



AVERTISSEMENT

Ce document est le fruit d'un long travail approuvé par le jury de soutenance et mis à disposition de l'ensemble de la communauté universitaire élargie.

Il est soumis à la propriété intellectuelle de l'auteur. Ceci implique une obligation de citation et de référencement lors de l'utilisation de ce document.

D'autre part, toute contrefaçon, plagiat, reproduction illicite encourt une poursuite pénale.

Contact : ddoc-theses-contact@univ-lorraine.fr

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

Code de la Propriété Intellectuelle. articles L 335.2- L 335.10

http://www.cfcopies.com/V2/leg/leg_droi.php

<http://www.culture.gouv.fr/culture/infos-pratiques/droits/protection.htm>

Un modèle d'apprentissage multimodal pour un substrat distribué d'inspiration corticale

THÈSE

présentée et soutenue publiquement le 10 novembre 2010

pour l'obtention du

Doctorat de l'université Henri Poincaré – Nancy 1
(spécialité informatique)

par

Thomas Girod

Composition du jury

Rapporteurs : Hugues Berry – CR INRIA (HDR) - Lyon
Hélène Paugam-Moisy – Pr - Université de Lyon

Examineurs : Frédéric Alexandre – DR INRIA - UHP Nancy
Vincent Chevrier – MDC (HDR) - UHP Nancy
Philippe Gaussier – Pr - Université de Cergy-Pontoise
Sylvain Contassot-Vivier – Pr - UHP Nancy

Remerciements

Ce travail n'aurait jamais vu le jour sans le concours de nombreuses personnes. Je tiens en particulier à remercier ici mon encadrant Frédéric Alexandre pour m'avoir guidé scientifiquement et soutenu moralement ; mon jury de thèse pour avoir pris le temps de lire et évaluer mon travail ; mon père pour m'avoir initié très tôt aux sciences ; mes camarades de CORTEX, MAIA et plus généralement du LORIA, pour tous les bons moments passés ensemble auprès de la machine à café ; *last but not least*, Léa pour son soutien quotidien, et tous les amis qui aiment prétendre que mon travail consiste à créer Terminator.

Une université ressemble beaucoup à un récif de corail. Elle offre des eaux calmes et des particules alimentaires aux organismes délicats mais merveilleusement conçus qui seraient incapables de survivre aux coups de boutoir du ressac de la réalité, où les gens posent des questions comme : "Ce que vous faites, ça sert à quelque chose ?" et autres absurdités.

Terry Pratchett, Ian Stewart et Jack Cohen. La Science du Disque-monde. 1999

Table des matières

Introduction

Partie I Structure et fonction

7

Chapitre 1

Echelle microscopique : le neurone

1.1	Biologie du neurone	9
1.1.1	Anatomie	9
1.1.2	Fonctionnement du neurone	10
1.2	Modélisation du neurone	14
1.2.1	Modèles de conductance électrique	15
1.2.2	Modèles à spike	17
1.2.3	Modèles à fréquence de décharge	18
1.3	Conclusion	22

Chapitre 2

Echelle mésoscopique

2.1	Anatomie et physiologie du cortex cérébral	24
2.1.1	Structure laminaire	24
2.1.2	Types de neurones	24
2.1.3	Connectivité	26
2.1.4	Colonne corticale	26
2.1.5	Aire corticale	27
2.1.6	Résumé	30

2.2	Schémas de connectivité	30
2.2.1	Structure hiérarchisée	30
2.2.2	Connectivité ordonnée	31
2.2.3	Connectivité bidirectionnelle	31
2.2.4	Connectivité intra-aire	32
2.3	Modélisation corticale	34
2.3.1	Unités simples	34
2.3.2	Champs neuronaux	34
2.3.3	Dynamique locale d'activité	36
2.3.4	Micro circuit	37
2.4	Modélisation de la connectivité corticale	38
2.4.1	Flux montant	39
2.4.2	Flux latéral	40
2.4.3	Flux descendant	41
2.5	Multimodalité	43
2.6	Conclusion	44

Partie II Apprentissage 47

Chapitre 3 Plasticité cérébrale
--

3.1	Plasticité synaptique	52
3.1.1	Propriétés de la plasticité synaptique	52
3.1.2	Phénomènes biologiques de plasticité synaptique	53
3.1.3	Spike-Timing Dependent Plasticity (STDP)	53
3.2	Plasticité intrinsèque	54
3.2.1	Propriétés de la plasticité intrinsèque	54
3.2.2	Mécanismes biologiques de la plasticité intrinsèque	54
3.3	Plasticité structurelle	55
3.3.1	Évolution dendritique et synaptique	55

3.3.2	Neurogénèse	55
3.4	Conclusion	56

Chapitre 4

L'apprentissage hebbien

4.1	Apprentissage hebbien classique	57
4.1.1	Le problème de la croissance infinie	57
4.1.2	Borne supérieure	58
4.1.3	Terme de fuite	58
4.1.4	Dépression synaptique	58
4.1.5	La compétition entre les synapses	58
4.2	Apprentissage hebbien bidirectionnel	59
4.2.1	Postsynaptic gating	59
4.2.2	Presynaptic gating	60
4.3	Conclusion	60

Chapitre 5

La règle d'apprentissage Bienenstock-Cooper-Munro (BCM)

5.1	Origine	61
5.2	Fonctionnement	63
5.2.1	Activité	63
5.2.2	Apprentissage	63
5.2.3	Seuil θ dynamique	64
5.3	Analyse	68
5.3.1	Stabilité dans le cas unidimensionnel	68
5.3.2	Stabilité dans le cas multidimensionnel	69
5.3.3	Conditions d'existence des points stables	70
5.3.4	Étude statistique	72
5.3.5	Codage en population	74
5.4	Applications	75
5.4.1	Binocularité et lésions	75

5.4.2	Apprentissage sur des images naturelles et sélectivité à l'orientation .	77
5.4.3	Auto-organisation corticale	79
5.5	Conclusion	81

Partie III Modèle et expérimentations 83

Chapitre 6 Jeux de données, protocoles de test et outils

6.1	Stimuli orthogonaux	85
6.2	Posture / position	86
6.3	Apprentissage en parallèle	87
6.4	Outil de modélisation	88
6.4.1	Description générale	88
6.4.2	Architecture	89
6.4.3	Simulation	90
6.4.4	Expérience	91
6.4.5	Visualisation	92
6.4.6	Avantages et inconvénients	93

Chapitre 7 Intégration du flux montant et du flux latéral

7.1	Modèle	95
7.2	Choix des paramètres	96
7.3	Simulation	96

Chapitre 8 Paramètres de BCM et τ adaptatif

8.1	Relation entre τ et η	99
8.2	τ adaptatif	100
8.3	Valeur initiale τ_0	104

Chapitre 9

Modèle d'apprentissage multimodal

9.1	Co-apprentissage et binocularité	107
9.1.1	Dominance et corrélation des entrées	108
9.1.2	Solution au problème de la dominance	110
9.2	Modulation et apprentissage guidé	110
9.2.1	Protocole de test	110
9.2.2	Modulation de l'activité	111
9.2.3	Modulation de l'apprentissage	113
9.3	Modulation et co-apprentissage	117
9.4	Résumé	119

Chapitre 10

Application : apprentissage de la relation posture/position

10.1	Apprentissage guidé	121
10.1.1	Relation posture - position	122
10.1.2	Relation position - posture	123
10.2	Co-apprentissage	125
10.2.1	Mesurer la similarité des sélectivités	127
10.3	Discussion	128

Conclusion et perspectives

Introduction

Il est assez ironique de constater que l'organe qui nous permet de comprendre le monde qui nous entoure échappe lui-même à notre compréhension. Le cerveau et les processus cognitifs dont il est le siège sont au coeur de nombreuses questions fondamentales. Des questions de natures scientifique et philosophique telles que la nature de la conscience et de la cognition, mais aussi de nature spirituelle telles que l'existence de l'âme et la notion de libre-arbitre. Des questions lourdes d'implications, au coeur même de ce qui nous définit en tant qu'individus, en tant que cultures et en tant qu'espèce. Le travail que nous présentons ici vise à contribuer à une meilleure compréhension des mécanismes biologiques qui sous-tendent la cognition. Parler de cognition et de biologie dans un cadre informatique peut sembler étrange ; pour cette raison, commençons par situer clairement la problématique.

L'intelligence Artificielle

La question de la cognition a été abordée au fil des siècles par de nombreux penseurs dans de nombreuses disciplines, donnant naissance à diverses théories explicatives des fonctions cognitives chez l'homme et chez l'animal.

La question de la cognition a été centrale au développement initial de l'informatique. En témoignent les travaux de pionniers tels qu'Alan Turing, posant d'une part les bases de l'informatique avec la machine de Turing (Turing, 1937), et réfléchissant d'autre part à la notion de cognition avec le test de Turing (Turing, 1950). Depuis ces travaux fondateurs, le champ de l'informatique s'est considérablement élargi - l'*Intelligence Artificielle* (IA) est aujourd'hui le domaine de l'informatique s'intéressant à la cognition.

L'IA symbolique

Les premiers travaux en IA se sont orientés vers une approche dite *symbolique* : les mécanismes du raisonnement sont modélisés par la manipulation de symboles à l'aide de règles logiques. Cette approche a permis la conception de programmes capables de résoudre des problèmes complexes tels que la preuve automatique de théorèmes mathématiques. On citera par exemple le *General Problem Solver* de Newell and Simon (1972).

L'IA symbolique permet la modélisation de mécanismes cognitifs de haut niveau, mais possède une limitation : les entrées du système (symboles et règles logiques) sont définies par un opérateur humain, et ses sorties (résultats) sont interprétées par un opérateur humain. Harnad (1990) parle de *Symbol Grounding Problem* : la tâche consistant à transcrire un environnement complexe, riche et dynamique tel que nous l'expérimentons tous les jours, en un ensemble de

symboles manipulables par des règles logiques est une tâche extrêmement difficile - ici laissée à l'opérateur humain. Il en résulte que le système est déconnecté de la réalité car incapable de construire ses propres représentations. Par conséquent, le système est incapable d'interagir en temps réel, et surtout d'associer une dimension sémantique aux symboles manipulés.

L'IA numérique

L'IA *numérique* est une approche alternative à l'IA symbolique, dont le parti pris est de se positionner à un niveau sub-symbolique. La cognition est modélisée par des traitements numériques plutôt que de manipulations sur des symboles. Le système peut alors se passer de l'étape de traduction des entrées sous forme symbolique, permettant une interface directe entre le système et l'environnement par le biais de capteurs et d'effecteurs.

Une hypothèse développée dans le paradigme de l'IA numérique est que la cognition est un phénomène *émergent*. Dans une perspective bio-inspirée, la cognition est *incarnée* et son développement *ascendant*. On part d'un corps en interaction avec son environnement par le biais de capteurs (sens) et d'effecteurs (motricité), et la cognition se construit sur cette base.

Le développement cognitif de l'agent passe par l'élaboration d'un modèle interne de son environnement lui permettant d'adopter des comportements adaptés à la situation. La complexité des comportements observés est très variable ; par exemple, il pourra s'agir d'un réflexe de fuite suite à un indice perceptif associé au danger, ou encore de l'anticipation permanente de l'évolution de l'environnement combinée à une planification des actions en constante réévaluation. Dans le contexte d'une cognition émergente, ces comportements ne sont pas issus de fonctions cognitives fondamentalement différentes - ils sont les reflets de modèles plus ou moins élaborés de l'environnement.

On peut alors se demander comment l'agent élabore ce modèle interne de l'environnement. Dans le vivant on distingue généralement l'inné, provenant du patrimoine génétique, et l'acquis, provenant de l'expérience quotidienne. L'IA numérique s'intéresse à l'acquis, par le biais de l'*apprentissage*. Le modèle de l'environnement s'élabore progressivement à partir du vécu, permettant une adaptation rapide de l'agent à son environnement. Ceci soulève la question des mécanismes de l'apprentissage, dont l'importance est centrale au développement de la cognition.

Le connexionnisme et les neurosciences

Au sein de l'IA numérique, l'IA *connexionniste* est l'approche consistant à s'inspirer du système nerveux en général, et du cerveau en particulier, pour concevoir des traitements numériques. En effet, le cerveau étant le substrat de la cognition dans le vivant, il semble logique de s'en inspirer pour modéliser une cognition émergente.

Par le biais des *neurosciences computationnelles*, l'IA connexionniste est en interaction forte avec les neurosciences, qui étudient le lien entre la cognition et le substrat biologique. Les échanges sont bidirectionnels : d'une part, les sciences du vivant apportent à l'informatique une inspiration pour concevoir de nouveaux traitements de l'information et d'autre part, l'informatique apporte la possibilité d'étudier le vivant par le biais de la simulation et des modèles numériques.

Notre travail se situe dans ce contexte, à mi-chemin entre l'informatique et la biologie. Par le développement de modèles numériques, nous cherchons à interpréter les fonctions cognitives en terme de traitement de l'information supporté par une architecture de calcul distribuée, dont les flux d'informations ainsi que les mécanismes de fonctionnement et d'apprentissage au sein de ces flux soient conçus et validés en référence à notre architecture cérébrale.

Les enjeux des neurosciences computationnelles

Les neurosciences computationnelles visent la réalisation d'une cognition émergente inspirée du cerveau, en vue de mieux comprendre le lien entre la cognition et le substrat biologique. Cet objectif passe par plusieurs éléments qui nous paraissent fondamentaux.

Modèle d'architecture

Lorsque l'on cherche à explorer une fonction cognitive complexe, c'est tout d'abord le réseau des structures cérébrales impliquées dans la fonction qu'il faut définir, en mettant également le doigt sur les principaux flux d'informations et les interactions avec le monde extérieur. Ensuite, chacune de ces structures cérébrales peut faire l'objet d'une modélisation, qui visera en particulier à définir sa fonction élémentaire (et le lien avec la micro-circuiterie locale) et ses mécanismes d'apprentissage (permettant de préciser le type de mémoire de cette structure). Une telle approche permet alors de définir une fonction cognitive complexe (telle que la navigation ou l'attention sélective) comme étant le résultat de l'interaction entre structures neuronales et des échanges entre plusieurs systèmes de mémoire.

Modèle de calcul

La base de la modélisation en neurosciences computationnelles est le modèle de calcul. L'effort de modélisation passe par l'établissement d'un modèle de codage et de propagation adapté à la structure neuronale considérée.

De plus, il faut définir la manière dont l'information est traitée par les neurones. Ces traitements doivent être biologiquement plausibles, c'est-à-dire avoir une correspondance avérée ou supposée avec des mécanismes biologiques. Étant donnée la nature distribuée du système nerveux, la plausibilité biologique nécessite généralement un modèle de calcul lui aussi distribué, *i.e.* un ensemble d'unités de calcul autonomes et interconnectées.

L'obstacle principal à la modélisation réside certainement dans la complexité extrême du sujet modélisé. Le moindre millimètre carré de tissu neuronal comporter des milliers de neurones, et la physiologie d'un seul neurone est suffisamment riche pour nécessiter des modèles complexes.

Face à une telle complexité et compte tenu de la capacité de calcul actuelle des ordinateurs, il est nécessaire de simplifier d'une manière ou d'une autre. Le modèle de calcul passe donc par le choix d'une échelle de modélisation, basée sur une certaine conception des mécanismes de la cognition. Dans la conception d'un modèle, on choisit de simplifier certains aspects et de se concentrer sur d'autres, jugés primordiaux pour atteindre notre but, la conception d'un modèle de calcul permettant l'émergence des traitements complexes qui sous-tendent les fonctions cognitives.

Modèle d'apprentissage

Un autre aspect central de la modélisation en neurosciences computationnelles est la notion d'apprentissage, c'est à dire l'évolution de la cognition en fonction des expériences. L'apprentissage permettant l'élaboration progressive d'un modèle de l'environnement, il joue un rôle central dans l'émergence des fonctions cognitives.

Les neurosciences ont mis en évidence le phénomène de plasticité cérébrale, l'évolution des structures neuronales par le biais d'un ensemble de mécanismes régissant l'évolution. Il est communément admis que ces mécanismes jouent un rôle majeur dans la capacité d'apprentissage - les neurosciences computationnelles modélisent donc l'apprentissage par le biais de la plasticité cérébrale.

Positionnement

Le but de notre travail est de concevoir un mécanisme de traitement et d'apprentissage multimodal, en gardant à l'esprit la contrainte de la plausibilité biologique. Nous nous intéresserons à la multimodalité dans un contexte cognitif, ce qui nous amènera à orienter notre travail vers le cortex cérébral.

En particulier, pour que nos travaux puissent être intégrés à d'autres modèles corticaux, en particulier ceux développés dans l'équipe Cortex, notre modèle devra se baser sur les spécificités structurelles et fonctionnelles du cortex cérébral.

Plus spécifiquement, il devra reposer sur des flux d'informations compatibles avec la micro-circuiterie locale du cortex et sur des mécanismes d'apprentissage regroupant les principales caractéristiques corticales (apprentissage implicite, robustesse au bruit, stabilité, etc). Commençons à évoquer les liens entre multimodalité et cortex.

Problématique : multimodalité et cortex cérébral

Notre corps nous amène à percevoir notre environnement sous des aspects divers : nous disposons de cinq sens clairement distincts, véhiculant chacun une perception spécifique de notre environnement. Malgré cela, nous expérimentons notre environnement comme un tout, et la représentation que nous en construisons est peuplée d'objets définis par des caractéristiques associées à plusieurs sens. Ainsi, une orange est à la fois une forme, une couleur, une texture, une odeur, un goût, un nom etc - on parle alors de modalités sensorielles. Mais on pourrait aussi parler de modalités motrices ; par exemple, dans la *théorie des affordances* (Gibson, 1977), un objet est défini par l'ensemble des actions que l'agent peut effectuer sur ce dernier. Ainsi, notre orange est aussi l'acte de l'attraper, de l'éplucher, d'en prononcer le nom etc.

L'*intégration multimodale* est la capacité à combiner toutes ces modalités de natures diverses en un *tout cohérent*, en une *expérience globalisante* de notre environnement. Naturellement, les diverses perceptions ne sont pas associées arbitrairement - il s'agit d'extraire du *sens* de ces associations, de renforcer des représentations par des perceptions multiples. Ainsi, voir les mouvements des lèvres d'une personne aide à "entendre" ce qu'elle dit.

L'intégration multimodale est donc une propriété importante de la cognition, dont on observe les effets dans diverses régions du cerveau, tant au niveau cortical (aires associatives frontales et

pariétales) qu’au niveau sous-cortical (colliculus, amygdale ...). Des travaux tels que Ghazanfar and Schroeder (2006) indiquent qu’au sein du cortex cérébral, l’intégration multi-sensorielle est un phénomène omniprésent, et pas simplement au niveau associatif. Reynaud (2002) présente une liste des régions du cerveau concernées par l’intégration multimodale.

En résumé, notre problématique est la suivante : comment la micro-circuiterie corticale permet-elle la mise en place d’une intégration multimodale et d’un apprentissage associé.

Plan

Pour traiter notre problématique, notre travail se divisera en trois grandes parties.

Dans la première partie, nous traiterons la question de la modélisation corticale. Nous aborderons dans un premier temps les concepts fondamentaux des modèles connexionnistes : neurones, synapses, et les modèles mathématiques et numériques associés. Nous verrons que cette échelle de modélisation dite *microscopique* n’est vraisemblablement pas adaptée à la modélisation du cortex cérébral et de ses fonctions cognitives complexes, mais qu’elle donne les clés pour comprendre le passage d’une micro-circuiterie à la fonction cognitive correspondante. Dans un deuxième chapitre, nous aborderons l’échelle de modélisation dite *mésoscopique*, plus adaptée à nos besoins. Nous verrons quels en sont les fondements biologiques, et quels types de modèles corticaux on en retire. Ces modèles seront la base sur laquelle nous construirons notre propre modèle par la suite.

Dans la deuxième partie, nous aborderons la question de l’apprentissage, qui se trouve au coeur de notre problématique. Nous verrons que la fonction d’apprentissage est associée aux phénomènes de plasticité cérébrale, et quels modèles numériques découlent de ces observations. Notre problématique visant la mise en place d’un apprentissage biologiquement plausible, adapté à la structure et à la fonction du cortex, robuste et adaptable à un contexte multimodal, nous nous intéresserons en particulier au modèle BCM, qui regroupe ces propriétés.

Dans la troisième partie, nous proposerons un modèle d’apprentissage multimodal positionné à l’échelle mésoscopique décrite dans la première partie, et se basant sur le modèle d’apprentissage BCM décrit dans la deuxième partie. Notre modèle apportera des améliorations sur la stabilité de BCM, et un mécanisme permettant l’apprentissage dans un contexte multimodal. De plus nous présenterons une bibliothèque simplifiant la construction de modèles et le déroulement de simulations. Enfin, nous testerons notre modèle sur un cas pratique, l’apprentissage d’une relation non-linéaire entre posture et position.

Première partie

Structure et fonction

Chapitre 1

Echelle microscopique : le neurone

La question du lien entre les fonctions cognitives et le système nerveux est au coeur des neurosciences. La compréhension de ce lien passe par l'étude de l'anatomie et de la physiologie du système nerveux. Rámon y Cajal a montré il y a un peu plus d'un siècle que le système nerveux est composé de cellules polarisées (neurones), interconnectées par le biais de synapses. Si le neurone est effectivement l'*élément atomique* composant le système nerveux, il semble logique d'en étudier le fonctionnement. Nous verrons qu'il peut se traduire en terme de signal, d'information et de traitements, permettant la conception de modèles mathématiques et numériques.

La modélisation à l'échelle du neurone est généralement qualifiée de *microscopique*. Nous verrons que cette échelle n'est vraisemblablement pas adaptée à la notion de multimodalité, ce qui justifiera notre positionnement à une échelle de représentation intermédiaire dite *mésoscopique*. Il nous semble cependant important de comprendre les mécanismes fondamentaux au niveau du neurone avant de se placer à un niveau supérieur.

1.1 Biologie du neurone

Il est communément admis que les neurones jouent un rôle fondamental dans la capacité de traitement de l'information du système nerveux. Dans un premier temps, nous allons étudier ce que la neurobiologie nous apprend de leur anatomie et de leur fonctionnement.

1.1.1 Anatomie

En règle générale, un neurone se compose de trois parties : un *arbre dendritique*, un *corps cellulaire* et un *axone*. Il existe cependant de nombreuses variations morphologiques, comme en témoigne la figure 1.1. Si l'anatomie et le fonctionnement varient d'un type de neurone à un autre, le principe fondamental reste le même. Nous ne détaillerons pas ici les différences anatomiques qui mènent à classer les neurones dans différentes catégories. L'anatomie générique que nous allons décrire se rapproche de celle du neurone moteur.

On peut se représenter un neurone comme une cellule possédant des entrées et une sortie. Les entrées du neurones sont les *dendrites*, qui forment ensemble l'*arbre dendritique*. Les dendrites ont de très nombreuses ramifications partant de divers points du corps cellulaire, et viennent

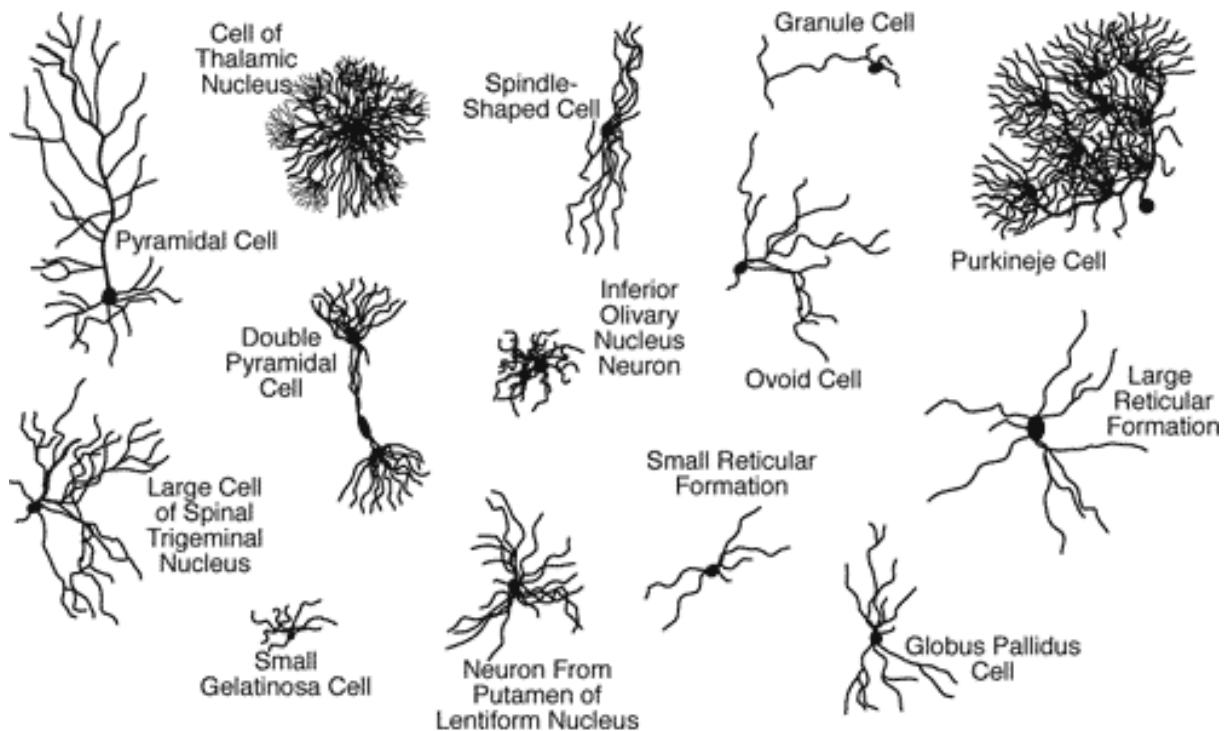


FIGURE 1.1 – Quelques types de neurones, d’après les dessins de Ramon y Cajal

former des connexions appelées *synapses* avec les axones des autres neurones. Un neurone possède en moyenne 10^4 connexions synaptiques. La sortie du neurone est l’axone, qui s’étend pour atteindre d’autres neurones. Contrairement aux dendrites, il prend racine en un unique point du corps cellulaire.

1.1.2 Fonctionnement du neurone

Comme pour toute cellule, la membrane du neurone est chargée électriquement. La présence d’ions non uniformément répartis à l’intérieur et à l’extérieur du neurone entraîne une différence de potentiel électrique au niveau de la membrane appelée *potentiel membranaire*. La forte variabilité du potentiel membranaire est une caractéristique particulière des neurones. Les variations du potentiel membranaire forment des signaux électriques capables de se propager d’un neurone à un autre, ce qui amène à conclure que l’activité électrique est un vecteur de transmission et de traitement de l’information dans le système nerveux (Kandel et al., 2000).

Se posent alors plusieurs questions que nous allons aborder :

- Quel mécanisme est à l’origine du potentiel membranaire ?
- A quoi ressemble la dynamique du potentiel membranaire ?
- Quelles sont les causes de ces variations ?
- Quelles sont les conséquences de ces variations ?
- Comment se propagent les variations de potentiel membranaire ?

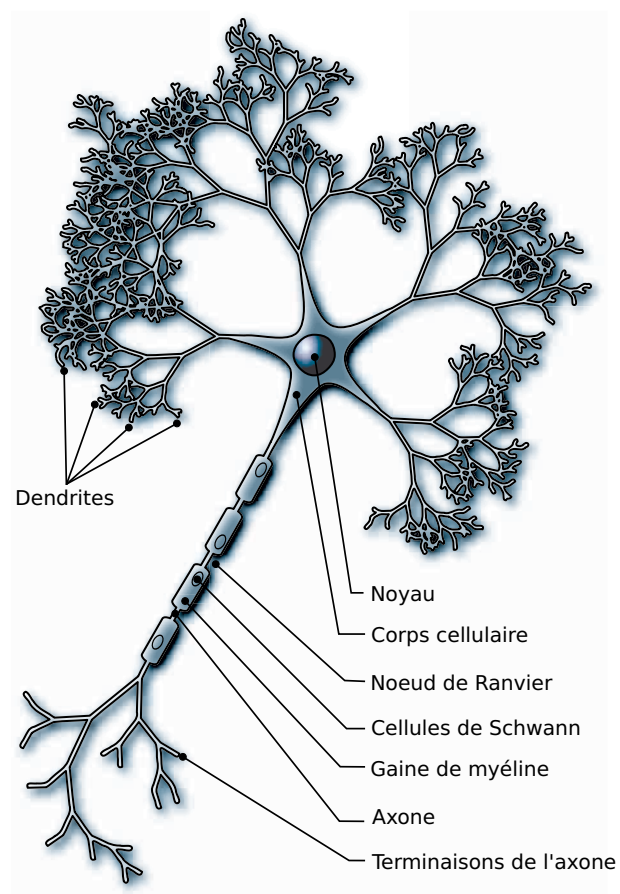


FIGURE 1.2 – Anatomie d'un neurone - Wikimedia Commons

Canaux ioniques

Le potentiel membranaire est causé par une différence de charge électrique entre l'intérieur et l'extérieur de la cellule. Cette différence de charge est principalement liée à la concentration interne et externe en ions calcium (Ca^{2+}), sodium (Na^{+}) et potassium (K^{+}). La membrane du neurone est équipée de protéines appelées *canaux ioniques*, dont la fonction est de faire transiter des ions au travers de la membrane (vers l'intérieur ou vers l'extérieur), altérant ainsi les concentrations ioniques, et donc le potentiel membranaire du neurone.

Les canaux ioniques sont majoritairement de deux types. Les canaux *voltage-dépendants* dont l'ouverture contrôlée par la modification du potentiel membranaire du neurone. Les canaux *chimio-dépendants*, dont l'ouverture est contrôlée par la présence d'une substance chimique spécifique sur un récepteur. L'action des canaux ioniques donne donc des propriétés électriques actives aux neurones (variations fortes de potentiel membranaire), tandis que les autres cellules ont des propriétés électriques sont passives (préservation d'un potentiel membranaire stable).

Variations du potentiel membranaire

Le potentiel membranaire d'un neurone est caractérisé tout d'abord par un *potentiel au repos*. Celui se situe typiquement autour de -65mV , bien que sa valeur puisse varier selon la nature du neurone considéré (Kandel et al., 2000). Cette polarisation est entretenue par l'action des canaux ioniques voltage-dépendants qui viennent réguler la concentration ionique à l'intérieur du neurone.

L'activité du neurone est caractérisée par ses variations autour du potentiel de repos. Lorsque le potentiel membranaire est supérieur au potentiel au repos, on parle de *dépolarisation*. Inversement, lorsque le potentiel membranaire est inférieur au potentiel au repos, on parle d'*hyperpolarisation*.

Transmission synaptique

La variation du potentiel membranaire est initialement causée par l'activité synaptique. Une synapse est une connexion entre un axone *pré-synaptique* et une dendrite *post-synaptique*. Une synapse est composée d'un ou plusieurs *boutons synaptiques*. La terminaison présynaptique d'un bouton synaptique contient des *vésicules de neurotransmetteurs*. La dépolarisation de la terminaison présynaptique provoque la fusion des vésicules de neurotransmetteurs avec la membrane, émettant ainsi les neurotransmetteurs dans la fente synaptique.

Du côté postsynaptique, les neurorécepteurs situés sur la membrane captent les neurotransmetteurs émis dans la fente synaptique. L'activation de ces canaux ioniques chimio-dépendants entraîne une variation du potentiel membranaire dans la terminaison post-synaptique du bouton synaptique. Cette variation peut être une dépolarisation ou une hyperpolarisation, selon le type de neurotransmetteur et le type de récepteur concerné.

La transmission électro-chimique que nous venons de décrire est extrêmement complexe, régie par de nombreux paramètres. Il existe un grand nombre de neurotransmetteurs, de capteurs, ainsi que d'autres mécanismes entrant en jeu (récepteurs métabotropes, second messagers ...). Il serait trop long et hors de propos d'aller jusqu'à ce niveau de détail.

Ainsi, on retiendra qu'une dépolarisation pré-synaptique provoque un événement chimique, qui entraîne à son tour une variation du potentiel membranaire post-synaptique. On parle d'activité électrochimique. Il est important de noter que l'efficacité de cette transmission est variable

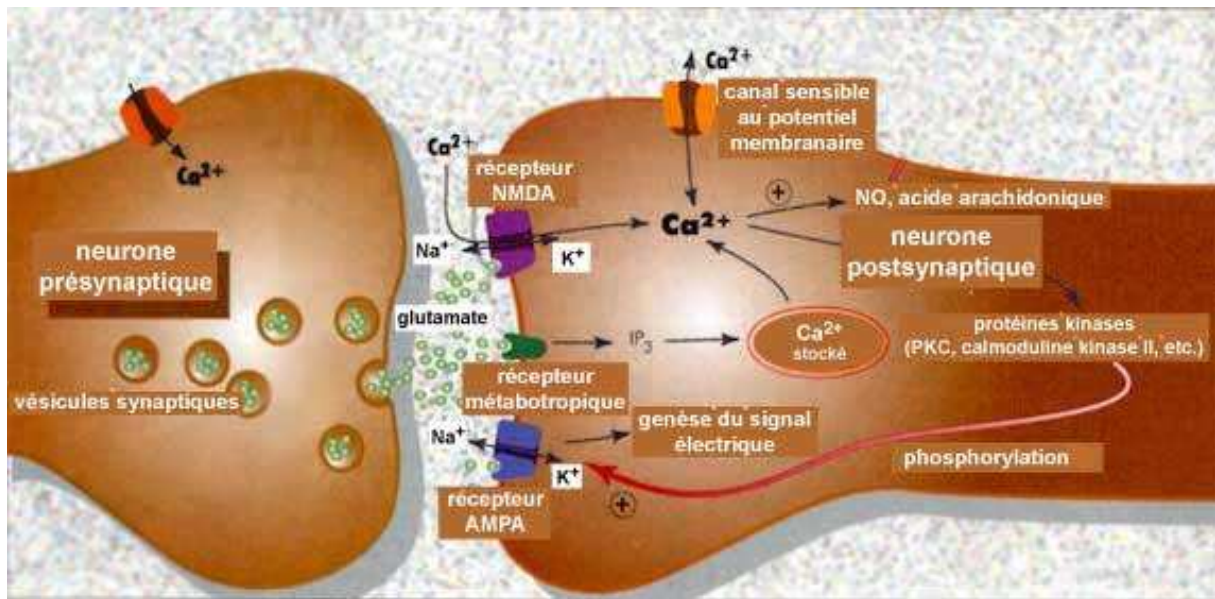


FIGURE 1.3 – Anatomie et fonctionnement d’une synapse, adapté de <http://lecerveau.mcgill.ca>

d’une synapse à une autre, et évolue généralement dans le temps. Cette propriété de *plasticité synaptique* joue un rôle fondamental dans la capacité d’apprentissage. Nous reviendrons sur ce sujet en détail dans la partie II.

Potentiel post-synaptique

Nous avons vu que l’activité électrochimique des synapses provoque une dépolarisation ou une hyperpolarisation des sites synaptiques. Une telle variation est appelée *potentiel post-synaptique* (PPS). On distingue les *potentiels post-synaptiques excitateurs* (PPSE) correspondant à une dépolarisation, des *potentiels post-synaptiques inhibiteurs* (PPSI) correspondant à une hyperpolarisation.

Un PPS est une variation extrêmement faible. Cependant, l’arbre dendritique d’un neurone possède en moyenne 10^4 connexions synaptiques. L’accumulation des PPS reçus sur toutes les synapses du neurone post-synaptique est à même de causer des variations conséquentes du potentiel membranaire. Lorsque la dépolarisation est suffisamment forte, le neurone émet un *potentiel d’action*.

Potentiel d’action

Chaque neurone possède un seuil nommé *seuil de décharge*, situé généralement autour de -55mV . Lorsque la dépolarisation du neurone atteint ce seuil, des canaux ioniques voltage-dépendants s’ouvrent, provoquant une forte dépolarisation transitoire atteignant $+30\text{mV}$ à $+40\text{mV}$. Le potentiel membranaire redescend ensuite, passant généralement par une phase d’hyperpolarisation avant de se stabiliser à nouveau au potentiel de repos. La fin de ce cycle d’activité est souvent accompagnée d’une *période réfractaire*, durant laquelle le neurone ne peut pas passer à l’état dépolarisé.

Ce pic de dépolarisation (*spike* dans la littérature anglo-saxonne) est appelé *potentiel d'action*. Les potentiels d'actions sont couramment perçus comme étant le support de communication électrique entre les neurones. Nous reviendrons plus tard sur la manière dont l'information peut être codée dans ces signaux électriques.

Propagation des potentiels d'action

L'émission d'un potentiel d'action entraîne une dépolarisation forte et transitoire à la base de l'axone. La dépolarisation se propage aux zones adjacentes de la membrane, entraînant le dépassement du seuil de décharge dans les zones adjacentes. Ceci provoque une répétition du potentiel d'action, qui se propage à son tour. La phase réfractaire qui suit l'émission d'un potentiel d'action évite que le potentiel d'action ne se propage vers sa source, donnant ainsi une propagation unidirectionnelle, parcourant l'axone.

Ce mécanisme de propagation du potentiel d'action est accéléré par la présence d'une *gaine de myéline* autour de l'axone, comme l'illustre la figure 1.2. La gaine de myéline forme des sections séparées par des *noeuds de Ranvier*. Entre deux noeuds de Ranvier, la myéline réduit la capacité électrique de la membrane, augmentant la distance à laquelle une dépolarisation locale se propage. La dépolarisation se propage alors jusqu'au noeud de Ranvier suivant, augmentant la distance entre un potentiel d'action et sa répétition. Ceci permet d'accélérer la vitesse de transmission du potentiel d'action.

Résumé

L'activité électrique des neurones est complexe, et perçue comme le vecteur principal de communication dans le système nerveux. Les neurones communiquent par transmission de potentiels d'action sur leurs axones, et codent donc de l'information au sein de ceux-ci. Schématiquement, on peut voir le neurone comme une cellule recevant des signaux d'entrées de plusieurs sources, intégrant ces signaux dans son activité électrique, et émettant à son tour des signaux à partir de son activité.

Par la suite, nous allons chercher à comprendre comment l'information est codée dans les potentiels d'action, et quel calcul le neurone effectue sur cette information.

1.2 Modélisation du neurone

La modélisation neuronale est un moyen d'étudier et de comprendre le fonctionnement des neurones. En construisant des modèles reproduisant les différents aspects de l'activité du neurone, on acquiert une compréhension des principes de fonctionnement du neurone.

Les modèles mathématiques permettent d'exprimer la dynamique de l'activité du neurone par un système d'équations différentielles. Lorsque l'étude formelle des modèles mathématiques est trop complexe, la simulation informatique permet de tester les modèles mathématiques et les comparer avec les observations effectuées sur le vivant.

Dans le domaine des neurosciences computationnelles, la connaissance obtenue par la modélisation est utilisée pour exprimer l'activité neuronale en terme de traitement de l'information. Cette passerelle entre les neurosciences et les sciences informatiques apporte à la fois un point

de vue différent au domaine des neurosciences et une source d'inspiration pour concevoir de nouvelles approches du traitement de l'information dans le domaine informatique.

Dans une optique de traitement de l'information, la modélisation neuronale se concentre sur les propriétés électriques du neurone, puisque celles-ci sont perçues comme étant le vecteur de traitement et de transmission de l'information dans le système nerveux. Les différents modèles interprètent différemment le lien entre activité électrique et information, mais s'accordent globalement sur le rôle central des potentiels d'action.

1.2.1 Modèles de conductance électrique

La fonction du neurone en terme de traitement de l'information est étroitement liée à son activité électrique. Les modèles de conductance électrique reproduisent le comportement électrique de la membrane d'un neurone - ou de toute autre cellule excitable - en l'exprimant sous la forme d'un circuit électrique.

Hodgkin-Huxley

Hodgkin and Huxley (1952) observèrent des potentiels d'actions par enregistrement intracellulaire sur un axone géant de calmar. Afin d'expliquer quantitativement les potentiels d'action observés, ils proposèrent un modèle mathématique de la génération de potentiels d'action basé sur les lois électriques. Ce modèle est remarquable car il fait intervenir des variables qui seront plus tard justifiées par la découverte des canaux ioniques.

Le modèle se base sur le système de quatre équations différentielles couplées :

$$C \frac{dV}{dt} = -g_K n^4 (V - E_K) - g_{Na} m^3 h (V - E_{Na}) - g_L (V - E_L) + I(t) \quad (1.1)$$

$$\tau_n(V) \frac{dn}{dt} = -[n - n_0(V)] \quad (1.2)$$

$$\tau_m(V) \frac{dm}{dt} = -[m - m_0(V)] \quad (1.3)$$

$$\tau_h(V) \frac{dh}{dt} = -[h - h_0(V)] \quad (1.4)$$

$$(1.5)$$

V est le potentiel membranaire du neurone. Le modèle introduit trois constantes g_K , g_{Na} et g_L , représentant respectivement les conductances des canaux calcium et sodium, et la conductance de la membrane causée par la fuite naturelle d'ions. Les canaux ioniques étant voltage-dépendants, leur conductance varie en fonction du potentiel électrique et du temps ; pour représenter cela, le modèle introduit trois variables n , m et h . De plus, les constantes E_x représentent les potentiels d'équilibre des différents canaux. Ce système peut être schématisé par un circuit électrique illustré par la figure 1.4.

Dérivés du modèle Hodgkin-Huxley

A la suite du modèle de Hodgkin-Huxley, d'autres modèles de conductance électrique ont vu le jour, afin de généraliser le principe à d'autres neurones n'ayant pas les mêmes canaux ioniques

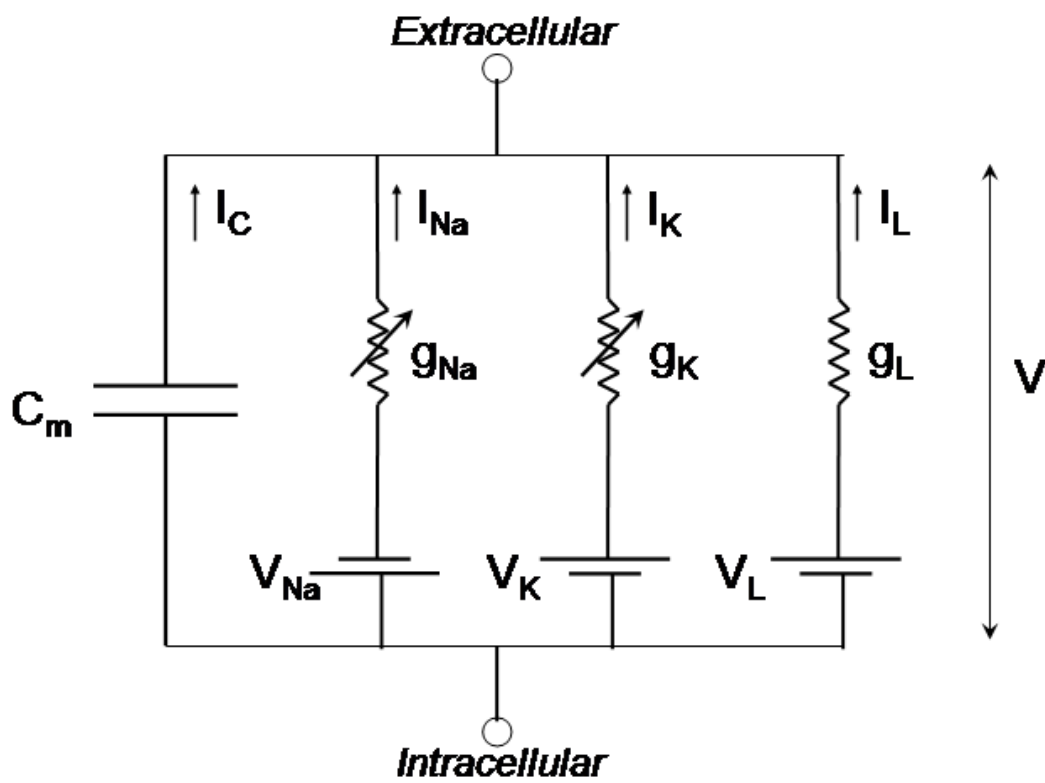


FIGURE 1.4 – Schématisation du modèle de Hodgkin-Huxley (Scholarpedia)

ou de simplifier la modélisation. Parmi les plus utilisés, on citera par exemple Morris-Lecar (Morris and Lecar, 1981) ou encore FitzHugh-Nagumo (Fitzhugh, 1961; Nagumo et al., 1962).

Modèles compartimentaux

Les modèles de conductance électrique que nous venons de présenter réduisent le neurone à un point caractérisé par une activité électrique. Or, le potentiel membranaire d'un neurone n'est pas uniforme : les potentiels post-synaptiques et les potentiels d'action sont l'exemple même d'une variation locale du potentiel membranaire.

Les modèles compartimentaux rendent compte de l'influence de la morphologie du neurone sur son activité électrique. Pour cela, l'arbre dendritique du neurone est divisé en compartiments inter-connectés, représentés par des "segments" de diamètre variable. Chaque compartiment possède son activité électrique locale, et les lois de conductance des câbles électriques sont adaptées pour représenter la propagation du potentiel électrique au sein du neurone.

Ces modèles permettent de modéliser de manière plus précise le cumul et la propagation des potentiels post-synaptiques au sein de l'arbre dendritique. En effet, il s'agit d'un phénomène complexe et fondamental au fonctionnement du neurone, puisqu'il entraîne la dépolarisation du neurone et l'émission des potentiels d'action.

1.2.2 Modèles à spike

Les modèles à spike se basent sur l'hypothèse suivante : ce qui importe dans l'activité d'un neurone, c'est la temporalité de l'émission des potentiels d'action. Pour un potentiel d'action donné, on ne s'intéresse pas au détail de la variation du potentiel membranaire, mais à l'instant auquel le potentiel d'action est émis.

Les modèles à spike représentent donc un neurone par un état interne continu variant en fonction des entrées, et des événements discrets représentant l'émission de potentiels d'action, transmis entre les neurones.

Leaky Integrate and Fire

Le modèle *Integrate and Fire* et son extension *Leaky Integrate and Fire* (LIF) décrivent le principe de base des modèles de neurone à spike. Le principe consiste à intégrer l'activité électrique dans le potentiel membranaire. Le dépassement du seuil de décharge entraîne l'émission d'un spike et la réinitialisation du potentiel membranaire. L'ajout d'un terme de fuite permet au neurone de se stabiliser au potentiel de repos en l'absence de stimulation extérieure. Soit :

$$\tau_m \frac{dV(t)}{dt} = -(V(t) - E_L) + RI(t) \quad (1.6)$$

où v est le potentiel membranaire, E_L une constante représentant le potentiel de repos, τ_m une constante de temps, I le courant entrant. Le terme $-(V(t) - E_L)$ introduit la fuite dans le mécanisme d'intégration. Lorsque le potentiel membranaire atteint le seuil de décharge ϑ , le neurone émet un spike, représenté par une activité unitaire qui est transmise aux autres neurones. Suite à l'émission du spike, le potentiel membranaire est réinitialisé à l'état de repos.

$$V(t) = E_L \quad (1.7)$$

Autres modèles

Il existe de nombreux autres modèles que nous ne citerons pas ici. Chaque modèle propose un compromis réalisme/complexité différent. La manière de représenter l'état interne varie, posant des hypothèses simplificatrices différentes. La suite de notre travail se concentrant sur une échelle de description plus élevée, nous n'entrerons pas dans le détail des choix de modélisation au niveau neuronal.

Détecteur de coïncidences

Une manière d'interpréter le traitement effectué par un neurone à spike à fuite est de le considérer comme un *détecteur de coïncidences* (Trappenberg, 2010) : l'arrivée simultanée de plusieurs stimuli sur certaines entrées du neurone entraîne son activation.

Soit un neurone LIF recevant des stimuli sur plusieurs entrées. Le plus souvent, l'arrivée d'un stimulus isolé entraîne une augmentation du potentiel membranaire v insuffisante pour déclencher l'émission d'un potentiel d'action. Avec la présence du terme de fuite, le potentiel membranaire rejoint rapidement son état au repos. L'arrivée d'un nouveau stimulus isolé reproduit alors le même phénomène, et le neurone n'est jamais suffisamment dépolarisé pour émettre un potentiel d'action.

L'émission d'un potentiel d'action dépend donc de l'arrivée concomitante de plusieurs stimuli dans un intervalle de temps restreint, ce qu'on pourrait apparenter à un "et logique". Suivant le profil de poids du neurone, toute combinaison de stimuli n'entraînera pas l'émission d'un potentiel d'action. On peut donc en conclure qu'étant donné un profil de poids, le neurone est à même de détecter la coïncidence de stimuli spécifiques, et que cette détection est exprimée par l'émission d'un potentiel d'action.

Résumé

Les modèles à spike proposent un modèle simplifié de l'activité électrique des neurones basé sur la temporalité des potentiels d'action. En effet, cette approche considère que le code neuronal peut être correctement approximé en considérant la temporalité d'événements discrets (les potentiels d'action) plutôt que la variation continue de l'activité électrique des neurones. Sur cette base, on obtient des neurones artificiels à même de détecter des coïncidences entre plusieurs entrées.

1.2.3 Modèles à fréquence de décharge

Les modèles à fréquence de décharge se basent sur l'hypothèse suivante : une grande partie de l'information codée par un neurone réside dans son niveau d'activité moyen, que l'on associe à sa fréquence d'émission de potentiels d'action - ou *fréquence de décharge*.

Ce choix est motivé par plusieurs observations. Depuis les observations de Adrian and Zotterman (1926), on sait que dans la plupart des systèmes sensoriels, la fréquence de décharge

augmente non-linéairement avec l'intensité du stimulus présenté. Cependant, la fréquence de décharge varie aussi entre plusieurs stimuli d'intensité similaire. Cette observation implique qu'un neurone peut développer une sensibilité particulière à certains stimuli - la fréquence de décharge permettant donc d'exprimer la sélectivité du neurone. Henry et al. (1974) observe ce phénomène chez le chat, où la présentation d'une barre dans le champ de vision du chat provoque dans un neurone une fréquence de décharge variable selon l'orientation de la barre.

L'activité d'un neurone

La valeur d'activité dans un modèle de neurone fréquentiel peut être associée de plusieurs manières au processus biologique d'émission de potentiels d'action. Certains modèles insistent sur un lien fort entre l'activité modélisée et des grandeurs numériques observables expérimentalement dans le vivant ; ce n'est cependant pas le cas de tous les modèles. Citons donc quelques définitions courantes de l'activité dans les modèles fréquents.

Moyenne temporelle La méthode la plus simple pour calculer la fréquence de décharge instantanée est la moyenne temporelle sur une fenêtre de temps glissante. Soit $y(t)$ la fréquence de décharge instantanée à l'instant t . On la calcule en considérant $s(t, t + \Delta t)$, le nombre de potentiels d'action émis dans une fenêtre temporelle de longueur Δt . Ainsi :

$$y(t) = s(t, t + \Delta t) / \Delta t \quad (1.8)$$

Le choix de la longueur de la fenêtre de temps a une forte incidence sur les résultats obtenus par cette approche. L'émission de potentiels d'action n'étant pas un phénomène régulier, une fenêtre de temps trop courte par rapport à l'écart moyen entre deux potentiels d'action donnera une mauvaise approximation de la fréquence de décharge instantanée. A l'inverse, une fenêtre de temps trop longue donnera une fréquence de décharge instantanée stable, mais lissera les variations ayant lieu à une échelle de temps très courte.

En général on choisit une fenêtre assez courte, avec Δt de l'ordre de 100 à 500ms. Cette méthode a été utilisée avec succès dans de nombreuses expériences. Cependant, on observe chez les organismes vivants des temps de réaction très courts, pouvant descendre en dessous des fenêtres de temps généralement considérées. Cette représentation de l'activité ne peut alors pas rendre compte de ces phénomènes - un argument avancé par Thorpe et al. (1996) pour justifier l'intérêt des modèles à spike.

Moyenne spatialisée Plutôt que de considérer les potentiels d'actions émis par un neurone, on considère les PA émis par une population de neurones proches et fortement connectés entre eux. Soit k le nombre de neurones mesurés, $s_i(t_1, t_2)$ le nombre de potentiels d'action du neurone i entre les instants t_1 et t_2 , avec $\Delta t = t_2 - t_1$:

$$y(t) = \frac{1}{k\Delta t} \sum_{i=1}^k s_i(t, t + \Delta t) \quad (1.9)$$

Avec une population suffisamment grande, cette approche permet d'obtenir une fréquence de décharge fiable sur une fenêtre de temps réduite, pouvant ainsi rendre compte de variations rapides d'activité. Cette amélioration de la résolution temporelle se fait cependant au prix d'une

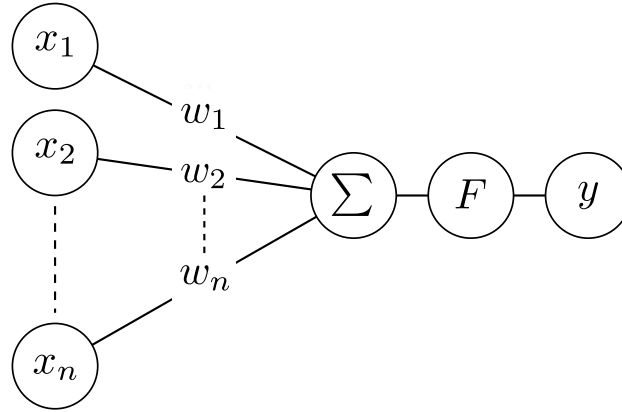


FIGURE 1.5 – Neurone formel

perte de résolution spatiale : on se place à l'échelle de la population de neurones plutôt que du neurone.

Le neurone formel

Les modèles fréquentiels varient principalement dans leur manière d'intégrer l'activité entrante pour obtenir la fréquence de décharge du neurone. L'approche la plus courante consiste à calculer la somme pondérée des entrées puis de faire usage d'une fonction de transfert F non-linéaire :

$$y(t) = F\left(\sum_{i=1}^n w_i(t)x_i(t) - \theta\right) \quad (1.10)$$

où x_i est l'activité à l'entrée i du neurone, w_i le poids synaptique associé à l'entrée i et θ une constante représentant le seuil d'activation du neurone. Le *neurone formel* décrit par McCulloch and Pitts (1943) est la première tentative de formalisation de l'action du neurone en terme de traitement de l'information. Il s'agit d'un cas particulier du calcul décrit précédemment pour lequel la fonction d'activation F est la fonction de Heaviside. Ainsi, l'activité du neurone est binaire, avec $y = 1$ si la stimulation reçue dépasse le seuil d'activation, et $y = 0$ autrement.

Neurone prototype

Une autre approche du neurone formel consiste à intégrer des entrées numériques en calculant une distance entre le vecteur d'entrées et le vecteur de poids. Suivant les modèles, le calcul de la distance peut varier.

$$y(t) = F\left(\sum_{i=1}^n (w_i(t) - x_i(t))^2\right) \quad (1.11)$$

La fonction d'activation utilisée est une fonction à base radiale, de sorte que l'activité exprime la proximité entre le vecteur d'entrée et le vecteur de poids :

$$F(x) = e^{-\beta x} \quad (1.12)$$

où $\beta > 0$ est une constante contrôlant la pente de la fonction radiale. Ainsi, le vecteur de poids représente un prototype, et l'activation exprime la distance de l'entrée à ce prototype. Cette méthode est couramment utilisée dans des modèles tels que les cartes auto-organisatrices de Kohonen (1982).

Interprétation entre terme de filtre

Considérons $\mathbf{x}$ le vecteur des entrées d'un neurone et $\mathbf{w}$ le vecteur des poids synaptiques associés. $\mathbf{x}(t)$ est le stimulus reçu par le neurone à l'instant t . La sortie $y(t)$ du neurone dépend du stimulus $\mathbf{x}(t)$ et du vecteur de poids $\mathbf{w}$: le neurone formel se base sur un produit scalaire tandis que le neurone prototype se base sur une distance. Dans les deux cas, la valeur du vecteur de poids $\mathbf{w}$ rend le neurone sensible à certains stimuli plutôt que d'autres ; on dit alors que le neurone se comporte comme un filtre, car il est capable de discriminer certaines configurations de ses entrées.

Résumé

Le codage en fréquence pose l'hypothèse que la modélisation de la fréquence de décharge des neurones est suffisante pour rendre compte des traitements principaux effectués par le système nerveux. Des auteurs tels que Richmond et al. (1990, 1987); Optican and Richmond (1987); Heller et al. (1995) se sont intéressés au codage de l'information dans les potentiels d'actions pour les aires corticales V1 et IT chez le primate. Selon ces études, la fréquence de décharge code en moyenne 73% de l'information pour V1, et 84% pour IT. On ajoutera cependant que la latence du premier potentiel d'action est importante : elle code en moyenne 35% de l'information dans V1, et 48% dans IT. Ces résultats vont dans le sens des travaux de Gautrais and Thorpe (1998), qui insistent sur l'importance de la latence du premier spike pour expliquer la rapidité de certains traitements visuels. Néanmoins, il ressort des nombreuses études sur le sujet que ces chiffres varient selon l'espèce, la région et la fonction étudiée - ce qui semble indiquer que les deux paradigmes de codage (fréquence et temporalité) co-existent. Les modèles à fréquence de décharge semblent donc pertinents au regard de la quantité substantielle d'information codée par cette grandeur.

En gardant à l'esprit les hypothèses simplificatrices posées par l'approche, le codage en fréquence s'appuie sur un formalisme mathématique simple (i.e. le neurone formel ou un dérivé), ce qui a facilité l'étude des traitements de l'information obtenus par la construction de réseaux de neurones artificiels. On citera par exemple les travaux de Rosenblatt (1958) sur le perceptron qui ont permis la mise au point de classifieurs linéaires en se basant sur des neurones artificiels. Par la suite, Hornik et al. (1989) a démontré que sous certaines conditions le modèle du perceptron est à même d'approximer toute fonction continue. De tels travaux apportent des hypothèses sur les traitements que pourraient effectuer les neurones biologiques sur l'information transmise dans le système nerveux.

1.3 Conclusion

Les travaux en biologie apportent une quantité importante de connaissances sur l'anatomie et la physiologie des neurones. Ces informations nous permettent de comprendre le fonctionnement d'un neurone, et faire un premier lien entre structure et fonction : un neurone intègre l'activité électrique en provenance de plusieurs sources pour produire une activité électrique en sortie. L'information est codée dans l'activité électrique, et le traitement effectué par le neurone est défini par le processus d'intégration et d'émission d'activité électrique. En considérant l'ensemble des connexions entrantes d'un neurone comme un espace d'entrée, on peut exprimer la fonction du neurone sous la forme d'une projection de son espace d'entrée vers son espace de sortie. Sur cette base, il est possible de construire des traitements simples, et la combinaison de plusieurs neurones pour former des réseaux permet la mise en place de traitements plus complexes.

Cependant, il semble difficile de parler de multimodalité à l'échelle du neurone. Un neurone ne fait vraisemblablement pas de distinction entre les potentiels d'actions reçus, quelle qu'en soit la provenance ou l'information transmise. Pour traiter la question de la multimodalité, il nous semble pertinent de nous placer à une échelle supérieure, raisonnant en terme de *voies neuronales*, de traitements d'information à l'échelle de la modalité.

La modélisation de structures corticales conséquentes à l'échelle du neurone se heurte à l'incroyable complexité du cerveau. Même avec des modèles de neurones relativement simples, la simulation de dizaine de milliers de neurones nécessite une puissance de calcul difficilement accessible. On citera par exemple le *Blue Brain Project* (Markram, 2006), qui vise la simulation d'une colonne corticale (voir chapitre 2) (environ 10000 neurones) à un haut niveau de détail et en temps réel. Hill and Markram (2008) ont annoncé la complétion de la phase I du projet, c'est à dire la simulation d'une colonne du cortex somatosensoriel du rat. Ceci aura néanmoins nécessité l'utilisation des 8192 processeurs du super-calculateur *Blue Gene*.

Face à un tel obstacle, il nous semble nécessaire de changer d'échelle de modélisation, pour se positionner à un niveau mésoscopique. Cette échelle se base sur la population de neurones comme unité fondamentale de traitement - pour la modélisation du cortex, cette unité est le plus souvent associée à la notion de colonne corticale. Les colonnes corticales sont assemblées pour former des cartes corticales, et la connectivité du réseau prend la forme de projections entre les cartes.

Les modèles à spike étant plus adaptés à la modélisation à l'échelle du neurone (un spike représentant un potentiel d'action), la suite de notre travail se focalisera sur un codage en fréquence de décharge.

Chapitre 2

Echelle mésoscopique

Nous avons montré dans le chapitre précédent qu'un neurone effectue des traitements sur les signaux nerveux en combinant les signaux reçus pour en émettre des nouveaux. Ce mécanisme permet la mise en place de traitements de l'information rudimentaires, et c'est en combinant un grand nombre de neurones que des traitements complexes peuvent être mis en place.

La complexité du cerveau est liée à deux facteurs : le grand nombre de neurones qui le composent et le degré de connectivité entre les neurones. Ainsi, le cerveau humain est composé d'environ 10^{11} neurones, et d'environ 10^{14} connexions synaptiques (Swanson, 1995). Devant une telle complexité, l'effort de modélisation passe par une simplification sur deux aspects : l'échelle de représentation et la connectivité.

Échelle de représentation La dimension du réseau nécessite un réajustement de l'échelle de représentation. Plutôt que d'étudier et modéliser le système nerveux à l'échelle du neurone, on change de granularité pour se positionner à l'échelle de la population localisée de neurones. D'un point de vue structurel, cette échelle dite *mésoscopique* est pertinente : Doya (1999) rappelle que chaque région du cerveau possède un certain degré de régularité dans sa structure. Cette régularité amène à concevoir la structure d'une région comme la reproduction en masse d'un circuit neuronal caractéristique. Ce circuit neuronal peut alors correspondre à la *population de neurones* sur laquelle la modélisation mésoscopique s'appuie. En particulier, le cas de la régularité du cortex cérébral est bien connu, puisqu'il est étudié depuis plus d'un siècle (Ramon y Cajal, 1909). L'étude du cortex cérébral a amené à formuler l'hypothèse de la *colonne corticale* comme étant son circuit fondamental (Mountcastle, 1957).

Connectivité La complexité de la connectivité au sein du cerveau en général et du cortex cérébral en particulier nécessite la mise en place d'un paradigme simplificateur, tout comme la population de neurones permet de considérer des réseaux à plus grande échelle. Ainsi, à l'échelle mésoscopique, on ne considère plus les connexions entre des neurones, mais des profils de connectivité entre des régions du cerveau, issus des travaux en neurobiologie sur les statistiques de connectivité dans le cerveau. Dans le cas du cortex cérébral, qui est au coeur de notre travail, la connectivité est souvent décrite à l'aide de trois flux d'informations auxquels des fonctions sont associées : le flux *montant*, le flux *latéral* et le flux *descendant*.

Dans ce chapitre, nous allons tout d'abord nous intéresser aux justifications biologiques de la modélisation mésoscopique du cortex cérébral. Par ce biais, nous introduirons les notions de

colonne corticale, micro-colonne, macro-colonne et aire corticale.

Les bases biologiques étant posées, nous verrons comment cette nouvelle unité de calcul peut être modélisée, en nous appuyant sur quelques modèles de la littérature.

Après quoi, nous nous intéresserons à la connectivité à l'échelle mésoscopique. Nous présenterons les principes généraux qui régissent la connectivité dans le cortex cérébral, puis nous aborderons un par un les trois flux d'informations : leurs bases biologiques, leurs fonctions, et les modèles qui en découlent.

2.1 Anatomie et physiologie du cortex cérébral

La neurobiologie nous renseigne sur le cortex cérébral de deux manières. D'une part, elle nous apporte des informations sur la structure du réseau neuronal - les types de neurones, leurs quantités, leurs interconnexions. D'autre part, elle nous apporte des informations sur les dynamiques d'activité du réseau en fonctionnement. La combinaison de ces informations nous amène à émettre des hypothèses sur l'organisation fonctionnelle du cortex cérébral.

L'étude du cortex cérébral par diverses techniques d'imagerie révèle une certaine régularité dans son organisation. En appliquant la technique d'imprégnation argentique mise au point par Camillo Golgi, Ramon y Cajal (1909) a été le premier à mettre en évidence l'organisation particulière des neurones dans le cortex cérébral. Depuis ces travaux précurseurs, les progrès dans le domaine de l'imagerie cérébrale ont permis le développement d'une vision plus précise de l'organisation du cortex cérébral.

2.1.1 Structure laminaire

Déplié, le cortex cérébral est une structure plane composée de six couches. Chaque couche est identifiable par sa composition particulière (fig. 2.1). On observe une variation de l'épaisseur pour certaines couches (en particulier la couche IV) en fonction de l'emplacement dans le cortex cérébral.

Les six couches du cortex cérébral sont classées dans trois catégories : *granulaire* (IV), *infra-granulaire* (V-VI) et *supra-granulaire* (I-III). La couche granulaire (IV) concentre la majorité des connexions entrantes. Il pourra s'agir d'afférences extra-corticales telles que la *voie thalamo-corticale* amenant l'information sensorielle depuis le thalamus, mais aussi d'afférences intra-corticales, en provenance d'autres régions du cortex cérébral. Les couches V-VI projettent des connexions vers de nombreuses régions extra-corticales. Les projections des couches I-III sont de type cortico-corticale.

2.1.2 Types de neurones

On dénombre une quarantaine de types de neurones dans le cortex cérébral, qu'on regroupe généralement dans deux catégories : les *neurones de projection* et les *interneurones* (environ 20% à 25% des neurones du cortex cérébral). Les neurones de projection sont majoritairement pyramidaux et excitateurs. Les interneurones sont majoritairement inhibiteurs et forment des connexions locales.

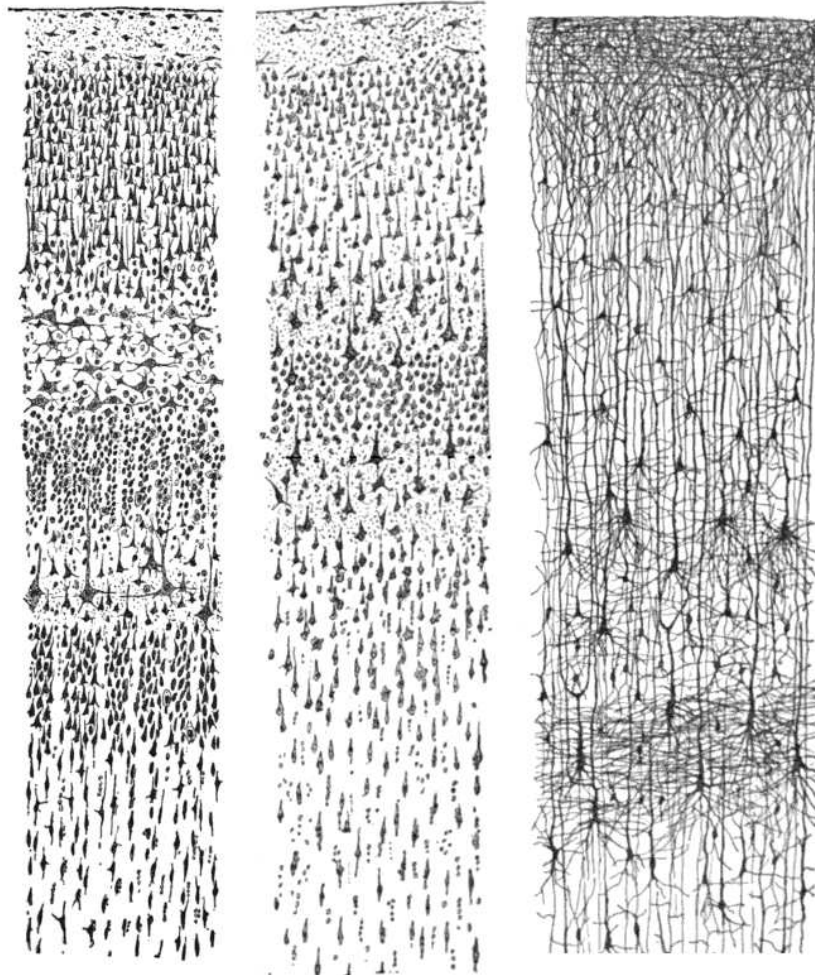


FIGURE 2.1 – Vue en coupe du cortex cérébral, dessin de Ramon-y-Cajal. A gauche : cortex visuel d'un adulte humain (teinture Nissl) ; au milieu : cortex moteur d'un adulte humain (teinture de Nissl) ; à droite : cortex d'un enfant humain âgé de six semaines (teinture de Golgi). La teinture de Nissl fait apparaître les corps cellulaires. Les variations de densité et de forme des corps cellulaires marquent une organisation laminaire. La teinture de Golgi fait apparaître les axones et les dendrites. On observe la disposition verticale des axones de neurones pyramidaux.

2.1.3 Connectivité

La connectivité au sein du cortex est majoritairement locale, s'étendant jusqu'à $\approx 50\mu m$. On distingue en particulier une forte connectivité verticale (Lorente de Nó and Fulton, 1949) : de nombreux neurones ont leur corps cellulaire situé dans une couche, tandis que les terminaisons dendritiques et axoniques traversent toute l'épaisseur du cortex cérébral.

Les connexions distantes sont plus rares et partent des couches supra-granulaires du cortex cérébral. A mesure que l'on monte vers la surface, les connexions sont de plus en plus distantes. De plus amples informations sur la connectivité du cortex cérébral peuvent être trouvées dans Thomson and Lamy (2007); Hagmann et al. (2008); Kandel et al. (2000).

2.1.4 Colonne corticale

Les observations anatomiques présentées précédemment indiquent une structure répétitive dans tout le cortex, mais ne permettent pas de délimiter clairement les dimensions d'un éventuel circuit cortical fondamental. La mise en parallèle des observations anatomiques et physiologiques permet cependant de poser la notion de colonne corticale, comme étant l'unité fonctionnelle fondamentale du cortex cérébral.

Activité cohérente verticalement

Mountcastle (1957) montre qu'en introduisant une électrode perpendiculairement à la surface du cortex cérébral, les neurones enregistrés sur le chemin de l'électrode ont des propriétés très similaires. D'une part, ils partagent le même *champ récepteur*, c'est à dire qu'ils ne réagissent qu'à des stimulations provenant d'un emplacement précis. Il pourra s'agir par exemple d'un endroit précis dans le champ de vision, ou encore d'une petite surface de peau. D'autre part, ils réagissent de la même manière aux stimuli locaux (même propriétés d'excitabilité). En revanche, si l'électrode est introduite de biais, les neurones enregistrés sur le chemin ont des propriétés différentes. Il écrit :

[...] These data, together with others upon the response latencies of the cells of different layers of the cortex to peripheral stimuli, support an hypothesis of the functional organization of this cortical area. This is that the neurons which lie in narrow vertical columns, or cylinders, extending from layer II through layer VI make up an elementary unit of organization, for they are activated by stimulation of the same single class of peripheral receptors, from almost identical peripheral receptive fields, at latencies which are not significantly different for the cells of the various layers. [...]

Cette observation est la base du concept de *colonne corticale* : un circuit cortical composé de neurones disposés perpendiculairement à la surface du cortex, formant l'unité fonctionnelle du cortex. Mountcastle estime à entre 300 et 500 μm le diamètre de cette colonne corticale. Cette idée rejoint la forte connectivité verticale décrite précédemment.

Micro-colonne et macro-colonne

En utilisant le même procédé, Hubel and Wiesel (1962, 1968) ont observé dans le cortex visuel du chat puis du macaque cette même invariance verticale. Ces travaux étendent la notion de colonne corticale proposée par Mountcastle. En effet, les auteurs observent deux phénomènes : la dominance oculaire et la sélectivité à l'orientation.

Dominance oculaire A l'échelle proposée par Mountcastle, de l'ordre de $500\mu m$, on observe des bandes dans lesquelles les neurones sont plus sensibles aux stimulations provenant d'un oeil.

Sélectivité à l'orientation On observe une sensibilité dépendant de l'orientation du stimulus visuel. En présentant une barre dans le champ de vision du chat et en faisant varier l'orientation de cette barre, on constate une variation de la réponse des neurones. Les auteurs constatent que cette sélectivité varie tout comme la dominance oculaire, mais à une échelle plus petite, de l'ordre de $50\mu m$. La figure 2.3 illustre la répartition de la sélectivité à l'orientation à la surface de l'aire V1.

Ceci amène à raffiner la proposition de Mountcastle, en faisant la distinction entre la *micro-colonne* ($\approx 50\mu m$) formée de 100 à 300 neurones et la *macro-colonne* ($\approx 300 - 500\mu m$) formée de 8000 à 24000 neurones. D'un point de vue fonctionnel, on peut dire que la macro-colonne est un ensemble de micro-colonnes recevant la même information dans la couche IV (même champ récepteur), mais n'effectuant pas le même traitement dessus (propriétés d'excitabilité différentes).

Limites de la notion de colonne corticale

En plus de 50 années de recherche sur la colonne corticale, de nombreuses expériences ont été menées, montrant que les propriétés d'une organisation en colonne corticale ne sont pas retrouvées chez tous les mammifères. Horton and Adams (2005) présentent une revue des travaux en neurobiologie sur la colonne corticale, et conclut que si la notion de colonne corticale est couramment utilisée, son existence biologique en tant que structure clairement isolée est sujette à débats.

S'il semble difficile de définir précisément les limites anatomiques de la colonne corticale, cette notion reste néanmoins intéressante pour expliquer l'organisation fonctionnelle du cortex cérébral. C'est d'ailleurs sur cette notion que s'appuient de nombreux travaux de modélisation corticale, dont nous présenterons quelques exemples dans la suite de ce chapitre.

2.1.5 Aire corticale

La notion d'aire corticale est associée à un effort de modélisation du cortex cérébral à une échelle plus large que la colonne corticale. Dans ce paradigme, le cortex cérébral est perçu comme étant composé de régions (les aires corticales) interconnectées, ayant chacune une fonction particulière.

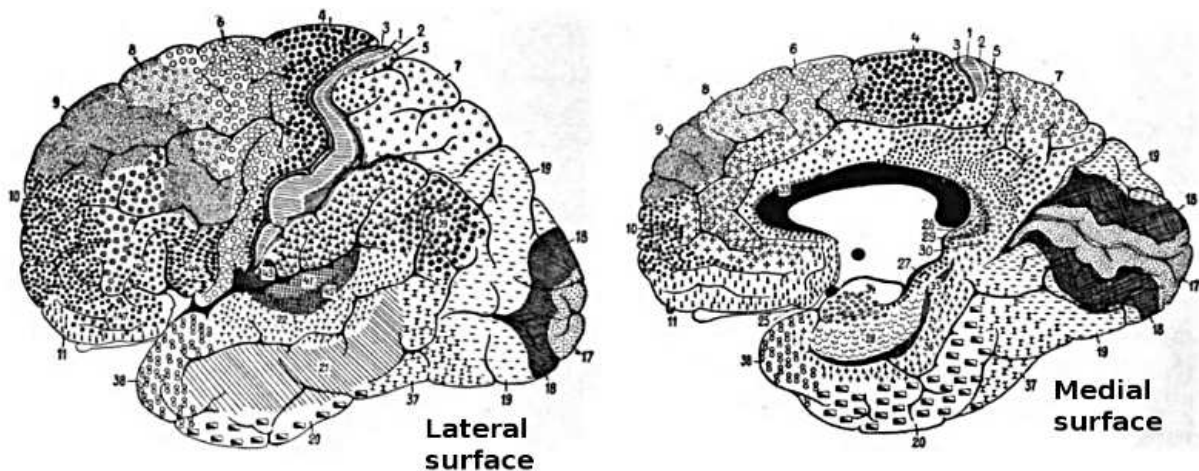


FIGURE 2.2 – Aires de Brodmann : division du cortex cérébral basée sur ses propriétés cytoarchitectoniques.

Subdivision anatomique et physiologie

Dans ses travaux sur la cytoarchitecture du cortex cérébral humain, Brodmann (1909) montre une variation structurale du cortex selon l'emplacement considéré : variation du nombre et de l'épaisseur des couches laminaires, différence d'arborisation dendritique, etc. Ce constat l'amène à proposer une subdivision du cortex en 52 aires, en se basant sur ces variations (fig 2.2).

Brodmann (1909) fait le parallèle entre ces observations anatomiques et les informations obtenues par l'étude des lésions cérébrales. Grâce à l'observation de leurs effets, on sait depuis longtemps que certaines régions du cortex cérébral jouent un rôle important dans des fonctions cognitives spécifiques. Dès 1861, Paul Broca a montré qu'une région située dans le lobe frontal postérieur gauche (aires 44/45 dans la classification de Brodmann, aujourd'hui nommée aire de Broca) joue un rôle primordial dans l'utilisation du langage. Ce parallèle amène Brodmann à associer une fonction cognitive à chaque aire du cortex cérébral. Les progrès de l'imagerie cérébrale durant les dernières décennies (tomographie par émission de positrons, magnéto-encéphalographie, électroencéphalographie) permettent de visualiser l'activité cérébrale *in vivo*, montrant effectivement des foyers d'activité différents selon la tâche cognitive effectuée. Cependant, les tâches cognitives mobilisent généralement de nombreuses régions - l'association entre une aire et une fonction est donc toute relative.

La variation structurale au sein du cortex cérébral est toute relative ; elle reste relativement faible en comparaison des différences structurales fortes qui existent entre les différentes régions sous-corticales du cerveau. Dans ce contexte, le fait d'observer une spécialisation fonctionnelle au sein du cortex cérébral implique que la structure corticale possède un certain degré de généricité, permettant la mise en place de traitements variés de l'information ; d'où l'intérêt particulier des neurosciences computationnelles pour ladite structure.

Catégorisation fonctionnelle

Dès 1870, John Hughlings Jackson propose une organisation hiérarchisée du cortex cérébral, qui a depuis été justifiée expérimentalement. Selon lui, le cortex cérébral possède trois types d'aires : (1) les aires sensorielles correspondant aux entrées, (2) les aires motrices correspondant aux sorties et (3) les aires associatives recevant l'information des aires sensorielles et motrices. Ces dernières permettent la mise en relations des modalités sensorielles et motrices. Cette mise en relation est à la base de la construction de représentations plus complexes, permettant à terme la mise en place de fonctions cognitives avancées.

Aires sensorielles Les aires sensorielles intègrent l'information sensorielle reçue par le biais des connexions extra-corticales entrantes (couche IV). À l'exception de l'olfaction, ces signaux transitent par le thalamus formant la *voie thalamo-corticale*.

Aires associatives Les aires associatives intègrent l'information en provenance d'aires corticales dites inférieures au sens de la vision hiérarchisée du traitement de l'information. Ces intégrations permettent la construction de représentations plus riches, combinant des informations issues de plusieurs modalités sensorielles et/ou motrices. On distingue les aires associatives *unimodales* et *multimodales*. Les aires unimodales se situent en périphérie des aires sensorielles et traitent l'information provenant de l'aire sensorielle adjacente uniquement. Les aires multimodales effectuent l'association de plusieurs modalités sensorielles.

Aires motrices Les aires motrices sont chargées du contrôle moteur, permis par la transmission de leur l'activité par le biais des connexions extra-corticales sortantes (couches V/VI). Les aires motrices supérieures se chargent de la planification des mouvements, tandis que les aires primaires s'occupent de leur exécution.

Organisation fonctionnelle au sein d'une aire

Nous avons précédemment cité les travaux de Hubel and Wiesel (1962, 1968) pour définir les notions de micro- et macro-colonne. Les auteurs ont observé que le cortex visuel réagit localement en fonction de l'orientation du stimulus présenté, permettant d'estimer la micro-colonne à environ 50 μm .

Cette observation n'est cependant pas la seule : la variation de la sélectivité à l'orientation au sein du cortex visuel primaire n'est pas aléatoire chez le chat et le primate. La pénétration en biais d'une électrode au sein du cortex visuel primaire montre une variation progressive de l'orientation préférentielle des neurones enregistrés sur son chemin. Ceci indique que les micro-colonnes voisines sont sélectives à des orientations similaires.

À l'aide d'une technique de coloration sensible au courant électrique, Blasdel and Salama (1986) présente une carte des orientations préférentielles dans l'aire V1 du macaque (fig. 2.3). Cette carte montre bien que la sélectivité à l'orientation varie en règle générale de façon progressive, avec quelques exceptions. On observe des cassures dans la carte, marquant un changement brusque de l'orientation préférentielle, qu'on appelle *pinwheels*.

L'organisation du cortex visuel primaire étant largement dépendante de l'environnement dans lequel l'individu grandit (Wiesel and Hubel, 1965), cette organisation fonctionnelle particulière

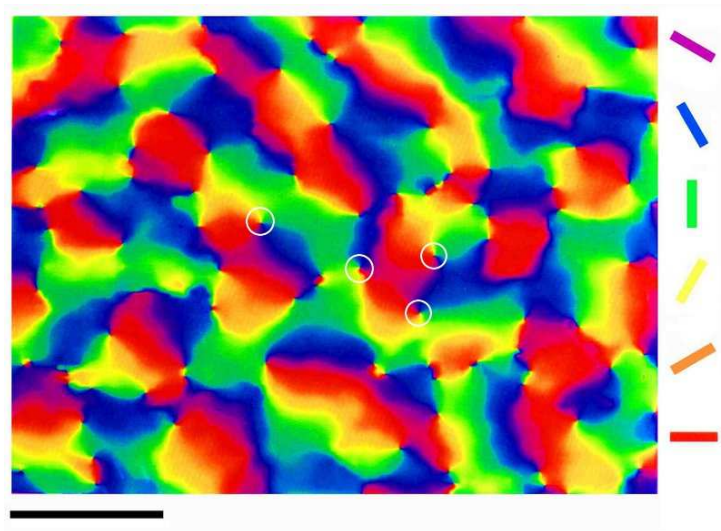


FIGURE 2.3 – Répartition des orientations préférentielles dans l'aire V1 du macaque, adapté de Blasdel and Salama (1986). Le trait noir représente 1mm sur la surface corticale. Les ronds blancs indiquent quelques exemples de *pinwheels*.

est certainement liée au développement du cortex visuel. Cet aspect sera traité plus en détail dans la partie II sur l'apprentissage.

2.1.6 Résumé

Les données anatomiques et physiologiques sur le cortex cérébral nous apportent des pistes pour la modélisation mésoscopique. La notion de colonne corticale pose un bon repère pour la granularité de la représentation, l'unité de calcul modélisée. L'organisation en macro-colonnes et en aires corticales nous oriente vers la construction de cartes corticales composées de colonnes disposées de façon à former un maillage bi-dimensionnel. Enfin, les statistiques de connectivité nous apportent des informations sur la manière de construire la connectivité entre les colonnes corticales du réseau.

2.2 Schémas de connectivité

La neurobiologie apporte de nombreuses informations sur les principes généraux gouvernant la connectivité au sein du cerveau en général, et du cortex cérébral en particulier. Ces informations statistiques permettent de raisonner à grande échelle en terme de schémas de connectivité.

2.2.1 Structure hiérarchisée

Les traitements complexes effectués au sein du cortex cérébral sont généralement mis en place par une hiérarchie de traitements simples, l'ensemble formant un *système fonctionnel*. Ainsi, le cortex visuel traite l'information sensorielle par une série de représentations de complexité

croissante : des propriétés telle que l'orientation locale et la couleur détectée dans l'aire V1 permettent la construction de contours dans l'aire V2, la détection du mouvement dans l'aire MT, puis la reconnaissance de formes etc.

L'étude des lésions cérébrales a amené à identifier deux voies de traitement dans le cortex visuel, que l'on nomme voie temporale et voie pariétale. Ungerleider et al. (1982) proposent de formaliser la fonction de ces deux voies de communication et de traitements en *what* et *where*. La voie temporale (*what pathway*) effectue la reconnaissance des objets perçus, hypothèse justifiée par l'activation de neurones réagissant à des stimuli de plus en plus complexes à mesure que l'on progresse dans la voie temporale. La voie pariétale (*where pathway*) traite quant à elle ce qui a trait au positionnement des objets dans l'espace.

Depuis sa proposition, l'interprétation *what/where* a été remise en question ; en particulier, Goodale and Milner (1992) argumentent en faveur d'une interprétation alternative de la fonction de la voie pariétale :

Our alternative perspective on modularity in the cortical visual system is to place less emphasis on input distinctions (e.g. object location versus object qualities) and to take more account of output requirements. [...] it is proposed that this distinction ('what' versus 'how') - rather than the distinction between object vision and spatial vision ('what' versus 'where') - captures more appropriately the functional dichotomy between the ventral and dorsal projections.

Selon les auteurs, la voie pariétale construit donc une représentation de la façon d'atteindre et de manipuler l'objet, plutôt qu'une représentation de son positionnement.

2.2.2 Connectivité ordonnée

On observe dans le cortex cérébral des projections axoniques (*pathways*) regroupant de nombreuses connexions et reliant les aires formant un système fonctionnel. Ces projections ont la particularité d'être ordonnées : se projetant d'un espace bidimensionnel (la carte corticale source) vers un autre espace bidimensionnel (la carte corticale destination), elles tendent à transposer la topologie de la source vers la destination.

Ainsi, on observe dans le cortex visuel un phénomène de *rétinotopie*, qui est la préservation de la topologie de l'espace visuel perçu par la rétine dans les aires corticales (et sous corticales) impliquées dans la vision. Le même phénomène de préservation de la topologie est observé dans le cortex somatosensoriel (somatotopie) et dans le cortex auditif (tonotopie).

2.2.3 Connectivité bidirectionnelle

Il faut cependant relativiser cette conception séquentielle du traitement de l'information dans les systèmes sensoriels. Si, effectivement, l'information sensorielle subit des traitements consécutifs pour permettre des représentations de plus en plus abstraites, la communication ne se fait pas pour autant à sens unique.

Trappenberg (2010) cite l'exemple de la boucle thalamo-corticale : la connectivité descendante entre les aires sensorielles primaires et le thalamus est très dense, estimée à dix fois la densité de la connectivité montante ! Cette observation amène à penser que le thalamus fait plus que relayer l'information sensorielle vers le cortex cérébral. De manière générale, on observe pour

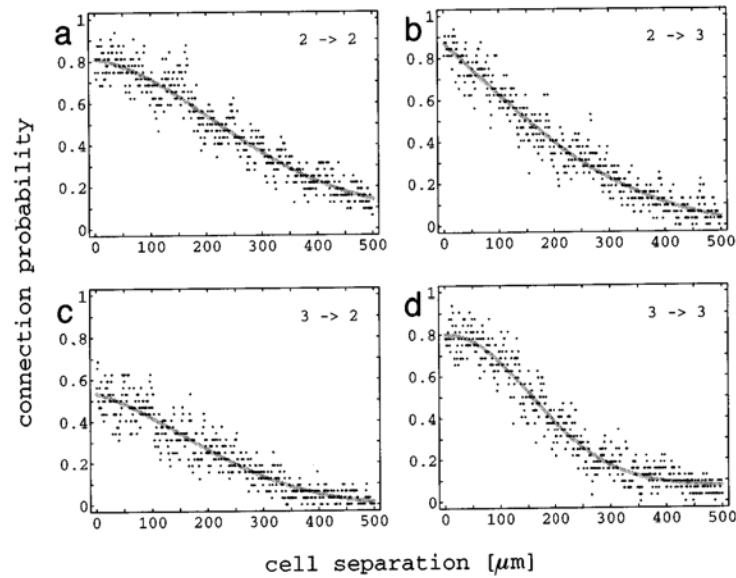


FIGURE 2.4 – Probabilité de connexion en fonction de la distance entre deux neurones pyramidaux. Mesures prises chez le rat, dans les couches corticales de II/III. Adapté de Hellwig (2000).

chaque projection d’une aire inférieure vers une aire supérieure, l’existence d’une projection dans le sens inverse. Bien souvent, la densité de cette *connectivité descendante* est comparable à celle de la *connectivité montante*. La présence de ces connexions “feedback” remet donc en question la vision séquentielle du traitement de l’information. Aujourd’hui encore, le rôle précis de cette connectivité est sujet à débat.

2.2.4 Connectivité intra-aire

Au sein d’une aire, on distingue globalement deux schémas de connectivité associés à deux familles de neurones. Les *interneurones* forment des connexions à courte distance, majoritairement inhibitrices. En revanche, les *neurones de projection* forment des connexions plus complexes.

Hellwig (2000) présente une étude de la connectivité corticale entre les neurones pyramidaux des couches II/III chez le rat, sur une distance de 0 à 500 μm . Les résultats montrent une probabilité de connexion décroissante en fonction de la distance (fig 2.4).

Stettler et al. (2002) présentent des observations similaires pour les neurones de projection dans le cortex visuel (V1) du macaque. Jusqu’à 400-500 μm , les auteurs observent une répartition homogène des connexions (fig 2.5 A). Au delà de cette limite, les connexions forment des paquets (fig 2.5 B). La disposition particulière des connexions au delà de 500 μm laisse penser qu’il s’agit de connexions vers des cibles spécifiques, contrairement aux connexions plus courtes dont la répartition indique des connexions non-spécifiques. On notera que ces ordres de grandeurs coïncident avec les dimensions proposées pour la macro-colonne, ce qui indiquerait une interaction non-spécifique entre les micro-colonnes voisines, et une interaction spécifique entre les colonnes plus distantes, appartenant à des macro-colonnes différentes.

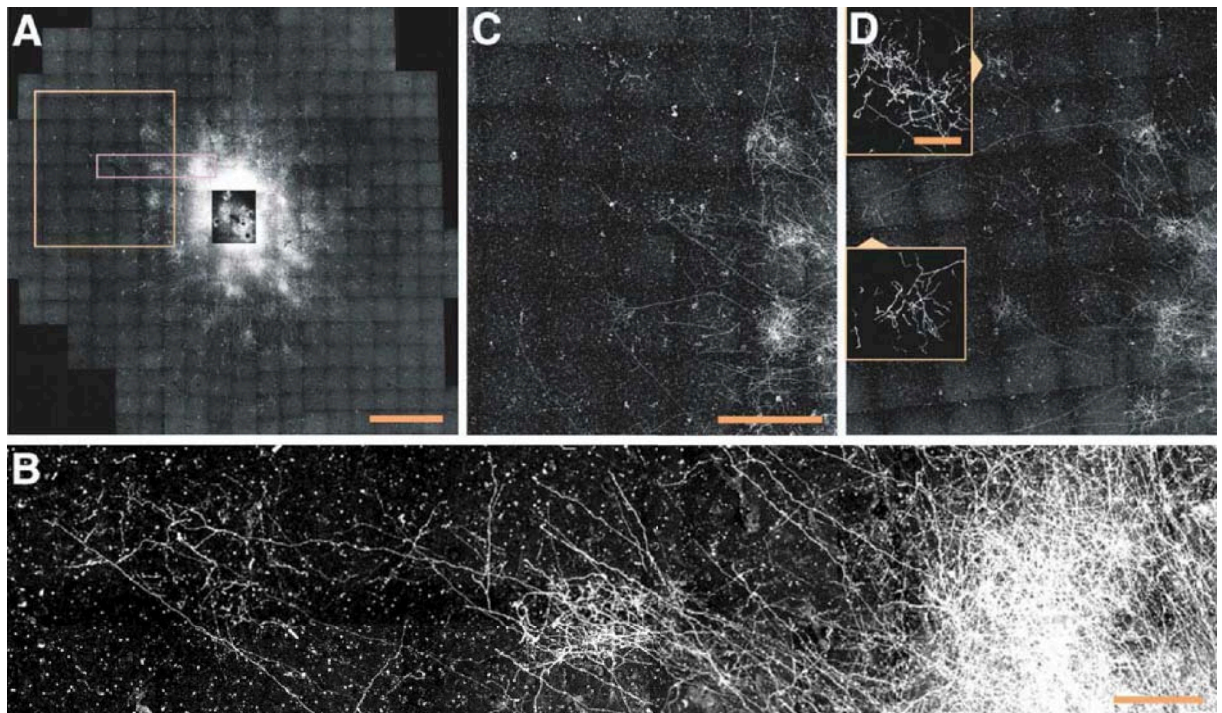


FIGURE 2.5 – Connectivité latérale dans V1 chez le macaque, adapté de Stettler et al. (2002). L'image A montre la distribution globalement uniforme des connexions à courte distance. Les images B,C et D montrent un regroupement des connexions en paquets à mesure que l'on s'éloigne du neurone observé.

2.3 Modélisation corticale

La section précédente nous a apporté les bases pour notre changement d'échelle de représentation : l'unité fondamentale d'un modèle mésoscopique du cortex cérébral n'est plus le neurone mais la population de neurones, à une échelle correspondant à la colonne corticale.

La modélisation corticale se base généralement sur la représentation du cortex à l'aide de cartes corticales. Une carte corticale est un regroupement d'unités de calcul formant un maillage généralement bidimensionnel, à l'image du cortex cérébral. La modélisation se positionnant à l'échelle mésoscopique, l'unité de calcul représente une population de neurones qu'on associe le plus souvent à la micro-colonne corticale, et une carte corticale est un ensemble de micro-colonnes appartenant à la même aire corticale, voire à la même macro-colonne suivant le modèle construit.

La complexité de l'unité de calcul est un facteur permettant de différencier les divers modèles corticaux. Nous allons à présent introduire plusieurs approches de la modélisation corticale de complexité variable en nous appuyant sur des modèles tirés de la littérature.

2.3.1 Unités simples

Une première approche consiste à simuler l'activité corticale à l'aide d'unités de calcul très simples. C'est par exemple le cas du modèle *Self-Organizing Maps* développé par Kohonen (1982). Dans ce modèle, chaque unité est un neurone prototype, et toutes les unités partagent un même flux montant.

$$y_i = F\left(\sum_{j=1}^n (w_{ij} - x_j)^2\right) \quad (2.1)$$

y_i est l'activité de l'unité i de la carte, $\mathbf{x}$ le vecteur d'entrées partagé par toutes les unités, $\mathbf{w}_i$ le vecteur de poids associé à l'unité i . $F(x)$ est une fonction radiale :

$$F(x) = e^{-\beta x} \quad (2.2)$$

On notera que le modèle SOM a donné naissance à de nombreux modèles dérivés tels que (Wiemer, 2000; Girod et al., 2007). Nous avons donc une modélisation discrète du cortex cérébral, où chaque unité effectue séparément un calcul simple. Nous parlerons plus tard de la connectivité latérale, qui permet les interactions entre les unités.

2.3.2 Champs neuronaux

Les champs neuronaux dynamiques (*Dynamic Neural Fields*, DNF) (Amari, 1977) approchent la modélisation du tissu cortical en s'inspirant de la notion de champ continu en physique. On exprime l'activité du champ neuronal par une fonction dépendante de l'espace et du temps, et le profil des connexions à l'intérieur de ce champ par une fonction dépendante de l'espace (position de la source et de la destination). L'équation 2.3 exprime un champ neuronal continu uni-dimensionnel.

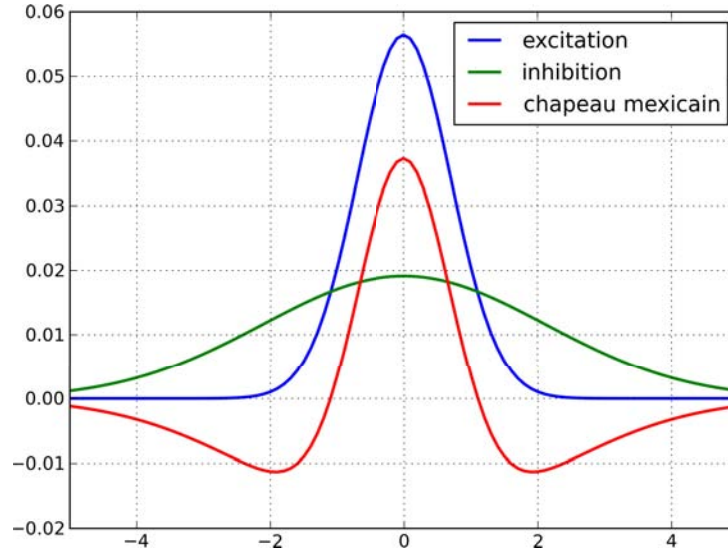


FIGURE 2.6 – Exemple de profil de poids en différence de gaussiennes, ou “chapeau mexicain”.

$$\tau \frac{\partial u(x, t)}{\partial t} = -u(x, t) + \int_y w(x, y) r(y, t) dy + I^{ext}(x, t) \quad (2.3)$$

$$r(x, t) = g(u(x, t)) \quad (2.4)$$

- $u(x, t)$ est le potentiel au point x et au temps t ,
- $w(x, y)$ est le poids synaptique reliant les points x et y ,
- $r(x, t)$ est l’activité sortante au point x et au temps t ,
- $g(u)$ est la fonction d’activation du champ,
- $I^{ext}(x, t)$ est la stimulation extérieure reçue au point x et à au temps t .

Précédemment, nous avons décrit la connectivité latérale intra-aire comme étant composée de connexions excitatrices et inhibitrices dont la densité décroît progressivement avec la distance. Ce profil de connectivité est couramment modélisé dans les champs neuronaux dynamiques, sous la forme d’une différence de gaussiennes. La première gaussienne représente le profil des connexions excitatrices en fonction de la distance, et la seconde le profil des connexions inhibitrices (fig 2.6). La gaussienne inhibitrice a une variance plus large que la gaussienne excitatrice, donnant à la distribution l’aspect d’un “chapeau mexicain”.

Le profil de connectivité en chapeau mexicain donne donc une excitation locale et une inhibition distante. A première vue, ceci semble contradictoire avec la connectivité observée dans le cortex cérébral, qui est composée de connexions inhibitrices locales et de connexions excitatrices distantes. Cette contradiction est expliquée par l’échelle observée : la connectivité inhibitrice est interne à la macro-colonne, permettant la compétition entre les micro-colonnes proches, tandis que les profils de poids des champs neuronaux représentent la connectivité telle qu’on peut la retrouver dans la couche III, reliant les macro-colonnes d’une aire corticale. A cette échelle, on retrouve donc une densité de connectivité excitatrice qui décroît avec la distance, tandis que l’inhibition atteint des régions distantes par propagation de proche en proche.

Les profils de poids latéraux basés sur une excitation locale et une inhibition globale ont fait l’objet d’études pour leurs propriétés de stabilité (Amari, 1977; Giese, 1998; Taylor, 1999). Dans

certaines conditions (bon équilibre entre l'inhibition et l'excitation), ce type de connectivité entraîne la formation de "bulles d'activité", c'est à dire d'une activité localisée et stable au sein de la carte. Les intérêts sont multiples :

- l'inhibition globale met en place une compétition au sein de la carte faisant émerger un "vainqueur" au sein de la population - par exemple l'unité recevant la stimulation externe la plus forte. Nous décrirons plus en détail ce mécanisme de "winner take all" dans la section 2.4.2.
- L'excitation locale ajoutée à l'inhibition globale permet la formation d'une bulle d'activité autour du vainqueur, capable de s'auto-entretenir grâce à la connectivité récurrente excitatrice. Une bulle d'activité peut rester stable après le retrait du stimulus à l'origine de sa formation - propriété qui peut être utilisée pour représenter une *mémoire de travail*.
- Une bulle d'activité formée est très résistante à l'ajout de bruit et de distracteurs. Rougier (2006) le montre avec l'expérience suivante. (1) Une bulle d'activité est formée dans le champ neuronal à l'aide d'une stimulation montante localisée. (2) si on déplace progressivement la stimulation montante, la bulle suit le mouvement. (3) si on ajoute du bruit sur la stimulation montante, la bulle se maintient sur le stimulus d'origine. (4) Si on ajoute des distracteurs (i.e. d'autres stimulations montantes localisées similaires à la première), la bulle se maintient sur le stimulus d'origine. Ces propriétés peuvent être utilisées pour construire des mécanismes d'attention visuelle (Rougier and Vitay, 2006).

Les champs neuronaux sont donc une modélisation continue de l'activité corticale, mettant en avant les effets de la connectivité latérale sur les dynamiques d'activité au sein du champ.

2.3.3 Dynamique locale d'activité

Dans les modèles précédents, nous avons décrit l'activité locale par une valeur numérique. Nous avons vu que la colonne corticale montre une activité cohérente. Certains modèles s'intéressent à reproduire plus en détail la dynamique d'activité au sein de la colonne. Un intérêt de cette approche est que les valeurs obtenues par le modèle peuvent être comparées à des enregistrements de potentiels de champ locaux (*Local Field Potential*, LFP).

Pour illustrer ce principe, prenons l'exemple du modèle de Wilson and Cowan (1972). Ce modèle s'appuie sur la présence de deux grandes familles de neurones dans la colonne corticale, à savoir les neurones pyramidaux (excitateurs) et les interneurones (inhibiteurs). Plutôt que de représenter l'activité de la colonne par une valeur numérique, le modèle simule l'interaction entre ces deux populations antagonistes, et au sein de chaque population (fig. 2.7).

L'activité d'une population est exprimée par la proportion de ses neurones émettant un potentiel d'action à l'instant t , donnant deux variables $E(t)$ et $I(t)$ - $E(t) = 1$ correspond à tous les neurones excitateurs actifs, $E(t) = 0$ correspond à l'activité au repos. La proportion de neurones actifs à l'instant t est calculée en fonction de la stimulation reçue par la population (fig. 2.7), mais aussi l'historique d'activité de la population. En effet, le modèle tient compte de la notion de période réfractaire des neurones : si à l'instant t la moitié de la population est activée, alors cette moitié ne pourra pas être activée avant $t + r$, r la durée de la période réfractaire.

Ainsi, ce modèle basé sur deux variables est à même de modéliser des dynamiques d'activité locale complexes, en plus des dynamiques globales au sein d'une carte.

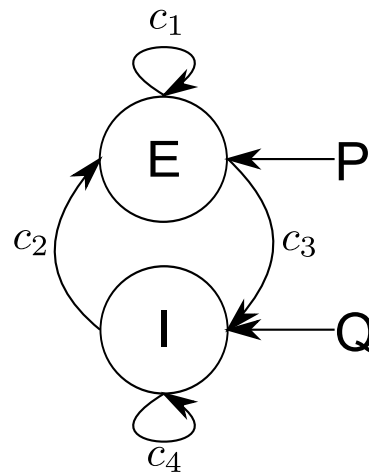


FIGURE 2.7 – Modèle de Wilson-Cowan : l'activité locale est modélisée par l'interaction d'une population excitatrice E et d'une population inhibitrice I . les paramètres c_1 à c_4 représentent le coefficient de connectivité entre une population source et une population destination. P et Q sont les stimulations externes reçues par les populations respectivement excitatrice et inhibitrice.

2.3.4 Micro circuit

Une autre approche de la modélisation de la colonne corticale consiste à construire un automate possédant une fonction bien précise, motivée par une théorie plus vaste du traitement de l'information dans le cortex cérébral.

On peut par exemple citer les travaux de Burnod (1989), qui propose un *automate colonne* représentant l'interaction entre les neurones pyramidaux et les interneurones au sein de la colonne corticale. Il écrit :

Compared with a cellular automaton, a columnar automaton has more functional and learning capacities. To take an analogy with computers, we could imagine a cell to correspond to one computer memory location, and a column as an area of computer memory in which may be stored different computer programs to perform different subfunctions.

On retrouve bien cette notion d'unité de calcul plus complexe et fonctionnellement plus riche que le neurone.

Structure

La colonne corticale de Burnod (1989) possède deux divisions correspondant respectivement aux couches inférieures (granulaire et infragranulaires) et supérieures (supragranulaires). Chaque division possède son propre état interne indépendant.

Cette séparation est liée aux connexions entrantes et sortantes du modèle. La division infragranulaire est la source et la destination des connexions extra-corticales, telles que la voie thalamo-corticale relayant l'information sensorielle. De son côté, la division supra-granulaire est la source et la destination des connexions intra-corticales, la connectivité au sein du cortex cérébral reliant les différentes régions plus ou moins distantes.

Activité

La colonne corticale de Burnod possède trois niveaux d'activité. L'activité de la division infra-granulaire est issue de l'intégration des stimulations de provenance extra-corticales. L'activité de la division supra-granulaire est issue de l'intégration des stimulations de provenance cortico-corticales. Le niveau d'activité global de la colonne est obtenu en combinant les activités des deux divisions.

Le modèle distingue trois états pour la colonne : *inactive*, *sensibilisée* et *activée*. En règle générale, la colonne corticale ne passe à l'état global actif que si les deux divisions sont stimulées. Lorsqu'une seule division est stimulée, la colonne est simplement sensibilisée.

La différence entre la sensibilisation et l'activation réside dans la propagation d'activité qui en découle. Lorsque la colonne est activée, elle émet une stimulation qui se propage par la voie extra-corticale. En revanche, si la colonne est seulement sensibilisée, la stimulation émise se propage par la voie intra-corticale.

Fonction

Le mécanisme de double activité est au centre du modèle. Considérons que l'activation de la division supra-granulaire d'une colonne représente un but à atteindre. Cette sensibilisation se propage à d'autres colonnes par le biais de la connectivité intra-corticale. Parmi les colonnes ainsi sensibilisées, certaines ont aussi leur division infra-granulaire stimulée. Il en résulte une activation de ces colonnes, qui provoque une stimulation extra-corticale (action) ayant pour effet de modifier l'environnement. Cette modification de l'environnement sensibilise de nouvelles colonnes, et le processus se poursuit jusqu'à ce que la colonne originale, représentant le but à atteindre, passe à l'état actif. Dans ce processus récursif nommé *arbre d'appels*, l'activation d'une colonne intermédiaire représente l'atteinte d'un but intermédiaire avant d'atteindre le but final.

Burnod (1989) propose donc une colonne corticale avec une fonction, à savoir la mise en relation d'une perception (activité infra-granulaire) et d'un but (activité supra-granulaire). Cette fonction est bien plus complexe que celle du simple neurone, et offre un paradigme permettant la mise en place de traitements complexes.

2.4 Modélisation de la connectivité corticale

La modélisation mésoscopique s'intéresse aux schémas de connectivité à grande échelle basés sur les statistiques de connectivité observées dans le vivant, plutôt que le détail de la connectivité entre les neurones. Une connexion entre deux unités de calcul représente donc un ensemble de connexions synaptiques entre deux populations de neurones, et le "poids" de cette connexion exprime le couplage entre les deux populations, dépendant à la fois de l'efficacité synaptique et de la densité des connexions.

Les modèles mésoscopique du cortex cérébral s'organisent en reprenant la disposition bidimensionnelle du cortex cérébral. Les colonnes sont regroupées pour former des *cartes corticales*. Cette organisation introduit une relation topologique entre les colonnes d'une carte corticale : on peut mesurer la distance entre deux colonnes.

En se basant sur cette notion de topologie du réseau, des connexions relient les colonnes. Les connexions sont généralement regroupées en projections intra- ou inter-cartes, les profils de poids de celles-ci étant déterminés par la relation topologique entre la source et la destination.

On formalise généralement la connectivité corticale par trois flux. *Le flux montant* apporte des informations de nature sensorielle en provenance de régions sous-corticale ou d'aires corticales inférieures. *Le flux latéral* permet l'interaction entre les colonnes au sein d'une aire corticale. *Le flux descendant* apporte une information de nature attentive en provenance d'une aire corticale supérieure. Nous allons maintenant nous intéresser à ce formalisme.

2.4.1 Flux montant

Le flux montant représente l'ensemble des connexions arrivant dans la couche IV d'une aire donnée. Ces connexions proviennent de régions sous-corticales ou d'aires corticales inférieures. Étant donné l'organisation hiérarchique du traitement de l'information dans le cortex cérébral, l'information véhiculée par le flux montant est généralement d'origine sensorielle.

Filtrage de l'information

Le rôle du flux montant est de véhiculer l'information sensorielle et de la filtrer afin d'en réduire la dimensionalité. Les travaux de Hubel & Wiesel (Hubel and Wiesel, 1962, 1968) sur le cortex visuel illustrent ce phénomène : à l'échelle de la micro-colonne, les neurones de l'aire visuelle V1 sont sélectif à l'orientation du stimulus visuel - l'intensité de l'activité de la colonne varie selon l'orientation d'un stimulus localisé dans une zone du champ de vision. Dans ce cas précis, on peut donc dire que de toute l'information contenue dans cette zone du champ visuel, la micro-colonne de V1 extrait une information sur l'orientation du stimulus.

Les profils de poids sur le flux montant sont donc définis de sorte que les unités de la carte corticale opèrent comme des filtres détectant la présence de certains stimuli dans le signal entrant, et cette présence est codée dans l'activité de l'unité.

Ici, la partie importante réside dans l'apprentissage des profils de poids, de sorte que les filtres soient adaptés à l'information transitant sur le flux montant. Nous parlerons de cet aspect dans la partie II sur l'apprentissage.

Codage de l'information au sein d'une carte corticale

Considérons une carte corticale composée de colonnes recevant la même information sur leur flux montant. Chaque colonne effectue un filtrage en fonction de son profil de poids, et code la présence d'un stimulus par son activité. L'activité de la carte dans son ensemble exprime donc la totalité de l'information extraite à partir du signal montant à un instant donné.

Lorsque l'information est ainsi codée collectivement par un ensemble d'éléments, on parle de *codage en population*. A mesure que la population croît, la quantité d'information représentable par la population croît exponentiellement. Foldiak and Endres (2008) explorent les propriétés du codage en population et distinguent plusieurs situations :

Codage local Si, pour tout signal entrant, la sortie est l'activation d'une seule unité de la population, la capacité d'encodage de la population est restreinte, proportionnelle au nombre d'unités la composant.

En revanche, l'information encodée est très simple à décoder, car il n'y a pas d'ambiguïté : l'activation d'une unité indique la reconnaissance d'une entrée spécifique.

Codage dense Si, pour tout signal entrant, la sortie est l'activation d'un sous-ensemble d'unités de la population, la capacité d'encodage de la population est immense, car elle croît exponentiellement avec la taille de la population : une population de N unités binaires peut coder 2^N entrées spécifiques.

En revanche, l'information encodée est plus difficile à décoder, car il y a beaucoup d'ambiguïté : une unité peut être active pour plusieurs entrées différentes, et il faut accéder à la sortie de toutes les unités pour être certain qu'il s'agit bien de la réponse à une entrée spécifique.

Codage "sparse" Il s'agit d'un compromis entre les deux approches précédentes : l'information est codée par l'activation d'un petit sous-ensemble de colonnes corticales. Cette approche augmente largement la capacité d'encodage de la carte, tout en gardant l'information assez facilement décodable.

Exemple : on souhaite reconnaître des formes géométriques colorées, à l'aide d'unités dont la sortie est binaire en fonction de l'entrée (active ou inactive selon la présence ou l'absence d'une caractéristique prédéterminée). Il y a 4 formes et 4 couleurs, soit 16 possibilités.

- Avec le *codage local*, il faut une unité pour chaque combinaison possible, soit 16 unités.
- Avec le *codage "sparse"*, on peut se contenter de 8 unités : une pour chaque forme, une pour chaque couleur. L'activation combinée de deux unités permet de connaître la forme et la couleur.
- Avec le *codage dense*, 4 unités suffisent à former 2^4 combinaisons d'activation, qui correspondent chacune à une forme. On notera au passage qu'aucune unité activée correspond à une détection.

Intuitivement, cet exemple montre bien que le codage local est très simple à lire : on associe une unité à une perception. Le codage "sparse" est simple à lire, à partir du moment où on a compris quelle groupe d'unités code quelle caractéristique. En revanche, le codage dense nécessite de connaître parfaitement la table d'association entre les perceptions et l'activité.

2.4.2 Flux latéral

Le flux latéral représente la connectivité locale au sein de l'aire corticale. L'interaction entre les colonnes corticales au sein d'une aire permet la mise en place de traitements complexes à l'échelle de la carte corticale.

Coopération et compétition latérale

La connectivité locale au sein du cortex cérébral est à la fois composée de connexions inhibitrices, par le biais des interneurones, et de connexions excitatrices par le biais des neurones de projection. Ces diverses connexions permettent des interactions variées entre les colonnes corticales voisines.

Prenons l'exemple d'une carte corticale composée de plusieurs colonnes. Toutes les colonnes partagent la même source d'information sur le flux montant, mais filtrent le signal différemment. Chaque colonne est donc la représentation d'un certain stimulus dans le signal montant.

De plus, il existe une connectivité latérale entre les colonnes de la carte, et cette connectivité peut être excitatrice ou inhibitrice.

Supposons que deux colonnes possèdent des connexions latérales inhibitrices réciproques. Un stimulus est présenté sur le flux montant et entraîne l'activation des deux colonnes. Alors, une inhibition réciproque se met en place entre les deux colonnes. La colonne la plus active possédant un plus grand pouvoir d'inhibition, elle va progressivement affaiblir l'activité de l'autre, pour finalement la dominer. La connectivité latérale inhibitrice met en place une *compétition* entre les colonnes, donc entre leurs représentations. Cette compétition impose une contrainte d'exclusion entre des représentations, de sorte qu'elles ne sont pas actives en même temps.

La connectivité latérale excitatrice a l'effet inverse. La stimulation d'une colonne par le flux montant se propage latéralement, facilitant l'activation d'autres colonnes. Cette *coopération* pose une contrainte de co-occurrence de plusieurs représentations au sein de la carte.

Ces phénomènes d'interactions latérales sont au centre de la modélisation par *champs neuronaux dynamiques*.

Résumé

La connectivité latérale permet d'altérer l'activité issue du flux montant pour lui donner une certaine cohérence à l'échelle de la carte corticale. Bien souvent, influencer l'activité d'une unité implique une influence sur les mécanismes d'apprentissage s'y déroulant. Ainsi, la connectivité latérale est souvent utilisée pour influencer l'apprentissage se déroulant sur la connectivité montante. L'apprentissage montant d'une unité n'est plus isolé, mais dépend de la carte corticale dans son ensemble, et du profil de la connectivité latérale. Nous étudierons ces mécanismes plus en détail dans la partie II sur l'apprentissage.

2.4.3 Flux descendant

Le flux descendant représente l'ensemble des connexions provenant des aires supérieures et arrivant les couches II/III de la colonne corticale. Le flux descendant a une influence modulatrice sur l'activité de la colonne corticale. Grossberg (2007) parle de *modulation sous le seuil d'activation* : L'activation de la colonne corticale est permise par la stimulation montante - le flux descendant ne fait que favoriser ou défavoriser l'activation.

Ce mécanisme de modulation est généralement associé au principe d'*attention*. Si on s'attend à percevoir un certain stimulus, cette attente se traduit par une sensibilisation (modulation favorisante) des colonnes corticales correspondant au stimulus concerné. Inversement, on peut sciemment ignorer certains stimuli pour se concentrer sur d'autres, jugés plus pertinents. Ce phénomène est très bien illustré par le *Cocktail Party Effect* (Arons, 1992) : dans un environnement bruyant ou beaucoup de personnes parlent en même temps, nous sommes capables de nous concentrer sur une conversation et d'ignorer le reste du bruit ambiant.

Ce principe est illustré par le modèle LAMINART (Grossberg, 2007), qui propose un micro-circuit cortical effectuant l'interface entre le signal pré-attentif (montant) et le signal attentif (descendant).

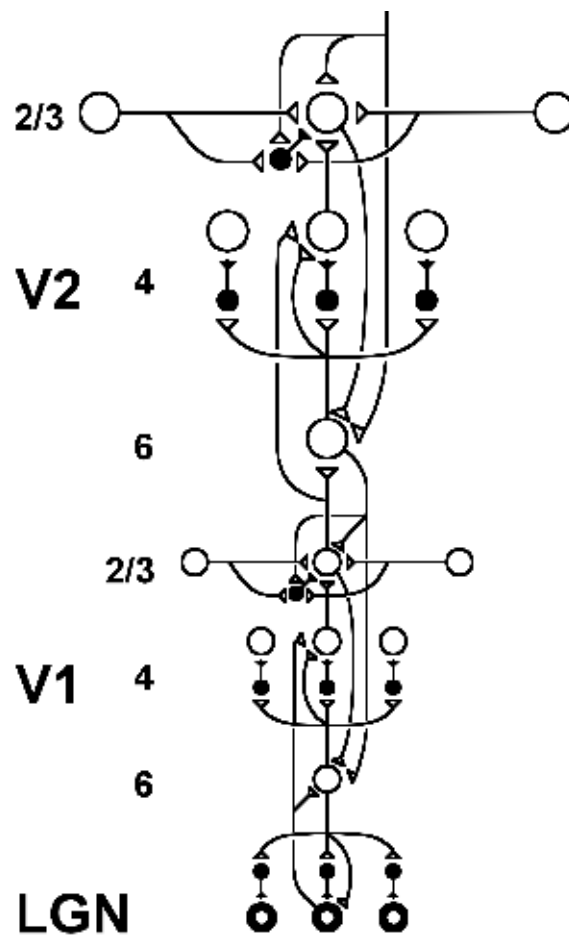


FIGURE 2.8 – Modèle de colonne corticale adapté de Grossberg (2007). Le schéma représente deux colonnes appartenant respectivement aux aires V1 et V2. Ce modèle propose un micro-circuit détaillé pour représenter la colonne corticale, composé de neurones pyramidaux excitateurs (ronds blancs) et d'interneurones inhibiteurs (ronds noirs).

LAMINART : Pre-attentive / attentive interface

Le modèle de Grossberg (2007) s'articule autour de l'intégration des trois flux montant, latéral et descendant au sein d'un micro-circuit inspiré de la structure laminaire du cortex (voir figure 2.8).

Dans ce modèle, la couche IV est le point de convergence des trois flux. Le flux montant dit "pré-attentif" arrive directement dans la couche IV, tandis que les autres transitent par les couches infra-granulaires, qui forment un circuit permettant la modulation de l'activité de la couche IV. Le flux descendant dit "attentif" est intégré dans les couches supra-granulaires, puis renvoyé vers la couche VI, participant ainsi à la modulation de l'activité de la couche IV, donc de l'activation de la colonne.

Ce mécanisme d'interface *pré-attentive/attentive* s'inspire de la *Théorie de Résonance Adaptative* (ART) (Carpenter and Grossberg, 2003). Le pré-attentif fait référence à l'information sensorielle, qui est traitée par la connectivité montante, mais qui ne fait pas de discrimination sur

l'information pertinente. L'attentif fait référence à l'information provenant d'autres régions du cerveau, orientant la perception vers ce qui est supposé pertinent.

2.5 Multimodalité

Notre travail étant centré sur la notion d'apprentissage multimodal, étudions la question de la multimodalité à l'échelle mésoscopique. Comme nous l'avons mentionné précédemment, il nous semble difficile d'appréhender l'intégration multimodale à l'échelle du neurone, qui ne fait vraisemblablement pas de distinction entre les potentiels d'action reçus selon la provenance ou l'information transmise. Notre passage à l'échelle mésoscopique nous amène à nous reposer la question de l'intégration multimodale au sein d'une structure plus complexe - dans notre cas, la colonne corticale.

Rappelons que la multimodalité n'est pas un phénomène exclusivement cortical : on trouve aussi des propriétés multimodales dans certaines structures sous-corticales telles que le colliculus supérieur (Stein and Meredith, 1993), les ganglions de la base (Chudler et al., 1995), le thalamus (Benedek et al., 1997) ou l'amygdale (Nishijo et al., 1988).

Les modèles multimodaux de la littérature sont majoritairement orientés vers l'intégration multisensorielle, en particulier dans le colliculus supérieur et s'appuient souvent sur la notion d'inférence bayésienne (Knill and Richards, 1996) pour modéliser la perception. N'Guyen (2010) fait un état de l'art du domaine et propose le premier modèle bayésien de l'intégration multisensorielle dans le colliculus supérieur. Notre travail quant à lui s'oriente plutôt vers la modélisation corticale biologiquement plausible.

Les travaux sur la modélisation de la multimodalité dans le cortex cérébral sont plus rares. On notera par exemple les travaux de Reynaud (2002), qui propose un modèle de la multimodalité basé sur une mémoire associative bidirectionnelle. Ce modèle a son intérêt mais ne cadre pas avec notre contrainte de plausibilité biologique : nous souhaitons construire un modèle qui puisse correspondre à un traitement cortical. On retiendra aussi les multiples itérations du modèle DARWIN d'Edelman (Almassy et al., 1998) qui s'intéressent à la multimodalité pour le développement de comportements chez un robot. Des modèles tels que Guigon (1993) ou Menard (2005) abordent la notion de multimodalité par le biais de l'intégration de flux multiples au sein de la colonne corticale. Nous avons vu que la colonne corticale est un circuit complexe recevant dans différentes couches des stimulations provenant de différentes régions. Cette connectivité corticale est le plus souvent représentée par trois flux - montant, latéral et descendant. Dans ce contexte, l'intégration multimodale réside dans la manière qu'à la colonne corticale d'intégrer séparément chaque flux pour ensuite les combiner.

Pour illustrer ce problème, étudions le modèle de Guigon (1993), qui s'intéresse en détail à l'intégration de plusieurs flux d'informations au sein de la colonne corticale.

Modèle de Guigon

Guigon (1993) présente un modèle de colonne corticale inspiré des travaux de Burnod (1989). Ce modèle conçoit la colonne corticale sous l'angle de la multi-modalité, de l'intégration de plusieurs flux d'informations. La généralité du modèle permet d'exprimer des interactions complexes, additives ou multiplicatives, entre un nombre arbitraire de modalités.

Le modèle sépare la colonne corticale en plusieurs étages nommés *divisions* $D^i, 1 \leq i \leq N$. Chaque division D^i est caractérisée par une sortie y^i , ainsi que par un ensemble de connexions entrantes $x_j^i, 1 \leq j \leq n_i$ partageant une même origine, et donc une même sémantique. On notera que si ces étages peuvent correspondre à des couches corticales, mais ce n'est pas nécessairement le cas. Chaque division calcule son activité à la manière d'un neurone formel, c'est à dire par le produit scalaire entre son vecteur d'entrée et son vecteur de poids.

$$u^i(t) = F\left(\sum_{j=1}^{n_i} w_j^i(t)x_j^i(t)\right) \quad (2.5)$$

où F est une fonction non-linéaire bornée en $[0; 1]$

La combinaison des multiples modalités se fait dans le calcul de la sortie de chaque division. Chaque division D^i calcule sa sortie y^i à partir de l'activité de toutes les divisions de la colonne, à l'aide de deux termes L^i et Q^i :

- L^i est un vecteur associant un coefficient à chaque division ;
- Q^i est une matrice associant un coefficient à chaque paire de divisions.

$$q^i(t) = \sum_{j=1}^N \sum_{k=j+1}^N Q_{jk}^i(t) u^j(t) u^k(t) \quad (2.6)$$

$$l^i(t) = \sum_{j=1}^N L_j^i(t) \cdot u^j(t) \quad (2.7)$$

$$y^i(t) = F(q^i(t) + l^i(t)) \quad (2.8)$$

Nous nous intéressons à ce modèle car il met en place un mécanisme d'intégration multimodale. Dans un premier temps, l'activité de chaque division est issue de ses entrées, mettant en place un filtre comme nous l'avons montré dans la section 1.2.3. Dans un second temps, le vecteur L^i et la matrice Q^i expriment des relations complexes entre les modalités convergeant vers une même colonne corticale : L^i exprime l'influence respective de chaque modalité tandis que Q^i exprime l'influence de la co-activation de deux modalités sur la sortie de la division.

Apprentissage et multimodalité

Dans le modèle que nous venons de décrire, l'intégration multimodale se base sur le choix de filtres adaptés à l'information traitée et de coefficients permettant la détection d'évènements corrélés. Cela ne nous dit pas comment ces filtres et ces coefficients sont acquis. Cette question de l'apprentissage est complexe, même pour le traitement d'une seule modalité. Nous l'aborderons en détail dans la partie II dédiée à l'apprentissage.

2.6 Conclusion

L'échelle mésoscopique propose un paradigme original de modélisation de l'activité cérébrale. Dans le cadre de la modélisation corticale, l'unité de calcul à laquelle on se réfère communément est la colonne corticale. Les colonnes corticales sont organisées en cartes bi-dimensionnelles,

introduisant une notion de topologie du réseau. Les cartes sont reliées par des projections de connexions dont les poids s'inspirent souvent des propriétés statistiques de la connectivité *in vivo*.

En tant qu'unité de calcul, la colonne corticale est une structure complexe dont la fonction est sujette à de nombreuses hypothèses. Ainsi, il existe de nombreux modèles de la colonne corticale, reflétant différentes visions de l'organisation du traitement de l'information au sein du cortex cérébral.

La connectivité corticale est quant à elle formalisée par trois flux nommés *montant*, *latéral* et *descendant*. Ce formalisme est globalement bien accepté, puisqu'on le retrouve dans de nombreux modèles de la littérature. De même, les rôles associés à ces flux sont similaires d'un modèle à l'autre : le filtrage de l'information pour le montant, la cohérence au sein d'une carte pour le latéral, et l'attention pour le flux descendant.

Le paradigme des trois flux d'information nous offre un cadre clair pour la suite de notre travail.

Dans un premier temps, les représentations doivent émerger sur le flux montant, dont le rôle est de filtrer l'information entrante. Ensuite, le flux latéral doit apporter une certaine cohérence aux représentations au sein d'une carte par le biais de son influence sur l'activité des colonnes corticales qui la composent. Ceci pose les bases du codage en population, une représentation distribuée au sein d'une carte, au lieu d'un ensemble de représentations indépendantes. Enfin, le flux descendant peut lui aussi apporter une modulation à l'activité de la carte. Cette modulation pourra elle aussi influencer l'apprentissage.

Ceci soulève deux questions : (1) Comment intégrer l'activité des différents flux au sein de la colonne corticale ? C'est le choix du circuit cortical utilisé. (2) Comment l'activité provenant des différents flux influence-t-elle l'apprentissage ayant lieu sur un flux ? Par exemple, comment le flux latéral influence-t-il l'apprentissage du flux montant ?

Deuxième partie

Apprentissage

La notion d'apprentissage est au coeur des neurosciences. Notre expérience quotidienne témoigne de l'évolution de nos fonctions cognitives. Chaque jour, nous apprenons, mémorisons, oublions en fonction de nos expériences. Ainsi, comprendre les mécanismes biologiques régissant les fonctions cognitives n'est qu'une partie de la question ; encore faut-il comprendre les mécanismes régissant l'acquisition et l'évolution de ces fonctions cognitives.

Dans la première partie, nous nous sommes intéressés à la relation entre les structures biologiques et les fonctions cognitives. Nous allons à présent nous intéresser à la relation entre la plasticité cérébrale - la capacité qu'à le cerveau de modifier sa structure - et la notion d'apprentissage. En effet, il est communément admis que les différentes formes de plasticité cérébrale - la *plasticité synaptique*, la *plasticité intrinsèque* et la *plasticité structurelle* - jouent un rôle important dans la capacité d'apprentissage.

La relation entre plasticité et apprentissage a été formalisée par Hebb (1949), donnant naissance à la notion d'apprentissage hebbien. On résume généralement ce principe en une phrase : *What fires together, wires together* - les neurones qui s'activent ensemble voient leurs connexions renforcées. Cette hypothèse a plus tard été vérifiée biologiquement, et est un pilier des algorithmes d'apprentissage dans les neurosciences computationnelles.

La notion d'apprentissage a été formalisée de nombreuses manières, dans divers champs s'intéressant à la cognition. Dans le domaine de l'IA connexionniste, on distingue généralement trois grands types d'apprentissage, que sont l'*apprentissage supervisé*, l'*apprentissage par renforcement* et l'*apprentissage non-supervisé*. Doya (1999) propose un parallèle entre ces différentes formes d'apprentissage et certaines régions du cerveau. Selon lui, le cervelet effectue un apprentissage supervisé, les ganglions de la base un apprentissage par renforcement et le cortex cérébral un apprentissage non-supervisé.

Apprentissage supervisé Le principe de l'apprentissage supervisé se base sur un agent apprenant et un superviseur. Le superviseur présente à l'agent (ici, un réseau de neurones artificiels) un "problème" dont il connaît à l'avance la solution. Le système retourne une solution généralement imparfaite. Le superviseur peut alors comparer sa solution avec la solution de l'agent, et renvoyer un "signal d'erreur" exprimant la distance (l'erreur) entre la solution de l'agent et la solution optimale. Ce signal d'erreur est alors utilisé par l'algorithme d'apprentissage de manière à réduire l'erreur à l'avenir.

Apprentissage par renforcement (Sutton and Barto, 1998) Woergoetter and Porr (2008) décrivent l'apprentissage par renforcement comme étant l'apprentissage par l'interaction avec l'environnement. L'agent apprend des conséquences de ces actions plutôt que par le biais d'un signal d'erreur explicite. L'agent reçoit un *signal de renforcement*, une valeur numérique de récompense exprimant le succès des actions entreprises. Le but de l'apprentissage par renforcement est de maximiser la somme des récompenses dans le temps.

Apprentissage non-supervisé L'apprentissage non-supervisé consiste à reconnaître des motifs récurrents dans les perceptions, indépendamment de toute considération de pertinence de l'information ainsi extraite. Barlow (1989) en donne une explication claire :

Completely nonredundant stimuli are indistinguishable from random noise. Thus, such a visual stimulus would look like a television set tuned between stations, and an auditory stimulus would sound like the hiss on an unconnected telephone line. [...] [R]edundancy is the part of our sensory experience that distinguishes it from noise; the knowledge it gives us about the patterns and regularities in sensory stimuli must be what drives unsupervised learning.

With this in mind one can begin to classify the forms the redundancy takes and the methods of handling it.

Du point de vue de la théorie de l'information, un motif récurrent dans le signal perceptif est une particularité de sa distribution statistique, une baisse de l'entropie. On parle donc de construction d'un modèle statistique de l'environnement perçu.

Il existe de nombreux modèles connexionnistes dérivés de l'apprentissage hebbien mettant en place de ces trois formes d'apprentissage. Par exemple, le perceptron (Rosenblatt, 1958) fait de l'apprentissage supervisé ; les *Self-Organizing Maps* (Kohonen, 1982) font de l'apprentissage non-supervisé ; Quoique moins nombreux, on trouve aussi des modèles hebbiens de l'apprentissage par renforcement (Pennartz, 1997). Dans le cadre de ce travail, nous nous intéressons à l'émergence de représentations multimodales, par conséquent à l'apprentissage non-supervisé dans le cortex cérébral.

Dans cette partie, nous aborderons d'abord les mécanismes biologiques de la plasticité cérébrale, qui jouent un rôle majeur dans l'apprentissage. Nous présenterons ensuite la notion d'apprentissage hebbien, ses variantes et les règles d'apprentissage qui en découlent. Enfin, nous nous intéresserons en détail à la règle d'apprentissage hebbienne BCM (Bienenstock et al., 1982), une règle permettant un apprentissage non-supervisé local. La suite de notre travail se basera sur cette règle d'apprentissage, à laquelle nous apporterons quelques modifications pour permettre la modulation de l'apprentissage, ce qui ouvrira la porte à des représentations complexes et distribuées.

Chapitre 3

Plasticité cérébrale

Citons la définition de la plasticité donnée par Konorski (1948) :

The application of a stimulus leads to changes of a twofold kind in the nervous system. [...] [T]he first property, by virtue of which the nerve cells *react* to the incoming impulse [...] we call *excitability*, and [...] changes arising [...] because of this property we shall call *changes due to excitability*. The second property, by virtue of which certain permanent functional transformations arise in particular systems of neurons as a result of appropriate stimuli or their combination, we shall call *plasticity* and the corresponding changes *plastic changes*.

La plasticité est donc le phénomène selon lequel un neurone recevant des stimuli appropriés voit ses propriétés d'excitabilité modifiées. Il s'agit donc d'une définition de la plasticité centrée sur le neurone.

Dans sa définition, Konorski parle de *stimuli appropriés ou de leur combinaison*. Il s'agit là d'un point important : la plasticité ne serait pas causée par n'importe quel stimulus. Ceci nous mène à une hypothèse fondamentale des neurosciences. Hebb (1949) propose le principe suivant pour expliquer l'apprentissage :

When an axon of a cell A is near enough to excite cell B or repeatedly or persistently takes part in firing it, some growth or metabolic change takes place in both cells such that A's efficiency, as one of the cells firing B, is increased.

Selon ce principe, le *stimulus approprié* de Konorski est l'activation concomitante d'un neurone présynaptique et d'un neurone post-synaptique. Celle-ci entraîne un renforcement de la relation entre les deux neurones, le neurone présynaptique devenant plus à même de provoquer l'activation du neurone post-synaptique. Ce principe, centré sur l'*interaction entre deux neurones*, propose un rôle fonctionnel à la plasticité cérébrale : permettre un *apprentissage associatif*.

Pour en comprendre les mécanismes biologiques, la plasticité cérébrale a fait l'objet de nombreuses études. On constate qu'elle prend de nombreuses formes et implique de nombreux mécanismes encore à l'étude aujourd'hui. Une présentation exhaustive étant hors du champ de ce travail, nous allons introduire les trois formes de plasticité cérébrale les plus étudiées :

- la *plasticité synaptique*, centrée sur l'évolution de l'efficacité des synapses ;
- la *plasticité intrinsèque*, centrée sur l'évolution de l'excitabilité des neurones ;

– la *plasticité structurelle*, centrée sur l'évolution de la structure du réseau neuronal.

Nous insisterons en particulier sur la plasticité synaptique, qui est la plus couramment modélisée.

3.1 Plasticité synaptique

L'hypothèse de la plasticité synaptique remonte aux travaux de Ramon y Cajal. Ce phénomène, d'une importance fondamentale pour le postulat de Hebb, implique que l'efficacité d'une synapse varie au cours du temps. L'existence de la plasticité synaptique a été vérifiée expérimentalement (Bliss and Lomo, 1973; Castellucci et al., 1978), mettant en avant plusieurs comportements et plusieurs mécanismes impliqués.

3.1.1 Propriétés de la plasticité synaptique

Les expériences sur la plasticité synaptique présentent deux critères pour décrire une modification synaptique : la nature de la modification (renforcement ou dépression) et la persistance de la modification (court terme ou long terme). On a alors quatre types de plasticité synaptique : *Long Term Potentiation* (LTP), *Long Term Depression* (LTD), *Short Term Potentiation* (STP) et *Short Term Depression* (STD).

Nature de la modification

Les premières preuves expérimentales de la plasticité synaptique remontent aux travaux de Bliss and Lomo (1973). Dans leur expérience, l'induction d'une stimulation à haute fréquence entraîne une augmentation de l'amplitude des potentiels post-synaptique.

Inversement, d'autres expériences ont montré que l'induction d'une stimulation à basse fréquence entraîne une diminution de l'amplitude des potentiels post-synaptique. Pour une revue de ces travaux, voir Linden and Connor (1995).

Persistance de la modification

Les expériences sur la limace des mers (Castellucci et al., 1978) ont montré une variabilité dans la persistance de la modification synaptique.

L'application d'un stimulus tactile bénin (pas de renforcement, qu'il soit positif ou négatif) sur les branchies de la limace des mers entraîne un réflexe de retrait. La répétition de ce stimulus entraîne un affaiblissement du réflexe de retrait. On nomme ce phénomène *habituation*.

On constate que la manière de stimuler entraîne une persistance plus ou moins longue de l'habituation. Une session de dix stimulations consécutives entraînera une habituation de quelques minutes, tandis que quatre sessions séparées de quelques heures à une journée entraîneront une habituation pouvant durer jusqu'à 3 semaines.

La plasticité semble donc opérer à différentes échelles temporelles. La modification à court terme (STP/STD) correspondrait à une adaptation du comportement sur le moment, tandis que la modification à long terme (LTP/LTD) serait liée à la consolidation d'une adaptation en vue

de la faire persister. On peut voir dans ce double mécanisme une conception de l'apprentissage : s'adapter rapidement, et plus tard consolider les adaptations pertinentes.

3.1.2 Phénomènes biologiques de plasticité synaptique

Les mécanismes biologiques entrant en jeu dans le phénomène de plasticité synaptique sont nombreux et complexe ; certains mécanismes ont cependant été mis à jour par les travaux en neurobiologie.

Court terme

La plasticité à court terme est généralement associée à une modification des propriétés de la synapse du côté présynaptique (Kandel et al., 2000). Après une période d'habituation (STD), on observe une réduction de la quantité de neurotransmetteurs émis lors de l'arrivée d'un potentiel d'action présynaptique. Inversement, la quantité de neurotransmetteurs augmente pour la sensibilisation (STP).

Long terme

La plasticité synaptique à long terme n'est pas encore comprise dans son intégralité. Les premières étapes du processus ont été isolées expérimentalement : l'augmentation de la concentration en ions Ca^{2+} associée au potentiel post-synaptique (voir section 1.1.2) active certaines enzymes dont l'action mène à une modification de l'efficacité synaptique. Le détail de ce renforcement est encore sujet à débats (Bear, 1996). Il semblerait que ce mécanisme basé sur les récepteurs NMDA permette à la fois la LTP et la LTD : une faible élévation de la concentration en ions Ca^{2+} provoquerait la LTD, tandis qu'une forte élévation provoquerait la LTP (Cumming et al., 1996). Les phénomènes de LTP et LTD étant centraux à la notion d'apprentissage, nous reviendrons dessus dans la suite de ce travail.

3.1.3 Spike-Timing Dependent Plasticity (STDP)

Dans le cadre d'expérimentations biologiques, les travaux de Markram et al. (1997); Bi and Poo (1998) ont démontré l'existence d'un lien entre la plasticité synaptique et le temps d'arrivée des potentiels d'action. En provoquant l'émission de potentiels d'actions dans deux neurones présynaptique et postsynaptique, on constate que l'arrivée répétée du PA présynaptique quelques millisecondes *avant* l'émission du PA postsynaptique cause la LTP de la synapse concernée. Inversement, l'arrivée répétée du PA présynaptique quelques millisecondes *après* l'émission du PA postsynaptique cause la LTD de la synapse concernée.

Il est possible d'interpréter ce mécanisme en considérant que si un PA arrive à un neurone quelques millisecondes avant que celui n'émette un PA à son tour, le PA émis par le neurone présynaptique a participé à l'activation du neurone postsynaptique - en conséquence de quoi la connexion entre les deux neurones est renforcée. En revanche, si le neurone postsynaptique n'émet pas de PA malgré la réception d'un PA provenant du neurone présynaptique, c'est que la communication entre les deux neurones n'était pas pertinente - en conséquence de quoi la connexion entre les deux neurones est affaiblie. On retrouve ici le principe fondateur de l'apprentissage hebbien : *What fires together wires together*.

3.2 Plasticité intrinsèque

La plasticité intrinsèque regroupe l'ensemble des mécanismes faisant varier l'excitabilité du neurone.

3.2.1 Propriétés de la plasticité intrinsèque

On associe généralement la plasticité intrinsèque à l'*homéostasie*, un phénomène de régulation visant à maintenir le niveau d'activité d'un neurone stable. Turrigiano and Nelson (2004) présente l'homéostasie comme étant un mécanisme nécessaire au bon fonctionnement des structures neuronales :

Over short timescales, the activity of a central neuron must fluctuate considerably, as these fluctuations carry information. Over long timescales, however, [...] forces that generate net increases or decreases in excitation over time will disrupt the function of central circuits if they are unopposed by homeostatic forms of synaptic plasticity.

Ici, le constat est qu'à stimulation présynaptique égale, le renforcement ou la dépression synaptique implique un changement dans la stimulation reçue par le neurone. Par des mécanismes variés - ici, une forme de plasticité synaptique est citée - l'homéostasie vise à normaliser l'activité moyenne du neurone quelle que soit l'évolution.

L'information étant codée dans les variations de l'activité électrique des neurones, cette régulation de l'activité permet de maintenir ses variations au sein d'une certaine plage, pour éviter la perte d'information par saturation ou affaiblissement du signal.

3.2.2 Mécanismes biologiques de la plasticité intrinsèque

La plasticité intrinsèque existe sous diverses formes. Desai et al. (1999) mettent en avant la plasticité de l'excitabilité dans un neurone de culture (neurone cortical pyramidal). L'expérience montre qu'empêcher un neurone de décharger pendant deux jours entraîne une forte chute de son seuil de décharge. Dans cette expérience, il s'agit donc bien d'une modification intrinsèque du neurone indépendante de l'activité présynaptique.

En présentant les bases biologiques possibles de l'homéostasie, Turrigiano and Nelson (2004) introduisent d'autres mécanismes biologiques assimilables à la plasticité intrinsèque, centrés autour de la notion de *plasticité synaptique homéostatique* - ou comment le fonctionnement des synapses peut mettre en place l'homéostasie.

L'article met en avant une régulation globale de l'efficacité des synapses excitatrices entrantes - permettant ainsi l'ajustement de la fréquence de décharge du neurone. On observe un phénomène de *synaptic scaling* : l'efficacité de toutes les synapses excitatrices est lentement modifiée pour compenser une augmentation ou une diminution d'activité. On est donc ici à la limite de la notion de plasticité intrinsèque, indépendante de l'activité présynaptique.

Pour expliquer le synaptic scaling, on peut citer la variation des *courants post-synaptique excitateurs miniatures* (mEPSC). Un mEPSC est la réaction d'un bouton synaptique suite à l'émission spontanée d'une vésicule de neurotransmetteurs. On observe qu'un changement de la

fréquence moyenne de décharge du neurone post-synaptique entraîne une modification proportionnelle de l'amplitude des mEPSC. Cette variation d'amplitude pourrait être causée par une modification du nombre de récepteurs et/ou de leur conductance. Le niveau d'activité moyen du neurone jouerait alors un rôle sur le mécanisme de renouvellement des récepteurs.

3.3 Plasticité structurelle

La plasticité structurelle est l'évolution de la structure du réseau neuronal. Cette évolution peut avoir lieu par le biais de l'évolution des terminaisons neuronales (axones et dendrites), la formation de nouveaux boutons synaptiques, ou encore la naissance de nouveaux neurones. Ici encore, les mécanismes mis en jeu sont nombreux, complexes et mal connus.

3.3.1 Évolution dendritique et synaptique

Lamprecht and LeDoux (2004) proposent une revue des mécanismes de plasticité structurelle associés aux épines dendritiques du neurone post-synaptique. Ces mécanismes sont généralement associés au renforcement et à la stabilisation de l'apprentissage synaptique.

1. On observe un ré-arrangement du cytosquelette au niveau des synapses, associé à l'augmentation de la concentration en ions Ca^{2+} . Cette concentration est, comme nous l'avons montré précédemment, une conséquence de l'activité synaptique. À terme, le ré-arrangement du cytosquelette mène à la formation de nouveaux boutons synaptiques, renforçant ainsi la connexion.
2. La majorité des synapses excitatrices se forment à partir des épines dendritiques. Or, la quantité d'épines dendritiques semble varier en fonction de l'activité du neurone. De même, la forme des épines dendritiques est variable et joue sur leur efficacité, ce qui pourrait être un autre mécanisme de modulation de la transmission synaptique.

3.3.2 Neurogénèse

La neurogénèse est un mécanisme permettant le renouvellement des cellules nerveuses. Reynolds and Weiss (1992) ont présenté les premières preuves que la neurogénèse existe chez l'adulte ; avant cela, on admettait généralement que les cellules nerveuses ne pouvaient pas se régénérer. Depuis, la neurogénèse chez l'adulte a été observée dans de nombreuses espèces, y compris chez l'homme (Eriksson et al., 1998). Chez la plupart des mammifères la neurogénèse chez l'adulte semble localisée au bulbe olfactif et à l'hippocampe (Aimone et al., 2007).

La neurogénèse est caractérisée par la génération de nouvelles cellules nerveuses au sein d'une population de cellules souches. Après leur naissance, les cellules nerveuses migrent pour s'intégrer à des structures neuronales pré-existantes. Cependant, un grand nombre de cellules ainsi créées meurent rapidement. À l'heure actuelle, la neurogénèse est un aspect de la plasticité cérébrale encore mal compris.

3.4 Conclusion

Nous avons vu que la plasticité cérébrale est un phénomène complexe, pouvant prendre de nombreuses formes. Les mécanismes biologiques impliqués sont encore aujourd'hui mal connus, mais certaines formes de plasticité ont fait l'objet de nombreuses recherches. La plasticité synaptique en particulier est la mieux comprise aujourd'hui, tant au niveau des mécanismes biologiques (comment varie l'efficacité synaptique) que des implications fonctionnelles (cadre théorique de l'apprentissage hebbien). Dans la suite de ce travail, nous allons donc nous concentrer sur la plasticité synaptique, qui a fait l'objet de nombreux travaux de modélisation et d'études théoriques en terme de traitement de l'information.

Chapitre 4

L'apprentissage hebbien

L'apprentissage hebbien (Hebb, 1949) est une notion fondamentale dans les travaux de modélisation de l'apprentissage. La grande majorité des modèles bio-inspirés de l'apprentissage dérivent de ce principe d'altération de l'efficacité synaptique en fonction des activités respectives du neurone présynaptique et du neurone postsynaptique.

Nous allons à présent nous intéresser à la formalisation de l'apprentissage hebbien dans les modèles informatiques des neurosciences computationnelles. Le domaine étant très vaste, il n'est pas envisageable de faire un tour exhaustif de toutes les règles d'apprentissage hebbiennes. Nous allons donc nous concentrer sur la règle fondamentale, les problèmes qu'elle soulève, et les variations qui en découlent.

4.1 Apprentissage hebbien classique

L'hypothèse de départ posée par Hebb (1949) est que l'activation concomitante d'un neurone présynaptique et d'un neurone post-synaptique entraîne un renforcement de l'efficacité des synapses reliant les deux neurones. La manière la plus simple de représenter l'apprentissage hebbien est de considérer l'équation suivante :

$$\frac{dw_{ij}}{dt} = \eta v_i \cdot v_j \quad (4.1)$$

où i et j sont deux cellules respectivement post- et présynaptiques, w_{ij} l'efficacité de la synapse reliant les deux cellules, v_i et v_j les fréquences de décharge respectives des deux cellules, η une constante représentant le taux d'apprentissage.

4.1.1 Le problème de la croissance infinie

Le mécanisme que nous venons de présenter a un énorme défaut : il ne possède aucun mécanisme de dépression synaptique. Chaque stimulation entraîne donc l'augmentation des poids synaptiques, sans aucune limite. Ce problème de croissance infinie des poids synaptiques dans l'apprentissage hebbien a fait l'objet de nombreux travaux, dont on trouvera une revue dans Abbott and Nelson (2000).

Les nombreuses solutions proposées au problème de la croissance infinie ont donné naissance à une grande variété de règles d'apprentissage dites hebbiennes. Gerstner and Kistler (2002) propose une classification des règles d'apprentissage hebbiennes sur laquelle nous nous basons pour les descriptions qui suivent.

4.1.2 Borne supérieure

Pour éviter la croissance infinie, une solution consiste à introduire une borne supérieure aux poids synaptiques.

$$\frac{dw_{ij}}{dt} = \eta v_i \cdot v_j \cdot (w_{max} - w_{ij}) \quad (4.2)$$

w_{max} est la borne supérieure, une constante strictement positive. Cette modification crée cependant un nouveau problème : à terme, tous les poids tendent vers la borne maximale, ce qui entraîne la perte de toute association apprise.

4.1.3 Terme de fuite

Une autre solution au problème de croissance infinie consiste à ajouter un terme de fuite proportionnel au poids synaptique.

$$\frac{dw_{ij}}{dt} = \eta v_i \cdot v_j \cdot (w_{max} - w_{ij}) - \iota \cdot w_{ij} \quad (4.3)$$

où ι est une constante positive proche de 0. En l'absence de stimulation, les poids synaptiques tendent vers 0. Quand le neurone est stimulé, l'apprentissage hebbien se déroule normalement, mais le terme de fuite gagne en importance. Pour qu'une connexion reste forte, elle doit donc participer régulièrement à l'activation du neurone post-synaptique.

4.1.4 Dépression synaptique

On constatera que le terme de fuite introduit la notion de dépression synaptique, absente de la formulation originale de Hebb. Outre son existence biologique avérée (chap. 3.1), la dépression synaptique (LTD) semble être nécessaire à la stabilisation de l'apprentissage hebbien. L'ajout d'un terme de fuite introduit la LTD de manière implicite, comme un phénomène permanent contrebalançant la LTP. Nous verrons que d'autres règles d'apprentissage modélisent la LTD de manière explicite, jouant un rôle dans l'apprentissage au même titre que la LTP.

4.1.5 La compétition entre les synapses

Le problème de la croissance infinie des poids peut être résolu en considérant que le renforcement d'une synapse se fait au détriment des autres. On peut, par exemple, normaliser la sommes des poids partageant le même neurone post-synaptique (Miller and MacKay, 1994). Cette approche permet une compétition *spatiale*, c'est à dire entre toutes les synapses d'un neurone à un instant donné.

Soit un neurone formel effectuant la somme de ses entrées v_i , pondérée par les poids correspondants w_i . En cas d'activation du neurone, la variation des poids synaptiques peut s'exprimer ainsi :

$$w_i(t+1) = \frac{w_i(t) + \eta v_i}{1 + \eta \sum_j v_j} \quad (4.4)$$

ηv_i est la variation du poids w_i selon la règle d'apprentissage 4.1, le terme post-synaptique valant 1 ; $\eta \sum_j v_j$ est la variation cumulée de tous les poids synaptiques. Le renforcement d'un poids synaptique entraîne une normalisation du vecteur de poids, qui se traduit par une dépression de tout le vecteur de poids.

La normalisation des synapses nécessite le partage d'informations non-locales (l'efficacité synaptique de chaque synapse), allant à l'encontre de la propriété de localité de l'apprentissage hebbien. On peut justifier un tel mécanisme en considérant que le renforcement des synapses se base sur une ressource limitée et partagée par toutes les synapses. Néanmoins, des apprentissages permettent une normalisation des poids sans passer par le calcul de termes non-locaux ; on citera par exemple la règle de Oja (1982), qui introduit un terme de dépression synaptique fonction de l'activité postsynaptique.

4.2 Apprentissage hebbien bidirectionnel

Les extensions à l'apprentissage hebbien visent à en enrichir le comportement : en plus de provoquer le renforcement synaptique (Long-Term Potentiation, LTP), la règle d'apprentissage doit aussi provoquer la dépression synaptique (Long-Term Depression, LTD) afin d'éviter la croissance continuelle des poids. On parle alors d'*apprentissage hebbien bidirectionnel*.

Les règles d'apprentissage hebbiennes se basent sur deux termes, représentant les activités pré- et post-synaptique. En se basant sur ce constat, la classification présentée dans (Gerstner and Kistler, 2002) distingue deux types d'apprentissage hebbien bidirectionnel, qualifiés de *presynaptic gating* et de *postsynaptic gating*.

4.2.1 Postsynaptic gating

Ces règles fonctionnent de telle sorte que le terme post-synaptique contrôle l'intensité de l'apprentissage, tandis que le terme présynaptique contrôle le signe de l'apprentissage (LTP/LTD). Le comportement général des règles de type *postsynaptic gating* est donc de renforcer les poids des synapses en concordance avec leur stimulation présynaptique. Ainsi, (Gerstner and Kistler, 2002) prend le cas de l'équation 4.5.

$$\frac{dw_{ij}}{dt} = v_i[v_j - w_{ij}] \quad (4.5)$$

Pour une stimulation v_j fixe, le point de stabilisation de cette règle d'apprentissage est $w_{ij} = v_j$. A stimulation présynaptique constante, une telle règle d'apprentissage converge vers un état dans lequel le vecteur des poids synaptiques est égal au vecteur d'entrées. Par extension, avec une entrée probabiliste, le vecteur de poids converge vers un prototype représentant la stimulation moyenne. Un tel mécanisme est largement utilisé dans les algorithmes d'apprentissage

non-supervisé tels que *Self-Organizing Map* (Kohonen, 1982), dans lequel les vecteurs de poids représentent des prototypes.

4.2.2 Presynaptic gating

Ces règles inversent le rôle des deux termes : le terme présynaptique contrôle l'intensité de l'apprentissage, tandis que le terme post-synaptique contrôle le signe de l'apprentissage. Le comportement général des règles de type *presynaptic gating* est d'effectuer une normalisation du vecteur des poids synaptiques.

$$\frac{dw_{ij}}{dt} = [v_\theta - v_i]v_j \quad (4.6)$$

L'équation 4.6 présente un exemple d'une telle règle : la constante $v_\theta > 0$ est appelée valeur cible. A stimulation fixe, l'apprentissage se stabilise lorsque $v_\theta = v_i$. Avec une activité post-synaptique $v_i = w_{ij}.v_j$, l'apprentissage se stabilise pour $w_{ij}.v_j = v_\theta$. Il s'agit donc d'une règle qui va opérer une normalisation du vecteur de poids indexée sur l'activité post-synaptique. Si l'exemple précédent rappelait l'apprentissage utilisé dans SOM, celui-ci rappelle la compétition spatiale introduite par la normalisation simple des poids présentée plus tôt.

4.3 Conclusion

L'apprentissage hebbien pose les bases de la modélisation de l'apprentissage dans le cerveau. Selon cette hypothèse, l'évolution de l'efficacité synaptique associée à l'activité des neurones pré- et postsynaptique est le mécanisme biologique fondamental par lequel l'apprentissage opère. La notion d'apprentissage hebbien a donné naissance à de nombreuses règles d'apprentissage dans le champ des neurosciences computationnelles, reproduisant divers types d'apprentissage (supervisé, par renforcement, non-supervisé).

En vue de concevoir un modèle d'apprentissage multimodal, nous allons maintenant nous intéresser au modèle BCM (Bienenstock et al., 1982), un modèle d'apprentissage hebbien bi-directionnel. Ce modèle permet d'acquérir de manière autonome la sélectivité à un stimulus particulier à partir du signal entrant.

Chapitre 5

La règle d'apprentissage Bienenstock-Cooper-Munro (BCM)

BCM (d'après le nom des auteurs, Bienenstock, Cooper et Munro) est un modèle d'apprentissage introduit par Bienenstock et al. (1982). Ce modèle possède certaines propriétés qui nous amènent à l'étudier plus en détail.

D'un point de vue biologique, BCM vise la modélisation corticale, à une échelle mésoscopique. A l'origine, le modèle a été développé pour rendre compte de phénomènes observés dans le cortex visuel primaire du chat et des primates, tels que la dominance oculaire, la réorganisation suite à une lésion, ou encore la sélectivité à l'orientation.

D'un point de vue modélisation, BCM se base sur un codage en fréquence représentant un niveau d'activité au sein d'une masse corticale. L'apprentissage mis en place est de type hebbien bidirectionnel, rendant compte de la LTP et de la LTD. Cet apprentissage à la particularité d'être totalement autonome : à partir d'une distribution d'entrée, l'unité acquiert la sélectivité à un stimulus particulier selon des critères statistiques dépendant uniquement des entrées.

5.1 Origine

BCM tire ses origines du modèle *Cooper-Liberman-Oja* (CLO), présenté dans Cooper et al. (1979). Les auteurs proposent un modèle de plasticité synaptique basé sur l'existence d'un seuil permettant de faire coexister les notions de renforcement (LTP) et de dépression (LTD) au sein d'un même mécanisme d'apprentissage. Ce modèle se positionne dans la lignée des travaux Nass and Cooper (1975), cherchant à expliquer le développement de la sélectivité à des stimuli structurés dans les neurones du cortex visuel.

$$\eta \dot{w} = \phi_{CLO}(c) \cdot d \quad (5.1)$$

Ce modèle est une variante de l'apprentissage hebbien. On notera quelques changements de notation : d et c sont respectivement les activités pré- et postsynaptiques, précédemment notées v_j et v_i . η est le taux d'apprentissage, $\dot{w}$ la variation du poids synaptique w en fonction du temps, $\phi_{CLO}(c)$ le terme post-synaptique, issu de la fonction ϕ_{CLO} .

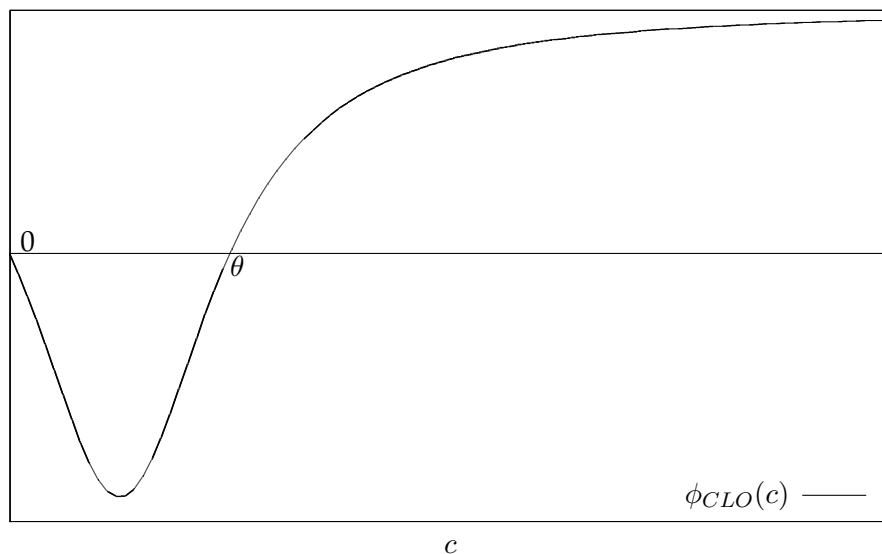


FIGURE 5.1 – La fonction ϕ_{CLO} détermine le signe du terme post-synaptique de l'apprentissage hebbien

La fonction ϕ_{CLO} introduit un seuil de modification θ (à ne pas confondre avec le seuil de décharge, souvent appelé θ dans la littérature) : si l'activité post-synaptique est supérieure à ce seuil, l'apprentissage est de type LTP ; si l'activité post-synaptique est inférieure à ce seuil, l'apprentissage est de type LTD. Il s'agit donc d'un apprentissage hebbien *bidirectionnel*, de type *presynaptic gating* selon la terminologie de Gerstner and Kistler (2002). La figure 5.1 présente la forme générale de la fonction ϕ_{CLO} .

Les auteurs constatent que ce modèle permet l'acquisition de la sélectivité à un stimulus particulier dans un environnement structuré, tandis qu'un environnement bruité entraîne l'absence de sélectivité. Ceci coïncide avec les observations de Buisseret and Imbert (1976), qui montrent une forte influence de l'environnement dans le développement des neurones du cortex visuel - en particulier le développement de la sélectivité à l'orientation observée par Hubel and Wiesel (1962).

Cependant, cette règle d'apprentissage n'est pas sans défauts.

- CLO ne s'intéresse qu'au cas d'une seule cellule, ignorant l'immense complexité du système visuel.
- Quoique plausibles, les stimuli utilisés pour ce modèle restent artificiels.
- La règle d'apprentissage manque de stabilité. Pour que la sélectivité se développe, le seuil θ doit être choisi précisément, de sorte qu'avec les poids synaptiques initiaux, un seul stimulus génère une activité supérieure à θ .
- La règle d'apprentissage ne résout pas toujours le problème classique de fuite des poids. Si le seuil est trop bas, la LTP prend le pas sur la LTD, poussant les poids vers l'infini. Inversement, si le seuil est trop haut, c'est la LTD qui prend le pas sur la LTP, poussant les poids vers 0.

La règle d'apprentissage *Bienenstock-Cooper-Munro* (BCM) introduite dans Bienenstock et al. (1982) reprend le principe de seuil θ marquant la séparation LTP/LTD, mais se base sur un seuil *adaptatif*. Plutôt que d'être fixe, θ est fonction de l'historique d'activité du neurone post-synaptique. Cette adaptativité permet de résoudre les problèmes d'instabilité de CLO.

5.2 Fonctionnement

5.2.1 Activité

Le modèle BCM se base sur une représentation fréquentielle de l'activité neuronale, à une échelle mésoscopique. L'activité instantanée de l'unité est une fonction de la somme pondérée des entrées :

$$c_i = \sigma\left(\sum_{j=0}^n w_{ij} \cdot d_j\right) \quad (5.2)$$

où c_i est l'activité instantanée de l'unité i , σ une fonction d'activation croissante, w le vecteur des poids synaptiques et d le vecteur des activités pré-synaptiques. On notera que c_i représente un niveau global d'activité au sein d'une masse neuronale. Bien qu'on puisse rapprocher cette valeur d'une fréquence de décharge spatialisée, le modèle ne vise pas la correspondance précise avec de telles données biologiques.

L'usage d'une fonction d'activation σ en sigmoïde est introduite par Law and Cooper (1994) dans le but d'éviter les instabilités causées par un environnement complexe et changeant. Par exemple, un environnement visuel réaliste comprendra des variations brusques d'intensité lumineuse, qui pourraient causer des déséquilibres si l'activité n'est pas bornée. La fonction de sortie se traduit généralement par une fonction de saturation :

$$\sigma(c) = \begin{cases} \sigma_{inf} & \text{si } c < \sigma_{inf} \\ \sigma_{sup} & \text{si } c > \sigma_{sup} \\ c & \text{autrement} \end{cases} \quad (5.3)$$

Sauf mention contraire, les exemples que nous présenter pour illustrer les propriétés de stabilité de BCM considèrent que σ est la fonction identité.

5.2.2 Apprentissage

Tout comme dans CLO (Cooper et al., 1979), BCM est un apprentissage hebbien de type *pre-synaptic gating*, dont le terme post-synaptique est issu d'une fonction ϕ . La variation de poids s'exprime sous la forme :

$$\dot{w}_{ij} = d_j \phi(c_i, \theta_i) \eta \quad (5.4)$$

où η est une constante représentant le taux d'apprentissage, d_j l'activité présynaptique et $\phi(c_i, \theta_i)$ le terme post-synaptique de l'apprentissage hebbien.

La fonction ϕ détermine le signe de l'apprentissage en fonction de la position de c_i par rapport au seuil θ_i . Il existe de nombreuses formes mathématiques de la fonction ϕ . Le détail importe peu, tant que les propriétés suivantes sont vérifiées :

- $\phi(c, \theta)$ s'annule pour $c = 0$ et $c = \theta$,
- $\phi(c, \theta) < 0$ pour $c \in]0; \theta[$,
- $\phi(c, \theta) > 0$ pour tout autre valeur de c .

On retient souvent la simple fonction parabolique 5.5 qui vérifie les propriétés énoncées. Blais (1998) se base sur cette forme de la fonction ϕ pour étudier la stabilité de l'apprentissage BCM.

$$\phi(c, \theta) = c(c - \theta) \quad (5.5)$$

5.2.3 Seuil θ dynamique

Dans Bienenstock et al. (1982), les auteurs avancent l'hypothèse d'un seuil LTP/LTD dynamique, fonction de l'historique de l'activité de la cellule. Ainsi, une stimulation forte et prolongée entraîne la hausse du seuil θ , favorisant alors la dépression synaptique ; inversement, une stimulation faible prolongée entraîne la chute du seuil θ , favorisant alors le renforcement synaptique.

Cette hypothèse, qui apporte au modèle sa stabilité, est importante. En effet, elle sera validée plus tard par Kirkwood et al. (1996), qui mettra en évidence un phénomène de décalage du seuil LTP/LTD en fonction de l'historique d'activité du neurone. Des enregistrements dans le cortex visuel des rats privés de lumière indiquent une facilitation de la LTP suite à un stimulus. Le phénomène s'inverse en exposant les rats à la lumière. Ce constat joue donc en faveur de la plausibilité biologique de la règle BCM.

On notera que des travaux plus récents (Senn et al., 2001; Shouval et al., 2002; Izhikevich and Desai, 2003; Toyozumi et al., 2005) ont établis un lien entre la règle d'apprentissage STDP (*Spike Timing Dependent Plasticity*) et BCM, lorsque la première est transposée vers le codage en fréquence. La STDP permet un apprentissage hebbien bidirectionnel dans le cadre plus fin des neurones à spike. Ce lien vient donc encore renforcer la plausibilité biologique de BCM comme mécanisme d'apprentissage.

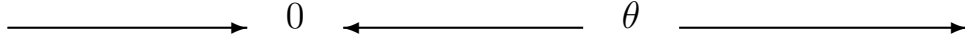
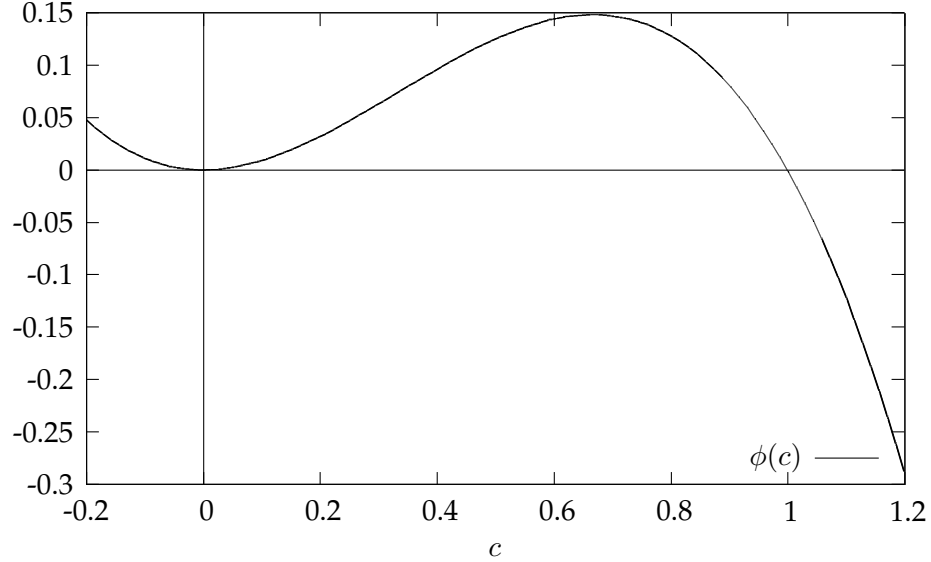
Pour étudier la stabilité de BCM, considérons un cas simplifié. Soit une unité recevant une entrée constante $d > 0$ par le biais d'une connexion de poids synaptique w . L'activité post-synaptique c est calculée selon l'équation 5.2 soit $c = w.d$ (σ = fonction identité). La variation du poids w est calculée d'après l'équation 5.4. Sur cette base, nous cherchons les points stables pour lesquels $\dot{w} = 0$. L'entrée étant non-nulle et constante, les points stables sont dépendant du terme post-synaptique.

Instabilité du seuil fixe

Soit une constante $\theta > 0$, comme dans Cooper et al. (1979). Trivialement,

$$\dot{w} = \phi(c, \theta) = 0 \iff \begin{cases} c = 0 \\ c = \theta \end{cases} \quad (5.6)$$

De plus, l'étude du signe de $\dot{w}$ (figure 5.2) nous apporte deux informations supplémentaires : (1) 0 est le seul point fixe *stable*. Si θ est un point fixe, la moindre déviation provoque une sortie du point fixe. (2) $c > \theta$ entraîne une fuite vers l'infini.

FIGURE 5.2 – Évolution des poids avec un seuil θ fixe, reproduit à partir de Cooper et al. (2004).FIGURE 5.3 – Profil de la fonction ϕ lorsque $\theta = c^2$. Pour une activité présynaptique strictement positive, le signe de $\phi(c)$ est aussi le signe de l'apprentissage, ce qui nous permet d'observer l'existence d'un point stable pour $c = 1$

Variation instantanée

Afin de résoudre ce problème d'instabilité, Bienenstock et al. (1982) proposent un seuil θ dynamique, fonction de l'activité post-synaptique c . Une relation supra-linéaire entre θ et c est proposée. Ainsi :

$$\theta = c^2 \quad (5.7)$$

$$\phi(c, \theta) = c^2 - c^3 \quad (5.8)$$

En considérant que l'activité présynaptique est toujours positive, le signe de la fonction ϕ détermine le signe de la variation du poids w (LTP/LTD) et par conséquent le signe de la variation de l'activité c . La fonction ϕ admet deux points fixes pour lesquels $\phi(c) = 0$, et donc $\dot{w} = 0$: $c = 0$ et $c = 1$ comme l'illustre la figure 5.3. Cependant, $c = 0$ est instable car la moindre déviation autour de $c = 0$ entraîne une augmentation de w éloignant de $c = 0$. $c = 1$ est donc l'unique point stable de l'apprentissage :

$$\begin{cases} 0 < c < 1 & \iff c^2 < c & \iff \dot{w} > 0 \\ c > 1 & \iff c^2 > c & \iff \dot{w} < 0 \end{cases} \quad (5.9)$$

Quel que soit le poids de départ, on converge toujours vers une valeur de w telle que $c = 1$. C'est cette relation supra-linéaire entre c et θ qui permet de stabiliser le système.

Variation progressive

L'introduction de la moyenne temporelle de l'activité dans le calcul de θ permet une évolution lente du seuil LTP/LTD, beaucoup plus plausible biologiquement.

$$\theta = E_{\tau}(c^2) \quad (5.10)$$

Le seuil θ est défini comme l'espérance de l'activité au carré sur une fenêtre de temps τ . On notera que cette formulation provient de Intrator (1990); elle vient remplacer la formulation originale de Bienenstock et al. (1982) qui était $\theta = E(c)^2$. Cette dernière est problématique dans le cas où $E(c) = 0$. Cette variation de la règle BCM originale est nommée *Quadratic BCM* (QBCM). La majorité des travaux qui ont suivi se basant sur cette formulation de BCM, on retiendra celle-ci.

Pour les besoins des analyses qui suivent, considérons que la fonction σ est linéaire et que nous connaissons à l'avance l'ensemble des stimuli utilisés pour l'apprentissage ainsi que leurs probabilités respectives. A tout instant t , on peut alors calculer le seuil θ à partir du profil de poids $\mathbf{w}(t)$:

$$\theta(t) = E_{\tau}(c^2) = \sum_i p_i [\mathbf{w}(\mathbf{t}) \cdot \mathbf{d}_i]^2 \quad (5.11)$$

où d_i est un stimulus, p_i sa probabilité associée, et $\mathbf{w}(t)$ est le vecteur de poids à l'instant t .

Cependant, dans le cas pratique, ces informations ne sont pas disponibles; le seuil θ est alors calculé à l'aide d'une convolution exponentielle :

$$\theta(t) = \frac{1}{\tau} \int_{-\infty}^t c^2(t') e^{-\frac{t-t'}{\tau}} dt' \quad (5.12)$$

où $\theta(t)$ est la valeur de θ à l'instant, $\tau > 0$ une constante d'intégration. τ détermine la vitesse de décroissance des poids utilisés dans la convolution, et donc l'importance relative de l'activité récente par rapport à l'activité ancienne. τ représente donc bien une fenêtre temporelle : plus τ est grand, plus θ évolue lentement. Un effet similaire serait obtenu en calculant la moyenne de c^2 au sein d'une fenêtre glissante de taille τ :

$$\theta(t) = \frac{1}{\tau} \int_{t-\tau}^t c^2(t') dt' \quad (5.13)$$

Mais l'intérêt majeur de la convolution exponentielle est qu'elle peut se traduire en une équation différentielle simple :

$$\dot{\theta} = \frac{1}{\tau} (c^2 - \theta) \quad (5.14)$$

Le calcul de θ peut alors se faire de manière itérative, sans avoir besoin de garder en mémoire l'activité de c durant une fenêtre de temps. De plus, ce type de calcul est bien connu puisqu'il correspond à un filtre passe-bas - le paramètre τ contrôlant la fréquence de coupure du filtre.

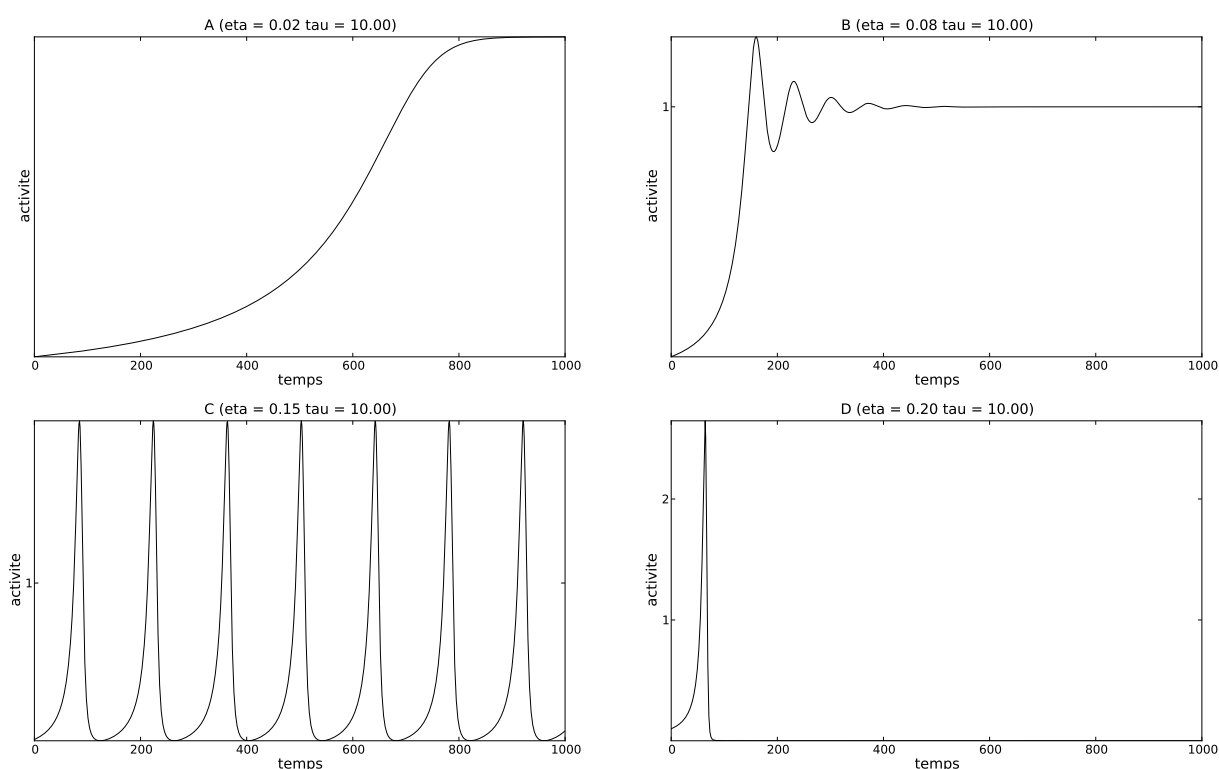


FIGURE 5.4 – Régimes d'activité en fonction des paramètres η et τ , pour une stimulation fixe strictement positive. La figure A représente le régime d'activité A : le système converge directement vers le point stable car l'apprentissage est suffisamment lent par rapport à l'évolution du seuil θ ; la figure B représente le régime d'activité B : l'apprentissage n'est pas suffisamment lent par rapport à l'évolution du seuil θ , entraînant une oscillation de l'activité - cependant la vitesse d'apprentissage reste dans les limites définies par $\alpha < 1$, et les oscillations s'atténuent progressivement jusqu'à atteindre la stabilité ; la figure C représente le régime C : pour la même raison que le régime B, l'apprentissage entraîne une oscillation de l'activité mais celle-ci n'est pas atténuée, maintenant le système dans un état instable ; la figure D représente une autre possibilité du régime C, ou le système se stabilise à une valeur inexploitable.

Stabilité et moyenne temporelle

L'usage d'une moyenne temporelle dans le calcul du seuil θ a pour effet de ralentir la variation de celui-ci, causant un délai entre les variations de c et θ . La variation du seuil θ étant garante de la convergence de l'apprentissage, ce ralentissement modifie la dynamique de convergence.

Les vitesses de variation respectives de c et θ déterminent le régime du système. Trois paramètres influencent ces vitesses de variation. (1) la constante d'intégration τ ; plus τ est grand, plus la variation de θ est lente. (2) le taux d'apprentissage η ; plus η est petit, plus la variation de w est lente, et par conséquent la variation de c . (3) l'activité pré-synaptique d ; plus l'activité pré-synaptique est faible, plus la variation de w est lente.

La figure 5.4 montre les différents régimes d'activité observés en fonction des paramètres τ et η . Blais (1998) présente une analyse détaillée des différents régimes possible dans le cas de la stimulation constante présenté précédemment. L'analyse mathématique mène à la conclusion qu'il existe trois régimes distincts liés à la valeur du rapport $\alpha = \eta\tau d^2$:

Régime A pour $6\alpha - 1 - \alpha^2 \leq 0$: le système converge directement vers le point stable, sans oscillations,

Régime B pour $\alpha < 1$ et $6\alpha - 1 - \alpha^2 > 0$: le système converge vers le point stable par oscillations atténuées,

Régime C pour $\alpha > 1$: le système est instable, les oscillations ne s'atténuent plus. Les outils mathématiques utilisés pour l'analyse ne sont pas adaptés.

On en conclut que si la condition $\eta\tau d^2 < 1$ est vérifiée, l'apprentissage converge vers le point de stabilité, directement ou par oscillations atténuées selon le rapport.

5.3 Analyse

Pour comprendre une règle d'apprentissage, il n'est pas suffisant d'en comprendre les mécanismes. Cette section analyse le comportement de BCM, afin de mieux appréhender ce qui est appris en terme d'information, et d'étudier plus généralement les propriétés de stabilité de cette règle d'apprentissage.

5.3.1 Stabilité dans le cas unidimensionnel

Précédemment, nous avons utilisé le cas unidimensionnel (une entrée, un poids synaptique) pour montrer qu'avec une stimulation constante et des paramètres respectant certains critères, l'apprentissage BCM entraîne une stabilisation des poids de sorte que $c = 1$. Restons dans le cas unidimensionnel, mais remplaçons la stimulation constante par une stimulation probabiliste pouvant prendre deux valeurs :

$$d = \begin{cases} P(0) = 1 - p \\ P(s) = p \end{cases} \quad (5.15)$$

$P(i)$ indique la probabilité d'apparition du stimulus i . Ce qui implique les valeurs suivantes pour c :

$$c = \begin{cases} P(0) = 1 - p \\ P(ws) = p \end{cases} \quad (5.16)$$

On peut alors simplifier l'écriture de θ selon la définition donnée par l'équation 5.11 :

$$\theta = E(c^2) = p(ws)^2 \quad (5.17)$$

On cherche les points fixes de l'équation d'apprentissage 5.4, c'est à dire les valeurs de w telles que $\dot{w} = 0$.

$$\dot{w} = \eta dc(c - \theta) \quad (5.18)$$

Si $d = 0$, $\dot{w} = 0$. Sinon, c'est le terme post-synaptique qui est déterminant. On s'intéresse donc au terme post-synaptique lorsque $d = s$.

$$\phi(c, \theta) = 0 \iff ws(ws - p(ws)^2) = 0 \quad (5.19)$$

$$\iff \begin{cases} w = 0 \\ w = \frac{1}{ps} \end{cases} \quad (5.20)$$

$$\iff \begin{cases} c = 0 \\ c = \frac{1}{p} \end{cases} \quad (5.21)$$

$$(5.22)$$

On retrouve donc deux points fixes, et l'étude du signe de $\phi(c, \theta)$ nous indique que $w = \frac{1}{ps}$ est le seul point stable. Dans le cas d'une telle stimulation, on converge donc vers un poids w inversement proportionnel à la probabilité d'apparition du stimulus, tel que c oscille entre 0 et $\frac{1}{p}$. Connaissant les probabilités des stimuli, le calcul de l'espérance de l'activité est trivial :

$$E(c) = p\frac{1}{p} + (1-p)0 = 1 \quad (5.23)$$

En conclusion, l'apprentissage mène à un poids w qui normalise l'activité moyenne de l'unité. Dans le cas présent, cette normalisation s'accompagne d'une oscillation de l'activité inversement proportionnelle à la probabilité d'apparition du stimulus non-nul.

Incidence de la fonction d'activité

Nous venons de voir que BCM se stabilise lorsque $c = \theta = \frac{1}{p}$. Si la fonction d'activité est choisie telle que $\sigma_{sup} < \frac{1}{p}$, on peut en conclure que l'égalité $c = \theta = \frac{1}{p}$ ne sera jamais accessible. Par conséquent, le choix de la borne supérieure peut empêcher la convergence vers certains points stables.

5.3.2 Stabilité dans le cas multidimensionnel

Le cas unidimensionnel permet d'étudier la stabilité, mais pas la sélectivité de BCM : en effet, il n'y a qu'un seul stimulus auquel l'unité peut réagir. Cooper et al. (2004) étend donc l'analyse de la convergence de BCM au cas bidimensionnel, puis à N dimensions.

Dans BCM, la compétition est temporelle, entre les stimuli. Les stimuli entraînant une activation suffisante (ie. $c > \theta$) voient leurs synapses actives renforcées, tandis que les stimuli entraînant une activation insuffisante (ie. $c < \theta$) voient leurs synapses actives affaiblies. Ce mécanisme renforce donc la discrimination entre les stimuli "forts" et les stimuli "faibles", amenant l'unité à être fortement active pour un sous-ensemble de stimuli, et faiblement au reste.

Pour illustrer ce mécanisme, considérons le cas bidimensionnel. Soit une unité possédant deux entrées exprimées par le vecteur $\mathbf{w}$ et recevant deux stimuli d^1 et d^2 linéairement indépendants. Les points fixes d'un tel système sont les vecteurs $\mathbf{w}$ tels que $\dot{\mathbf{w}} = 0$ pour d^1 et d^2 . Les points stables de ce système sont les vecteurs $\mathbf{w}$ tels que l'ajout d'un léger bruit sur $\mathbf{w}$ n'entraîne pas la divergence vers un autre point fixe.

Dans ce cas, BCM admet deux points fixes w^1 et w^2 , tous deux orthogonaux à un stimulus (figure 5.5). On peut étudier la stabilité de ces points fixes en calculant l'apprentissage moyen

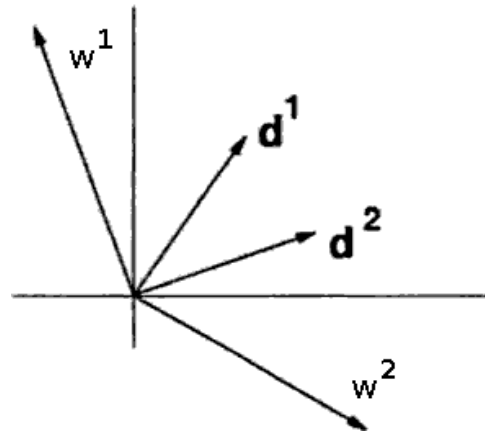


FIGURE 5.5 – Solutions stables m^1 et m^2 pour une entrée bidimensionnelle composée de deux stimuli d^1 et d^2 . Les axes représentent à la fois l'échelle des entrées et des poids. Figure adaptée de Cooper et al. (2004).

causé par les stimuli en fonction du vecteur de poids. Cette analyse est illustrée par la figure 5.6. Plus généralement, pour N stimuli linéairement indépendants dans un espace à N dimensions, il existe N points stables, chacun étant orthogonal à $N - 1$ stimuli.

5.3.3 Conditions d'existence des points stables

Les études de stabilité présentées précédemment se basent sur un certain nombre d'hypothèses que nous allons expliciter.

Changement de signe des poids synaptiques

Nous avons vu que les points fixes correspondent à des vecteurs de poids orthogonaux à tous les stimuli sauf un. Soit deux stimuli équiprobables :

$$\begin{aligned} \mathbf{d}_1 &= (1, 0.5) \\ \mathbf{d}_2 &= (0.5, 1) \end{aligned}$$

Dans ce cas, les points stables obtenus par l'apprentissage BCM sont :

$$\begin{aligned} \mathbf{w}_1 &= \left(-\frac{4}{3}, \frac{8}{3}\right) \\ \mathbf{w}_2 &= \left(\frac{8}{3}, -\frac{4}{3}\right) \end{aligned}$$

Ceci illustre une propriété importante de BCM : *certain points stables ne sont atteignables que si les poids synaptiques peuvent changer de signe*. Habituellement, le signe d'un poids synaptique

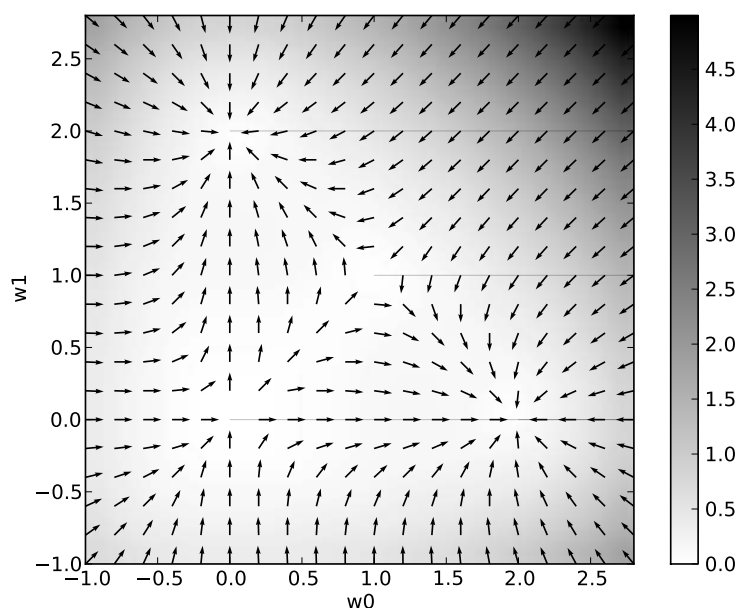


FIGURE 5.6 – Illustration du cas bidimensionnel. Ce graphe représente l'espace des vecteurs poids $\mathbf{w} = [w_0, w_1]$. Soumis à deux stimuli équiprobables $d^1 = [1, 0]$ et $d^2 = [0, 1]$, on peut calculer pour tout vecteur poids $\mathbf{w}$ l'apprentissage moyen $\dot{\mathbf{w}}$, représenté ici par les vecteurs. Pour plus de lisibilité, les vecteurs sont normalisés - le dégradé représente $|\dot{\mathbf{w}}|$, la norme de l'apprentissage moyen. On observe la présence de quatre points fixes, dont deux points stables $m^1 = [0, 2]$ et $m^2 = [2, 0]$ situés aux coeurs de bassins d'attractions. Les deux autres points fixes sont $\mathbf{w} = [0, 0]$ et $\mathbf{w} = [1, 1]$, ceux-ci sont instables.

permet d'exprimer la nature de la connexion : excitatrice (positif) ou inhibitrice (négatif). Or, la neurobiologie nous montre que la nature d'une synapse ne change pas avec l'apprentissage. La justification avancée par Bienenstock et al. (1982) est de dire que BCM se positionne à une échelle supérieure au neurone. Ainsi, une connexion dans le modèle BCM représente un ensemble de synapses excitatrices et inhibitrices entre deux masses neuronales. L'apprentissage peut alors faire varier la balance entre l'excitation et l'inhibition.

Stabilité du seuil θ

Nous avons montré précédemment que les points stables sont associés à une activité oscillant entre $c = 0$ pour les stimuli non sélectifs, et $c = \frac{1}{p}$ pour le stimulus sélectif de probabilité p . La démonstration de la stabilité se base sur l'équation 5.11 pour calculer le seuil θ comme une espérance "instantanée". Ce calcul n'est possible que si l'on connaît l'ensemble des stimuli présentés à l'unité, et leurs probabilités respectives. Dans le cas pratique où ces informations ne sont pas connues, θ est calculé itérativement selon la méthode itérative décrite par l'équation 5.14.

Pour toute stimulation variable, le seuil θ n'est donc jamais totalement fixe, ce qui implique que l'apprentissage ne s'arrête jamais. Au lieu de cela, les poids varient autour des points stables, et l'amplitude des variations est corrélée avec l'instabilité du seuil θ .

Nous avons dit précédemment que la méthode itérative pour calculer θ correspond à un filtre passe-bas. Cela signifie que les variations de c se répercutent en des variations de θ d'autant plus affaiblies que τ est grand. De même, plus l'amplitude des variations de c est grande, plus θ est instable. Ce phénomène est fortement accentué par le fait que θ se base sur c^2 .

A mesure que l'unité devient sélective, l'activité c oscille de plus en plus, jusqu'à atteindre les limites $c = \frac{1}{p}$ et $c = 0$ comme nous l'avons démontré dans la section 5.3.1. La figure 5.7 montre cet affinement de la sélectivité, sa conséquence sur l'activité c et sur la stabilité du seuil θ .

Plus le seuil θ est instable, plus les poids varient autour des points stables. Cette variabilité peut remettre en question la viabilité d'un point stable - l'apprentissage peut amener l'unité à quitter le point stable pour converger vers un autre, alors que la distribution statistique des stimuli n'a pas changé. On peut donc conclure que tous les points stables ne sont pas forcément accessibles. L'accessibilité d'un point stable dépend du paramètre τ , mais aussi de la probabilité d'apparition du stimulus correspondant. Moins un stimulus est probable, plus le point stable associé est éloigné, et plus la convergence vers ce dernier cause la déstabilisation du seuil θ .

Il s'agit là d'une limitation du modèle BCM : l'intensité de l'instabilité causée par l'apprentissage dépend d'un paramètre qu'on ne peut connaître *a priori* : la probabilité d'apparition du stimulus pour lequel l'unité est sélective. En conséquence, il devient nécessaire de choisir un paramètre τ arbitrairement grand pour prévenir l'instabilité, quitte à ce que cela soit inutile. Nous adresserons ce problème dans la suite de notre travail, dans le chapitre 8.

5.3.4 Étude statistique

La *malédiction de la dimensionalité* décrite par Bellman (1961) est un problème bien connu dans le domaine des neurosciences computationnelles. Soit une distribution de données dans un espace E ; la quantité de données nécessaire pour caractériser correctement cette distribution croît exponentiellement avec le nombre de dimensions de l'espace. Les données sont trop éparées,

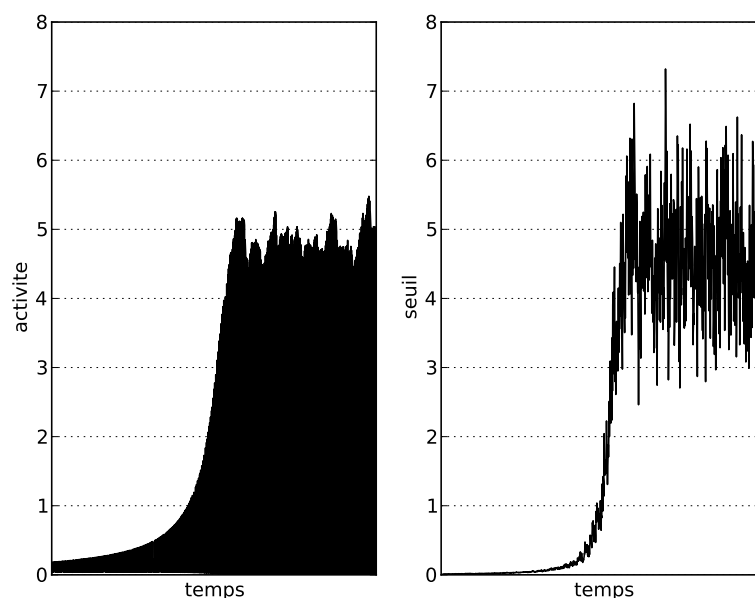


FIGURE 5.7 – Évolution de l’activité (gauche) et du seuil θ (droite) pour une unité soumise à cinq stimuli orthogonaux et équiprobables, sur 60000 itérations. Le point stable théorique se situe à $c = \theta = \frac{1}{p} = 5$.

et donc difficilement exploitables. Pour pouvoir exploiter les données d’un espace de grande dimension, on les projette dans un espace de plus faible dimension, sur lequel il est possible de travailler. Cependant, cette étape de projection altère la distribution, et entraîne donc une perte d’information plus ou moins grande. Dans ce cas, comment choisir une ou plusieurs projections de manière à atteindre un espace suffisamment petit pour être exploitable, sans pour autant perdre les spécificité de la distribution initiale ?

La *poursuite de projection* est une méthode de réduction de la dimensionalité recherchant des projections “intéressantes”. Cette méthode se base sur le constat fait par Diaconis and Freedman (1984) : lorsqu’on projette des données en grande dimension vers un espace de faible dimension, on obtient le plus souvent une distribution normale. Donc, l’information réside dans les projections qui divergent fortement de la distribution normale.

Un neurone artificiel possède un espace d’entrée de grande dimension (la taille du vecteur d’entrée) et le projette vers un espace de sortie de faible dimension (l’activité sortante du neurone). Les poids du neurone définissent la projection de l’espace d’entrée vers l’espace de sortie effectuée par le neurone. Les poids du neurone étant obtenus par l’apprentissage, on peut conclure que ce dernier détermine le traitement statistique (projection) effectué par le neurone sur les données de son espace d’entrée. Il convient donc d’étudier les règles d’apprentissage sous cet angle statistique.

Intrator and Cooper (1992) et Cooper et al. (2004) utilisent la poursuite de projection comme cadre d’analyse pour la règle d’apprentissage BCM, afin de qualifier l’apprentissage en termes statistiques, et permettre la comparaison avec d’autres règles d’apprentissage. L’analyse porte sur la sensibilité aux moments d’ordres 1 à 4 de la distribution d’entrée. L’analyse en composante principale est sensible aux moments d’ordre 1 (espérance) et d’ordre 2 (variance) de la distribution. En revanche, l’étude montre que la règle BCM permet de développer une sensibilité aux moments d’ordre 3 (asymétrie) et d’ordre 4 (aplatissement) de la distribution.

Les auteurs argumentent que l'analyse en composante principale permet une bonne préservation de l'information, car on minimise la perte d'information à la reconstruction du signal. Mais toute information contenue dans les moments d'ordre supérieur est perdue. Dans le cadre d'une tâche de catégorisation, ces moments peuvent être pertinents. BCM permet donc la mise en place de traitements statistiques sophistiquées tout en restant dans un cadre biologiquement plausible.

Un autre aspect de BCM mis en avant par les auteurs est sa fonction de réponse particulière. Un parallèle est fait avec les travaux de Barlow (1961) sur le codage optimal de l'information. Spécifiquement, Barlow avance qu'un neurone ne devrait pas se limiter à détecter l'occurrence d'un événement ; son signal de sortie devrait aussi transmettre une information sur la probabilité d'apparition de l'événement en question. Or, comme nous l'avons montré précédemment, un point stable de BCM est caractérisé par une réponse maximale valant $\frac{1}{p}$, p étant la probabilité d'apparition du stimulus sélectif.

5.3.5 Codage en population

BCM permet à une unité d'acquérir un filtre, qui entraîne l'activation de l'unité lorsqu'un événement spécifique est perçu dans le signal entrant. Cependant, l'information contenue dans le signal se réduit rarement à un seul événement. Le codage en population, que nous avons introduit dans la section 2.4.1 consiste à utiliser plusieurs unités partageant les mêmes stimuli afin d'acquérir un ensemble de filtres sur un signal donné. La diversité des filtres acquis augmente la quantité d'information extraite des stimuli.

Par exemple, si l'apprentissage correspond à une analyse en composante principale, faire en sorte que chaque unité extrait une composante différente augmente l'information obtenue. Un tel traitement nécessite une forme d'organisation au sein de la population qui découle d'un mécanisme permettant l'interaction entre les unités.

Or, BCM effectue un apprentissage totalement autonome pour chaque unité. Il faut donc mettre en place un mécanisme d'interaction au sein de la population, qui permette aux unités de la population de s'organiser pour améliorer le codage de l'information.

Interaction latérale dans BCM

Afin de permettre un codage en population, Cooper et al. (2004) étudient le cas d'une population d'unités possédant une connectivité latérale fixe. Ainsi :

$$c_i = \mathbf{w}_i \cdot \mathbf{d} + \sum_j L_{ij} c_j \quad (5.24)$$

Où c_i est l'activité de l'unité i de la population, w_i le vecteur de poids associé à l'unité i , $\mathbf{d}$ le stimulus entrant et L la matrice de connectivité latérale au sein de la population.

Les auteurs effectuent plusieurs analyses sur ce type de réseau. L'approximation de l'interaction latérale à l'aide d'un champ moyen montre que pour un tel réseau, les unités gardent les mêmes points stables, avec les mêmes profils de réponse, mais avec des vecteurs de poids différents. Une deuxième analyse plus générale est développée par Castellani et al. (1999). Cette nouvelle analyse vient confirmer les résultats de l'analyse par champ moyen, et ajoute que l'interaction latérale ne crée pas de points stables supplémentaires.

De plus l'analyse montre bien que des poids latéraux positifs amènent les unités à être sélectives au même stimulus, tandis que des poids négatifs amènent les unités à être sélectives à des stimulus différents. Néanmoins, l'interaction latérale n'empêche pas la population de converger vers une certaine combinaison de sélectivités. Une inhibition latérale ne rendra pas impossible la convergence vers un état où deux unités seraient sélectives au même stimulus. Au lieu de cela, l'interaction latérale *modifie la probabilité de converger vers un état stable*. L'inhibition latérale *favorise* donc la diversité des sélectivités.

Compétition par inhibition réciproque

Dans le chapitre 2, nous avons étudié l'interaction latérale au sein d'une carte d'unités, par le biais des travaux sur les champs neuronaux dynamiques (Amari, 1977). Nous avons vu que l'inhibition latérale permet la mise en place d'une compétition au sein de la carte, qui influence l'activité des unités qui la composent.

Or, lorsqu'un apprentissage hebbien a lieu, la modification de l'activité d'une unité implique la modification de son apprentissage. Pour BCM, une unité recevant une inhibition ou une excitation issue de la compétition latérale pourra voir son activité passer d'un côté du seuil θ à l'autre, et donc inverser l'apprentissage (passer de la LTP à la LTD, et réciproquement). Plus généralement, altérer l'activité de l'unité par le biais d'une source à part permet de favoriser la LTP ou la LTD à un instant donné.

Il s'agit là d'une idée importante pour la suite : *altérer l'activité de l'unité par le biais d'une source extérieure permet d'orienter l'apprentissage effectué par BCM*. Nous verrons par la suite qu'un signal excitateur coïncident avec un stimulus particulier augmente la probabilité que l'apprentissage a de converger vers la sélectivité à ce stimulus.

Auto-organisation

L'analyse de Castellani et al. (1999) se conclut par l'hypothèse que des profils de poids latéraux plus complexes pourraient permettre de mettre en place des phénomènes d'auto-organisation tel qu'on en observe dans le cortex cérébral. Nous étudierons cet aspect plus en détail dans la section 5.4.3.

5.4 Applications

Depuis sa création, le modèle BCM a été utilisé pour reproduire un certain nombre d'observations faites dans le vivant, et en particulier sur le cortex visuel primaire.

5.4.1 Binocularité et lésions

Les expériences de Hubel and Wiesel (1962) mettent en évidence la présence de bandes de dominance oculaire dans le cortex visuel primaire du chat : on constate que les neurones de l'aire V1 développent en général un biais en faveur d'un œil, quand bien même l'aire V1 reçoit l'information visuelle des deux yeux par le biais des projections thalamo-corticales. De plus, la déprivation sensorielle par suturation d'un œil entraîne une accentuation de la dominance en faveur de l'œil fonctionnel.

A l'aide de BCM, Bienenstock et al. (1982) reproduit ces phénomènes de dominance oculaire, et d'adaptation aux lésions. Une unité reçoit sa stimulation de deux vecteurs d'entrées représentant les deux rétines. On adapte le calcul de l'activité de BCM :

$$c = \mathbf{w}_l \cdot \mathbf{d}_l + \mathbf{w}_r \cdot \mathbf{d}_r \quad (5.25)$$

c l'activité de l'unité, $\mathbf{d}_l$ et $\mathbf{d}_r$ les vecteurs d'entrées correspondant aux deux rétines, $\mathbf{w}_l$ et $\mathbf{w}_r$ les vecteurs de poids associés aux deux entrées. Afin de reproduire les résultats observés par les biologistes, plusieurs cas de figures sont étudiés :

Stimuli corrélés Les deux rétines reçoivent des stimuli identiques ; l'unité développe la sélectivité sans la dominance oculaire. De plus, on constate que les deux rétines deviennent sélectives au même stimulus, malgré des poids initiaux différents.

Déprivation binoculaire Après une période d'apprentissage en conditions normales (stimuli corrélés), la suture des deux yeux est simulée en remplaçant les stimuli par du bruit sur les deux rétines. Cette déprivation binoculaire fait décroître progressivement la sélectivité, pour amener les poids dans un état instable. Une déprivation partielle (ajout d'un bruit corrélé sur les deux rétines) entraîne un affaiblissement de la sélectivité.

Déprivation monoculaire Après une période d'apprentissage en conditions normales, on simule la suture d'un oeil. Cette modification entraîne le renforcement de l'influence de la rétine stimulée, et un affaiblissement de l'influence de la rétine bruitée.

Suture inversée Après une première phase de déprivation monoculaire, on inverse les rôles : la rétine bruitée est à nouveau stimulée, et inversement. On observe que la dominance oculaire s'inverse progressivement, pour favoriser la rétine valide.

Stimuli décorrélés Le strabisme est simulé en présentant sur les deux rétines des stimuli tirés de la même distribution, mais totalement décorrélés. Dans la majorité des cas, l'unité développe une sélectivité monoculaire, supprimant l'influence de l'autre rétine. Dans de rares cas, l'unité développe la sélectivité sur les deux rétines, mais pas nécessairement au même stimulus.

BCM permet donc de reproduire certains phénomènes observés dans le vivant tels que le développement de la dominance oculaire et la réorganisation suite à une lésion - cette dernière se traduisant par un changement fort de la distribution d'entrées.

Discussion : cause de la dominance oculaire

Dans le modèle de la binocularité proposé ici, l'activité provenant des deux rétines est intégrée à l'activité de l'unité sans faire de distinction. La fusion des deux rétines donnerait un modèle strictement équivalent, où $\mathbf{d}_l$ et $\mathbf{d}_r$ forment un unique vecteur d'entrée, et $\mathbf{w}_l$ et $\mathbf{w}_r$ forment un unique vecteur de poids. Dans ce contexte, cette conception est justifiée, si on considère que la couche IV du cortex visuel primaire reçoit les connexions transmettant les signaux des deux rétines. On peut alors considérer que l'intégration des potentiels d'action dans l'activité, et par conséquent l'apprentissage ne font pas de distinction entre les deux sources.

Ce constat fait, étudions la compétition mise en place par BCM à cette échelle globale, sans faire une distinction entre les deux entrées. Spécifiquement, considérons chaque paire de stimuli

comme un seul stimulus. Soit S l'ensemble des n stimuli présentés par tirage uniforme à l'unité sur ses deux rétines.

Toute stimulation appartient à l'ensemble $S \times S$. Si les stimuli présentés sur les deux rétines sont corrélés, il y a n stimulations possibles, de probabilités $\frac{1}{n}$. En revanche, si les deux entrées sont décorrélées, le nombre de stimulations possibles s'élève à n^2 , de probabilités $\frac{1}{n^2}$.

Nous avons montré dans la section 5.3.3 que pour une constante d'intégration τ donnée, la convergence vers un point fixe entraîne une déstabilisation progressive du seuil θ . Si cette déstabilisation est trop forte, le point fixe perd sa stabilité, rendant impossible toute convergence de l'apprentissage vers cet état.

Or, ce phénomène est lié à la probabilité d'apparition du stimulus associé au point fixe concerné. Dans le cas des stimuli décorrélés, la probabilité d'apparition d'un stimulus tiré dans $S \times S$ est de $\frac{1}{n^2}$. Même avec un nombre relativement restreint de stimuli, la probabilité d'apparition devient très vite extrêmement faible. Les points fixes associés sont alors rendus inaccessibles par l'instabilité du seuil θ .

Dans ce contexte, l'apprentissage tend vers une sélectivité monoculaire, donnant naissance au phénomène de dominance oculaire. En effet, la probabilité d'apparition d'un stimulus tiré dans S est de $\frac{1}{n}$. A probabilité plus haute, point fixe plus proche, ce qui implique un seuil θ plus stable.

5.4.2 Apprentissage sur des images naturelles et sélectivité à l'orientation

Jusqu'ici, nous avons présenté des expériences se basant exclusivement sur des stimuli simplifiés. Law and Cooper (1994) propose d'appliquer BCM sur des stimuli visuels réalistes. Les stimuli sont des vignettes extraites d'un ensemble d'images naturelles. Dans un souci de simplicité, les pré-traitements ayant lieu avant l'arrivée au cortex visuel primaire sont résumés à une uniformisation de l'intensité du stimuli. Cette uniformisation est obtenue à l'aide d'un profil de poids en différence de gaussiennes, comme le montre la figure 5.8. Ce pré-traitement a pour effet de lisser l'image tout en préservant les contours.

Au niveau cortical, le modèle d'unité utilisé est la variante QBCM, présenté dans la section 5.2.

$$c = \sigma(\mathbf{d} \cdot \mathbf{w}) \quad (5.26)$$

σ est une fonction sigmoïde simplifiée, bornant l'activité dans l'intervalle $[0, M]$.

$$\sigma(c) = \begin{cases} 0 & \text{si } c < 0 \\ M & \text{si } c > M \\ c & \text{autrement} \end{cases} \quad (5.27)$$

La figure 5.9 représente les résultats obtenus en reproduisant l'expérience. La rétine est un cercle de 15 pixels de diamètres, Les vignettes sont tirées d'une image naturelle.

L'usage de la fonction de transfert σ est justifiée dans Cooper et al. (2004), qui reprend l'expérience :

A necessary condition for such receptive field to develop is the use of a non-symmetric sigmoidal function [...]. If a linear or symmetric sigmoidal function is used

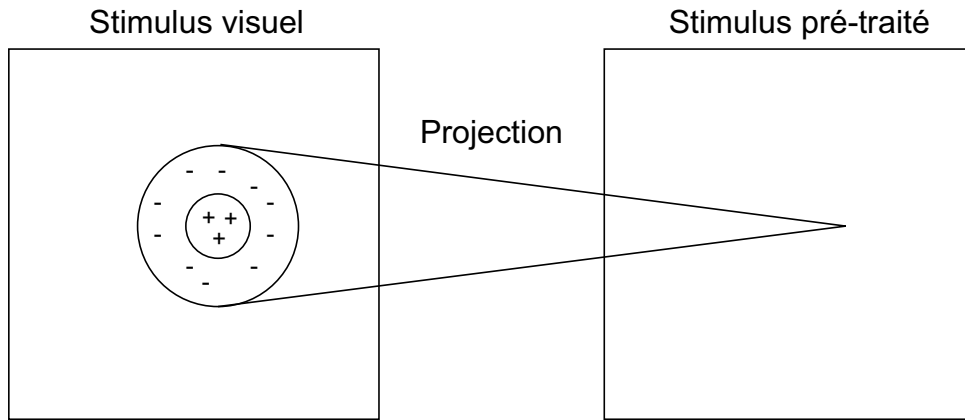


FIGURE 5.8 – Pré-traitement effectué sur les stimuli visuels. Le rond dans la carte de gauche correspond au champ récepteur associé à une unité de la carte de droite. Le profil de poids de cette projection est translaté pour chaque unité de la carte de droite. Ce type de traitement correspond à une convolution 2D entre l'image d'entrée et un noyau correspondant au profil de poids en différence de gaussiennes.

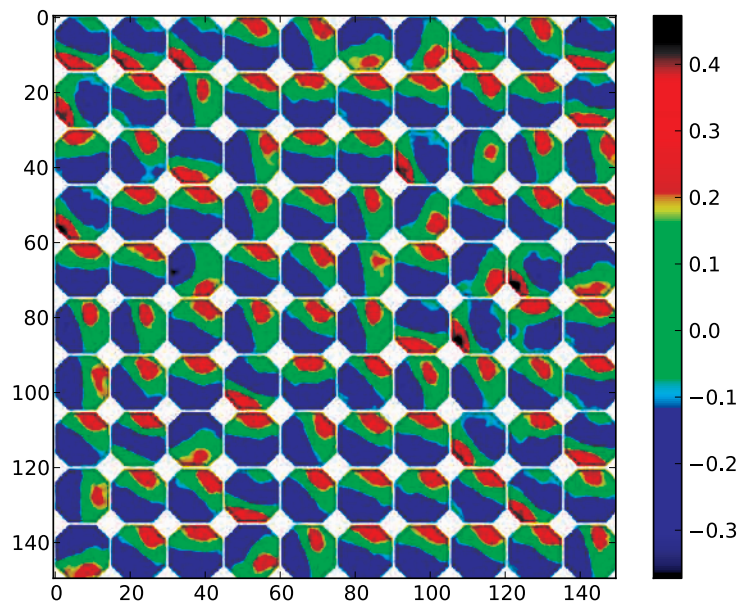


FIGURE 5.9 – Poids synaptiques des 100 unités constituant la carte corticale, après apprentissage sur un jeu d'images naturelles. Chaque vignette correspond aux poids d'une unité. $M = 20$, $\tau = 10^3$, $\eta = 10^{-6}$.

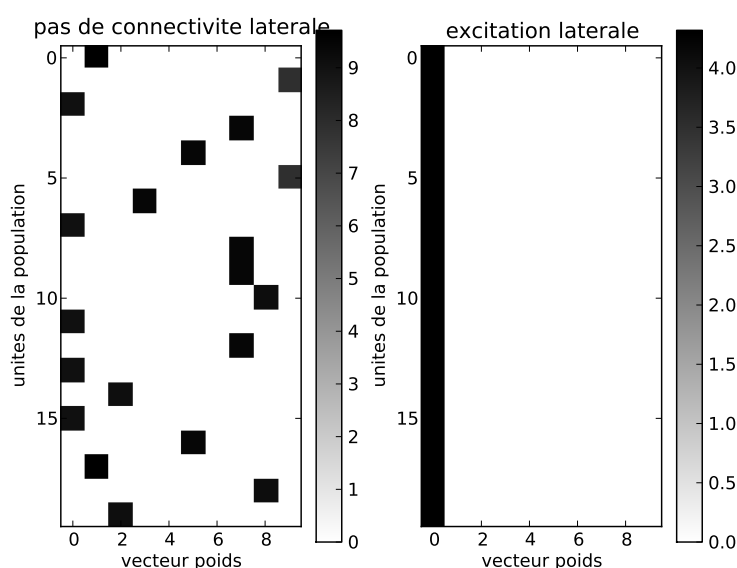


FIGURE 5.10 – Influence de la connectivité latérale excitatrice sur l'apprentissage de la population. Soit deux populations d'état et de poids initiaux identiques. Les deux populations sont soumises aux mêmes stimuli. A gauche, les 20 unités de la population sont indépendantes. A droite, on ajoute une connectivité latérale excitatrice uniforme. Dans les deux figures, chaque ligne représente le vecteur de poids d'une unité après apprentissage. On constate que l'excitation latérale entraîne la convergence de toutes les unités vers la même sélectivité.

then orientation selective receptive field do not develop. This is because the odd moments of these natural images are essentially zero, due to symmetry in the natural images. Since BCM depends on the third moment [...], if a linear neuron is used it would not develop selective RF's. The non-symmetric neuron breaks the symmetry in the images, resulting in oriented RF's.

5.4.3 Auto-organisation corticale

Nous avons montré précédemment (section 5.3.5) que la mise en place d'une inhibition latérale uniforme au sein d'une population permet une compétition entre les unités de la population. Cette compétition ne remet pas en cause l'existence des points stables des unités, mais favorise la convergence de la population vers un ensemble de sélectivités distinctes. Inversement, la mise en place d'une connectivité latérale excitatrice favorise la convergence de toute la population vers une même sélectivité. La figure 5.10 illustre la répartition des sélectivités avec (droite) et sans (gauche) connectivité latérale excitatrice.

Partant de ce constat, il semble logique qu'une connectivité latérale non uniforme permette une forme d'auto-organisation au sein de la population, à la manière des cartes auto-organisatrices de Kohonen (1982). Cette hypothèse est étudiée sous différentes formes dans plusieurs travaux.

Cooper et al. (2004) proposent une première expérience dans un environnement simplifié, où la population et l'entrée du réseau sont des vecteurs. Les stimuli présentés se basent sur une distribution gaussienne dont le centre est positionné aléatoirement sur le vecteur d'entrée. La sélectivité à une gaussienne est aisément reconnaissable : la position du poids le plus élevé correspond au centre du stimulus sélectif. L'activité est calculée en deux étapes. Dans un premier

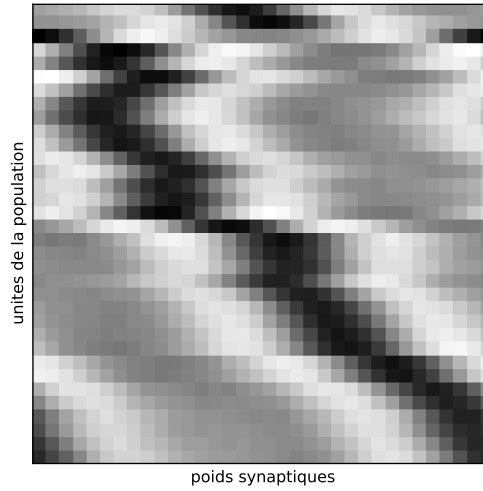


FIGURE 5.11 – Auto-organisation dans le cas 1d. Chaque ligne représente les poids synaptiques d’une unité de la carte. On constate que les unités proches au sens de la topologie de la carte ont des profils de poids similaires, ce qui indique la présence d’un phénomène d’auto-organisation lors de l’apprentissage.

temps, l’activité issue de la stimulation montante est calculée :

$$c_i^0 = \sum_j w_{ij} d_j \quad (5.28)$$

Puis, cette activité est propagée latéralement dans la population et vient s’ajouter à l’activité montante issue de la stimulation montante.

$$c_i = \sigma \left(c_i^0 + \sum_k L_{ik} \sigma(c_k^0) \right) \quad (5.29)$$

ou L est la matrice des poids latéraux, et σ est une fonction sigmoïde bornant l’activité.

Les poids latéraux de la matrice L sont ensuite déterminés en fonction de la distance entre les unités interconnectées. Le profil choisi est la différence de gaussiennes, ou “chapeau mexicain”. Ce profil particulier est couramment utilisé dans les travaux sur l’auto-organisation corticale (von der Malsburg, 1973). Les unités proches ont des connexions réciproques excitatrices, favorisant l’apprentissage de sélectivités similaires ; les unités éloignées ont des connexions réciproques inhibitrices, favorisant l’apprentissage de sélectivités distinctes.

La figure 5.11 illustre le résultat de cet apprentissage pour une carte unidimensionnelle de 34 unités - chaque ligne correspond au profil de poids d’une unité. On constate que les profils de poids de deux unités proches au sens de la topologie de la carte sont généralement similaires plutôt que d’être disposés aléatoirement comme dans la figure 5.10 (gauche). Un phénomène d’auto-organisation a donc bien eu lieu lors de l’apprentissage.

Le cas de l’auto-organisation dans une carte à deux dimensions est aussi étudié. Sur une tâche de sélectivité à l’orientation, on observe une organisation des sélectivités qui reflète les résultats de Mountcastle (1957), à savoir une variation globalement progressive de la sélectivité à l’orientation, avec des brisures nettes dans certaines régions (“pinwheels”).

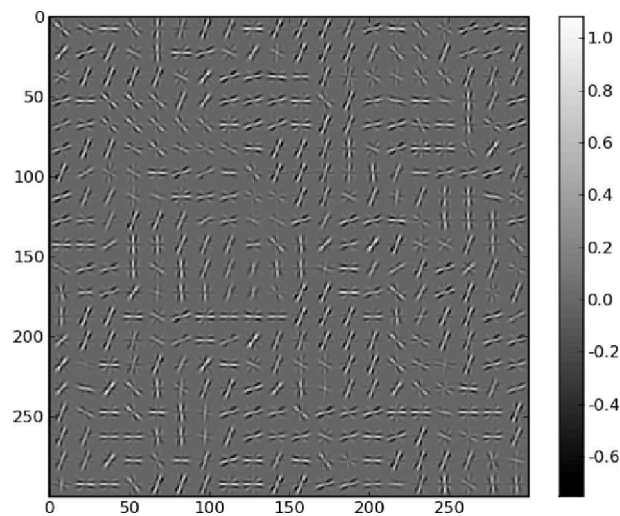


FIGURE 5.12 – Auto-organisation de la sélectivité à l’orientation dans une carte à deux dimensions.

Auto-organisation basée sur une connectivité latérale complexe

Shouval et al. (2000) étendent le cas de l’auto-organisation à une carte à deux dimensions. Plutôt que de considérer un profil de poids en différence de gaussienne pour la connectivité latérale, les auteurs proposent de reproduire une connectivité plus proche de celle observée dans le cortex visuel primaire. Spécifiquement, la connectivité latérale utilisée s’appuie sur la spécificité des connexions excitatrices. Les neurones qui partagent la même sélectivité à l’orientation tendent à être plus fortement connectés entre-eux ; c’est le biais d’*iso-orientation*. De plus, les connexions latérales d’un neurone se font principalement le long de l’axe correspondant à son orientation préférentielle. C’est le biais de *co-alignement*.

L’apprentissage sur la connectivité montante développe des sélectivités dont la répartition au sein de la carte satisfait les contraintes engendrées par la connectivité latérale. Il en résulte une organisation globale de la carte très proche de celle observée dans le cortex visuel primaire.

Discussion

La connectivité latérale au sein d’une population fonctionne comme une contrainte sur la convergence des unités. Cette contrainte dépend du profil de poids utilisé, et permet d’obtenir des phénomènes d’auto-organisation.

Il est important de se rappeler que la connectivité latérale est fixe. Elle vient moduler l’activité issue de la connectivité montante. Il est intéressant de constater que cette simple modulation de l’activité est suffisante pour orienter la convergence de l’unité vers un point fixe plutôt qu’un autre. Ce point sera développé dans la suite de notre travail.

5.5 Conclusion

BCM est un modèle de calcul et d’apprentissage simple parfaitement adapté au calcul distribué. Chaque unité effectue ses calculs de manière totalement autonome. La règle d’apprentis-

sage effectue un apprentissage non-supervisé permettant d'acquérir une sélectivité à certaines spécificités de la distribution du signal entrant.

De plus, l'ajout d'une connectivité latérale au sein d'une population d'unités BCM permet l'introduction de mécanismes de collaboration/compétition. Ces interactions entre les unités permettent la mise en place d'un codage en population et d'un phénomène d'auto-organisation au sein de la carte corticale.

Cependant, les expériences sur la binocularité nous montrent que la capacité de BCM à combiner plusieurs modalités semble limitée par la compétition entre les modalités entraînant un phénomène de dominance. Dans la suite de notre travail, nous étudierons plus en détail ce phénomène de dominance et proposerons une variante de BCM plus adaptée à l'apprentissage multimodal. Nous poserons ainsi les bases d'une unité de calcul à partir de laquelle nous pourrions faire de l'apprentissage non-supervisé multimodal basé sur un substrat distribué.

Enfin, nous avons vu que le choix du paramètre τ est d'une importance capitale pour la stabilité et la vitesse de l'apprentissage, mais qu'il n'est pas possible d'en déterminer la valeur *a priori* - trop petit, l'apprentissage est rapide mais instable - trop grand, l'apprentissage stable mais extrêmement lent. Nous proposerons donc un mécanisme pour contourner ce problème.

Troisième partie

Modèle et expérimentations

Chapitre 6

Jeux de données, protocoles de test et outils

L'activité de modélisation doit se concevoir avec des éléments qui peuvent paraître annexes mais sont cependant essentiels. D'une part, l'évaluation du modèle doit satisfaire certains critères (reproductibilité, fiabilité) et il importe de définir dans cette perspective les jeux de données et protocoles de tests qui seront utilisés pour les premières évaluations du modèle, avant la mise sur pied d'applications plus liées au comportement, permettant d'évaluer sa validité cognitive, comme il sera présenté au dernier chapitre.

D'autre part, l'évaluation nécessite des outils informatiques de simulation de qualité (fiabilité des calculs, généricité, simplicité d'usage, performances). L'utilisation de logiciels bien établis est généralement une garantie de qualité ; cependant, si les simulateurs de réseaux à spikes ne manquent pas, nous n'en avons pas trouvé de satisfaisant pour la modélisation mésoscopique et le codage en fréquence. En conséquence, nous avons développé la bibliothèque de simulation NADA, que nous présenterons en fin de chapitre.

6.1 Stimuli orthogonaux

Le jeu de données des stimuli orthogonaux consiste en N vecteurs de taille N orthogonaux et identiques à un décalage près. On peut construire ce jeu de données à l'aide de la matrice identité de taille N - chaque ligne est alors un stimulus du jeu de données.

Si le tirage des stimuli se fait selon une distribution uniforme ainsi que l'initialisation des poids, l'apprentissage BCM a une probabilité égale de devenir sélectif à n'importe lequel des N stimuli.

Les travaux de Blais (1998) montrent qu'il existe un point stable pour chaque stimulus lorsque les stimuli sont linéairement indépendants. C'est naturellement le cas des stimuli orthogonaux. De plus, il est aisé de caractériser la sélectivité à de tels stimuli. Représentons un ensemble de stimuli par une matrice où chaque ligne est un stimulus :

$$stim = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad (6.1)$$

Si la probabilité d'apparition des stimuli vaut p_1 , p_2 et p_3 pour les trois stimuli, alors les points stables sont définis par les vecteurs poids suivants :

$$poids = \begin{pmatrix} \frac{1}{p_1} & 0 & 0 \\ 0 & \frac{1}{p_2} & 0 \\ 0 & 0 & \frac{1}{p_3} \end{pmatrix} \quad (6.2)$$

où chaque ligne est le vecteur de poids associé à un stimulus. Avec de tels points stables, il est aisé d'analyser la convergence de l'apprentissage. Pour cela, nous proposons les calculs suivants :

Soit $\mathbf{w}$ le vecteur de poids obtenu après apprentissage. L'indice s du poids le plus élevé du vecteur de poids nous permet d'obtenir l'indice du stimulus auquel l'unité est la plus sélective.

$$s = \operatorname{argmax}(\mathbf{w}) \quad (6.3)$$

Nous avons dit qu'au point stable, le vecteur de poids est nul sauf pour une valeur valant $\frac{1}{p}$. En conséquence, on mesure la qualité de la sélectivité par le calcul suivant :

$$q = \frac{\max(\mathbf{w})}{\sum_i w_i} \quad (6.4)$$

Un autre aspect de la sélectivité est l'intensité avec laquelle l'unité répond au stimulus sélectif. Pour cela, il suffit de prendre le poids le plus élevé du vecteur de poids :

$$i = \max(\mathbf{w}) \quad (6.5)$$

Nous nous baserons par la suite sur ces trois mesures pour quantifier l'apprentissage dans les expériences utilisant le jeu de données des stimuli orthogonaux.

Naturellement, il s'agit là d'un jeu de données totalement artificiel, qui ne peut être considéré comme représentatif des stimulations auxquelles un organisme vivant est soumis. Mais l'intérêt d'un tel jeu de données est l'environnement contrôlé qu'il fournit, nous permettant d'étudier les propriétés statistiques et de convergence du modèle.

6.2 Posture / position

Pour tester l'apprentissage multimodal sur un exemple plus complexe, nous proposons de mettre en place un protocole expérimental dans lequel le modèle apprend sur deux modalités différentes mais corrélées, car issues du même phénomène.

Pour mettre en place un tel protocole, nous choisissons de nous inspirer d'une tâche de contrôle moteur. Soit un bras disposant de deux articulations, bougeant dans un espace bidimensionnel (figure 6.1). A partir de ce bras, on construit deux modalités : (1) la posture du bras, exprimée à l'aide des angles des deux articulations et (2) la position de la "main", exprimée dans le repère spatial bidimensionnel. L'intérêt de ce protocole expérimental est que la relation entre l'espace des postures et l'espace des positions n'est pas *bijection* : on peut associer plusieurs postures à une même position. De plus la relation entre les deux espaces est non-linéaire. Par exemple,

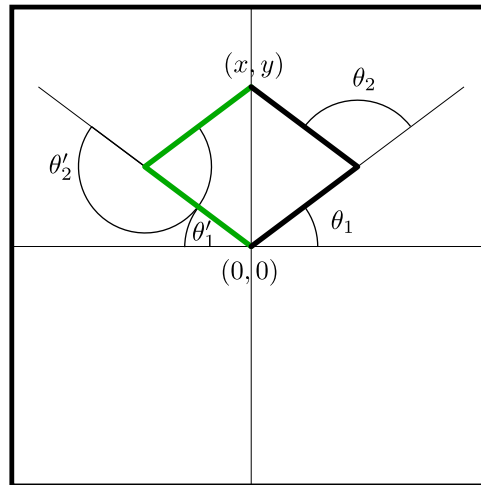


FIGURE 6.1 – Un bras possédant deux articulations, bougeant dans un espace bidimensionnel. A partir de la posture du bras représentée par le couple (θ_1, θ_2) , on calcule la position de la main (x, y) . Le figure montre qu’une position peut correspondre à plusieurs postures. Ici le bras vert représente une posture alternative donnant la même position.

une légère variation de posture pourra entraîner ou non une forte variation de position suivant la posture initiale du bras.

A partir des informations de posture et de position, on construit deux profils d’activité qui serviront de stimuli au modèle. Les entrées du modèle étant des cartes corticales, les stimuli prennent la forme d’une bulle d’activité dont la position dans la carte correspond à la posture (θ_1, θ_2) ou à la position (x, y) (voir figure 6.2).

Le protocole consiste à déterminer une posture en choisissant aléatoirement les deux angles dans une distribution uniforme sur l’intervalle $[0; 2\pi[$, puis à calculer les coordonnées x, y correspondant à l’extrémité du bras. A partir de ces données, on génère les bulles d’activité correspondantes dans les deux cartes. On dispose alors de deux modalités issues d’un même phénomène, mais dont la relation de corrélation n’est pas triviale.

Rappelons que notre but n’est pas de construire un modèle crédible de l’apprentissage du contrôle moteur - nous utilisons ce problème pour sa simplicité de mise en oeuvre et pour la relation entre les deux modalités.

6.3 Apprentissage en parallèle

Dans la majorité des simulations des chapitres suivants, plusieurs unités partagent un même champ récepteur et apprennent donc à partir des mêmes stimulations. En l’absence de connectivité latérale entre les unités, le seul paramètre différenciant les unités est leur vecteur de poids initial.

Les vecteurs de poids initiaux et l’ordre de présentation des stimuli étant déterminés aléatoirement selon une distribution uniforme, un protocole d’apprentissage tel que celui des stimuli orthogonaux mène à une répartition des sélectivités statistiquement uniforme entre les N stimuli. Sur cette base, il est possible d’étudier les effets des diverses modifications apportées au modèle sur les statistiques de convergence de la population.

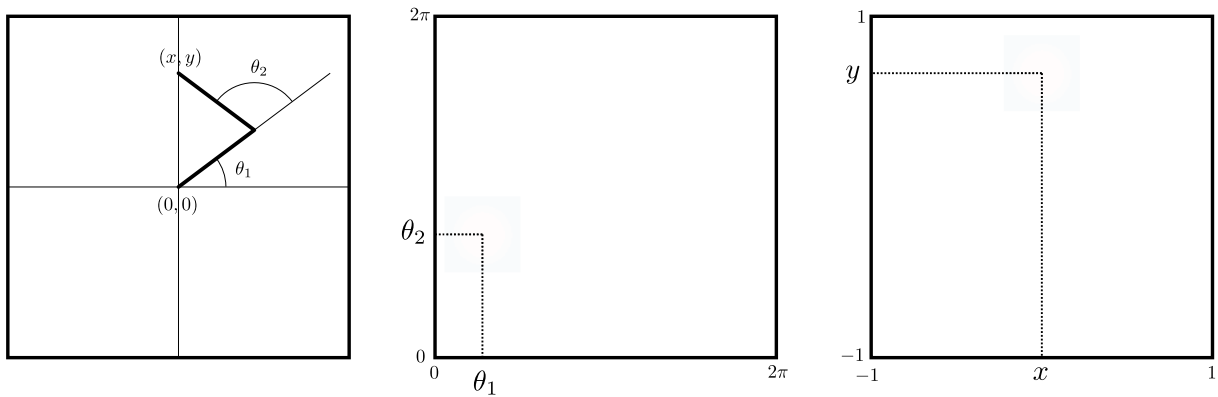


FIGURE 6.2 – Construction des stimuli de posture et de position à partir des couples (θ_1, θ_2) et (x, y) . A gauche, le bras avec ses informations de posture et de position. Au milieu, le stimulus représentant la posture - la bulle d'activité (en rouge) est centrée aux coordonnées (θ_1, θ_2) . A droite, le stimulus représentant la position - la bulle d'activité est centrée aux coordonnées (x, y) .

6.4 Outil de modélisation

Notre travail de modélisation et de simulation a naturellement nécessité un simulateur. Une recherche nous a amené au constat que la grande majorité des simulateurs dans le domaine des neurosciences computationnelles sont conçus pour simuler des réseaux de neurones à spikes, et par conséquent inadaptés à notre travail. Deux outils correspondaient à nos besoins.

Topographica (<http://topographica.org/Home/index.html>), dont l'usage n'a pas été possible pour des raisons techniques (nombreuses dépendances, conflits avec les packages installés sur le système).

DANA (<http://dana.loria.fr/>), que nous avons utilisé un temps, mais fini par abandonner pour cause d'instabilité du code (nombreuses réécriture du cœur du simulateur, modifications régulières de l'API).

Nous avons donc développé une bibliothèque adaptée à la modélisation mésoscopique et au prototypage rapide nommée NADA. Comme son nom l'indique, ce travail s'inspire des travaux de Nicolas Rougier sur le simulateur DANA, et a été développé avec la participation de Mathieu Lefort. Le code source de NADA est disponible à l'adresse suivante : <http://github.com/jiyunatori/nada>

6.4.1 Description générale

NADA est une bibliothèque développée en *Python*. Le modèle interprété du langage permet un développement rapide et incrémental des modèles, et des simulation interactives. La lenteur généralement associée à *Python* est compensée par un usage massif de la bibliothèque de calcul scientifique *Numpy*, qui offre une interface vers des structures de données et des routines de calcul fortement optimisées écrites en C et FORTRAN.

6.4.2 Architecture

Pour correspondre à une échelle de modélisation mésoscopique, NADA représente les réseaux de neurones comme un ensemble de cartes neuronales (classe `Group`) reliées entre elles par des projections de connexions (classe `Link`).

Cartes neuronales

La classe `Group` permet d’instancier des cartes neuronales. Une carte est définie par sa géométrie et le type des unités qui la compose. Par exemple :

```
g = Group(shape=(10,10), fields=['ext','c','t'])
```

Cette ligne de code définit une carte bidimensionnelle de 100 unités, disposées en une matrice de 10x10, au sein de laquelle chaque unité contient trois champs nommés *ext*, *c* et *t*. Les champs permettent de stocker les valeurs numériques flottantes décrivant l’état de l’unité.

Notre approche de la modélisation se concentrant sur les cartes neuronales plutôt que les unités isolées, les données sont stockées en associant un vecteur à chaque champ. En effet, on accède bien plus souvent aux valeurs d’un champ pour toute la carte qu’aux valeurs de tous les champs pour une unité.

Ainsi, la classe `Group` rend les champs accessibles en lecture et en écriture par le biais d’une syntaxe simple bénéficiant de toute l’expressivité de la classe `ndarray` de la librairie Numpy :

```
g['ext'] = 1 # écrit 1 dans le champ 'ext' sur toute la carte
g['c'] # retourne un numpy.ndarray 10x10 des valeurs du champ 'c'
g['t'][0,0] = 0 # écrit dans le champ 't' de l'unité (0,0)
g['t'][0] = 1 # écrit 1 dans le champ 't' de la première ligne
g['t'][:,0] = 1 # écrit 1 dans le champ 't' de la première colonne
```

Projections de connexions

La classe `Link` permet d’instancier des projections de connexions entre deux cartes neuronales. Une projection est définie par sa carte source et sa carte destination). Ainsi :

```
source = Group(shape=(10,10), fields=['ext','c','t'])
destination = Group(shape=(10,10), fields=['ext','c','t'])
projection = Link(source=(source,'c'),target=(destination,'ext'))
```

Il est important de noter ici que pour définir la projection, il n’est pas suffisant de donner une source et une destination. Le rôle d’une projection est de propager une activité depuis la source vers la destination. Ainsi, la source est accompagnée du nom d’un champ, dont le contenu sera propagé par la projection ; de plus, la destination est accompagnée du nom d’un champ dans lequel l’activité propagée sera stockée.

En interne, le profil de poids d’une projection est représenté par une matrice $M \times N$, avec M la taille de la carte destination et N la taille de la carte source.

6.4.3 Simulation

Jusqu'ici, nous avons montré comment définir un réseau. Nous allons maintenant voir les mécanismes de simulation permettant l'évolution de l'activité dans le réseau et l'évolution du réseau lui-même par l'apprentissage.

La simulation se déroule en itérations successives. Chaque itération est divisée en trois étapes : la *propagation*, l'*évaluation* et l'*apprentissage*.

Tous ces aspects sont paramétrables dans NADA, car basés sur des fonctions écrites par l'utilisateur.

Propagation

La *propagation* est l'étape qui consiste pour chaque projection à prendre l'activité de sa source, y appliquer son profil de poids et écrire le résultat dans sa destination.

Dans de nombreux modèles, la propagation est simulée par une somme pondérée des entrées, ce qui se traduit par un produit scalaire entre le vecteur d'entrées et le vecteur de poids de l'unité. On peut généraliser ce principe à des cartes et des projections à l'aide d'un produit matriciel entre la matrice représentant le profil de poids de la projection et le vecteur d'activité de la source. Ce calcul est effectué par défaut dans la classe `Link`. Cependant, il est possible de redéfinir la fonction de propagation en modifiant l'attribut `Link.propagateFunction`.

```
def weightedSum(link, source):  
    return np.dot(link['W'], source)  
  
projection.propagateFunction = weightedSum
```

En appelant la méthode `Link.propagate()`, on exécute la propagation pour la projection concernée. En appelant la méthode `Group.propagate()`, on exécute la propagation pour toutes les projections entrantes de la carte concernée.

Evaluation

L'*évaluation* est l'étape qui consiste pour chaque carte à faire évoluer son état interne à partir des stimulations propagées. Sur le même principe que pour la propagation, l'utilisateur peut définir ces calculs par le biais d'une fonction :

```
def posture_eval(state):  
    tau = 100  
    state['act'] = state['ext'] + state['int_c']  
    state['int_t'] += (state['int_c']**2 - state['int_t']) / tau  
    state['out'] = 1 / (1 + np.exp(-state['int_c'] + 1))  
  
posture.evalFunction = posture_eval
```

Ici, la fonction définie prend un paramètre qui est une référence à l'état interne de la carte. Toutes les opérations effectuées au sein de cette fonction viennent donc actualiser l'état interne de la carte.

En appelant la méthode `Group.eval()`, on exécute l'évaluation de la carte concernée.

Apprentissage

Pour paramétrer l'apprentissage, on définit une fonction prenant trois paramètres correspondant aux états interne de la projection, la source et la destination :

```
def apprentissage(link, source, target):
    eta = 0.001
    pre      = source['c']
    post     = target['c'] * (target['c'] - target['t'])
    link['W'] += pre * post * eta

projection.learnFunction = apprentissage
```

Cet exemple de code montre une implémentation possible de la règle d'apprentissage BCM. Ici encore, on peut accéder aux états en lecture comme en écriture, ce qui permet d'effectuer l'actualisation. Tel quel, il y a cependant un problème. Soit une source, une destination et une projection :

```
# x.shape retourne la geometrie de x
source.shape # (5,7)
destination.shape # (2,3)
projection.shape # (2,3,5,7)
```

En interne, Group stocke les données dans des vecteurs, et Link dans des matrices :

```
# x.state donne un acces direct aux donnees stockees
source.state.shape # (35,)
destination.state.shape # (6,)
projection.state.shape # (6,35)
```

Sous cette forme, toute opération entre la source et la destination sera un échec, car les dimensions ne correspondent pas. Pour éviter cela, on redimensionne source et target avant de les passer à la fonction d'apprentissage :

```
source.asSource.shape # (1,35)
destination.asTarget.shape # (6,1)
(source.asSource * destination.asTarget).shape # (6,35)
```

Ainsi, le redimensionnement de source et target permet une écriture simple des règles d'apprentissage - l'utilisateur n'a pas à se soucier de la géométrie de l'entrée et de la sortie.

En appelant la méthode `Link.learn()`, on exécute l'apprentissage pour la projection concernée. En appelant la méthode `Group.learn()`, on exécute l'apprentissage pour toutes les projections entrantes de la carte concernée.

6.4.4 Expérience

Après avoir défini la simulation comme nous l'avons indiqué, il ne reste plus à l'utilisateur qu'à écrire la portion de code lançant l'exécution.

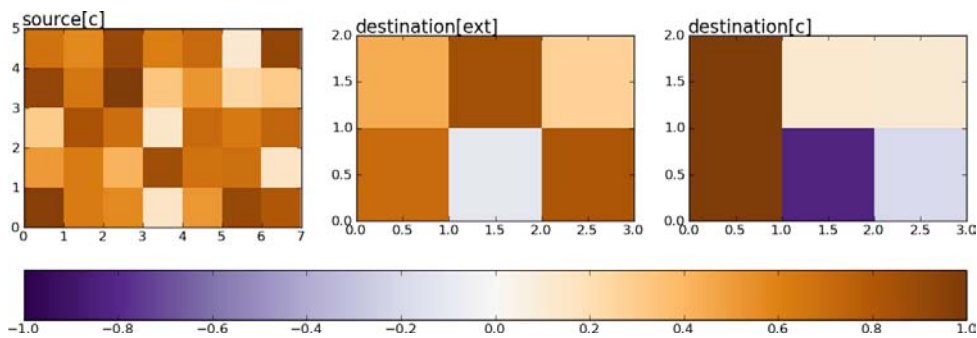


FIGURE 6.3 – Affichage des valeurs contenues dans trois champs appartenant à deux cartes : `source['c']`, `destination['ext']` et `destination['c']`. Chaque carré correspond à une unité de la carte.

```
epoch = 1000
for i in xrange(epoch):
    source['c'] = np.random.random(source.shape)
    destination.propagate()
    destination.eval()
    destination.learn()
```

6.4.5 Visualisation

De par l'usage de Python et Numpy, il est possible d'utiliser divers outils pour visualiser les modèles simulés ou construire des interfaces graphiques. En particulier, la bibliothèque *Matplotlib* offre de nombreuses possibilités dans ce domaine, telles que l'affichage de courbes et d'images à partir d'objets `numpy.ndarray`.

Sur cette base, NADA fournit quelques outils de visualisation rudimentaires permettant d'observer l'état du réseau interactivement, dont voici un exemple de code :

```
v = View([(source, 'c'), (destination, 'ext'), (destination, 'c')])
v.show()
```

Ici, on crée une vue du réseau qui affiche les champs `source['c']`, `destination['ext']` et `destination['c']` sur une même ligne. La liste de listes permet de contrôler la disposition des cartes sur l'affichage. Cet affichage permet deux choses :

- l'affichage du contenu d'un champ sous la forme d'une carte colorée, comme l'illustre la figure 6.3. Ici, chaque carré correspond à une unité de la carte ;
- l'affichage interactif des poids d'une projection comme l'illustre la figure 6.4. En cliquant sur une unité dans une carte, les cartes remplacent l'affichage des valeurs contenues dans les champs par les poids des éventuelles connexions allant vers l'unité cliquée.

Pour un usage simple, il est conseillé d'utiliser l'interpréteur *ipython*, en particulier avec l'option `ipython -pylab`. Cette dernière effectue automatiquement le chargement de la bibliothèque *matplotlib*, et gère l'affichage dans un fil d'exécution séparé. Ainsi, il est possible d'avoir la visualisation ouverte tout en gardant la main dans l'interpréteur pour faire tourner la simulation.

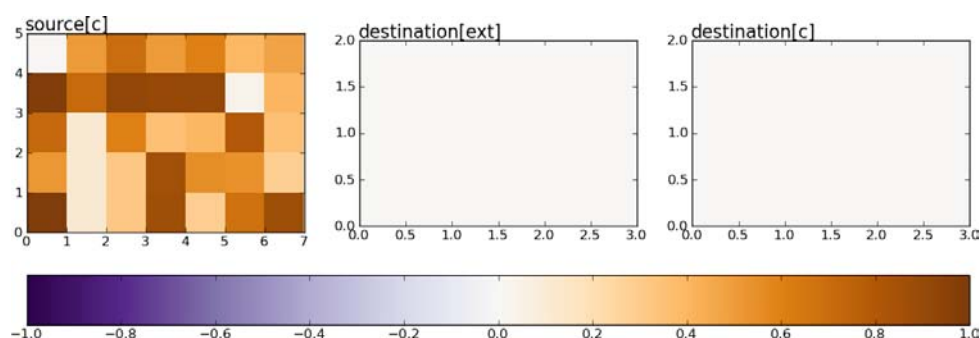


FIGURE 6.4 – Même affichage que la figure 6.3, mais on a cliqué sur le carré correspondant à l’unité (0,0) de la carte `destination`. Il en résulte que le contenu des cartes est remplacé par les poids des éventuelles connexions allant vers l’unité `destination[0,0]`. Le blanc indique une absence de connexion ou une connexion nulle. L’affichage indique qu’il existe une projection allant de `source` vers `destination` dont les poids sont affichés.

6.4.6 Avantages et inconvénients

La principale faiblesse de NADA est liée à un problème central de Python, à savoir le support médiocre du parallélisme, qui rend difficile l’exploitation des processeurs multi-coeurs aujourd’hui largement répandus. La parallélisation permettrait un gain de performance conséquent.

Les évolutions de la bibliothèque Numpy pourraient apporter ce gain de performance de manière transparente. En effet, Numpy se base sur les bibliothèques de calcul scientifique *BLAS* et *LAPACK*, dont certaines implémentations sont parallélisées. Nos essais sur ce point n’ont cependant pas été concluants, car les opérations les plus couramment utilisées par NADA ne semblent pas encore bénéficier de cette effort de parallélisation.

Dans le cadre de la simulation mésoscopique, les intérêts de Python sont cependant multiples.

Tout d’abord, la simulation s’exécute dans l’interpréteur. En conséquence, il est possible de l’arrêter à tout moment et de la relancer. Ceci permet par exemple d’étudier l’état du réseau en milieu de simulation, ou de modifier des paramètres.

De plus, Python dispose d’un très grand nombre de bibliothèques permettant d’interfacer les simulations avec des composants externes. Par exemple, la bibliothèque PIL (Python Image Library) permet de manipuler aisément des images et ainsi construire des simulations utilisant des stimuli visuels.

Enfin, les performances de Python ne sont certes pas aussi bonnes que celles d’un code natif, mais nos simulations se traduisant par des calculs matriciels, Numpy permet d’obtenir des performances acceptables. Ceci donne à NADA une grande simplicité d’utilisation tout en maintenant des performances acceptables.

Chapitre 7

Intégration du flux montant et du flux latéral

Dans la section 5.3.5 nous avons introduit les travaux sur l'usage de BCM dans le codage en population. A l'aide d'une connectivité latérale fixe entre plusieurs unités effectuant un apprentissage sur le flux montant, on met en place des dynamiques de compétition (inhibition) et de coopération (excitation).

Il est cependant important de noter que dans les simulations, cette intégration se traduit par un calcul en deux temps : d'abord on calcule l'activité montante, et ensuite on propage l'activité latéralement. D'un point de vue biologique, ce calcul semble un peu artificiel : il revient à réinitialiser l'activité à chaque pas de temps, et à ignorer la simultanéité du montant et du latéral.

Dans Girod and Alexandre (2008), nous avons proposé une modification de BCM visant à contourner cet aspect artificiel. L'approche développée consiste à calculer l'activité c itérativement, à partir des entrées montantes et latérales, tout comme le seuil θ est calculé itérativement à partir de l'activité.

7.1 Modèle

Soit une carte d'unités avec une connectivité latérale fixe L et une connectivité montante W . Toutes les unités partagent le même champ récepteur sur la projection montante. On calcule l'activité des unités comme suit :

$$\tau_c \frac{\partial c_i}{\partial t} = -c_i + \sigma \left(\sum_j W_{ij} d_j + \sum_k L_{ik} s_k \right) \quad (7.1)$$

avec $\tau_c > 1$ une constante d'intégration, $\sigma(x)$ une fonction sigmoïde bornant l'activité et s_k l'activité sortante de l'unité k :

$$s = \frac{1}{1 + e^{-c\lambda + \lambda}} \quad (7.2)$$

La sortie des unités est une fonction sigmoïde bornée entre 0 et 1, avec une pente λ . Nous avons vu que les points stables de BCM se situent pour des profils de poids faisant osciller l'activité de sorte que $\bar{c} \approx 1$. Pour cette raison, nous avons choisi une fonction de sortie centrée en 1 : pour les stimuli non-sélectifs, la sortie est ≈ 0 , et pour le ou les stimuli sélectifs, la sortie est ≈ 1 .

Le calcul du seuil θ ne change pas :

$$\tau_\theta \frac{\partial \theta}{\partial t} = -\theta + c^2 \quad (7.3)$$

ici, on a renommé la variable τ du modèle classique en τ_θ pour la différencier de τ_c .

7.2 Choix des paramètres

Dans l'algorithme BCM classique, les paramètres étaient choisis de manière à respecter l'inégalité suivante :

$$\eta \tau d^2 < 1 \quad (7.4)$$

ce qui signifie que pour qu'il y ait convergence, il faut que la variation des poids (dépendant de η et d^2) soit suffisamment lente par rapport à l'évolution du seuil θ (dépendant de τ). Sinon, on entre dans des régimes oscillatoires.

Or, l'introduction d'un calcul itératif de c a une répercussion sur la vitesse de variation du seuil θ . Intuitivement, le ralentissement de la variation de l'activité c se répercute sur θ - donc la variation de θ est plus lente avec notre modification. Ceci implique donc de ralentir l'apprentissage.

Expérimentalement, on compare notre modèle avec τ_c et τ_θ différenciés, avec le modèle classique tel que $\tau = \tau_c \tau_\theta$. On constate que la variation du seuil θ est beaucoup plus lente pour le modèle classique que pour notre intégration en deux étapes. Ceci nous permet de proposer la relation suivante pour obtenir un taux d'apprentissage suffisamment lent par rapport à τ_c et τ_θ :

$$\eta \tau_c \tau_\theta d^2 < 1 \quad (7.5)$$

7.3 Simulation

Pour tester cette modification du modèle original, nous nous sommes inspirés de l'expérience présentée dans la section 5.4.3 sur l'auto-organisation corticale. Le but était d'observer le développement de l'auto-organisation grâce à la connectivité latérale, avec notre modification du calcul de l'activité.

Contrairement à l'expérience originale, nous avons choisi de tester le développement de l'auto-organisation dans un cas bidimensionnel : l'apprentissage de la sélectivité à l'orientation au sein d'une carte de 20x20 unités. Toutes les unités partagent un même champ récepteur de dimension 15x15, sur lequel 32 filtres de Gabor d'orientation différentes sont présentés selon un tirage aléatoire uniforme.

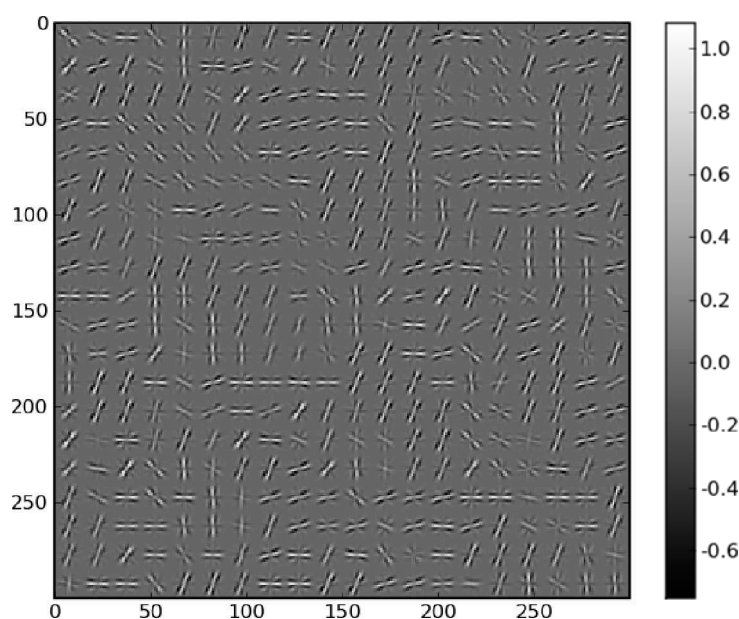


FIGURE 7.1 – Auto-organisation au sein d’une carte 2D, avec une intégration itérative des flux montant et latéral dans l’activité de la carte. La figure représente les poids synaptiques des 20x20 unités d’une carte 20x20. Chaque vignette représente les poids d’une unité.

Au sein de la carte, une connectivité latérale fixe en différence de gaussiennes est mise en place. On considère une carte torique, de sorte que le profil de poids est le même pour toutes les unités, au décalage près. De plus, les poids de l’excitation et de l’inhibition sont normalisés de manière à s’annuler au total.

Contrairement au modèle classique, l’activité de la carte nécessite plusieurs itérations avant de se stabiliser ; pour cette raison, les stimuli sont présentés pendant 30 itérations. L’apprentissage est néanmoins effectué à chaque itération.

La figure 7.1 montre le résultat de l’apprentissage par l’affichage des poids montants de chacune des 20x20 unités de la carte. Malgré le calcul itératif de l’activité et la présentation prolongée des stimuli, les profils de poids obtenus sont similaires à des filtres de gabor, indiquant un développement de la sélectivité à l’orientation. De plus, la répartition des sélectivités dans la carte montre les signes d’un phénomène d’auto-organisation : les unités proches tendant à être sélectives à des orientations similaires et forment des variations généralement progressives.

Chapitre 8

Paramètres de BCM et τ adaptatif

Ce chapitre est une contribution au modèle BCM indépendante des questions d'apprentissage multimodal. Il est néanmoins possible d'utiliser le mécanisme que nous allons décrire dans le cadre du modèle que nous développerons par la suite.

Le modèle BCM utilise deux paramètres pour contrôler la vitesse et la stabilité de l'apprentissage : le taux d'apprentissage η contrôle la vitesse à laquelle les poids synaptiques varient en fonction de l'activité ; la taille τ de la fenêtre sur laquelle la valeur du seuil θ , l'espérance de l'activité au carré, est calculée.

Dans la section 5.2.3, nous avons présenté l'analyse de stabilité effectuée par Blais (1998), qui démontre que si la condition suivante est vérifiée, l'apprentissage tend vers la stabilité (directement, ou par oscillations atténuées) :

$$\tau\eta d^2 < 1 \tag{8.1}$$

où τ contrôle la vitesse d'évolution du seuil θ , η est le taux d'apprentissage et d est l'activité présynaptique. Nous devons donc choisir nos paramètres de simulation de sorte que cette condition soit respectée.

De plus, le paramètre τ doit être choisi de sorte que θ soit une approximation correcte de l'espérance de l'activité au carré. Pour cela, la fenêtre doit être suffisamment grande pour contenir un échantillon représentatif de l'ensemble des stimuli présentés à l'unité.

Enfin, le dernier critère entrant en jeu dans la détermination des paramètres est le temps de simulation. En effet, si le choix d'un τ élevé assure une certaine stabilité au seuil θ et à l'apprentissage, la simulation devient néanmoins extrêmement longue. Les paramètres optimaux sont donc ceux qui respectent les conditions de stabilité tout en permettant d'avoir des temps de simulation à la mesure de nos ressources de calculs.

8.1 Relation entre τ et η

La condition définie par Blais (1998) est le rapport entre la vitesse d'évolution des poids (dépendant de η et de d) et la vitesse d'évolution du seuil θ (dépendant de τ).

Dans notre modèle, considérons que la stimulation issue d'une unité est toujours bornée sur l'intervalle $[0; 1]$. On peut alors ignorer le terme d^2 de la condition de stabilité. L'apprentissage stable le plus rapide est alors :

$$\eta < \frac{1}{\tau} \quad (8.2)$$

Il faut cependant prendre en compte le fait que l'unité reçoit N entrées, et que l'apprentissage a lieu sur chaque connexion entrante. On introduit alors le terme N dans le calcul de η :

$$\eta = \frac{1}{\tau N} \quad (8.3)$$

Dans ce cas, on a un apprentissage stable, quel que soit le nombre d'entrées. Cependant, la convergence se fait par oscillations atténuées. Blais (1998) définit une autre région de l'espace de paramètres pour laquelle la convergence a lieu, mais sans les oscillations :

$$\tau \eta d^2 \leq 3 - 2\sqrt{2} \quad (8.4)$$

Pour plus de stabilité, on peut donc calculer η comme suit :

$$\eta = \frac{1}{\tau N} (3 - 2\sqrt{2}) \quad (8.5)$$

8.2 τ adaptatif

Nous avons montré dans la section 5.3.3 que le calcul itératif de du seuil θ introduit une certaine instabilité - à mesure que l'apprentissage converge vers un point stable, l'activité oscille puisque ce dernier est associé à une activité variant entre 0 et $\frac{1}{p}$, p étant la probabilité d'apparition du stimulus sélectif. De par la nature de son calcul, les oscillations se répercutent dans la dynamique du seuil θ :

$$\dot{\theta} = \frac{1}{\tau} (c^2 - \theta) \quad (8.6)$$

L'instabilité forte de l'activité c , inhérente au développement de la sélectivité, se reporte donc nécessairement sur le seuil θ , comme nous l'avons illustré dans la figure 5.7.

Pour compenser ce phénomène, on choisit généralement une valeur très grande pour θ . Par exemple, Law and Cooper (1994) définit $\tau = 1000$ afin d'avoir une stabilité suffisante pour effectuer un apprentissage sur des images naturelles. L'inconvénient de cette approche est qu'un τ très grand implique un taux d'apprentissage η très petit, et donc un temps de simulation très long.

Autre problème, l'instabilité de θ dépend de la probabilité p d'apparition du stimulus auquel l'unité devient sélective. En effet, si $p = \frac{1}{2}$, l'apprentissage mène l'activité à converger entre 0 et 2. L'instabilité de θ qui en découle est bien plus faible que si $p = \frac{1}{10}$, par exemple. On ne peut donc pas quantifier *a priori* l'instabilité qui sera causée par l'apprentissage, et c'est pour cette raison que l'on choisit une valeur arbitrairement grande pour τ , quitte à rendre la simulation inutilement longue.

Il est cependant possible de contourner ce problème en rendant le paramètre τ adaptatif. En effet, à mesure que l'apprentissage converge, le seuil θ tend progressivement vers $\frac{1}{p}$. La valeur du seuil θ est donc un indicateur de sa propre instabilité. L'idée est donc de ralentir la vitesse de variation du seuil θ en fonction de sa valeur. Le seuil θ variant en fonction de l'activité c au carré, et le but étant le ralentissement progressif de la variation du seuil, nous proposons la formule suivante :

$$\tau = \tau_0(\theta^2 + 1) \quad (8.7)$$

Avec $\tau_0 \geq 1$, la valeur minimale de τ . Ainsi, à mesure que l'apprentissage converge et que le seuil θ augmente, sa variation se fait plus lente pour empêcher le développement de l'instabilité. Naturellement, le même terme doit être introduit dans la règle d'apprentissage afin de préserver le rapport entre τ et η :

$$\eta = \eta_0(\theta^2 + 1) \quad (8.8)$$

Pour illustrer les effets du seuil adaptatif, on effectue un apprentissage sur 10 stimuli orthogonaux équiprobables, et on trace l'évolution de l'activité c et du seuil θ , avec un τ fixe ou adaptatif. L'apprentissage est censé converger vers un point stable pour lequel $c = 10$ pour un stimulus, et $c = 0$ pour les autres.

Pour le cas τ fixe, on choisit $\tau = 10$. Pour la cas τ adaptatif, on choisit $\tau = 10(1 + \theta^2)$. Dans le second cas, τ va donc varier entre 10 au commencement de la simulation, et $10 \cdot 10^2 = 1000$ après convergence.

La figure 8.1 trace l'évolution de l'activité et du seuil θ pour $\tau = 10$ et $\tau = 10(1 + \theta^2)$. On peut faire une première observation sur la forme générale de l'apprentissage : la convergence se fait en deux phases caractéristiques : on a tout d'abord une phase de convergence lente, avec un seuil θ très faible ; puis une phase de convergence rapide, durant laquelle la sélectivité se renforce, le seuil θ montant par la même occasion.

On constate qu'avec un τ fixe, l'instabilité du seuil θ croît fortement lors de la phase de convergence rapide. Dans le cas présent, l'instabilité du seuil θ causée par un paramètre τ trop faible est telle que l'apprentissage ne peut converger jusqu'au point stable - au lieu de cela, la réponse maximale de l'unité sature autour de 5 pour le stimulus sélectif, au lieu de 10.

Avec un τ adaptatif en revanche, l'instabilité du seuil θ n'explose pas pendant la deuxième phase de la convergence. Cette réduction de l'instabilité par le ralentissement progressif au fur et à mesure de la convergence permet à l'apprentissage d'atteindre le point stable ; contrairement au cas $\tau = 10$, où un τ trop faible choisi *a priori* empêche toute convergence au delà d'un certain seuil.

On peut aussi faire la comparaison en prenant $\tau = 1000$ pour la simulation avec τ fixe, ce qui correspond ici à la valeur du τ adaptatif après convergence. La figure 8.2 illustre ce cas. On constate que l'apprentissage possède une bonne stabilité et atteint le point stable. Cependant, la simulation est extrêmement longue, en particulier la première phase de la convergence. On n'atteint le point stable qu'après 8 millions d'itérations environ.

En conclusion, le τ adaptatif permet d'obtenir la stabilité nécessaire à la convergence vers un point stable éloigné, tout en évitant la lenteur inhérente à un τ fixe très large. De plus, l'instabilité du seuil θ dépend de la proximité du point stable vers lequel l'apprentissage converge (*ie.*

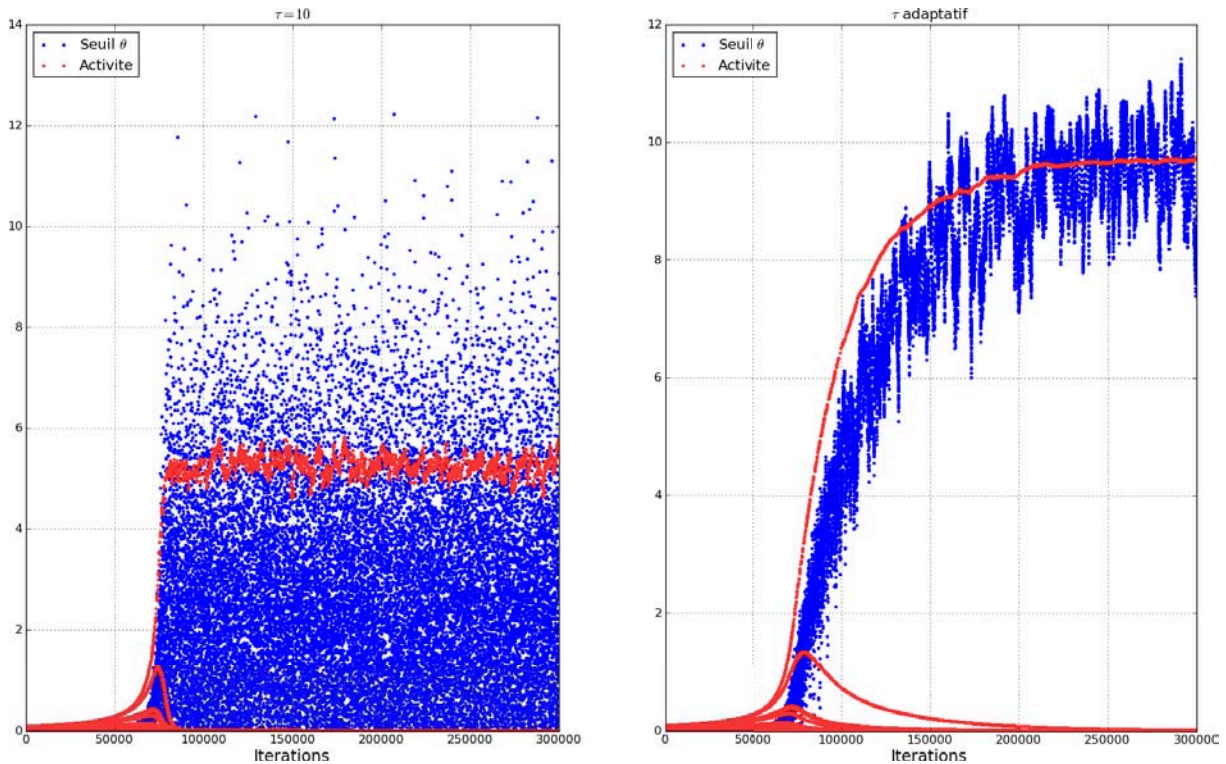


FIGURE 8.1 – Apprentissage sur 10 stimuli orthogonaux, avec $\tau = 10$ (gauche) et $\tau = 10(1 + \theta^2)$ (droite). En bleu, l'évolution du seuil θ . En rose, l'évolution de l'activité. Cette dernière illustre bien la différenciation progressive qui se fait entre les différents stimuli présentés : après convergence, la réponse est forte pour un stimulus et faible pour les autres. Sur la figure de gauche, on observe une forte instabilité du seuil θ augmentant à mesure que l'apprentissage converge. A droite, le τ adaptatif permet de réduire cette instabilité, permettant la convergence de l'apprentissage vers le point stable $\theta \approx \frac{1}{p}$.

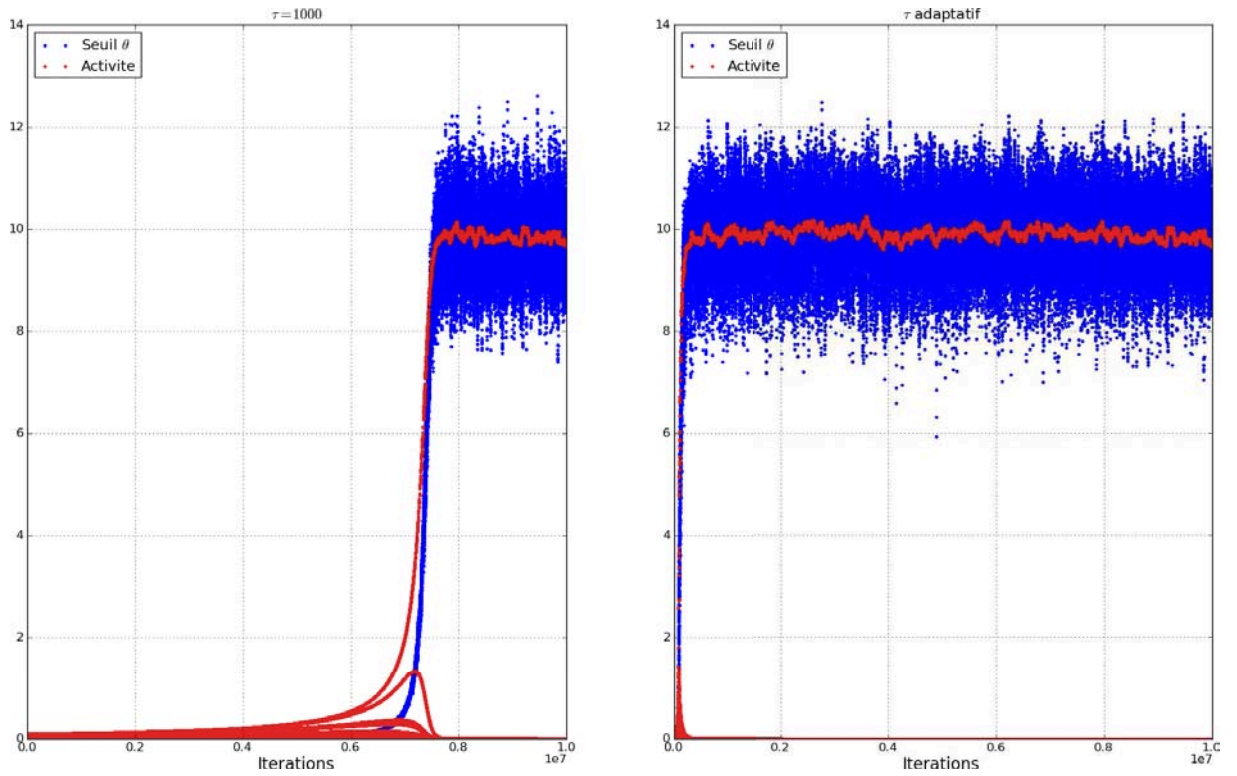


FIGURE 8.2 – Apprentissage sur 10 stimuli orthogonaux, avec $\tau = 1000$ (gauche) et $\tau = 10(1 + \theta^2)$ (droite). Les points bleus marquent les variations du seuil θ . Les points roses marquent les variations de l'activité. Dans les deux cas l'apprentissage converge correctement vers le point stable $\theta \approx \frac{1}{p}$, mais on constate que le τ réduit grandement le temps de convergence de l'apprentissage.

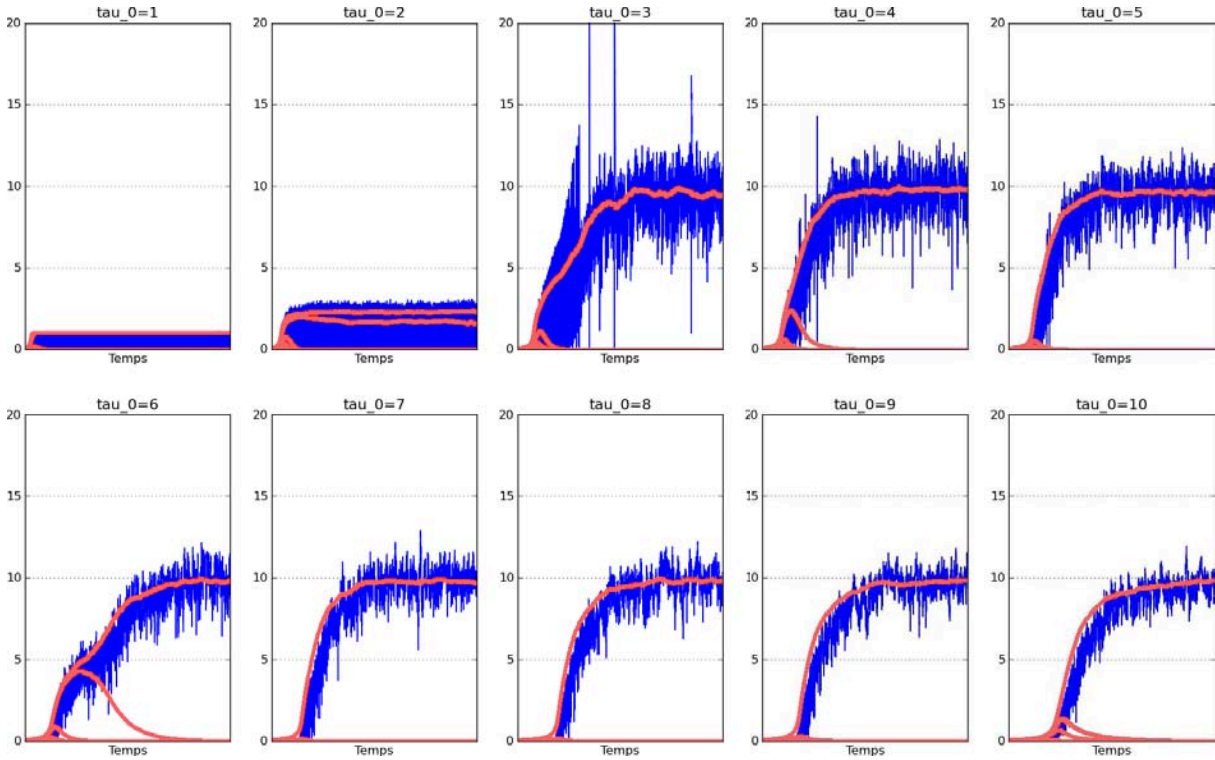


FIGURE 8.3 – Exemples de convergences de l’apprentissage sur 10 stimuli orthogonaux et équiprobables, en fonction de la valeur de τ_0 . Comme sur les figures 8.1 et 8.2 le bleu représente l’évolution du seuil θ et le rose l’évolution de l’activité.

de la probabilité associée à ce point stable). La convergence vers un point stable proche (forte probabilité associée) ne nécessite pas de ralentir la simulation avec un τ très grand - mais il s’agit d’une information qu’on ne peut pas connaître *a priori*, à moins de travailler sur des jeux de données totalement artificiels. Le τ adaptatif permet de “ralentir” la simulation uniquement selon la nécessité, pendant son exécution.

8.3 Valeur initiale τ_0

Nous avons proposé un τ adaptatif permettant de contourner le problème de l’instabilité du seuil liée à la convergence, et l’analyse de Blais (1998) permet de déterminer η en fonction de τ . Cependant, nous avons introduit un nouveau paramètre, τ_0 , la valeur minimale de τ .

Le paramètre τ_0 est déterminant en début de simulation, lorsque les poids sont faibles et l’unité non-sélective. Dans ce cas, le seuil θ est généralement très proche de 0, ce qui implique que $\tau \approx \tau_0$. A mesure que θ augmente, τ_0 devient négligeable dans le calcul de τ .

La figure 8.3 reprend la simulation précédente, à savoir un apprentissage sur 10 stimuli orthogonaux équiprobables, et fait varier la valeur de τ_0 . On observe différents comportements selon la valeur de τ_0 .

Il nous semble difficile de déterminer une valeur optimale pour τ_0 . La figure 8.3 montre clairement que si τ_0 est trop faible, la sélectivité ne se développe pas. La sortie de cet état non-sélectif initial dépend visiblement de τ_0 . On pourrait comparer le cas de figure où τ_0 est trop petit à la

variation instantanée du seuil : $\theta = c^2$. Dans ce cas, l'apprentissage joue son rôle de normalisation des poids de sorte que l'activité moyenne $\bar{c} \approx 1$, mais la sélectivité ne peut se développer.

Dans l'exemple donné, il semble possible de développer la sélectivité à un stimulus (c'est à dire d'atteindre un point stable) avec des valeurs de τ_0 assez faibles (e.g. $\tau_0 = 3$). Il est cependant difficile de généraliser à partir de cet exemple et de proposer une règle pour choisir τ_0 . La convergence de l'apprentissage en fonction de τ_0 n'est probablement pas la même suivant la taille du vecteur d'entrées et la probabilité des stimuli présentés.

Enfin, un τ_0 trop faible implique que le calcul de θ ne peut pas être une approximation correcte de l'espérance de l'activité au carré. Mais la qualité de cette approximation importe surtout lorsque l'apprentissage converge vers les points stables - en début de simulation, on en est encore très loin.

Chapitre 9

Modèle d'apprentissage multimodal

Comme nous l'avons dit précédemment, l'intégration multimodale est la capacité à élaborer des représentations cohérentes à partir de plusieurs modalités sensorielles et/ou motrices. Le but de notre travail est de concevoir un modèle d'apprentissage permettant l'élaboration de telles représentations.

Dans le chapitre 2, nous avons présenté le modèle de Guigon (1993). Ce modèle propose un mécanisme d'intégration multimodale en deux étapes : dans un premier temps, chaque modalité intègre séparément sont activité, puis les activités sont combinées à l'aide des coefficients du vecteur L (associant un poids à chaque modalité) et de la matrice Q (associant un poids à chaque paire de modalités).

L'apprentissage fonctionne lui aussi sur deux niveaux : dans un premier temps, chaque modalité effectue séparément un apprentissage non-supervisé sur ses entrées, puis un second apprentissage a lieu sur les coefficients de L et Q , de sorte que ceux-ci reflètent les relations de corrélation entre l'activité des différentes modalités.

Notre approche s'inspire de ce principe d'intégration distincte de plusieurs modalités au sein de chaque unité de calcul. Cette approche suppose qu'au sein d'une colonne corticale, plusieurs modalités peuvent être intégrées séparément, avec un apprentissage distinct pour chaque modalité.

Pour ce qui est de l'apprentissage, nous optons pour une voie différente : nous mettons l'accent sur la notion de *co-apprentissage*. Chaque modalité apprend une sélectivité à un stimulus particulier, mais toutes les modalités s'influencent réciproquement de sorte que les modalités sont sélectives à des stimuli généralement co-occurents.

9.1 Co-apprentissage et binocularité

Notre but est de faire du co-apprentissage entre plusieurs modalités. Dès l'article fondateur (Bienenstock et al., 1982), le modèle BCM en propose un exemple par le biais du modèle de la binocularité, que nous avons présenté dans la section 5.4.1. Dans ce modèle, une unité reçoit deux projections, provenant de deux vecteurs sources représentant des rétines.

$$c = \mathbf{d}_l \cdot \mathbf{w}_l + \mathbf{d}_r \cdot \mathbf{w}_r \quad (9.1)$$

avec c l'activité de l'unité, $\mathbf{d}_l$ et $\mathbf{d}_r$ les deux rétines, $\mathbf{w}_l$ et $\mathbf{w}_r$ les vecteurs de poids correspondant.

En se basant sur ce modèle, les auteurs font varier les stimuli présentés sur les deux rétines pour reproduire les expériences de Hubel and Wiesel (1962) (voir section 5.4.1). Les auteurs remarquent plusieurs propriétés de l'apprentissage BCM :

- lorsque les mêmes stimuli sont présentés sur les deux rétines, l'unité développe la même sélectivité sur les deux projections, quand bien même les poids initiaux seraient différents ;
- lorsqu'une rétine est bruitée, les poids de la projection associée faiblissent, causant la dominance de l'autre projection ;
- lorsque les deux rétines présentent des stimuli décorrélés, l'unité développe une *dominance oculaire* : une projection est faiblement sélective et joue peu dans l'activité de l'unité tandis que l'autre est fortement sélective.

Ces expériences nous amènent à poser la question suivante : qu'en est-il du cas où les deux entrées sont partiellement corrélées ? Nous proposons donc de répondre à cette question.

9.1.1 Dominance et corrélation des entrées

Pour étudier le comportement de l'apprentissage binoculaire en fonction de la corrélation des deux entrées, nous mettons en place le protocole expérimental suivant.

On se base sur le jeu de données des stimuli orthogonaux. Dix stimuli présentés sur deux entrées en parallèle. On présente le même stimulus sur les deux entrées avec une probabilité p . Sinon, on présente deux stimuli différents. Ainsi, on peut faire varier la corrélation entre les deux entrées tout en gardant une distribution uniforme pour chaque entrée séparée. L'apprentissage se déroule sur 1000 unités en parallèle.

La figure 9.1 illustre la dominance en fonction de la corrélation. Soit q_i^{max} et q_i^{min} les sélectivités des projections de l'unité i , de sorte que $q_i^{min} \leq q_i^{max}$. La moyenne des q_i^{min} et q_i^{max} sur la population nous donne la sélectivité dominante moyenne et la sélectivité dominée moyenne pour une corrélation des entrées donnée.

Pour calculer la dominance, on se base sur la formule suivante :

$$d_i = 1 - \frac{q_i^{min}}{q_i^{max}} \quad (9.2)$$

où d_i est la dominance de l'unité d'indice i . Ici encore, la moyenne sur la population nous donne la dominance moyenne pour une corrélation des entrées donnée.

La figure 9.1 montre bien que le phénomène de dominance se met très vite en place à mesure que la corrélation des entrées décroît. Même avec 80% de chance de présenter deux stimuli identiques (corrélation 0.8), on observe une forte dominance d'une projection sur l'autre.

Notre but étant de faire de l'apprentissage entre des entrées partiellement corrélées, ce phénomène de dominance d'une modalité sur l'autre nous apparaît comme problématique - en effet, nous ne souhaitons pas que l'apprentissage en vienne systématiquement à ignorer toutes les modalités entrantes sauf une.

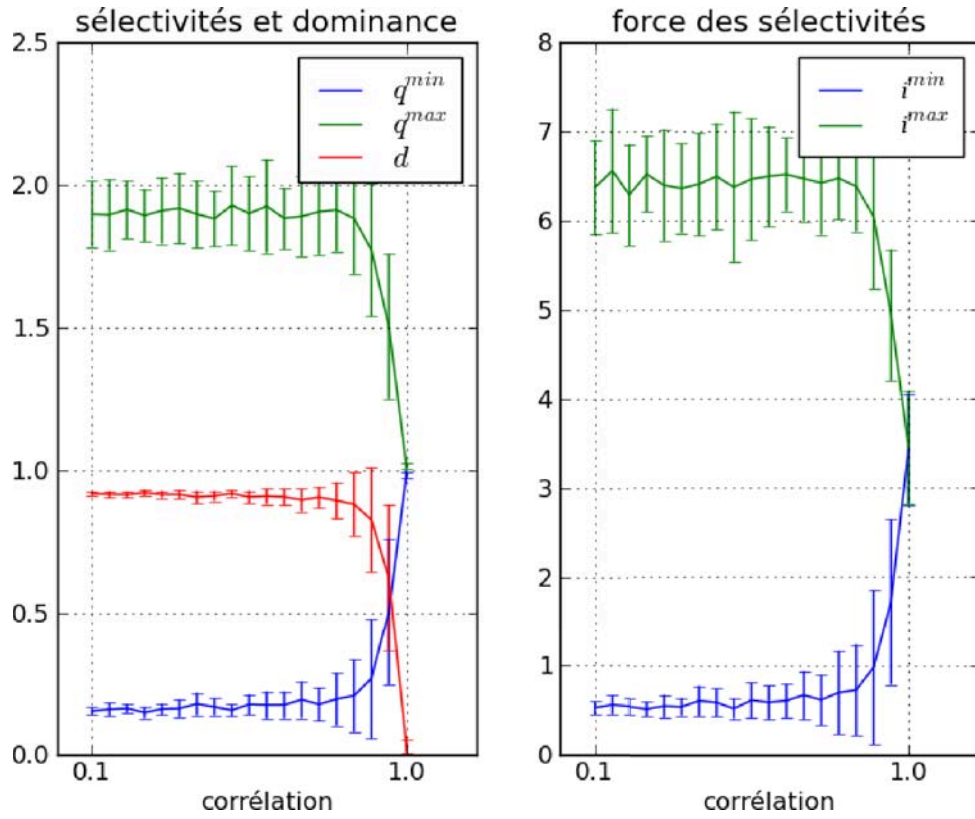


FIGURE 9.1 – Dominance en fonction de la corrélation entre les deux projections de l'unité. q^{min} est la qualité moyenne de la sélectivité de la projection dominée, q^{max} la qualité moyenne de la sélectivité de la projection dominante, d la dominance moyenne, i^{min} l'intensité moyenne de la sélectivité de la projection dominée et i^{max} l'intensité moyenne de la sélectivité de la projection dominante. Ces valeurs sont calculées pour 1000 unités - les intervalles sont les écarts types.

9.1.2 Solution au problème de la dominance

Dans le modèle BCM, le mécanisme de seuil adaptatif permet la normalisation des poids, de sorte que l'activité moyenne de l'unité reste proche de $\bar{c} = 1$. Ce mécanisme s'accompagne d'une compétition entre les entrées permettant le développement de la sélectivité.

Dans les simulations *BCM binoculaire*, aucune distinction n'est faite entre les deux modalités correspondant aux deux rétines. Il en résulte que la normalisation des poids et la compétition se font à une échelle globale, mettant en concurrence les différentes modalités. Il en découle le phénomène de dominance observé lorsque les entrées sont partiellement corrélées.

Une façon d'éviter ce phénomène de dominance serait de mettre en place un mécanisme de normalisation indépendant pour chaque modalité, tout en gardant la propriété de co-apprentissage. Pour cela, nous proposons de traiter chaque modalité séparément, en calculant une activité et un seuil adaptatif propre à chaque modalité. Ainsi, la normalisation s'opère de sorte que l'activité moyenne issue de chaque modalité reste proche de $\bar{c}_i = 1$.

Afin de réintroduire le co-apprentissage, on propose un mécanisme de modulation. La modulation est un phénomène selon lequel une source tierce peut altérer la réponse d'un neurone aux stimulations qu'il reçoit. Salinas and Thier (2000) discute de la plausibilité biologique de tels mécanismes et des intérêts en terme de traitement de l'information, en particulier dans le contexte multimodal.

Ce mécanisme de modulation permet d'orienter l'apprentissage ayant lieu sur chaque modalité. En calculant le terme de modulation à partir de l'activité de toutes les modalités, on introduit une influence réciproque permettant le co-apprentissage.

Nous allons donc nous intéresser à la manière d'introduire un terme de modulation dans le modèle BCM, de manière à orienter son apprentissage.

9.2 Modulation et apprentissage guidé

Nous avons vu précédemment que par le biais de la connectivité latérale, il est possible d'influencer l'apprentissage effectué par BCM de manière à favoriser la convergence vers une certaine sélectivité. Ce mécanisme est utilisé pour le codage en population et l'auto-organisation.

Cette combinaison d'un apprentissage autonome et d'une supervision faible est l'idée que nous développons dans Girod and Alexandre (2009). Il s'agit d'un principe simple qui nous semble adapté à l'apprentissage cortical. En effet, il reste non-supervisé tout en permettant des interactions telles que la modulation de l'apprentissage montant par l'activité du flux descendant.

On peut voir plusieurs manières d'introduire une modulation de l'apprentissage dans le modèle BCM. L'approche utilisée par l'introduction de la connectivité latérale consiste à moduler l'activité de l'unité, ce qui aura une influence sur l'apprentissage. Une autre approche consiste à introduire la modulation directement dans la règle d'apprentissage. Nous étudierons les deux approches.

9.2.1 Protocole de test

Pour illustrer et comparer les différentes formes de modulation que nous allons introduire, nous mettons en place deux protocoles expérimentaux. Nous nous basons sur le protocole des

stimuli orthogonaux auquel nous ajoutons le terme de modulation. Le principe consiste à présenter une modulation favorisante lorsque le premier stimulus du jeu de stimuli est présenté. Après convergence, on calcule la proportion d'unités sélectives au premier stimulus. En l'absence de modulation, cette proportion doit valoir environ $\frac{1}{N}$, N étant le nombre de stimuli présentés.

Le premier protocole expérimental consiste à présenter la modulation favorisante à chaque fois que le premier stimulus est présenté, mais à faire varier l'intensité de la modulation d'une simulation à l'autre.

Le deuxième protocole expérimental consiste à garder la même intensité de modulation d'une simulation à l'autre, mais à faire varier la probabilité d'apparition de la modulation favorisante lorsque le premier stimulus est présenté. Par la suite, on appellera cette probabilité *fiabilité* de la modulation.

Dans les deux protocoles, nous analyserons les résultats de l'apprentissage en déterminant pour chaque unité le stimulus sélectif s et l'intensité de la sélectivité i . Ces données nous permettront de calculer la proportion d'unités sélectives au stimulus favorisé en fonction de l'intensité ou de la probabilité de la modulation. L'expérience sera reproduite dix fois pour chaque jeu de paramètres afin de calculer une moyenne et un écart type.

9.2.2 Modulation de l'activité

La modulation de l'activité consiste à introduire un terme de modulation m dans le calcul de l'activité c . On peut distinguer la modulation additive (9.3) et la modulation multiplicative (9.4).

$$c = m + \sum_{i=0}^n w_i \cdot d_i \quad (9.3)$$

$$c = m \sum_{i=0}^n w_i \cdot d_i \quad (9.4)$$

Dans le cas où on fait usage d'une fonction d'activation σ pour le calcul de l'activité, la modulation de l'activité peut être vue comme une variation de la fonction d'activation - la modulation additive correspondant à une variation du seuil, et la modulation multiplicative correspondant à une variation du gain.

Étude de stabilité

Étudions les effets de la modulation de l'activité sur les points stables de BCM en reprenant les analyses de la section 5.3. Nous avons montré qu'à stimulation constante, la règle d'apprentissage converge vers une valeur de w telle que $c = 1$. La modulation d'activité ne change pas cette propriété, uniquement la valeur de w pour laquelle $c = 1$:

$$c = dw \iff w = \frac{1}{d} \quad (9.5)$$

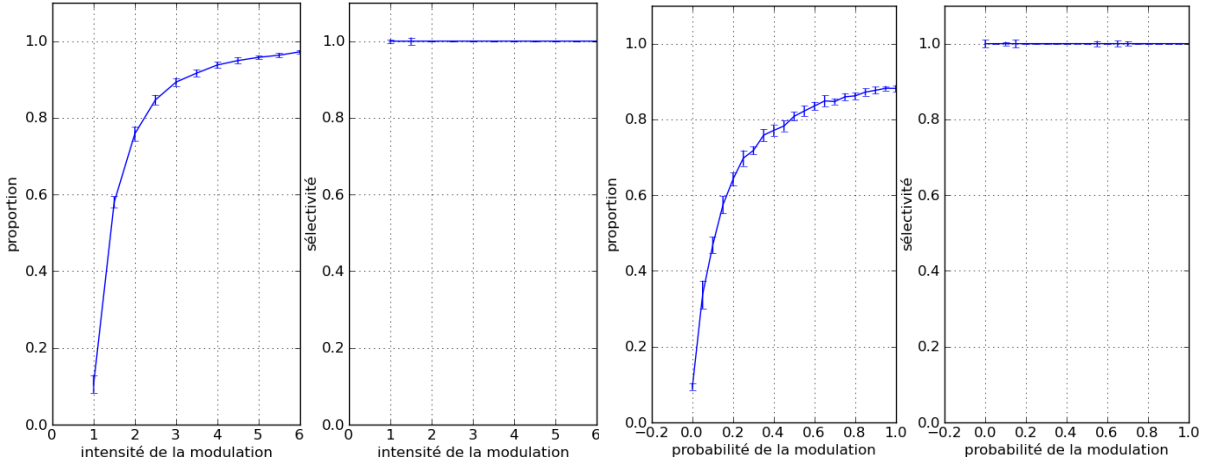


FIGURE 9.2 – Modulation multiplicative : proportion d’unités sélectives au stimulus favorisé en fonction de l’intensité de la modulation (gauche) ou de la probabilité d’apparition de la modulation (droite)

$$c = dw + m \iff w = \frac{1 - m}{d} \quad (9.6)$$

$$c = dwm \iff w = \frac{1}{dm} \quad (9.7)$$

avec c l’activité, d l’activité présynaptique, w le poids synaptique et m la modulation.

L’autre analyse est celle de la stabilité dans le cas unidimensionnel, avec un stimulus variant entre $d = 0$ et $d = s$ avec la probabilité $P(s) = p$. Dans ce cas, on constate que l’apprentissage converge vers une valeur de w telle que l’activité c oscille entre 0 et $\frac{1}{p}$, soit $w = \frac{1}{ps}$.

Cette même analyse transposée au cas d’une modulation m constante (additive ou multiplicative) donne les résultats suivants :

$$c = dw \iff w = \frac{1}{ps} \quad (9.8)$$

$$c = dw + m \iff w = \frac{1}{ps} - m \quad (9.9)$$

$$c = dwm \iff w = \frac{1}{psm} \quad (9.10)$$

On constate que dans les trois cas, l’apprentissage se stabilise à une valeur de w de sorte que $c = 0$ ou $c = \frac{1}{p}$, en prenant en compte le terme de modulation m .

Effets de la modulation multiplicative

La figure 9.2 (gauche) présente les résultats obtenus avec le premier protocole expérimental : on fait varier l’intensité de la modulation pour observer son incidence sur la convergence de l’apprentissage. L’expérience montre bien une corrélation positive entre l’intensité de la modulation favorisante et le biais en faveur du premier stimulus.

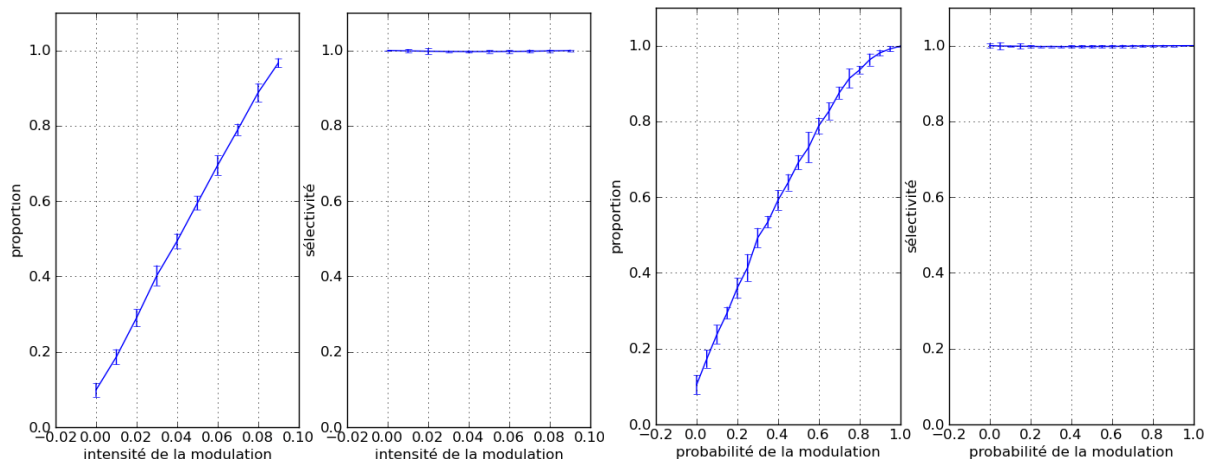


FIGURE 9.3 – Modulation additive : proportion d’unités sélectives au stimulus favorisé en fonction de l’intensité de la modulation (gauche) ou de la probabilité d’apparition de la modulation (droite)

La figure 9.2 (droite) présente les résultats obtenus avec le deuxième protocole expérimental : on fait varier la probabilité d’apparition de la modulation favorisante pour observer son incidence sur la convergence de l’apprentissage. Ici, l’intensité de la modulation reste donc fixe avec $m = 3$, mais la fiabilité de la modulation varie. L’expérience montre bien une corrélation positive entre la probabilité d’apparition de la modulation favorisante et le biais en faveur du premier stimulus, même pour une probabilité faible.

Ces expériences montrent que la modulation multiplicative de l’activité est un moyen valide pour guider l’apprentissage dans le modèle BCM.

Effets de la modulation additive

La figure 9.3 (gauche) illustre le biais causé par la modulation additive en fonction de son intensité. On observe une corrélation positive entre l’intensité de la modulation et le biais en faveur du premier stimulus.

La figure 9.3 (droite) illustre le biais causé par la modulation additive en fonction de la probabilité d’apparition de la modulation. Ici, l’intensité de la modulation additive est fixe à $m = 0.1$. On observe là encore une corrélation positive entre la probabilité d’apparition de la modulation et le biais en faveur du premier stimulus.

9.2.3 Modulation de l’apprentissage

La modulation de l’apprentissage consiste à introduire un terme de modulation dans la règle d’apprentissage plutôt que dans le calcul de l’activité de l’unité. Au sein de la règle d’apprentissage, la valeur de l’activité c par rapport au seuil θ détermine le signe de l’apprentissage. La modulation a donc pour but de faire varier la valeur de l’activité, de sorte à favoriser la LTP ou la LTD selon la nature de la modulation (favorisante ou défavorisante). On introduit le terme de modulation dans la fonction ϕ . Comme pour la modulation de l’activité, on peut imaginer une modulation additive :

$$\phi(c, \theta, m) = c(c + m - \theta) \quad (9.11)$$

ou multiplicative :

$$\phi(c, \theta, m) = c(cm - \theta) \quad (9.12)$$

Afin d'illustrer les effets de la modulation de l'apprentissage, on reprendra les mêmes expériences que pour la modulation de l'activité : une modulation facilitatrice est transmise au neurone lors de la présentation du premier stimulus.

Étude de stabilité

Pour étudier la stabilité de l'apprentissage, nous nous sommes intéressés à la fonction ϕ . Celle-ci détermine le signe de l'apprentissage, et nous avons montré que les points stables existent lorsque $\phi(c, \theta) = 0$ quelle que soit l'entrée. Sans modulation nous avons vu que le seul point stable existant est $c = 1$. Pour étudier les effets de l'ajout d'un terme de modulation dans la règle d'apprentissage, on reprend le cas de la variation instantanée du seuil, $\theta = c^2$.

Modulation multiplicative Pour la modulation multiplicative avec $m > 0$, on obtient :

$$\phi(c) = c(cm - c^2) \iff c^2(m - c) \quad (9.13)$$

Sous cette forme $\phi(c)$ s'annule donc pour $c = 0$ et $c = m$. Comme pour BCM sans modulation, le signe de $\phi(c)$, illustré par la figure 9.4 montre que l'apprentissage admet un point stable $c = m$.

Tout comme pour la modulation d'activité, on peut étendre l'analyse au cas de la stabilité dans le cas unidimensionnel, avec un stimulus variant entre $d = 0$ et $d = s$ avec la probabilité $P(s) = p$. On constate alors que l'apprentissage se stabilise sur $w = \frac{m}{p}$.

Modulation additive Pour la modulation additive, on obtient :

$$\phi(c) = c(c + m - c^2) \quad (9.14)$$

Contrairement à la modulation multiplicative, la modulation additive introduit des modifications importantes dans la stabilité. La figure 9.5 montre que pour $m = 2$, on obtient deux points stables. Plus généralement, on obtient :

modulation	points stables
$m > -\frac{1}{4}$	$c = \frac{1}{2}(1 + \sqrt{4m + 1})$ et $c = \frac{1}{2}(1 - \sqrt{4m + 1})$
$m = -\frac{1}{4}$	$c = \frac{1}{2}$
$m < -\frac{1}{4}$	$c = 0$

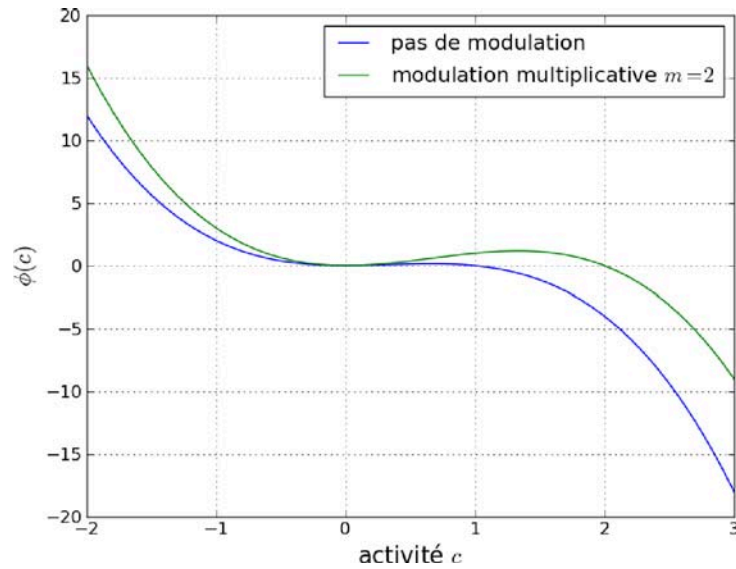


FIGURE 9.4 – Comparaison de la fonction $\phi(c)$, avec et sans modulation multiplicative. Comme pour la figure 5.3, $\theta = c^2$. Avec un stimulus constant strictement positif, $\phi(c)$ contrôle le signe de l'apprentissage. On constate que l'ajout du terme de modulation n'ajoute pas de point stable, mais déplace le point stable de $c = 1$ à $c = m$.

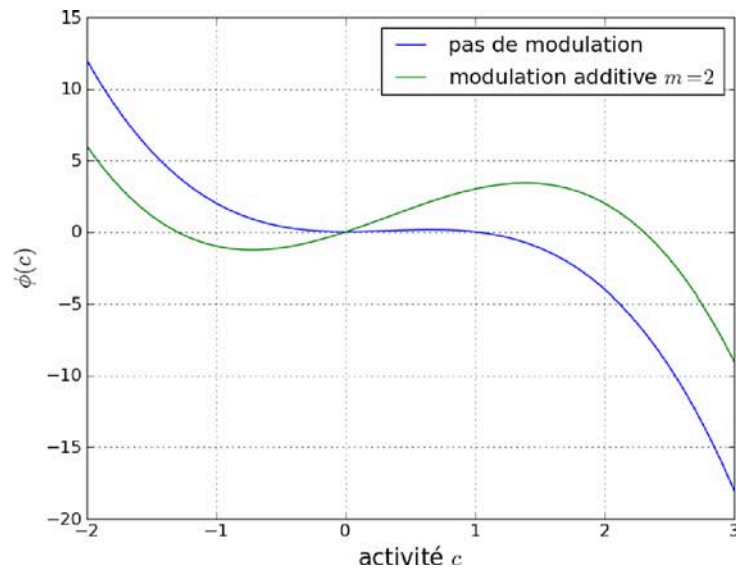


FIGURE 9.5 – Comparaison de la fonction $\phi(c)$ avec et sans modulation additive. Contrairement au cas de la modulation multiplicative illustré précédemment (fig 9.4), l'ajout du terme de modulation additive introduit un deuxième point stable négatif.

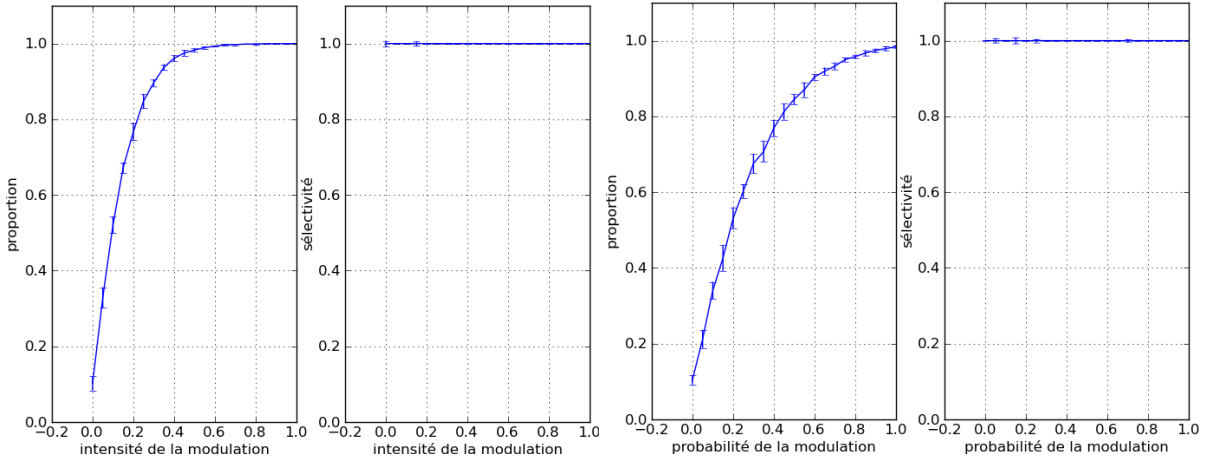


FIGURE 9.6 – Modulation additive de l'apprentissage : proportion d'unités sélectives au stimulus favorisé en fonction de l'intensité de la modulation (gauche) ou de la probabilité d'apparition de la modulation (droite)

La modulation additive est donc plus problématique : soit elle ajoute un second point stable à l'apprentissage, soit elle déplace le point stable à $c = 0$, ce qui impliquerait une convergence vers des poids nuls.

Dans le cas de la stabilité dans le cas unidimensionnel, avec un stimulus variant entre $d = 0$ et $d = s$ avec la probabilité $P(s) = p$, les points stables deviennent :

$$c = \frac{1}{2p}(1 + \sqrt{4pm + 1}), 4pm > -1 \quad (9.15)$$

$$c = \frac{1}{2p}(1 - \sqrt{4pm + 1}), 4pm > -1 \quad (9.16)$$

Effets de la modulation additive

La figure 9.6 illustre le biais causé par la modulation en fonction de l'intensité de la modulation (à gauche) et en fonction de la probabilité d'apparition d'une modulation fixe $m = 0.5$. On observe bien une corrélation positive entre la modulation et le biais en faveur du stimulus favorisé. Malgré la modification du nombre de points stables indiquée par l'analyse, la qualité de la convergence ne semble pas affectée dans le cas présent.

En l'état, il est difficile de généraliser à partir de ce contexte simple. Lorsque l'apprentissage se fait sur des stimuli riches, l'évolution des poids amène généralement à des variations de l'activité allant vers des valeurs négatives. Il se pourrait que dans certains cas de figures, l'apprentissage soit gêné par ce point stable négatif.

Effets de la modulation multiplicative

La figure 9.7 illustre le biais causé par la modulation en fonction de l'intensité de la modulation (à gauche) et en fonction de la probabilité d'apparition d'une modulation fixe $m = 2$. On

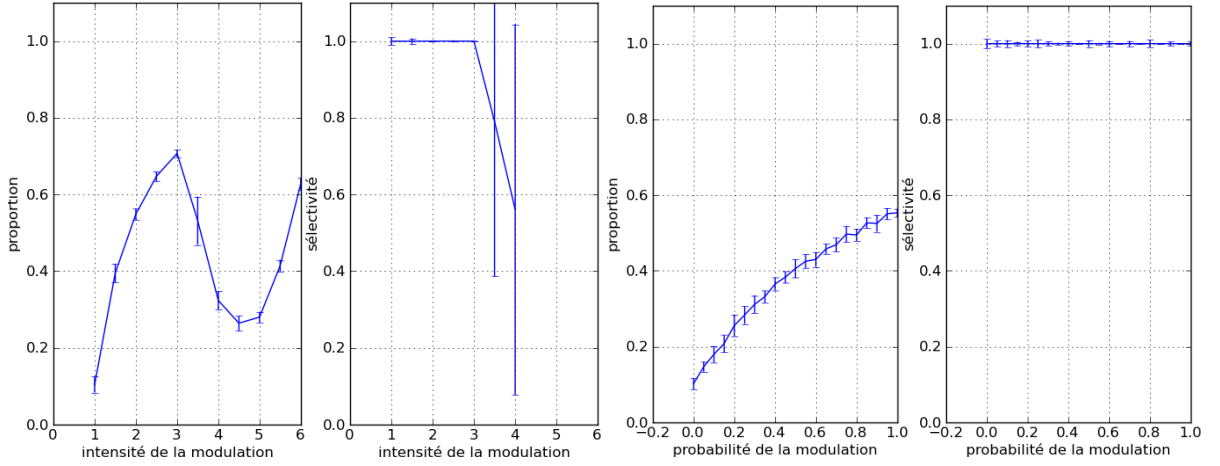


FIGURE 9.7 – Modulation multiplicative de l’apprentissage : proportion d’unités sélectives au stimulus favorisé en fonction de l’intensité de la modulation (gauche) ou de la probabilité d’apparition de la modulation (droite)

observe un début de biais en faveur du stimulus favorisé, puis une cassure provoquée par des problèmes de convergence lorsque la modulation est trop forte. Il semble que ce problème soit lié au fait que l’introduction de la modulation invalide la condition de stabilité $\eta\tau d^2 < 1$ introduite par Blais (1998) (voir chap 5). En conséquence, une modulation trop forte peut causer des instabilités telles qu’on les observe lorsque les paramètres τ et η sont mal choisis, ce qu’on observe dans le cas présent.

9.3 Modulation et co-apprentissage

Notre approche du co-apprentissage entre plusieurs modalités se base sur le calcul d’un terme de modulation à partir de l’activité des modalités, qui sera ensuite appliqué aux modalités afin d’orienter l’apprentissage. Nous avons présenté deux manières d’introduire une modulation dans le modèle BCM. La première consiste à biaiser indirectement l’apprentissage en modulant l’activité de l’unité. La seconde consiste à biaiser directement la règle d’apprentissage. Dans les deux cas, nous avons observé le comportement de la modulation pour le cas additif et le cas multiplicatif.

Dans le cas du co-apprentissage, la modulation de l’activité nous semble problématique car elle introduit une récurrence dans le calcul de l’activité. En effet, la modulation est calculée à partir de l’activité des modalités, qui est elle-même modifiée par le terme de modulation. En revanche, la modulation sur l’apprentissage n’a pas ce problème, puisque l’activité des unités n’est pas modifiée. Ceci nous amène à préférer un mécanisme de modulation de l’apprentissage. De plus, nous avons vu que la modulation multiplicative de l’apprentissage possède certains risques d’instabilité, ce qui nous amène à nous orienter sur la modulation additive de l’apprentissage.

Sur ces considérations, nous proposons donc le modèle suivant. Soit une unité recevant plusieurs modalités. c_i est l’activité issue de la i -ème modalité et un seuil adaptatif θ_i est calculé pour chaque modalité. s est l’activité totale de l’unité :

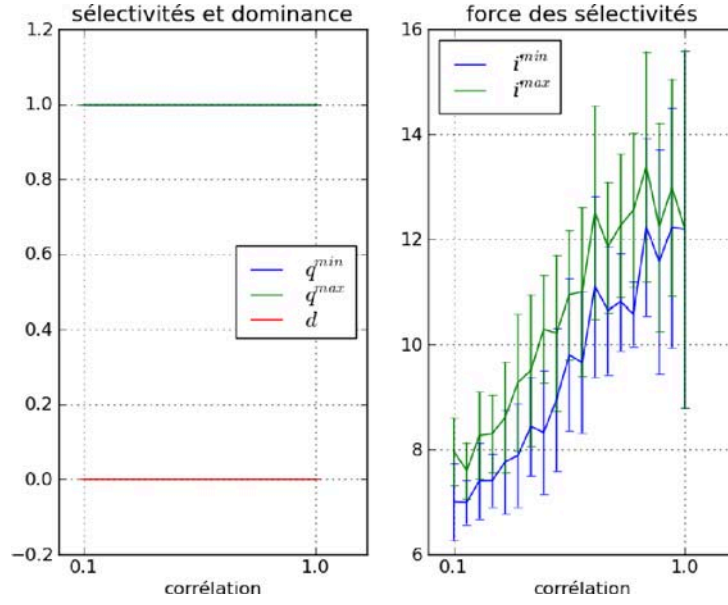


FIGURE 9.8 – Dominance en fonction de la corrélation entre les deux projections de l'unité, dans le cas de l'apprentissage indépendant. q^{min} est la qualité moyenne de la sélectivité de la projection dominée, q^{max} la qualité moyenne de la sélectivité de la projection dominante, d la dominance moyenne, i^{min} la force moyenne de la sélectivité de la projection dominée et i^{max} la force moyenne de la sélectivité de la projection dominante.

$$s = \sum_i c_i \quad (9.17)$$

Pour introduire la modulation, nous modifions le terme postsynaptique ϕ de la règle d'apprentissage :

$$\phi(c_i, \theta_i, s) = c_i(s - \theta_i) \quad (9.18)$$

Ceci est équivalent à une modulation additive de l'apprentissage :

$$\phi(c_i, \theta_i, s) = c_i(c_i + m - \theta_i) \quad (9.19)$$

avec

$$m = \sum_{j \neq i} c_j \quad (9.20)$$

Pour illustrer les effets de cette nouvelle règle d'apprentissage, reprenons l'expérience de la section 9.1.1, mettant en avant l'effet de dominance oculaire allant de paire avec la décorrélation des deux modalités.

La figure 9.8 (gauche) montre que quel que soit le degré de corrélation entre les deux modalités, la qualité de la sélectivité est élevée sur les deux modalités. La figure 9.8 (droite) montre que la décorrélation des entrées fait baisser la force de la sélectivité. Néanmoins, l'écart entre la force de deux modalités reste relativement faible.

9.4 Résumé

Nous proposons un modèle basé sur l'intégration séparée de plusieurs modalités au sein d'une unité de calcul. L'apprentissage multimodal est obtenu en effectuant un apprentissage non-supervisé sur chaque modalité, et en mettant en place une influence réciproque entre les apprentissages par le biais d'un mécanisme de modulation. L'influence réciproque permet de favoriser le développement de sélectivités à des événements co-occurents, développant ainsi des représentations multimodales.

Chapitre 10

Application : apprentissage de la relation posture/position

Pour tester notre modèle d'apprentissage multimodal, nous l'appliquons au protocole posture/position décrit dans la section 6.2. Le but est de faire le lien entre la posture du bras et la position de la main. Rappelons qu'il ne s'agit pas ici de faire un modèle réaliste du contrôle moteur, mais d'apprendre une relation de corrélation non-triviale entre deux modalités.

10.1 Apprentissage guidé

Notre modèle d'apprentissage multimodal se base sur la notion d'apprentissage guidé : chaque modalité effectue un apprentissage non-supervisé, et toutes les modalités s'influencent réciproquement pour biaiser leurs apprentissages respectifs en faveur d'un tout cohérent.

Pour illustrer ce mécanisme sur la relation posture-position, nous allons procéder par étapes. Dans un premier temps, nous allons apprendre sur une modalité et utiliser la seconde pour en guider l'apprentissage.

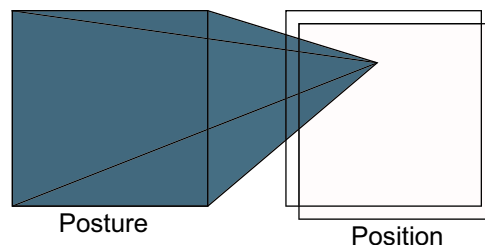


FIGURE 10.1 – Réseau posture-position : les deux cartes reçoivent une activité externe représentant la posture et la position respectivement, et l'activité issue de la carte de posture se propage vers la carte de position par le biais de la projection sur laquelle l'apprentissage a lieu.

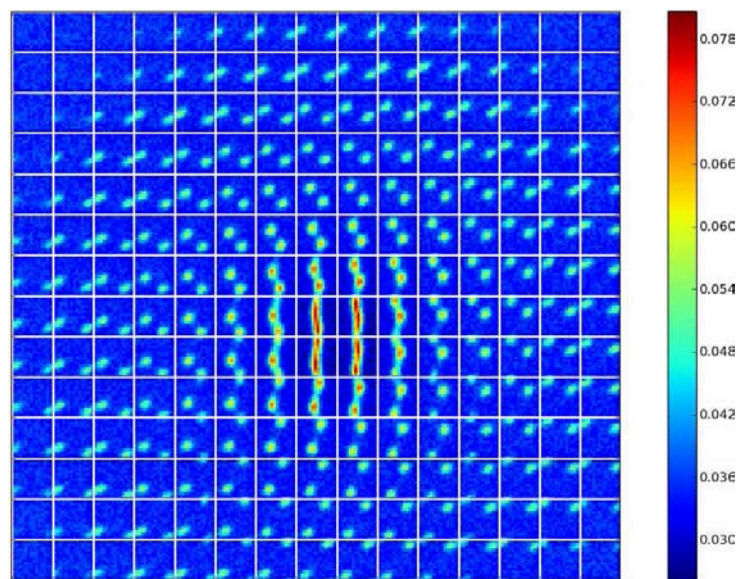


FIGURE 10.2 – Poids de la projection posture-position après apprentissage. L'image est construite en mettant cote à cote les champs récepteurs des 15x15 unités de la carte de position. Chaque vignette illustre donc la sélectivité à la posture acquise par l'unité à cet emplacement de la carte de position.

10.1.1 Relation posture - position

On connecte la carte de posture à la carte de position à l'aide d'une projection complète dont les poids sont initialisés aléatoirement selon une distribution uniforme et sur laquelle on effectue l'apprentissage. Chaque unité de la carte de position développe donc la sélectivité à une certaine posture, sous l'influence de l'information de position présente localement.

La figure 10.2 montre le résultat de cet apprentissage. L'apprentissage BCM classique voudrait que chaque unité développe la sélectivité à une certaine posture, c'est à dire une position de la bulle d'activité dans la carte de posture. Mais dans le cas présent, on constate que les profils de poids indiquent des sélectivités différentes. Ceci est dû à l'influence de la modulation de position sur l'apprentissage. En effet, on distingue quatre profils de poids types, qui correspondent à quatre cas particuliers de la relation posture-position :

1. les profils de poids des vignettes en périphérie sont uniformes, indiquant par là une absence de sélectivité - quand bien même ces unités partagent le même champ récepteur que les autres unités. Ces unités correspondent à des positions que le bras ne peut jamais atteindre, même totalement déplié. Ainsi le biais issu de la position défavorise le développement de la sélectivité ;
2. les profils en périphérie plus proches correspondent aux positions atteinte uniquement lorsque le bras est tendu. Dans ce cas, il n'existe qu'une seule posture pour atteindre la position. Ainsi le biais issu de la position favorise le développement de la sélectivité à la posture unique correspondant à la position ;
3. les profils plus proches sont caractérisés par un profil de poids composé de deux zones de poids plus élevés, ce qui implique deux sélectivités distinctes pour une même unité. Ceci est dû au fait que ces positions sont atteintes par deux postures. Ainsi le biais issu

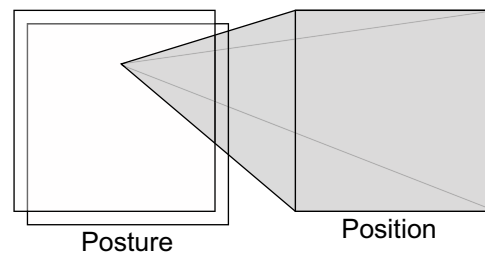


FIGURE 10.3 – Réseau position-posture : les deux cartes reçoivent une activité externe représentant la posture et la position respectivement, et l’activité issue de la carte de position se propage vers la carte de posture par le biais de la projection sur laquelle l’apprentissage a lieu.

de la position favorise le développement de la sélectivité pour aux deux postures correspondant à la position ;

4. enfin, les profils de poids du centre sont marquées par une “bande” verticale de poids plus élevés. Ceci est dû au fait que la position centrale est atteinte lorsque le bras est totalement replié sur lui même, et ce quelque soit l’angle de la première articulation. Ainsi le biais issu de la position favorise le développement de la sélectivité à toute posture permettant d’atteindre la position centrale.

Nous retiendrons de cette expérience que notre algorithme oriente l’apprentissage de manière à ce que la sélectivité sur l’information de posture soit cohérente avec l’information de position présente localement. De plus, si plusieurs postures sont cohérentes avec la position, la sélectivité se développe pour chacune de ces postures.

10.1.2 Relation position - posture

Si on reproduit la même expérience, mais en inversant le sens de la projection, de sorte que les connexions vont de la carte de position vers la carte de posture, chaque unité de la carte de posture développe une sélectivité à la position sur la projection position-posture, sous l’influence de l’information de posture présente localement.

La figure 10.4 montre le résultat de cet apprentissage. Contrairement au cas précédent, la relation qui est apprise ici associe *une* position à *une* posture. Il en résulte que les profils de poids des unités sont composés d’une unique zone de poids plus élevés, donc d’une unique sélectivité.

Il y a cependant une certaine variabilité dans les profils de poids qui, ici aussi, est causée par l’influence sur l’apprentissage de l’information locale - ici une information de posture. On constate que certains profils de poids ont leurs poids forts concentrés en un point, indiquant une sélectivité à une position bien spécifique, tandis que d’autres ont leurs poids forts plus étalés, formant un arc de cercle.

Cette diversité est due à la non-linéarité de la relation position-posture, introduite dans la section 6.2 : suivant l’angle formé par la deuxième articulation (bras étendu ou bras replié), le déplacement de la main causé par une variation d’angle sur la première articulation est très différent (voir figure 10.5).

Ainsi, la bande verticale centrale correspond à un ensemble de postures pour lesquelles le bras est replié. Par conséquent, quelque soit l’angle de la première articulation, la position obtenue

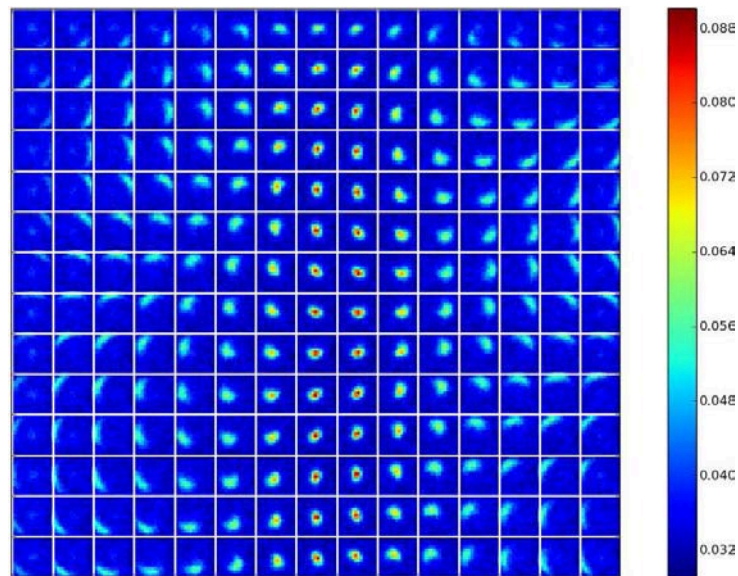


FIGURE 10.4 – Poids de la projection position-posture après apprentissage. L'image est construite en mettant cote à cote les champs récepteurs des 15x15 unités de la carte de posture. Chaque vignette illustre donc la sélectivité à la position acquise par l'unité à cet emplacement de la carte de posture.

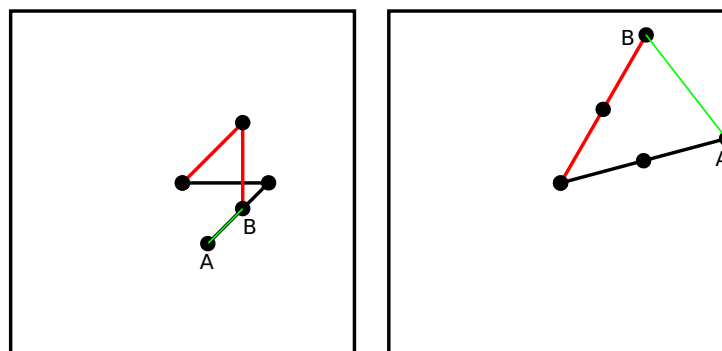


FIGURE 10.5 – Déplacement de la main en fonction du repliement du bras. *A* est la position initiale de la main, *B* la position finale. Pour un même mouvement angulaire de 45 degrés, le déplacement de la main avec le bras partiellement replié (à gauche) est plus court qu'avec le bras totalement déplié (à droite).

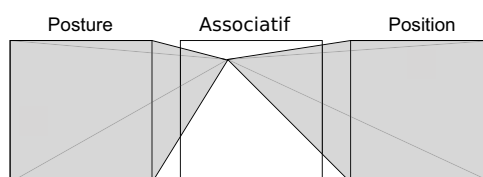


FIGURE 10.6 – Réseau co-apprentissage : les deux cartes de posture et de position reçoivent les activités externes ; l'activité issue de ces deux cartes se propage vers la carte associative par le biais des deux projections sur lesquelles l'apprentissage a lieu. Une modalité latérale inhibitrice est ajoutée à la carte associative pour diversifier les apprentissages.

est la main au centre.

A mesure que l'on s'écarte de cette verticale au centre, le bras est de plus en plus étendu. En conséquence, pour deux stimuli de posture presque identiques (légère variation de la première articulation), les stimuli de position obtenus sont très différents. La modulation issue de l'information de posture tend donc à favoriser le développement de la sélectivité pour un ensemble de positions plus grand - la disposition en arc de cercle étant naturellement causée par la construction du bras.

10.2 Co-apprentissage

Dans une troisième expérience, on projette l'information de posture et l'information de position dans une troisième carte dite associative (figure 10.6). Ainsi, les unités de la carte associative apprennent à partir de deux modalités reçues en parallèle, comme dans le cas des expériences sur la binocularité présentées dans la section 5.4.1. Pour favoriser un apprentissage diversifié au sein de la carte associative, on ajoute une inhibition latérale au sein de la carte associative, afin de mettre en place une compétition entre les unités.

Notre apprentissage multimodal doit faire en sorte que pour chaque unité, les sélectivités apprises sur les deux modalités forment un tout globalement cohérent - en d'autres termes, si une unité est sélective à une certaine position, elle doit être aussi sélective aux postures qui permettent d'obtenir cette position. De plus, notre modèle d'apprentissage multimodal doit résoudre le problème de dominance (voir section 9.1) : l'intensité de la sélectivité (c'est à dire l'amplitude de la réponse maximale) doit être comparable entre les deux modalités.

Soit P un ensemble de postures obtenu en échantillonnant l'espace des postures par angles de 0.1 radians. Chaque posture est un couple d'angles (θ_1, θ_2) , et $|P| = 63^2$.

Pour comparer la sélectivité sur les deux modalités de chaque unité, on calcule pour chaque posture (θ_1, θ_2) les stimuli de posture et de position correspondant, puis la propagation d'activité via les deux projections. En sortie, on obtient pour chaque unité deux *profils de réponses* représentant la sélectivité de l'unité par le biais de la modalité de posture et de la modalité de position. Un profil de réponse indique à quel sous-ensemble de P une unité répond fortement. Si les deux profils de réponses d'une unité sont similaires, alors l'unité a développé la même sélectivité sur les deux modalités.

La figure 10.7 présente quelques exemples de profils de réponses obtenus avec notre modèle d'apprentissage multimodal. On distingue clairement une similarité entre les profils de réponse de posture et de position de chaque unité.

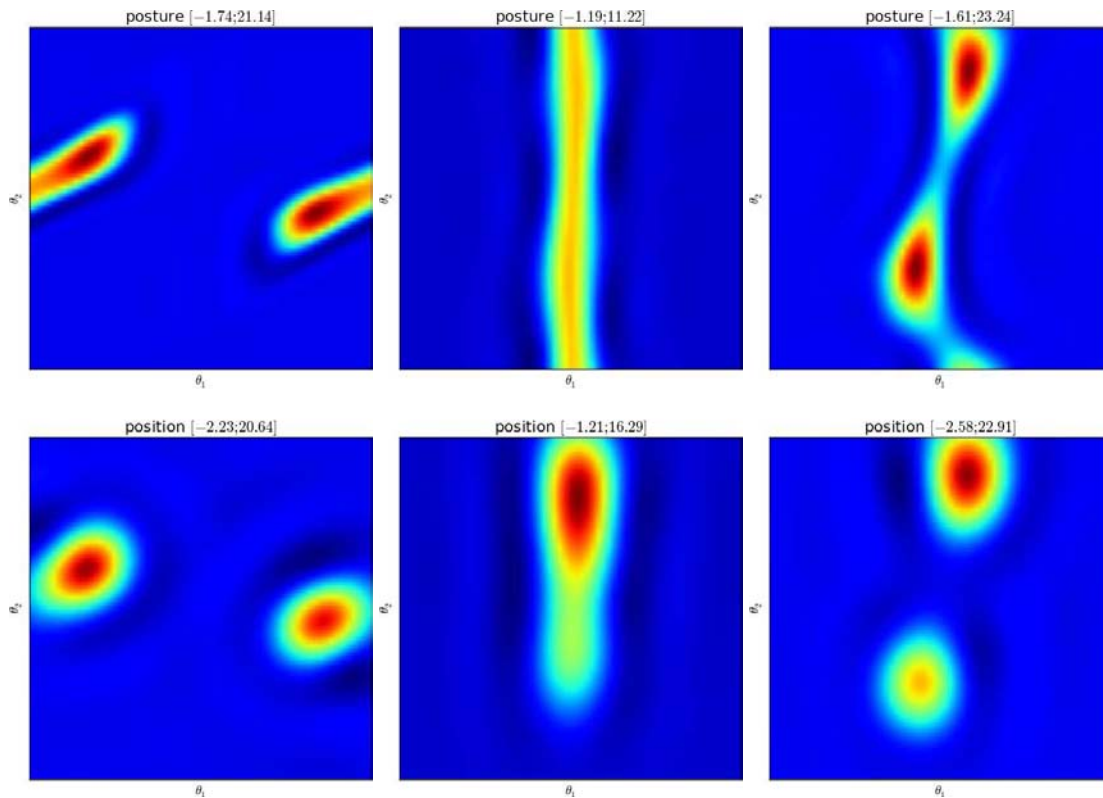


FIGURE 10.7 – Profils de réponse de trois unités, après un apprentissage utilisant notre modèle d'apprentissage multimodal. Les deux axes représentent les angles θ_1 et θ_2 des articulations, formant ainsi l'espace des postures possibles pour le bras. La couleur indique l'intensité de la réponse sur une modalité. Entre crochets, les valeurs minimales et maximales du profil. Chaque colonne correspond à une unité : en haut le profil de réponse de la modalité de posture ; en bas le profil de réponse de la modalité de position.

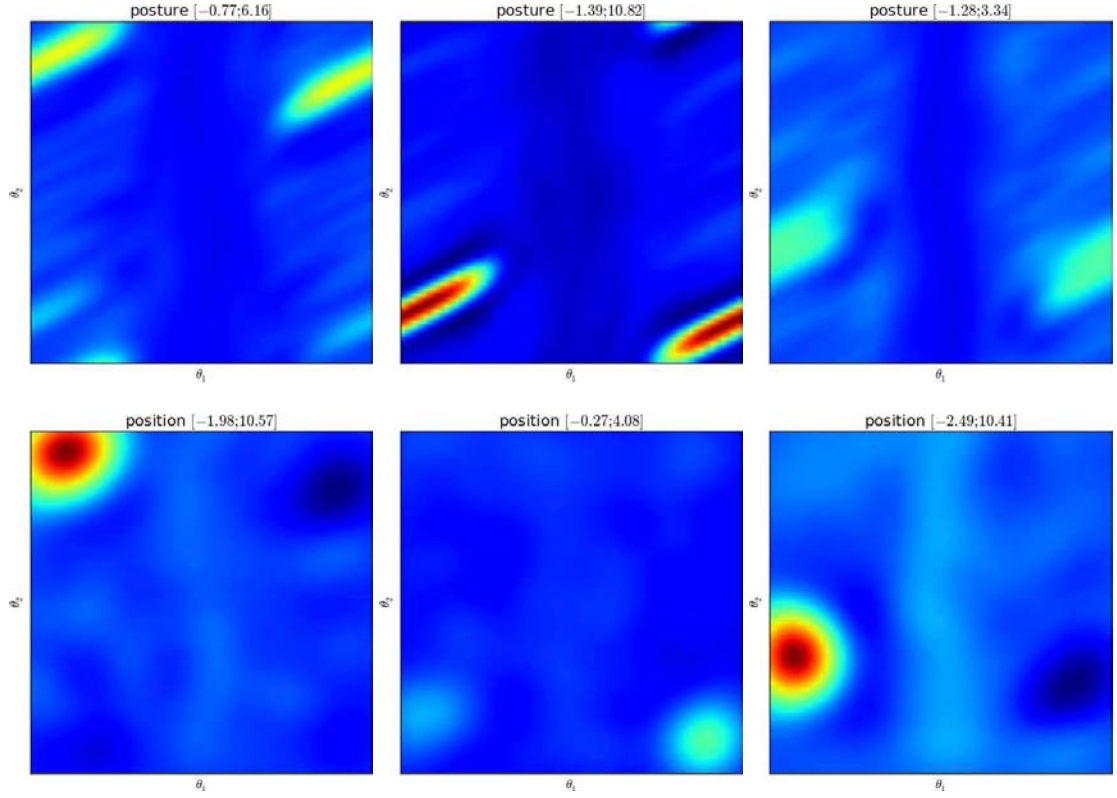


FIGURE 10.8 – Même expérience que la figure 10.7, mais en utilisant une règle BCM classique pour faire un apprentissage multimodal.

Pour faire une comparaison, nous reproduisons la même expérience mais en utilisant le modèle BCM classique dit “binoculaire”, dans lequel les deux modalités ne sont pas différenciées - et donc en compétition :

$$c = \mathbf{w}_{pos} \cdot \mathbf{s}_{pos} + \mathbf{w}_{pst} \cdot \mathbf{s}_{pst} \quad (10.1)$$

$$\dot{\theta} = \frac{1}{\tau}(c^2 - \theta) \quad (10.2)$$

$$(10.3)$$

La figure 10.8 montre quelques exemples de sélectivités obtenues avec ce type d'apprentissage. On constate ici des différences plus fortes entre les profils de réponse de posture et de position de chaque unité.

10.2.1 Mesurer la similarité des sélectivités

Pour pouvoir affirmer que notre modèle permet bien d'obtenir une représentation multimodale plus cohérente - c'est à dire des profils de réponse plus proches sur les deux modalités - on introduit une mesure de la similarité entre les profils de réponse. Pour cela, nous calculons la distance euclidienne entre les profils de réponse centré réduits des deux modalités :

$$r'_X = \frac{r_X - \bar{r}_X}{\sigma_{r_X}} \quad (10.4)$$

r_X est le profil de réponse de la modalité X , $\bar{r}_X$ sa réponse moyenne et σ_{r_X} son écart type. Le calcul de la distance euclidienne entre les profils centrés réduits donne les résultats suivants :

modèle	distance	écart type
BCM classique	5028.79	986.96
BCM multimodal	1239.02	791.05

Ces résultats montrent que, conformément à nos buts, notre modèle d'apprentissage multimodal produit des profils de réponse bien plus similaires que l'apprentissage BCM classique - en moyenne 4 fois plus proches au sens de la métrique utilisée. La cohérence entre les sélectivités développées sur les deux modalités est donc meilleure.

10.3 Discussion

Nous avons présenté un modèle d'apprentissage multimodal basé sur la règle d'apprentissage BCM. Notre modèle se base sur une unité de calcul capable d'intégrer séparément plusieurs modalités et d'effectuer un apprentissage non-supervisé sur chacune de ces modalités, tout en favorisant le développement de sélectivités à des stimuli synchrones.

D'un point de vue cognitif, cela signifie qu'une unité tendra à élaborer des représentations multimodales, pour peu que la synchronie soit un bon indicateur de la relation entre plusieurs perceptions - ce qui est souvent le cas.

De plus, notre modèle ne se limite pas nécessairement à deux modalités - il est possible d'en intégrer un plus grand nombre, construisant ainsi des représentations plus riches. Il est aussi envisageable de grouper un grand nombre de perceptions en représentations de complexité croissante, en construisant une hiérarchie de cartes corticales. L'activité d'une carte devient alors une modalité qu'on intègre à d'autres à un autre niveau de représentation.

Conclusion et perspectives

Conclusion

La problématique initiale de ce travail était de mettre en place un mécanisme d'apprentissage multimodal, basé sur un modèle biologiquement plausible inspiré du cortex cérébral.

Après avoir présenté une description au niveau microscopique, nous avons expliqué pourquoi nous choisissons finalement l'échelle mésoscopique, couramment utilisée pour modéliser le cortex cérébral. Notre modèle consiste en un ensemble d'unités de calcul correspondant à des colonnes corticales, regroupées en cartes bidimensionnelles, interconnectées par des projections de connexions formant trois voies de communications : la voie montante, la voie latérale et la voie descendante. Les projections permettent de différencier les modalités au sein du modèle.

Sur cette base, nous avons développé un modèle d'apprentissage non-supervisé dont la particularité est de favoriser le développement de sélectivités cohérentes entre les différentes modalités (*i.e.* les projections) intégrées par une unité. Ainsi, la co-occurrence répétée de plusieurs stimuli sur des modalités différentes est interprétée comme l'appartenance de ces stimuli à un même phénomène. Les sélectivités acquises par l'unité sur ses modalités entrantes représentent alors différents aspects d'un même phénomène.

Notre modèle d'apprentissage multimodal s'appuie sur l'idée qu'au sein d'un micro-circuit comme la colonne corticale, plusieurs modalités peuvent être intégrées séparément, avec un apprentissage non-supervisé distinct sur chaque modalité. A partir de cette hypothèse, nous avons introduit la notion d'*apprentissage guidé* : l'apprentissage non-supervisé sur une modalité est modulé par un signal externe permettant de biaiser la convergence en faveur (ou en défaveur) de certains stimuli. En basant le signal de modulation sur l'activité totale des modalités, on met en place une influence réciproque entre les apprentissages favorisant l'acquisition de sélectivités à des événements co-occurents.

Outre la capacité d'apprentissage multimodal, notre modèle a l'intérêt d'être simple : la micro-circuiterie de la colonne corticale est réduite à l'intégration séparée de chaque modalité, la mise en commun des activités et l'apprentissage. Les voies montantes, latérales et descendantes sont traitées comme des modalités distinctes, préservant la simplicité et la généralité du modèle.

De plus, il est possible de définir des modalités fixes, qui ne sont pas sujettes à l'apprentissage. Dans ce cas, la modalité fixe agit comme une contrainte sur l'apprentissage des autres modalités, biaisant leurs convergences en faveur d'un ensemble de sélectivités qui serait cohérent avec son activité. Appliqué à la connectivité latérale (intra-carte), ce principe permet d'obtenir des propriétés telles que le codage en population, ou encore l'auto-organisation. Par exemple, avec une connectivité latérale inhibitrice uniforme au sein d'une carte, le meilleur moyen de réduire

les incohérences entre la modalité latérale et les autres modalités est de faire en sorte que toutes les unités soient sélectives à un stimulus différent - *i.e.* ne soient pas actives en même temps.

Limitations

Notre modèle hérite du calcul simple (somme pondérée) et non-borné de l'activité propre à BCM. L'amplitude de la réponse d'une unité dépend directement de la probabilité du stimulus sélectif. Un tel calcul de l'activité peut s'avérer problématique dans le cadre d'un réseau récurrent : des poids trop forts risquent de former des boucles au sein du réseau provoquant l'explosion de l'activité. La règle d'apprentissage ne normalisant pas directement les poids - elle détermine les poids de sorte que $\bar{c} = 1$ - de telles boucles peuvent se former avec l'apprentissage.

Il est possible de contenir ce problème en bornant la sortie des unités, mais l'explosion d'activité serait alors remplacée par un phénomène de saturation. Plutôt que de résoudre artificiellement ce problème, l'adaptation des principes fondamentaux de BCM (seuil θ adaptatif) à un modèle d'activité plus robuste de l'activité nous semblerait préférable.

Perspectives

Notre modèle a introduit un principe d'apprentissage guidé, permettant à plusieurs apprentissages non-supervisés de s'influencer réciproquement par le biais d'une modulation. Dans notre modèle, la modulation provient de l'activité combinée des modalités, mais il serait envisageable d'utiliser une autre source de modulation. La modulation pourrait, par exemple, être contrôlée par un mécanisme dopaminergique, pour modéliser un apprentissage par renforcement basé sur les ganglions de la base.

Un autre point à noter est que notre modèle est capable d'apprentissage multimodal lorsque les différents stimuli sont synchronisés. L'assouplissement de cette contrainte temporelle, de sorte que l'apprentissage multimodal puisse se faire entre des stimuli non synchronisés serait une amélioration certaine. À terme, le but serait d'être à même de proposer un mécanisme d'apprentissage multimodal *spatio-temporel*, plutôt qu'uniquement spatial. Cependant, l'apprentissage non-supervisé spatio-temporel - capable d'apprendre à reconnaître des motifs dans l'espace (formes) et dans le temps (motifs) - est toujours, à notre connaissance, une question ouverte.

Les neurosciences computationnelles forment un domaine de recherche jeune. Beaucoup de questions sont posées, peu de réponses sont apportées. Si l'approche ascendante permet la formation de représentations basées sur l'environnement, le processus reste très coûteux en calculs. Dès lors, il est difficile d'étudier l'émergence de représentations abstraites, par manque de puissance. Nous sommes encore loin de rejoindre le niveau de représentation servant de socle à l'IA symbolique.

Depuis le point de départ de l'IA, la création d'une *intelligence artificielle*, alter ego de l'Homme, la manière dont nous concevons l'intelligence a beaucoup évolué. Si la reproduction de l'intelligence humaine est toujours notre but, l'approche ascendante nous dit que le développement cognitif est intimement lié à notre expérience de l'environnement par le l'intermédiaire de notre corps. En conséquence, rien ne garantit l'émergence d'une intelligence anthropomorphe que

nous pourrions comprendre. Notre but deviendrait alors la compréhension des mécanismes *permettant* notre intelligence.

Bibliographie

- LF Abbott and SB Nelson. Synaptic plasticity : taming the beast. *Nature neuroscience*, 3 :1178–1183, 2000.
- ED Adrian and Y Zotterman. The impulses produced by sensory nerve endings. ii. the response of a single end-organ. *J. Physiol*, 61 :151–171, 1926.
- J. B. Aimone, S. Jessberger, and F. H. Gage. Adult neurogenesis. *Scholarpedia*, 2(2) :2100, 2007.
- N Almassy, GM Edelman, and O Sporns. Behavioral constraints in the development of neuronal properties : a cortical model embedded in a real-world device. *Cerebral Cortex*, 8(4) :346, 1998. ISSN 1047-3211.
- S Amari. Dynamics of pattern formation in lateral-inhibition type neural fields. *Biological Cybernetics*, 27(2) :77–87, 1977.
- B Arons. A review of the cocktail party effect. *Journal of the American Voice I/O Society*, 12(7) : 35–50, 1992.
- HB Barlow. Possible principles underlying the transformations of sensory messages. *Sensory communication, contributions*, page 217, 1961.
- HB Barlow. Unsupervised learning. *Neural Computation*, 1(3) :295–311, 1989.
- MF Bear. A synaptic basis for memory storage in the cerebral cortex. *Proceedings of the National Academy of Sciences*, 93(24) :13453–13459, 1996.
- RE Bellman. *Adaptive control processes - A guided tour*. Princeton University Press, Princeton, New Jersey, U.S.A., 1961.
- G Benedek, J Perenyi, G Kovacs, L Fischer-Szatmari, and YY Katoh. Visual, somatosensory, auditory and nociceptive modality properties in the feline supragenulate nucleus. *Neuroscience*, 78(1) :179–189, 1997.
- G Bi and M Poo. Synaptic modifications in cultured hippocampal neurons : dependence on spike timing, synaptic strength, and postsynaptic cell type. *Journal of Neuroscience*, 18(24) : 10464, 1998.
- EL Bienenstock, LN Cooper, and PW Munro. Theory for the development of neuron selectivity : Orientation specificity and binocular interaction in visual cortex. *Journal of Neuroscience*, 2, 1982.

- BS Blais. *The Role of the Environment in Synaptic Plasticity : Towards an Understanding of Learning and Memory*. PhD thesis, Brown University, 1998.
- GG Blasdel and G Salama. Voltage-sensitive dyes reveal a modular organization in monkey striate cortex. *Nature*, 321(6070) :579–585, 1986.
- TVP Bliss and T Lomo. Synaptic plasticity in the hippocampus. *Journal of Physiology*, 1973.
- K Brodmann. Beiträge zur histologischen lokalisation der grosshirnrinde : dritte mitteilung : Die rindenfelder der niederen affen. *Journal für Psychologie und Neurologie*, 4 :177–226, 1909.
- P Buisseret and M Imbert. Visual cortical cells : their developmental properties in normal and dark reared kittens. *The Journal of Physiology*, 255(2) :511, 1976.
- Y Burnod. *An adaptive neural network : the cerebral cortex*. Masson editeur, Paris, France, France, 1989. ISBN 0-13-019464-6.
- G.A. Carpenter and S. Grossberg. Adaptive resonance theory. *The handbook of brain theory and neural networks*, 2 :87–90, 2003.
- GC Castellani, N Intrator, H Shouval, and LN Cooper. Solutions of the bcm learning rule in a network of lateral interacting nonlinear neurons. *Network : Computation in Neural Systems*, 10 (2) :111–121, 1999.
- VF Castellucci, TJ Carew, and ER Kandel. Cellular analysis of long-term habituation of the gill-withdrawal reflex of aplysia californica. *Science*, 202(4374) :1306–1308, 1978.
- EH Chudler, K Sugiyama, and WK Dong. Multisensory convergence and integration in the neostriatum and globus pallidus of the rat. *Brain research*, 674(1) :33–45, 1995.
- LN Cooper, F Liberman, and E Oja. A theory for the acquisition and loss of neuron specificity in visual cortex. *Biological Cybernetics*, 1979.
- LN Cooper, N Intrator, BS Blais, and HZ Shouval. *Theory of cortical plasticity*. World Scientific, 2004.
- JA Cummings, RM Mulkey, RA Nicoll, and RC Malenka. Ca²⁺ signaling requirements for long-term depression in the hippocampus. *Neuron*, 16(4) :825–833, 1996.
- NS Desai, LC Rutherford, and GG Turrigiano. Plasticity in the intrinsic excitability of cortical pyramidal neurons. *Nature Neuroscience*, 2(6), 1999.
- P Diaconis and D Freedman. Asymptotics of graphical projection pursuit. *The Annals of Statistics*, pages 793–815, 1984.
- K. Doya. What are the computations of the cerebellum, the basal ganglia and the cerebral cortex ? *Neural networks*, 12(7-8) :961–974, 1999.
- PS Eriksson, E Perfilieva, T Bjork-Eriksson, AM Alborn, C Nordborg, DA Peterson, and FH Gage. Neurogenesis in the adult human hippocampus. *Nature Medecine*, 4(11) :1313–1317, 1998.
- R Fitzhugh. Impulses and physiological states in theoretical models of nerve membrane. *Biophysical Journal*, 1(6) :445–466, 1961.

-
- P Foldiak and D Endres. Sparse coding. *Scholarpedia*, 3(1) :2984, 2008.
- J. Gautrais and S. Thorpe. Rate coding versus temporal order coding : a theoretical approach. *Biosystems*, 48(1-3) :57–65, 1998.
- W Gerstner and WM Kistler. *Spiking neuron models : single neurons, populations, plasticity*. Cambridge university press, 2002.
- A.A. Ghazanfar and C.E. Schroeder. Is neocortex essentially multisensory ? *Trends in Cognitive Sciences*, 10(6) :278–285, 2006.
- J.J. Gibson. The theory of affordances. *Perceiving, acting, and knowing : Toward an ecological psychology*, pages 67–82, 1977.
- MA Giese. *Dynamic neural field theory for motion perception*. Kluwer Academic Publishers Norwell, MA, USA, 1998.
- T Girod and F Alexandre. Mécanisme d’auto-organisation corticale : un modèle basé sur la règle de plasticité synaptique bcm. In *Deuxième conférence française de Neurosciences Computationnelles, "Neurocomp08"*, 2008.
- T Girod and F Alexandre. Effects of a modulatory feedback upon the bcm learning rule. In *CNS*, 2009.
- T Girod, L Bougrain, and F Alexandre. Toward a robust 2d spatio-temporal self-organization. In *ESANN*, 2007.
- M.A. Goodale and A.D. Milner. Separate visual pathways for perception and action. *Trends in neurosciences*, 15(1) :20–25, 1992.
- S Grossberg. Towards a unified theory of the neocortex : Laminar cortical circuit for vision and cognition. Technical report, Boston University, 2007.
- E Guigon. *Modélisation des propriétés du cortex cérébral : comparaison entre les aires visuelles, motrices et préfrontales*. PhD thesis, Ecole centrale des arts et manufactures, 1993.
- P Hagmann, L Cammoun, X Gigandet, R Meuli, CJ Honey, and VJ Wedeen. Mapping the structural core of human cerebral cortex. *PLOS Biology*, 6(7) :1479–1493, 2008.
- S. Harnad. The symbol grounding problem. *Physica D : Nonlinear Phenomena*, 42(1-3) :335–346, 1990.
- DO Hebb. *The organization of behaviour ; a neuropsychological theory*. Wiley, New York, 1949.
- J Heller, JA Hertz, TW Kjaer, and BJ Richmond. Information flow and temporal coding in primate pattern vision. *Journal of Computational Neuroscience*, 2(3) :175–193, 1995.
- B Hellwig. A quantitative analysis of the local connectivity between pyramidal neurons in layers 2/3 of the rat visual cortex. *Biological Cybernetics*, 82(2) :111–121, 2000.
- GH Henry, B. Dreher, and PO Bishop. Orientation specificity of cells in cat striate cortex. *Journal of Neurophysiology*, 37(6) :1394, 1974.
- S Hill and H Markram. The blue brain project. In *Engineering in Medicine and Biology Society, 2008. EMBS 2008. 30th Annual International Conference of the IEEE*. IEEE, 2008.

- A Hodgkin and A Huxley. A quantitative description of membrane current and its application to conduction and excitation in nerve. *Journal of Physiology*, 117 :500–544, 1952.
- K Hornik, S Maxwell, and H White. Multilayer feedforward networks are universal approximators. *Neural networks*, 2(5) :359–366, 1989.
- JC Horton and DL Adams. The cortical column : a structure without a function. *Philosophical Transactions of the Royal Society B : Biological Sciences*, 360(1456) :837, 2005.
- DH Hubel and TN Wiesel. Receptive fields, binocular interaction and functional architecture in the cat's visual cortex. *The Journal of Physiology*, 160(1) :106, 1962.
- DH Hubel and TN Wiesel. Receptive fields and functional architecture of monkey striate cortex. *Journal of Physiology*, 195(1) :215–243, 1968.
- N Intrator. A neural network for feature extraction. Technical report, Brown University, 1990.
- N. Intrator and L.N. Cooper. Objective function formulation of the bcm theory of visual cortical plasticity : Statistical connections, stability conditions. *Neural Networks*, 5(1) :3–17, 1992.
- EM Izhikevich and NS Desai. Relating stdp to bcm. *Neural Computation*, 15(7) :1511–1523, 2003.
- ER Kandel, J Schwartz, and T Jessell. *Principles of Neural Science*. McGraw-Hill Education, 2000.
- A Kirkwood, MG Rioult, and MF Bear. Experience-dependent modification of synaptic plasticity in visual cortex. *Nature*, 381 :526–528, June 1996.
- DC Knill and W Richards. *Perception as Bayesian inference*. Cambridge Univ Pr, 1996.
- T Kohonen. Self-organized formation of topologically correct feature maps. *Biological Cybernetics*, 1982.
- J Konorski. *Conditioned Reflexes and Neuron Organization*. Cambridge University Press, 1948.
- R Lamprecht and J LeDoux. Structural plasticity and memory. *Nature Reviews Neuroscience*, 5 : 45–54, 2004.
- C C Law and L N Cooper. Formation of receptive fields in realistic visual environments according to the bienenstock, cooper, and munro (bcm) theory. *Proc Natl Acad Sci U S A*, 91(16) : 7797–801, August 1994. ISSN 0027-8424.
- DJ Linden and JA Connor. Long-term synaptic depression. *Annual Review of Neuroscience*, 18 (1) :319–357, 1995.
- R Lorente de Nó and JF Fulton. *Physiology of the nervous system*, chapter 15 - Cerebral cortex : architecture, intracortical connections, motor projections. Oxford University Press, 3 edition, 1949.
- H. Markram. The blue brain project. *Nature Reviews Neuroscience*, 7(2) :153–160, 2006.
- H Markram, J Lubke, M Frotscher, and B Sakmann. Regulation of synaptic efficacy by coincidence of postsynaptic aps and epsps. *Science*, 275(5297) :213, 1997.
- WS McCulloch and W Pitts. A logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematical Biophysics*, 5 :115–133, 1943.

-
- O Menard. *Mécanismes d'inspiration corticale pour l'apprentissage et la représentation d'asservissements sensori-moteurs en robotique*. PhD thesis, Université Henri-Poincaré, Nancy 1, 2005.
- KD Miller and DJC MacKay. The role of constraints in hebbian learning. *Neural Computation*, 6(1) :100–126, 1994.
- C Morris and H Lecar. Voltage oscillations in the barnacle giant muscle fiber. *Biophysical journal*, 35(1) :193–213, 1981.
- VB Mountcastle. Modality and topographic properties of single neurons of cat's somatic sensory cortex. *Journal of Neurophysiology*, 1957.
- J Nagumo, S Arimoto, and S Yoshizawa. An active pulse transmission line simulating nerve axon. *Proceedings of the IRE*, 50 :2061–2070, 1962.
- MM Nass and LN Cooper. A theory for the development of feature detecting cells in visual cortex. *Biological cybernetics*, 19(1) :1–18, 1975.
- A. Newell and H.A. Simon. *Human problem solving*. Prentice-Hall Englewood Cliffs, NJ, 1972.
- S N'Guyen. *Mise au point du système vibrissal du robot-rat Psikharpx et contribution à la fusion de ses capacités visuelle, auditive et tactile*. PhD thesis, Université Pierre et Marie Curie, 2010.
- H Nishijo, T Ono, and H Nishino. Topographic distribution of modality-specific amygdalar neurons in alert monkey. *Journal of Neuroscience*, 8(10) :3556, 1988.
- E Oja. A simplified neuron model as a principal component analyzer. *Journal of Mathematical Biology*, 1982.
- LM Optican and BJ Richmond. Temporal encoding of two-dimensional patterns by single units in primate inferior temporal cortex. iii. information theoretic analysis. *Journal of Neurophysiology*, 57(1) :162, 1987.
- CMA Pennartz. Reinforcement learning by hebbian synapses with adaptive thresholds. *Neuroscience*, 81(2) :303–319, 1997.
- S Ramon y Cajal. *Histologie du système nerveux de l'homme et des vertébrés*. Maloine, Paris, pages 774–838, 1909.
- E Reynaud. *Modélisation connexionniste d'une mémoire associative multimodale*. PhD thesis, Institut National Polytechnique de Grenoble, 2002.
- BA Reynolds and S Weiss. Generation of neurons and astrocytes from isolated cells of the adult mammalian central nervous system. *Science*, 255(5052) :1707–1710, 1992.
- BJ Richmond, LM Optican, M Podell, and H Spitzer. Temporal encoding of two-dimensional patterns by single units in primate inferior temporal cortex. i. response characteristics. *Journal of Neurophysiology*, 57(1) :132, 1987.
- BJ Richmond, LM Optican, and H Spitzer. Temporal encoding of two-dimensional patterns by single units in primate primary visual cortex. i. stimulus-response relations. *Journal of Neurophysiology*, 64(2) :351, 1990.

- F. Rosenblatt. The perceptron : A probabilistic model for information storage and organization in the brain. *Psychological review*, 65(6) :386–408, 1958.
- NP Rougier. Dynamic neural field with local inhibition. *Biological cybernetics*, 94(3) :169–179, 2006.
- NP Rougier and J Vitay. Emergence of attention within a neural population. *Neural Networks*, 19(5) :573–581, 2006.
- E Salinas and P Thier. Gain modulation : A major computational principle of the central nervous system. *Neuron*, 27(1) :15–21, July 2000.
- W. Senn, H. Markram, and M. Tsodyks. An algorithm for modifying neurotransmitter release probability based on pre-and postsynaptic spike timing. *Neural Computation*, 13(1) :35–67, 2001.
- HZ Shouval, DH Goldberg, JP Jones, M Beckerman, and LN Cooper. Structured long-range connections can provide a scaffold for orientation maps. *Journal of Neuroscience*, 20(3) :1119, 2000.
- HZ Shouval, MF Bear, and LN Cooper. A unified model of nmda receptor-dependent bidirectional synaptic plasticity. *Proceedings of the National Academy of Science*, 99 :10831–10836, August 2002.
- BE Stein and MA Meredith. *The merging of the senses*. MIT Press, 1993.
- DD Stettler, A Das, J Bennett, and CD Gilbert. Lateral connectivity and contextual interactions in macaque primary visual cortex. *Neuron*, 36(4) :739–750, 2002.
- RS Sutton and AG Barto. *Reinforcement learning : An introduction*. The MIT press, 1998.
- LW Swanson. Mapping the human brain : past, present, and future. *Trends in neurosciences*, 18 (11) :471–474, 1995.
- JG Taylor. Neural ‘bubble’ dynamics in two dimensions : foundations. *Biological Cybernetics*, 80 (6) :393–409, 1999.
- AM Thomson and C Lamy. Functional maps of neocortical local circuitry. *Frontiers in neuroscience*, 1(1) :19–42, 2007.
- S. Thorpe, D. Fize, and C. Marlot. Speed of processing in the human visual system. *nature*, 381 (6582) :520–522, 1996.
- T Toyozumi, JP Pfister, K Aihara, and W Gerstner. Generalized bienenstock–cooper–munro rule for spiking neurons that maximizes information transmission. *Proceedings of the National Academy of Sciences of the United States of America*, 102(14) :5239, 2005.
- TP Trappenberg. *Fundamentals of Computational Neuroscience*. Oxford University Press, 2 edition, 2010.
- A.M. Turing. On computable numbers, with an application to the entscheidungsproblem. *Proceedings of the London Mathematical Society*, 2(1) :230, 1937.
- AM Turing. Computing machinery and intelligence. *Mind*, 59(236) :433–460, 1950.

-
- GG Turrigiano and SB Nelson. Homeostatic plasticity in the developing nervous system. *Nature Neuroscience*, 5, 2004.
- LG Ungerleider, M Mishkin, et al. Two cortical visual systems. *Analysis of visual behavior*, 549 : 586, 1982.
- C von der Malsburg. Self-organization of orientation sensitive cells in the striate cortex. *Kybernetik*, 14(2) :85–100, 1973.
- JC Wiemer. *Learning topography in neural networks : towards a better understanding of cortical topography*. PhD thesis, Bochum university, 2000.
- TN Wiesel and DH Hubel. Comparison of the effects of unilateral and bilateral eye closure on cortical unit responses in kittens. *Journal of Neurophysiology*, 28(6) :1029, 1965.
- HR Wilson and JD Cowan. Excitatory and inhibitory interactions in localized populations of model neurons. *Biophysical journal*, 12, 1972.
- F. Woergoetter and B. Porr. Reinforcement learning. *Scholarpedia*, 3(3) :1448, 2008.

Résumé

Le domaine des neurosciences computationnelles s'intéresse à la modélisation des fonctions cognitives à travers des modèles numériques bio-inspirés. Dans cette thèse, nous nous intéressons en particulier à l'apprentissage dans un contexte multimodal, c'est à dire à la formation de représentations cohérentes à partir de plusieurs modalités sensorielles et/ou motrices.

Notre modèle s'inspire du cortex cérébral, lieu supposé de la fusion multimodale dans le cerveau, et le représente à une échelle mésoscopique par des colonnes corticales regroupées en cartes et des projections axoniques entre ces cartes. Pour effectuer nos simulations, nous proposons une bibliothèque simplifiant la construction et l'évaluation de modèles mésoscopiques.

Notre modèle d'apprentissage se base sur le modèle BCM (Bienenstock-Cooper-Munro), qui propose un algorithme d'apprentissage non-supervisé local (une unité apprend à partir de ses entrées de manière autonome) et biologiquement plausible. Nous adaptons BCM en introduisant la notion d'apprentissage guidé, un moyen de biaiser la convergence de l'apprentissage BCM en faveur d'un stimulus choisi. Puis, nous mettons ce mécanisme à profit pour effectuer un co-apprentissage entre plusieurs modalités. Grâce au co-apprentissage, les sélectivités développées sur chaque modalité tendent à représenter le même phénomène, perçu à travers différentes modalités, élaborant ainsi une représentation multimodale cohérente dudit phénomène.

Mots-clés: neurosciences computationnelles, multimodalité, apprentissage guidé, BCM.

Abstract

The field of computational neuroscience is interested in modeling the cognitive functions through biologically-inspired, numerical models. In this thesis, we focus on learning in a multimodal context, ie the combination of several sensitive/motor modalities.

Our model draws from the cerebral cortex, supposedly linked to multimodal integration in the brain, and modelize it on a mesoscopic scale with 2d maps of cortical columns and axonic projections between maps. To build our simulations, we propose a library to simplify the construction and evaluation of mesoscopic models.

Our learning model is based on the BCM model (Bienenstock-Cooper-Munro), which offers a local, unsupervised, biologically plausible learning algorithm (one unit learns autonomously from its entries). We adapt this algorithm by introducing the notion of guided learning, a mean to bias the convergence to the benefit of a chosen stimuli. Then, we use this mechanism to establish correlated learning between several modalities. Thanks to correlated learning, the selectivities acquired tend to account for the same phenomenon, perceived through different modalities. This is the basis for a coherent, multimodal representation of this phenomenon.

Keywords: computational neuroscience, multimodality, unsupervised learning, BCM.

