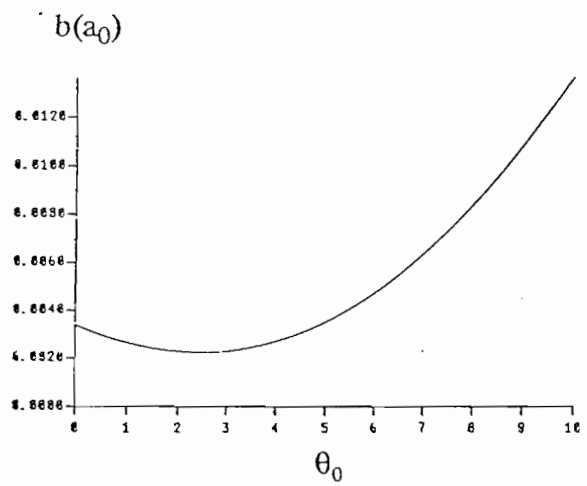
$$\theta_1 = 45^\circ$$

$$\theta_0 \text{ optimal} = 2.5^\circ$$



$$\theta_1 = 30^\circ + \theta_0$$

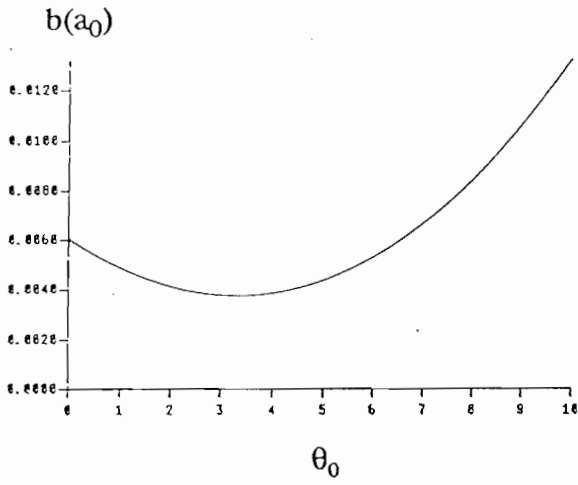
$$\theta_1 \text{ optimal} = 1.5^\circ$$



$$\theta_1 = 45^\circ + \theta_0$$

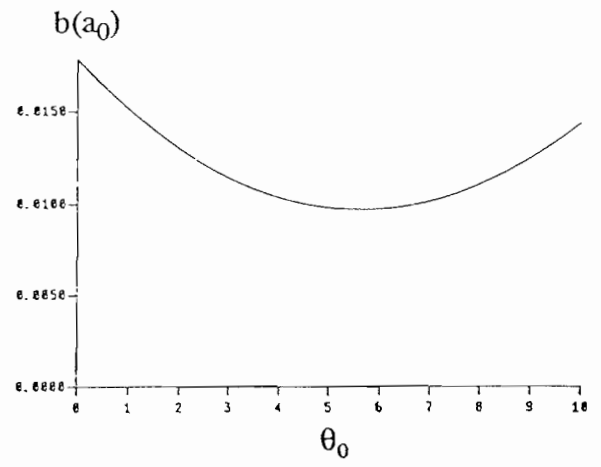
$$\theta_0 \text{ optimal} = 2.4^\circ$$

Fig. V-5-a : $b(a_0)$ pour $\lambda_2 = 0.6328\mu m$



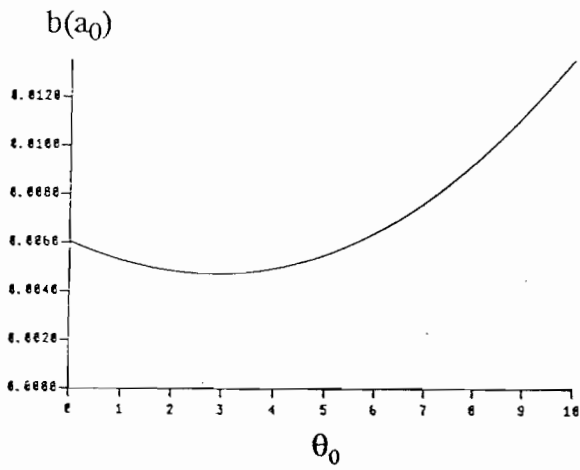
$$\theta_1 = 30^\circ$$

$$\theta_0 \text{ optimal} = 3.3^\circ$$



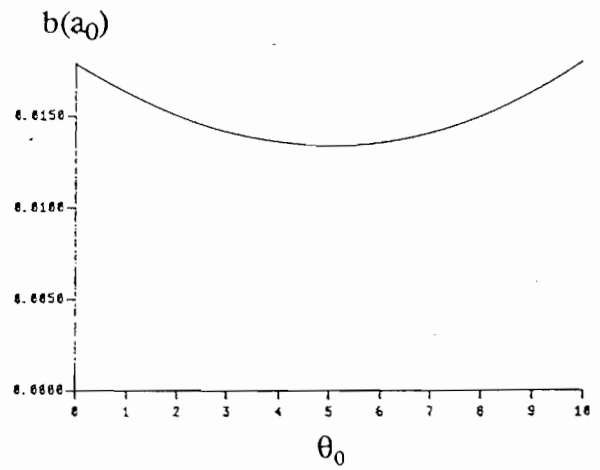
$$\theta_1 = 45^\circ$$

$$\theta_0 \text{ optimal} = 5.6^\circ$$



$$\theta_1 = 30^\circ + \theta_0$$

$$\theta_0 \text{ optimal} = 2.8^\circ$$

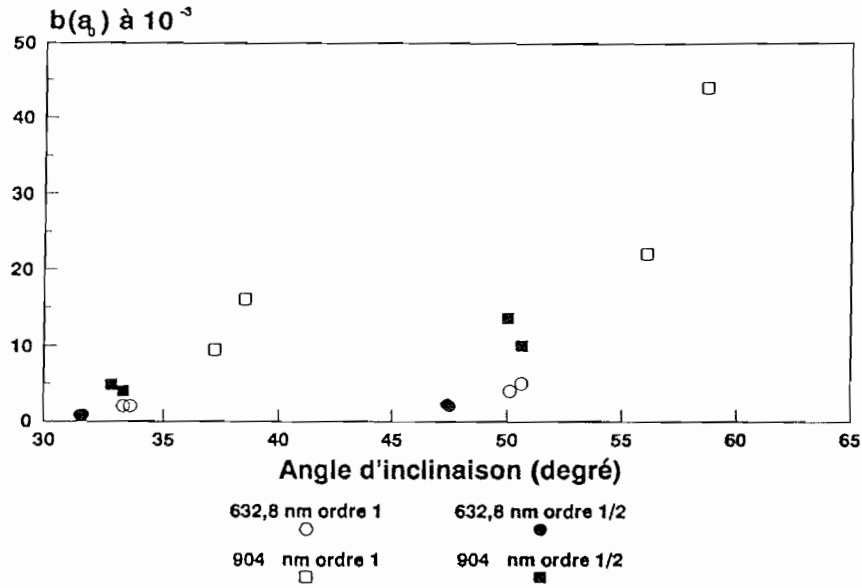


$$\theta_1 = 45^\circ + \theta_0$$

$$\theta_0 \text{ optimal} = 5.0^\circ$$

Fig. V-5-b : $b(a_0)$ pour $\lambda_2 = 0.904 \mu m$

Nous résumons les résultats dans le graphe suivant qui présente les valeurs minimales de $b(a_0)$ en fonction des différentes paramètres :



Nous remarquons que ce minimum est d'autant plus grand que la différence entre les longueurs d'onde est grande. Et il est plus grand pour le plan z_T^1 que pour le plan $z_T^{1/2}$. Par conséquent nous pouvons déduire que la meilleure correction se fait pour la plus faible longueur d'onde, pour le plus faible angle entre faisceau objet et référence, et pour les demi-distances de Talbot.

V-5 SIMULATION

Nous avons simulé l'enregistrement et la restitution de l'hologramme pour l'objet qui nous a servi tout le long de cette étude, à savoir un réseau carré de pixels bidimensionnel infini. La période du réseau est de $300\ \mu\text{m}$, chaque pixel est un carré de $50\ \mu\text{m}$.

La longueur d'onde d'enregistrement est de $488\ \text{nm}$.

Pour produire l'image restituée nous avons utilisé la formule du paragraphe V-2-6 sans développement des racines.

La sommation des termes de la série de Fourier est comprise entre -28 et $+28$.

L'angle de l'onde d'éclairage est variable, tandis que celui du faisceau de référence est fixé à 30° par rapport à la normale de la plaque holographique.

Cette simulation est faite pour plusieurs longueurs d'onde : $632.8\ \text{nm}$, $670\ \text{nm}$, $775\ \text{nm}$, $850\ \text{nm}$ et $904\ \text{nm}$. Nous ne présentons que les résultats les plus significatifs.

L'angle du faisceau de restitution est calculé pour chaque longueur d'onde selon la condition de Bragg pour les hologrammes de phase épais.

Chaque image présente la tache d'un motif restitué dans un carré de $300\text{ }\mu\text{m}$ de côté. Elle est codée sur $128 * 128$ pixels et 256 niveaux de gris.

Comme critère de qualité de l'image restituée, nous avons adopté le rapport de l'énergie contenue dans un carré de $50\text{ }\mu\text{m}$ de côté par rapport à l'énergie totale contenue dans l'image. Cette énergie est calculée à partir de la somme des valeurs des pixels. Ce rapport est désigné par l'intégrale de recouvrement ou IR.

V-1- Restitution avec la longueur d'onde 488 nm

Nous avons commencé tout d'abord à simuler la restitution idéale de l'hologramme avec la même longueur d'onde que l'enregistrement dont nous présentons les résultats par la suite pour la distance de Talbot et la demi-distance de Talbot. D'après ces simulations, il est clair que la restitution de l'image à la demi-distance de Talbot est meilleur que celle à la distance de Talbot.

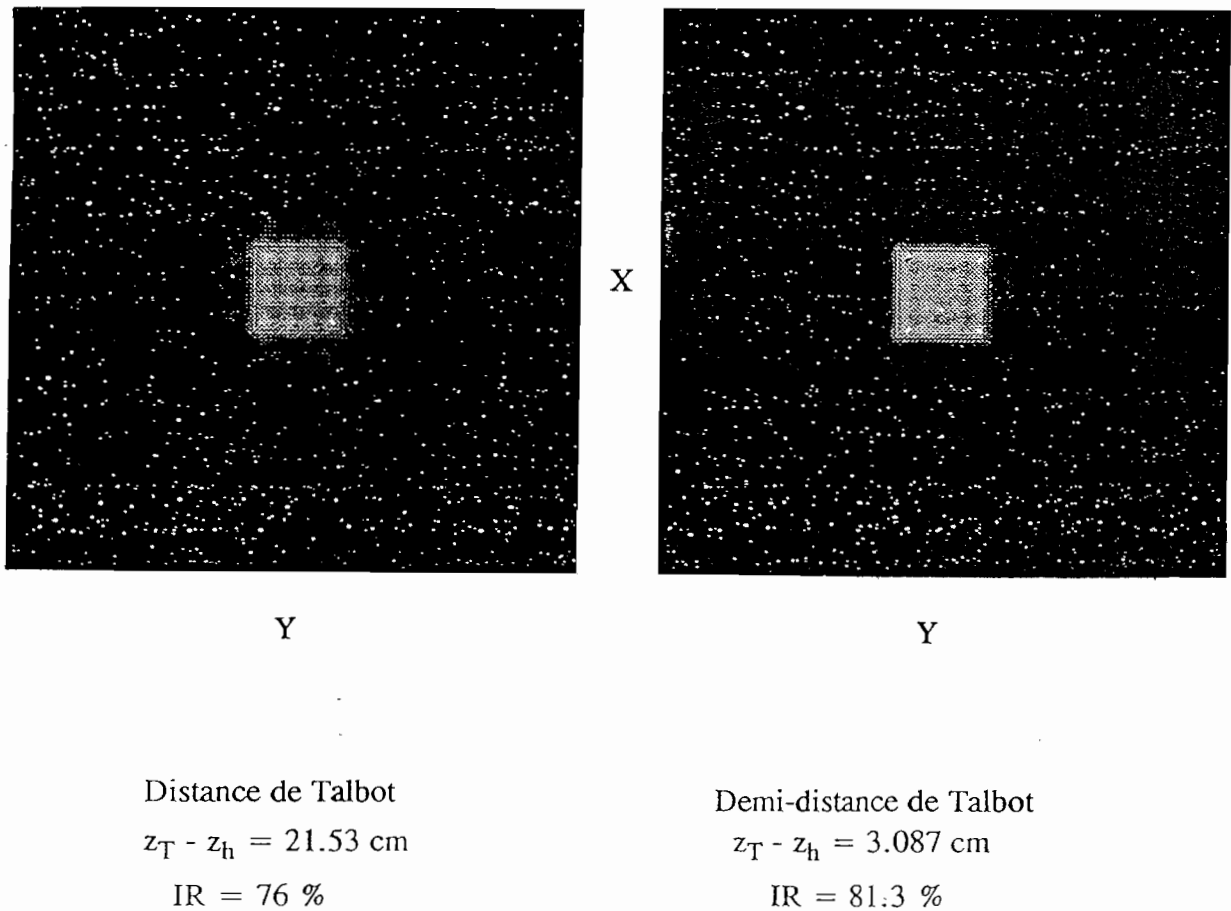
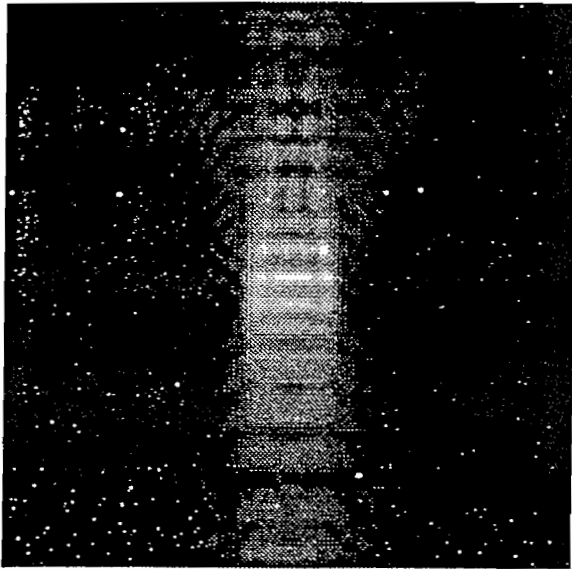


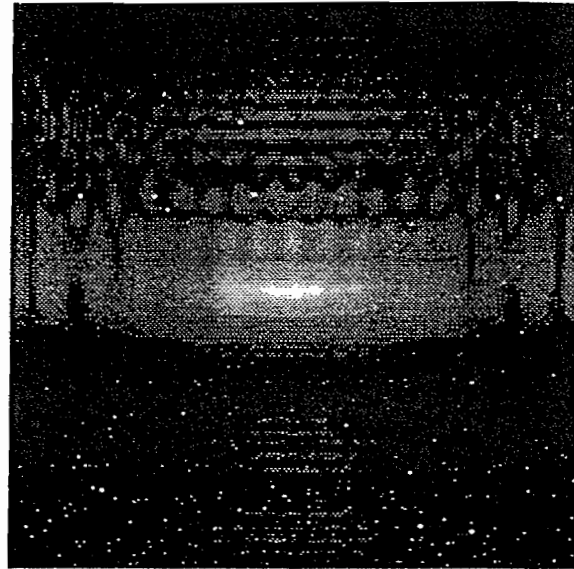
Fig. V-6 : Restitution à $\lambda_2 = 0.488\text{ mm}$, $\theta_2 = 30^\circ$.

V-5-2 Effet d'astigmatisme

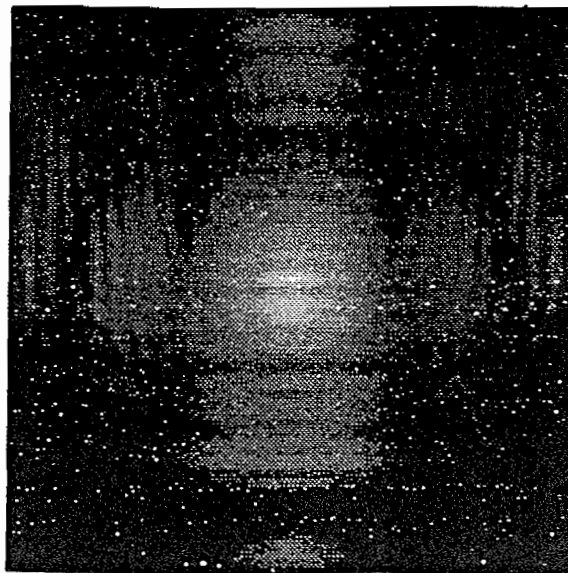
L'effet d'astigmatisme nous a paru le plus net pour une restitution avec la longueur d'onde 904 nm. Nous présentons les résultats des simulations pour cette longueur d'onde dans les plans z_T^i , z_T^j et dans un plan située à mi-distance entre les deux plans précédents.



Plan focal y
 $z_T - z_h = 11.2942 \text{ cm}$
 IR = 23 %



Plan focal x
 $z_T - z_h = 10.6653 \text{ cm}$
 IR = 18 %



Plan intermédiaire
 $z_T - z_h = 10.9797 \text{ cm}$
 IR = 19 %

Fig V-7 : Effet d'astigmatisme.

V-5-3 Correction dans le plan de Talbot

Nous avons simulé la correction en changeant l'inclinaison de l'onde d'éclairage pour différentes longueurs d'onde. Nous présentons les résultats pour les deux longueurs d'onde 632.8 nm et 904 nm. Pour la dernière nous avons présenté différentes inclinaison de l'onde d'éclairage pour voir l'effet de l'angle d'éclairage sur la correction.

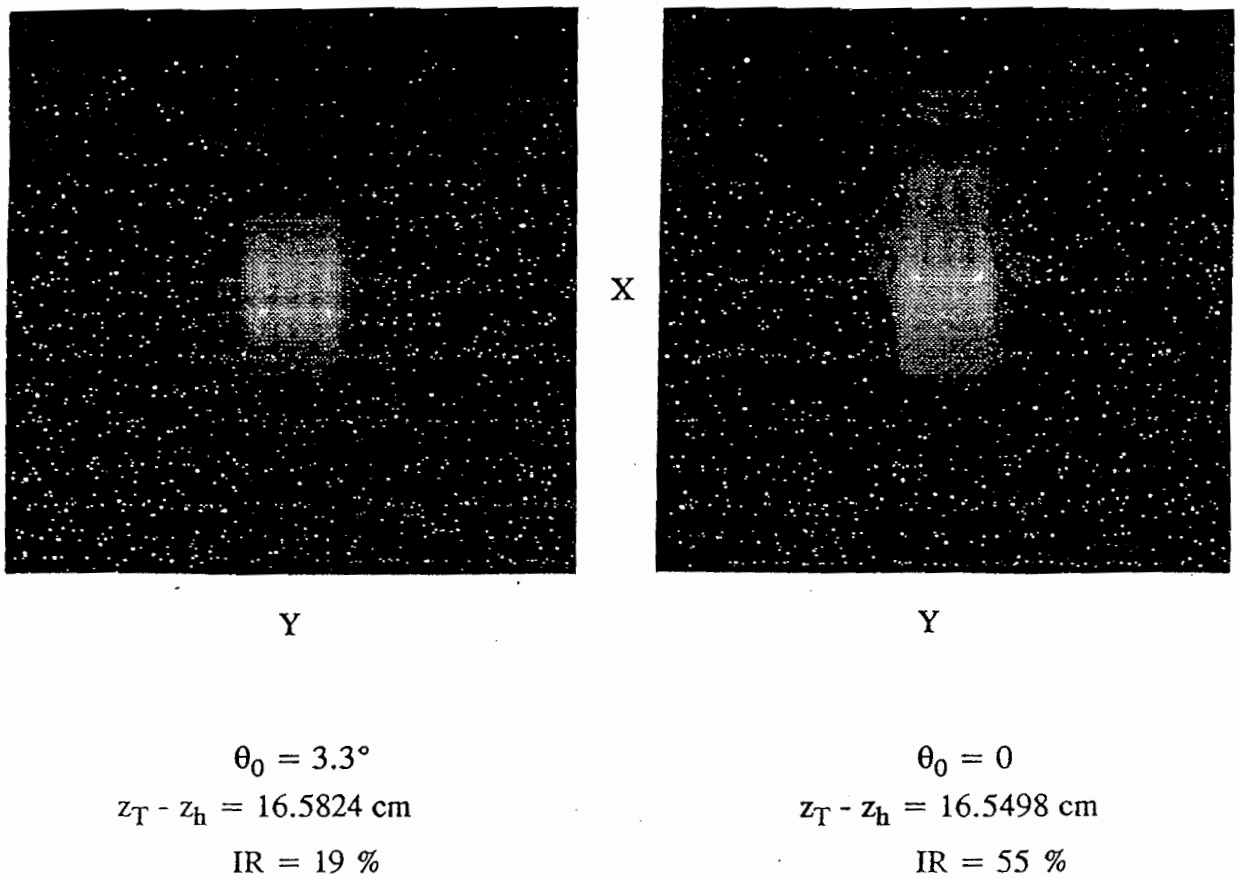


Fig. V-8 (a) : Correction dans le plan de Talbot
 $\lambda_2 = 0.6328 \mu\text{m}$



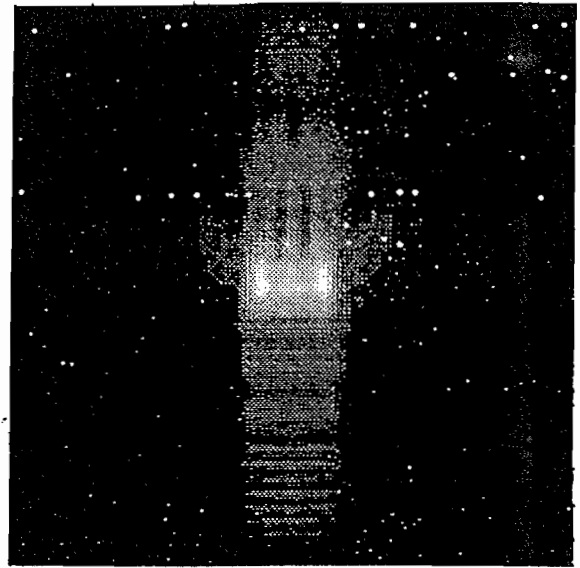
Y

$$\theta_0 = 7.2^\circ$$

$$z_T - z_h = 11.5420 \text{ cm}$$

$$\text{IR} = 44 \%$$

X

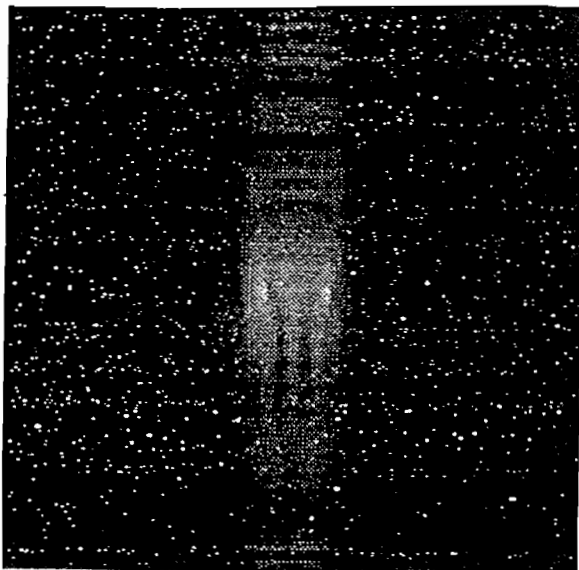


Y

$$\theta_0 = 5^\circ$$

$$z_T - z_h = 11.5244 \text{ cm}$$

$$\text{IR} = 43 \%$$



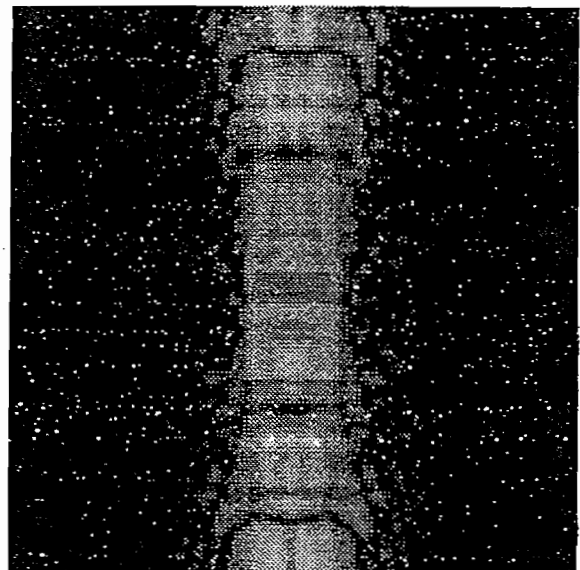
Y

$$\theta_0 = 10^\circ$$

$$z_T - z_h = 11.4919 \text{ cm}$$

$$\text{IR} = 40 \%$$

X



Y

$$\theta_0 = -5^\circ$$

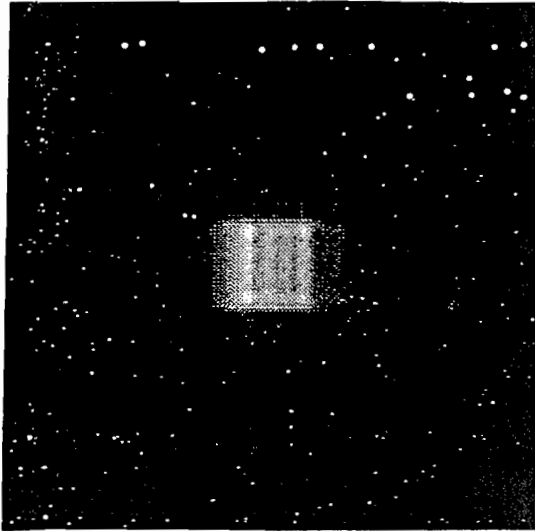
$$z_T - z_h = 10.7953 \text{ cm}$$

$$\text{IR} = 12 \%$$

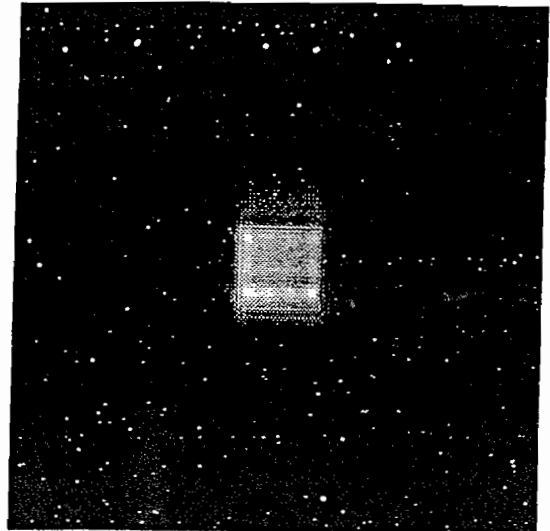
Fig. V-8 (b) : Résultats dans le plan de Talbot avec différents angles
 $\lambda_2 = 0.904 \mu\text{m}$

V-4 Correction dans le demi-plan de Talbot

Nous avons refait les mêmes simulations que précédemment mais à la demi-distance de Talbot.

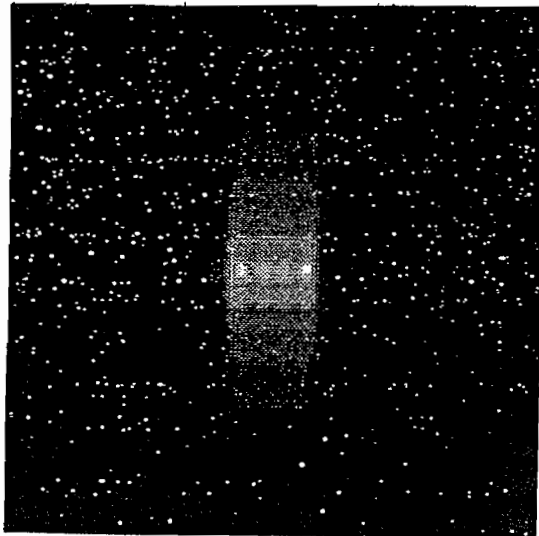


$$\begin{aligned}\theta_0 &= 1.6^\circ \\ z_T - z_h &= 11.4919 \text{ cm} \\ \text{IR} &= 76.06 \%\end{aligned}$$

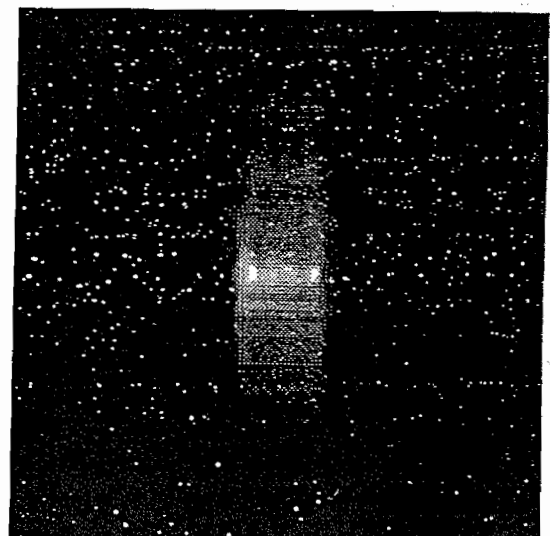


$$\begin{aligned}\theta_0 &= 0^\circ \\ z_T - z_h &= 2.3733 \text{ cm} \\ \text{IR} &= 75.93 \%\end{aligned}$$

Fig. V-9 : Correction dans le demi plan de Talbot
 $\lambda_2 = 0.6328 \mu\text{m}$



$$\begin{aligned}\theta_0 &= 3.3^\circ \\ z_T - z_h &= 1.6341 \text{ cm} \\ \text{IR} &= 40 \%\end{aligned}$$



$$\begin{aligned}\theta_0 &= 0^\circ \\ z_T - z_h &= 1.6196 \text{ cm} \\ \text{IR} &= 12 \%\end{aligned}$$

Fig. V-10 : Correction dans le demi plan de Talbot
 $\lambda_2 = 0.904 \mu\text{m}$

V-6 EFFET DE LA CORRECTION DES ABERRATIONS

Comme mentionné dans le paragraphe précédent, le critère de qualité pour l'image restituée est l'intégrale de recouvrement IR entre la tache restituée et le carré de $50 \times 50 \mu\text{m}$.

Nous présentons dans les tableaux suivants les valeurs de cet intégrale en fonction de la longueur d'onde de restitution pour des images à éclairage normale du masque et des images corrigées.

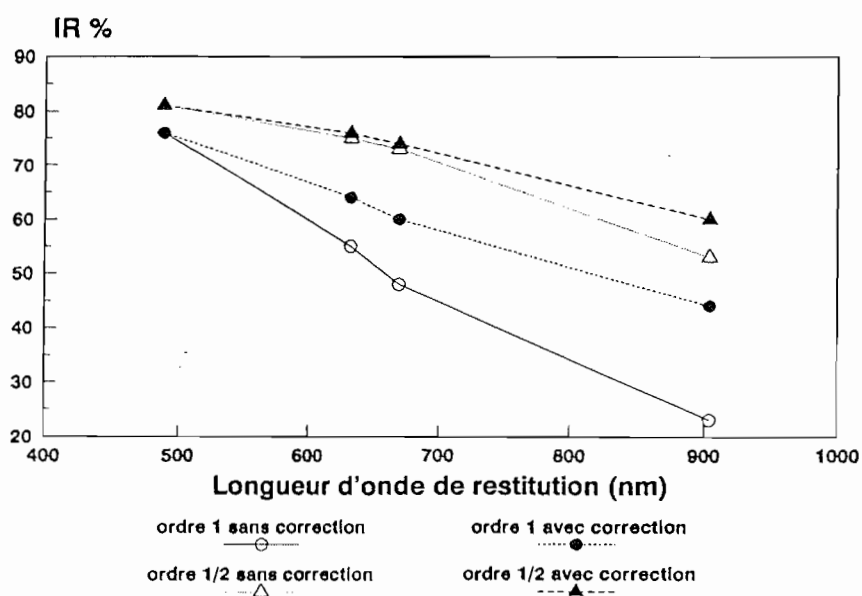
Plan de Talbot d'ordre 1

λ nm	488	632.8		670		904	
θ_2	0.0°	0.0°	3.3°	0.0°	3.9°	0.0°	7.2°
IR	76%	55%	64%	48%	60%	23%	44%

Plan de Talbot d'ordre 1/2

λ nm	488	632.8		670		904	
θ_2	0.0°	0.0°	1.6°	0.0°	1.9°	0.0°	3.3°
IR	81%	75%	76%	73%	74%	53%	60%

Le graphe suivant présente les résultats des simulations et les variations de IR en fonction de la longueur d'onde de restitution :



Nous remarquons la décroissance de IR avec l'écart de la longueur d'onde de restitution par rapport à celle de l'enregistrement. Cette décroissance est moins forte pour la demi distance de Talbot que pour la distance de Talbot. L'effet de la correction est plus net pour les grandes longueurs d'onde et dans le plan de Talbot (IR est pratiquement doublé pour 904 nm dans le plan de Talbot). Dans tous les cas, nous remarquons une nette amélioration de IR avec la correction des aberrations.

V-7 AUTRE POSSIBILITE POUR LA CORRECTION DES ABERRATIONS

Nous avons déterminé le terme d'aberration par une certaine fonction $\tilde{A}(\mu, \nu)$ dans le plan de Fourier de telle manière que la transformée de Fourier de la tache restituée est égale à :

$$\tilde{m}(\mu, \nu) \cdot \tilde{A}(\mu, \nu)$$

si l'objet initial a une transformée de Fourier égale à :

$$\tilde{m}(\mu, \nu) \cdot \tilde{A}^{-1}(\mu, \nu)$$

telle que $\tilde{A}^{-1}$ est la fonction inverse de $\tilde{A}$, à ce moment après la restitution la transformée de Fourier du motif restitué sera $\tilde{m}(\mu, \nu)$ et le motif objet voulu sera correctement restitué. Ceci revient à fabriquer un masque avec des formes de trous donnée par :

$$m(x, y) * \text{TF}^{-1}(\tilde{A}^{-1}(\mu, \nu))$$

au lieu de $m(x, y)$.

Connaissant l'expression de $\tilde{A}$ nous pouvons calculer $\tilde{A}^{-1}$, ensuite calculer sa transformée de Fourier inverse et convoluer le résultat avec la forme de trou voulue.

Un tel masque peut être à des niveaux de gris différents et peut contenir des termes de phase, ce qui rend sa fabrication assez difficile.

V-8 CONCLUSION

A partir d'une modélisation mathématique de l'enregistrement et de la restitution de l'hologramme, nous avons montré que dans l'approximation de Fresnel cette modélisation reconstruit parfaitement l'image du masque. Les aberrations interviennent en développant le calcul au delà de cette approximation (développement des racines au deuxième ordre). En

isolant les termes responsables de ces aberrations, nous avons proposé une méthode simple de corriger ces aberrations en inclinant le masque par rapport au faisceau objet d'un certain angle. Les simulations sur ordinateur ont montré l'efficacité de cette correction en prenant comme critère la concentration de l'énergie lumineuse sur la surface d'une photodiode ($50 \times 50 \mu\text{m}^2$).

TROISIEME PARTIE

REALISATION MICROELECTRONIQUE ET CARACTERISATION

CHAPITRE VI

CIRCUIT ELECTRONIQUE INTEGRE (VLSD)

CHAPITRE VI

CIRCUIT ELECTRONIQUE INTEGRE (VLSI)

Le chapitre III a présenté les différentes fonctions à implanter sur le circuit électronique dans le cadre de la réalisation de l'automate cellulaire "Gaz sur réseau".

Nous avons fabriqué un premier prototype de ce circuit contenant seulement un processeur élémentaire (une cellule), avec ses différentes unités. Le système d'initialisation et de lecture est tout électronique. Le circuit est équipé des photodiodes intégrées pour la réception des instructions optiques.

Dans ce chapitre, nous présentons la conception, la réalisation et le test de ce circuit.

VI-1 LES FONCTIONS DU CIRCUIT

Les fonctions du prototype, que nous avons présentées dans le chapitre III, sont évidemment les mêmes que pour le circuit final. Le prototype a comme but de vérifier l'implantation de ces fonctions électroniquement. Certaines de ces fonctions seront modifiées sur le prototype. Nous allons par la suite discuter ces différentes fonctions.

VI-1-1 Présentation de l'état du gaz

Deux mémoires à 7 bits sont implantées sur le circuit et représentent la configuration des particules :

- "mémoire d'entrée" : qui représente la configuration des particules incidentes. "1" logique dans un bit de mémoire représente la présence d'une particule dans la direction correspondante, et "0" logique représente l'absence de cette particule.
- "mémoire de sortie" : qui représente la configuration des particules résultant de la collision des particules précédentes.

Les différents bits de ces deux mémoires sont distribués sur une maille hexagonale qui reflète les directions de propagation dans le gaz sur réseau.

VI-1-2 L'initialisation

Il s'agit d'affecter la mémoire d'entrée par la configuration des particules incidentes voulue. Dans la version finale, cette initialisation se fait optiquement. Compte-tenu que ce circuit de test ne contient qu'une seule cellule, une initialisation électronique est possible en utilisant un registre à décalage à 7 bascules pour affecter la mémoire d'entrée.

VI-1-3 Le traitement

Nous avons vu au chapitre III que le traitement au niveau de la cellule, par la méthode de substitution symbolique, passe par deux étapes : la reconnaissance et la substitution.

La reconnaissance est en fait une comparaison séquentielle entre le contenu de la "mémoire de correspondance", qui se trouve à l'extérieur du circuit, et celui de la mémoire d'entrée. La communication pour faire cette comparaison se fait par voie optique : le contenu de la mémoire de correspondance est projeté sur les photodiodes du circuit, et nous comparons l'état d'éclairage de ces photodiodes à l'état du contenu de la "mémoire d'entrée". Ainsi, chaque bit de la mémoire d'entrée est lié à une photodiode, et le tout est connecté à une "unité de reconnaissance" qui indique la correspondance ou non des deux états.

La substitution est une copie, conditionnée par l'état de l'"unité de reconnaissance", de l'état des photodiodes dans les bits de la "mémoire de sortie".

Des impulsions de séquençement sont générées électroniquement pour synchroniser le fonctionnement du circuit entre les étapes de reconnaissance et de substitution.

VI-1-4 La propagation

Le circuit est constitué d'une seule cellule, les cellules voisines sont simulées par le registre à décalage. La propagation se fait donc entre cette cellule et les bascules du registre.

VI-1-5 La lecture des résultats

Une lecture électronique des résultats est possible en utilisant les registres à décalage. De cette façon, nous pouvons lire séquentiellement le contenu de la "mémoire de sortie" après le traitement.

Nous envisageons d'utiliser des modulateurs opto-électroniques à base des puits quantiques multiples pour une lecture parallèle des bits de la "mémoire de sortie". Compte-tenu des caractéristiques spéciales de ces modulateurs, les plots de connexion seront fabriqués spécialement, et sont appelés : "sortie modulateurs".

VI-2 CONCEPTION DU CIRCUIT ELECTRONIQUE

Nous distinguons dans le circuit entre le cœur qui représente la cellule de traitement, et le périphérique constitué du registre à décalage et les plots de "sortie modulateurs".

VI-2-1 Conception de la cellule : coeur du circuit

La cellule est constituée de :

- 7 parties identiques que nous appelons "unité particule" (car chacune de ces parties correspond à une particule dans le modèle du "gaz sur réseau").
- Et d'une "unité de reconnaissance" commune à toutes les parties.

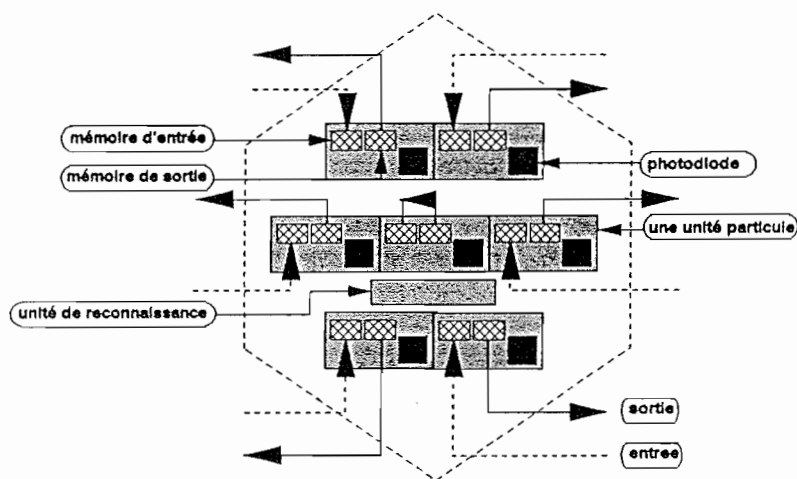


Fig VI-1 : Schéma de la cellule.

VI-2-1-a Unité de reconnaissance

L'unité de reconnaissance décide si la configuration projetée sur les photodiodes est exactement la même que celle contenue dans la "mémoire d'entrée" et donne, ou non, l'ordre d'enregistrer la configuration substituante dans la "mémoire de sortie".

Cette unité est constituée d'un point mémoire dont l'entrée est directement liée à une capacité C (voir fig. VI-6).

Le mécanisme de reconnaissance consiste à charger au préalable la capacité C (avec le transistor T0 et le signal Φ_0) ce qui impose le niveau logique "1" à l'entrée du point mémoire. Cette capacité est connectée aux photodiodes par un système de transistors commandé par le contenu de la "mémoire d'entrée" de telle manière que si l'éclairage des photodiodes ne correspond pas au contenu de cette mémoire la capacité C sera déchargée à travers les photodiodes et le niveau logique à l'entrée du point mémoire tombe à "0". Dans le cas contraire la capacité C gardera sa charge initiale. L'étape de reconnaissance finit par

l'enregistrement dans le point mémoire (avec le signal Φ_3) le niveau logique imposée par cette capacité.

VI-2-1-b Unité particule

L'"unité particule" est constituée de deux points mémoire (l'un est un bit de la "mémoire d'entrée", l'autre un bit de la "mémoire de sortie"), d'une photodiode et d'un système de transistors pour communiquer avec l'"unité de reconnaissance".

i- Point mémoire

Un point mémoire est réalisé en technologie VLSI par deux inverseurs [44], et un système de deux transistors complémentaires commandés avec un signal Φ . Ce signal permet l'enregistrement la donnée comme le montre la figure suivante :

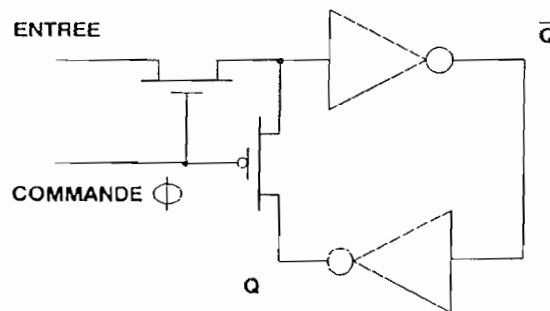


Fig VI-2 : schéma d'un point mémoire.

ii- Photodiode

Une photodiode est une diode polarisée en inverse. Son schéma équivalent est constitué d'une capacité interne, et une source de courant qui décharge la capacité sous l'effet d'un éclairage optique (Fig. VI-3) (une description détaillée de son fonctionnement est donnée dans le chapitre VII). La capacité interne de la photodiode est préalablement chargée à travers le transistors T4 avec le signal Φ_4 (voir fig. VI-6). Cette photodiode est connectée aux deux points mémoires et à l'unité de reconnaissance comme c'est expliqué par la suite.

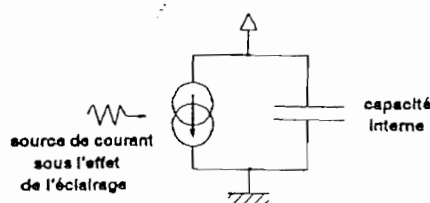


Fig. VI-3 : Schéma d'une photodiode.

iii- Transistors de connexion

Nous distinguons un système pour la reconnaissance, et un autre pour la substitution.

La reconnaissance : Après une précharge de la capacité C de l'unité de reconnaissance, ce système a comme rôle de décharger cette capacité C à travers les photodiodes si l'éclairage de l'ensemble de celles-ci ne correspond pas au contenu de la "mémoire d'entrée", c'est à dire si la reconnaissance a échoué. Pour cela, nous connectons la capacité C aux photodiodes au moment d'éclairage de celles-ci.

Cette connexion peut se faire avec un simple transistor comme indiqué dans la figure suivante :

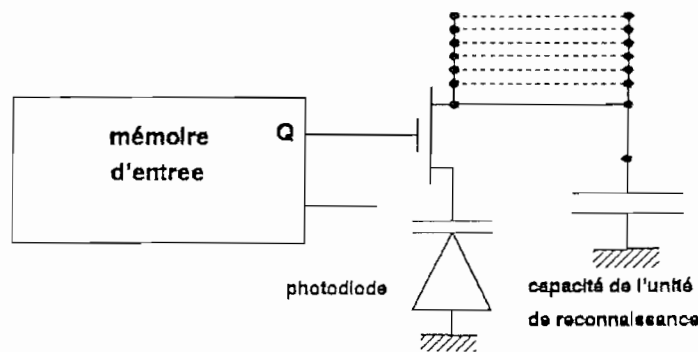


Fig VI-4 : connexion simple pour la reconnaissance.

Dans ce cas, la capacité C serait connectée aux photodiodes dont le bit de la "mémoire d'entrée" contiendrait "1" logique. Si l'une de ces photodiodes est éclairée la capacité C serait déchargée à travers cette photodiode, et la reconnaissance n'aurait pas lieu. Donc, il faut éclairer les photodiodes par la configuration complémentaire à la configuration à reconnaître : l'éclairage d'une photodiode reconnaît un "0" dans le bit correspondant de la mémoire d'entrée.

Dans ce cas, il y a risque de reconnaître une fausse configuration comme dans l'exemple suivant :

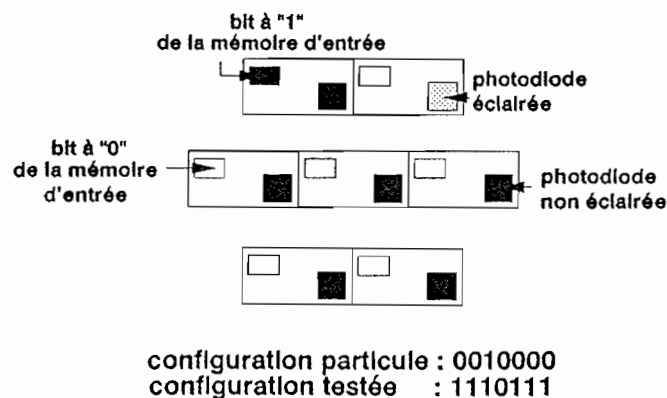


Fig. VI-5

Par conséquent, il est nécessaire de faire le test aussi bien sur la configuration présente dans la "mémoire d'entrée" que sur la configuration complémentaire. C'est à dire d'utiliser les sorties Q et $\bar{Q}$ des points mémoire d'entrée, et le système de connexion pour la reconnaissance aura la forme de la figure VI-6 :

Sous l'action du signal Φ_1 le transistor T1 est fermé, si "1" bit est à un logique, le transistor T3 est fermé et la capacité C est connectée à la photodiode correspondante. Nous éclairons les photodiodes par la configuration complémentaire à celle à reconnaître (diode éclairée reconnaît un "0" dans le bit correspondant de la mémoire d'entrée). Si la configuration à reconnaître correspond au contenu de la mémoire d'entrée, la capacité C n'est pas déchargée. Nous recommençons avec les configurations complémentaires et avec le signal Φ_2 et le transistor T2. A la suite de ces deux étapes nous finissons la reconnaissance par l'enregistrement de l'état de la capacité C dans le point mémoire de l'"unité de reconnaissance" avec le signal Φ_3 .

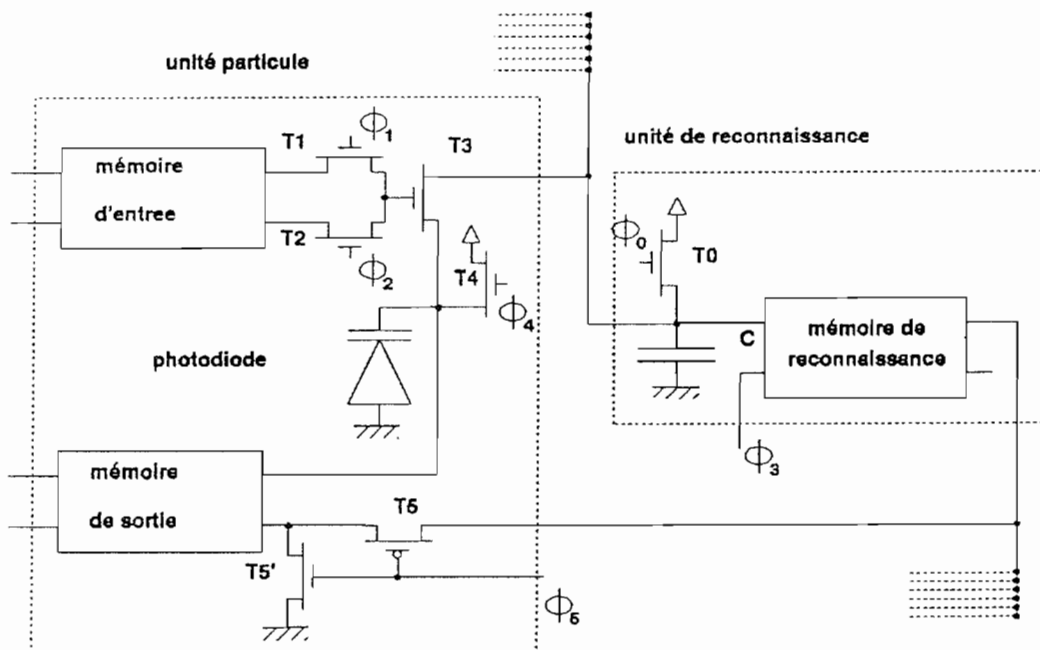


Fig VI-6 : Schéma électrique de la cellule.

Substitution : Ce système a comme rôle de permettre l'enregistrement dans la mémoire de sortie de l'état d'éclairage des photodiodes à condition que l'unité de reconnaissance soit au niveau logique "1". Ce système est constitué de deux transistors T5 et T5' de deux natures différentes : l'une N l'autre P (voir fig VI-6). L'enregistrement dans la mémoire de sortie est déclenché par le signal Φ_5 . Quand ce signal est au niveau haut le transistor T5' est

-la sortie de chaque bit de la "mémoire de sortie" est connectée à une entrée du premier bascules du registre correspondant. En fait, le premier point bascule de chaque registre a deux entrées : celui du registre précédent, et celui du bit de la "mémoire de sortie".

Le transfert entre le registre à décalage et les mémoire de la cellule se fait par un seul signal A.

VI-2-2-b Sortie modulateurs

Le circuit sera connecté à un circuit de modulateurs opto-électroniques à puits quantiques multiples fabriqués en technologie AsGa pour la lecture optique du résultat du traitement.

Ces modulateurs fonctionnent entre des tensions variant entre 0 et 10 V (voire 15 V) sous un courant de fuite d'environ 100 μ A .

Le circuit de test est fabriqué sur une puce de Silicium et travaille avec des tensions 0 et 5 V (7 V maxi) avec un courant de quelques micro-ampères, ce qui est insuffisant pour les modulateurs opto-électroniques. Donc, une connexion directe entre les deux circuits est impossible, il faut un adaptateur de tension pouvant débiter des courants suffisants pour le circuit des modulateurs.

Compte-tenu de cette inadéquation de tension de fonctionnement entre les deux circuits, nous avons été menés à étudier des plots spécialisés pour les sorties modulateurs AsGa en attendant que ceux-ci fonctionnent sous 5 V, et qu'ils soient moins gourmands en courant !.

Ces adaptateurs peuvent être réalisés directement sur le circuit Silicium comme le montre la figure VI-8.

Le système des transistors T1 et T2 (Fig. VI-8) sert à élever transitoirement la tension du niveau logique "1" de 5 V à 10V (voire plus en cas de nécessité), avec une inversion de logique.

Les transistors T3 et T1 sont montés en miroir de courant. Nous contrôlons ainsi, par la tension V imposée sur la plaque "contrôle du courant", le courant dans T3 et donc dans T1 (courant qui contrôle le plot de sortie) (voir annexe B).

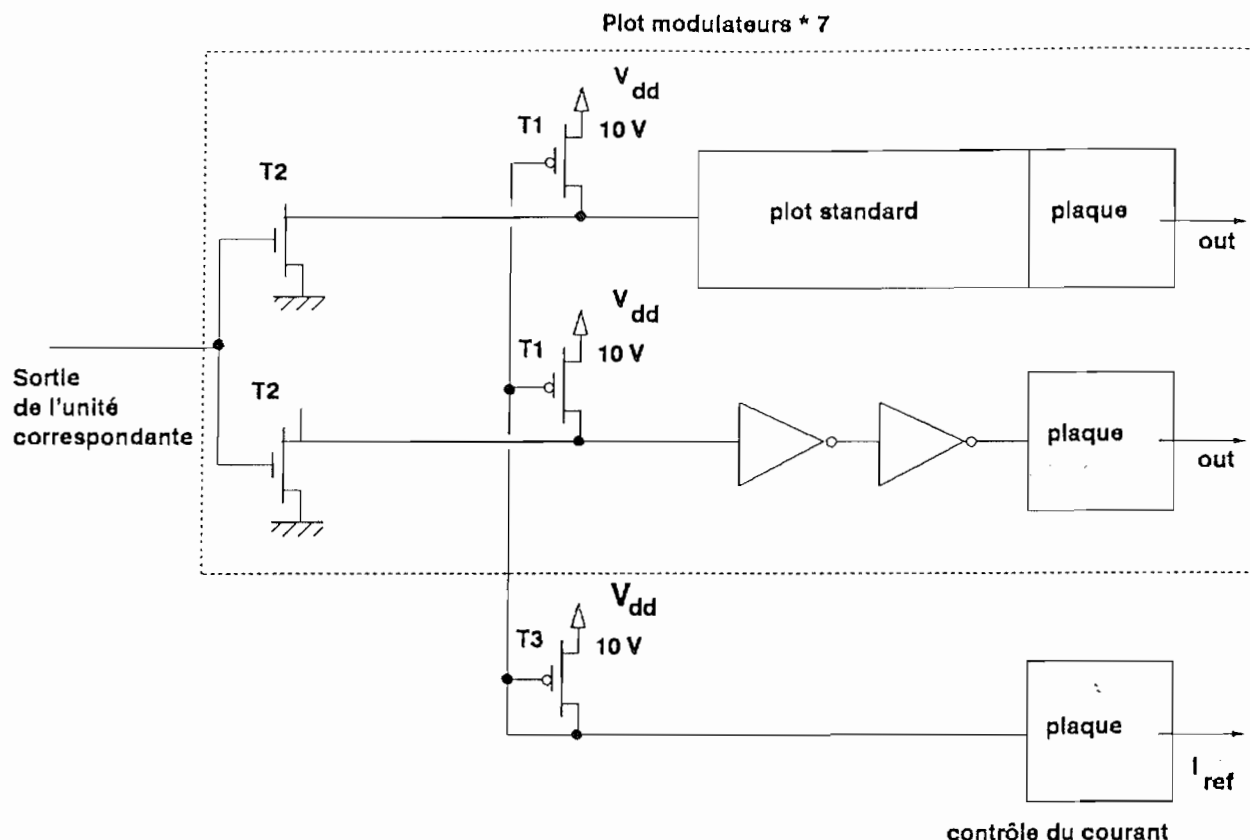


Fig. VI-8 : Plots de sortie modulateurs.

VI-3 DESSIN ET REALISATION DU CIRCUIT

Le circuit est réalisé sur une puce de Silicium de dopage P en circuit intégré VLSI technologie 2 μm .

La conception et les dessins des plans ont été faits avec le programme CAO MAGIC sur Apollo.

Le fonctionnement des différentes parties du circuit a été simulé sur le programme Spice. Malgré les différents essais, le programme SPICE n'a pas accepté de simuler le fonctionnement du circuit global à cause de ses dimensions. Le nombre de transistors du circuit est de :

133	pour le circuit coeur
63	pour les plots modulateurs
49	pour le registre à décalage
245	au total

Les différents plans de conception sont présentés par la suite.

VI-3-1 Point mémoire

Comme nous avons vu, le point mémoire est constitué de deux inverseurs et deux transistors pour enregistrer l'information. Chaque inverseur est réalisé conformément au schéma suivant dans la technologie VLSI [44] :

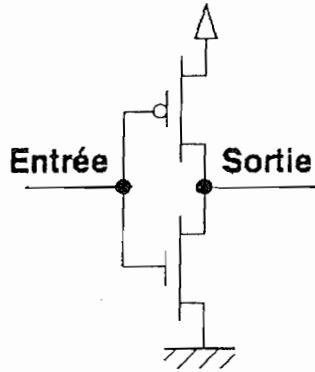


Fig VI-9 : Schéma d'un inverseur.

VI-3-2 Unité particule

Cette unité contient le point mémoire entrée, le point mémoire sortie, le système de connexion et la photodiode. La photodiode est réalisée par implantation d'une couche dopée N^+ sur le substrat dopé P . Les dimensions de cette photodiode sont de $50 * 50 \mu m^2$. Pour éviter le transport de charge entre la photodiode et le reste du circuit, celle-ci est entourée d'un anneau de garde (dopée P^{++}) et portée à la masse.

VI-3-3 La cellule

La cellule contient 7 unités particule et l'unité de reconnaissance. La capacité C de l'unité de reconnaissance est en fait la capacité interne d'une diode polarisée en inverse (donc de même nature que les photodiodes). Cette capacité est estimée à 56 fF. Une couche métallique couvre cette capacité pour la protéger contre la lumière.

Des plots de test sont connectés aux différentes parties du circuit : Les sorties des points mémoires, l'entrée et la sortie du point mémoire de l'unité de reconnaissance. Ces plots sont faits de plaques métalliques pouvant recevoir des pointes de test.

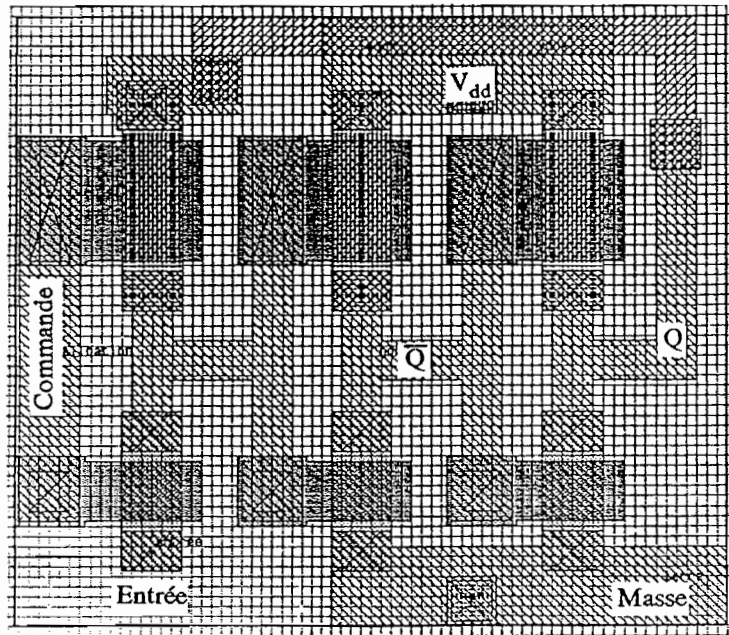
VI-3-4 Le registre à décalage

Chaque registre est constitué de deux bascules (Fig.VI-10 b), le premier a deux entrées informations (registre et mémoire de sortie) et leurs commandes (A et $R1$ respectivement), le deuxième a une simple entrée (registre précédent) avec sa commande ($R2$).

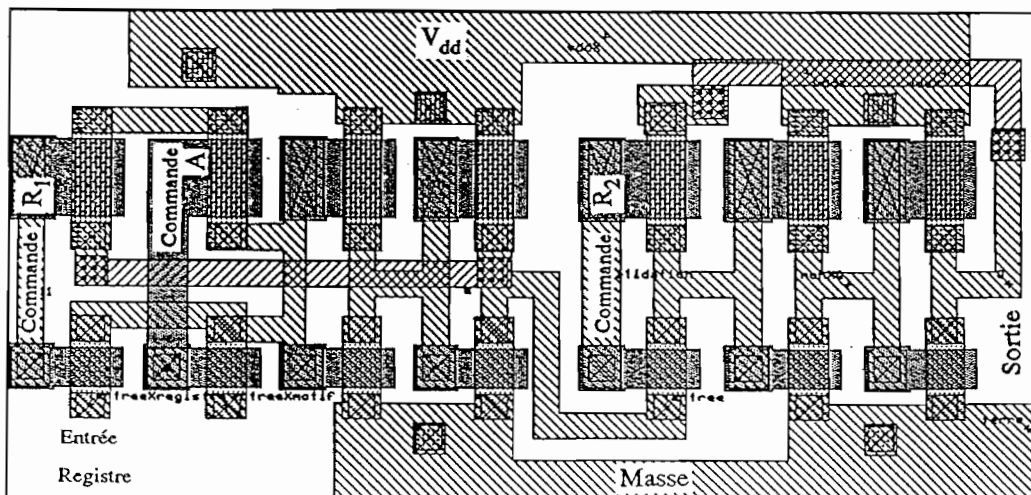
VI-3-5 Le plot sortie modulateur

Nous avons implanté sur le circuit deux sortes de plots de sortie modulateurs : Les plots standards et les plots fabriqués (Fig.VI-10 e). Les plots standards sont bien protégés mais sont conçus pour travailler à 7 V maximum. Pour les plots fabriqués, les dimensions des transistors sont faites de telle manière qu'ils supportent des courants allant jusqu'à 1 mA et des tensions jusqu'à 10 V.

Les figures VI-10 montrent les dessins de conception de ce circuit.

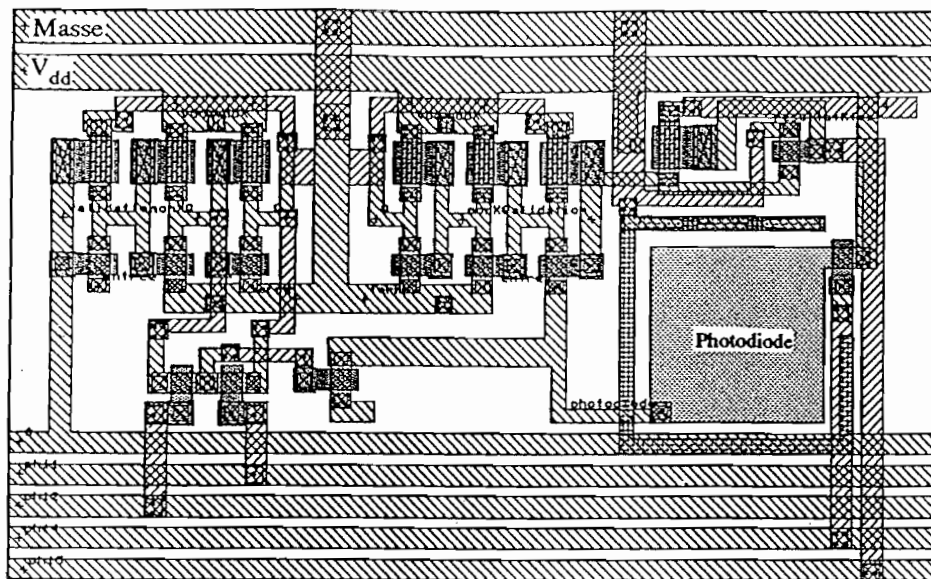


(a) Point mémoire.

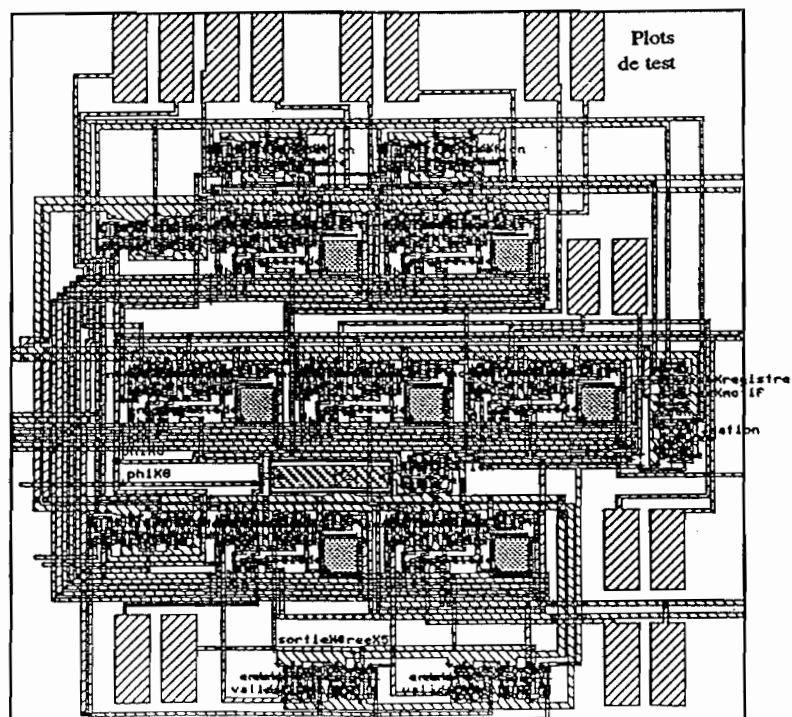


(b) Registre à décalage

Fig. VI-10

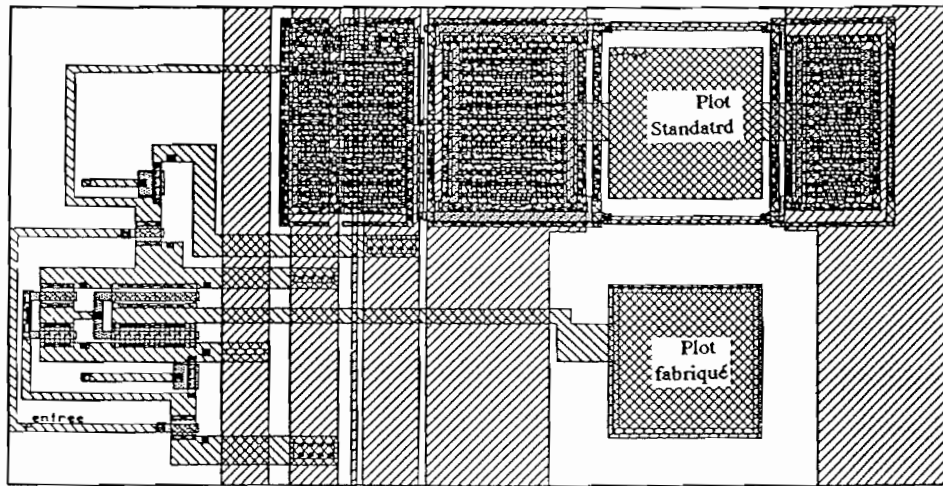


(c) Unité particule.

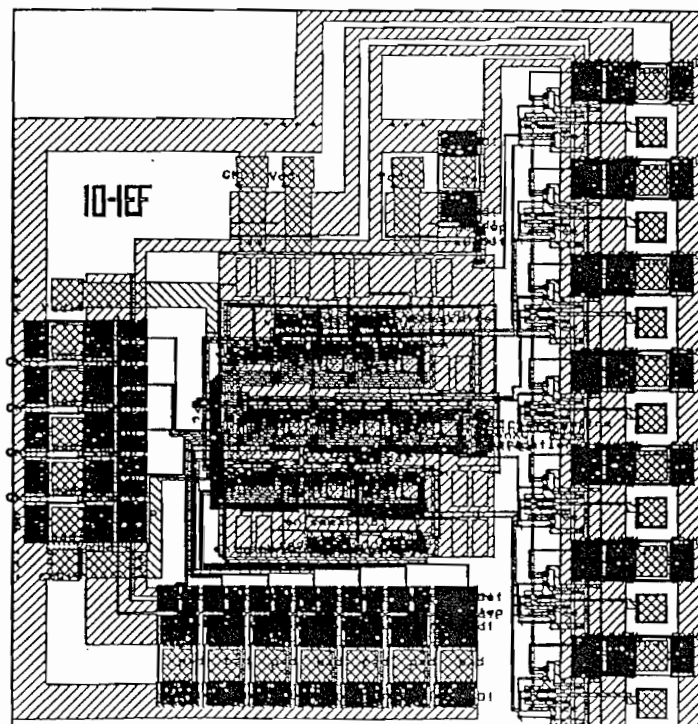


(d) Cœur du circuit.

Fig. VI-10



(e) Sortie modulateurs.



(f) Circuit global.

Fig. VI-10

VI-4 TEST DU CIRCUIT

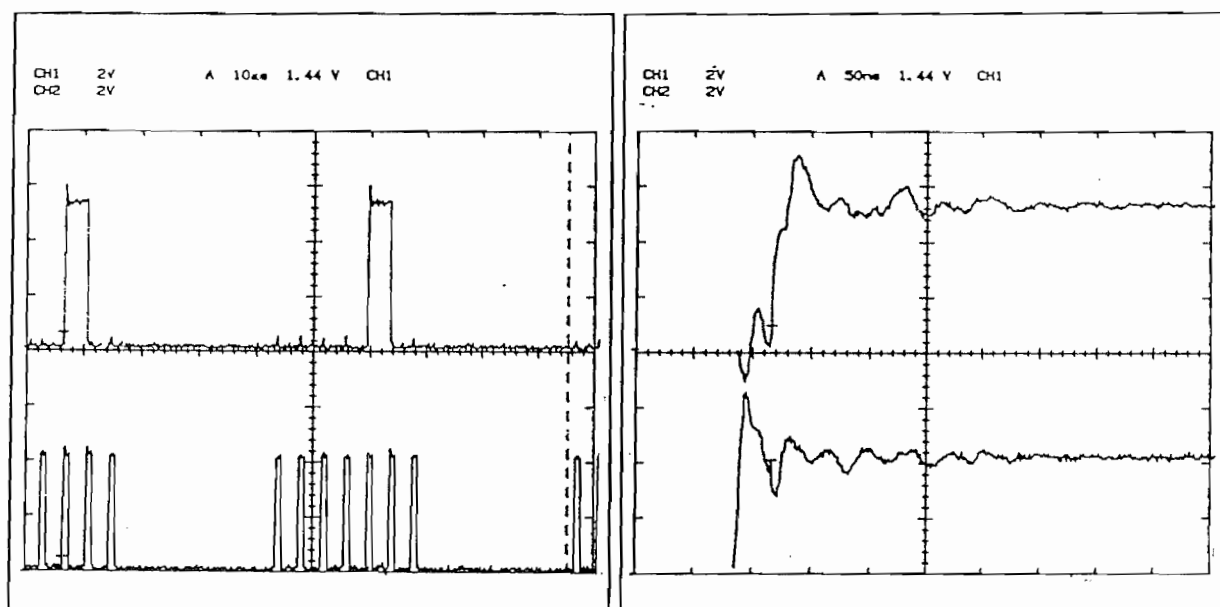
Nous avons mené sur le circuit plusieurs tests, dont nous présentons les résultats par la suite.

VI-4-1 Test du registre à décalage

C'est le premier test du circuit. Il s'agit d'introduire un mot et de le retrouver à la sortie.

le registre a été vérifié avec différents mots. La sortie est le complément du mot d'entrée, ceci est dû aux plots de sortie qui complémentent l'information.

Nous avons observé le temps de réponse du circuit en comparant la sortie (à un bit au niveau haut) avec l'impulsion R2, voir fig. VI-11, et nous estimons un temps de réponse pour les mémoires d'environ 50 ns (Ce temps de transition est très approximatif car les conditions de test n'étaient pas optimisées en ce qui concerne les impédances des plots de test et le temps de transition des impulsions électroniques envoyées.



(a) temps de base 10 ms

(b) temps de base 50 ns

Fig. VI-11 : Test du registre à décalage.

VI-4-2 Test du circuit principal

VI-4-2-a Mémoires d'entrée

Le but de ce test est de vérifier le chargement des mémoires d'entrée des différentes unités à partir du registre à décalage. Compte-tenu que les sorties des mémoires d'entrée sont liées à des plots de test, nous avons pointé chaque plot par une pointe pour relever la tension de la mémoire, et nous avons introduit un mot dans le registre à décalage, puis copié dans la mémoire d'entrée par le signal "A". L'éventuel changement d'état d'un point mémoire est relevé par la pointe.

Nous avons, ainsi, vérifié les mémoires d'entrée 1-6.

VI-4-2-b Unité de reconnaissance

Nous avons vérifié le fonctionnement du point mémoire. Nous chargeons la capacité C avec le signal (Φ_0) actionné pendant 10 μ s. Nous envoyons ensuite le signal d'enregistrement au point mémoire et nous vérifions la transition du point mémoire au niveau logique "1" avec une pointe sur le plot de test lié à la sortie de ce point mémoire. Ensuite, nous laissons la capacité se décharger par le courant de fuite, et nous vérifions la transition à "0" de la sortie de la "mémoire de reconnaissance" en actionnant le signal (Φ_3).

Pour caractériser le temps de chargement et de déchargement de la capacité, nous avons commencé par l'utilisation de la pointe sur les plots de test. Comme cette pointe introduit des courants de fuite importants, nous avons obtenu des faux résultats. Nous avons estimé les caractéristiques de ce circuit par le test global du circuit.

VI-4-2-c Test des mémoires de sortie

Les sorties de la mémoire de sortie sont liées à des plots de test et au registre à décalage. Le but du test est de vérifier l'écriture dans les mémoires de sortie à partir de l'éclairage des photodiodes. Pour cela, nous forçons l'unité de reconnaissance à "1" logique en chargeant sa capacité et déclenchant tout de suite l'enregistrement dans le point mémoire de reconnaissance par Φ_3 . Ainsi, l'écriture dans les mémoires de sortie sera permise. Pour l'écriture, nous chargeons les capacités internes des photodiodes, ce qui impose des "1" logique à l'entrée des mémoires de sortie. Par éclairage d'une photodiode, la tension à l'entrée du point mémoire correspondant tombe à 0 V. Nous enregistrons ensuite, par le signal (Φ_5), et nous lisons la mémoire de sortie par le registre à décalage.

VI-4-2-d Test de fonctionnement global du circuit

Compte-tenu des problèmes imposés par les pointes et les sondes de l'oscilloscope, nous avons procédé à la simulation du fonctionnement global du circuit pour le caractériser. Nous n'éclairons éventuellement qu'une seule photodiode par le faisceau d'un laser He-Ne focalisé avec un objectif de microscope.

a) Description de la méthode du test :

Nous décrivons la méthode du test dans un cas simple où la reconnaissance se fait d'une façon certaine. Nous mettons le circuit dans l'obscurité et nous maintenons les signaux Φ_1 et Φ_2 au niveau logique "0", à ce moment la capacité C de l'unité de reconnaissance est isolée des photodiodes donc ne doit se décharger que par sa propre courant de fuite qui est très faible. Nous préchargeons cette capacité C et au bout d'un certain temps nous déclenchons l'enregistrement dans le point mémoire de l'unité de reconnaissance par le signal Φ_3 . L'unité de reconnaissance présente "1" logique, signe de reconnaissance, si nous n'avons pas attendu trop longtemps avant le signal Φ_3 . Pour saisir la réponse de reconnaissance, nous préchargeons les capacités internes des photodiodes par le signal Φ_4 et nous éclairons une photodiode qui est la seule à présenter une tension nulle à l'entrée du bit correspondant de la mémoire de sortie. Nous déclenchons ensuite l'enregistrement dans la mémoire de sortie avec le signal Φ_5 et nous lisons son contenu par le registre à décalage. Si la reconnaissance est faite nous retrouvons l'image d'éclairage de la photodiode à la sortie de la mémoire de sortie, et en cas de non reconnaissance nous retrouvons l'ancien contenu de la mémoire de sortie. C'est ainsi que nous avons vérifié que le courant de fuite de la capacité C est très faible, ainsi que la procédure d'écriture dans la mémoire de sortie.

b) Test de décharge de la capacité C à travers les photodiodes en obscurité :

Nous forçons la mémoire de sortie à "111111" par une procédure identique à celle décrite précédemment sans aucun éclairage des photodiodes. Nous chargeons ensuite "111111" dans la mémoire d'entrée. Nous maintenons le signal Φ_1 au niveau logique "0", tandis que le signal Φ_2 est actionné au niveau logique "1" pendant un certain temps τ_2 caractéristique du temps de décharge de la capacité C (si elle a lieu). Dans ces conditions, la capacité C ne doit pas être connectée aux photodiodes, et la reconnaissance doit avoir lieu à n'importe quel temps τ_2 .

Après le temps τ_2 nous déclenchons l'enregistrement dans la mémoire de reconnaissance avec Φ_3 , et nous saisissons la réponse de reconnaissance comme c'est fait précédemment en éclairant une photodiode et en enregistrant ensuite dans la mémoire de sortie. La reconnaissance a été saisie pour des temps τ_2 jusqu'à 300 μ s (limitation de dispositif utilisé).

Cette expérience montre encore une fois que la capacité C se décharge très lentement par son courant de fuite.

Nous avons recommencer la même expérience mais en maintenant le signal Φ_2 au niveau "0" et en actionnant le signal Φ_1 au niveau "1" pendant un temps τ_1 . Dans ce cas la capacité de reconnaissance C est connectée aux photodiodes, toujours dans l'obscurité, pendant le temps τ_1 . En principe la reconnaissance doit être saisie comme précédemment car les photodiodes sont dans l'obscurité.

L'expérience a montré qu'il suffit moins que $\tau_1 = 1 \mu s$ pour ne pas pouvoir saisir la reconnaissance. Deux explications peuvent justifier ce résultat :

- Ou bien le courant de fuite des photodiodes est très important. Une explication peu probable car les photodiodes ont la même structure que la capacité C .
- Ou bien la charge de la capacité C se distribue sur les capacités internes des photodiodes (phénomène de division de tension) pour que la tension aux bornes de la capacité C chute presque instantanément en dessous de la tension du seuil.

Pour trancher entre ces deux explications, nous avons précharger les capacités internes des photodiodes par le signal Φ_4 avant d'actionner le signal Φ_1 . Ainsi le phénomène de division de tension n'aurait pas lieu, tandis que si les courants de fuite des photodiodes étaient importants, ils auraient déchargé leurs capacités internes. En recommençant ainsi l'expérience, nous saisis la reconnaissance pour des temps τ_1 allant jusqu'à $300 \mu s$ comme c'était attendu, ce qui confirme le phénomène de division de tension.

c) Influence de la division de tension sur le fonctionnement de l'algorithme :

Selon le déroulement de l'algorithme, la reconnaissance se fait en deux étapes (voir § VI-2-1-b) : avec la configuration à reconnaître (projetée simultanément avec Φ_1) et la configuration complémentaire (projetée avec Φ_2). Pour éviter une fausse non reconnaissance par division de tension, il faut précharger les capacités internes des photodiodes avec le signal Φ_4 .

A ce moment, un autre problème se pose par la division de tension dans le sens inverse, c'est à dire charger la capacité C à partir des photodiodes, ce qui constitue une nouvelle source d'erreur dans le cas où il y a une décharge de la capacité " C " en première étape et pas en deuxième, et nous aurons ainsi une fausse connaissance car si la capacité C se décharge en première étape c'est que la reconnaissance ne doit pas avoir lieu. Nous avons vérifié ce phénomène, et nous avons constaté que la charge des photodiodes suffit pour élever la tension de C au niveau logique "1".

Pour résoudre ce problème, il suffit de remarquer que la décharge de C en première étape et pas en deuxième se présente lorsque le motif à reconnaître est inclu dans le motif existant.

A ce moment, il faut défiler les motifs à reconnaître en commençant par le plus simple (où il n'y a pas de particules présentes), et en finissant par le plus compliqué (où les sept particules sont présentes). Il peut y avoir plusieurs reconnaissances dans un noeud, mais c'est la dernière qui sera la bonne, et c'est elle qui déterminera la configuration de substitution finale.

d) Test de la décharge de la capacité C à travers des photodiodes éclairées

Nous avons recommencer la même expérience que b) avec préchargement des capacités des photodiodes mais en éclairant une photodiodes d'une façon continue avec la lumière d'un laser HeNe de puissance $0.1 \mu\text{W}$. Avec le signal Φ_1 maintenu au niveau logique "0" et le signal Φ_2 à "1" pendant τ_3 , la reconnaissance devrait avoir lieu car la capacité C est déconnectée des photodiodes. Nous avons saisi la reconnaissance jusqu'à $\tau_3 = 120 \mu\text{s}$.

En inversant les fonctions des signaux Φ_1 et Φ_2 , la reconnaissance ne doit pas avoir lieu car la capacité C est connectée à la photodiodes éclairée. En fait, nous avons obtenu la non reconnaissance en moins de $1 \mu\text{s}$.

Nous avons recommencé l'expérience précédente mais en chargeant un "0" dans le bit de la mémoire d'entrée correspondant à la photodiode éclairée. A ce moment nous devons avoir la reconnaissance en actionnant le signal Φ_1 et la non reconnaissance avec le signal Φ_2 (le résultat contraire à la précédente). Nous avons saisi la reconnaissance avec Φ_1 jusqu'à 280 ms , et la non reconnaissance avec Φ_2 en moins de $1 \mu\text{s}$.

e) Résultats du test

En prenant précaution de précharger les capacités internes des photodiodes avant d'éclairer les photodiodes et d'actionner les signaux Φ_1 et Φ_2 , le circuit répond bien aux fonctions demandées à condition que l'éclairage des photodiodes pendant la période de reconnaissance ne dure pas plus que $120 \mu\text{s}$.

VI-4-3 Test des plots de sortie AsGa

Nous avons branché le plot de contrôle du courant conformément au schéma de la figure suivante, et nous avons vérifié la transition entre 0 et 10 V pour les deux sortes des plots (standards et fabriqués), à partir du chargement de la mémoire de sortie. Pour le fonctionnement statique des plots, nous avons utilisé comme charge une résistance variable, cette résistance nous permet de déterminer le courant maximal débité par le plot pour une tension de fonctionnement donnée. Le courant maximal débité sous 10 V est de $100 \mu\text{A}$ pour les deux sortes des plots.

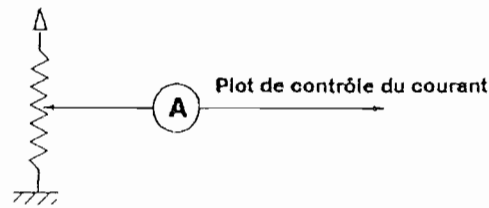


Fig VI-12: Branchement du plot de réglage.

Pour le fonctionnement dynamique, nous avons simulé le comportement des puits quantiques par une capacité de $C = 10 \text{ pF}$ en parallèle avec une résistance de $R_1 = 15 \text{ M}\Omega$. Une résistance de $R_2 = 2.2 \text{ k}\Omega$ est branchée en série pour relever les courants débités pendant les transitions.

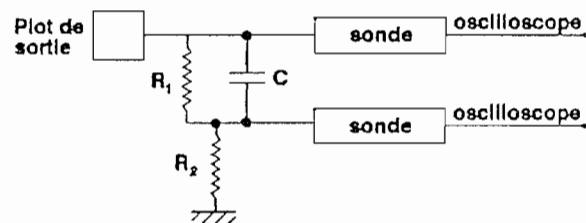


Fig VI-13 : Etude dynamique des plots de sortie.

Nous avons relevé les courbes des figures VI-14 et VI-15. Les courbes de la figure VI-15 sont en fonction des tensions de réglage variables. Nous constatons que plus la tension de réglage est élevée plus le temps de remontée jusqu'à 10 V est long. Ce temps varie entre 150 ns pour une tension de réglage de 5 V, jusqu'à 200 ns pour une tension de réglage de 8.9 V, et au delà de cette valeur, la transition ne se manifeste pas. C'est dû essentiellement au blocage du transistor "T1" (voir fig VI-8).

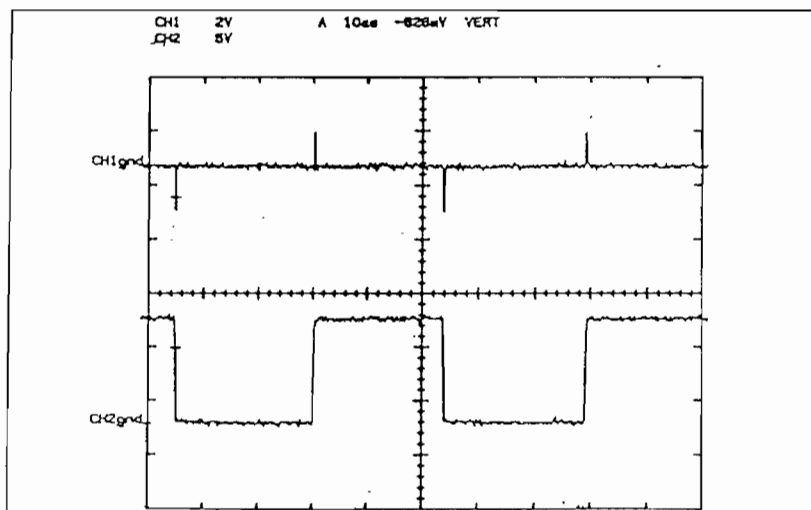
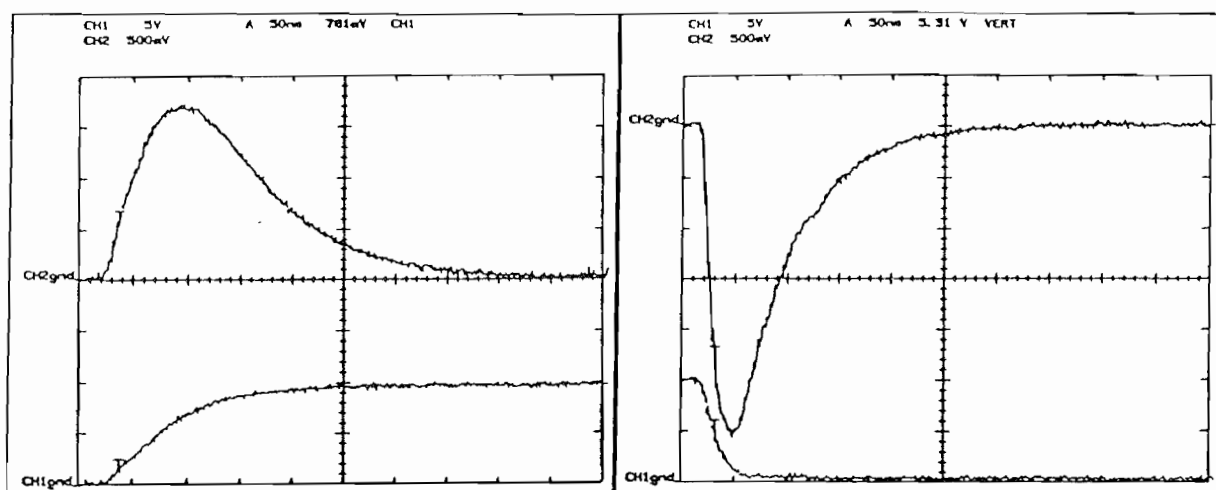
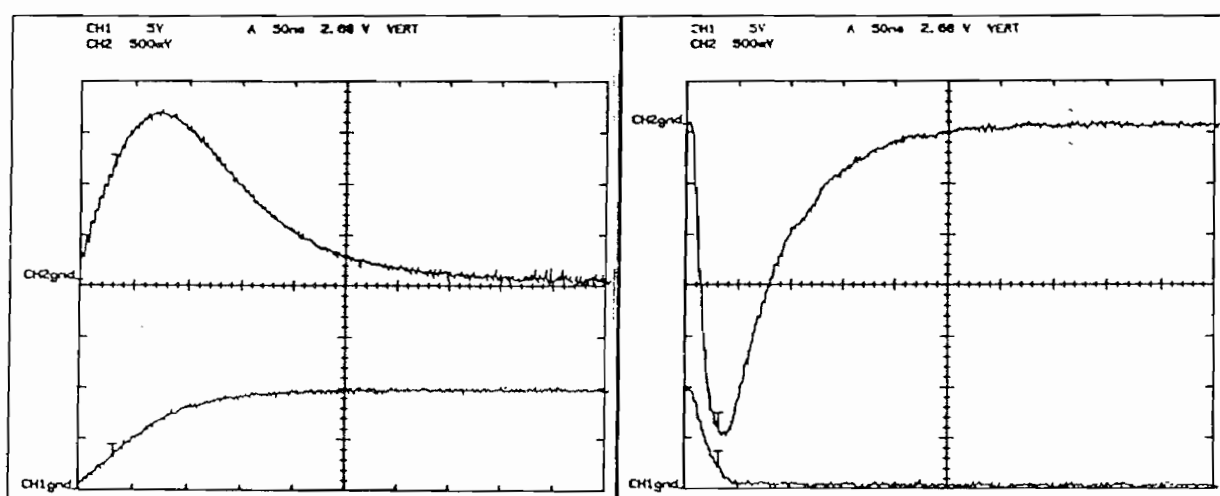


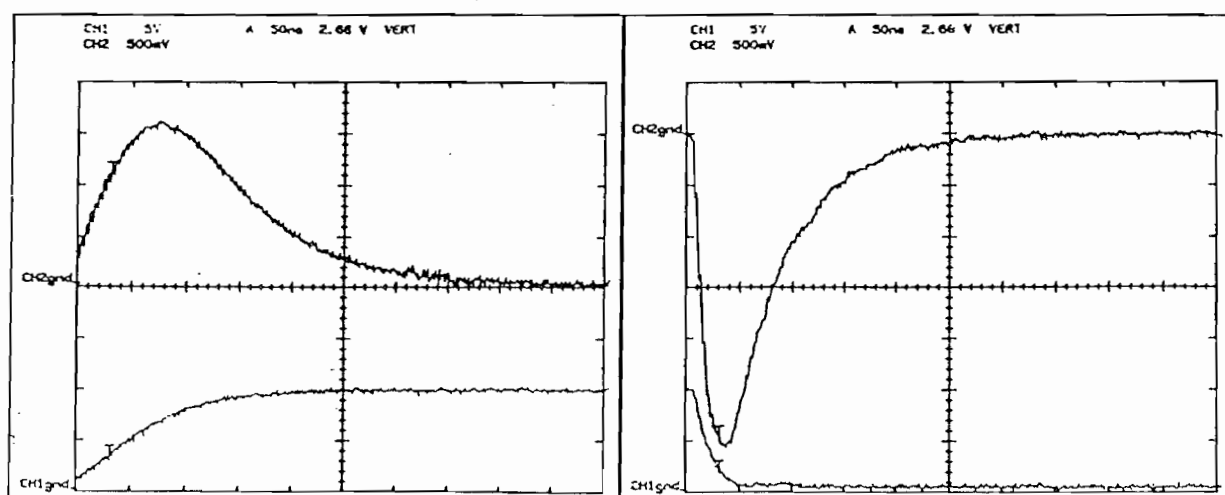
Fig. VI-14 : Tension et courant des plots de sortie modulateurs.



(a) Niveau logique : 10 V. Tension de réglage : 5.03 V

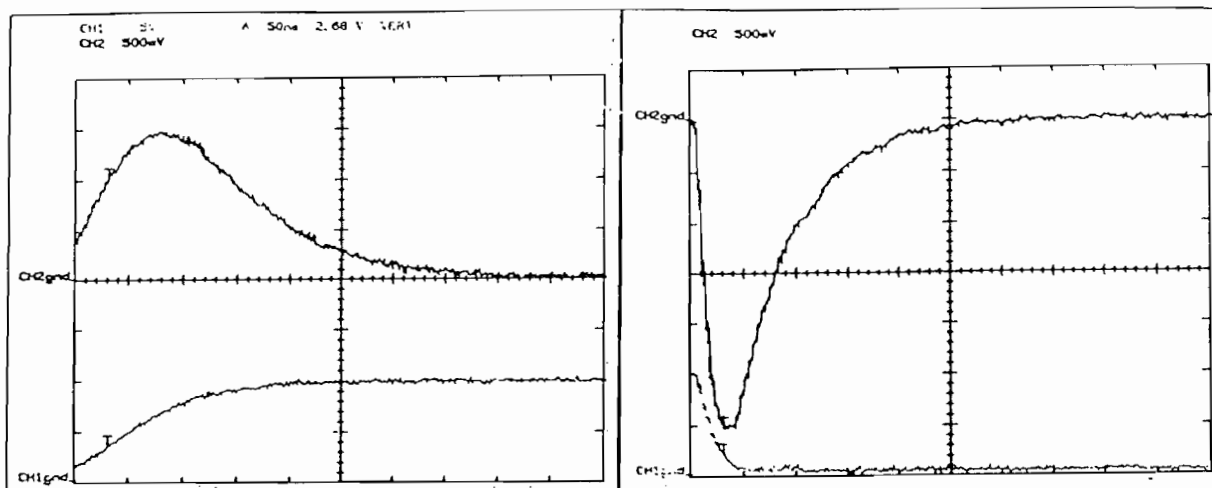


(b) Niveau logique : 10 V. Tension de réglage : 6.68 V

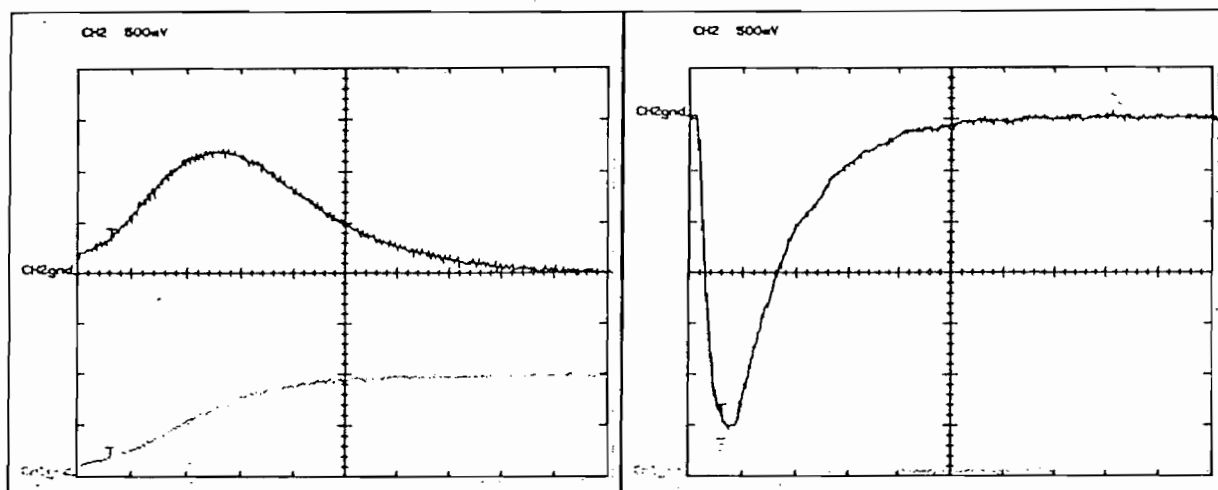


(c) Niveau logique : 10 V. Tension de réglage : 7.5 V

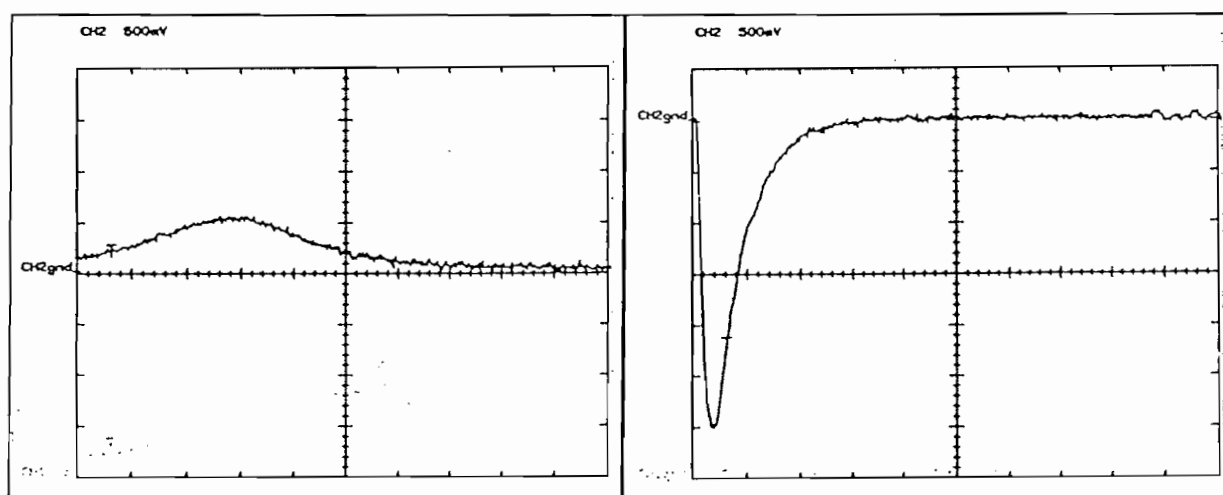
Fig. VI-15 : Tension et courant des plots de sortie modulateurs.



(d) Niveau logique : 10 V. Tension de réglage : 8.02 V

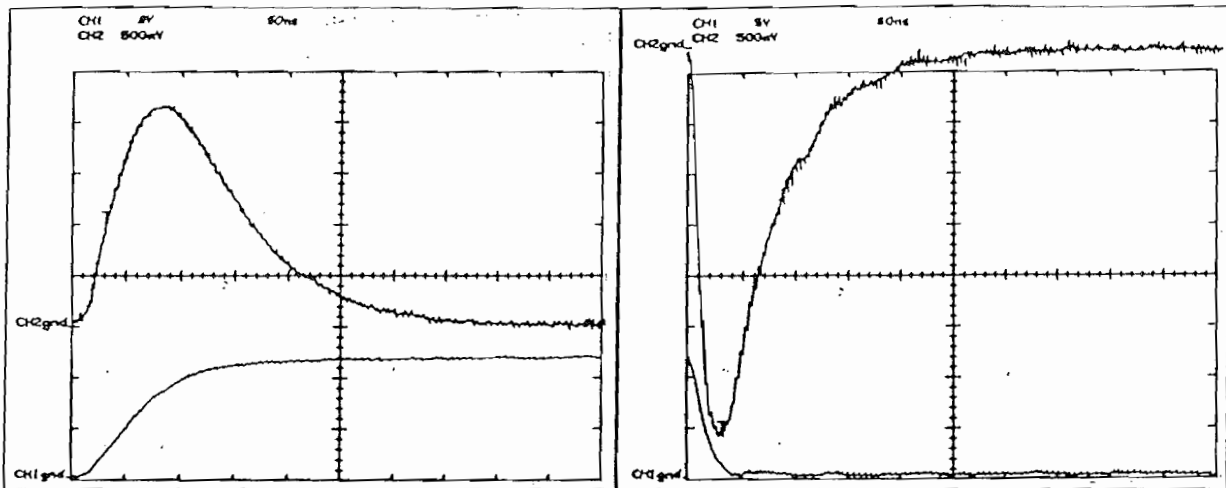


(e) Niveau logique : 10 V. Tension de réglage : 8.5 V

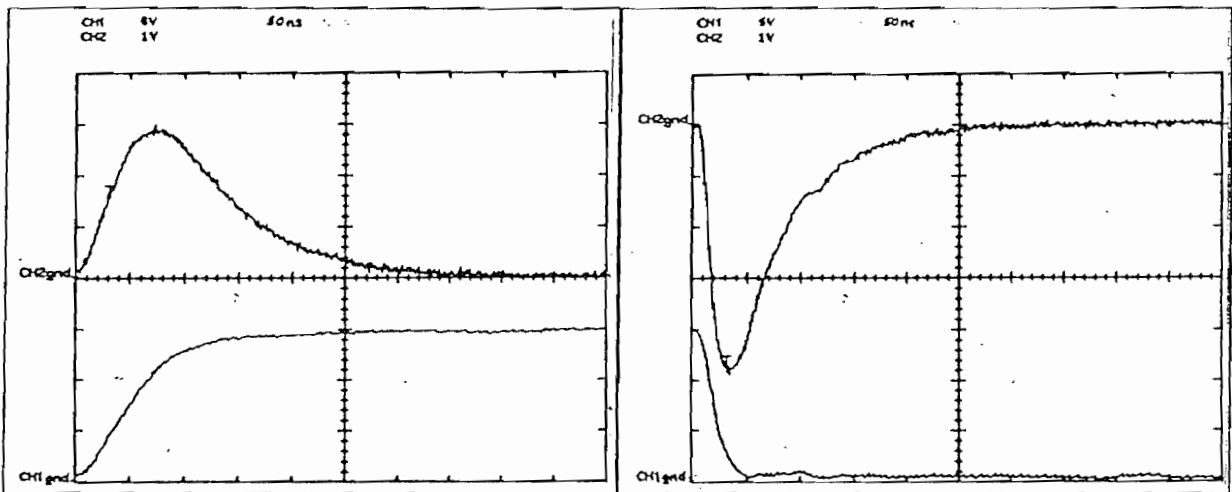


(f) Niveau logique : 10 V. Tension de réglage : 8.9 V

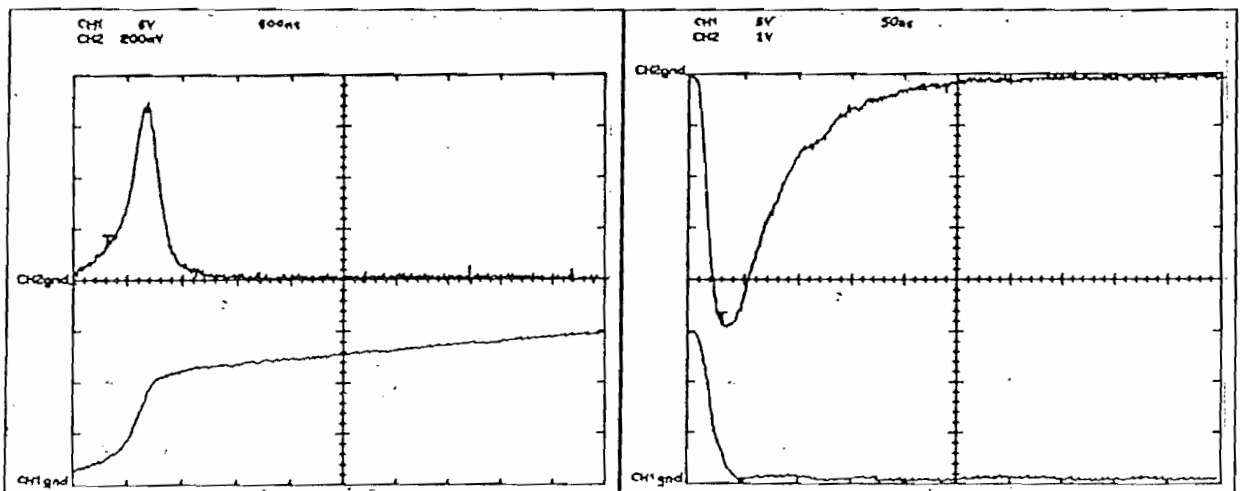
Fig. VI-15 : Tension et courant des plots de sortie modulateurs.



(g) Niveau logique : 12 V. Tension de réglage : 7.5 V.



(h) Niveau logique : 15 V. Tension de réglage : 10.26 V



(i) Niveau logique : 15 V. Tension de réglage : 13.27 V.

Fig. VI-15 : Tension et courant des plots de sortie modulateurs.

En fait, plus la tension de réglage est élevée, plus le temps de transition est lent et les courants débités sont faibles. Ceci s'explique par le fait que la tension de réglage coïncide le canal des transistors T1 (fig. VI-8) et les courants débités par ces transistors (par conséquent par les plots) sont faibles, ce qui implique un temps plus lent pour charger les capacités de charge.

Les courbes de la figure VI-15 montrent des impulsions assez importantes en courant, de valeur maximale variable en fonction de la tension de réglage, entre 800 μA et 230 μA pendant des durées d'environ 600 ns.

Les courbes des figures VI-14 (g,h et i) montrent la possibilité de fonctionnement sous 12 et 15 V .

Ce test a été réalisé pour les plots fabriqués. En faisant le même test pour les plots standards, ces plots n'ont pas pu commuter à cause des faibles courants débités.

En utilisant des sondes avec une haute impédance, ce qui enlève l'ambiguïté sur la mesure des courants avec des sondes ordinaires, nous avons refait les mesures des courants avec une résistance R_2 de 736 Ω , et nous avons obtenu des courbes comparables à celles obtenues précédemment. Nous présentons dans le tableau suivant les résultats des nouvelles mesures.

Tension de réglage en V		Temps de transition en ns	Valeur maximale en μA
7.07	montée	150	680
	descente	50	1500
7.77	montée	150	540
	descente	50	1500
8.2	montée	450	210
	descente	150	1500

Ce test montre bien le fonctionnement des plots sous des tension de 10 à 15 V en débitant des courant allant jusqu'à 1.5 mA impulsionnels. Le temps de commutation est fonction de la tension de réglage. La tension de réglage doit être le plus faible pour les commutations rapides.

VI-5 PERSPECTIVES

Nous proposons par la suite plusieurs modifications qui sont censées d'améliorer les performances des prochains circuits du projet.

VI-5-1 Problème de division de tension

Nous avons relevé dans le paragraphe VI-4-2-d le problème de division de tension entre la capacité de l'unité de reconnaissance et les capacités internes des photodiodes. Ce problème est dû au fait que les deux sortes de capacité sont faites de la même technologie, donc la valeur de la capacité est proportionnelle à la surface. La surface actuelle de la capacité de reconnaissance est de $4500 \mu\text{m}^2$ tandis que celle d'une photodiode est de $2500 \mu\text{m}^2$, le rapport est très favorable pour la division de tension.

Pour éviter ce problème de division de tension entre les capacités, Il faut augmenter la valeur de la capacité du circuit de reconnaissance pour être 30 fois plus grande que la capacité d'une photodiode, à ce moment la division de tension diminue la tension de la capacité de reconnaissance d'environ 1V dans le cas où toutes les photodiodes sont déchargées, ce qui est tolérable. Compte tenu de la surface d'une photodiode qui est $2500 \mu\text{m}^2$, la nouvelle surface de la capacité du circuit de reconnaissance doit être de $75000 \mu\text{m}^2$.

Vue la valeur de cette surface de la capacité de l'unité de reconnaissance, il est peut être plus efficace d'envisager une autre conception de fonctionnement électronique de chaque processeur élémentaire. Dans ce qui suit, nous proposons des modifications dans le processeur élémentaire en ce qui concerne la procédure de reconnaissance.

VI-5-2 Changement de la procédure de reconnaissance

La comparaison entre le contenu de la mémoire d'entrée et l'état d'éclairage de la photodiode correspondante se fait à travers une porte XOR suivi d'une porte NOR pour déterminer la correspondance entre les 7 mémoires d'entrée et l'état d'éclairage des 7 photodiodes, la sortie de cette porte est connectée directement à l'entrée du point mémoire de l'unité de reconnaissance (voir Fig VI-16).

Les avantages de ce nouveau système sont les suivants :

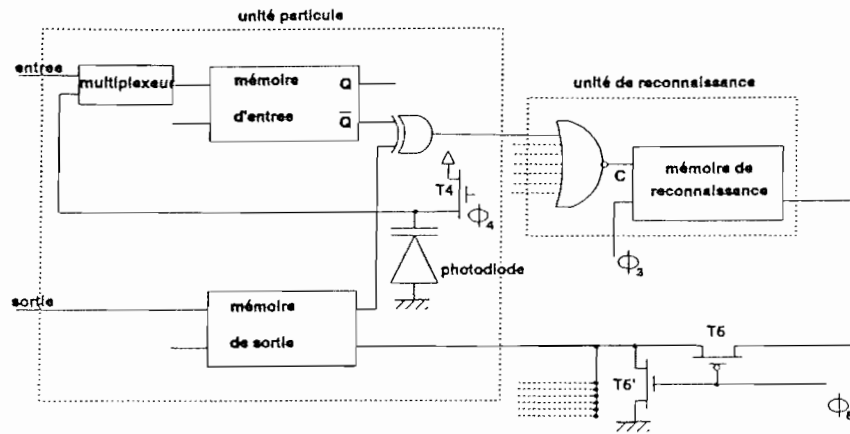


Fig. VI-16 : Schéma modifié du circuit.

- Elimination des deux phases Φ_1 et Φ_2 dans la procédure de reconnaissance en plus de l'élimination de la phase de chargement de la capacité de l'unité de reconnaissance. Ceci diminue, de trois, le nombre des commandes électriques aussi bien que le nombre de plots.
- Simplifier la procédure de reconnaissance qui consiste seulement en une seule projection de la configuration à reconnaître au lieu de la projection de la configuration et de la configuration complémentaire.
- Racourcir le temps de reponse des photodiodes à cause du rapport des valeurs des capacités à décharger.

Il est certain que l'inconvénient de cette proposition est de rendre plus compliquées les fonctions électroniques et d'augmenter le nombre de transistors par processeur élémentaire. Pour examiner cette complexité avec l'augmentation éventuelle de la surface occupée, nous avons établi des schémas électronique des portes XOR et NOR [44].

La porte XOR se compose de 6 transistors selon le schéma (Fig VI-17) indiqué la surface totale occupée par chaque porte est de $86 \times 43 \mu\text{m}^2$, cette surface est suffisamment petite pour pouvoir implémentée dans chaque unité sans modifier la surface de chaque unité.

La porte NOR à 7 entrées se compose de 14 transistors comme l'indique la figure VI-18. La surface occupée est de $94 \times 57 \mu\text{m}^2$ qui est encore plus petite que la surface occupée par le condensateur actuelle de l'unité de reconnaissance.

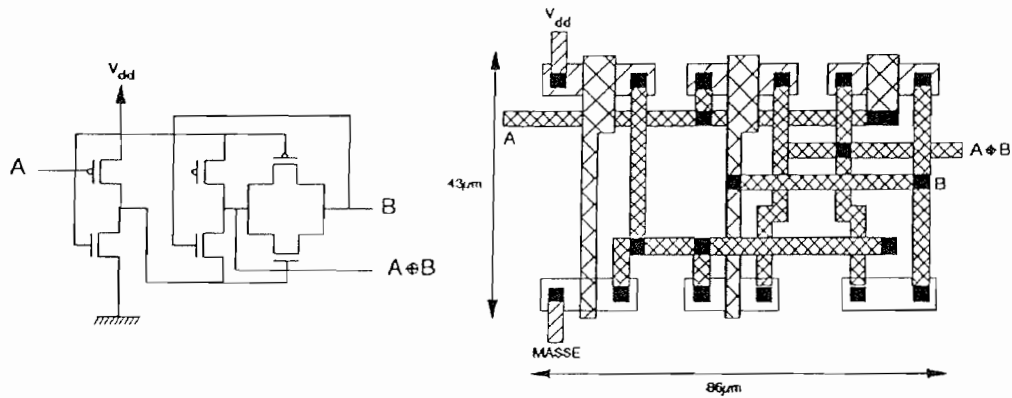


Fig. VI-17 : Porte XOR.

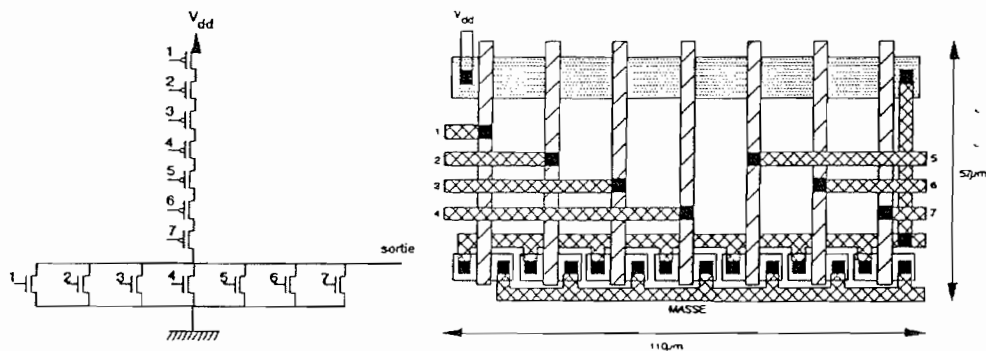


Fig. VI-18 : Porte NOR.

Le bilan de ces modifications est :

- simplifier la procédure de reconnaissance et des commandes optique et électroniques.
- Aucune augmentation de la surface occupée.
- Augmentation du nombre des transistors par processeur élémentaire égale à :

$$(7 \cdot (6-3)) + 14 - 1 = 34 \text{ transistors.}$$

VI-5-3 Modification pour l'étape de substitution

Si la sortance (fan out) le permet, nous pouvons implanter le système des transistors T5 et T5' en raison d'un par cellule au lieu d'être un par unité comme c'est le cas pour le premier circuit. Ceci permettra de réduire de 12 le nombre de transistors par cellule (voir Fig. VI-16).

VI-5-4 Initialisation du circuit

Nous envisageons d'initialiser le circuit (affecter les mémoires d'entrée par des valeurs déterminées) par voie optique en utilisant les mêmes photodiodes. Un multiplexeur à deux entrées commandées par un signal "s" [44] permet de connecter les mémoires d'entrée soit aux mémoires de sortie des cellules voisines pour le fonctionnement en cours du circuit, soit directement aux photodiodes pour l'initialisation (voir Fig VI-19). Pour l'initialisation nous chargeons les capacités des photodiodes par le signal Φ_4 , ensuite nous éclairons les photodiodes par la configuration complémentaire à initialiser, et en fin, nous enregistrons dans les mémoires d'entrée par le signal A.

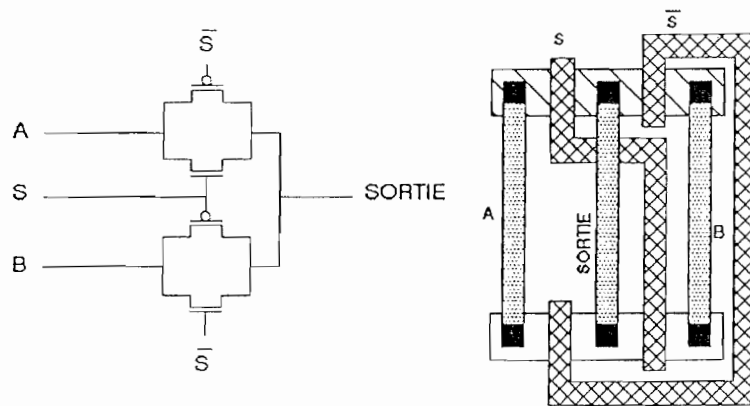


Fig. VI-19 : Schéma d'un multiplexeur.

VI-6 CONCLUSION

Nous avons vérifié le bon fonctionnement des différentes parties du circuit. En revanche, le fonctionnement global du circuit présente quelques phénomènes imprévus, comme la division de tension entre la capacité de l'unité de reconnaissance et les capacités internes des photodiodes, qui est due essentiellement à la faible valeur de la capacité de l'unité de reconnaissance. Il faut envisager de mettre une capacité de plus grande valeur dans les prochaines réalisations ou de modifier le fonctionnement électronique de reconnaissance.

Les plots fabriqués semblent plus adaptés pour la connexion au circuit AsGa que les plots standards.

Des modifications sont proposées pour améliorer les performances du circuit.

CHAPITRE VII

ETUDE OPTO-ELECTRONIQUE

CHAPITRE VII

ETUDE OPTO-ELECTRONIQUE

Le circuit intégré VLSI a deux sortes d'entrée : électronique et optique. Les entrées électronique sont étudiées, expérimentées et adaptées au fonctionnement du reste du circuit aussi bien sur le plan énergétique que sur celui de la fréquence de travail. Les entrées opto-électroniques, les photodiodes par exemple, utilisent une source d'énergie différente dont il faut préciser les caractéristiques nécessaires pour qu'elle soit adaptée au fonctionnement du reste du circuit électronique.

Dans ce chapitre, nous présentons une modélisation du comportement des photodiodes utilisées dans les circuits intégrés en fonction de l'éclairage optique. Cette modélisation prend en compte la profondeur de pénétration des photons en fonction de la longueur d'onde optique utilisée, le taux de génération des porteurs (électrons-trous) et leurs longueurs de diffusion. Les différents facteurs régissant le comportement des porteurs générés détermineront le photo-courant en fonction de la puissance et de la longueur d'onde optiques. Des cas particuliers sont traités pour pouvoir comparer la modélisation à quelques résultats expérimentaux déjà établis.

VII-1 CONSTITUTION ET FONCTIONNEMENT DES PHOTODIODES

VII-1-1 Composition

Une photodiode est une diode polarisée en inverse (mode bloqué). Les photons génèrent, par ionisation, des couples électron-trou qui créent un photo-courant traversant la photodiode dans le sens inverse [45]. Ce photo-courant est fonction d'une part de la longueur d'onde et de la puissance optique, et d'autre part des caractéristiques des photodiodes comme l'épaisseur et le dopage des différentes régions ainsi que la tension appliquée.

Les photodiodes utilisées sont dites perpendiculaires parce que le plan de la jonction est perpendiculaire au faisceau lumineux. Nous distinguons deux types de photodiodes : N et P; nous décrivons ci-dessous leur composition et leurs caractéristiques.

Photodiode type N

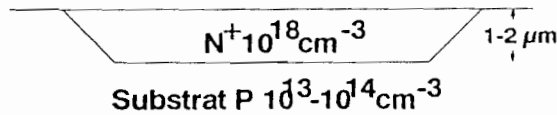


Fig VII-1 : Structure photodiode N.

Elle est constituée d'une région dopée N^+ (dopage de 10^{18} cm^{-3}) d'épaisseur 1-2 μm directement implantée sur le substrat de dopage P (10^{13} - 10^{14} cm^{-3})

Photodiode type P

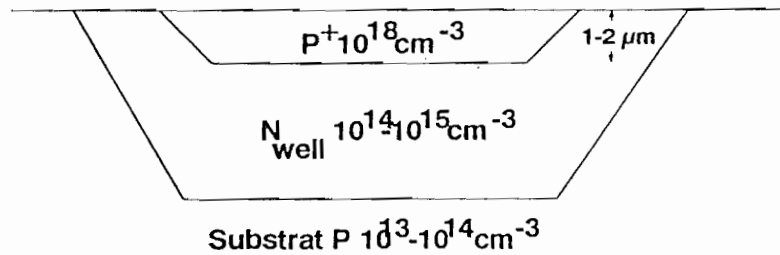


Fig. VII-2 : Structure d'une photodiode P.

Elle est constituée d'une région dopée P^+ (dopage de 10^{18} cm^{-3}) d'épaisseur 1-2 μm implantée sur un caisson dopé N (10^{14} - 10^{15} cm^{-3}) et d'épaisseur 10 μm sur le même substrat P (10^{13} - 10^{14} cm^{-3}).

VII-1-2 Caractéristiques des photodiodes

Plusieurs caractéristiques des photodiodes sont déterminantes pour la sensibilité :

a) La zone de charge d'espace

C'est la zone déplétée des porteurs libres par suite de la diffusion de ces porteurs, il y règne un champ électrique intense qui sépare tout couple électron-trou généré dans cette zone.

La largeur de cette zone est fonction du dopage des différentes régions ainsi que de la tension appliquée aux bornes de la photodiode. Pour une polarisation inverse la formule qui donne l'épaisseur de cette zone est :

côté P

$$W_P = \sqrt{\frac{2\varepsilon \cdot (V_D - V_R)}{q \cdot N_A \cdot \left(1 + \frac{N_A}{N_D}\right)}}$$

côté N

$$W_N = \sqrt{\frac{2\varepsilon \cdot (V_D - V_R)}{q \cdot N_D \cdot \left(1 + \frac{N_D}{N_A}\right)}}$$

tel que V_D est le potentiel de diffusion interne de la jonction P-N, V_R est la tension extérieure appliquée, ε la permittivité du semiconducteur, q la charge élémentaire, N_D le taux du dopage côté N et N_A le taux du dopage côté P.

Le tableau suivant donne les valeurs de W_P et de W_N en fonction des dopages utilisés sous la tension de polarisation inverse $V_A = 5 \text{ V}$:

	Type N		Type P	
	$N_D = 10^{18}$		$N_A = 10^{18}$	
	$p_P = 10^{13}$	$p_P = 10^{14}$	$n_N = 10^{14}$	$n_N = 10^{15}$
$W_P \text{ } \mu\text{m}$	25	8	$2.5 * 10^{-3}$	$7.9 * 10^{-4}$
$W_N \text{ } \mu\text{m}$	$2.5 * 10^{-4}$	$7.9 * 10^{-4}$	8	2

c- La capacité interne de la photodiode

A la suite de la diffusion des porteurs libres, il existe une charge négative côté P et une charge positive côté N, ce qui est équivalent à une capacité dont la valeur est donnée en fonction des dopages et de la tension inverse appliquée V_R par la formule :

$$C = \sqrt{\frac{q \cdot N_D \cdot \varepsilon}{2 \cdot \left(1 + \frac{N_D}{N_A}\right) \cdot (V_D - V_R)}}$$

Compte-tenu des valeurs numériques pour le Silicium:

$$\begin{aligned} q &= 1,6 * 10^{-19} \text{ C} \\ \epsilon &= 10^{-12} \text{ F/cm} \\ N_D &= 10^{14} \text{ cm}^{-3} \\ N_A &= 10^{18} \text{ cm}^{-3} \end{aligned}$$

$$\text{nous trouvons } C = \frac{2.8}{\sqrt{V_D - V_R}}$$

$$\text{pour } V_A = 5 \text{ V}, C = 1.25 \text{ nF/cm}^2$$

d- Le coefficient d'absorption

Le flux de la lumière à une certaine profondeur x dans la matière est donné par la loi de Beer-Lambert :

$$\Phi(x) = \Phi_0 \cdot \exp(-\alpha \cdot x)$$

tel que α est le coefficient d'absorption.

La figure VII-3 donne le coefficient d'absorption du Silicium [45] :

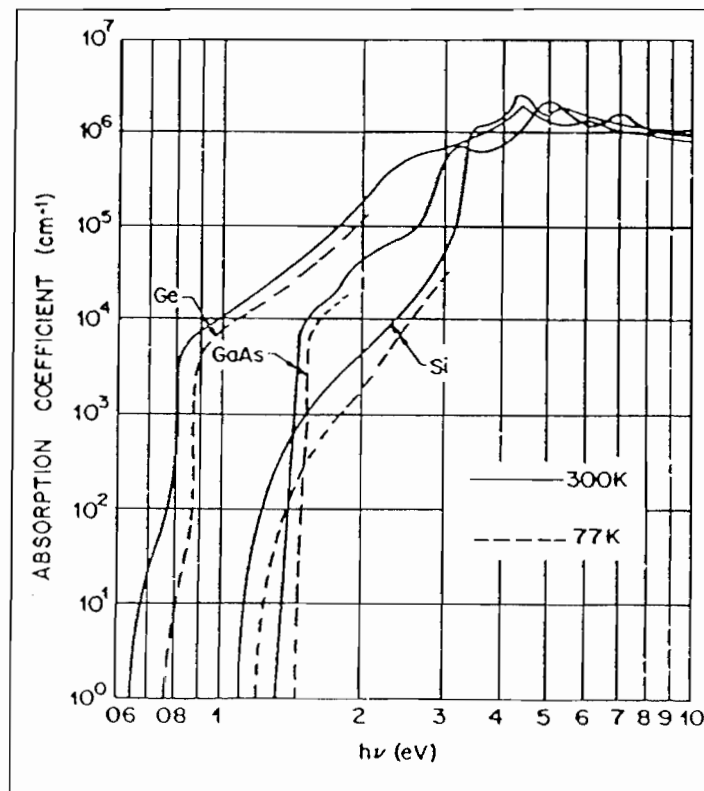


Fig. VII-3 : Coefficient d'absorption du Silicium.

A 300 °K et pour la longueur d'onde de 0.6328 μm , nous estimons le coefficient d'absorption à :

$$\alpha = 4 \cdot 10^3 \text{ cm}^{-1}$$

e- La longueur de diffusion des porteurs

Les porteurs créés par effet photo-ionisation ont une durée de vie déterminée avant de se recombiner; Cette durée de vie leur permet de parcourir une certaine distance qui dépend de la mobilité des porteurs, nous appelons cette distance *la longueur de diffusion*. Cette longueur est déterminante pour le photo-courant car c'est elle qui détermine si les porteurs générés hors la zone de charge d'espace seront séparés ou non.

La longueur de diffusion est différente selon que le porteur est un trou (h) ou électron (e). Elle est donnée en fonction du coefficient de diffusion D_k ($k = h$ ou e) du porteur en question et de la durée de vie τ_k de ce porteur par la relation :

$$L_k = \sqrt{D_k \cdot \tau_k}$$

Le coefficient de diffusion est, à son tour, fonction de la mobilité du porteur μ_k selon la relation :

$$D_k = \mu_k \cdot \frac{kT}{q}$$

La figure suivante donne la mobilité des électrons et des trous dans le Silicium à 300 °K en fonction de la concentration d'impuretés [46], ([47]).

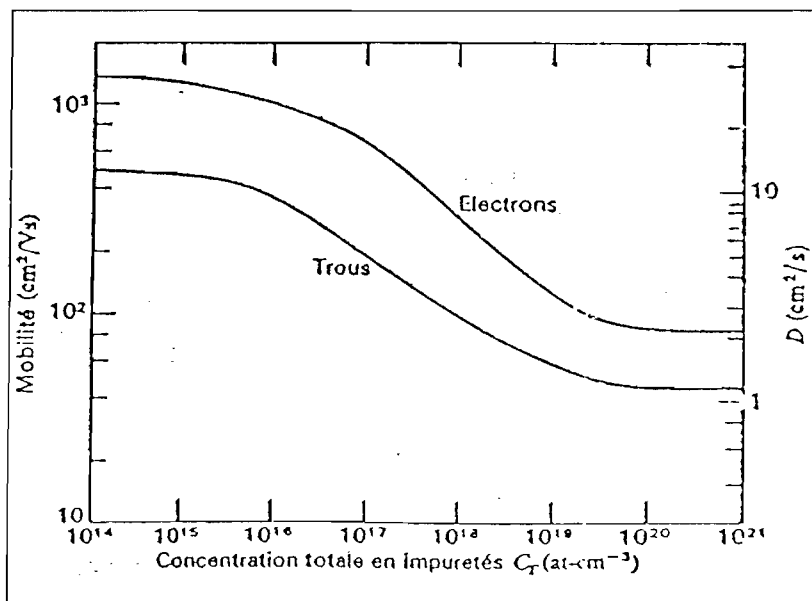


Fig VII-4 : Mobilité des porteurs en fonction de la concentration.

Le tableau suivant donne les valeurs de la mobilité et du coefficient de diffusion à 300 °K pour les concentrations utilisées

Dopage	μ_e $\text{cm}^2\text{V}^{-1}\text{s}^{-1}$	μ_h $\text{cm}^2\text{V}^{-1}\text{s}^{-1}$	D_e cm^2s^{-1}	D_h cm^2s^{-1}
10^{18} cm^{-3}	250	120	8	4
10^{15} cm^{-3}	1020	350	22	7
10^{14} cm^{-3}	1150	500	30	11

La durée de vie des porteurs minoritaires dépend de leurs natures (électrons ou trous), de la nature du dopage et de la concentration du dopage. Le tableau suivant présente les estimations de la durée de vie des porteurs minoritaires dans les régions des photodiodes étudiées :

Région	Dopage cm^{-3}	$\tau_{\text{minoritaires}}$
N	10^{14} - 10^{15}	10 -100 ns
P+	10^{18}	1 -10 ns
P	10^{13} - 10^{14}	10 -100 ms
N+	10^{18}	10 -100 ns

e- Profondeur de la jonction P-N

La profondeur de la jonction est déterminante pour la sensibilité de la photodiode car elle précise la proportion de l'énergie lumineuse efficace pour le photo-courant compte-tenu de l'absorption de cette énergie dans la couche supérieure.

Les photodiode sont réalisées par diffusion de dopants dans la matrice du Silicium, la profondeur de la diffusion dépend de la nature des photodiodes : pour les photodiodes N la diffusion se fait directement sur le substrat P ; par contre, pour les photodiodes P, la diffusion se fait dans un caisson déjà diffusé. Les profils de la concentration des dopants en fonction de la profondeurs sont présentés dans la figure suivante :

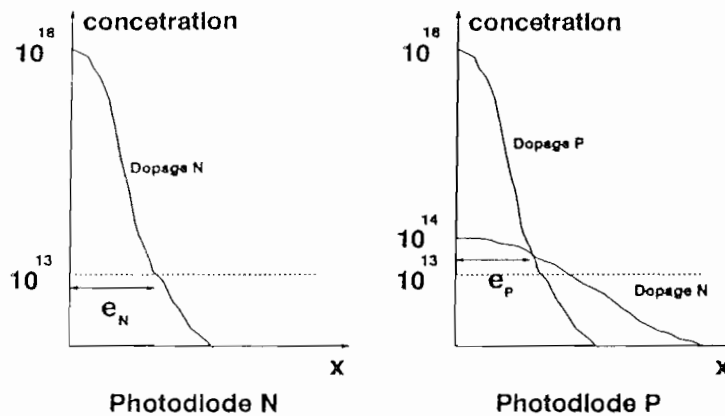


Fig VII-5 : Profondeur de diffusion du dopage.

Nous déduisons de cette figure que la jonction de la photodiode N est plus profonde que celle de la photodiode P.

f- Courant de fuite

Le courant de fuite est le courant de saturation qui traverse la diode en cas de polarisation inverse, ce courant est faible et constant quelque soit la valeur de la tension de polarisation. L'importance de ce courant se manifeste au moment où la photodiode n'est pas éclairée : L'éclairage de la photodiode a pour but, en général, d'amener la tension à ses bornes à une certaine valeur, par déchargement de sa capacité interne. Cette nouvelle tension exprime l'information portée par l'éclairage. Dans l'obscurité, aucune information n'est reçue par la photodiode, par conséquent la tension aux bornes de la photodiode ne doit pas changer, mais le courant de fuite décharge cette capacité, donc porte une information parasite dont l'importance est directement liée à la valeur de ce courant de fuite.

La formule qui donne la valeur de la densité du courant de fuite (où courant de saturation) est la suivante :

$$J_s = q \cdot n_i^2 \left(\frac{D_e}{L_e \cdot N_A} + \frac{D_h}{L_h \cdot N_D} \right)$$

n_i : La concentration des porteurs libre dans le Silicium intrinsèque.

q : La charge élémentaire

Nous avons estimé les valeurs de ce courant pour des photodiodes de surface égale à $40 \times 40 \mu\text{m}^2$:

Photodiode N : $I_{SDN} = 3 \cdot 10^{-14}$ Ampère
 Photodiode P : $I_{SDP} = 2.93 \cdot 10^{-14}$ Ampère

La première valeur est en accord avec les résultats expérimentaux [48]. Nous remarquons qu'il n'y a pas de différence sensible entre les deux sorte des photodiodes.

VII-2 MODELISATION DU RENDEMENT QUANTIQUE DES PHOTODIODES

Un des facteurs déterminants de la réponse d'une photodiode est le rendement quantique η défini comme le rapport entre le nombre des porteurs électriques participant au photocourant et le nombre de photons incidents [45].

Le mécanisme de génération du photocourant dans la photodiode est le suivant :

Les photons incidents sont absorbés dans le Silicium selon la loi de Beer-Lambert présenté précédemment. Chaque photon absorbé génère un couple électron-trou tant que l'énergie du photon est supérieur à la largeur de la bande interdite du semi-conducteur :

$$h\nu > E_G \text{ ou } \lambda < 1.24/E_G$$

ce qui donne pour le Silicium $\lambda < 1.1 \mu\text{m}$.

Le taux de génération de ces porteurs électriques dépend de la profondeur x de leur absorption, et est désigné par $G(x)$. Pour que les porteurs générés participe au photocourant, ils doivent être séparés sinon ils se recombinent après un temps caractéristique de leurs durées de vie. Cette séparation n'est possible que par le champ électrique dans la zone de charge d'espace. La portion des porteurs générés qui participera au photo-courant est désigné par $R(x)$.

VII-2-1- Calcul de $G(x)$

Le nombre de photons absorbés à la profondeur x de la surface est la différence du flux lumineux entre les couches x et $x+dx$ d'où :

$$\begin{aligned} G(x).dx &= -(\Phi(x + dx) - \Phi(x)) \\ &= -d\Phi/dx \, dx \\ &= \Phi_0 \cdot \alpha \cdot \exp(-\alpha.x) \, dx \end{aligned}$$

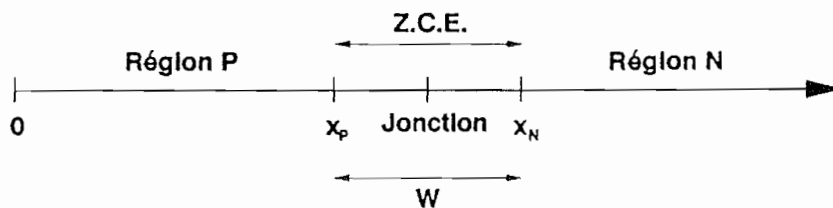
d'où :

$$G(x) = \Phi_0 \cdot \alpha \cdot \exp(-\alpha \cdot x)$$

VII-2-2- Calcul de $R(x)$

Ce calcul est fait pour une photodiode P, le résultat reste le même pour une photodiode N en changeant les porteurs minoritaires et les régions.

Pour le coefficient $R(x)$, nous distinguons deux régions : la zone de charge d'espace limitée entre les abscisses x_p et x_N , et la zone dans les deux régions limitant la zone précédente. Le comportement des porteurs générés par photo-ionisation est différent dans ces deux zones :



- Dans la zone de charge d'espace : Compte-tenu du champ électrique intense qui règne dans cette région, tous les électron-trous générés seront séparés et atteignent les bords de la photodiode. par conséquent le coefficient :

$$R(x) = 1 \quad \text{pour} \quad x_p < x < x_N$$

- En dehors de la zone de charge d'espace les électrons et les trous sont séparés dans la mesure où la diffusion permet aux porteurs minoritaires d'atteindre la zone de charge d'espace. La probabilité d'atteindre cette zone dépend de la longueur de diffusion par la relation (pour une photodiode de type P) :

$$R(x) = \exp((-x_p + x)/L_{ep}) \quad \text{pour} \quad 0 < x < -x_p$$

$$R(x) = \exp(x_N - x)/L_{enN}) \quad \text{pour} \quad x_N < x$$

VII-2-3- Calcul de la densité du photo-courant

Nous divisons la photodiode en trois régions : la régions de surface (P) qui est la plus dopée, la zone de charge d'espace (ZCE) et la région profonde (N) que nous supposons infinie. Compte-tenu que la région superficielle est beaucoup plus dopée que la région profonde, la zone de charge d'espace se trouve presque exclusivement dans la région profonde.

Dans le calcul suivant k désigne les porteurs minoritaires.

La densité de courant des porteurs minoritaires de la région superficielle est :

$$\begin{aligned}
 J_{ep} &= -q \cdot \int_0^{x_p} G(x) \cdot R(x) dx \\
 &= -q\alpha\Phi_0 \int_0^{x_p} \exp(-\alpha x) \cdot \exp\left[\frac{x - x_p}{L_{ep}}\right] dx \\
 &= -q\alpha\Phi_0 \left\{ \frac{L_{ep}}{1 - \alpha L_{ep}} \exp(-\alpha x_p) - \exp\left(-\frac{x_p}{L_{ep}}\right) \right\}
 \end{aligned}$$

La densité de courant des porteurs minoritaires de la région profonde est :

$$\begin{aligned}
 J_{hN} &= -q \int_{x_N}^{\infty} G(x) \cdot R(x) dx \\
 &= -q\alpha\Phi_0 \int_{x_N}^{\infty} \exp(-\alpha x) \cdot \exp\left(-\frac{x - x_N}{L_{eN}}\right) dx \\
 &= -q\alpha\Phi_0 \cdot \frac{L_{ep}}{1 + \alpha L_{eN}} \cdot \exp(-\alpha x_N)
 \end{aligned}$$

La densité de courant de la zone de charge d'espace est :

$$\begin{aligned}
 J_{ZCE} &= -q \int_{x_p}^{x_N} G(x) dx \\
 &= -q\Phi_0 \cdot \exp(-\alpha x_p) \cdot (\exp(-\alpha W) - 1)
 \end{aligned}$$

VII-2-4 Evaluation du rendement quantique

Nous avons vu que le rendement quantique est défini comme le rapport du nombre de charges élémentaires collectées par la jonction au nombre de photons incidents. Son expression est donnée en fonction des grandeurs macroscopiques par la relation :

$$\eta = J_{ph}/q\Phi_0$$

J_{ph} est la densité de photocourant ($A \cdot cm^{-2}$);

q est la charge élémentaire (Coulomb);

Φ_0 est le flux incident ($photons \cdot cm^{-2} \cdot s^{-1}$)

D'après le calcul précédent du photocourant, l'expression générale du rendement quantique est donnée par :

$$\eta = \left[\frac{L_{eP}}{1 - \alpha L_{eP}} \left[\exp\left(-\frac{x_P}{L_{eP}}\right) - \exp(-\alpha x_P) \right] \right. \\ \left. + \exp(-\alpha x_P) \cdot [1 - \exp(-\alpha W)] \right. \\ \left. + \frac{L_{hN}}{1 + L_{hN}} \cdot \exp(-\alpha x_N) \right]$$

A partir de la formule précédente nous avons évalué le rendement quantique en tenant compte des valeurs suivantes pour le coefficient de diffusion et pour la durée de vie :

Photodiode N :

Pronfondeur de la jonction $x_N = 2\mu\text{m}$
 Largeur de la Z.C.E. $W_N = 8\mu\text{m}$

Région N

Coefficient de diffusion $D_{hN} = 4 \text{ cm}^2\text{s}^{-1}$
 Durée de vie des trous $\tau_{hN} = 10 \text{ ns}$

Région P

Coefficient de diffusion $D_{eP} = 30 \text{ cm}^2\text{s}^{-1}$
 Durée de vie des trous $\tau_{hP} = 100 \mu\text{s}$

Photodiode P :

Pronfondeur de la jonction $x_P = 1,1.5,2 \mu\text{m}$
 Largeur de la Z.C.E. $W_P = 2 \mu\text{m}$

Région P

Coefficient de diffusion $D_{eP} = 8 \text{ cm}^2\text{s}^{-1}$
 Durée de vie des trous $\tau_{eP} = 1 \text{ ns}$

Région N

Coefficient de diffusion $D_{hN} = 7 \text{ cm}^2\text{s}^{-1}$
 Durée de vie des trous $\tau_{hN} = 100 \text{ ns}$

Les résultats du calcul sont donnés dans le tableau suivant :

	Photodiode N	Photodiode P		
Profondeur de la jonction	2 μm	1 μm	1.5 μm	2 μm
Rendement quantique η	0.54	0.55	0.48	0.41

Nous constatons que la valeur du rendement quantique pour la longueur d'onde 0.6328 μm est d'environ 0.5 pour les deux types des photodiodes, par conséquent la sensibilité est presque la même pour les deux types de photodiode.

VII-3 INFLUENCE DE LA CAPACITE INTERNE DES PHOTODIODES SUR LE TEMPS DE REPONSE

Comme la réponse des photodiodes est caractérisée par la variation de la tension à ces bornes, sa capacité interne joue un rôle important sur ce temps de réponse. En fait, les expériences menées l'IEF [49] ont montré que le temps de réponse des photodiodes de types P est 3 fois plus grand que celui des photodiodes de type N. Cette différence ne peut pas être expliquée par la différence de sensibilité qui est presque identique pour les deux photodiodes. Cette différence peut être expliquée par la différence de la capacité interne des deux sorte des photodiodes utilisées.

Le calcul tient en compte les concentrations suivantes des impuretés dans les deux types des photodiodes :

En utilisant la formule de la capacité interne des photodiodes nous avons :

pour une photodiode N :

$$C_N = \sqrt{\frac{q\epsilon \cdot 10^8}{1 + \frac{10^8}{10^3}} \cdot \frac{1}{2(V_D - V_R)}} = \sqrt{\frac{q\epsilon \cdot 10^3}{2(V_D - V_R)}}$$

pour une photodiode P

$$C_P = \sqrt{\frac{q\epsilon \cdot 10^8}{1 + \frac{10^8}{10^4}} \cdot \frac{1}{2(V_D - V_R)}} = \sqrt{\frac{q\epsilon \cdot 10^4}{2(V_D - V_R)}}$$

d'où

$$\frac{C_P}{C_N} = \sqrt{10} = 3.16$$

et compte-tenu que le temps de réponse est proportionnel à la capacité, le temps de réponse des photodiodes P est 3 fois plus grand que le temps de réponse des photodiodes N [49].

VII-4 EVALUATION DE LA SENSIBILITE DES PHOTODIODES EN FONCTION DE λ

Nous définissons la sensibilité S comme le rapport entre le courant délivré par la photodiode en Ampère et la puissance optique reçue en Watt. La relation qui relie la sensibilité au rendement quantique est :

$$S = \frac{q}{h\nu} \eta = \frac{q}{h \cdot c} \lambda \cdot \eta$$

d'où :

$$S = 8.06 \cdot 10^5 \cdot \lambda \cdot \eta$$

Pour évaluer la sensibilité des photodiodes en fonction de la longueur d'onde, nous sommes partis de la courbe suivante qui donne le rendement quantique en fonction de la longueur d'onde [45].

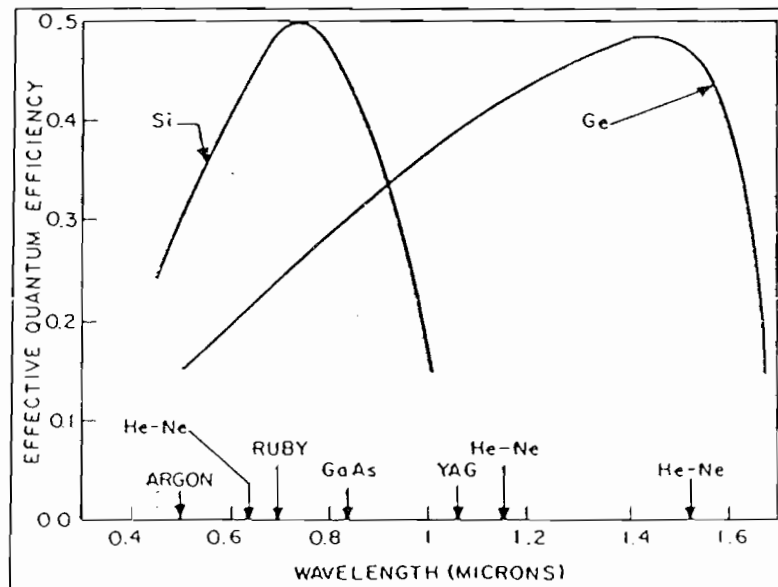


Fig. VII-6 : Rendement quantique.

La coupure de la courbe du rendement quantique côté infrarouge est due à l'absorption fondamentale en liaison directe avec la largeur de la bande interdite du matériau. Celle côté

ultra-violet est due à l'absorption superficielle que nous avons modélisée dans le paragraphe VII-2-4.

A partir de la courbe ci-dessus, nous avons modélisé le rendement quantique pour le Silicium en fonction de la longueur d'onde par regression polynomiale d'ordre 4, en imposant une valeur de 0.5 et une dérivée nulle à $\lambda = 0.74 \mu\text{m}$.

Le résultat de cette interpolation est présenté dans la figure suivante, cette interpolation nous permet de faire une modélisation prédictive du comportement de la sensibilité en fonction de la longueur d'onde comme elle est présentée sur la même figure :

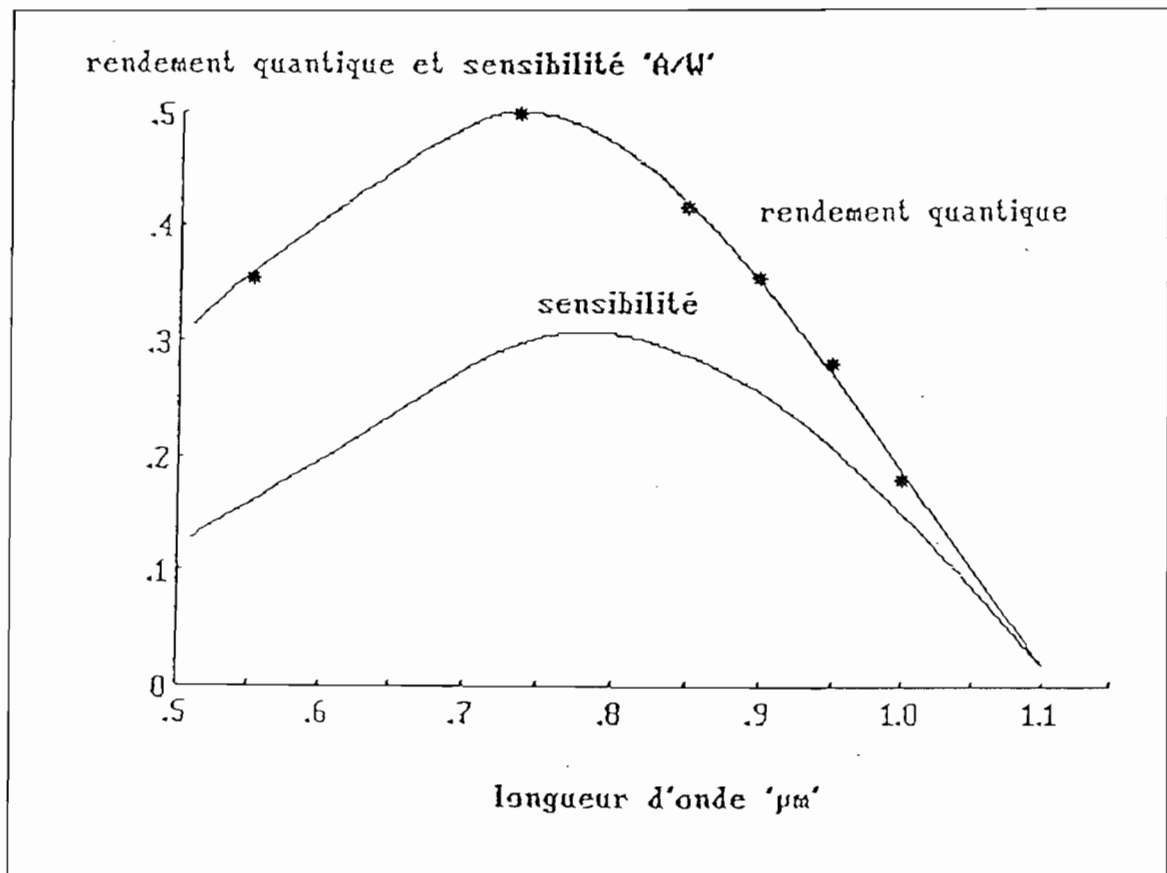


Fig. VII-7 : Sensibilité des photodiodes Silicium.

VII-5 ETUDE DE CAS

Nous prenons comme exemple d'application une rétine particulière (circuit horloge optique de l'IEF) [49]. Ce circuit est composé d'une matrice de cellule sur lesquelles sont distribuées optiquement les tops d'horloge en utilisant un système de deux photodiodes N et P. Les deux photodiodes sont connectées conformément au schéma suivant :

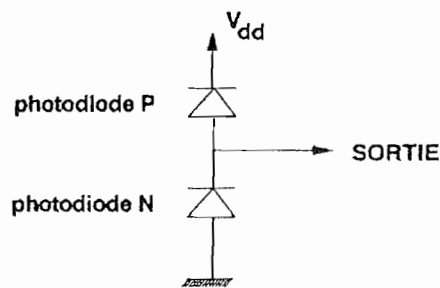


Fig. VII-8 : Montage des photodiodes.

Nous considérons le schéma électrique équivalent suivant pour ce circuit :

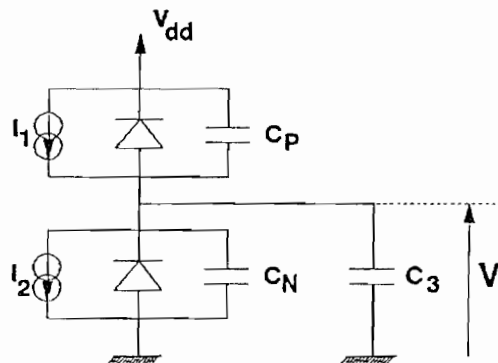


Fig. VII-9 : Schéma électrique.

Les deux photodiodes sont éclairées alternativement. Le temps de montée (ou de descente) de la tension V doit être compatible avec le temps de réponse des composants électronique. Et nous évaluons la puissance optique nécessaire pour assurer un tel temps de réponse.

Nous considérons que les deux photodiodes ont la même sensibilité pour toutes les longueurs d'onde, mais deux capacités internes différentes et deux courants de fuite différents :

Photodiode N :	$C_N = 10 \text{ fF}$
	$I_{SDN} = 0.3 * 10^{-14} \text{ A}$
Photodiode P :	$C_P = 30 \text{ fF}$
	$I_{SDP} = 2.0 * 10^{-14} \text{ A}$

Le reste du circuit est représenté par une capacité en parallèle :

$$C_3 = 50 \quad \text{fF}$$

L'équation qui régit la variation de la tension V :

$$\frac{dV}{dt} = \frac{I_1 - I_2 + I_{SDN}}{C_N + C_P + C_3}$$

VII-5-1 Puissance optique nécessaire

Pour éviter les transitions aléatoires de l'électronique commandée par les impulsions d'horloge, il faut que la pente de la variation de tension dans la région de seuil de l'électronique (située autour de 2.5 V et de largeur d'environ 100 mV) soit supérieure à 100 mV/ns [50] :

$$dV/dt > 10^8 \text{ V/s}$$

Les courants I_1 et I_2 sont des photocourants dépendants de la puissance optique P_1 et P_2 . Ces courants, générés par l'éclairage optique, sont considérés comme nettement plus grands que les courants de fuite des diodes.

En tenant compte des valeurs des capacités, la condition pour assurer la valeur précisée précédemment de la pente de la tension devient :

$$P_1 * S_1 > 10^{-5} \text{ Ampère}$$

Connaissant la courbe de la sensibilité S en fonction de la longueur d'onde, nous pouvons déduire la puissance minimale à chaque longueur d'onde pour éclairer une photodiode.

La figure suivante donne la puissance minimale nécessaire pour éclairer une photodiode en fonction de la longueur d'onde optique, la puissance donnée par la courbe permet à la tension aux bornes de la photodiode de passer la région de seuil de largeur de 100 mV en moins de 1 ns.

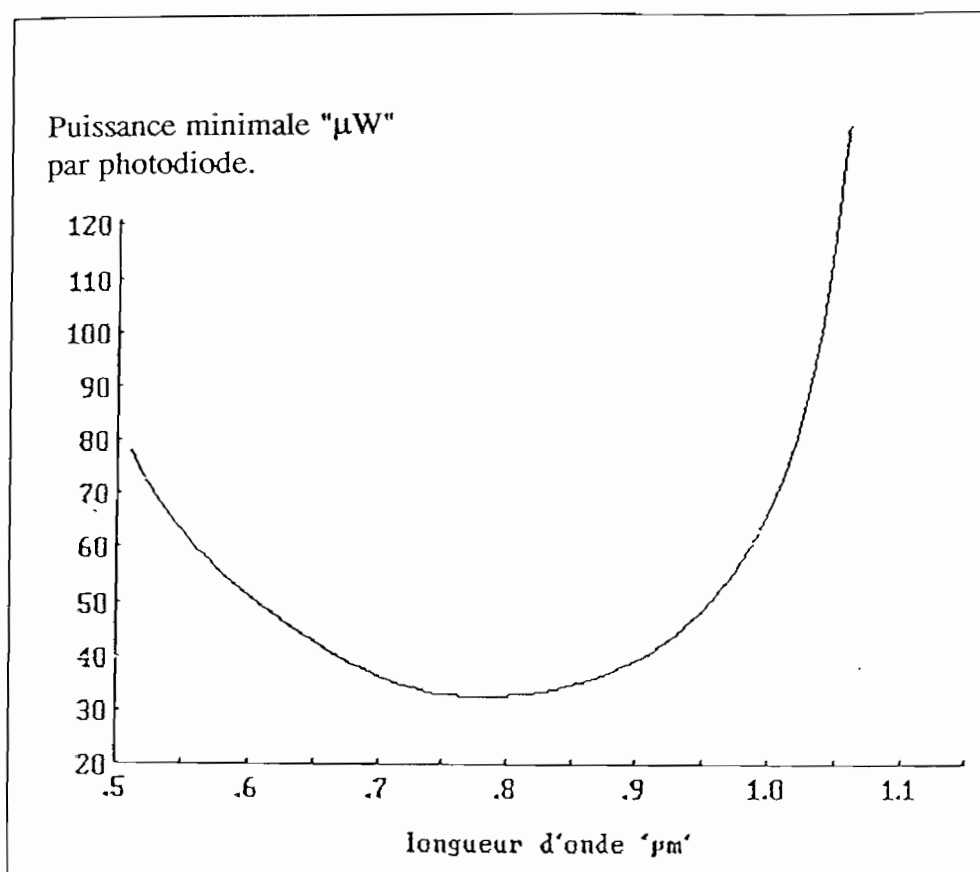


Fig. VII-10 : puissance optique.

VII-5-2 Temps de montée de la puissance optique

La puissance optique nécessite un certain temps T_1 pour monter de sa valeur minimale 0.01% de P_1 jusqu'à sa valeur maximale P_1 . Ce temps T_1 doit être suffisamment petit pour que la tension V n'atteigne pas la région de seuil. En considérant que la tension V part de la valeur 0 V et que le seuil est situé à 2V, nous obtenons la condition suivante pour le temps T_1 :

$$S * P_1 * T_1 < .36 * 10^{-12} \text{ C}$$

ce qui correspond à une certaine valeur de charge électrique pour mener la tension V à la valeur 2 V.

La figure suivante donne le temps de montée maximal de la puissance optique en fonction de la longueur d'onde et de la puissance utilisée à partir de sa valeur minimale.

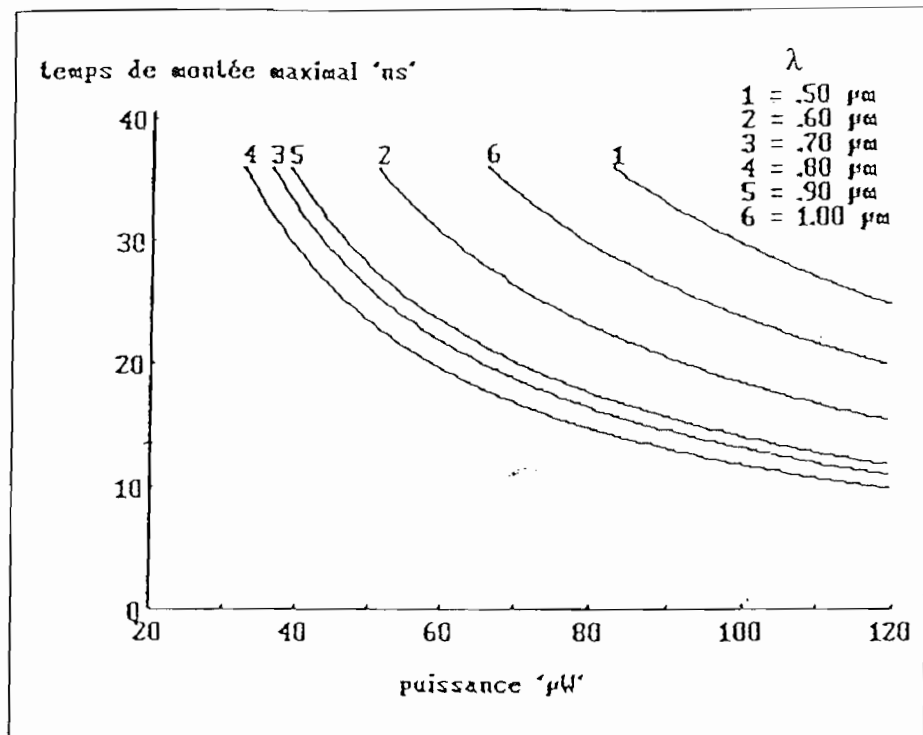


Fig VII-11 : Temps de montée.

VII-5-3 Largeur de l'impulsion optique

Pour que la tension V commute entre les deux valeurs logiques : 0 et 5 V, il faut une énergie optique minimale pour générer la quantité de courant nécessaire. En utilisant une puissance optique donnée, il faut que l'impulsion optique dure un certain temps T_2 . Considérant que l'excursion de la tension V est de 5V nous obtenons la condition suivante pour le temps T_2 :

$$S * P_1 * T_2 > .5636 * 10^{-12} \text{ Coulomb}$$

Ce qui correspond à environ 1.761.250 paires électron-trou.

La figure suivante donne la largeur minimale de l'impulsion optique en fonction de la longueur d'onde et de la puissance optique utilisées :

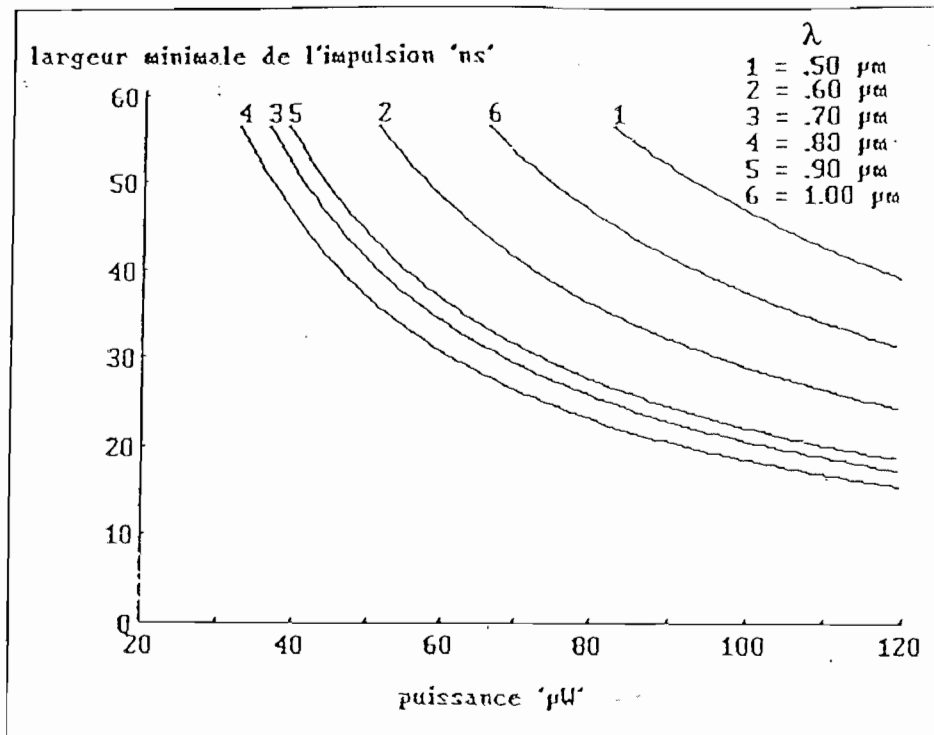


Fig VII-12 : Largeur de l'impulsion.

VII-6 CONCLUSION

Cette étude nous a permis de mieux comprendre les différents facteurs qui déterminent le photo-courant en fonction de la longueur d'onde et de la puissance optique utilisée.

Nous avons étudié particulièrement le circuit de l'horloge optique et nous avons donné les caractéristiques optiques nécessaires en fonction de la longueur d'onde : la puissance optique à utiliser, le temps de montée maximal et la largeur minimale.

CHAPITRE VIII

REALISATION EXPERIMENTALE

CHAPITRE VIII

REALISATION EXPERIMENTALE

Dans ce chapitre, nous présentons un montage de démonstration, pour le "Gaz sur Réseau" contenant le circuit électronique, le système d'éclairage avec une seule diode laser, et le circuit des modulateurs opto-électroniques. L'hologramme de Talbot n'est pas utilisé car le circuit ne contient qu'une seule cellule.

VIII-1 CIRCUIT DE COMMANDE

Le fonctionnement du circuit électronique exige des signaux électriques aussi bien que des signaux optiques. Dans ce paragraphe, nous analysons les différentes séquences des signaux électriques nécessaires au fonctionnement du circuit électronique, ainsi que la réalisation d'une carte électronique pour la génération de telles séquences.

VIII-1-1 Séquence des signaux de commande

Le rôle de ces signaux est l'initialisation, la lecture électronique du circuit et la synchronisation entre les séquence de reconnaissance et de substitution d'une part et l'éclairage des photodiodes d'autre part.

VIII-1-1-a Initialisation et lecture du circuit

Il s'agit d'introduire un mot de 7 bits dans le registre à décalage. Pour cela, il faut générer les impulsions " R_1 ", " R_2 " alternativement comme le montre le schéma ci-dessous, en même temps que la présence les différents bits sur le plot "ENTREE" du circuit. A la fin de cette étape, nous activons le signal "A" pour transférer le contenu du registre dans les mémoires d'entrée du circuit.

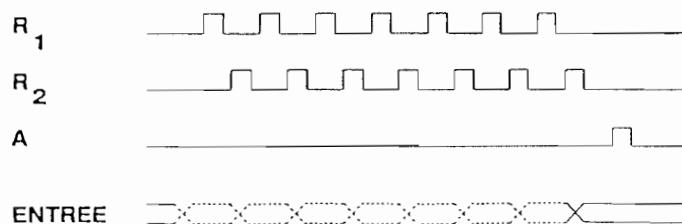


Fig. VIII-1 : Séquences d'initialisation.

Pour la lecture électronique du résultat il suffit de déclencher le signal "A" pour transférer le contenu des mémoires de sortie dans le registre à décalage. Ensuite, les séquences R_1 et R_2 sont inversées par rapport à l'initialisation. A ce moment les bits du registre se retrouvent séquentiellement sur le plot "SORTIE" du circuit.

VIII-1-1-b Séquence pour la reconnaissance

Il faut charger au préalable la capacité "C" de l'unité de reconnaissance avec le signal Φ_0 , ainsi que les capacités internes des photodiodes avec le signal Φ_4 . Pour la reconnaissance du motif enregistré dans les mémoires d'entrée, il faut fermer les transistors T1 avec le signal Φ_1 (Fig. VI-6) en même temps qu'il faut éclairer les photodiodes par le motif à reconnaître. Avant de recommencer la deuxième étape de reconnaissance avec le motif complémentaire, il est nécessaire de recharger les capacités internes des photodiodes avec le signal Φ_4 pour éviter l'effet de division de tension décrit dans le chapitre VI (§ VI-4-2-d-c). La reconnaissance du motif complémentaire est la même que précédemment mais avec le signal Φ_2 .

Les séquences pour la reconnaissance sont présentées dans la figure VIII-2.

VIII-1-1-c Séquences pour la substitution

Un chargement préliminaire des capacités internes des photodiodes avec le signal Φ_4 impose au préalable le niveau logique "1" à toutes les entrées des mémoires de sortie. Ensuite, le signal Φ_5 déclenche l'enregistrement dans les mémoires de sortie, cet enregistrement, nous le rappelons, est conditionné par le niveau logique "1" de l'unité de reconnaissance.

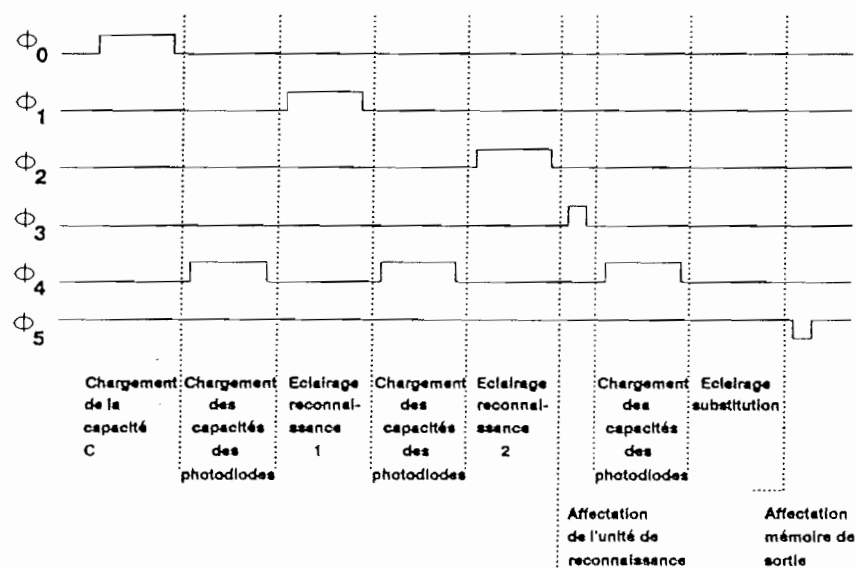


Fig. VIII-2 : Séquence de synchronisation.

VIII-1-2 La carte de commande

Cette carte génère tous les signaux nécessaires au fonctionnement du prototype : l'initialisation du circuit électronique, les séquences de synchronisation, la lecture électronique des résultats et l'allumage des diodes laser.

La carte est constituée des circuits intégrés logiques. La figure VIII-3 montre le schéma de cette carte. Dans la suite nous donnons une brève description de sa conception et de ses fonctions.

a) Système d'initialisation et de lecture

Ce système est constitué d'un registre à décalage à entrée parallèle qui présente séquentiellement la configuration initiale du circuit sur le plot d'entrée "ENTREE". Cette configuration est choisie par les interrupteurs "S1" et validée par le bouton "B1". La sortie des différents bits est synchronisée avec les signaux R_1 et R_2 .

Les signaux R_1 , R_2 et A sont générés par le circuit PAL1 dont le fonctionnement est déclenché par le bouton "B2" séquencé par une horloge extérieure.

Le circuit 7 inverse la chronologie des signaux R_1 , R_2 et A pour la lecture électronique des résultats.

b) Séquences de synchronisation

Ces séquences ($\Phi_0 - \Phi_5$) sont générées par le circuit PAL2. Le déclenchement se fait sur le signal "A" du PAL1 qui détermine la fin de l'initialisation. La suite des signaux $\Phi_0 - \Phi_5$ est répétée autant de fois qu'il y a de configurations à reconnaître et à substituer. Après le défilement de toutes les configurations le circuit PAL2 s'arrête par un système de blockage.

c) Allumage des diodes laser

Les différentes configurations d'allumage des diodes laser sont enregistrées sur le circuit EPROM. Les adresses des différentes configurations sont générées par un système de compteurs binaires dont l'horloge est le signal Φ_4 du circuit PAL2. Le signal de validation sortie est donné par une combinaison logique des signaux du PAL2 pour synchroniser cet allumage avec les séquences générées. La sortie des données est connectée à des circuits tampons permettant l'adaptation du courant pour une commande directe soit des diodes laser ou soit au travers d'une carte spéciale.

d) Système de blockage

Ce système permet (circuits n° 2,12,13,15,16) d'arrêter le défilement des séquences du circuit PAL2 au bout d'un seul cycle ou des 128 cycles, et de déclencher la lecture électronique des résultats par le circuit PAL1.

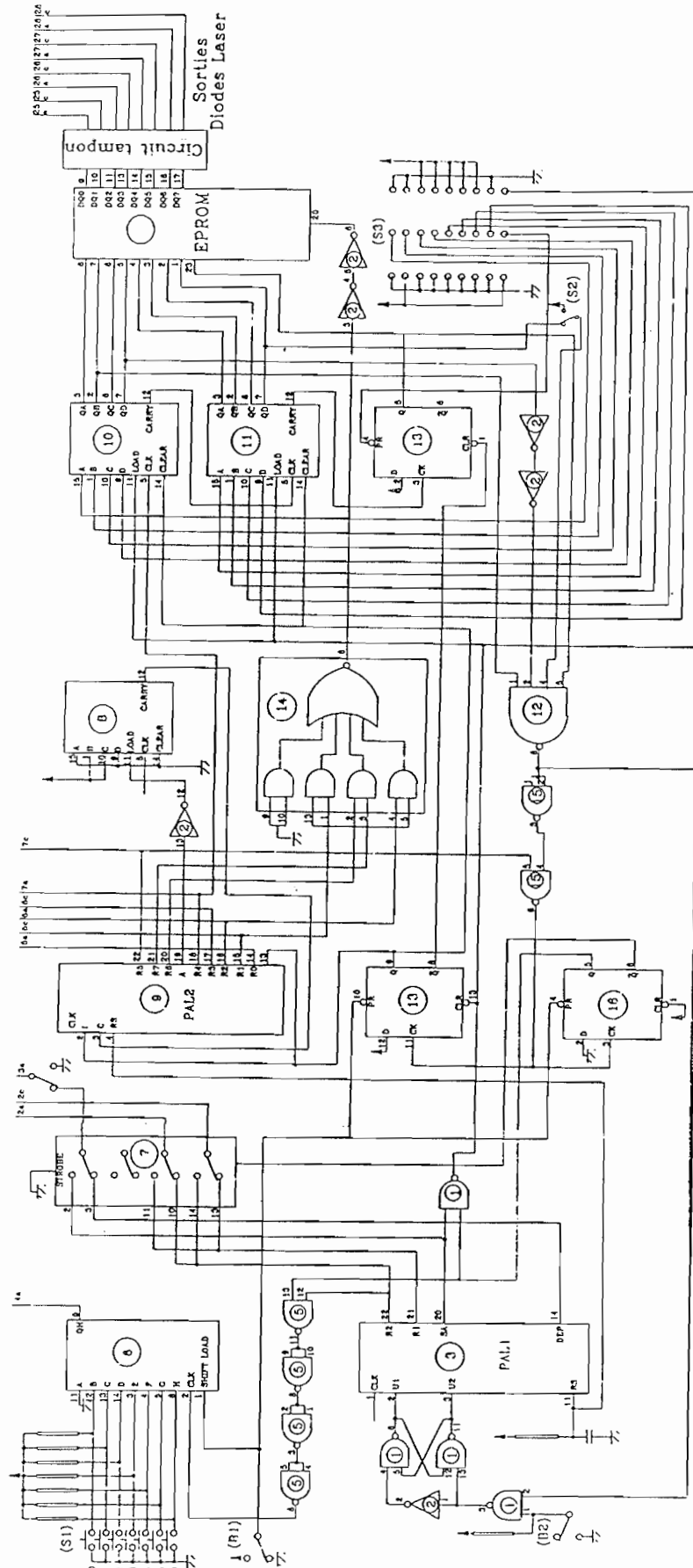


Fig. VIII-3 : Schéma du circuit de commande.

Les interrupteurs "S2" permettent le choix entre un seul cycle ou 128 cycles. Les connecteurs "S3" déterminent le choix de l'adresse de départ pour la lecture du contenu de l'EPROM, ce qui offre le choix de faire une seule opération de reconnaissance et de substitution par une configuration choisie.

VIII-2 DIODE LASER D'ECLAIRAGE

VIII-2-1 Caractéristiques de la diode laser à utiliser

L'éclairage des photodiodes se fait par une matrice de diodes laser modulables en intensité selon les séquences de commande présentées dans le paragraphe précédent. Le choix des diodes laser est dû essentiellement à leur puissance et à leur facilité de modulation en intensité par simple modulation électrique de leur courant.

Une certaine condition est imposée en ce qui concerne la longueur d'onde de ces diodes lasers. En fait, dans le projet final de l'automate, il est envisagé de poser un circuit de modulateurs opto-électronique, à base de puits quantiques multiples GaAs/GaAlAs, au dessus du circuit Silicium. Le circuit AsGa (plus particulièrement le matériau de substrat) doit être transparent pour les signaux d'éclairage des photodiodes, pour cela, il faut que la longueur d'onde d'éclairage des photodiodes λ_{ph} soit supérieure à celle correspondant à l'énergie de bande interdite de l'AsGa située à $\lambda_{AsGa} = 0.87 \mu m$.

D'autre part, cette longueur d'onde λ_{ph} doit être absorbée par le silicium pour générer le photo-courant. Par conséquent, cette longueur d'onde doit être inférieure à celle correspondante à la bande interdite du Silicium située à $\lambda_{Si} = 1.1 \mu m$.

Il en résulte que la longueur d'onde d'éclairage des photodiodes doit vérifier la double condition :

$$0.85 \mu m < \lambda_{ph} < 1.1 \mu m$$

La puissance de la diode laser et la durée d'éclairage sont déterminées par la sensibilité des photodiodes (voir Chapitre VII).

VIII-2-2 Diode laser utilisée

La diode laser utilisée dans le montage est fabriquée en collaboration entre le laboratoire LAAS de Toulouse et le CNET Bagneux. Elle est à puits quantiques multiples GaInAs/GaAs, et présente les caractéristiques suivantes :

Longueur de cavité : 500 μm
 Courant de seuil : # 8 mA

Longueur d'onde : 0.995-1 μm
Fonctionnement en régime impulsionnel :
Largeur d'impulsion : 400 ns
taux de répétition : 1 kHz
Courant maximal : 100-150 mA

La courbe de la puissance optique en fonction du courant est présentée dans la figure suivante :

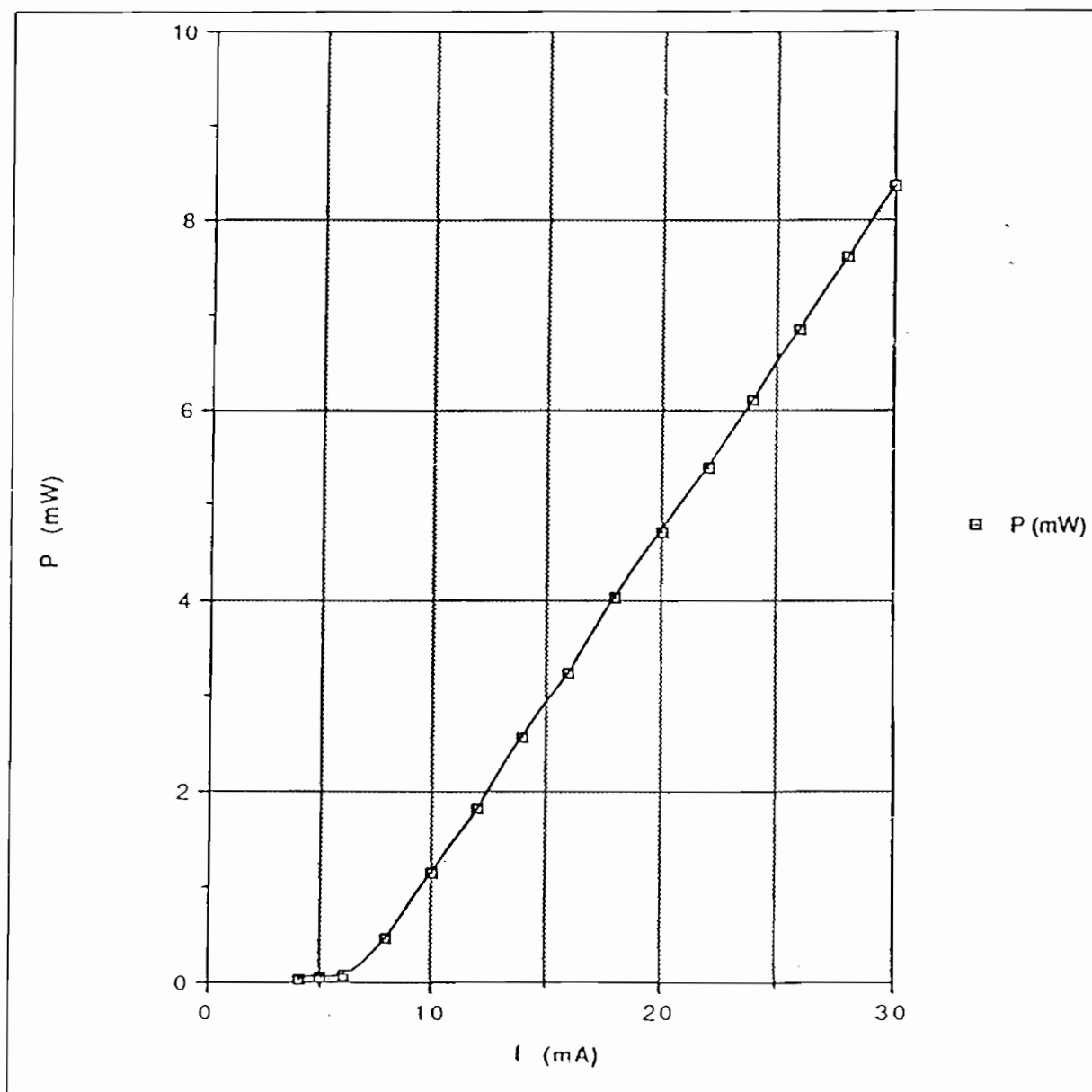


Fig. VIII-4 : Caractéristique de la diode laser.

VIII-2-3 Carte de modulation de la diode laser

Compte-tenu des précautions à prendre pour la modulation du courant de la diode laser utilisée, une carte de modulation est utilisée pour contrôler la modulation, la durée et la stabilité du courant de cette diode.

Cette carte reçoit la commande du début d'impulsion sous forme d'un signal TTL. À l'aide d'un monostable réglable, nous pouvons déterminer la durée de l'impulsion du courant qui varie entre 0 et 200 ns. Le courant de la diode est généré par un système de deux transistors (T1 et T2) (Fig. VIII-5) à grande bande passante (800 MHz) qui assure un temps de montée du courant d'environ 2 ns. Le niveau du courant est réglable à l'aide d'un résistor entre 0 et 120 mA.

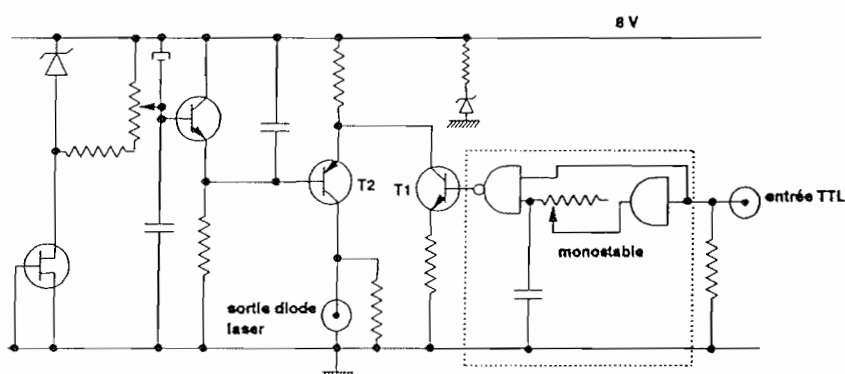


Fig. VIII-5 : Schéma du circuit de modulation de la diode laser.

VIII-3 MODULATEURS OPTO-ELECTRONIQUE

Nous utilisons pour la lecture optique du résultats des modulateurs opto-électronique travaillant en reflexion fabriqués au laboratoire Central de Recherche de Thomson C.S.F.

VIII-3-1 Principe de fonctionnement

Chaque pixel de ces modulateur a la structure suivante (voir Fig. VIII-6) :

- Une couche GaAs de 10 nm d'épaisseur suivit d'une couche $\text{Ga}_{0.3}\text{Al}_{0.3}\text{As}$ de 993.3 nm d'épaisseur.
- Une zone de puits quantiques multiples composée de 75 périodes chacune de deux couches : GaAs (épaisseur 10 nm) et $\text{Ga}_{0.7}\text{Al}_{0.3}\text{As}$ (épaisseur 10 nm).
- Une zone de réflecteur de Bragg (en dessous de la zone précédente), composée de 12 période de deux couche : $\text{Ga}_{0.9}\text{Al}_{0.1}\text{As}$ (épaisseur 60.6 nm) et AlAs (épaisseur 72.1 nm) avec maximum de réflexion à une longueur d'onde d'environ 860 nm.

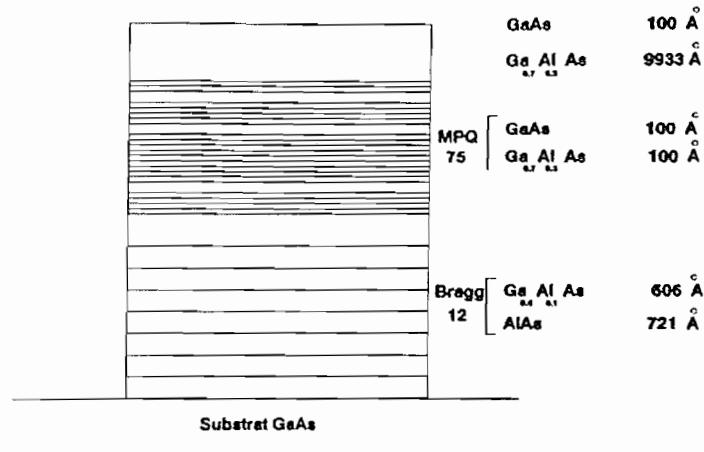


Fig VIII-6 : Structure d'un pixel des modulateurs.

En appliquant un champ électrique, La réflexion du réflecteur de Bragg est très peu modulée. Par contre l'absorption, à la longueur d'onde 860 nm est fortement modulée par effet Stark [51].

Donc, par modulation de la tension appliquée aux bornes de cette structure, nous modulons l'absorption des puits quantiques multiples, et par la suite l'intensité lumineuse réfléchie par le système.

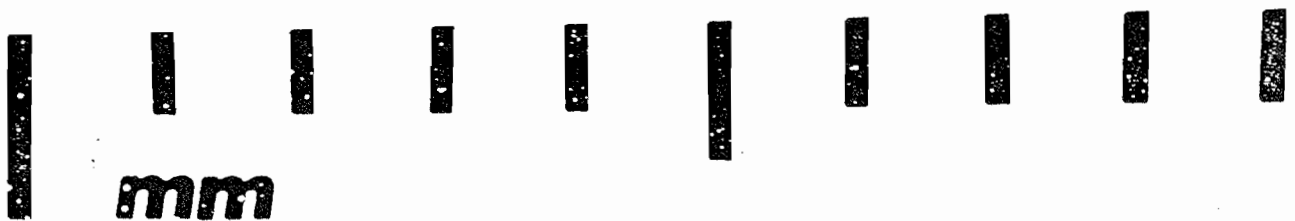
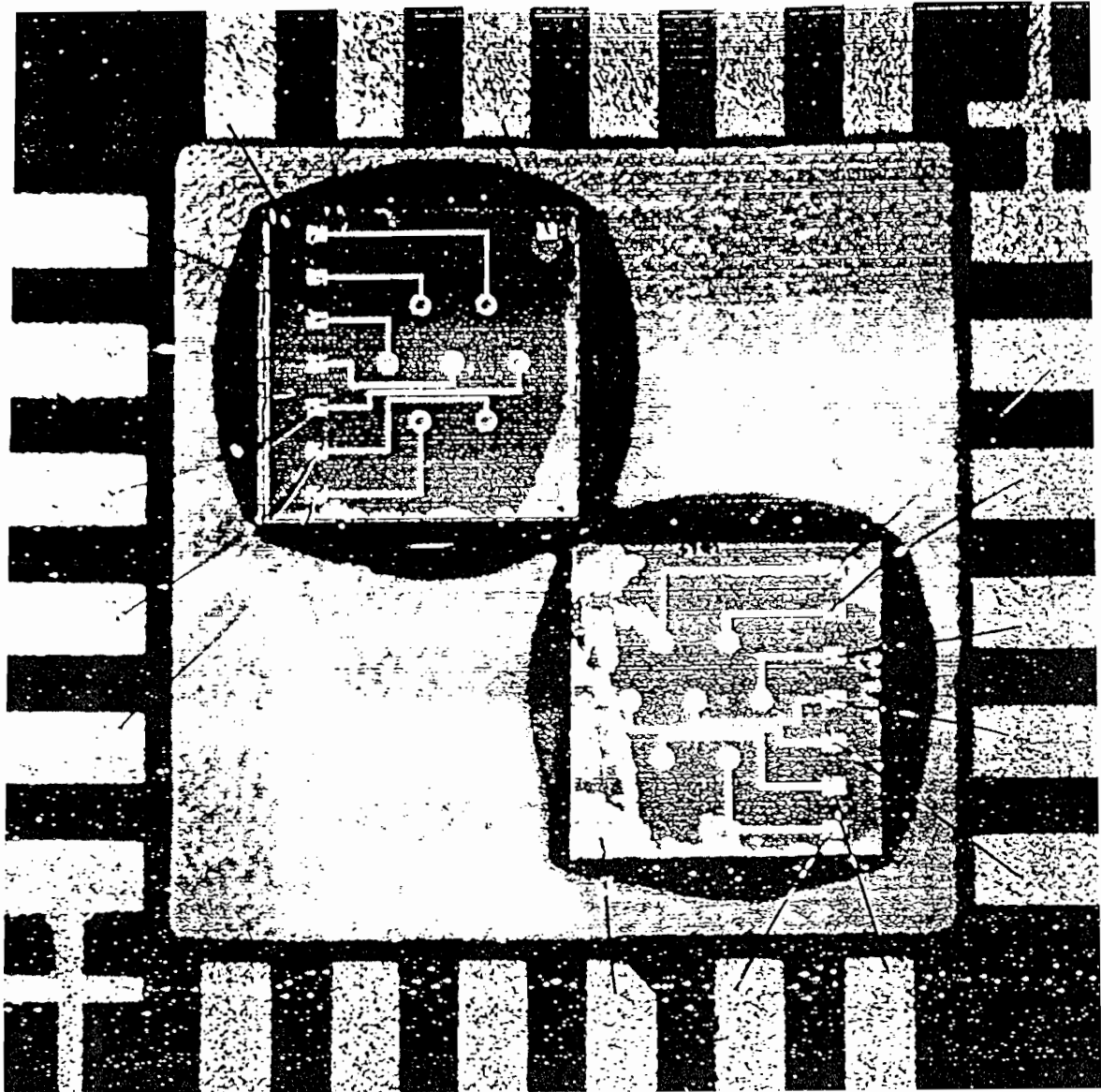
VIII-3-2 Caractéristiques du circuit utilisé

Les modulateurs utilisés se présentent sous forme d'un circuit comportant deux matrices chacune constituée de 7 pixels sous forme hexagonale (Fig. VIII-7). Le diamètre de chaque pixel est de 200 μm .

Les caractéristiques de ces modulateurs sont présentées par la suite (D'après les tests effectués à Thomson) :

Longueur d'onde d'éclairage: $\lambda = 860 \text{ nm}$
 Tension de modulation : V entre 0 et 20 V
 Contraste en réflexion : > 90%

La figure VIII-8 montre l'effet de la modulation de la tension sur l'intensité réfléchie en fonction de la longueur d'onde. La figure VIII-9 montre le contraste obtenu par réflexion en fonction de la longueur d'onde : le meilleur contraste est à $\lambda = 860 \text{ nm}$.



THOMSON-CSF

LABORATOIRE CENTRAL DE RECHERCHES

Fig VIII-7 : Circuit des modulateurs à puits quantiques multiples

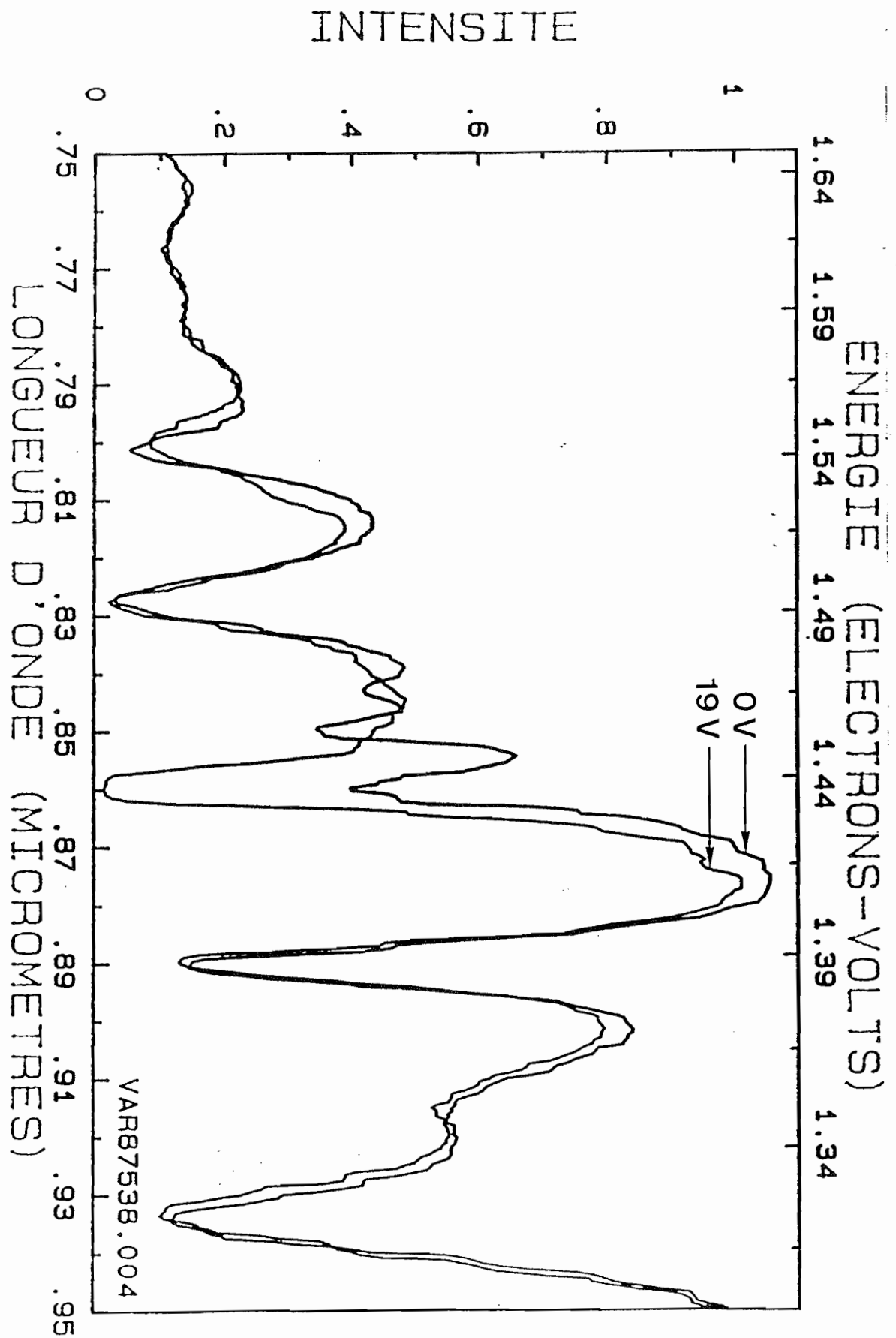


Fig. VIII-8 : Effet de la modulation de la tension sur l'intensité réfléchie.

05-28-1991-VAR875/3/8-PAS=.5nM-FENTES=2.4nM-CUBE-RG645-1*40V-19volts
 05-28-1991-VAR875/3/8-PAS=.5nM-FENTES=2.4nM-CUBE-RG645-1*40V-0volt-T: 22xC

ENERGIE (ELECTRONS-VOLTS)

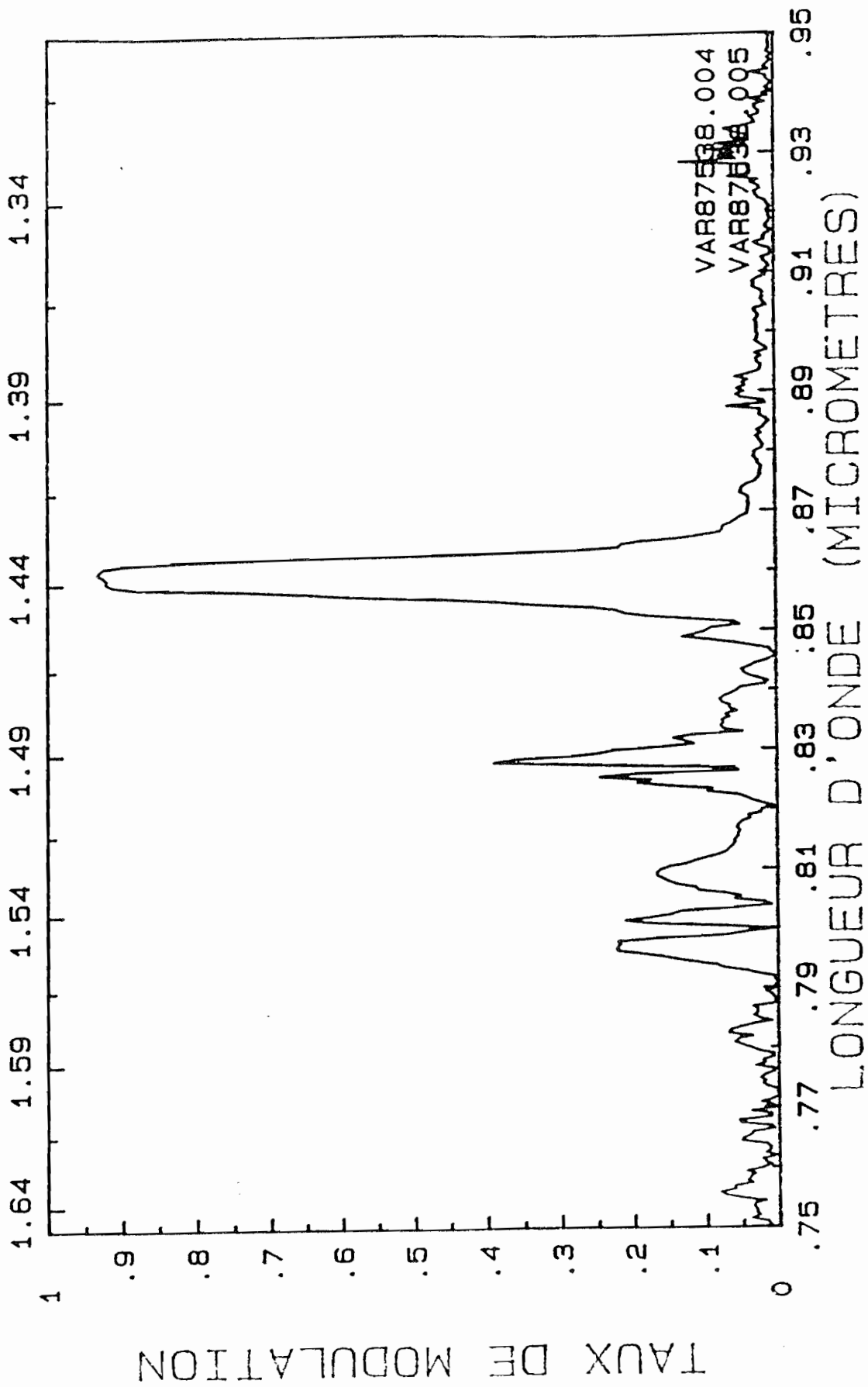


Fig. VIII-9 : Le contraste de modulation.

VIII-3-4 Mise en fonction des modulateurs

Nous avons mis en fonction les modulateurs de Thomson en éclairant un pixel par la lumière blanche à l'aide d'un microscope. La lumière réfléchie est projetée sur une caméra CCD devant laquelle est posé un filtre interférentiel à 860 nm. En modulant la tension de ce pixel entre 0 et 20 V, nous avons remarqué la modulation de l'intensité réfléchie par le pixel sur un écran moniteur vidéo avec un très bon contraste [Exposition Physique 91].

VIII-3-4 Utilisation dans le prototype

Le circuit des modulateurs projette optiquement le contenu de la mémoire de sortie du circuit électronique sur une caméra CCD. La connexion entre le circuit des modulateurs et le circuit Silicium se fait extérieurement. Les modulateurs sont éclairés par une source de longueur d'onde $\lambda = 0.860 \mu\text{m}$, l'onde réfléchie sur les modulateurs est projetée sur la caméra à l'aide d'une lame séparatrice.

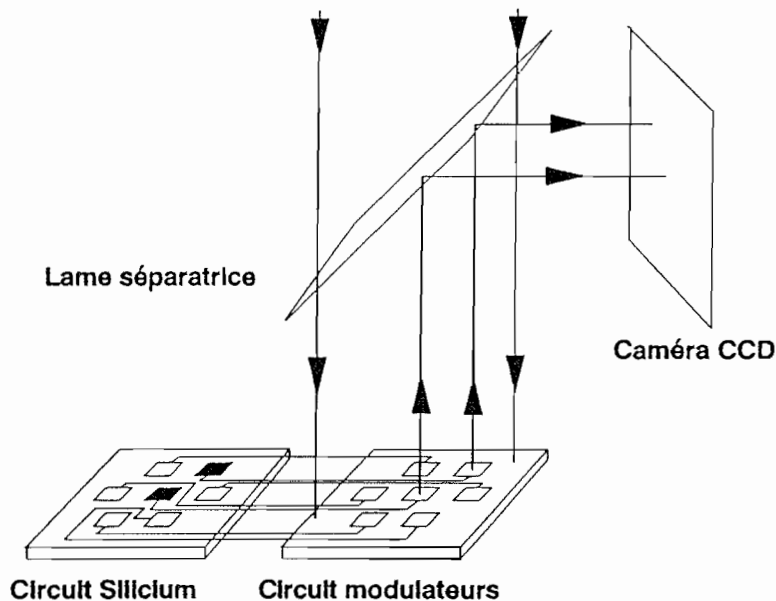


Fig. VIII-10 : Utilisation du circuit modulateurs opto-électronique.

Dans les étapes avancées du projet, les modulateurs seront pausés directement sur le circuit Silicium, la connexion électrique entre les deux circuits se fait avec des billes d'Indium.

VIII-4 MONTAGE EXPERIMENTAL

Ce montage consiste à éclairer une seule photodiode du circuit électronique et vérifier la reconnaissance et la substitution avec une configuration à un seul bit [Salon de Physique 91].

Le circuit est monté sur un système à trois déplacements selon x, y et z. L'alignement du circuit par rapport à la diode laser se fait à l'aide d'une caméra CCD et un cube séparateur de polarisation (voir Fig. VIII-11).

L'éclairage de la diode laser se fait à partir du circuit de commande qui envoie une impulsion à un générateur de fonction qui, durant cette impulsion, commande le courant de la diode laser à une cadence de 1 MHz. c'est ainsi que nous envoyons au circuit l'énergie lumineuse suffisante pour basculer les photodiodes.

Le résultat du traitement est visualisé sur un système de 7 diodes électroluminescentes qui sont connectées à la mémoire de sortie du circuit à la place des modulateurs opto-électronique.

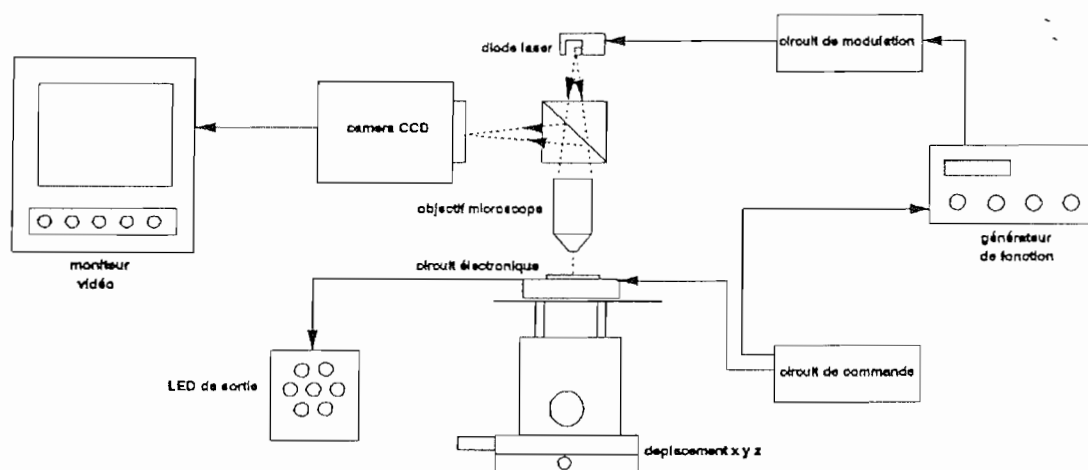


Fig. VIII-11 : Montage de démonstration.

Une étude d'un montage utilisant des fibres optiques pour éclairer les 7 photodiodes du circuit est présentée dans l'annexe C. L'éclairage des 7 photodiodes permettra un test dans un montage complet.

CHAPITRE IX

CARACTERISATION ET ARCHITECTURE DES AUTOMATES CELLULAIRES

CHAPITRE IX

CARACTERISATION ET ARCHITECTURE DES AUTOMATES CELLULAIRES

Ce chapitre présente un bilan général de l'étude concernant l'automate cellulaire en projet. Nous étudierons les caractéristiques de cet automate "Gaz sur Réseau". Ensuite, nous comparons les architectures tout électroniques et les architectures opto-électroniques.

IX-1 CARACTERISATION DE L'AUTOMATE EN PROJET

Nous présentons par la suite les caractéristiques essentielles de l'automate cellulaire "Gaz sur Réseau" en tenant compte de la réalisation actuelle et des perspectives.

Ces caractéristiques concernent l'encombrement de l'automate, la fréquence de travail, la puissance électrique consommée et le nombre de connexions possibles à réaliser.

IX-1-1 Surface occupée sur le circuit électronique

A l'état actuel du circuit, l'implantation d'une cellule occupe une surface de 1 mm^2 (sans le registre à décalage qui ne serait pas reproduit dans la version finale). Ceci implique que sur une puce de 1 cm^2 , il est possible d'implanter de 100 cellules.

Il est important de signaler que la technologie utilisée pour la réalisation de ce circuit est de $\Lambda = 2 \text{ }\mu\text{m}$ (La largeur minimale des différentes zones dopées ou métalliques). A la réalisation de ce circuit, nous avons utilisé une largeur minimale des zones égale à $4 \text{ }\mu\text{m}$ pour vérifier la possibilité de fonctionnement sous 10 V, le test a montré l'inutilité de cette largeur. La réalisation du même circuit avec les dimensions permises par la technologie, ainsi que la réduction de la surface de la photodiode à $25 * 25 \text{ }\mu\text{m}^2$, permet de gagner un facteur 2 sur les dimensions du circuit, donc un facteur 4 sur la surface occupée ce qui permet d'implanter $4 * 10^2$ ($20 * 20$ cellules) cellules sur une puce de 1 cm^2 . Ce nombre est largement en dessous du nombre nécessaire pour le traitement de l'algorithme "Gaz sur Réseau" estimé dans le chapitre III à $1.3 * 10^5$ cellules. Ceci revient à utiliser une puce 3000 fois plus grande, ce qui est pratiquement impossible, ou bien utiliser 3000 puces avec leur système de connexion électronique et optique ce qui occupe un volume considérable.

Ainsi nous voyons que le parallélisme complet de l'algorithme par la méthode proposée demande un encombrement considérable. Une solution plus compacte consiste à utiliser un

certain nombre de circuits ($10 * 10$ par exemple), et à traiter sur ce système les différentes parties du gaz d'une façon séquentielle.

Remarque : Il existe des technologies microélectroniques plus fines (en dessous de $1 \mu\text{m}$), ceci permet de gagner un certain facteur (4 pour une technologie de $1 \mu\text{m}$) mais reste du même ordre de grandeur du nombre de circuits à utiliser.

IX-1-2 Fréquence de fonctionnement

Pour estimer la fréquence de fonctionnement du circuit, il faut prendre compte du temps de basculement des transistors et du temps de charge des capacités contenues dans le circuit (la capacité de l'unité de reconnaissance et la capacité interne des photodiodes). La décharge photo-électrique des photodiodes dépend de la puissance optique des diodes laser utilisées.

IX-1-2-a temps de basculement des transistors

Nous avons estimé le temps de montée et de descente des transistors du circuit à 50 ns (voir chapitre VI). Le temps effectif peut être plus petit compte-tenu des capacités introduites par les points du test. La fréquence de fonctionnement du circuit est évaluée d'environ 100 MHz.

IX-1-2-b Temps de charge des capacités

Ce temps dépend de la valeur de ces capacités et du courant débité par les transistors de charge.

La capacité de l'unité de reconnaissance et celle des photodiodes sont constituées de la capacité de diffusion d'une région dopée N^+ sur le substrat P. Cette capacité est estimée à 1.25 nF/cm^2 . Compte-tenu des dimensions de ces différentes capacités leurs valeurs sont estimées comme suit :

Unité de reconnaissance	$150 * 30 = 4500 \mu\text{m}^2$	$C_c = 56.25 \text{ fF}$
Photodiode	$50 * 50 = 2500 \mu\text{m}^2$	$C_p = 31.25 \text{ fF}$

Le courant délivré par le transistor de charge (Type N) est donné par la formule :

$$I = \beta \frac{V_{DD}^2}{2}$$

où β est le facteur de transconductance donné comme suit [44,45,52] :

$$\beta = \frac{\mu\epsilon}{t_{ox}} \left(\frac{W}{L} \right)$$

- μ : est la mobilité des électrons estimée à $500 \text{ cm}^2\text{V}^{-1}\text{s}^{-1}$
 ϵ : La permittivité relative du Silicium estimée à 10^{-12} F/cm
 t_{ox} : L'épaisseur de la couche d'oxyde de la porte du transistor estimée à 50 nm .
 W : La largeur du transistor $6 \mu\text{m}$.
 L : La longueur du transistor $6 \mu\text{m}$.

Le courant maximal débité par le transistor de charge est donc estimé à :

$$I_{\text{max}} = 10 \mu\text{A}$$

d'où le temps nécessaire pour charger la capacité C de l'unité de reconnaissance (la plus grande) :

$$\tau = \frac{C_c}{I_{\text{max}}} \cdot V_{\text{DD}}$$

estimé à 30 ns , ce qui est de même ordre de grandeur que le temps de basculement des transistors.

IX-1-2-c temps d'exécution d'une itération

Compte-tenu des temps de fonctionnement calculés précédemment, et du diagramme des implusions donné dans la figure VIII-2 du chapitre VIII pour un cycle de reconnaissance-substitution, le temps pour exécuter un cycle est estimé à 500 ns en travaillant avec une horloge de 20 MHz de fréquence (Une majoration du temps de réponse des différentes parties). Une itération est l'exécution de 128 cycles, ce qui demande un temps égal à $64 \mu\text{s}$, ce qui revient à exécuter environ 15000 itérations par seconde. Il est évident que cette fréquence de fonctionnement est indépendante du nombre de cellules utilisées.

IX-1-3 Puissance dissipée

Nous distinguons entre la puissance dissipée dans le basculement des transistors dans le circuit électronique VLSI, et la puissance dissipée par voie optique pour la distribution des instructions.

IX-1-3-a Puissance électronique

Cette puissance est directement liée à la consommation en courant des transistors du circuit électronique. Nous distinguons entre la puissance statique et la puissance dynamique de ces transistors.

- La puissance statique est celle dissipée par le transistor hors période de commutation. En technologie CMOS le courant théorique débité par le transistor est nul en temps de repos quelque soit le niveau logique du transistor (ceci grâce à la combinaison des transistors N et P utilisé par les inverseurs, unité de base des points mémoires utilisés). Toutefois, il reste le courant de fuite dû au courant inverse dans les jonctions p-n qui est donné par la relation (chapitre VII) :

$$i_0 = i_s \{ \exp(qV/kT) - 1 \}$$

Ce courant est souvent estimé entre 0.1 et 0.5 nA [45]. La puissance statique dissipée est donnée par :

$$P_s = \sum_{\text{transistors}} i_0 \cdot V_{DD}$$

- La puissance dynamique dissipée dans les transistors correspond au courant débité dans les transistors lors de la commutation entre les niveaux logiques. Cette puissance est donnée par la relation [44] :

$$P_d = \sum_{\text{transistors}} \frac{C_L \cdot V_{DD}^2}{t_p}$$

où C_L est la capacité de charge de chaque transistor, $t_p = 1/f_p$ où f_p est la fréquence de travail.

La capacité de charge C_L est composée de deux capacités :

- C_g la capacité entre le "grille" du transistor et le substrat. Cette capacité dépend de la nature du transistor et est évaluée à 35 fF pour les transistors N et 68 fF pour les transistors P. Le circuit contient pratiquement autant de transistors N et P, ce qui permet de prendre une moyenne de cette capacité de 50 fF sans commettre une grande erreur.

- C_s la capacité de sortie qui contient tous les fils de connexion, les photodiodes ... etc. Cette capacité est évaluée à 50 fF.

La capacité de charge aura la valeur moyenne $C_L = 100$ fF par transistor.

La puissance totale dissipée par un transistor :

$$P_t = P_s + P_d$$

En considérant une tension d'alimentation $V_{DD} = 5$ V, et une fréquence de travail $f_p = 2$ MHz (La plupart des transistors commutent une fois par cycle de reconnaissance-substitution), nous trouvons la puissance électronique dissipée :

$$P_t = (25 \cdot 10^{-10} + 5 \cdot 10^{-6})N \cdot 5 \cdot N \mu W$$

où N est le nombre de transistors.

Dans notre cas le nombre de transistors dans le cœur du circuit est égal à 133, donc la puissance dissipée dans le cœur du circuit est égale à 665 μW .

IX-1-3-b Puissance optique

La puissance optique dépend de la capacité de chaque photodiode et du temps demandé pour sa décharge. Nous commençons par évaluer la puissance optique nécessaire pour chaque photodiode pour aboutir à la puissance optique nécessaire pour l'automate global.

La puissance optique nécessaire pour décharger une photodiode menée jusqu'à la tension V_{DD} pendant un temps τ est donnée par la relation :

$$P_{opt} = \frac{C_p \cdot V_{DD}}{S \cdot \tau}$$

où :

- C_p est la capacité interne de la photodiode estimé à 30 fF.
- S est la sensibilité de la photodiode estimée pour la longueur d'onde $\lambda = 1 \mu m$ (longueur d'onde des diodes laser utilisées (voir ch.VIII)) à $S = 0.15$ A/W.

Pour un temps de décharge compatible avec le reste du circuit de 50 ns, la puissance optique nécessaire est de :

$$P_{opt} = 20 \mu W$$

Cette puissance doit être disponible au niveau de la photodiode. Pour la puissance à la sortie de la diode laser il faut multiplier cette puissance par un facteur allant jusqu'à 10 compte-tenu des pertes dans les dispositifs optique (divergence du faisceau, réflexions partielles sur les dioptries, efficacité de diffraction de l'hologramme ...etc) ce qui nous amène à une puissance à la sortie de la diode laser d'environ 200 μW . La diode laser est allumée pendant 50 ns et dans le cas extrême 3 fois par cycle de reconnaissance-substitution. Par conséquent, la puissance optique moyenne par diode laser est de 60 μW . La conversion électrique-optique dans les diodes laser se fait avec un rendement maximal de 30%, ce qui implique que la puissance totale dissipée par diode laser est de 200 μW . Nous avons besoin de 7 diodes laser pour le fonctionnement du circuit, donc la puissance totale dissipée par voie optique est de 1400 μW .

IX-1-3-c Evaluation de la puissance totale nécessaire au fonctionnement

La puissance électronique et optique dissipée par cellule est d'environ 2 mW. Cette puissance est proportionnelle au nombre de cellules utilisées, par exemple en implantant 10^2 cellules sur une puce de 1 cm^2 la puissance totale nécessaire pour le fonctionnement d'un tel circuit est de 0.2 W.

IX-1-4 Nombre de connexions (entrée-sortie)

La notion de connexion ici est restreinte à l'échange d'informations entre les processeurs du circuit et l'extérieur soit à l'entrée soit à la sortie. Dans notre automate seules les connexions optiques sont considérées car les autres entrées électroniques ne servent qu'aux signaux de séquençement sans échange d'informations.

Remarque : Dans la version préliminaire du circuit électronique présentée dans le chapitre VI, deux signaux électroniques sont considérées comme connexions : "ENTREE" et "SORTIE" du registre à décalage. Nous n'en tenons pas en compte car ils ne font pas partie de la conception finale de l'automate.

Si le nombre de connexions est défini dans le cadre de l'échange d'informations, le critère d'évaluation est le nombre d'informations échangées par unité de temps. Nous utilisons comme critère le "GIBP : Gate Interconnect Bandwidth Product" introduit par Guifoyle [53] qui est le produit entre le nombre des ports d'interconnexion et leur fréquence d'échange d'information.

Dans notre cas, le circuit dispose de 7 ports d'interconnexions qui sont les photodiodes. Chaque port reçoit 3 informations (d'un bit) par cycle de reconnaissance-substitution de durée 500 ns, ce qui donne une fréquence d'échange d'informations de 6 MHz, et la valeur de GIBP est égale à $42 \cdot 10^6 \text{ bits.s}^{-1}$. Nous avons vu que l'implantation de $4 \cdot 10^2$ cellules de même architecture est possible sur une puce de 1 cm^2 ce qui donne :

$$\text{GIBP} = 1.68 \cdot 10^{10} \text{ bits.s}^{-1}$$

Pour un circuit électronique à connexions purement électroniques, le nombre de connexions dépend des plots de connexion utilisés. Si nous prenons, comme exemple, les plots utilisés dans notre circuit, ils ont une largeur de $171 \text{ }\mu\text{m}$. Le nombre de plots maximal que nous pouvons implanter sur une puce de 1 cm^2 est de 233. pour la même fréquence d'échange d'information, la valeur du GIBP est égale à

$$\text{GIBP} = 4.66 \cdot 10^8 \text{ bits.s}^{-1}$$

L'avantage de l'optique dans ce cas revient à ce que la connexion se fait sur la surface et pas sur le périmètre du circuit. Cet avantage prend plus de valeur en diminuant le pas de la technologie Λ , car le GIBP croît dans le cas opto-électronique comme $1/\Lambda^2$, tandis que dans le cas électronique, il croît plus faiblement que $1/\Lambda$ à cause du problème de soudure sur les plots électroniques du circuit qui conservent pratiquement les mêmes dimensions.

IX-1-5 Capacité de calcul

Nous traitons ici la capacité de calcul d'une puce contenant $4 \cdot 10^2$ cellules de traitement fonctionnant à une fréquence (cycle de reconnaissance-substitution) de 2 MHz. Pour cette puce la capacité de calcul est de $8 \cdot 10^8$ sites par seconde. Compte-tenu que chaque cellule exécute 3 opérations par cycle, la vitesse de calcul est de $2.4 \cdot 10^9$ opérations par seconde. Cette vitesse de calcul est largement au dessus des puces existantes actuellement.

IX-2 COMPARAISON ENTRE ARCHITECTURES OPTOELECTRONIQUES ET TOUT ELECTRONIQUES

D'après l'analyse faite précédemment sur l'automate cellulaire "Gaz sur Réseau", il est apparu que l'avantage de l'optique est dû essentiellement à sa capacité à distribuer les instructions en parallèle sur un grand nombre de processeurs élémentaires. Nous présentons par la suite une comparaison entre une architecture toute électronique et une architecture

hybride optoélectronique pour la distribution d'un signal sur un certain nombre de processeurs élémentaires [54].

Nous considérons le cas de la distribution séquentielle (bit par bit) sur un ensemble de $N \times N$ processeurs élémentaires implantés sur une puce électronique CMOS en technologie $\Lambda = 1 \mu\text{m}$.

Dans la solution tout électronique, un plot d'entrée fournit un signal qui est divisé en N lignes pour alimenter les N^2 processeurs élémentaires (P.E.). Dans la version optoélectronique, chaque P.E. est muni d'une photodiode et le plot d'entrée est remplacé par une diode laser qui éclaire un illuminateur de tableau qui forme l'image de la diode laser sur les N^2 photodiodes (voir fig. XI-1).

Cette comparaison concerne essentiellement la fréquence, l'énergie et la surface occupée.

Nous considérons $N \times N = 100 \times 100$ P.E., chacun de $d \times d = 100 \times 100 \mu\text{m}^2$ de surface.

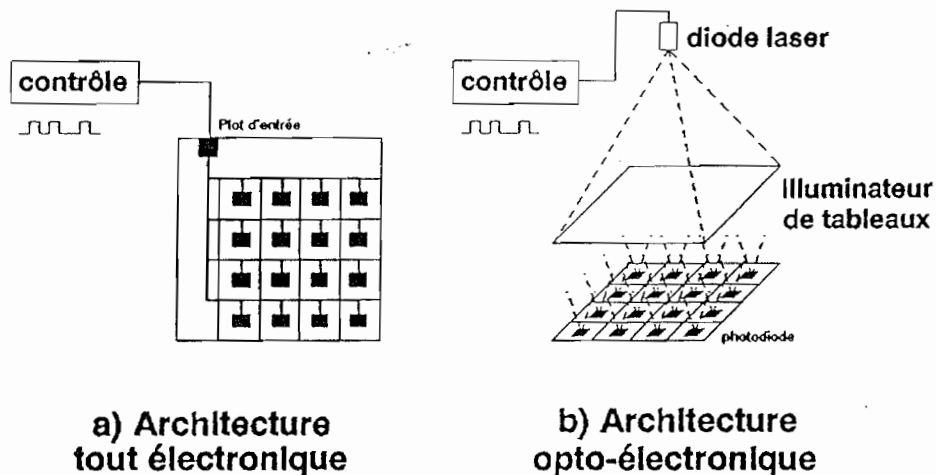


Fig. IX-1 : Distribution des instructions.

IX-2-1 Fréquence et consommation d'énergie

La capacité moyenne par processeur élémentaire constitue un facteur essentiel pour déterminer la fréquence de travail aussi bien que la consommation de l'énergie.

Pour la solution tout électronique, les différentes capacités à prendre en compte sont :

- La capacité du plot d'entrée, qui de l'ordre de 400 fF à répartir entre les N^2 P.E.

- La capacité de chacune des N lignes de longueur $N.d$. Pour une longueur de $10000\text{ }\mu\text{m}$. Elle est de l'ordre de 700 fF , à répartir entre les N P.E.
- la capacité de grille du transistor d'arrivée de chaque P.E. qui est de l'ordre de 20 fF .

La capacité moyenne pour chaque P.E. est de 30 fF .

Dans le cas optoélectronique, les lignes de connexion et le plot d'entrée disparaissent et sont remplacées par la capacité de la photodiode. Cette capacité dépend de la surface de la photodiode. A cause d'un problème d'alignement optique la surface d'une photodiode ne peut pas être en dessous de $10\text{ }\mu\text{m}^2$; ce qui correspond à une capacité de 10 fF . La sortie de la photodiode est connectée à un transistor dont la capacité de grille est de 20 fF , ce qui ramène la capacité moyenne par P.E. à 30 fF .

Par conséquent, les deux solutions sont pratiquement équivalentes du point de vue fréquence de travail et consommation de l'énergie au niveau de la puce électronique. Le taux de conversion électrique-optique pour les diodes laser qui est au maximum de 30% défavorise la solution opto-électronique.

IX-2-2 Surface de Silicium

Dans la solution tout électronique, la surface de la ligne de connexion concernant chaque P.E. est de $100\text{ }\mu\text{m}^2$. Tandis que pour la solution optoélectronique, cette surface est de $10\text{ }\mu\text{m}^2$. Par conséquent, pour la technologie considérée $\Lambda = 1\text{ }\mu\text{m}$, la solution optoélectronique est favorable par rapport à la solution tout électronique. Néanmoins, il faut noter que cette surface décroît comme Λ^{-2} dans le cas tout électronique, par contre, elle est constant dans le cas optoélectronique.

IX-2-3 Autres considérations

La solution optoélectronique est plus compliquée à la réalisation à cause de l'alignement optique. Par contre elle est plus tolérante pour les défauts locaux qui, si elles affectent un fil de connexion, mettent en défaut toute la ligne des P.E. alimentée par ce fil.

La discussion précédente était portée sur la distribution d'un seul signal. Dans le cas de la distribution de plusieurs signaux (7 dans notre automate), l'architecture toute électronique se confronte à un problème de croisement des différents fils de connexion ce qui demande plusieurs niveaux de couches métalliques pour une grande surface consacrée à ce croisement. Dans ce cas l'architecture optoélectronique semble plus favorable car le croisement des faisceaux optiques n'aboutit pas à un échange d'informations.

IX-3 CONCLUSION

Le choix entre architectures toute électronique ou opto-électronique ne semble pas évident dans le cas général comme le montre la comparaison précédente. Ce choix dépend de chaque application.

Pour l'automate cellulaire "Gaz sur Réseau", le choix d'une distribution optique d'instructions est favorable à la solution toute électronique, cet avantage est dû à la distribution en parallèle des mots de 7 bits.

L'implantation de l'algorithme avec une architecture optoélectronique semble prometteur, malgré le nombre modeste de cellules implantées sur une puce, du point de vue capacité et rapidité de calcul qui est nettement supérieure aux puces connues actuellement. Le point fort de cet automate réside dans son parallélisme et le nombre d'entrée-sortie réalisable.

CONCLUSION GENERALE

CONCLUSION GENERALE

Nous avons présenté dans ce mémoire l'intérêt que porte l'optique pour la réalisation des architectures parallèles. En choisissant l'algorithme "Gaz sur Réseau" comme exemple de démonstration, nous avons démontré la possibilité de son traitement sur le modèle des automates cellulaires.

Après avoir présenté le schéma général de l'automate cellulaire optoélectronique, nous avons étudié et réalisé (ou utilisé) les différents composants optiques, électroniques et optoélectronique de l'automate cellulaire. Ces composants sont les suivant :

- Hologramme illuminateur de tableaux à effet Talbot : Nous avons étudié l'effet Talbot et la possibilité de l'utiliser pour réaliser un hologramme illuminateur de tableaux. La réalisation de cet hologramme est faite sur gélatine bichromatée avec une bonne efficacité de diffraction. Nous avons mené une étude détaillée des aberrations chromatiques et nous avons proposé une méthode simple pour les corriger. Les simulations sur ordinateur confirme l'efficacité de la méthode proposée.

- Circuit microélectronique : En collaboration avec l'I.E.F., nous avons conçu, réalisé et testé un circuit (VLSI) contenant une seule cellule de traitement. Une étude particulière a été faite pour une connexion directe entre ce circuit et les modulateurs optoélectroniques pour la sortie parallèle des résultats.

- Modulateurs optoélectronique : Ces modulateurs ont été conçus et fabriqués à Thomson L.C.R. Ils sont à base de puits quantiques multiples.

Ces différents composants de l'automate cellulaire optoélectronique sont séparées. Les caractéristiques et la fonctionnalité de chacun a été testé à part. Rassembler ces différentes parties dans un montage fonctionnel est nécessaire pour la continuité du projet.

Nous avons estimé les caractéristiques du montage final du projet et sa capacité de calcul qui est très prometteur et comparable à ceux des grandes machines.

Les applications de l'optique dans l'ordinateur, surtout pour les connexions, peut aussi bien être par la propagation dans l'espace libre, comme c'est le cas traité dans ce mémoire, que par le système planaire en utilisant l'optique guidée.

ANNEXES

Annexe A

Traitement des gélamines bichromatée

Les plaques utilisées sont des "Kodak 649 F" avec le traitement suivant :

I- Préparation des plaques

- Fixation F5 filtré 15 minutes pour enlever l'argent.
- Rinçage à l'eau distillée pendant 1 heure. On enlève l'antihalo en essuyant le côté verre avec un papier imbibé de méthanol.
- Deux bains de méthanol de 10 minutes chacun.
- Sécher à la tournette à 2000 tours par minute 20 secondes.

II- Bichromatage

- Préparer une solution de bichromate d'ammonium (5-10%) en diluant 20 grammes de ce bichromate dans 150 ml d'eau à une température supérieure à 40° et compléter la solution à 200 ml.
- Filtrage de la solution.
- Porter la solution à 40° dans un bain à eau chauffée.
- Ajouter une goutte de Kodak photoflo pour 100 ml de solution.
- Tremper la plaque dans la solution en agitant régulièrement.
- Tournette à 2000 tours par minute pendant 10 secondes.
- Essuyage des bords et du dessous.
- Tournette pendant 10 secondes.
- Essuyage des bords et du dessous à l'alcool.
- Séchage dans un four à 30° pendant une heure
- Conserver dans une boîte avec actigel.
- Utilisation après 12 heures.

III- Développement

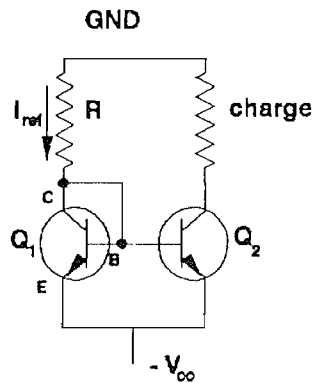
- Rinçage à l'eau courante pendant 8 heures.
- Rinçage à l'eau distillée pendant 8 minutes.
- Isopropanol 50% pendant 2 minutes.

- Isopropanol 90% pendant 2 minutes.
- Isopropanol 100% pendant une durée supérieure à 10 minutes.
- Sortir les plaques du bain d'alcool le plus uniformément possible, ou laisser évaporer l'alcool.

Annexe B

Miroir de courant

Le miroir de courant est réalisé par un transistor identique à celui dont nous voulons mesurer le courant du collecteur. La connexion des deux transistors est donnée dans le schéma suivant :



Le transistor monté en miroir de courant (Q_2) a son émetteur connecté à sa base.

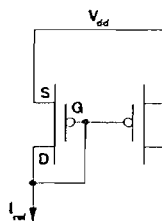
Les deux transistors étant identiques, ils ont le même coefficient β . D'autre part, ils sont sous la même tension V_{BE} , donc ils ont le même courant de base I_B . Nous pouvons écrire les équations suivantes :

$$I_{ref} = I_{C1} + 2I_B \quad \text{avec} \quad I_{C1} = \beta \cdot I_B \quad \text{d'où} \quad I_{ref} = (\beta + 2) \cdot I_B$$

$$I_{C2} = \beta I_B = \beta \cdot (I_{ref} / (\beta + 2)) \approx I_{ref} \quad \text{car} \quad \beta \gg 2$$

Ainsi I_{ref} est presque identique au courant du collecteur du transistor Q_2 . Ce courant est fixé par $(V_{cc} - V_{BC})/R$.

Le même raisonnement est valable pour les transistors MOS.



Annexe C

Puissance lumineuse à la sortie de la fibre

Puissance de la diode laser

C-1 INTRODUCTION

L'objectif de cette étude est l'évaluation de la puissance de la diode-laser pour éclairer les photodiodes des circuits électroniques VLSI. L'exemple étudié est le circuit de l' "Horloge Optique" qui contient 64 processeurs élémentaires sous forme d'une matrice 8*8, dans chaque processeur élémentaire nous devons éclairer deux photodiodes successivement pour générer l'horloge optique (voir Chapitre VII § V-5).

Le montage optique consiste à utiliser des fibres optiques monomodes. A la sortie de la fibre, le faisceau lumineux est collimaté pour éclairer l'hologramme de Talbot qui transforme le faisceau en petits points selon le réseau carré (voir figure C-1).

L'étude commence par la modélisation du profil de l'énergie à la sortie de la fibre. Nous pouvons ainsi évaluer la puissance utile pour éclairer l'hologramme par rapport à la puissance à l'intérieur de la fibre.

En considérant les différentes pertes dues au taux d'injection dans la fibre, au montage optique et à l'efficacité de l'hologramme, et sachant la puissance nécessaire pour éclairer une photodiode (voir Chapitre VII), nous pouvons évaluer la puissance de la diode laser à utiliser en fonction de la longueur d'onde. L'étude est faite pour trois longueurs d'onde 0.675 μm , 0.790 μm et 0.904 μm .

Nous finissons cet annexe par l'étude de la disposition des fibres pour éclairer deux réseaux séparés de photodiodes.

C-2 PROFIL D'ENERGIE A L'INTERIEUR DE LA FIBRE

La résolution des équations de Maxwell dans les fibres optiques à saut d'indice donne un profil d'énergie en fonctions de Bessel assez compliqué. Ce profil peut être approximé par un profil gaussien dont la largeur ω_0 est donnée par la relation suivante [55] :

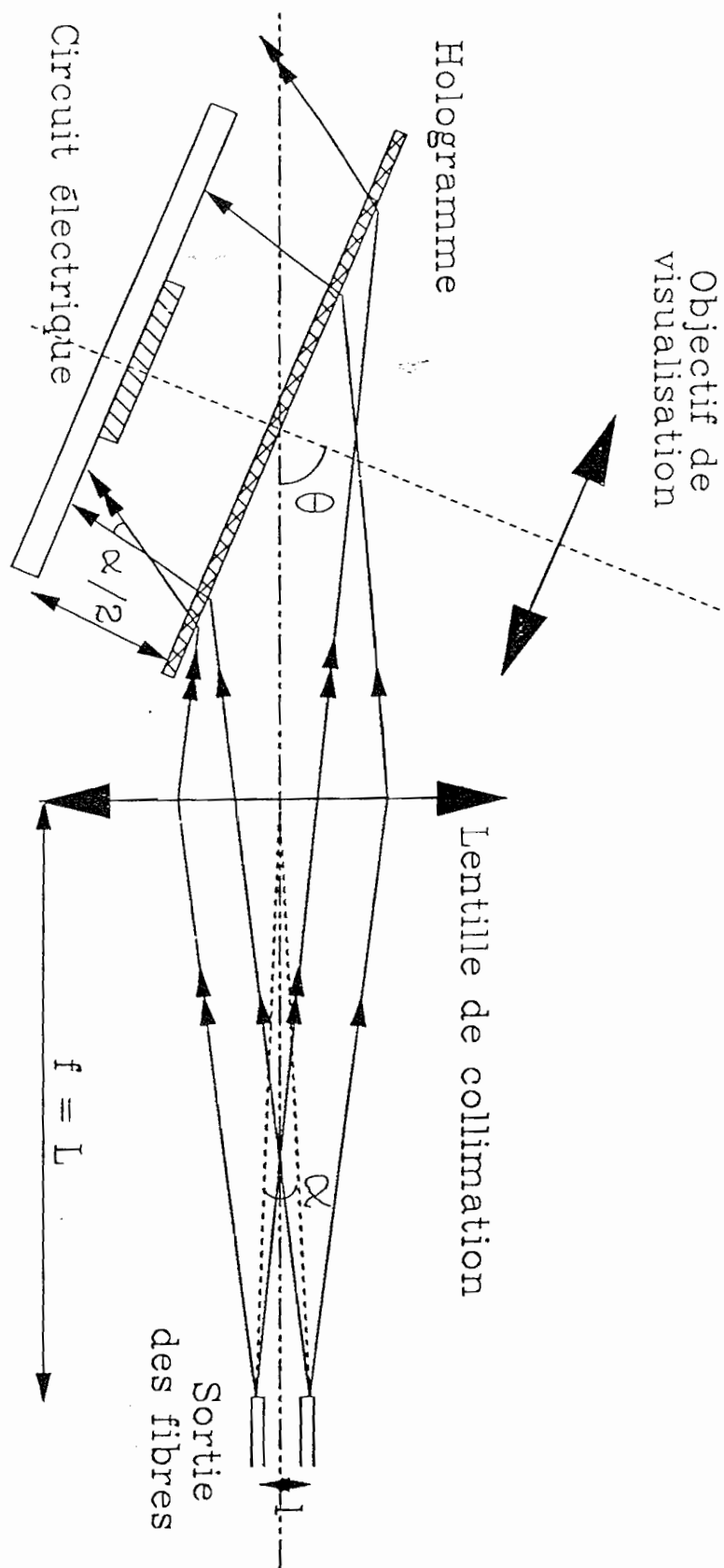


Fig. C-1 : Montage d'éclairage en utilisant des fibres optiques.

$$\frac{\omega_0}{a} = 0.65 + \frac{1.619}{V^{3/2}} + \frac{2.879}{V^6}$$

tel que a est le rayon du coeur de la fibre, et V est la fréquence normalisée donnée par la relation :

$$V = 2\pi \cdot \frac{(\text{O.N.}) \cdot a}{\lambda} = \sqrt{8} \cdot a \cdot n_{\text{coeur}} \cdot \frac{\sqrt{\Delta}}{\lambda}$$

$$\text{L'ouverture numérique : O.N.} = \sqrt{n_{\text{coeur}}^2 - n_{\text{gaine}}^2}$$

$$\Delta = \frac{n_{\text{coeur}} - n_{\text{gaine}}}{n_{\text{coeur}}}$$

La fréquence normalisée détermine si la fibre est monomode ou multimode pour la longueur d'onde utilisée selon le critère :

$$V \leq 2.405 \text{ la fibre est monomode}$$

ainsi nous pouvons considérer que le champ électrique à l'intérieur de la fibre a la forme :

$$E = E_0 \cdot \exp\left(-\frac{r^2}{\omega_0^2}\right)$$

et la puissance optique dans la fibre :

$$W_0 = E_0^2 \int_0^{2\pi} \int_0^{\infty} \exp\left(-2 \cdot \frac{r^2}{\omega_0^2}\right) \cdot r \cdot dr \cdot d\theta = \frac{\pi}{2} E_0^2 \cdot \omega_0^2$$

C-3 PROFIL DE L'ENERGIE OPTIQUE A LA SORTIE DE LA FIBRE

Nous déterminons par la suite le profil de l'onde gaussienne à la sortie de la fibre dans un plan parallèle à celle-ci situé à une distance L . Les paramètres de cette onde gaussienne sont fonctions de l'onde gaussienne à l'intérieur de la fibre. Nous considérons que la largeur de la gaussienne à l'intérieur de la fibre est ω_0 , à la proximité de la fibre ω_1 et à la distance l de la fibre ω_L .

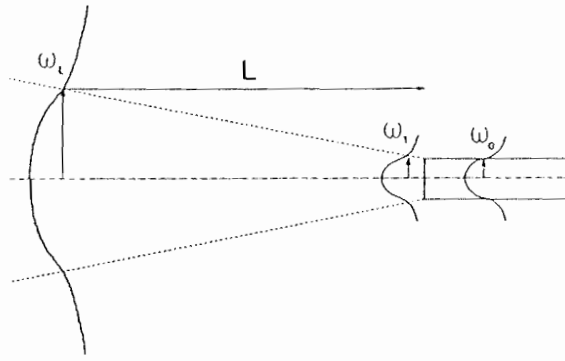


Fig.C-2 : Profil d'énergie à la sortie de la fibre.

Pour calculer le profil de l'énergie à n'importe quel endroit à la sortie de la fibre, nous utilisons la propagation des ondes gaussiennes en nous servant du calcul matriciel [56] (voir annexe D).

C-3-1 Profil d'énergie à la sortie de la fibre

En utilisant la matrice d'interface pour les ondes gaussiennes, nous trouvons l'expression suivante pour le champ électrique juste à la sortie de la fibre :

$$E = E_0 \frac{\omega_0}{\omega_1} \exp\left(-\frac{r^2}{\omega_1^2}\right)$$

avec

$$n_0 \cdot \omega_0^2 = n_1 \cdot \omega_1^2$$

Et nous avons la nouvelle amplitude :

$$E_1 = E_0 \frac{\omega_0}{\omega_1}$$

C-3-2 Profil d'énergie à une distance l de la fibre

L'onde gaussienne sort avec une largeur ω_1 et se propage dans le même milieu. La largeur de cette onde gaussienne ω_L à une distance L peut être calculée en utilisant la matrice de propagation de la fibre. Nous trouvons l'expression suivante pour le champ électrique :

$$E = E_1 \frac{\omega_1}{\omega_L} \exp\left(-\frac{r^2}{\omega_L^2}\right)$$

avec

$$\omega_L^2 = \omega_1^2 \left[1 + \left(\frac{z\lambda}{\pi\omega_1^2 n} \right)^2 \right]$$

C-4 BILAN ENERGETIQUE POUR UTILISER UNE DIODE LASER FIBREE

Le montage d'éclairage est proposé dans la figure C-1 . Ce montage utilise des diodes laser couplées à des fibres optiques. La lumière à la sortie de la fibre est collimatée avec une lentille à la distance l de la fibre. Le faisceau collimaté éclaire l'hologramme de Talbot qui distribue la lumière aux photodiodes du circuit arrangées sous forme d'une matrice carrée $8 * 8$. Le pas est le même sur les deux axes perpendiculaires et est égal à $300 \mu\text{m}$.

Nous allons estimer par la suite les différentes pertes de lumière dans les différentes étapes du montage.

La puissance de la diode laser est appelé par la suite W_{DL}

C-4-1 Puissance lumineuse nécessaire pour éclairer une photodiode

Nous avons étudié dans le Chapitre VII la sensibilité des photodiodes utilisées dans les circuits VLSI, et la puissance optique nécessaire pour le bon fonctionnement du circuit de l'"Horloge Optique". La figure VII-10 (§ VII-5-1) présente la puissance nécessaire pour faire commuter la tension des photodiodes utilisées sur ce circuit de 5 V en quelques dizaines de nano-secondes

C-4-2 Pertes par reflexion sur le silicium

Le faisceau lumineux projet sur les photodiodes de silicium est en partie réfléchi. Compte-tenu que la puissance optique W_{ph} est la puissance absorbée par la photodiode, la puissance projetée W_{Si} doit être :

$$(1-R_{Si}) \cdot W_{Si} = W_{ph}$$

tel que R_{Si} est le coefficient de reflexion sur le silicium sous incidence normale donné par :

$$R_{Si} = \left[\frac{n_{Si} - 1}{n_{Si} + 1} \right]^2$$

En considérant $n_{Si} = 3.4$, nous trouvons $R_{Si} = 0.30$

Remarque : Dans ce calcul, nous avons négligé la partie imaginaire de l'indice de réfraction n_{Si} qui est en général plus faible que la partie réelle.

C-4-3 Défaut de focalisation

Nous appelons défaut de focalisation tout défaut qui envoie de l'énergie optique en dehors de la surface à éclairer. Ceci revient à faire la différence entre l'intégrale du profil de la tâche formée sur la surface totale et l'intégrale de recouvrement entre cette tâche et la surface à éclairer.

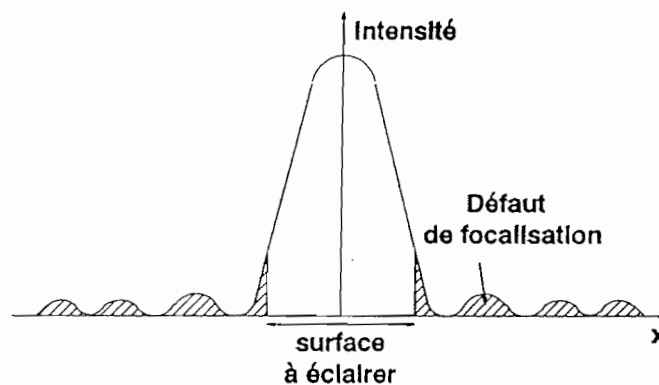


Fig. C-3 : Défaut de focalisation.

Dans notre exemple nous utilisons un hologramme de Talbot comme un illuminateur de tableau pour éclairer les photodiodes. L'expérience a montré que les tâches restituées ont une forme semblable au produit de sinc au lieu d'être des carrés parfaits, les pieds de ces sinc dépassent les bords des photodiodes, par conséquent une certaine proportion de l'énergie optique diffracté par l'hologramme est envoyée en dehors de la photodiode. Nous appelons cette proportion le défaut de focalisation que nous estimons dans le cas actuel de l'hologramme de Talbot :

$$D_{foc} = 10 \%$$

Dans notre application, nous avons la relation :

$$W_{Si} = D_{foc} \cdot W_{diff}$$

C-4-4 Nombre de photodiodes à éclairer

Dans le cas du circuit de l' "Horloge Optique" le nombre de photodiode à éclairer est égal à $N = 64$, par conséquent la puissance lumineuse utile diffractée par l'hologramme doit être de :

$$(1 - D_{foc}) \cdot W_{diff} = N \cdot W_{Si}$$

C-4-5 Efficacité de l'hologramme

Nous appelons efficacité de l'hologramme le rapport entre la puissance optique incidente sur l'hologramme, et celle diffractée par l'hologramme dans l'ordre 1 pour reconstruire l'image.

Pour mesurer cette efficacité, nous mesurons l'éclairement normal incident avant l'hologramme à l'aide d'un wattmeter, et l'éclairement de la lumière diffractée dans l'ordre 1 après l'hologramme de telle manière que :

$$\eta_{holo} = W_{diff} / W_{incident}$$

Nous trouvons une efficacité pour l'hologramme de Talbot de :

$$\eta_{holo} = 70 \%$$

C-4-6 Surface éclairée sur l'hologramme

La structure de l'hologramme de Talbot nous permet de le considérer comme une matrice de microlentilles. Chaque microlentille a une largeur égale à la période du réseau de l'objet Talbot utilisé pour enregistrer l'hologramme, dans notre cas cette période est égale à $300 \mu m$. Par conséquent une photodiode reçoit la lumière diffractée par la microlentille correspondante, c'est à dire d'une surface de $300 \cdot 300 \mu m^2$ de l'hologramme. Cette propriété nous permet d'éclairer juste la surface de l'hologramme utile pour éclairer notre matrice de photodiodes, dans le cas du circuit de l' "Horloge Optique" cette surface est un carré de côté égal à $c = 8 \cdot 300 \mu m = 2.4 mm$.

En tenant compte de la géométrie du faisceau gaussien qui éclaire l'hologramme et des considérations énergétiques, la surface éclairée de l'hologramme est plus grande que celle utilisée pour éclairer les photodiodes, d'où la notion de surface utile sur l'hologramme. En conséquence, il y aura une partie de la lumière projetée sur l'hologramme qui ne servira pas à éclairer les photodiodes donc sera perdue.

Les considérations énergétiques imposent, comme nous l'avons vu, une puissance optique minimale pour éclairer une photodiode W_{ph} . En tenant compte de toutes les pertes par efficacité de diffraction, par défauts de focalisation et par réflexion sur le silicium, l'éclairement minimale I_{min} sur la surface utile de l'hologramme doit vérifier la relation :

$$I_{min} \cdot c^2 \cdot \eta_{holo} \cdot (1 - D_{foc}) \cdot (1 - R_{Si}) = N \cdot W_{ph}$$

La restitution de l'hologramme se fait par un faisceau incliné d'un angle θ qui est fonction de la longueur d'onde de restitution. Nous considérons un faisceau gaussien circulaire de largeur ω_L dont le champ électrique est donné par :

$$E(r, 0) = \frac{\omega_1}{\omega_L} E_1 \cdot \exp\left(-\frac{r^2}{\omega_L^2}\right)$$

Ce faisceau est, en fait, la lumière sortant de la fibre optique et collimée par une lentille.

La projection de ce faisceau sur un plan incliné d'un angle θ (le plan de l'hologramme) fait une tâche elliptique dont les axes principaux sont donnés en fonction de ω_L par les relations suivantes :

$$\begin{aligned}\omega_{La} &= \omega_L \\ \omega_{Lb} &= \omega_L / \cos \theta\end{aligned}$$

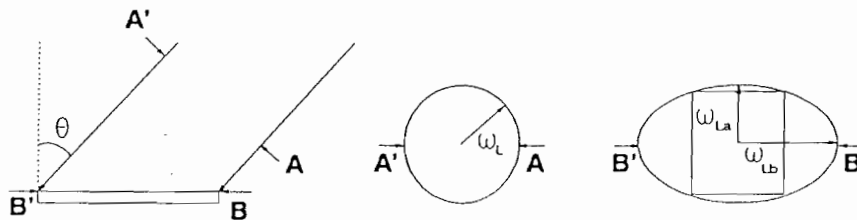


Fig. C-4 : Eclairage de l'hologramme.

Et le champ électrique sur ce plan est donné par :

$$E(x, y) = \frac{\omega_1}{\omega_L} E_1' \cdot \exp\left(-\frac{(x^2 \cdot \cos^2 \theta + y^2)}{\omega_L^2}\right)$$

Pour trouver l'expression de E_1' , nous écrivons la condition de conservation de l'énergie :

$$\int_0^{2\pi} \int_0^\infty E_1'^2(r, \theta) \cdot dr \cdot d\theta = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} E_1'^2(x, y) \cdot dx \cdot dy$$

et nous trouvons

$$E_1' = \sqrt{\cos \theta} \cdot E_1$$

L'intensité est, par conséquent, donnée par la relation :

$$I(x, y) = \frac{\omega_1^2}{\omega_L^2} \cdot E_1^2 \cdot \cos \theta \cdot \exp\left(-\frac{(x^2 \cos^2 \theta + y^2)}{\omega_L^2}\right)$$

Et l'intensité minimale sur le carré de côté c aura la valeur :

$$I_{\min} = \frac{\omega_1^2}{\omega_L^2} \cdot E_1^2 \cdot \cos \theta \cdot \exp\left(-\frac{c^2 (1 + \cos^2 \theta)}{2\omega_L^2}\right)$$

Nous considérons la puissance d'origine celle dans la fibre optique W_{fibre} qui est égale, comme nous l'avons vu, à :

$$W_{\text{fibre}} = \frac{\pi}{2} E_0^2 \cdot \omega_0^2 = \frac{\pi}{2} E_1^2 \cdot \omega_1^2$$

Nous rapportons la puissance lumineuse réellement utile au niveau de l'hologramme W_{holo} à cette puissance d'origine, nous obtenons la rapport suivant :

$$\frac{W_{\text{holo}}}{W_{\text{fibre}}} = \frac{2 \cdot c^2}{\pi \cdot \omega_L^2} \cdot \cos \theta \cdot \exp\left(-\frac{c^2 (1 + \cos^2 \theta)}{2\omega_L^2}\right)$$

Ce rapport nous permet de calculer la puissance optique dans la fibre en fonction de puissance nécessaire pour éclairer l'hologramme.

D'autre part, nous remarquons que ce rapport est directement fonction de ω_L qui est à son tour fonction des caractéristiques de la fibre optique et de la distance L entre la lentille de collimation et la fibre. En étudiant ce rapport en fonction de la distance L nous trouvons une courbe comme celle présentée dans la figure suivante :

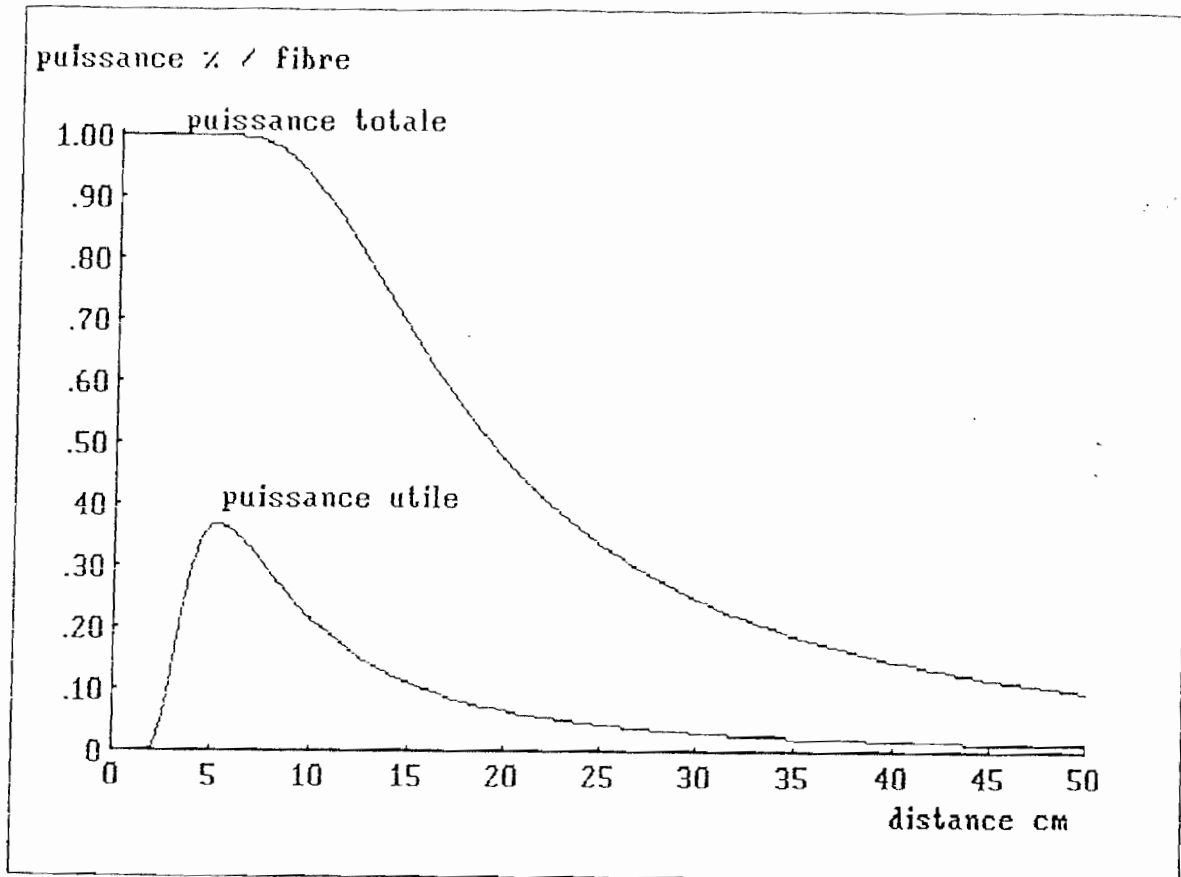


Fig. C-5 : Puissance optique à la sortie de la fibre optique.

Une telle courbe nous indique qu'il y a une distance optimale pour placer la lentille de collimation par rapport à la fibre. Cette distance optimale donne la focale de la lentille de collimation à utiliser.

Le diamètre de la pupille de la lentille de collimation D sera donnée en fonction du côté du carré à éclairer c et de l'angle θ selon la relation :

$$D = c \cdot \sqrt{2 \cdot (1 + \cos^2 \theta)}$$

Les applications numériques seront données dans le tableau C-1 en fonction des longueurs d'onde étudiées.

C-4-7 Pertes par reflexion sur les dispositifs optiques

Nous utilisons une lentille de collimation du faisceau lumineux sortant de la fibre à la distance L avant d'éclairer l'hologramme. Sur cette lentille nous perdons par reflexion sur les deux faces environ 8 % de la puissance optique sortant de la fibre. Ces pertes sont désignées par R_{ref}

C-4-8 Pertes par reflexion à l'interface de la fibre

Comme il y a une différence de l'indice de refraction entre l'air et le cœur de la fibre nous perdons environ 4 % de la puissance optique à la sortie de la fibre.

C-4-9 Pertes dans la fibre

L'atténuation dans les fibres est de l'ordre de 0.3 dB/km, et comme la longueur de la fibre sera d'environ d'un mètre, nous pouvons négliger les pertes dans les fibres tout en faisant attention que les fibres ne présentent pas de fortes courbures (les courbures ne font pas partie des 0.3 dB/km).

C-4-10 Taux d'injection dans la fibre

La puissance émise par la diode laser est en partie injectée dans la fibre optique. Ce taux d'injection dépend du connecteur mais on peut l'estimer à 30 %. Ainsi, la puissance dans la fibre est :

$$W_{\text{fibre}} = 0.30 * W_{DL}$$

C-4-11 Relation entre la puissance de la diode laser et la puissance de chaque photodiode

$$W_{DL} \cdot 0.27 \cdot \frac{W_{\text{holo}}}{W_{\text{fibre}}} \cdot \eta_{\text{holo}} \cdot (1 - D_{\text{foc}}) \cdot (1 - R_{\text{Si}}) = N \cdot W_{\text{ph}}$$

C-5 LA PUISSANCE IMPULSIONNELLE DE LA DIODE LASER

Nous avons étudié le bilan énergétique de ce montage pour trois longueurs d'onde 0.675 μm , 0.790 μm et 0.904 μm . Nous avons choisi les diamètres du cœur de la fibre optique pour qu'elle soit monomode. Nous avons modélisé le profil de l'énergie à l'intérieur de la

fibre pour en déduire la divergence de la lumière sortant. Ensuite nous avons étudié le facteur $W_{\text{holo}}/W_{\text{fibre}}$ en fonction de la distance de la fibre L et nous avons obtenu ainsi le diamètre et le focal optimal pour la lentille de collimation (voir tableau C-1).

C-6 DETERMINATION DE L'ESPACEMENT ENTRE LES FIBRES OPTIQUES

Dans le circuit de l' "Horloge optique" il faut éclairer deux photodiodes pour chaque cellule de la matrice avec deux sources différentes, donc avec deux diodes laser. L'espacement sur le circuit entre les deux photodiodes est de $50 \mu\text{m}$.

Pour éclairer ces deux réseaux de photodiodes, il suffit d'incliner les faisceaux lumineux éclairant l'hologramme d'un certain angle α (voir figure C-1), selon les lois de conjugaison des hologrammes les deux images restituées seront, en première approximation, inclinées du même angle α . (en première approximation, car si la porteuse est très inclinée on sort du domaine paraxial)

Pour incliner les deux faisceaux, il suffit de séparer les deux fibres optiques d'une certaine distance L derrière la lentille de collimation.

C-6-1 Distance du plan de restitution de l'hologramme

Selon les lois de conjugaison des hologramme, la distance du plan de restitution d est fonction de la longueur d'onde d'enregistrement λ_1 et celle de restitution λ_2 selon la relation :

$$d_1 \cdot \lambda_1 = d_2 \cdot \lambda_2$$

L'enregistrement de l'hologramme est fait à longueur d'onde $0.488 \mu\text{m}$ avec une distance du plan de restitution $d_1 = 15.35 \text{ mm}$. La relation précédente nous permet de déterminer la distance de restitution d_2 selon la longueur d'onde de restitution utilisée λ_2 .

C-6-2 Angle entre les deux faisceaux

La distance entre les deux réseaux des photodiodes étant p , l'angle entre les deux faisceaux sera :

$$\alpha = p / d_2$$

C-6-3 Distance entre les deux fibres optiques

Les deux faisceaux sortant des fibres optiques sont collimatés par une lentille de distance focale f . L'angle entre les deux faisceaux est fonction de la distance focale de la lentille f et de la séparation l entre les deux fibres selon la relation :

$$\alpha = l / f$$

Remarque : La distance p entre les deux réseaux des photodiodes étant très faible ($50 \mu\text{m}$), l'angle entre les deux faisceaux pourrait être extrêmement faible au point que la distance entre les deux fibres l soit difficile à réaliser pratiquement. Dans ce cas, nous pouvons adopter un angle plus grand pour que la déviation de l'image soit d'une période plus la distance p , ce qui désigné dans le tableau C-1 par (+période).

C-7 CONCLUSION

Nous avons déterminé les différents facteurs pour le montage d'éclairage du circuit de l'"Horloge optique" en utilisant des diodes laser fibrées et l'holgramme de Talbot. Les puissances des diodes laser semblent très grandes pour pouvoir utiliser les diodes laser habituelles. Des diodes laser impulsionnelles de grandes puissances sont commercialisées mais à des prix plus élevés. Une autre solution qui peut diminuer les pertes est d'utiliser directement les diodes laser sans les fibres optiques, mais la distance entre les deux diodes sera trop grande.

		Longueur d'onde		
		0.675 μm	0.790 μm	0.904 μm
FIBRE OPTIQUE	diamètre du coeur μm	4	4	5
	ouverture numérique	0.1	0.1	0.1
	distance focale de la lentille cm	3	3	2.5
SYSTEME DE COLLIMATION	diamètre de la lentille mm	2.72	2.64	2.52
	angle les deux faisceaux rad	0.0045	0.0053	0.006
	distance entre les fibres mm	0.135	0.159	0.15
	angle entre les deux faisceaux (+période)rad	0.0315	0.0371	0.042
	distance entre les fibres (+période) mm	0.945	1.113	1.05
HOLOGRAMME	angle d'éclairage de l'hologramme	56.97°	63.7°	71.12°
	$W_{\text{holo}}/W_{\text{fibre}}$	18.8%	16.5%	13%
	distance du plan de restitution mm	11.1	9.48	8.28
BILAN ENERGETIQUE	puissance de la diode laser mW	103	110	191
	puissance sur le circuit mW	2.24	2.11	2.88
	puissance sur une photodiode μW	35	33	45

Tableau C-1

Annexe D

Propagation des ondes gaussiennes

Calcul matriciel

La forme générale d'une onde gaussienne est donnée par l'expression :

$$E(r, \theta, z) = E \frac{\omega_0}{\omega(z)} \left\{ \exp \left[-i (kz - \eta(z)) - r^2 \left(\frac{1}{\omega^2(z)} + \frac{i k}{2R(z)} \right) \right] \right\}$$

$$\omega^2(z) = \omega_0^2 \left[1 + \left(\frac{z\lambda}{\pi\omega_0^2 n} \right)^2 \right]$$

$$R(z) = z \left[1 + \left(\frac{\pi\omega_0^2 n}{\lambda z} \right)^2 \right]$$

$$\eta(z) = \tan^{-1} \left(\frac{\lambda z}{\pi\omega_0^2 n} \right)$$

Nous définissons le rayon de courbure imaginaire $q(z)$:

$$\frac{1}{q(z)} = \frac{1}{R(z)} - i \frac{\lambda}{\omega_0^2 n}$$

Suivant le calcul matriciel le rayon imaginaire change selon la relation :

$$q_2 = \frac{Aq_1 + B}{Cq_1 + D}$$

A, B, C, D sont les éléments de la matrice :

$$\begin{bmatrix} A & B \\ C & D \end{bmatrix}$$

Nous donnons la matrice de propagation dans les deux cas :

Interface diélectrique :

$$\begin{bmatrix} 1 & 0 \\ 0 & n_1/n_2 \end{bmatrix}$$

Propagation à la distance d :

$$\begin{bmatrix} 1 & d \\ 0 & 1 \end{bmatrix}$$

Annexe E

Participation aux congrès

- a) 1990 International Topical Meeting on Optical Computing.
Kobe, JAPON, 8-12 Avril 1990.
- b) Quatrième Colloque National de Visualisation et de Traitement d'Image en
Mécanique des Fluides.
Lille, 29 Mai-1 Juin 1990.
- c) Optics in Complex Systems.
Garmisch-Partenkirchen, R.F.A., 5-10 Août 1990.
- d) Sixième Journée d'Etude sur les Fonctions Optiques dans l'Ordinateur.
Strasbourg, 6-7 Septembre 1990
- e) Septième Journée d'Etude sur les Fonctions Optiques dans l'Ordinateur.
Brest, 19-20 Septembre 1991.
- f) 75^{ème} Exposition de Physique.
Villepinte, 18-22 Novembre 1991.
- g) WOIT Conference.
Edinburg, U.K., 15-17 Décembre 1991.

ARRAY ILLUMINATOR HOLOGRAM BASED ON TALBOT EFFECT

I. SEYD-DARWISH, P. CHAVEL, J. TABOURY, Y. MALET
INSTITUT D'OPTIQUE (CNRS), ORSAY, FRANCE

I- ARRAY ILLUMINATOR

1) WHY ?

To broadcast one signal optically to many sites on an electro-optic or non-linear-optical two dimensional devices.

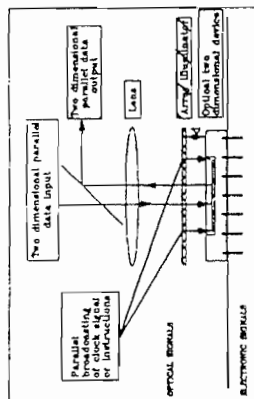


Figure 1
Motivation

ESD, PCT, J. F. DARWISH
100 - 15
Garmisch - Partenkirchen
Germany

I- ARRAY ILLUMINATOR (CONT'D)

2) REQUIREMENTS

- Focus light onto pixels ($\rightarrow$).
- No light outside pixels.
- High efficiency.
- Array illuminator should be transparent to other signals ($\rightarrow$).
- For mechanical stability keep distance d between array and the 2D device short

For system integration } illuminator and the 2D device short

II- OUR SOLUTION : THE TALBOT HOLOGRAM

1) THE TALBOT EFFECT

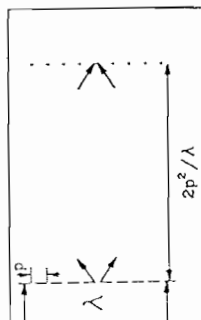


Figure 2

2) TALBOT HOLOGRAM RECORDING

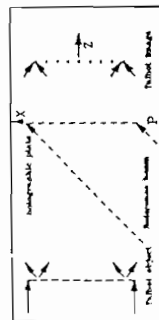


Figure 3

3) TALBOT HOLOGRAM RECONSTRUCTION

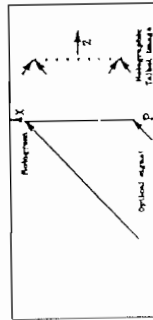


Figure 4

- efficiency (c) can be high (thick phase medium) see III.
- Transparent to other signal (d) (Bragg-mismatched).

III- EFFICIENCY OPTIMIZATION

- 100 % theoretical efficiency for thick hologram can be reached for a given ratio of reference to object beams (experimentally: $\eta > 90\%$).
- In Talbot hologram, η is smaller because object beam in the plane P is not uniform: large intensity dynamic range, so $\eta_{\text{min}} < \eta < \eta_{\text{max}}$.
- Optimum for η and $(\eta_{\text{max}} - \eta_{\text{min}})$ is reached in the minimum dynamic range plane.

4) Fresnel diffraction calculations:

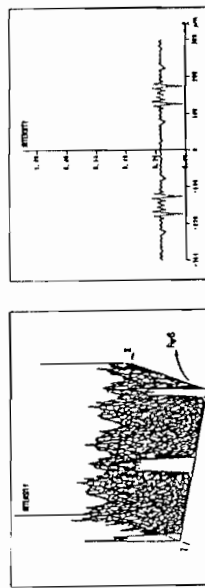


Figure 5 : PERSPECTIVE VIEW OF FRESNEL DIFFRACTION
P1 OR P2

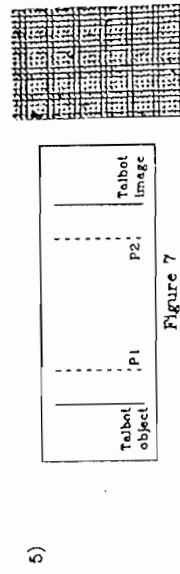


Figure 6

Choose P2 for minimum d (c)

- Interpretation : at P2, we effectively have a set of spherical waves converging onto the pixels and overlapping slightly.

IV- RESULTS

1) EXPERIMENTAL SETUP

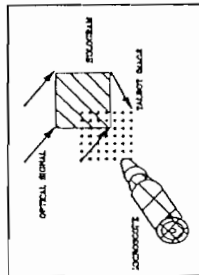


Figure 8 (a)



Figure 8 (b)

Parameters : Pixel size $50 \times 50 \mu\text{m}^2$
Pixel spacing $300 \mu\text{m}$
Wavelength $\lambda = 0.488 \mu\text{m}$
Talbot distance $= 36.9 \text{ cm}$
Distance $d = 3.1 \text{ cm}$

2) ACHIEVEMENTS

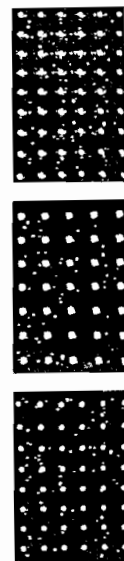


Figure 9 (a)

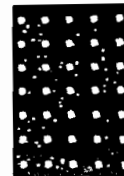


Figure 9 (b)

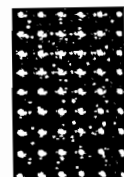


Figure 9 (c)

$\eta = 3\%$ $\eta = 70\%$

NOTE:

Compared with holographic lenslet arrays, this technique has the advantage of global recording. Also, pointlike defects are eliminated by the Talbot effect. The method is simple and has slightly higher diffraction efficiency than Damman gratings.

AUTOMATE CELLULAIRE OPTO-ELECTRONIQUE

POUR LE GAZ SUR RESEAU

I. SEYD-DARWISH P. CHAVEL J. TABOURY S. MALLICK

Institut d'Optique Théorique et Appliquée (Unité associée CNRS)
BP. 147, 91403 ORSAY CEDEX

F. DEVOS R. REYNAUD T. MAURIN

Institut d'Electronique Fondamentale (Unité associée CNRS)
Université de Paris XI Bât. 220, 91405 ORSAY CEDEX

Résumé

Nous décrivons un automate cellulaire opto-électronique pour le traitement de l'algorithme de simulation des écoulements bidimensionnels des fluides "Gaz sur Réseau". Cet algorithme exige un grand nombre d'itérations. L'automate proposé fait le traitement en parallèle sur un grand nombre de processeurs implantés sur une puce électronique avec des dispositifs optiques d'entrée-sortie, ce qui diminue le temps de calcul.

1/ INTRODUCTION

Il a été montré la possibilité de simuler l'écoulement bidimensionnel des fluides par le modèle appelé "Gaz sur réseau" <1>. Dans ce modèle les particules sont restreintes à occuper les noeuds d'un réseau régulier, et à se déplacer à une vitesse constante entre les noeuds voisins. L'évolution du fluide est obtenue par les collisions des particules au niveau des noeuds du réseau. Ces collisions sont régies par des lois spécifiques pour vérifier les propriétés macroscopiques du fluide, comme la conservation de la masse, de la quantité de mouvement et de l'énergie.

Cette méthode a l'avantage de fournir un outil simple pour étudier l'écoulement des fluides en évitant la résolution des équations de Navier-Stokes.

Dans ce modèle, l'aspect cellulaire du traitement permet de concevoir des

automates cellulaires parallèles spécialisés. La machine R.A.P.I est l'une des premières machines connues pour cette application <2>.

Nous proposons une réalisation opto-électronique d'un tel automate cellulaire. Nous décrivons un dispositif optique qui assure le traitement parallèle des sites de collision, et une visualisation parallèle des résultats.

2/ PRINCIPE ET ETAPES DE FONCTIONNEMENT DU GAZ SUR RESEAU

La théorie montre qu'il faut utiliser un réseau hexagonal <1>. A l'instant "t", chaque noeud voit arriver au maximum 6 particules, avec éventuellement une particule immobile en son centre. Ces particules entrent en collision, puis se propagent vers les noeuds voisins.

La présence de parois ou d'obstacles à l'écoulement du fluide est simulée par des lois particulières de choc sur la surface des parois et par des zones interdites pour la présence des particules.

Le traitement au niveau de chaque noeud consiste à reconnaître la configuration des particules incidentes à l'instant "t" et à lui substituer la configuration correspondante des particules après collision. Cette reconnaissance se fait par comparaison de cette configuration avec les 128 (=2⁷) configurations possibles. Une fois que ce traitement est fait pour tous les noeuds du réseau, les particules résultantes se propagent vers les noeuds voisins et une itération est ainsi terminée (fig. 1).

(b)

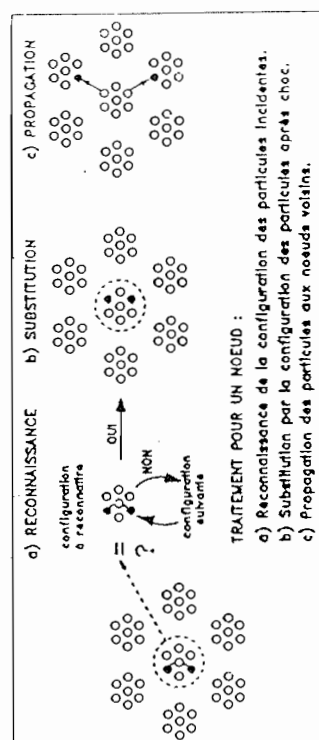


Figure 1: "Gaz sur réseau": principes de l'algorithme

Un traitement sur ordinateur classique entièrement séquentiel se fait itérativement : au cours d'un balayage des nœuds, la configuration de chaque nœud est associée par une table de transcodage au contenu du nouvel état qui est inscrit dans des registres des nœuds voisins.

En fait, ce traitement au niveau de chaque nœud est indépendant de l'état des nœuds voisins, ce qui permet de faire le traitement pour tous les nœuds en parallèle.

Nous décrivons par la suite un automate cellulaire semi-parallèle : parallèle pour les nœuds du réseau et séquentiel pour les configurations possibles, ce qui doit contribuer à accélérer le calcul global en supprimant le balayage.

3/ REALISATION OPTO-ELECTRONIQUE

L'état du gaz est présenté sur un circuit électronique composé de plusieurs processeurs distribués sur un réseau hexagonal. Chaque processeur est associé à un nœud du réseau et composé de 7 unités correspondant aux particules présentes en chaque nœud. Chaque unité contient une mémoire d'entrée à 1 bit qui indique la présence ou l'absence de particule incidente selon la direction correspondante, d'une mémoire de sortie qui contient la particule résultante du choc et d'une photodiode. Une unité de reconnaissance à 1 bit est reliée à chaque processeur pour décider de la reconnaissance, ou non, de la configuration des particules incidentes. Chaque mémoire de sortie est liée à la mémoire d'entrée correspondante dans une cellule voisine, ce qui constitue le réseau de connexion pour la propagation des particules (fig. 2).

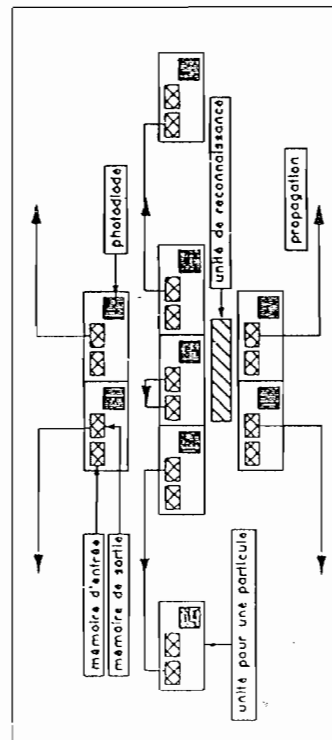


Figure 2: Schéma d'un processeur

Nous projetons en parallèle sur l'ensemble des processeurs et séquentiellement dans le temps les 128 configurations possibles des particules incidentes. Ces configurations éclairent les photodiodes de chaque processeur. Chaque configuration de particules incidentes est suivie de la configuration de substitution correspondante. Pour chaque processeur, l'état d'éclairage des photodiodes est comparé au contenu de l'ensemble des mémoires d'entrée, en cas de correspondance parfaite l'unité de reconnaissance se met au niveau logique '1' : signe de reconnaissance. Le niveau logique '1' de l'unité de reconnaissance permet à l'image de substitution d'être transférée dans les mémoires de sortie. Une fois les différentes possibilités avec leurs substitutions projetées, un signal électrique global transfère le contenu des mémoires de sortie dans les mémoires d'entrée correspondantes des cellules voisines et une itération est alors terminée.

La projection des images se fait par un système holographique illuminé par 7 sources lumineuses (diodes laser) réparties en hexagone qui émettent les différentes configurations des particules. Le système holographique fournit autant d'image de ces sources qu'il y a de processeurs et distribue ces images aux différents processeurs.

Nous initialisons le système par la projection de l'image de la configuration initiale des particules sur les photodiodes du circuit par un autre système optique.

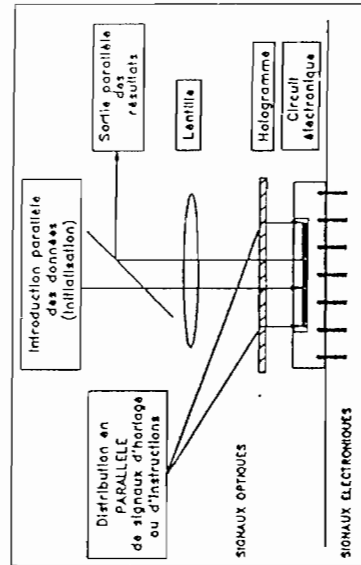


Figure 3: Schéma final de l'automate cellulaire opto-électronique

L'état du gaz à la fin d'une itération sera lu en utilisant un circuit AsGa à puits quantiques multiples couplés avec les mémoires de sortie de chaque cellule. Ce circuit projette une image de l'état de l'ensemble des mémoires de sortie (donc de l'état du gaz) sur une caméra CCD. Transférée à un ordinateur, cette image sera traitée sans interrompre le fonctionnement de l'automate (fig. 3).

4/ REALISATION

Nous avons réalisé un circuit de validation contenant un seul processeur. Le test sur ce circuit révèle une vitesse de fonctionnement supérieure à 10.000 itérations par seconde.

En même temps, nous développons une procédure pour réaliser le système holographique de distribution des signaux.

5/ CONCLUSION

Nous avons présenté le principe de réalisation d'un automate cellulaire opto-électronique spécialisé pour le traitement du "gaz sur réseau". Cet automate permet un traitement parallèle au niveau de tous les processeurs en même temps. Une visualisation parallèle de toutes les cellules est proposée en utilisant des circuits à puits quantiques multiples.

REFERENCES

- <1> U. FRISCH, B. HASSLACHER, Y. POMEAU.
"Lattice gaz automata for Navier-Stokes equation", Phys.Rev.Lett.
56,1505-1508, (1986).
- <2> A.CLOUQUEUR, D. D'HUMIERES.
"R.A.P.1, a cellular automaton machine for fluid dynamics", Complex
Sys. 1, 585 (1987).

(c)

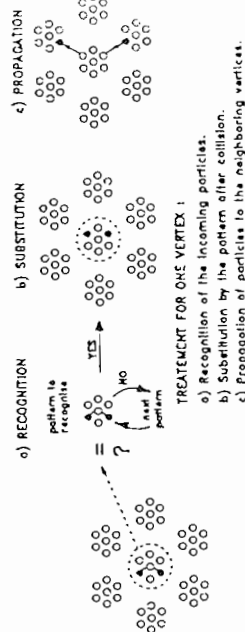


Fig 2 : "Gas lattice" : Algorithm principle.

Input memories, each connected to a photodiode. The match pattern is projected optically onto the photodiodes. If the pattern corresponds to the memories content, one bit in the processor is set to high logical level (recognition sign), otherwise it stays low. The high logical signal allows the substitute pattern, which is projected onto the photodiodes just after the match pattern, to be transferred to seven 1-bit output memories. Recognition and substitution are performed sequentially for all possible incoming particle configurations, and in parallel for the processors. Then one iteration is achieved by electrically propagating the output memory states to the input memories of neighboring cells, and the circuit is ready for another iteration.

2. OPTICAL DEVICES

2.1. Array illuminator

For the parallel input of instructions, we need an optical device to distribute the same optical patterns to all the processors. We have developed an array illuminator hologram based on the Talbot effect.

2.2. Output optical device

For the parallel output, we plan to use multiple quantum wells arranged in the same lattice form. Each MQW pixel is connected to an output memory of the circuit, its reflection state is driven by the memory voltage. In this way, an image of the memories state can be formed, quasi instantaneously onto another parallel electronic circuit to analyse the results.

3. ACHIEVEMENT

We have implemented and tested a first electronic circuit containing one processor corresponding to one vertex. The second circuit will contain several processors connected to their neighbors.

We have realized an array illuminator hologram for a square matrix of $30 \times 50 \mu\text{m}^2$ spots with period of $300 \mu\text{m}$ on dichromated gelatine. The efficiency is about 70%.

4. CONCLUSION

This project is aimed at illustrating the high potential throughput of massively parallel one-chip processors boosted by optical functions.

REFERENCES

1. U. Frisch, B. Hasselbacher and Y. Pomeau, "Lattice-Gas Automata for the Navier-Stokes Equation", *Phys. Rev. Lett.*, vol. 56, pp. 1505-1508, (1986).
2. J. Taboury, G. Chauve, P. Chavel, "An optical approach to lattice gas automata", *Proc. SPIE*, vol. 963, Topical Meeting "Optical Computing 88", pp 680-686, TOULON (FRANCE), (1988).
3. I. Seyd-Darwish, P. Chavel, J. Taboury, "Array illuminator hologram based on Talbot effect", *Topical Meeting "Optical Computing 90"*, KOBE (JAPAN), (1990).

OPTO-ELECTRONIC AUTOMATA FOR LATTICE-GAS

I. Seyd-Darwish, P. Chavel, J. Taboury
Institut d'Optique Théorique et Appliquée (CNRS)
BP 147, 91403 Orsay cedex, FRANCE

F. Devos, R. Raynaud, T. Maurin
Institut d'Électronique Fondamentale (CNRS)
Université de Paris XI, Bât. 220, 91405 Orsay Cedex, FRANCE

1. ABSTRACT

We describe an opto-electronic parallel cellular automaton for a particular application: "Lattice-gas". We use an array illuminator hologram for parallel distribution of optical informations to an electronic circuit, and suggest to use an array of MQW modulators as parallel output of the results.

2. INTRODUCTION

The electronic interconnection is the limiting factor in highly parallel computing. So optical techniques are proposed to overcome these electronic limitations. We are investigating an opto-electronic parallel processor with electronic non linear operations which offers high performances in speed and integration: optical devices are used to distribute parallel informations to the electronic chip, and MQW array to readout the results in parallel (Fig1).

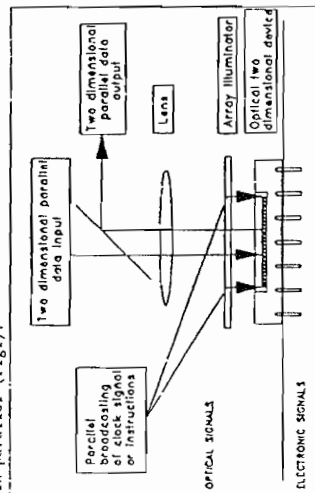


Fig 1 : Opto-electronic cellular automaton.

2. APPLICATION : LATTICE GAS

As a demonstration, we have chosen a cellular automaton specialized for the "lattice-gas" hydrodynamic algorithm. This algorithm is aimed to simulate the Navier-Stokes equation in the case of two dimensional aerodynamic flows using the statistical behaviour of particles in a gas. In this model, the gas particles move on a regular lattice (e.g. hexagonal lattice) with a constant speed, collisions between particles occur at the vertices of this lattice. Many billions of individual collision operations are required to obtain meaningful results.

We plan to perform the processing in parallel using a large number of processors distributed in an hexagonal lattice. One processor is situated at each collision site, its memory contains information on the incoming particles in the form of a 7-bit word. Symbolic substitution is used to perform the collision (Fig2).

3. ELECTRONIC CIRCUIT

The processors are implemented on a Si VLSI. Each processor contains seven 1-bit

OPTIMIZATION OF AN ARRAY ILLUMINATOR HOLOGRAM

I. SEYD-DARWISH, Ph. LEGENDRE, P. CHAVEL and J. TABOURY
 Institut d'Optique (unité associée au CNRS), ESO/
 Université de Paris Sud, BP. 147, F-91403 Orsay Cedex,
 France

Résumé - L'hologramme d'un réseau de points sources peut constituer un "illuminateur de réseau" utilisant efficacement l'énergie d'un laser pour la répartir, par exemple, entre les différents sites d'une matrice d'éléments non linéaires. Nous étudions les paramètres d'enregistrement qui permettent d'atteindre d'une part la meilleure efficacité et d'autre part le niveau d'aberration minimal en cas de changement de longueur d'onde entre l'enregistrement et la restitution.

Abstract The hologram of an array of point sources may be used as an array illuminator. We investigate the recording parameters leading to the optimal efficiency and to minimal aberrations in the case of reconstruction at a different wavelength.

1 - INTRODUCTION

The development of optical computing is related with the advent of new generations of components in the form of matrices of nonlinear optical or optoelectronic arrays. Examples include arrays of self-electrooptic effect devices, of n-i-p-i and p-n-p-n structures, of nonlinear etalons made of compound semiconductors in bulk form or in multiple quantum wells; similar geometries of components are found in electrically addressed semiconductor optical shutter arrays as well as in the much more common photodetector arrays. Another case is that of bulk materials where no pixels are defined at the fabrication stage but where it is desirable to concentrate light onto pixels upon utilization either to avoid crosstalk effects or to reach high intensity levels.

The need then arises to conserve light energy when illuminating these arrays and to avoid illuminating the inactive area between the pixels. Direct illumination by the plane or spherical wave coming from a laser often results in a considerable waste of energy: as an example, we have investigated optical clock distribution to an array of processing elements on a silicon chip; the processing elements were arranged in a square array of pitch 300 μm and each element was fitted with a 50 x 50 μm square photodetector for the optical clock signal. If a uniform optical clock beam is sent onto the chip (fig 1a), only (50 x 50 / 300 x 300) = 3 % of the light are useful, the remaining 97 % are wasted and may generate perturbation of the processor. While this example might be an extreme case, a similar situation is found when illuminating nonlinear optical element arrays.

This is why a new family of passive optical elements, named array illuminators (AI), has appeared recently. In an AI, as suggested on figure 1b, an illuminating wave (W), usually a laser beam, impinges on the component (array A) and is split in a

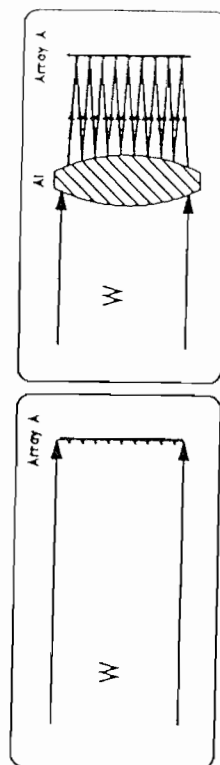


Fig. 1 (a)

Fig. 1 (b)

number of spherical wavelets, each converging onto one pixel. Several approaches have already been published in the literature and a review can be found in reference [1]. Important examples include

- the array of microlenses - possibly holographic microlenses -, where each lens corresponds to one pixel in the array; large arrays with a good efficiency (90 % or more) have been made;
- the Damann grating [2,3], which is a binary phase grating with a structure tailored to a uniform repartition of light in some set of order and a minimal diffraction efficiency in all other orders. Arrays illuminating uniformly some 20 x 20 orders with a total efficiency of the order of 65 % have been reported.

In this communication, we present an alternative holographic approach to array illumination. The principle is described in section 2. A comparison with microlens arrays and Damann gratings is presented in section 3. Sections 4 and 5 are devoted to the issues of diffraction efficiency optimization and wavelength change aberration correction, with experimental results for the former. Concluding remarks summarize the advantages and limitations of the method.

2 - PRINCIPLE

2a - General case:

Figure 2a depicts a naive AI composed of a mask with holes and a lens imaging the holes of the mask onto the pixel arrays. Except that it does not perturbate the area between the pixels, this AI has the same major limitation as direct illumination (figure 1a), namely, most of the energy is wasted: in this case, it is absorbed in the opaque parts of the mask.

Our idea is to use the setup of figure 2a not to illuminate the array, but to record an hologram using a plane (or spherical) beam coming from one side as the reference wave and the light from the mask as the object wave (figure 2b). The hologram

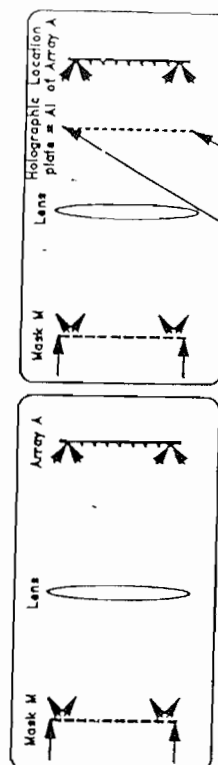


Fig. 2 (a)

Fig. 2 (b)

plane is placed between the lens and the image so that the lens is not needed upon reconstruction. The reconstruction setup is sketched on figure 3, where the AI hologram is suitably illuminated from the side and the diffracted image is projected onto the pixels. The direct "undiffracted" light through the hologram, labelled "order 0", may create a perturbation : to avoid it, either the hologram should have such a high diffraction efficiency that order 0 is negligible, or the geometry should allow order 0 to fall outside the useful area, as suggested in the figure.

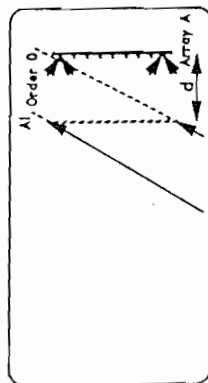


Fig. 3

2b) Case of a periodic array:

The above description applies to any distribution of pixels on array A, although, as will be discussed below, the maximal diffraction efficiency attainable will depend on that distribution. In practice, however, the arrangement of pixels in the array is often arranged in a periodic manner over a square, rectangular or hexagonal grid. While the principle of figure 2 still applies, the lens is not needed in that case since the Talbot self-imaging effect allows to form the image of the mask without a lens. This situation is shown on figure 4 and correspond to our experimental arrangement for the results described in section 4.

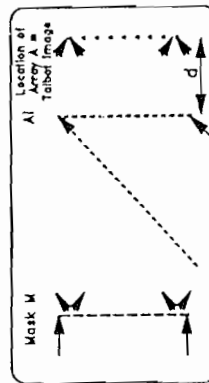


Fig. 4

3.- COMPARISON WITH OTHER ARRAY ILLUMINATORS:

Our holograms are well suited to a simultaneous recording of the complete beam forming the pixel images. The setup of figure 2b, not the one of figure 4, may be used for a sequential recording of the various pixels, but there is no clear advantage in so doing, so that we consider our principle as a collective fabrication process. Large number of pixels, in the order of 100 x 100 or more, can be recorded in one step.

In the case of Dammann gratings, the synthesis process increases considerably in complexity with the number of pixels : the calculation gets heavy and the feature dimensions get critical. For small arrays (from 2 x 1 to typically perhaps 20 x 20), no large difference should be expected between the performances of our

procedure of figure 2b and of Dammann gratings, although the fabrication is considerably different ; also, for small arrays, the setup of figure 4 is not useable because the Talbot effect, being valid for periodic structures only, is then affected with important border effects. These remarks suggest that our AI are preferable to Dammann gratings for large arrays.

Lenslet arrays and holographic lenslet arrays, on the other hand, can be manufactured in quite large arrays. In addition, in an holographic lenslet array, each holographic lenslet is the hologram of only one object point and can therefore have a very good diffraction efficiency, which is not possible with our approach because, as will be discussed in the next section, we are recording holograms of complex waves rather than just plane or spherical waves. The following two advantages of our AI compared to lenslet arrays, however, should be mentioned : collective fabrication, as already mentioned, and lower limitation by diffraction, as will be discussed now.

For the sake of compactness, ruggedness and stability, it is often desirable to keep the distance d between the array illuminator AI and the pixel array A as small as possible. Some device or architecture constraints, however, may request that an empty space be left in front of array A. Diffraction then comes into play. For example, assuming a periodic, square pixel arrangement with a spacing p on array A and an arrangement of individual lenslets with the same spacing on the AI, it is clear that the illuminating spots will be of typical half width $\lambda d/p$, with λ being the wavelength. This should be smaller than p by a factor α of the order of 2 to 10 if the AI is to serve its purpose, whence the condition $p = \sqrt{\alpha \lambda d}$. Using for example $d = 10$ cm, $\alpha = 5$ and $\lambda = 1 \mu\text{m}$, the pixel spacing should be $700 \mu\text{m}$, which is quite impractical for a large array. In our approach, light reaching each pixel of array A does not come from one small portion of the hologram, but may be diffracted by the complete hologram. Therefore diffraction hardly limits the size of spots, that is instead determined by the geometrical (or Talbot) image of the mask pixels. Very small sizes are feasible.

4.- DIFFRACTION EFFICIENCY :

For an hologram to reach an efficiency close to 100 %, one obvious necessary condition applies : it should be a phase hologram so as to avoid absorption. Accurate control of the phase profile, then, can be achieved by two means :

- adequate shaping of the surface
- or Bragg effect in a thick, index modulated medium.

Our AIs are natural holograms in essence. Photoresists and dichromated gelatin are adequate as phase media ; we have investigated the latter, which offers the possibility of a strong Bragg effect.

An additional necessary condition, however, applies to ensure that a hologram illuminated by a reconstruction wave of roughly constant modulus can be a phase hologram and show high efficiency : the diffracted wave should have constant modulus over the hologram plane, otherwise unavoidably in some regions part of the incoming energy must be absorbed or sent into other diffraction orders.

It is known that the above conditions are sufficient for holograms of plane waves or spherical waves of relatively small aperture, and that can be met by the use of dichromated gelatin. Our object wave, however, is far from a plane or spherical wave since it is the sum of many spherical waves converging onto the pixels of

array A and interfering with each other in the hologram plane. Therefore

- it is necessary to make sure that the above interference pattern should be as close as possible to a constant intensity with a phase modulation,
- even then, it is not guaranteed that a 100 % diffraction efficiency can be obtained.

Two sets of physical parameters are available to make the interference pattern as uniform as possible : the location of the hologram (parameter d of figure 3) and the relative phases of the pixels. We have investigated the optimization of parameter d in the configuration of a Talbot image (figure 4) by numerically maximizing the following uniformity parameter in the Fresnel diffraction region of mask M :

$$\left(\frac{1}{p_0} \int_0^p |E(x)|^2 dx \right)^2 \frac{1}{p_0} \int_0^p |E^2(x)|^2 dx$$

where $E(x)$ is the intensity at abscissa x.

Our parameters were the following : a square array of 90 x 90 square pixels of side 50 μm , spaced 300 μm apart, $\lambda = 488 \text{ nm}$ (a wavelength well suited for recording in dichromated gelatin), angle between object and carrier wave $\approx 30^\circ$, plates of thickness 13 μm made by dichromating Kodak 649 F plates. The corresponding Talbot distance is $p^2/\lambda = 184 \text{ mm}$ (using the fractional Talbot image of order 1/2, which is identical to the object).

The plane of maximal uniformity shows the intensity profile of figure 5 and is located 31 mm from the Talbot image along with three other intensity profiles in planes

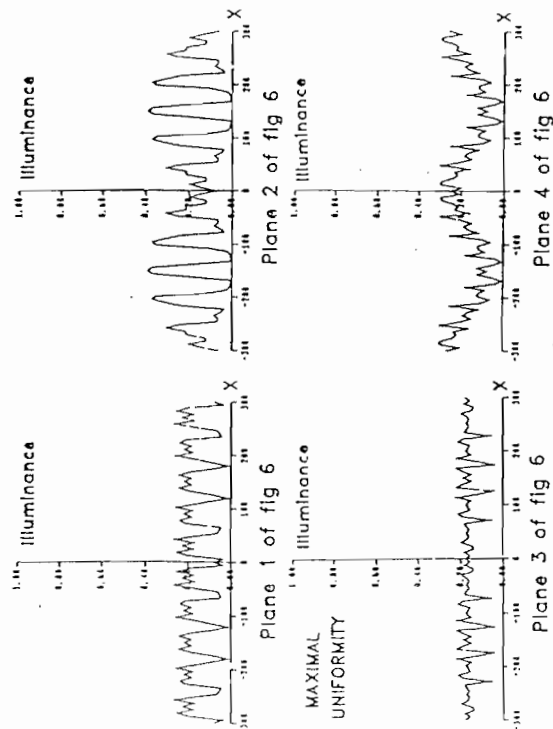


Fig. 5

nearby. The detailed calculations will be published elsewhere. An easy physical interpretation can be given for the existence of a plane of maximal uniformity. Let us consider the approximately spherical waves converging onto the pixels of array A (figure 6). Since the pixels are square, they are not spherical waves, but are instead close to the well known diffraction pattern of a square aperture, i.e. they have a sinc amplitude profile. The sinc squared intensity profiles are sketched at the pixels on the figure. The main lobes of the different pixels intersect on plane 3. It is clear that at plane 3 or closer to the pixels (plane 4), we have effectively independent waves and the Fresnel diffraction intensity profile is far from uniform, showing the individual lobes. Too far from plane 3 (plane 1), there is considerable overlap between the nodes and interference fringes appear, therefore the intensity is far

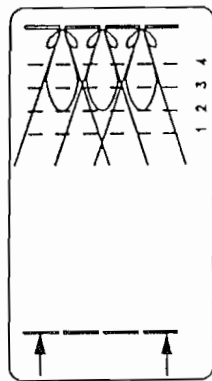


Fig. 6

from uniform also. Plane 2 is an optimum where the main lobes slightly interfere, but the interference contrast is still weak. Note that because of this overlap our hologram cannot be considered as an array of independent beam like an holographic lenslet array.

Figure 7 shows our experimental results : figure 7a is mask M, figure 7b is its Talbot image, figure 7c is the Talbot image as reconstructed by the hologram. Some noise, presumably due to the exact structure of the fringes in the thickness of the hologram, is visible. However, as already mentioned, figure 7b corresponds to only 3 % of the energy illuminating mask M, while figure 7c corresponds to 70 % of the energy illuminating the hologram (slightly more than most Dammann gratings). Figure 8 shows a microscope objective projecting a magnified image of part of the array onto a screen. The AI hologram and the reconstructing beam are visible.

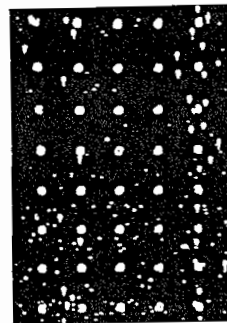


Fig. 7 (a)

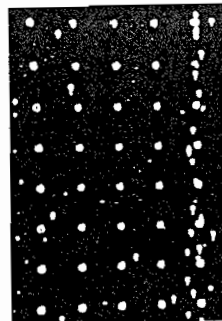


Fig. 7 (b)

(d)

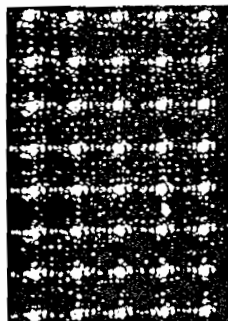


Fig. 7 (c)



Fig. 8

5 - CHANGE OF WAVELENGTH:

Another problem that deserves consideration in the case of natural holograms on dichromated gelatin is the change of wavelength between recording and reconstruction, since dichromated gelatin is sensitive in the blue and near ultraviolet region of the spectrum and most applications of AIs are in the red and near infrared region.

The change of wavelength generates aberrations that severely distort the image. It is necessary to investigate these aberrations and their minimization. We have analyzed the aberrations using the following as free parameters the incidence angles not only of the carrier wave, but also of the object wave onto the hologram (see figure 9). The difference between the two angles cannot be much smaller than 30° , otherwise the thickness effect in the $13 \mu\text{m}$ gelatin would not be strong enough to reach a good diffraction efficiency. Of course, the angles upon reconstruction are fixed by Bragg's law and by the ratio of recording to

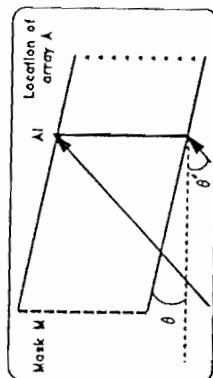


Fig. 9

The analysis will be published elsewhere. The main results are that with a reconstruction wavelength of 633 nm and the same pixel size and spacing as above, aberrations are hardly noticeable even without using the free parameters. Up to about 800 nm, aberrations are noticeable but can be reduced to a tolerable level by using the free parameters. Above 800 nm, the aberration minimization procedure is not very effective. This investigation will be pursued with other choices of parameters so as to find out if good AI holograms can be made in the critical region of 800 to about 1100 nm.

6 - CONCLUSION:

The above discussion has shown that there is no universal solution to array illumination because of the variety of pixel sizes, pixel spacing, pixel numbers, wavelength and tolerable distance between array illuminator and pixel array. Our AI holograms can reach an efficiency of 70 % and are adequate for large pixel counts. One chief advantage is simultaneous recording of all pixels. Also, it is worthwhile to note that because of the Bragg effect, the array illuminator is transparent to beams of other incidence or wavelength, allowing access of the array by multiple independent signals.

¹ N. Streibl, J. Mod. Opt. **36** (1989) 1559.

² H. Dammann and E. Klotz, Opt. Acta **24** (1977) 565.

³ U. Krackhardt and N. Streibl, Opt. Commun. **74** (1989) 31.

1 - INTRODUCTION :

Nous étudions différentes architectures pour la réalisation de machines à substitution symbolique : un grand nombre de processeurs élémentaires (PE) travaillent en parallèle. L'état de chaque PE est défini par un mot de N bits contenu dans une mémoire locale. Toute la machine évolue de façon synchrone, chaque PE suivant des règles d'évolution identiques. Ces règles sont des fonctions logiques a priori arbitraires par lesquelles l'état du PE évolue en tenant compte de l'état de P bits appelés ses voisins, et qui peuvent appartenir au même PE ou à d'autres PE. L'état des P bits voisins constitue un motif à reconnaître. Les règles d'évolution associent à chaque motif à reconnaître un motif de N bits à substituer, et qui constitue le nouvel état du PE. Il y a (au maximum) 2^P motifs possibles pour le motif présent dans le voisinage, et donc 2^P règles d'évolution.

Les PE considérés sont petits et peu puissants, l'intérêt d'une machine parallèle de substitution symbolique résulte de son parallélisme, qui est de type SIMD (une seule instruction diffusée à tous les PE, qui contiennent des données différentes). Nous nous proposons d'examiner différents modes de parallélisme massifs, de décrire pour certains des architectures possibles électroniques et optoelectroniques adaptées à la distribution des "instructions" que constituent les motifs à reconnaître et à substituer, puis à examiner les avantages et contraintes du recours à l'optique.

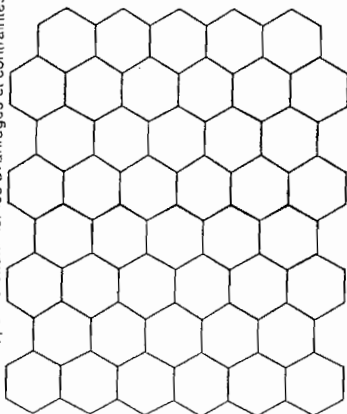


Fig 1a - le réseau hexagonal

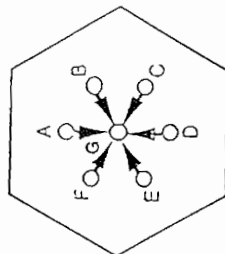


Fig 1b - la maille contient 7 noeuds, eux-mêmes disposés en hexagone.

Pour illustrer notre propos, nous choisissons de traiter le cas de la parallélisation d'un algorithme de "gaz sur réseau", destiné aux simulations d'hydrodynamique. Le cas envisagé, dont nous avons déjà discuté certains aspects^{2,3}, est une simplification du modèle décrit dans la référence 4. Il est loin de constituer un cas général du gaz sur réseau et ne constitue même pas un modèle complet du gaz pur dans un espace bidimensionnel, mais il en conserve le principe de base et se prête bien à la présente discussion.

Dans ce modèle, le plan est divisé en mailles hexagonales (fig 1a). Chaque maille contient un motif de sept noeuds qu'il est commode de représenter comme un hexagone centré ABCDEFG. En chacun des 7 noeuds peut se trouver une molécule, dont la vitesse est déterminée par sa position dans la maille (fig 1b). Dans les réalisations parallèles, un PE à $N = 7$ bits est chargé de décrire l'état et l'évolution de la maille et il y a autant de PE que de mailles. L'évolution peut être décrite de la façon décrite par la figure 2 :

- on compare l'état du PE au motif à reconnaître, c'est à dire que le "voisinage" est constitué par le PE lui-même ($P = N$) ;
- si l'état et le motif sont identiques, on effectue la substitution, le motif substitué représentant le résultat d'une "collision" qui s'est produite au noeud de la maille,
- puis on propage le motif résultant de la substitution vers les six mailles immédiatement voisines pour simuler le mouvement des molécules suivant la vitesse indiquée.

Seyd Darwish et al.

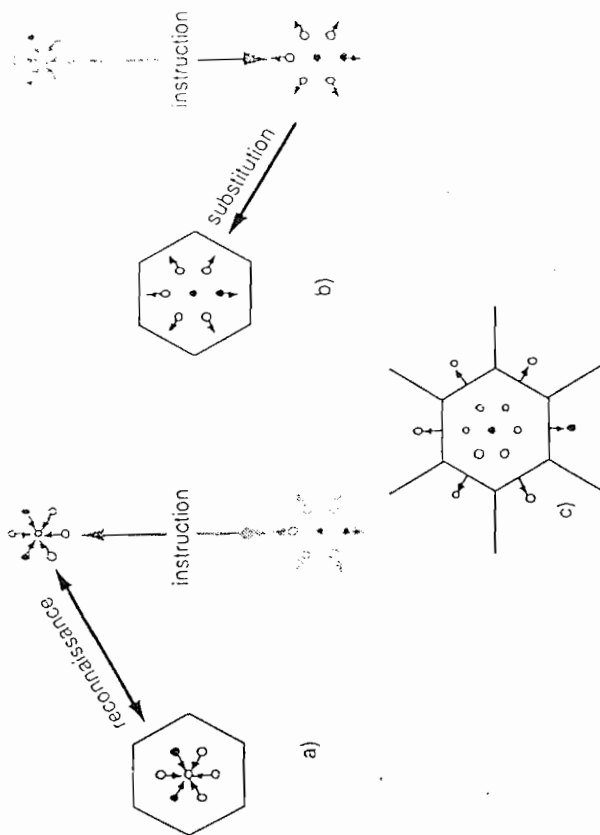


Fig. 2 Reconnaissance, substitution et propagation.

II - DIFFÉRENTS NIVEAUX DE PARALLÉLISME MASSIF

2.1 - Machines séquentielles :

Rappelons que dans une machine séquentielle, il n'y a pas de PE mais un seul processeur central : le tableau représentant l'état du réseau est examiné maille par maille. En chaque maille, on détermine le résultat de la collision par recours à une table de correspondance. La lourdeur du calcul nécessaire dans les cas pratiques justifie le désir de paralléliser le processus. Il est clair que le temps nécessaire pour une itération est proportionnel au nombre de mailles du réseau.

2.2 - Parallélisme pour les données :

Dans la version la plus compacte d'un processeur massivement parallèle qui profile du caractère spatialement invariant du traitement à effectuer, chaque PE contient son état de 7 bits et reçoit les "instructions", c'est à dire les 7 bits du motif à reconnaître suivis des 7 bits du motif de substitution, séquentiellement en 14 cycles par une seule voie de communication. Le temps nécessaire pour une itération est alors égal 14 cycles pour chacun des 2^7 motifs possibles, c'est à dire $2N \cdot 2N$ cycles.

2.3 - Parallélisme pour les données et pour les instructions :

On peut gagner un facteur N sur le temps précédent en communiquant les N bits du motif à reconnaître simultanément sur N lignes, puis les N bits du motif à substituer de même. Il faut alors dans chaque PE N unités de reconnaissance bit à bit, et l'augmentation de complexité du processeur correspond exactement au gain en parallélisme et en temps.

2.4 - Parallélisme pour les données, les instructions et le programme :

Enfin, dans la version la plus parallèle, on effectue simultanément dans chaque PE la reconnaissance des 2^N motifs possibles et la substitution du motif correspondant. L'ensemble de l'itération sur les 2^N instructions de 2N bits, qui constitue le cœur du "programme" du processeur cellulaire, est alors parallélisée à son tour. Ce type de parallélisme est souvent cité comme un but à atteindre pour la reconnaissance de

Seyd Darwish et al.

(c)

formes tout optique multicanaux. Par contre, la complexité du PE résultant est extrêmement lourde pour les approches tout électroniques ou optoélectroniques parallèles au niveau de la puce.

Dans ce qui suit, nous comparons les solutions tout électroniques et optoélectroniques pour les processeurs des paragraphes 2.2. et 2.3 qui correspondent à nos travaux actuels, la part confiée à l'optique dans les solutions optoélectroniques étant la diffusion des instructions.

III - AVANTAGES ET HANDICAPS DES SOLUTIONS OPTOÉLECTRONIQUES.

3.1 - Principe :

La figure 3 évoque les connexions de distribution d'instructions dans le cas 2.2 (parallélisme pour les données uniquement). Cette figure suppose une symétrie carrée et non hexagonale comme dans la discussion précédente. Dans la version tout électronique (fig. 3a), un plot d'entrée fournit un signal qui est divisé en K branches (ou lignes) de K connexions, où K^2 est le nombre total de PE. Dans la version optoélectronique (fig. 3b), chaque PE est muni d'une photodiode et le plot d'entrée est remplacé par une diode laser qui éclaire un "illuminateur de tableau" de façon que son image démultipliée K^2 fois se forme sur les photodiodes.

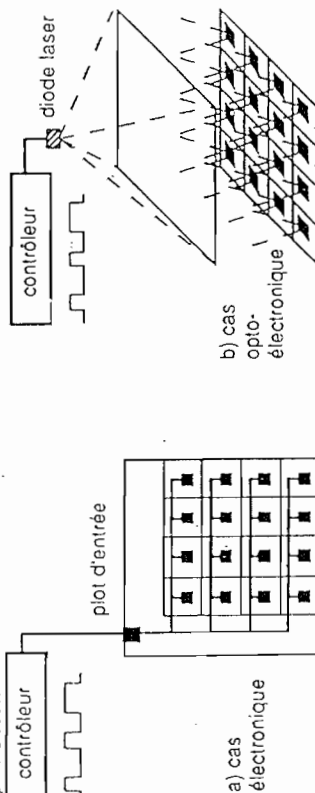


Figure 3. Diffusion des instructions.

Dans le cas du §2.3, il suffit de remplacer le plot d'entrée et ses K^2 connexions par N plots d'entrées avec K^2 connexions chacun, la diode laser par N diodes laser et la photodiode de chaque PE par N photodiodes disposées suivant une figure homothétique de celle des diodes laser.

Dans ce qui suit, nous envisageons des puces réalisées en CMOS de résolution de l'ordre de $1\ \mu\text{m}$ et comportant $K \times K = 100 \times 100$ PE d'environ $100 \times 100\ \mu\text{m}^2$. Ces chiffres ne sont que des ordres de grandeur puisque la taille exacte des PE, et par conséquent le nombre qu'il est raisonnable d'intégrer dépend de l'architecture envisagée (optoélectronique ou optique, parallélisme du §2.2 ou du §2.3). Les ordres de grandeur utilisés pour les capacités sont empruntés à des publications récentes^{5,6,7}.

3.2 - Consommation d'énergie :

Calculons l'énergie électrique nécessaire pour amener à chaque PE un bit d'instruction sous la tension $V = 5$ volts : $E = CV^2$, où C tient compte de l'ensemble des capacités qui interviennent ; envisageons d'abord le cas tout électronique.

- La capacité du plot d'entrée, de l'ordre de 400 fF, est à répartir entre les K^2 processeurs ; pour K de l'ordre de 10000 μm , elle est de l'ordre de 700 fF ;
- la capacité de grille du transistor d'arrivée sur chaque PE est de l'ordre de 20 fF.

Au total, on arrive pour le cas tout électronique à une valeur de C de l'ordre de 30 fF. Dans le cas optoélectronique, les capacités de plot et de ligne disparaissent mais sont remplacées par la capacité de la photodiode. Plus la photodiode est petite, plus sa capacité est faible, mais plus les problèmes d'alignement et de lumière parasite l'ont en dehors de la surface appropriée sont importants. Un choix optimiste correspond à une photodiode carrée de côté $3\ \mu\text{m}$, dont la capacité est 10 fF, ce qui mène encore à environ 30 fF au total. Les rendements de conversion de la diode laser et des photodiodes et les pertes de l'illuminateur de réseau⁸ interviennent donc en délaissant de l'optique sans qu'il y ait de gain en énergie à attendre par ailleurs.

Dans cette discussion sur l'énergie, rien n'est à changer pour le cas des architectures du §2.3 avec distribution parallèle des N bits d'instruction.

3.3 - Surface de silicium :

Dans la solution tout électronique, chacune des K lignes de connexion occupe au moins K fois la largeur minimale imposée par les règles de dessin de la technologie utilisée. En tenant compte de la surface nécessaire pour le plot et en répartissant l'aire totale trouvée entre les K^2 PE, on arrive à environ $100\ \mu\text{m}^2$. Du côté optoélectronique, seule intervient la photodiode, d'aire $10\ \mu\text{m}^2$. On constate donc un certain avantage, modéré toutefois par le fait que les connexions considérées ne concernent qu'une fraction modeste de la surface totale du PE, supposée de l'ordre de $100 \times 100\ \mu\text{m}^2$. Cependant, ce gain est N fois plus accusé dans le cas de la distribution parallèle des N bits d'instruction.

3.4 - Quelques critères non quantitatifs :

Si la solution optoélectronique procure un certain gain en surface de silicium, la compacité globale et sa facilité de mise en oeuvre subissent l'entrave de l'illuminateur de tableau et de son alignement. Nous avons présenté par ailleurs le type d'illuminateur que nous avons retenu⁹. Le développement de technologies de microoptique plus performantes s'impose pour résoudre cette difficulté.

D'un autre côté, un point à souligner en faveur de l'optoélectronique est la tolérance accrue aux défauts locaux. Un défaut sur l'une des lignes de connexions électroniques rend inutilisable toute la ligne de PE. Un défaut local sur une photodiode n'affecte qu'un PE. Si l'application envisagée et le nombre total de PE permettent une certaine tolérance aux défauts locaux, la solution optoélectronique est privilégiée. La diminution du nombre de niveaux d'interconnexions nécessaires dans le cas de la distribution parallèle de N bits d'instructions souligne encore cet avantage.

IV - CONCLUSIONS :

Pour la diffusion des signaux d'instruction, l'avantage apporté par l'optoélectronique à l'intégration massivement parallèle sur puce de processeurs de substitution symbolique n'est pas décisif si on ne prend en compte que l'énergie. Il devient plus net si l'on prend en compte la surface de silicium et la fiabilité, et l'on retrouve donc bien sûr les conclusions que l'on peut également tirer pour la distribution de signaux d'horloge¹⁰. Le gain est sensible dans le cas des PE relativement complexes à distribution en parallèle de tous les bits d'instruction.

Remarquons pour terminer que notre discussion s'est limitée à un des rôles de l'optique pour le fonctionnement des puces massivement parallèles, la diffusion d'instruction ; son avantage pour l'entrée des données (et non des instructions) par projection d'une image sur les photodétecteurs de la puce est indiscutable puisque les connexions électroniques ne peuvent en aucune façon fournir un tel flot d'entrées en parallèle. L'argument analogue pour la sortie des données sera applicable à si l'on parvient à associer de façon efficace une technologie de traitement électronique dense comme celle du silicium à une technologie de modulation de lumière. Le programme MOTS (matrices optoélectroniques pour le traitement du signal) du Ministère de la Recherche et de la Technologie, dans le cadre duquel s'inscrivent en partie les travaux que nous avons décrits, est un des essais dans ce sens.

RÉFÉRENCES :

- 1 - A. Huang, Proc. 10th IOCC, Cambridge (Mass., USA), IEEE, pp 13-17 (1983).
- 2 - J. Taboury, C. Chauve, P. Chavel, OC'88, Toulon, SPIE Proc 953, 681-686 (1988).
- 3 - I. Seyd-Darwish, J. Taboury, P. Chavel, IO'15, Garmisch-Partenkirchen (RFA), SPIE Proc. 1319, 173-174 (1990).
- 4 - U. Frisch, B. Hasslacher, Y. Pomeau, Phys. Rev. Lett. 56, 1505-1508 (1986).
- 5 - M. R. Feldman, S.C. Esener, C.C. Guest, S.H. Lee, Appl. Opt. 27, 1742-1751 (1988).
- 6 - H. Rodriquez, thèse, Université de Paris Sud, 1988.
- 7 - J. Millman, Microelectronics, Mc Graw Hill, 1989.
- 8 - I. Seyd-Darwish, thèse, Université de Paris Sud, 1991.
- 9 - I. Seyd-Darwish, Ph. Legendre, P. Taboury, P. Chavel, 6èmes journées d'étude sur les fonctions optiques dans l'ordinateur, Strasbourg, septembre 1990, Annales de Physique, sous presse.
- 10 - Th. Maurin, T. Porcher, S. Reynaud, D. Berchandy, F. Devos, 5èmes journées d'étude sur les fonctions optiques dans l'ordinateur, Toulouse, octobre 1989, actes, pp 181-186.

ARCHITECTURE D'UN SYSTEME OPTO-ELECTRONIQUE PARALLELE.

Groupe "PHYSIQUE DES IMAGES" de
l'Institut d'Optique Théorique et Appliquée - URA-14
Unité de Recherche Associée au CNRS
J. Taboury, P. Chavel et I. Seyd-Darwish

Le développement des circuits intégrés et des circuits imprimés ou hybrides impose l'intégration d'un grand nombre de connexions électriques qu'il est de plus en plus difficile de satisfaire lorsque l'on cherche à mettre en oeuvre sous une forme réellement optimum et efficace des algorithmes de traitement parallèles. Par ailleurs, l'industrie électronique spécialisée dans la fabrication et l'intégration de systèmes électroniques de pointe de type avionique, recherche de plus en plus une alternative optique pour résoudre d'innombrables problèmes de compacté.

Dans les thèmes de recherche du groupe de "Physique des images" de l'Institut d'Optique Théorique et Appliquée (I.O.T.A.) nous avons abordé cette réflexion depuis plusieurs années.

La collaboration de notre laboratoire avec d'autres organismes publics et privés tels que l'Institut d'Electronique Fondamental (I.E.F.) d'Orsay, le L.A.A.S. de Toulouse, le C.N.E.T. de Bagnoux et le Laboratoire Central de Corbeville de la (Thomson CSF), nous a permis de développer les principes d'une architecture opto-électronique dont nous présentons certains éléments à titre d'illustration.

Principe de la micro architecture parallèle proposée.

Le principe d'une telle architecture repose sur l'idée que le nombre de connexions électriques d'un circuit intégré classique est fondamentalement limité par la technologie utilisée. En effet, les connexions électriques sont couramment assurées par des fils d'or entre la périphérie de la puce et les picots métalliques du boîtier.

Une solution d'optimisation consiste à fournir au circuit un moyen de communiquer avec l'extérieur au travers d'un réseau supplémentaire et dense de connexions optiques pouvant assurer un fonctionnement parallèle à sa logique électronique. Ces connexions ne sont pas localisées sur la périphérie du circuit VLSI, mais directement sur une fraction de sa surface par l'intermédiaire de photodiodes et de modulateurs.

- Distribuer les cycles d'horloge en parallèle sur l'ensemble des unités microscopiques de calcul,
- fournir à l'ensemble de ces unités de calcul élémentaires les données sous une forme matricielle (image) et non vectorielle ainsi que
- ressortir dans le même format les résultats d'un calcul ou d'une itération en parallèle : tel est l'objectif que nous nous sommes fixés.

Le schéma de la figure 1 donne l'allure générale du principe expérimental retenu.

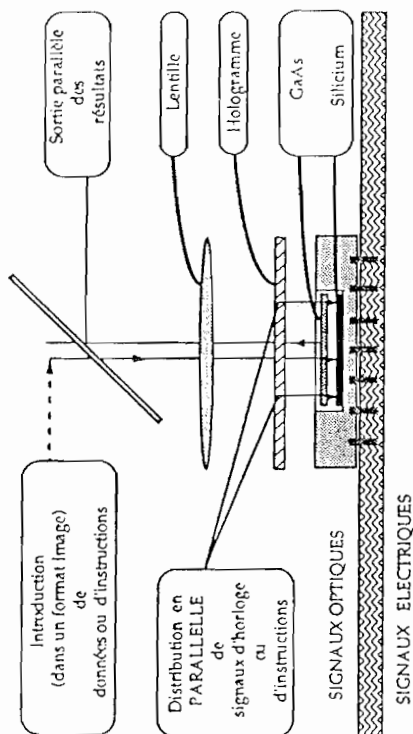


Figure 1

La troisième dimension optique

Ce principe a été développé pour permettre de tester sous une forme autonome les différentes fonctionnalités de cette architecture. La partie logique associée à l'algorithme utilisé² a été réalisée en collaboration avec l'IEF³ sur silicium en technologie CMOS. Malheureusement, le gap indirect du silicium ne confère pas à cette substance la propriété d'émettre de la lumière, on ne peut que recevoir au travers d'un réseau de photodiode un certain nombre de données codées optiquement sans pouvoir redistribuer sous une forme équivalente la matrice des résultats calculés par le réseau des micro-unités logiques. Aussi avons nous choisi:

1. Thèse de Iyad Seyd Darwish IOTA-Paris XI dec91
2. Références relatives à l'algorithme utilisé : "Lattice gas automata for Navier-Stokes equation", U. Frisch, B. Hasslacher, Y. Pomeau, Phys. Rev. Lett. 56,1505-1508, (1986). "An Optical approach to lattice gas automata", J. Taboury, C. Chauve et P. Chavel, Optical Computing 88, Toulon. "Automate cellulaire opto-électronique pour le gaz sur réseau", Colloque National de Visualisation et de Traitement des Images, LILLE, 29 mai - 1 juin 1990, IMFL-ONERA. "Opto-electronic automata for lattice-gas" I. Seyd Darwish, P. Chavel, J. Taboury et Y. Maliet, ICO-15, Garmisch-Partenkirchen (RFA), 5-10 août 1990, ICO, (1990), 173-174.

3. La collaboration de notre groupe avec le groupe "AXIS" de l'I.E.F. a permis la réalisation d'un circuit prototype en technologie CMOS muni de 7 photodiodes. Une seule unité de calcul a été implantée sur le circuit CMOS afin de valider certain concept technologique. En effet, la partie logique fonctionne sous une alimentation classique de 5 volts alors que la matrice accrue de modulateurs à puits quantiques fabriquée par la Thomson LCR requiert une tension de modulation comprise entre 0 et 20 volts pour assurer un taux de modulation correct. La version actuelle ne convertit les signaux TTL (0-5V) que sur une dynamique 0 - 10 volts.

(f)

d'allier au silicium les propriétés modulantes d'un réseau de Bragg à puits quantiques en GaAs développé dans le cadre du programme national MOTS⁴. Une matrice de 7 modulateurs répartis sur une maille hexagonale et fonctionnant par réflexion a été réalisée pour cette application par le Laboratoire Central de Recherche de THOMSON-CSF. Dans une version ultérieure, la matrice de modulateurs devra être connectée directement sur le circuit CMOS au travers d'un réseau de micro-billes d'indium selon une technologie déjà éprouvée pour les caméras CCD-IR.

L'existence de deux gaps différents pour le silicium et l'arséniure de gallium est un avantage pour un tel système si l'on a recourt à deux longueurs d'onde : la première pour l'écriture au travers du composant en GaAs des instructions et des données sur le circuit en silicium, l'autre, accordé sur la longueur d'onde d'utilisation du modulateur pour la sortie en parallèle des résultats. Si la seconde est imposée par la structure à puits quantiques du modulateur (860nm), la première longueur d'onde doit être supérieure à celle du gap de GaAs et a été choisie à 999 nm.

L'énergie nécessaire à l'écriture des données et instructions est fournie par une diode laser fonctionnant à 999 nm et fabriquée par le LAAS de Toulouse⁵.

La distribution en parallèle des signaux communs aux unités logiques du réseau implantés sur le circuit CMOS est effectuée grâce à l'hologramme d'une matrice de micro-lentilles⁶. Cet hologramme épais forme autant d'images d'une source ponctuelle qu'il y a d'unités logiques sur le circuit. Dans notre exemple 7 photodiodes sont nécessaires par site de calcul, aussi chaque micro-lentille en regard d'un site forme l'image de 7 diodes laser émettant sur la longueur d'onde pour laquelle l'hologramme présente une efficacité maximum. La particularité de cet hologramme repose sur le fait que le faisceau de lecture des résultats à la deuxième longueur d'onde ($\lambda_2=860\text{nm}$) ne présente aucun accord de phase avec le réseau épais. Dans ces conditions, l'hologramme se comporte comme une simple lame à face parallèle pour ce deuxième faisceau. La matrice de micro-lentilles est en fait réalisée d'une façon originale en ayant recourt à l'effet Talbot. Lorsque l'on éclaire une matrice de petits trous répartis aux noeuds d'un réseau régulier par un faisceau laser, il se forme une image de cette même matrice à une distance déterminée par la longueur d'onde et le pas du réseau. En plaçant la plaque holographique à une distance particulière de ce plan image, on peut ainsi enregistrer l'état d'un front d'onde correspondant à quelques détails près à celui que fournirait une matrice de micro-lentilles. L'efficacité de ce composant atteint couramment 80%. Son principal avantage réside dans la parfaite homogénéité en amplitude du front d'onde holographié. Un calcul d'aberration rendu nécessaire par le fait que la longueur

4. Programme MOTS : "Matrice optoélectronique pour le traitement du signal" programme (sélecteur en optique soutenu par le MRT mené en collaboration avec l'OTIA, IEF, CNET Bagnieux, LAAS de Toulouse et le Laboratoire Central de Recherche de THOMSON-CSF.

5. Le LAAS de Toulouse a fabriqué récemment une diode laser fonctionnant à 999 nm. Les caractéristiques de cette diode laser ont été imposées pour une application toute différente de la notre, elles correspondent aux impératifs liés à l'amplification de signaux optiques au sein de fibres optiques dopées en terre rares. Ces recherches sont effectuées au CNET de Bagnieux.

6. "Optimisation of an array illuminator hologram", I. Seyd Darwish, P. Chavel, J. Taboury et Y. Mallot, 6èmes journées d'études sur l'optique dans l'ordinateur Strasbourg, 6-7 septembre 1990, Annales de Physique, Colloque n°1, vol.16, février 1991, 103-110.

d'onde d'utilisation ($\lambda_2 > 860\text{nm}$) est beaucoup plus grande que la longueur d'onde utilisée à l'enregistrement ($\lambda = 488\text{nm}$ ou 514nm) a permis d'en optimiser la fabrication⁷.

Perspectives.

Une telle architecture présente, bien entendu, des avantages et des inconvénients. Si le principal avantage est sans nul doute l'augmentation du nombre potentiel de connexions sur le monde extérieur par l'utilisation rationnelle d'une matrice de photodiodes intimement liée au circuit logique et donc la possibilité de traiter en parallèle l'ensemble des données, un inconvénient majeur subsiste actuellement dans le fait qu'une telle architecture ne peut être utilisée en cascade. La solution serait de disposer la matrice de modulateurs sur la face opposée à celle des photodiodes pour pouvoir utiliser la même longueur d'onde en entrée et en sortie. Les difficultés technologiques rencontrées pour développer un tel composant ne peuvent être actuellement résolues. Une alternative serait de transposer toute la logique du traitement directement sur le substrat de GaAs et de répartir judicieusement photodiodes et modulateurs sur les deux faces du substrat.

En ce qui concerne l'évolution des composants optiques nécessaires au développement de telles architectures, la matrice de micro-lentilles (sous forme holographique pour la longueur d'onde λ_1 (lame à face parallèle pour λ_2) souffre non seulement de l'inévitable difficulté liée à la précision requise par les contraintes des alignements optiques, mais aussi de l'accès unilatéral des photodiodes et des modulateurs. Outre la solution proposée précédemment, l'alternative qui potentiellement peut résoudre cette difficulté que de nombreux systèmes ne peuvent supporter pour une simple raison de fiabilité, repose sur le développement de cartes hybrides à plusieurs niveaux de pistes transparentes, électriques et optiques comprenant des guides d'onde planaires, des coupleurs à réseau et des matrices de micro-lentilles intégrées. Une telle solution permettrait d'approcher la solution monolithique comme le schéma de principe de la figure suivante le laisse supposer.

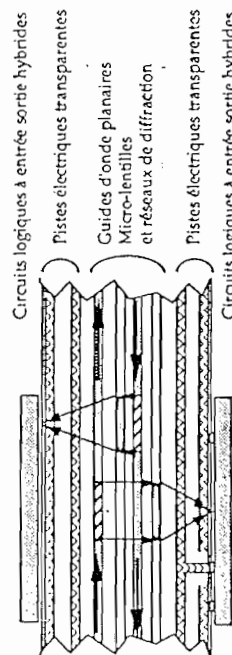


Figure : 2

7. L'hologramme présenté sur le stand a été enregistré à 488nm pour une utilisation à 633nm.

performs simple operations like memorizing, comparing and copying binary words. The interconnection between neighbouring EPs are made electronically because they are short interconnections. The required instructions that are performed simultaneously by all EPs are projected optically using modulated light sources and an array illuminator (3). The initialization of the EPs is made optically by projection of an image of the data directly onto the EPs (2). The results are read out in parallel using another chip made of optoelectronic modulators formed of pixels connected to the EPs and project optically an image of their contents onto an external component (4) [2].

3- LATTICE-GAS ALGORITHM

"Lattice-Gas" is an algorithm for the simulation of the Navier-Stokes equation in the case of aerodynamic flows using the statistical behaviour of particles in a gas [3]. Although three dimensional versions exist, we consider here a two-dimensional case.

In this model, the gas particles are restricted to move only on a hexagonal lattice. Their speed modulus can take only two values, zero and unity. Unity corresponds to a motion between two cells covered in the time of one iteration of the algorithm. Collisions between the particles occur at the lattice vertices. These collisions follow laws that respect the conservation of the particle number, the momentum and the energy of the gas. The laws are given in a table (lookup table, LUT) and the collision at each vertex is performed by symbolic substitution using that table.

With the symbolic substitution the execution of the algorithm contains three steps at each vertex (Fig. 2) :

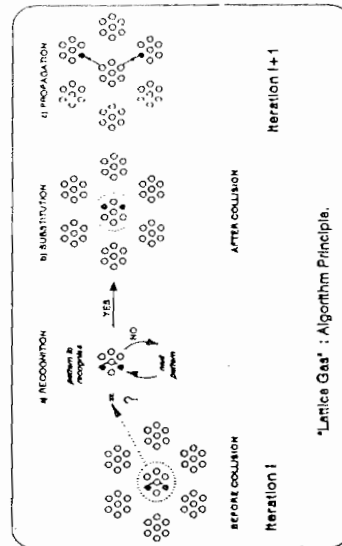


Fig. 2: Treatment steps.

Recognition step: A pattern of incoming particles is present at each vertex. The first step consists of recognizing this pattern by comparing it sequentially with all the possible patterns. As we use a hexagonal lattice, the maximum number of incoming particles at each vertex is 7 (including the one that may be present at the vertex), so we have $2^7 = 128$ comparisons to make.

Substitution step: When an incoming particle pattern is recognized at a vertex, we substitute it by the corresponding outgoing particle pattern resulting from the collision.

OPTICAL INPUT AND OUTPUT FUNCTIONS FOR A CELLULAR AUTOMATON ON A SILICON CHIP

Ilyad SEYD-DARWISH, Pierre CHAVEL, Jean TABOURY
Institut d'Optique Théorique et Appliquée (CNRS), Orsay, France.

Francis DEVOS, Thierry MAURIN, Roger REYNAUD
Institut d'Electronique Fondamentale (AXIS Group, CNRS), Orsay, France
With the contribution of : Thomson C.S.F. (L.C.R.) and LAAS (Toulouse).

1- INTRODUCTION

Parallel architectures for cellular automata are adapted to optical computing. Electronic implementations have the interconnection as limiting factor. Our purpose in the present research is to couple the performances of electronics which offers nonlinear components with high integration and speed, and those of optics which allows the implementation of dense and rapid interconnections (specially for the input and output), all that in a compact optoelectronic cellular automaton setup. The specific case that we selected for the purpose of illustration is the "Lattice-Gas" [1] automaton.

In this article we present the architecture of the cellular automaton implementation. After a brief description of the "Lattice-Gas" algorithm, we introduce the implementation of the different parts of this automaton.

2- CELLULAR AUTOMATON ARCHITECTURE

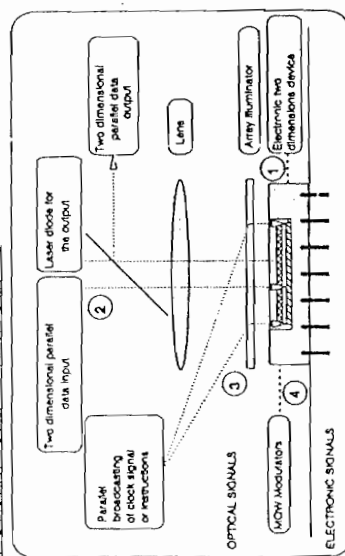


Fig. 1 : Architecture of the cellular automaton.

The cellular automaton contains a number of elementary processors (EPs) distributed on a regular lattice. These EPs perform nonlinear operations in parallel, the same for all the EPs, each one is connected to some of their neighbours. It is therefore one SIMD machine (single instruction, multiple data).

In our case (Fig.1), The EPs are implemented electronically on a silicon chip (1). Each EP

Propagation step: After 128 recognition/substitution steps, all vertices contain their outgoing particle patterns. We then propagate these particles to the neighbouring vertices and they will constitute the new incoming particles pattern for the next iteration.

Those three steps constitute one iteration. To have meaningful results of "Lattice-Gas", we need many thousands of iterations.

4. IMPLEMENTATION ON CELLULAR AUTOMATON

The electronic circuit contains EPs which contain the distribution of the particles and perform the collision. The LUT is memorized in an external memory named the "correspondence memory". This memory contains all the 128 possible incoming particle patterns, each pattern is followed by the appropriate outgoing particle patterns. The contents of this memory is projected optically and sequentially onto all the EPs of the electronic circuit.

4.1 ARRAY ILLUMINATOR BASED ON TALBOT EFFECT

The array illuminator (AI) broadcasts the same instruction to all the EPs on the electronic circuit. These instructions are words of 7 bits.

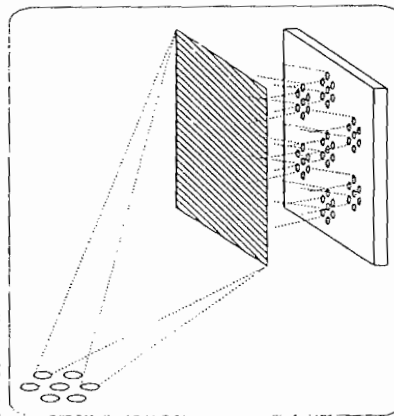


Fig. 3 : Array illuminator.

This AI illuminated by 7 sources images them onto each EP. As shown in figure 1, in addition to the instruction broadcast, two other optical signals have to access to the chips through the AI. For this reason, a diffractive AI with wavelength or angular selectivity is necessary (Fig.3). We have implemented a new process to make a diffractive AI; it is based on the Talbot effect [4].

a) Array illuminator principle

The Talbot effect [5] produces an image of regular gratings, illuminated by coherent light, at periodic distances multiple of the Talbot distance ($2p^2/\lambda$) (p is the grating period and λ is the wavelength). In addition, a shifted version of the grating image is reproduced at half the Talbot

distance.
This effect could be used directly for an array illuminator. However the efficiency is then very poor : for example, in the case of a square grating of $300 \mu\text{m}$ with square holes of $50 \times 50 \mu\text{m}^2$, only 3% of efficiency can be expected, just because of the duty cycle of the grating. These numerical values correspond to a typical test case that we developed in reference [4].
To improve this efficiency we record an hologram at some appropriate distance between the grating and its first Talbot image.

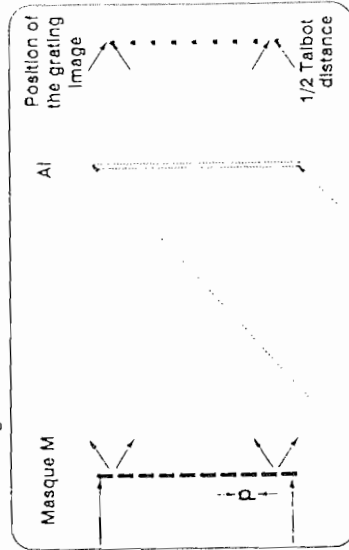


Fig.4 : Holographic array illuminator principle.

The position of this holographic plate is optimized according to the holographic techniques. To have the maximum gain of efficiency we have to use thick phase holograms. One important condition to have maximum efficiency is that the wave recorded at the plate shows a uniform intensity and all information must be encoded in phase modulation. This condition is hardly verified with a wave diffracted from a grating such as the one described above. We have studied the diffracted wave from such a grating in the Fresnel approximation to minimize the intensity variation at some diffraction plane where the holographic plate will be placed. Figure 5 shows the intensity variation at the optimal plane located, in our test case, at 3.1 cm from the half Talbot distance using a wavelength of 488 nm.

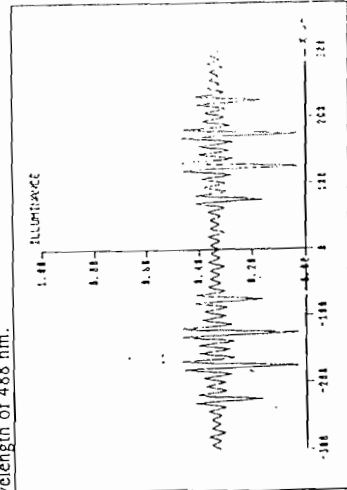


Fig. 5 : Minimum intensity modulation plane.

(g)

(i.e. the maximal number of particles at each vertex). Each unit represents one direction of the hexagonal lattice. One unit is composed of one input memory of 1 bit, one output memory of 1 bit and one photodiode. In addition, the EP contains one recognition unit.

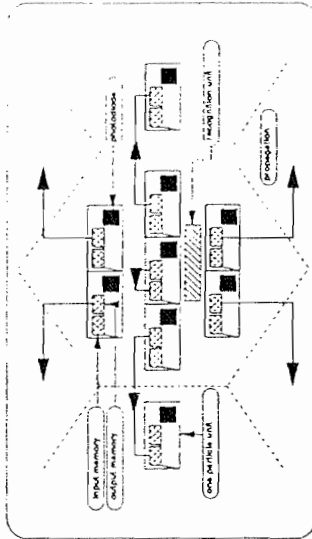


Fig. 7 : Electronic circuit.

In the set of the 7 units, the incoming particles pattern is first memorized. In the recognition step, the recognition pattern is projected on the 7 photodiodes. The EP compares the input memory contents with the illumination state of the photodiodes. In the case of an exact match, the recognition unit is set to logical level "1". Otherwise this unit stays at logical level "0". For the substitution step, the substitution pattern is projected just after the recognition step. If the recognition unit is at logical level "1" the illuminating state of the photodiode is transferred to the output memory, otherwise no copy occurs and the EP waits for the next pattern.

We repeat the projection of recognition and substitution patterns sequentially 128 times. At the end of this cycle the output memory contains the substitution pattern corresponding to the initial incoming pattern in the input memory. The propagation step is made by copying the contents of each output memory onto the input memory of the neighbouring EPs.

Interface problem

Two interface problems are encountered with the electronic circuit :

a) Input interface : The response time of the integrated photodiodes must be adapted to the operation frequency of the Silicon components.

We have modeled the sensitivity of the integrated photodiode knowing that the structure of the photodiode is a reverse biased P-N junction with an internal capacity. Absorption of photons generates electron-hole pairs which discharge this capacity and switch the voltage of the photodiode between two logical levels. The commutation time depends on the photodiode sensitivity as a function of illuminating wavelength and of the optical power.

Imposing a commutation speed of 100 mV/ns, we have deduced the optical power appropriate for the photodiode illumination light (Fig. 8 c).

We have recorded a thick phase hologram on dichromated gelatin using the grating described above with 90-90 holes. The recording wavelength was the 488 nm of Argon laser. The holographic plate was located at the optimal plane. The reconstruction of this hologram with the same wavelength reproduces the grating image with 70% of efficiency.

The study shows that the structure of the hologram is similar to a microlenses hologram. Bragg condition for the diffracted wave offers transparency for the other wavelengths.

The advantage of this technique is the simultaneous recording of a large number of microlenses.

b) Change of wavelength

The recording wavelength is imposed by the DCG sensitivity. In the automaton setup we use laser-diode for reasons of compacty. As the wavelength of these lasers is situated in the near infrared, the hologram will present chromatic aberrations.

We have conducted a study to minimize the influence of these aberrations by a simple change of the recording conditions, specially the incidence angle of the diffracted wave on the hologram plate. We take the energy concentration in the $50 \times 50 \mu\text{m}^2$ detector area as the correction criterion.

The simulation shows that for a given angle the influence of the aberration can be minimized. Figure 6 shows the simulation of this correction compared with the case of no correction. We observe that the energy concentration is practically doubled.

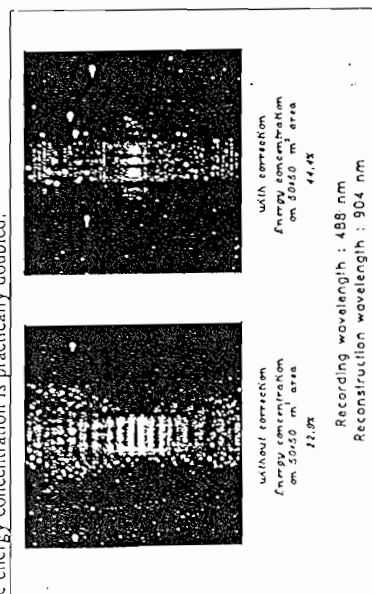


Fig. 6 : aberration correction.

4-2 ELECTRONIC CIRCUIT

The electronic demonstration circuit contains EPs distributed on an hexagonal lattice, each hexagone represents a collision vertex in the "Lattice-Gas" model. The function performed by each EP is reduced to the strict minimum i.e. memorizing of particles patterns (incoming particles and resulting particles from the collision), comparison and copying. In this way we can reduce the area occupied by each EP so as increase the number of these processors on one chip.

Figure 7 shows the schematic diagram of one elementary processor which is composed of 7 units

(g)

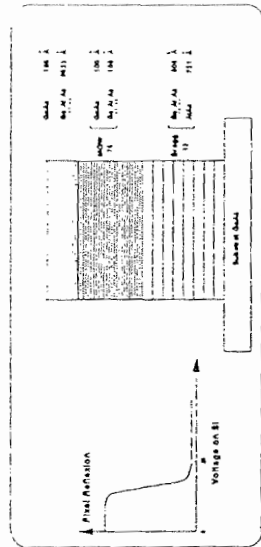


Fig. 10: MQW modulators.

In the final setup of the automaton the GaAs circuit containing the modulators would be put on the Silicon circuit and bonded by spherical Indium balls.

5- CONCLUSION

We have studied the implementation of a parallel optoelectronic automaton for a specific algorithm, the "Lattice-Gas". We have studied and implemented three components for this automaton :

- a microelectronic circuit : for the treatment
- a Talbot effect hologram : for instructions broadcast.
- Optoelectronic modulators : for parallel readout.

This automaton exploits the performances of the electronics for nonlinear operation and the optics for the input and output interconnection.

A further step in our plane is to connect all these components in a complete setup. We have estimated the performances of this setup for 2 μ m CMOS technology (french MCP process) :

- 15000 iteration per second
- 0.2W power dissipation for 10^2 EPs
- High connectivity about $1.68 \cdot 10^{10}$ bits/s/cm²

Some generalization of the use of such an automaton for other algorithms using symbolic substitution is possible. It can be made by simply changing of the contents of the lookup memory.

References

- [1] I. Seyd-Darwish, Thèse de doctorat en sciences, Université Paris XI, Orsay, France (1991).
- [2] I. Seyd-Darwish, P. Chavel, J. Taboury, F. Devos, T. Reynaud, T. Maurin, Optics in Complex Systems, Garmisch, Germany, SPIE 1312, 173 (1990).
- [3] U. Firsch, B. Hasslacher and Y. Pomeau, Phys. Rev. Lett. 56, 4505 (1986).
- [4] I. Seyd-Darwish, P. Chavel, J. Taboury, Annales de Physique, Colloque n°1, Supplément au n°1, 16, 103 (1991).
- [5] F.A. Talbot, Philos. Mag. 2, 401 (1836).

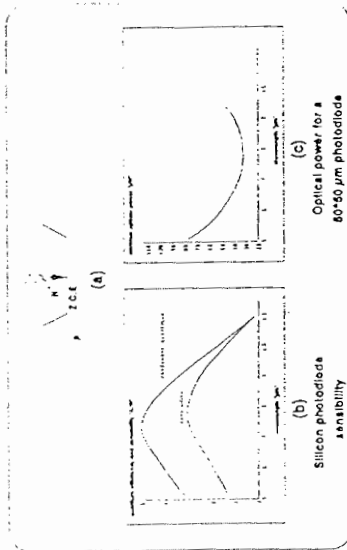


Fig. 8 : Integrated photodiode sensitivity.

Output interface : Optoelectronic modulators implemented on a GaAs circuit (see below) are used for parallel output. For a good contrast between logical levels "0" and "1", these modulators require a 10 V voltage change and a drive current of about 100 μ A. The Silicon electronic circuit works with a logical level "1" of about 5 V and a current of a few μ A only. We have implemented specific pads which can increase the logical level "1" from 5 to about 10 V using two transistors (T1, T2) (Fig. 9) and are protected with two inverters. The delivered current is controlled by image current transistor T3.

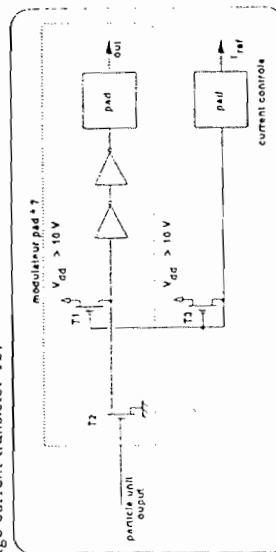


Fig. 9 : Modulator output pads.

4-3 OPTOELECTRONIC MODULATORS

Optics offers the possibility of parallel output. We use MQW optoelectronic modulators distributed on an hexagonal lattice, arranged in units of 7 pixels each. One pixel is composed of Multiple Quantum Wells above Bragg reflectors. Applying an electric field modifies the reflexion state of the pixel at one determined wavelength.

Each pixel reads the contents of one output memory of one EP by one to one connection between the pixels and the 1 bit output memories. The voltage of each output memory drives the reflexion state of the corresponding optoelectronic modulator pixel.

BIBLIOGRAPHIE

REFERENCES

Première partie

- [1] J. Wang, "Processeurs cellulaires optiques : architectures et réalisations", thèse de l'Université de Paris XI (1989).
- [2] P. Chavel, J. Taboury, "Binary optical cellular automata : concept and architectures", Proc SPIE, **CR35**, 245 (1990).
- [3] J. Taboury, J.M. Wang, P. Chavel, F. Devos and P. Garda, "Optical cellular processor architecture. 1: Principles", Appl. Opt.**9**, 1643 (1988).
- [4] J. Taboury, J.M. Wang, P. Chavel and F. Devos, "Optical cellular processor architecture. 2 : Illustration and system considerations", Appl. Opt. **28**, 3138 (1989).
- [5] S. Wolfram, "Theory and application of cellular automata", Word Scientific, Singapore, 1986.
- [6] J.W. Goodman et al., "Optical interconnections for VLSI systems", Proc. IEEE, **72**, 850 (1984).
- [7] B.D. Clymer and J.W. Goodman, "Timing uncertainty for receivers in optical clock distribution for VLSI", Opt. Eng. **27**, 944 (1988).
- [8] J.P.Boon et S.Yip, "Molecular Hydrodynamics", McGraw-Hill (1980).
- [9] D.C.Rapoport et F.Clementi, "Eddy formation in obstructed fluid flow : a molecular-dynamics study", Phys. Rev. Lett. **57**,695 (1986).
- [10] J.Hardy et Y.Pomeau, "Thermodynamics and hydrodynamics for a modeled fluid", J.Math.Phys. **13**, 1042 (1972);

J.Hardy, O.de Pazzis et Y. Pomeau, "Time evolution of two-dimensional model system. I. Invariant states and time correlation functions", J.Math.Phys. **14**, 1746 (1973);

J.Hardy, O.de Pazzis et Y. Pomeau, "Molecular dynamics of a classical lattice gas : Transport properties and time correlation function", Phys. Rev. **A13**, 1949 (1976).
- [11] U. Frisch, B. Hasslacher and Y. Pomeau, "Lattice-gas automata for Navier-Stokes equation", Phys. Rev. Lett. **56**, 1505 (1986).

- [12] D. d'Humières, Y. Pomeau et P. Lallemand, "Simulation d'allées de Von Karman bidimensionnelles à l'aide d'un gaz sur réseau", C.R. Acad. Sc. Paris, t. 301, S.II(20), 1391 (1985);

D. d'Humières, Y. Pomeau et P. Lallemand, "Ecoulement d'un gaz sur réseau dans un canal bidimensionnel : développement du profil de Poiseuille", C.R. Acad. Sc. Paris, t. 302, S.II(16), 983 (1985).
- [13] D. d'Humières, P. Lallemand et U. Frisch, "Lattice gas models for 3D hydrodynamics", Europhys. Lett. 2, 291 (1986);

D. d'Humières et P. Lallemand, "2-D and 3-D hydrodynamics on lattice gases", Helvetica Physica Acta 59, 1231 (1986);

J.P. Rivert et U. Frisch, "Simulation d'écoulement tridimensionnels par la méthode des gaz sur réseau : premiers résultats", C. R. Acad. Paris II 305, 751 (1987).
- [14] D. d'Humières, P. Lallemand et Y. Pomeau, "Simulation de l'hydrodynamique bidimensionnelle à l'aide d'un gaz sur réseau", Bulletin de la S.F.P. n°25, 14 (1986).
- [15] A. Huang, "Parallel algorithms for optical digital computers", in Technical Digest, IEEE Tenth International Optical Computing Conference, pp 13-17 (1983).
- [16] A. Clouqueur et D.d'Humières, "R.A.P.1, un Réseau d'Automates Programmables", Rapoport interne ENS (1986);

A. Clouqueur et D.d'Humières, "R.A.P.1, a cellular automaton machine for fluid dynamics", Complex Sys. 1, 585 (1987).
- [17] J. Taboury, C. Chauve, P. Chavel, "An optical approach to lattice gas automata", Optical Computing 88, Proc. SPIEE 963, 680 (1988).
- [18] I.Seyd-Darwish, P. Chavel, J. Taboury, S. Mallick, F. Devos, R. Reynaud, T. Maurin, "Automate cellulaire opto-électronique pour le gaz sur réseau", Proc. Colloque Nationale de Visualisation et Traitement d'Images en Mécanique des Fluides, 173 (1990).
- [19] I. Seyd-Darwish, P. Chavel, J. Taboury, F. Devos, R. Reynaud, T. Maurin, "Opto-electronic automata for lattice-gas", Optics in Complex Systems, SPIE 1319, 173 (1990).

Deuxième partie

- [20] M.J. Murdocca, A. Huang, J. Jahns and N. Streibl, Appl. Optics 27, 1651 (1988).

- [21] N. Streibl, "Beam shaping with optical array generators", J. Mod. Opt. **36**, 1559 (1989).
 - [22] L. Prod'homme et P. Chavel, "Réseaux de microlentilles réfractives", Rapport interne I.O.T.A. (1991).
 - [23] I.N. Ozerov et al., "Shaping the contours of dies for manufacturing of arrays having spherical elements", Sov. J. Opt. Techn. **48**, 49 (1981).
 - [24] M. Oikawa et al., "High numerical aperture planar microlens with swelled structure", Appl. Opt. **29**, 4077 (1990).
 - [25] Z.L. Liao et al., "Surface emitting In Ga AsP/InP laser with low threshold current and high efficiency", Appl. Phys. Lett. **46**, 115 (1985).
 - [26] G.D. Khoe et al., "Plasma CVD prepared SiO₂/Si₃N₄ graded index lenses ingrated in windows of laser package", Proc. 7th. Europ. Conf. Opt. Com (Copenhagen 1981).
 - [27] N.F. Borelli et al., "Planar gradient index structure", Proc. 4th topical Meeting on gradient index optical systems, (Kobe 1984).
 - [28] Z.D. Popvic et al., "Technique for monolithic fabrication of microlens arrays", Appl. Opt. **27**, 1281 (1988).
 - [29] R. Artzner, communication au colloque Opto 90.
 - [30] M.R. Taghizadeh et al., "Design and construction of holographic optical elements fot optical photonic switching applications", in Technical Degest of Topical Meeting on Photonic Switching (Optical Society of America, Washington,DC, 1987), paper FB2.
 - [31] A.W. Lohmann et F. Sauer, "Holographic telescope arrays", Appl. Opt. **27**, 3003 (1988).
 - [32] H. Dammann and K. Görtler, "High-efficiency in-line multiple imaging by means of multiple phase holograms", Opt. Commun. **3**, 312 (1971).
- H. Dammann and Klotz, "Generation of faultless multi-pinhole masks by means of spatial filtering", Opt. Commun. **13**, 268 (1975).
- J. Jahns and al., "Damman gratings for laser beam shaping", Opt. Eng. **28**, 1265 (1989).
- U. Krackhard and N. Sreibl, "Design of Dammann-gratings for array generation", Opt. Commun. **74**, 31 (1989).

- J. Turunen and al., "Stripe-geometry two-dimensional Dammann gratings", *Opt. Commun.* **74**, 245 (1989).
- [33] B. ROBERTSON, J. TURUNEN, H. ICHIKAWA, J. M. MILLER, M. R. TAGHIZADEH, A. VASARA, "Holograms in dichromated gelatin", *Appl. Opt.* **30**, 3711 (1991).
- [34] Hellesoy Ostein, "Array Generation by Means of Binary Phase Diffraction Grating", Rapport de stage, IOTA (Orsay), The Norwegian Institute of Technology (NTH Norvège), CNET (Bagneux) (1991).
- [35] A. W. Lohmann and al., "Array illuminator based on phase contrast", *Appl. Opt.* **27**, 2915 (1988).
- [36] A. W. Lohmann, "An array illuminator based on the Talbot-effect", *Optik* **79**, 41 (1988).
- [37] I. Seyd-Darwish, Ph. Legendre, P. Chavel and J. Taboury, "Optimization of an array illuminator hologram", *Annale de Physique, Colloque n°1, supplément au n°1*, **16**, 103 (1991).
- [38] F. A. Talbot, *Philos. Mag.* **9**, 401 (1836).
- [39] P. Chavel and T. C. Strand, "Range measurement using Talbot diffraction imaging of gratings", *Appl. Opt.* **23**, 862 (1984).
- [40] H. M. Smith, "Basic holographic principles", *Holographic Recording Materials, Topics in Applied Physics*, vol. 20, Springer Verlag (1977).
- [41] H. Kogelnik, "Coupled wave theory for thick hologram gratings", *The Bell System Technical Journal*, **48**, 2909 (1969).
- [42] B. J. Chang and C. D. Leonard, "Dichromated gelatin for the fabrication of holographic optical elements", *Appl. Opt.* **18**, 2407 (1979).
- J. Chang, "Dichromated gelatin holograms and their applications", *Opt. Eng.* (1980).
- [43] C. Bainier, "Etude des gélamines argentiques et bichromatées, application à l'holographie", Thèse de l'Université de Franche-Comté (1984).

Troisième partie

- [44] N. Weste et K. Eshraghian, "Principles of CMOS VLSI Design : A systems perspective", Addison Wesley (1988).

- [45] A. Vapaille et R. Castagné, "Dispositifs et circuits intégrés semiconducteurs : Physique et technologie", Dunod (1987).
- [46] A. S. Grove, "Physique et technologie à semi-conducteurs", Dunod, 116 (1971).
- [47] A. Ambroziak, "Semiconductor photoelectric devices", ILIFFE, 9 (1968);

C. Jacobi and al., "A review of some charge transport properties in Silicon", Solid State Electronics, **20**, 77 (1977).
- [48] J. C. Rodier, "Caractérisation d'une photodiode : Mesure de la sensibilité", Rapport interne IOTA (1989).
- [49] T. Maurin, T. Porcher, R. Reynaud, D. Berschandy et F. Devos, "Circuit intégré à contrôle et génération d'horloge par voie optique", Colloque International : Fonctions optiques dans l'ordinateur, 181 (1989 Toulouse).
- [50] T. Maurin, "Horloge optique : ordre de grandeur", Rapport interne IEF, (1990).
- [51] J. P. Pocholle, "Propriétés optiques des matériaux semiconducteurs à puits quantiques et applications dans le domaine du traitement du signal", Actes de l'Ecole d'Eté Optoélectronique, 205 (1989 Cargèse).
- [52] J. Millman, "Microelectronics : Digital and Analog Circuits and Systems", Mc Graw Hill, (1984).
- [53] P.S. Guilfoyle, "Digital optical computer fundamentals, implementation, and ultimate limits", Proc. SPIE, **CR35**, 288 (1991).
- [54] I. Seyd-Darwish, Ph. Lalanne, P. Chavel et J. Taboury, "Comparaison entre des architectures optoélectroniques et tout électroniques pour la distribution d'instruction dans un automate cellulaire", 7ème Journée des Fonctions Optiques dans l'Ordinateur, (Brest 1991).
- [55] J.P. Pocholle, "Caractéristiques de la propagation guidée dans les fibres optiques monomodes", L'optique guidée monomode et ses applications, Masson, Thomson-CSF, p 881-976.

- [56] A. Yariv, "Introduction to optical electronics", Holt Rinehart Winston, p. 18.