

Université Louis Lumière Lyon 2
THÈSE de Doctorat de Sciences Cognitives -Mention Neuropsychologie
Présentée et soutenue publiquement le 24 Septembre 2004
Par Nadia BEN BOUTAYAB

INTERACTIONS DES FACTEURS VISUELS ET LEXICAUX AU COURS DE LA RECONNAISSANCE DES MOTS ÉCRITS :

Réalisée sous la direction du Dr. Tatjana A. NAZIR
à L'Institut des Sciences Cognitives *Directeur: Pr. Yves BURNOD CNRS UMR 9075*

JURY: Pr Marc BRYSSBAERT, Professeur, University of London (United Kingdom) Pr Ram FROST, Professeur, Hebrew University (Israël) Dr Jonathan GRAINGER, Directeur de recherche CNRS, Université Aix en Provence Dr Tatjana NAZIR, Institut des Sciences Cognitives, Lyon, directeur de la thèse

Table des matières

Introduction . .	1
Chapitre 1 Reconnaissance visuelle des mots et modèles de lecture . .	3
1-1 Identité des lettres et codage de la position .	3
1-1-1 Codage absolu de la position des lettres .	4
1-1-2 Codage relatif de la position des lettres . .	4
1-2 Accès au lexique, sélection et voisinage orthographique . .	5
1-2-1 Voisinage orthographique : taille versus fréquence du voisinage .	5
1-2-2 Voisinage orthographique : faits expérimentaux .	6
1-3 Modèles de la reconnaissance visuelle des mots . .	7
1-3-1 Le modèle à activation vérification (AV) . .	7
1-3-2 Le modèle à activation interactive (AI) . .	8
1-3-3 Stratégies de lecture et décision lexicale . .	9
Chapitre 2 Le phénomène de l'effet de la position du regard dans les mots .	11
2-1 Position du regard dans les mots et efficacité de la reconnaissance .	11
2-1-1 Enregistrements des mouvements oculaires .	11
2-1-2 Le paradigme de la position variable du regard . .	12
2-1-3 Robustesse du phénomène .	13
2-2 Asymétrie de l'EPR : plusieurs interprétations .	14
2-2-1 Asymétrie de la lisibilité des lettres .	14
2-2-2 Spécialisation hémisphérique . .	15
2-2-3 Facteurs attentionnels, direction de lecture et apprentissage perceptif . .	15
2-2-4 Contraintes lexicales et structure informative des mots . .	16
Chapitre 3 Contraintes lexicales et Lisibilité des lettres . .	19
3-1 Contrainte orthographique .	19
3-1-1 L'effet de supériorité du début du mot .	19
3-1-2 Contrainte lexicale et position du regard dans le mot .	20

3-2 Structure morphologique . .	21
3-2-1 Notion de structure morphologique .	21
3-2-2 Morphologie et position du regard dans le mot .	22
3-3 Lisibilité des lettres et position du regard dans le mot .	24
3-3-1 Lisibilité des caractères romans . .	24
3-3-2 Asymétrie de la lisibilité des lettres : dominance hémisphérique et direction de lecture. . .	25
3-4 Un modèle mathématique « visuel » (Nazir et al, 1991 ; 1998) .	26
3-4-1 Probabilité de reconnaissance d'un mot .	26
3-4-2 Limites du modèle MPLI . .	27
Chapitre 4 Effet de la position du regard et modèles mathématiques . .	29
4-1 Le modèle de Clark et O'Regan (1999) .	29
4-1-1 Estimation de l'information visuelle .	29
4-1-2 Mesure d'ambiguïté du mot .	30
4-2 Le modèle de Stevens et Grainger (2003) .	30
4-2-1 Données empiriques de lisibilité des lettres . .	30
4-2-2 Fréquence positionnelle des lettres .	31
4-3 Une proposition de Kajii (2000) . .	32
4-3-1 Étapes menant à la reconnaissance d'un mot. .	32
4-3-2 Probabilité d'identifier un mot. .	32
4-3-3 Codage de l'identité et la position des lettres. .	33
4-4 Le modèle CLIP .	33
4-4-1 Probabilité d'occurrence d'un CLIP (C_i) . .	34
4-4-2 Probabilité d'identifier une lettre (L_i) . .	34
4-4-3 Contrainte lexicale et probabilité d'identifier un mot ($P(W C_i)$) . .	35
4-5 Évaluation du modèle CLIP (Kajii & Osaka, 2000) .	35
4-5-1 Probabilité d'occurrence des CLIPs . .	35
4-5-2 Mesure de la contrainte lexicale . .	36
4-5-3 Reconnaissance du mot et position du regard . .	36

4-5-4 Objectifs de la contribution expérimentale .	38
Chapitre 5 EPR, lettres, position et lexique .	41
5-1 Identification perceptive de mots et de pseudomots : Expérience 1 . .	41
5-1-1 Méthode .	42
5-1-2 Résultats et discussion . .	43
5-2 Simulations des données de l'expérience 1 .	44
5-2-1 Estimation du paramètre visuel .	45
5-2-2 Probabilité d'occurrence des CLIPs . .	45
5-2-3 Simulation 1 : Nombre de « voisins CLIP » de même longueur que le stimulus .	46
5-3 Analyses supplémentaires .	48
5-3-1 Identification des lettres et codage de la position . .	48
5-3-2 Analyse qualitative des erreurs .	49
5-3-3 Discussion .	50
5-4 Simulations : Nombre de « voisins CLIP » de plus haute fréquence .	51
5-4-1 Simulation 2 : « de 3 à 10 lettres » . .	51
5-4-2 Simulation 3 : « de 5 à 7 lettres » . .	52
5-4-3 Simulation 4 : Erreurs d'inférence lexicale . .	52
Chapitre 6 EPR, stratégies et tâches . .	57
6-1 Identification perceptive et Décision lexicale : Expérience 2 .	57
6-1-1 Méthode .	57
6-1-2 Résultats 1 : Pourcentages de bonnes réponses . .	58
6-1-3 Résultats 2 : Analyse qualitative des erreurs dans la tâche d'identification perceptive .	60
6-2 Comparaison Identification/Décision lexicale : données empiriques .	61
6-2-1 Comparaison des données moyennées à travers les cinq positions du regard .	61
6-2-2 Comparaison des données en fonction de la position du regard dans le stimulus .	61
6-3 Comparaison Identification/Décision lexicale : Simulations via le modèle CLIP . .	62
6-3-1 Estimation du paramètre visuel (br) .	62

6-3-2 Probabilités d'occurrence des CLIPs .	62
6-3-3 Simulation 5 : Identifications correctes des mots et des pseudomots .	63
6-3-4 Simulation 6 : Erreurs d'inférence lexicale . .	63
6-3-5 Simulation 7 : Décisions lexicales correctes pour les mots et les pseudomots .	64
Chapitre 7 EPR, asymétrie et familiarité visuelle .	67
7-1 Identification perceptive et format horizontal / vertical : Expérience 3 .	68
7-1-1 Méthode .	68
7-1-2 Résultats 1 : Identification des mots et des pseudomots . .	69
7-1-3 Résultats 2 : Identification des lettres individuelles .	70
7-1-4 Résultats 3 : Analyse qualitative des erreurs . .	71
7-2 Lisibilité des lettres et format de présentation : Expérience 4 . .	72
7-2-1 Méthode .	72
7-2-2 Résultats et discussion . .	73
7-2-3 Simulation des données de l'expérience 3 . .	75
7-3 Décision lexicale et format horizontal / vertical : Expérience 5 . .	76
7-3-1 Méthode .	76
7-3-2 Résultats .	76
Chapitre 8 Discussion générale et Conclusion .	81
8-1 Résumé . .	81
8-2 Limites du modèle . .	82
8-2-1 Choix des CLIPs . .	82
8-2-2 Champ du « voisin CLIP » .	83
8-2-3 CLIP-total ou identification perceptive totale .	84
8-3 Le modèle CLIP et les autres modèles visuo-lexicaux .	85
8-4 Acquisition de la lecture .	85
8-5 Déficits de lecture acquis .	86
Références .	87
Annexes . .	97

Annexe 1 . .	97
Annexe 2 : Modification du modèle de Nazir et al. (1991 ; 1998) .	98
Annexe 3 : : stimuli de l'expérience 1 . .	99
Annexe 4 : stimuli de l'expérience 2 .	99
Annexe 5 : stimuli de l'expérience 3 .	99
Annexe 6 : Expérience pilote x .	99

Introduction

Lorsque la lecture est correctement acquise, les phrases d'un texte sont comprises sans difficulté et de manière continue. Le regard saute d'un point à un autre des lignes, permettant ainsi d'extraire l'information visuelle critique qui va conduire à l'identification des mots.

Les modèles actuels de la lecture silencieuse considèrent généralement que la présentation d'un mot écrit provoque d'abord l'activation d'un ensemble limité de candidats lexicaux formellement proches du stimulus, puis la sélection d'un seul d'entre eux (e.g., McClelland & Rumelhart, 1981 ; Paap, Newsome, McDonald & Schvaneveldt, 1982 ; Segui, 1991). Autrement dit, lire, c'est récupérer parmi des milliers de mots, la représentation du mot stockée dans le lexique mental.

Pour autant, notre manière d'extraire l'information visuelle influence directement l'efficacité de la reconnaissance des mots. En particulier, la facilité avec laquelle un lecteur expert reconnaît un mot présenté brièvement dépend de la localisation de son regard dans le mot (e.g. O'Regan, Lévy-Schoen, Pynte, & Brugailière, 1984). La reconnaissance du mot est maximisée lorsque le regard se situe légèrement à gauche du centre et, diminue à mesure qu'on s'éloigne de cette position optimale, vers le début ou vers la fin du mot (phénomène de l'effet de la position du regard).

Dans ce contexte, la reconnaissance d'un mot écrit ne peut être envisagée du seul point de vue lexical et un rôle prépondérant doit être attribué aux facteurs visuels. L'objectif de ce travail est de mieux déterminer comment les *facteurs lexicaux* et les *facteurs visuels* interagissent au cours de la reconnaissance des mots écrits. Cette

interrogation sera abordée par le biais de l'investigation du phénomène de l'effet de la position du regard au cours de la lecture.

Nous allons dans un premier temps donner un aperçu de la reconnaissance visuelle des mots selon les modèles actuels de la lecture. Puis, nous nous pencherons plus précisément sur le phénomène de l'effet de la position du regard dans le mot en exposant les différentes théories qui lui sont attribuées, en particulier celles ciblant le point de vue visuel et le point de vue lexical. Une réelle avancée vers une tentative de réconciliation de ces deux théories a été réalisée grâce au développement de modèles dits « visuo-lexicaux », lesquels permettront d'avancer un certain nombre d'arguments concernant l'origine de l'effet de la position du regard dans le mot.

Chapitre 1 Reconnaissance visuelle des mots et modèles de lecture

1-1 Identité des lettres et codage de la position

Selon une conception largement admise aujourd'hui, la reconnaissance visuelle des mots dans les écritures alphabétiques passe par le traitement des lettres qui les composent (e.g., Bowers, 2000 ; Evett & Humphreys, 1981 ; Grainger & Jacobs, 1996 ; McClelland & Rumelhart, 1981 ; Paap, Newsome, McDonald & Schvaneveldt, 1982 ; Rayner & Pollatsek, 1989), bien que certains auteurs continuent de défendre l'idée d'une reconnaissance holistique des mots dans la lecture (e.g., Allen, Wallace & Weber, 1995 ; Cattell, 1886 ; Healy & Cunningham, 1992). On considère en général que le lecteur extrait une *représentation abstraite* de l'identité de chaque lettre, indépendante de la casse, de la taille ou de la typographie (Adams, 1979 ; Besner, 1989 ; McClelland & Rumelhart, 1981 ; Paap et al, 1982 ; Paap, Newsome & Noel, 1984 ; Rayner, McConkie & Zola, 1980).

Toutefois, au cours de la mise en correspondance des lettres avec la représentation orthographique stockée dans le lexique mental, l'information concernant la *position* des lettres doit également être codée. L'importance du codage de la position est évident dans les langues alphabétiques comme le français ou l'anglais qui possèdent de nombreux

mots anagrammes. En effet, nous sommes capables de distinguer des mots différents contenant les mêmes lettres comme « crier et cirer » via le codage de la position des lettres. De même, nous sommes capables de reconnaître que « crire » n'est pas un mot de la langue française. Toutefois, la question concernant la manière d'encoder la position des lettres n'est pas encore complètement résolue.

1-1-1 Codage absolu de la position des lettres

Selon une première hypothèse, l'identité abstraite de la lettre ainsi que sa position dans la séquence sont codées en même temps, de telle sorte que l'unité de codage ne serait pas la lettre elle-même, mais la lettre à une position donnée (McClelland & Rumelhart, 1981, voir aussi Paap et al., 1982). Par exemple, dans le modèle à activation interactive (McClelland & Rumelhart, 1981), chaque lettre est traitée par un canal indépendant qui extrait et analyse les traits visuels de la lettre dans une position spécifique. Ainsi, il existerait des canaux séparés représentant le « a » en première position, le « a » en deuxième position et ainsi de suite. Selon ce modèle, l'information concernant la position des lettres est codée automatiquement et, une lettre donnée est toujours identifiée à une position spécifique dans le mot ; cette lettre va ensuite activer tous les mots contenant cette lettre à cette position spécifique. Par exemple, lorsque l'unité correspondant à la lettre « t » en position initiale s'active, elle va augmenter le niveau d'activation de toutes les unités mots possédant la lettre « t » en première position (par exemple, en anglais « TAKE », « TIME », « TRIP », « TRAP ») et, va au contraire diminuer le niveau d'activation des unités mots ne possédant pas la lettre « t » en position initiale (e.g., « ABLE », « CART »). Un tel codage de chaque lettre selon une position spécifique dans la séquence est toutefois peu économique puisqu'il faut postuler l'existence d'un alphabet mental pour chacune des positions du mot.

1-1-2 Codage relatif de la position des lettres

Des arguments forts à l'encontre d'un codage absolu de la position des lettres au cours de la perception des mots écrits ont été obtenus via le paradigme d'amorçage orthographique¹. En manipulant la position des lettres partagées par l'amorce et la cible (position identique ou différente), certains auteurs ont estimé l'influence de la position absolue vs relative des lettres dans la séquence. Ces effets sont mesurés relativement à une condition contrôle où l'amorce et la cible ne partagent aucune lettre.

Humphreys, Evett et Quinlan (1990) furent les premiers à montrer que les effets d'amorçage varient en fonction à la fois du nombre et de la position des lettres partagées par l'amorce (non-mot) et la cible (mot) dans une tâche d'identification perceptive. Un recouvrement orthographique plus important produit un effet de facilitation plus grand,

¹ La technique d'amorçage consiste à présenter un mot *cible* précédé par un stimulus *amorce* contrôle, et à comparer cette condition contrôle avec des conditions de partage partiels orthographiques, phonologiques, morphologiques ou sémantiques. La tâche du participant porte sur le mot cible. Les effets d'amorçage sont évalués par le degré de facilitation (i.e., l'amélioration des performances) ou le degré d'inhibition (i.e., la détérioration des performances) du mot cible par rapport à l'effet de l'amorce contrôle.

mais seulement lorsque les lettres partagées occupent la même position dans l'amorce et dans la cible. Des résultats comparables ont été obtenus par Perressotti et Grainger (1999) dans la tâche de décision lexicale. Ces auteurs ont montré que l'amorce « BSLCRN » (qui préserve la position absolue des lettres B, L, C et N) ne facilite pas plus les performances de reconnaissance du mot cible « BALCON » que l'amorce « BCLN » qui préserve la position relative des lettres seulement. De plus, l'insertion de symboles non-alphabétiques pour indiquer la position absolue des lettres (« B-L-C-N ») entraînent des effets de facilitation identiques à ceux obtenus avec l'amorce « BCLN ».

L'ensemble de ces travaux montrent clairement que les lettres sont codées selon leurs positions relatives dans le mot, plutôt qu'en fonction de leurs positions absolues (voir aussi, Perressotti & Grainger, 1995 ; Grainger & Jacobs, 1991).

Les lettres du mot ainsi activées provoquent alors l'activation d'un ensemble limité de *candidats lexicaux* formellement proches du stimulus puis la *sélection* d'un seul de ces candidats (voir par exemple, McClelland & Rumelhart, 1981 ; Paap, et al., 1982 ; Segui, 1991). Les modèles de lecture ont tenté d'identifier les candidats possibles, de modéliser la façon dont s'effectue la sélection d'un seul d'entre eux et de déterminer dans quelle mesure la classe des candidats lexicaux influence la reconnaissance du mot.

1-2 Accès au lexique, sélection et voisinage orthographique

1-2-1 Voisinage orthographique : taille versus fréquence du voisinage

Les travaux de Coltheart, Davelaar, Jonasson et Besner (1977) sont à l'origine des études où l'impact du voisinage orthographique est étudié en faisant varier la taille du voisinage (N). Les mots étant construits à partir d'un nombre restreint de lettres, un mot donné partage la plupart du temps des lettres avec d'autres mots. La classe des mots orthographiquement *voisins* du mot cible rassemble selon la définition traditionnelle l'ensemble des mots de même longueur que ce dernier, qui diffèrent par une seule lettre. Par exemple, le mot « chaise » a deux voisins orthographiques (N=2), les mots « chaire » et « chasse », tandis que le mot « blanc » n'a qu'un seul voisin, « flanc ». Dans leur expérience princeps, Coltheart et al. (1977) ont montré que les non-mots possédant une valeur de N élevée tendaient à donner lieu à des latences de décisions lexicales plus élevées et moins précises que les non-mots avec un N faible. Toutefois, aucun effet de voisinage orthographique ne fut trouvé pour les mots. Ces travaux furent à l'origine de nombreuses études publiées sur l'effet du voisinage orthographique (voir Andrews, 1997 et Mathey, 2001 pour des revues récentes). Ces études ont manipulé essentiellement deux indices liés au voisinage orthographique : la taille ou la densité du voisinage (i.e., le nombre de voisins orthographiques d'un mot donné) et, la fréquence du voisinage (i.e., la présence de voisins orthographiques de plus haute fréquence que le mot cible).

1-2-2 Voisinage orthographique : faits expérimentaux

Les expériences publiées relatives à l'effet du voisinage orthographique appuient généralement l'hypothèse selon laquelle le voisinage d'un mot influence l'accès au lexique ; les données qu'elles rapportent forment néanmoins un ensemble complexe, dépendant essentiellement des tâches expérimentales, des indices de voisinage utilisés et de la langue considérée. Pour autant, les résultats obtenus dans les tâches de lecture habituelles peuvent être résumés comme suit.

De manière générale, les résultats obtenus dans la **tâche d'identification perceptive**² montrent un effet inhibiteur de la taille et de la fréquence du voisinage en français (Bozon & Carbonel, 1996 ; Ferrand & Grainger, 2001 ; Grainger & Segui, 1990 ; Grainger & Jacobs, 1996 ; Ziegler, Rey & Jacobs, 1998), en espagnol (Carreiras, Perea & Grainger, 1997) et, en anglais (Snodgrass & Mintzer, 1993 ; mais voir Sears, Hino & Lupker, 1995 pour des effets facilitateurs ou nuls). Ainsi, les mots ayant au moins un voisin de plus haute fréquence et/ou possédant de nombreux voisins entraînent des pourcentages d'identification moins élevés que les mots n'ayant pas de voisins plus fréquents et/ou très peu de voisins. Cependant bien qu'un effet inhibiteur soit manifeste dans la tâche d'identification perceptive, l'effet du voisinage sur les mots est plus controversé dans la **tâche de décision lexicale** où le sujet doit décider le plus rapidement et le plus correctement possible si la suite de lettres présentée sur un écran d'ordinateur correspond ou non à un mot de sa langue. L'effet de fréquence du voisinage est généralement inhibiteur en français (Bozon & Carbonel, 1996 ; Ferrand & Grainger, 2001 ; Grainger, O'Regan, Jacobs & Segui, 1989 ; Grainger & Segui, 1990 ; Grainger, O'Regan, Jacobs & Segui, 1992 ; Grainger & Jacobs, 1996), en espagnol (Carreiras et al., 1997) et, en anglais (Paap & Johansen, 1994 ; Perea & Pollatsek, 1998 ; mais voir Forster & Shen, 1996, pour des effets facilitateurs ou nuls). En revanche, l'effet de la taille du voisinage est plutôt facilitateur en français (Bozon & Carbonel, 1996 ; Grainger & Jacobs, 1996) et en anglais (Andrews, 1989 ; 1992 ; Forster & Shen, 1996 ; Sears et al., 1995 ; Ziegler & Perry, 1998) ; bien que certains auteurs aient obtenu des effets nuls (Carreiras et al., 1997 ; Coltheart et al., 1977 ; Ferrand & Grainger, 2001 ; Grainger et al., 1989 ; Paap et al., 1994).

En bref, dans la tâche d'identification perceptive, les effets de fréquence et de densité du voisinage sont inhibiteurs ; des effets facilitateurs et inhibiteurs, respectivement exercés par la densité et la fréquence du voisinage sont observés dans la tâche de décision lexicale. Les études portant sur les mouvements oculaires produits au cours de la lecture, fournissant des mesures complémentaires de celles des temps de réaction et des erreurs sont compatibles dans l'ensemble avec l'hypothèse de l'influence inhibitrice du voisinage orthographique lors d'un processus « naturel » d'identification (Grainger et

² La tâche d'identification perceptive consiste à présenter très brièvement un mot de manière à ce que toute l'information visuelle soit disponible mais seulement pendant quelques fractions de secondes ne permettant pas généralement l'identification complète du stimulus. La tâche du sujet est d'identifier le mot. Des variantes de cette tâche sont utilisées pour dégrader le mot comme le démasquage progressif ou la présentation fragmentée.

al., 1989 ; Perea & Pollatsek, 1998 ; Pollatsek, Perea & Binder, 1999).

Les effets de voisinage orthographique témoignent de l'intervention des mots voisins au cours de la reconnaissance des mots écrits. La nature facilitatrice ou inhibitrice des effets du voisinage orthographique a des implications théoriques importantes pour les modèles de la lecture.

1-3 Modèles de la reconnaissance visuelle des mots

Les modèles d'accès au lexique de type *activation vérification* (Paap et al., 1982 ; Paap & Johansen, 1994) et *activation interactive* (McClelland & Rumelhart, 1981) prévoient explicitement l'intervention du voisinage orthographique, tel que définit par Coltheart et al. (1977) lors de la reconnaissance des mots³.

1-3-1 Le modèle à activation vérification (AV)

Le **modèle AV** (Paap et al., 1982) comprend deux types d'unités, des « unités lettres » stockées dans un *alphabet mental* et des « unités mots » stockées dans le *lexique mental*. Chaque lettre de l'alphabet est représentée par une liste de traits visuels. Le niveau d'activation d'une lettre est déterminé par sa probabilité de confusion avec les autres lettres, calculée à partir de données empiriques recueillies pour chaque lettre et à chaque position. Les lettres codant la position spécifique vont être activées plus ou moins fortement dans l'alphabet en fonction de leur degré de « confusabilité ». Ainsi la lettre « P » en première position du mot « PORE » va activer l'unité lettre « P » plus que n'importe quelle autre lettre. D'autres lettres, visuellement similaires, comme le « R, G, A, B, etc » seront activées mais dans une moindre mesure.

Dans ce modèle, l'identification d'un mot se fait en deux étapes. Dans la première étape, dite d'*activation*, l'information visuelle issue du mot stimulus conduit à la sélection d'un petit ensemble de candidats lexicaux. La classe des candidats est constituée par la représentation du stimulus (par exemple le mot anglais « PORE ») et par celle de ses voisins orthographiques (« PORK », « GORE », « BORE », « POKE », etc). Puis, dans une seconde étape dite de *vérification*, les représentations lexicales des candidats, engendrées lors de l'activation sont comparées successivement à la représentation du stimulus, par ordre décroissant de fréquence d'usage des mots⁴ (voir aussi Forster, 1976 pour une proposition similaire). Une décision est prise pour chaque candidat vérifié. La reconnaissance du mot a lieu lorsque le critère de décision est atteint. Dans le cas contraire, le candidat est rejeté et la vérification porte sur le candidat suivant. Par

³ Ces modèles ont été développés à l'origine pour rendre compte de « l'effet de supériorité du mot » selon lequel le report d'une lettre est supérieur lorsque cette lettre est présente dans un mot plutôt qu'isolément (e.g., Reicher, 1969).

⁴ La fréquence d'usage des mots ou fréquence d'occurrence est une mesure objective permettant d'estimer le nombre de fois qu'un lecteur a rencontré un mot donné.

conséquent, seuls les candidats lexicaux plus fréquents seraient susceptibles d'interférer avec la reconnaissance du mot cible. Plus précisément, le degré d'interférence du voisinage orthographique sur la reconnaissance du mot cible, c'est à dire l'effet inhibiteur devrait être proportionnel au nombre de voisins plus fréquents.

1-3-2 Le modèle à activation interactive (AI)

Selon le **modèle AI** développé initialement par McClelland et Rumelhart (1981, voir aussi Rumelhart & McClelland, 1982), le processus de reconnaissance d'un mot écrit implique trois niveaux de traitement, opérant entièrement en parallèle : les « traits visuels distinctifs » (e.g., lignes horizontales, verticales et diagonales), les « lettres » et les « mots ». Les lettres sont codées en fonction de leurs positions à l'intérieur du mot et traitées simultanément. Les différentes unités sont interconnectées à l'intérieur d'un même niveau et à travers les niveaux adjacents. Les connexions sont excitatrices entre deux unités compatibles (par exemple, la lettre « t » en première position et les mots « TAKE » et « TIME ») et inhibitrices entre deux unités incompatibles (par exemple, les mots « TAKE » et « TIME »). Les connexions entre les unités lexicales sont inhibitrices. Plus précisément, au niveau des mots, il existerait un mécanisme d'inhibition lexicale qui entraînerait une inhibition mutuelle de tous les candidats lexicaux actifs afin de permettre la reconnaissance du mot cible, ce qui introduit la notion de compétition entre les différentes unités mots. La représentation qui excède la première son seuil d'activation sera identifiée comme la cible. La compétition entre les mots retarde par conséquent l'identification et un effet inhibiteur du voisinage émerge.

Selon ce modèle, la perception d'un mot est influencée non seulement par les informations de bas niveau (niveau des lettres) mais aussi par les informations de plus haut niveau (niveau des mots). Lorsqu'un stimulus est présenté, les caractéristiques des lettres détectées activent simultanément les unités des lettres qui les contiennent et inhibent les autres unités de lettres. De même, les lettres activées activent les unités des mots qui les contiennent et inhibent les autres unités lexicales. En retour, les mots activés activent les lettres dont ils sont composés et inhibent les autres lettres, il s'agit de la *réverbération*. Le niveau d'activation d'un mot est ainsi déterminé par l'excitation en provenance des lettres, par l'inhibition latérale entre les mots ainsi que par sa fréquence d'usage ; un mot s'active d'autant plus que sa fréquence d'usage est élevée. La force de l'inhibition latérale est proportionnelle au niveau d'activation du mot : plus une représentation lexicale est activée plus sa force inhibitrice est grande. Ce jeu des activations et inhibitions se poursuit jusqu'à ce que l'unité lexicale la plus activée atteigne le seuil critique d'activation : le mot est alors reconnu. C'est donc le produit final de deux phénomènes antagonistes, l'activation et l'inhibition, qui détermine la vitesse de reconnaissance d'un mot.

Les données expérimentales montrant un effet du voisinage orthographique sont globalement en accord avec les modèles de la reconnaissance des mots qui postulent l'activation d'un ensemble de candidats lexicaux formellement proches du stimulus mot. Toutefois, au-delà du consensus actuel concernant l'influence des voisins orthographiques sur la reconnaissance des mots, les mécanismes en jeu restent à

préciser. Les modèles de lecture acceptent généralement l'hypothèse selon laquelle les tâches expérimentales qui nécessitent l'accès à une représentation lexicale unique mettent en œuvre un processus commun d'identification lexicale (e.g., Jacobs & Grainger, 1992 ; Grainger & Jacobs, 1996). Cependant, des processus non lexicaux spécifiques aux tâches utilisées sont également susceptibles d'intervenir lors du traitement d'un stimulus.

Nous allons examiner dans quelle mesure les modèles AV et AI décrits plus haut, ainsi que leurs extensions plus récentes (Paap & Johansen, 1994 ; Grainger & Jacobs, 1996), prennent en compte les processus spécifiques aux tâches expérimentales. Les effets de fréquence et de densité du voisinage obtenus dans les tâches de reconnaissance des mots seront discutés en référence à ces deux grands cadres théoriques.

1-3-3 Stratégies de lecture et décision lexicale

Les modèles décrits ci-dessus sont globalement compatibles avec l'effet inhibiteur de la fréquence du voisinage observé dans les tâches de décision lexicale et d'identification perceptive et, en particulier avec l'effet inhibiteur du nombre de voisins des pseudomots dans la tâche de décision lexicale. Dans le modèle AV (Paap et al., 1982), seuls les voisins plus fréquents d'un stimulus mot contribuent à retarder sa reconnaissance. Dans le modèle AI (McClelland & Rumelhart, 1981), l'intervention d'un mécanisme d'inhibition lexicale lors du processus d'identification d'un stimulus mot permet de rendre compte de l'effet inhibiteur du voisinage orthographique. Par ailleurs, le modèle AV prévoit que les pseudomots activent des représentations lexicales selon un mécanisme similaire à celui des mots. D'après Coltheart et al., (1977), déterminer qu'un stimulus n'est pas un mot revient à établir qu'il est absent du lexique. Cette décision peut être obtenue après vérification de toutes les entrées lexicales sélectionnées. Ainsi, tout voisin orthographique activé, quelle que soit sa fréquence, doit être vérifié et rejeté, ce qui retarde d'autant plus le rejet du pseudomot qu'il possède de mots voisins. Dans le modèle AI, les résultats obtenus pour les pseudomots, dans la tâche de décision lexicale sont généralement attribués à l'activité globale du lexique (niveau mots). Plus un pseudomot produit d'activation dans le lexique, plus il est difficile à rejeter.

Cependant, il est beaucoup plus difficile de réconcilier les modèles AV et AI avec l'effet facilitateur de la densité du voisinage sur la décision lexicale des mots (e.g., Andrews, 1989 ; 1992). Une modification apportée au modèle AV par Paap et Johansen (1994) et au modèle AI par Grainger et Jacobs (1996) consiste à attribuer l'effet facilitateur de la densité du voisinage à un processus propre à la tâche. Ces auteurs postulent que la réponse de décision lexicale est donnée soit sur la base du traitement de vérification lexicale « profond », soit lorsque l'activation produite dans le lexique atteint un seuil critique. Ainsi, en décision lexicale, les participants pourraient ne pas effectuer un traitement complet des stimuli et répondre prématurément sur la base de l'activité globale du lexique, dans ce cas un effet facilitateur de la densité du voisinage émerge. En revanche, dans d'autres circonstances (non précisées par ces auteurs), lorsque les participants effectuent un traitement plus profond des stimuli, un effet inhibiteur de la fréquence du voisinage se produit.

En résumé, la définition traditionnelle du voisinage orthographique prend en compte uniquement les mots voisins qui possèdent exactement le même nombre de lettres que le stimulus. Cette définition est compatible avec les modèles d'accès au lexique privilégiant l'existence d'un processus analytique où les mots sont reconnus sur la base des unités sous-lexicales, les lettres. Selon un postulat commun à ces modèles, la longueur du mot est codée et toutes les lettres d'un mot le sont en fonction de leur position respective à l'intérieur du mot. La représentation du mot cible requiert alors d'être sélectionnée parmi l'ensemble des candidats lexicaux activés. Deux grandes conceptions théoriques du processus de sélection lexicale sont envisagées. (1) En accord avec l'hypothèse d'*inhibition lexicale*, les modèles AI supposent que les mots voisins entrent en compétition pour la reconnaissance du mot cible, c'est à dire qu'ils s'inhibent mutuellement. La compétition entre les mots retarde par conséquent l'identification du mot cible et un effet inhibiteur du voisinage émerge (McClelland & Rumelhart, 1981, voir aussi Grainger & Jacobs, 1996 ; Jacobs & Grainger, 1992). (2) Selon le mécanisme de *vérification sérielle* des modèles AV, les candidats lexicaux sont vérifiés un par un par ordre décroissant de fréquence d'usage des mots. Seuls les candidats lexicaux plus fréquents seraient susceptibles d'interférer avec la reconnaissance du mot cible (Paap et al., 1982). Des extensions récentes des modèles AV et AI ont été élaborées afin de rendre compte des effets spécifiques liés à la tâche de décision lexicale.

Dans l'ensemble, l'hypothèse d'inhibition lexicale semble fournir la meilleure explication des effets de voisinage orthographique lors de la reconnaissance des mots écrits.

Chapitre 2 Le phénomène de l'effet de la position du regard dans les mots

2-1 Position du regard dans les mots et efficacité de la reconnaissance

2-1-1 Enregistrements des mouvements oculaires

Les mouvements des yeux durant la lecture sont caractérisés par une succession de *saccades* et de *fixations*. Les saccades permettent de progresser dans le texte et les fixations sont nécessaires pour extraire l'information. Déjà, Erdmann et Dodge (1898) avaient souligné que les fixations du regard étaient localisées presque exclusivement sur les mots plutôt que sur les espaces inter mots et, en grande majorité vers le centre des mots. Ces premiers travaux ont reçu de considérables attentions par la suite. En particulier Dunn-Rankin (1978) et Rayner (1979) ont montré que dans des conditions « normales » de lecture de textes, les lecteurs tendent à fixer *préférentiellement* une position légèrement à gauche du centre des mots, du moins pour les lecteurs anglais et français.

O'Regan (1981) a le premier suggéré qu'il était fortement avantageux de fixer légèrement à gauche du centre d'un mot afin de le reconnaître avec la plus grande efficacité. Il a ainsi proposé que des fixations supplémentaires, témoignant de l'échec de l'identification du mot lors de la première fixation, devraient être nécessaires si la fixation initiale dans le mot est éloignée de cette position qu'il a appelé la « *position du regard convenable* ». Cette prédiction a été confirmée par la suite par plusieurs auteurs (McConkie, Kerr, Redix, Zola & Jacobs, 1989 ; O'Regan, Lévy-Schoen, Pynte, & Brugaillère, 1984 ; Radach & Kempe, 1993, Underwood, Clew & Everatt, 1990 ; Vitu, 1991 ; Vitu, O'Regan & Mittau, 1990). Ainsi, la probabilité de *re-fixer* un mot, dans des conditions de lecture écologiques, est minimale lorsque le regard se pose initialement vers le centre du mot et augmente de manière systématique à mesure que l'œil dévie de cette position « convenable ». Des données comparables ont été obtenues en lecture de mots isolés. La durée totale de fixation d'un mot présenté isolément est minimale lorsque le sujet fixe initialement la position convenable et diminue à mesure que la fixation initiale dans le mot s'en éloigne (voir Figure 1, O'Regan et al., 1984).

Figure 1. Durée de fixation totale de mots de 9 lettres en fonction de la position initiale du regard dans le mot lors d'une tâche de comparaison de mots (Expérience 2, O'Regan et al, 1984). Le temps de fixation était minimal lorsque les sujets fixaient la 5^e lettre des mots et augmentait à mesure que le regard s'éloignait de cette position "convenable".

L'ensemble de ces travaux démontre clairement que la lecture de mots est fortement influencée par la position du regard initiale dans le mot. Dans le but d'examiner plus directement la relation entre la position préférée initialement observée par Dunn-Rankin (1978) et Rayner (1979) et la position du regard convenable, Vitu et al. (1990) ont comparé les probabilités de refixation (ainsi que les durées de fixation) dans une condition de lecture de mots isolés et dans une condition de lecture de textes. Un effet de la fixation initiale du regard est observé dans les deux conditions cependant, cet effet est plus faible en condition de lecture de textes et la position optimale se situe davantage vers le centre des mots. Ces observations ont amené ces auteurs à conclure que la position préférée du regard observée en lecture « naturelle » soit la conséquence d'une stratégie employée afin d'optimiser la reconnaissance des mots (voir O'Regan, 1990 pour une synthèse). Toutefois, une vision opposée a récemment été proposée par Nazir (2000 ; 2003), suggérant que la position préférée du regard serait, non pas la conséquence mais la cause de l'observation d'une position optimale lors du traitement des mots isolés.

2-1-2 Le paradigme de la position variable du regard

L'existence d'un effet de la position du regard dans le mot a été répliquée avec des mesures autres que celles des mouvements oculaires. Les latences et les erreurs mesurées dans des tâches de dénomination et de décision lexicale tendent à être minimales lorsque les yeux fixent légèrement à gauche du centre du mot. Plus précisément, la Figure 2 présente un effet de la position du regard dans le mot typique obtenu dans une tâche de décision lexicale pour des mots français de 5 et 7 lettres (d'après Nazir, 1993). Les performances ont été mesurées via la *technique de la position*

variable du regard (voir Figure 2, panel de droite). Comme le montre le graphique de gauche de la figure, les pourcentages de reconnaissance d'un mot sont optimaux lorsque la position du regard est légèrement à gauche du centre du mot, et diminuent progressivement à mesure que l'œil dévie de cette *position optimale* (« optimal viewing position », OVP). De plus, les fixations du regard situées sur le début du mot (positions 1 et 2) tendent à donner lieu à de meilleures performances de reconnaissance que les fixations situées sur la fin du mot (positions 4 et 5). Cette variation des performances en fonction de la position du regard dans le mot porte le nom générique d'*effet de la position du regard* dans le mot (EPR par la suite). La courbe de l'EPR dans les mots présente classiquement la forme d'un « J-inversé », c'est à dire une asymétrie à gauche, du moins dans les langues romanes⁵ comme le français, l'allemand ou l'anglais par exemple.

Figure 2. À gauche, courbes classiques de l'effet de la position du regard pour des mots de 5 lettres (symboles pleins) et des mots de 7 lettres (symboles vides). À droite, illustration de la technique de la position variable du regard pour un mot de 5 lettres. Cette méthode consiste à demander au sujet de fixer un point situé au centre de l'écran d'un ordinateur. Il est ensuite remplacé par un mot présenté très brièvement, pendant un temps inférieur à une saccade oculaire de manière à contrôler la position du regard dans le mot. Le mot apparaît à différentes positions par rapport au point de fixation (Reproduit à partir de Nazir, 1993).

2-1-3 Robustesse du phénomène

L'EPR qui tire son origine des études portant sur les mouvements oculaires a donné lieu à un grand nombre de travaux dans le domaine de la reconnaissance visuelle des mots au cours de cette dernière décennie. Ces études ont indiscutablement reproduit un EPR dans les mots à travers différentes tâches de lecture (dénomination, identification perceptive, décision lexicale, comparaison de mots, etc), utilisant différentes variables dépendantes (latences, précisions des réponses, comportements oculaires), pour des mots de différentes longueurs, de différentes fréquences et dans plusieurs langues: *français* (Aghababian & Nazir, 2000 ; Brysbaert, Vitu & Schroyens, 1996 ; Montant, Nazir & Poncet, 1998 ; Nazir & Montant, 1996; Nazir, 1993 ; Nazir, Jacobs & O'Regan, 1998; Nazir, O'Regan & Jacobs, 1991, O'Regan, 1981 ; O'Regan et al., 1984; Pynte, 1996 ; Stevens & Grainger, 2003 ; Vitu, 1991 ; Vitu et al., 1990) *allemand* (Nazir, Heller & Sussman, 1992 ; Radach & Kempe, 1993), *anglais* (McConkie, Kerr, Redix, & Zola, 1988 ; McConkie et al., 1989 ; Zola, 1984), *hollandais* (Brysbaert & d'Ydewalle, 1988 ; Brysbaert & Meyers, 1993; Brysbaert, 1994; Brysbaert et al., 1996), *arabe* (Farid & Grainger, 1996), *hébreu* (Deutsch & Rayner, 1999) et, *japonais* (Kajii & Osaka, 2000).

L'EPR compte parmi les facteurs les plus robustes de la reconnaissance visuelle des mots, du moins dans les systèmes d'écritures alphabétiques. Il reste néanmoins un challenge pour les modèles actuels de la lecture qui ne considèrent pas l'influence de la variation de la position du regard dans le mot (cf. O'Regan & Jacobs, 1992).

⁵ Les langues romanes constituent l'ensemble des langues dérivées du latin (anglais, espagnol, français, italien, etc.).

L'interprétation communément admise de l'EPR tire son origine des propriétés du système visuel. En raison de la diminution de la concentration des cônes au niveau de la rétine, l'acuité visuelle chute de manière dramatique au sein même de la fovéa⁶. L'acuité est maximale à la fixation et diminue linéairement avec l'augmentation de la distance séparant l'objet du centre de la fixation de l'œil (Levi, Klein, Aitsebaomo, 1985; Olzak & Thomas, 1986; Weymouth, Hines, Acres, Raaf & Wheeler, 1928). La signification pratique de ce fait est qu'un objet est d'autant mieux perçu qu'il est proche du point de fixation de l'œil. Ainsi, la probabilité d'identifier une lettre individuelle placée directement sous la fixation est égale à 1 et diminue de manière linéaire à mesure qu'on s'éloigne de la fixation (Taylor & Brown, 1979, voir aussi Bouma, 1970 ; Townsend, Taylor & Brown, 1971). Toutefois, l'hypothèse simple en terme d'acuité visuelle « pure » prédit une courbe de l'EPR dans le mot parfaitement symétrique avec une position optimale au centre du mot où un maximum de lettres serait alors visibles (voir McConkie et al., 1989).

Au-delà du consensus actuel concernant le rôle majeur des limitations de l'acuité visuelle, l'origine de l'asymétrie de la fonction de l'EPR reste encore une énigme. Différentes théories ont été proposées pour rendre compte du phénomène de l'EPR dans les mots, chacune proposant un facteur critique qui, combiné aux effets de l'acuité visuelle explique l'asymétrie de l'EPR.

2-2 Asymétrie de l'EPR : plusieurs interprétations

2-2-1 Asymétrie de la lisibilité des lettres

Dans le but de dépasser l'apparente incohérence relevée entre les effets liés à l'acuité visuelle et l'asymétrie de la fonction de l'EPR dans les mots, Nazir et al. (1991 ; 1992 ; 1998) ont proposé que l'EPR dans les mots serait liée à l'*asymétrie de la lisibilité des lettres* à droite et à gauche de la fixation. En mesurant la lisibilité des lettres à différentes excentricités à gauche et droite de la fixation, Nazir et collaborateurs (1991 ; 1998) ont montré que la diminution des performances de perception des lettres dépendait non seulement de la distance par rapport à la fixation mais également du champ visuel stimulé. Ainsi, pour une excentricité donnée, les lettres à droite de la fixation sont plus facilement identifiées que les lettres à gauche de la fixation (voir aussi Bouma, 1973 ; Kajii & Osaka, 2000). L'asymétrie de l'EPR dans les mots serait due aux variations de la lisibilité des lettres constitutives du mot : la lisibilité totale des lettres, proportionnelle à la probabilité de reconnaissance du mot est ainsi maximale lorsque le regard se pose légèrement à gauche du centre.

D'autres recherches, ont également mentionné un avantage du champ visuel droit pour la perception du matériel orthographique (e.g., Bradshaw & Nettleton, 1983; Bryden,

⁶ La fovéa est définie comme les 2-3 degrés d'angle visuel autour du point de fixation (1 degré de champ visuel équivaut à une distance de 1 cm vue à 57 cm de distance).

1982; voir aussi, Chiarello, Lui, & Shears, 2001 pour une revue) dont l'origine a été principalement attribuée à deux facteurs, les *différences fonctionnelles entre les deux hémisphères* et les *habitudes liées à la direction de lecture*.

2-2-2 Spécialisation hémisphérique

L'asymétrie de la fonction de l'EPR a été attribuée à l'avantage du champ visuel droit pour la perception du matériel orthographique et ce, au sein même de la zone fovéale (e.g., Brysbaert & d'Ydewalle, 1988 ; Brysbaert et al., 1996; Shillcock, Ellison & Monaghan, 2000 ; Whitney, 2001 ; voir aussi Brysbaert, 2004 pour une revue récente). Étant donné que chaque hémichamp visuel se projette entièrement sur le cortex visuel contralatéral, les lettres présentées à droite de la fixation bénéficieraient d'un traitement plus efficace, du fait de leur transfert direct vers l'hémisphère gauche, dominant pour le langage.

L'hypothèse de la dominance hémisphérique pour le langage a directement été testée par Brysbaert (1994) à travers la comparaison de l'EPR chez des lecteurs hollandais « dominant gauche » (i.e., présentant une dominance hémisphérique classique à gauche pour le langage) et « dominant droit » (i.e., présentant une dominance hémisphérique non habituelle à droite pour le langage). Les résultats obtenus à l'issue de cette étude ont mis en évidence une asymétrie de la fonction de l'EPR moins prononcée (i.e., une courbe de l'EPR davantage symétrique) chez les sujets « dominant droit » comparée aux sujets « dominant typique », suggérant que bien que la dominance hémisphérique joue un rôle dans l'asymétrie des champs visuels droite/gauche, elle ne rend pas entièrement compte du phénomène de l'EPR dans les mots. Toutefois, ces données sont à prendre avec précaution en raison de la difficulté à déterminer avec certitude la dominance hémisphérique des sujets pour le langage (évaluée dans cette étude au travers d'une série de batteries de tests de la latéralité). Par ailleurs, les sujets dont les systèmes de traitement du langage sont localisés dans l'hémisphère droit sont rares ; de récentes études utilisant les techniques d'imagerie cérébrale (Knecht et al., 2000, 2003 ; Pujol et al., 1999) ont révélé qu'environ 95% des sujets droitiers et 75% des sujets gauchers traitent le langage principalement avec l'hémisphère gauche et seulement 20% des gauchers avec l'hémisphère droit.

2-2-3 Facteurs attentionnels, direction de lecture et apprentissage perceptif

Étant donné que les mots sont lus de gauche à droite dans les langues alphabétiques romanes, l'attention serait davantage portée vers la droite que vers la gauche de la fixation. Certains auteurs ont par conséquent suggéré que l'asymétrie de l'EPR serait liée aux habitudes de lecture et notamment à la direction de lecture, laquelle contraindrait notre manière de percevoir les lettres et les mots (e.g., Kirsner & Schwartz, 1986 ; Nazir, 2000). Ainsi, il a été montré que la taille de l'*empan perceptif* (i.e., la région à partir de laquelle on extrait l'information visuelle « utile » lors d'une seule fixation) diffère selon le sens de la lecture des mots. L'empan perceptif est plus large à droite qu'à gauche de la fixation chez les lecteurs anglais, tandis qu'il est plus large à gauche qu'à droite de la

fixation chez les lecteurs hébreux (Pollatsek, Bolozky, Well & Rayner, 1981 ; Rayner & Pollatsek, 1989).

Selon une conception originale proposée par Nazir (2000 ; 2003 ; Nazir et al., 1998), l'apprentissage perceptif, lié aux habitudes de lecture serait à l'origine de l'asymétrie de la courbe de l'EPR dans les mots. Comme nous l'avons vu plus haut, les yeux des lecteurs, en condition de lecture « naturelle » tendent à se poser le plus souvent à gauche du centre des mots et la probabilité avec laquelle le regard se pose ailleurs dans les mots décroît avec l'éloignement de cette position préférée (Dunn-Rankin, 1978 ; McConkie et al., 1988 ; Rayner, 1979 ; Vitu et al., 1990)⁷. Plus précisément, la corrélation observée entre le patron de mouvements des yeux en lecture et les courbes de reconnaissance du mot en fonction de la position du regard a conduit ces auteurs (Nazir et al., 1998 ; 2000 ; 2003) à formuler l'idée d'un apprentissage perceptif au cours de la lecture. L'asymétrie de la fonction de l'EPR est ainsi expliquée comme résultant de la fréquence avec laquelle les mots sont vus dans une région particulière du champ visuel.

Dans l'objectif de tester l'hypothèse liée aux habitudes de lecture, Farid et Grainger (1996) ont montré à l'issue d'une étude trans-linguistique menée auprès de locuteurs francophones et arabophones, que les lecteurs français présentaient une asymétrie de la fonction de l'EPR classique avec une position optimale à gauche du centre des mots ; en revanche, chez les lecteurs de la langue arabe, la fonction de l'EPR est davantage symétrique avec un maximum au centre des mots, suggérant que la direction de lecture influence la forme de la courbe de l'EPR mais n'explique néanmoins pas la totalité du phénomène de l'EPR (voir aussi Deutsch & Rayner, 1999 pour des résultats similaires en hébreu).

2-2-4 Contraintes lexicales et structure informative des mots

Une vision plus robuste concernant l'influence des contraintes lexicales sur la reconnaissance des mots postule que l'EPR serait lié à la structure informative des mots. Dès l'introduction du paradigme de la position variable du regard, O'Regan et al. (1984) ont souligné le fait que généralement plus de mots partagent les mêmes lettres finales que de mots ne partagent les mêmes lettres initiales, du moins en anglais et en français (e.g., Clark & O'Regan, 1999 ; O'Regan et al., 1984). Ainsi, en raison des plus fortes contraintes lexicales exercées sur les premières lettres comparativement aux dernières lettres, il est plus avantageux de fixer le début d'un mot que la fin d'un mot. Toutefois, l'hypothèse de la contrainte lexicale comme unique contributeur à l'asymétrie de l'EPR est remise en cause par les mots dont l'informativité se situe à la fin de la séquence, lesquels ne présentent pas d'asymétrie inversée à droite de la fonction de l'EPR mais, une position optimale au centre des mots (e.g., Brysbaert et al., 1996 ; O'Regan et al., 1984).

En résumé, le phénomène de l'EPR dans les mots ne peut être expliqué par l'asymétrie gauche droite de la lisibilité des lettres ou par la distribution de l'information lexicale dans le mot uniquement ; plutôt l'EPR dans les mots semble tirer son origine de la

⁷ La distribution des sites d'atterrissage au cours de la lecture « normale » serait le résultat de facteurs de bas niveau de type visuo-oculomoteur (Nazir et al., 1998 ; voir Ducrot & Pynte, 2002 pour une proposition similaire).

combinaison de plusieurs facteurs incluant les facteurs visuels, les habitudes de lecture, la distribution de l'information lexicale dans le mot et probablement les différences fonctionnelles hémisphériques lors du traitement du langage (e.g., Brysbaert et al., 1996).

Un premier pas vers une réconciliation de ces principales théories a récemment été réalisé par le développement de modèles qui considèrent que l'EPR dans les mots est dû à une combinaison de la lisibilité des lettres et des facteurs lexicaux (Clark & O'Regan, 1999 ; Kajii & Osaka, 2000 ; Stevens & Grainger, 2003).

Chapitre 3 Contraintes lexicales et Lisibilité des lettres

3-1 Contrainte orthographique

3-1-1 L'effet de supériorité du début du mot

Dans les langages comme l'anglais et le français, on reconnaît mieux un mot à partir de ses premières lettres qu'à partir de ses dernières lettres (e.g., Clark & O'Regan, 1999 ; Eriksen & Eriksen, 1974 ; Grainger & Jacobs, 1993 ; Humphreys et al., 1990 ; Inhoff & Tousman, 1990 ; Lima & Pollatsek, 1983 ; O'Regan et al, 1984 ; Pynte, 1996). En utilisant le paradigme d'amorçage Inhoff & Tousman (1990) ont montré que les latences de décision lexicale et de dénomination du mot cible anglais « DINNER » étaient plus rapides lorsque ce dernier était précédé par l'amorce « DINXXX » plutôt que par l'amorce « XXXNER ». De même, Grainger et Jacobs (1993) ont montré que les 4 lettres de l'amorce « TABL% » tendaient à donner lieu à des latences de décision lexicale plus courtes sur le mot cible « table » que les amorces « %ABLE » ou même « TA%LE ».

L'interprétation principalement attribuée à cet avantage du début du mot repose sur la

notion de contrainte lexicale ou « *informativité* » des mots. La contrainte lexicale est liée au nombre de mots pouvant être inféré à partir d'une information sensorielle limitée. Par exemple, si le participant a identifié uniquement les quatre premières lettres « TABL » et qu'il sait qu'elles forment les premières lettres d'un mot de 5 lettres, alors une seule inférence est possible : le mot « table ». D'un autre côté, les quatre dernières lettres « ABLE » d'un mot de 5 lettres sont dites moins contraignantes (ou moins informatives) car plusieurs mots « cable », « fable », « sable », « table » peuvent alors être inférés. D'un point de vue statistique, la première partie des mots est généralement « plus contraignante » que la seconde partie des mots, du moins en français et en anglais (e.g., Clark & O'Regan, 1999 ; O'Regan et al, 1984).

3-1-2 Contrainte lexicale et position du regard dans le mot

Étant donné la différence de structure orthographique véhiculée par les deux parties d'un mot, certains chercheurs ont suggéré que la reconnaissance d'un mot serait facilitée lorsque le regard se pose sur la partie informative du mot car celle-ci bénéficie alors d'une meilleure acuité visuelle et permet ainsi de mieux distinguer le mot donné des représentations lexicales voisines (e.g., Brysbaert et al., 1996 ; O'Regan, 1981 ; O'Regan et al., 1984 ; Pynte, 1996, voir la Figure 3 pour une illustration).

Figure 3. Illustration de la contrainte lexicale. Les quatre premières lettres « TABL » du mot « table » sont plus contraignantes que les quatre dernières lettres « *ABLE », lesquelles sont présentes dans plusieurs mots.*

La première étude ayant manipulé la contrainte lexicale et la position du regard dans le mot a été réalisée par l'équipe d'O'Regan et collaborateurs (1984). Ces auteurs ont sélectionné sur la base d'un comptage des entrées lexicales à partir de dictionnaires, des mots français « à début informatif » et « à fin informative ». Par exemple le mot à début informatif « COCCINELLE » constitue le seul mot de 10 lettres qui commence par les six premières lettres « COCCIN » ; tandis que le mot à fin informative « ARCHITECTE » est le seul mot de 10 lettres qui se termine par les six dernières lettres « ITECTE ». Les résultats de cette étude sont présentés dans la Figure 4. Comme on peut le voir sur la figure, les mots à début informatif ont donné lieu à des temps de fixation plus courts, lorsqu'ils étaient fixés sur la première moitié des mots en comparaison avec les fixations sur la seconde moitié des mots. Toutefois, bien que les durées de fixation dans les mots à fin informative sont inférieures à celles observées pour les mots à début informatif lorsque l'œil se pose sur la fin des mots, la tendance ne s'inverse pas complètement. Les temps de fixation enregistrés pour les mots à fin informative se sont avérés similaires pour les fixations sur le début ou sur la fin des mots (O'Regan et al., 1984, Expérience 3 ; voir aussi Pynte, Kennedy & Murray, 1991 et Pynte, 1996, pour des résultats similaires).

Figure 4. Durée totale de fixation des mots à début informatif et des mots à fin informative en fonction de la position des yeux dans les mots (début vs fin) lors d'une tâche de jugement de relation sémantique (le sujet doit décider si deux mots sont liés sémantiquement comme par exemple « fraise-framboise », d'après O'Regan et al., 1984,

Expérience 3).

Des résultats similaires ont été observés plus récemment, dans une étude en hollandais par Brysbaert et collaborateurs (1996). Ces auteurs ont estimé l'informativité de mots de 5 lettres dans une tâche de complétion de mots. Les mots à début informatif étaient hautement prédictibles à partir des trois premières lettres ; à l'inverse, les mots à fin informative étaient hautement prédictibles à partir des trois dernières lettres. Ils ont montré que les mots à début informatif étaient mieux identifiés lorsque les yeux fixaient le début des mots comparativement aux fixations sur la fin des mots. Cependant, bien que les mots à fin informative tendaient à donner lieu à de meilleures performances d'identifications que les mots à début informatif lorsque les yeux fixaient la fin des mots, la courbe de l'EPR des mots à fin informative ne présente pas une asymétrie inversée (voir la Figure 5).

Figure 5. Pourcentages de réponses correctes des mots à début informatif et à fin informative en fonction de la position du regard dans le mot (fixation sur la 1^o lettre, 3^o lettre ou 5^o lettre, d'après Brysbaert et al., 1996, Expérience 3).

3-2 Structure morphologique

3-2-1 Notion de structure morphologique

La notion de contrainte lexicale est étroitement liée à la structure morphologique des mots, en particulier, les racines des mots constituent généralement les parties informatives des mots. Les mots français peuvent se répartir en *mots simples* et *mots construits*. Les mots simples ou *monomorphémiques* sont des mots qui ne sont pas décomposables ; les mots construits ou *polymorphémiques*, par contre, peuvent être décomposés en éléments significatifs plus petits, les *morphèmes*. Par exemple, le mot « fleur » est un mot simple, par contre les mots « boisson, buvable, imbuvable, pourboire » sont construits à partir du mot simple « boire ». En français, on estime que 75% des mots contenus dans le lexique d'un lecteur « normal » sont polymorphémiques (Rey-Debove, 1984).

La morphologie dérivationnelle étudie la formation des mots nouveaux grâce à l'ajout de morphèmes à des mots existants. Il existe deux types de mots nouveaux, les mots *composés* comme « pourboire » et les mots *dérivés* comme « boisson, buvable ». Les mécanismes dérivationnels de la formation des mots dérivés sont la suffixation, la préfixation et l'infixation. Les affixes (ou morphèmes) liés sont des préfixes, des suffixes ou des infixes. Les préfixes s'ajoutent au début d'un morphème de base ou racine (par exemple « rechanter »), les suffixes s'ajoutent à la fin de la racine (par exemple « chanteur ») enfin, les infixes (par exemple « chantonner ») sont mêlés à la racine des mots.

En français, comme pour la plupart des langues indo-européennes⁸, les formes suffixées sont beaucoup plus fréquentes que les formes préfixées (Greenberg, 1966). Par ailleurs, les mots composés sont beaucoup moins fréquents en français que les mots suffixés. La structure des mots complexes varie à travers les langues et certaines langues sont morphologiquement plus riches que d'autres. Dans les langues isolantes comme le chinois, les mots tendent à être monomorphémiques. Dans les langues flexionnelles, comme l'anglais, le français ou l'hébreu, les mots sont souvent composés de plusieurs morphèmes. Parmi les langues flexionnelles, on distingue les langues *concaténatives* et les langues non concaténatives. Dans les langues concaténatives comme l'anglais, le français ou l'italien par exemple, la préfixation et la suffixation de morphèmes est le procédé morphologique le plus répandu. Par définition, les processus de préfixation et de suffixation entraînent une concaténation linéaire des unités morphémiques ; tandis que l'infixation est non concaténative puisque l'intégrité du morphème de base est modifiée. C'est précisément le cas des langues non concaténatives comme l'hébreu ainsi que la plupart des langues sémitiques comme l'arabe, où la formation morphologique repose principalement sur l'infixation. Ainsi, la plupart des mots hébreux sont constitués par l'entrelacement de deux morphèmes, un squelette de consonnes qui correspond à la racine et un patron phonologique.

Les langues concaténatives et non concaténatives se distinguent par la combinaison des unités morphologiques, ce qui a certainement des implications sur la manière dont les mots complexes sont représentés et traités au cours de la lecture (e.g., Frost, Forster, & Deutsch, 1997).

3-2-2 Morphologie et position du regard dans le mot

Comme nous venons de le voir, la plupart des mots français sont complexes, c'est-à-dire qu'ils sont composés d'une racine et d'un préfixe ou d'un suffixe. Comme les affixes sont répétés dans plusieurs mots, ils fournissent généralement moins « d'informativité » que les racines morphologiques. De ce fait, les fixations situées sur la racine des mots tendraient à donner lieu à de meilleures performances de reconnaissance des mots que les fixations situées sur les affixes.

Holmes et O'Regan (1987) ont comparé dans une tâche de compréhension de phrases l'EPR des mots suffixés et des mots préfixés. La logique sous jacente à cette approche prévoyait que la position optimale du regard serait décalée vers la gauche du centre lorsque la racine se situe au début des mots (comme dans le mot suffixé « **mangeur** ») et décalée vers la droite lorsque la racine se situe à la fin des mots (comme dans le mot préfixé « **relancer** »). Contrairement à leurs prédictions, la forme des courbes de l'EPR s'est révélée similaire dans les mots suffixés et préfixés (voir aussi Beauvillain, 1996 ; Hyöna & Pollatsek, 1998 ; Inhoff, Briihl & Schwartz, 1996 pour des études mesurant le comportement oculaire).

Ainsi, la morphologie des mots français et anglais semble peu contraindre la

⁸ Les langues indo-européennes sont issues de l'Indo-Europe, non directement attestées mais reconstituées par comparaison de différentes langues à l'origine desquelles elles se trouvent.

reconnaissance des mots. La prise en compte de la structure linguistique de la langue considérée peut résoudre certaines incohérences relatives à l'effet de la contrainte lexicale sur la reconnaissance des mots (voir Andrews, 1997 pour une proposition similaire concernant le voisinage orthographique). Ce dernier point peut-être illustré au travers d'une étude bilingue menée par Farid et Grainger (1996). À l'issue d'une tâche d'identification perceptive, ces auteurs ont montré que l'EPR n'était pas modulé par la position de la racine dans les mots français : les mots suffixés et les mots préfixés présentent tous deux une asymétrie classique à gauche (voir le graphique de gauche de la Figure 6). En revanche, une interaction entre la structure morphologique des mots et la position du regard est observée dans le cas des mots arabes. Comme présentés dans la Figure 6, les pourcentages de reconnaissance des mots suffixés arabes sont plus élevés lorsque les fixations sont situées sur le début des mots que lorsqu'elles sont situées sur la fin des mots; le patron de performances inverse est obtenu pour les mots préfixés, mettant en évidence une asymétrie inversée à la fin dans les mots préfixés arabes.

Figure 6. Pourcentages d'identifications correctes de mots de 7 lettres en fonction de la lettre fixée (1st représente la 1^o lettre et 7th, la dernière lettre) pour des mots préfixés et suffixés français et arabes (d'après Farid & Grainger, 1996, Expérience 1).

De même, Deutsch et Rayner (1999) ont manipulé la localisation de la racine dans des mots hébreux au cours d'une tâche d'identification perceptive. Les résultats ont révélé une dépendance entre la position de la racine dans le mot et la forme des courbes de l'EPR. Plus précisément, lorsque la racine est au début du mot, les performances d'identifications sont meilleures lorsque les fixations du regard sont situées au début du mot que lorsqu'elles sont situées à la fin du mot. À l'inverse, lorsque la racine est au centre du mot, le patron de résultats opposé est obtenu : de meilleures performances sont relevées lorsque le regard se pose sur la fin du mot que lorsqu'il se pose sur le début du mot. Enfin, bien qu'un effet de la position du regard persiste lorsque la racine est dispersée dans le mot, les performances d'identification lorsque les fixations sont sur le début du mot sont similaires aux fixations sur la fin des mots (voir Figure 7).

Figure 7. Pourcentages d'identifications correctes de mots hébreux de 7 lettres présentés brièvement (28-42 ms) en fonction de la lettre fixée (1st représente la 1^o lettre et 7th, la dernière lettre) et de la position de la racine dans le mot : au début (panel de gauche), au milieu (panel du milieu) ou dispersée (panel de droite, d'après Deutsch & Rayner, 1999, Expérience 2).

L'ensemble de ces résultats suggère que l'opérationnalisation de la contrainte lexicale des mots serait spécifique à la structure linguistique de la langue considérée. La contrainte lexicale semble se traduire par la notion de similarité orthographique dans les langues romanes et par la notion de structure morphologique dans les langues sémitiques comme l'hébreu ou l'arabe. Comme le soulignent Farid et Grainger (1996), la morphologie fournit plus d'informations critiques au processus de reconnaissance des mots arabes que des mots français. La contrainte lexicale mesurée par la structure orthographique du mot semble être plus appropriée dans le cas de la langue française.

Dans les langues romanes, la notion de contrainte lexicale d'un mot fait référence au nombre de mots partageant les mêmes premières lettres ou les mêmes dernières lettres que ce dernier, s'apparentant ainsi fortement à la notion de voisinage orthographique. Les effets de voisinage orthographique utilisent les indices de densité et de fréquence du voisinage par substitution d'une lettre pour lesquels la longueur des mots et la position des lettres sont déterminantes (voir chapitre1), en accord avec le codage des unités lexicales (mots) et sous-lexicales (lettres) postulé dans les modèles AV et AI. Parallèlement au voisinage orthographique, l'effet de la contrainte lexicale est de nature inhibiteur sur la reconnaissance visuelle des mots, suggérant l'existence d'effets liés à d'autres formes de similarité orthographique.

Bien que les données soient en faveur d'une influence de la contrainte lexicale sur l'EPR dans les mots, l'hypothèse de la contrainte lexicale comme unique contributeur à l'asymétrie de l'EPR est remise en cause par les mots dont l'informativité se situe à la fin de la séquence, lesquels ne présentent pas d'asymétrie inversée à droite de la fonction de l'EPR mais une position optimale au centre des mots (e.g., Brysbaert et al., 1996 ; Holmes & O'Regan, 1987 ; O'Regan et al., 1984). La prise en compte d'autres facteurs, tels que la variation de la lisibilité des lettres en fonction de la position du regard devrait permettre de résoudre certaines incohérences.

3-3 Lisibilité des lettres et position du regard dans le mot

3-3-1 Lisibilité des caractères romans

Bien que la probabilité d'identifier une lettre présentée isolément chute de manière linéaire avec l'augmentation de la distance à partir du centre de la fixation de l'œil (voir chapitre précédent), lorsque la lettre cible est comprise dans une séquence de lettres, sa lisibilité n'est plus soumise aux seules lois de l'acuité visuelle. Une lettre contenue dans une séquence est en effet mieux perçue dans le champ visuel droit que dans le champ visuel gauche (Bouma, 1973 ; Kajii & Osaka, 2000; Nazir et al., 1991 ; 1992 ; 1998). Par ailleurs, les lettres externes (i.e., la première et la dernière lettre) sont mieux perçues que les lettres internes en raison de la réduction du *masquage latéral*, phénomène principalement attribué à l'interaction entre les traits des lettres (Bouma, 1970 ; Estes, 1972 ; Estes, Allmeyer & Reder, 1976 ; Huckauf, Heller & Nazir, 1999 ; Mason, 1982 ; Massaro & Cohen, 1994 ; Prinzmetal, 1992 ; Townsend, Taylor & Brown, 1971).

Dans une série d'études psychophysiques, Nazir et collaborateurs ont mesuré les variations de lisibilité des lettres à différentes positions sur la rétine (Nazir et al., 1991; 1992 ; 1998). La Figure 8 présente une fonction de lisibilité-type obtenue pour des lettres présentées dans des séquences de « k » et fixées soit sur la première lettre (la séquence est alors présentée dans le champ visuel droit ; CVD) soit sur la dernière lettre (la

séquence est alors présentée dans le champ visuel gauche ; CVG, Nazir et al., 1991).

Figure 8. Scores de reconnaissance de lettres cibles contenues dans des séquences de 9 « k » présentées brièvement, fixées sur la dernière lettre (présentation dans le champ visuel gauche) ou sur la première lettre (présentation dans le champ visuel droit). Les symboles vides représentent les performances dans le champ visuel droit et les symboles pleins, les performances dans le champ visuel gauche (d'après Nazir, et al., 1991).

Comme l'indique la Figure 8, les performances diminuent linéairement avec l'augmentation de l'excentricité rétinienne de la lettre à partir du point de fixation dans le CVD et dans le CVG. Cependant, la pente des deux fonctions diffère d'un facteur de 1.8 ; ainsi, lorsque les performances diminuent de 10% pour une excentricité donnée dans le CVD, elles diminuent de 18% pour la même excentricité dans le CVG.

Bien que l'origine de cette asymétrie des champs visuels sur la perception des lettres n'ait jamais été expliquée par Nazir et al. (1991), certaines données dans le domaine de la reconnaissance des mots sont en faveur d'une influence de la dominance hémisphérique gauche pour le langage et des habitudes liées à la direction de lecture sur l'asymétrie des champs visuels (voir chapitre 2).

3-3-2 Asymétrie de la lisibilité des lettres : dominance hémisphérique et direction de lecture.

Une étude récente menée par notre équipe a montré que la dominance hémisphérique et les habitudes de lecture contribuaient toutes deux à l'asymétrie des champs visuels sur la perception des lettres (Nazir, Ben-Boutayab, Decoppet, Deutsch & Frost, 2004). Plus précisément, nous avons comparé le patron de lisibilité des lettres dans des langages dont la direction de lecture est opposée. Des chaînes de caractères identiques en alphabet roman (lettre 'x') ou hébreu (lettre 'aleph') ont été brièvement exposées à des lecteurs dont la langue natale était soit l'anglais, soit l'hébreu. Un des caractères de la chaîne était une lettre cible (par ex. un 'a') qui devait être identifiée par le participant. Les chaînes étaient présentées sur l'écran de sorte que les participants fixaient la première ou la dernière lettre d'une chaîne de 5 lettres (ex: xxxax ou xxxax). Les résultats montrent que l'effet du champ visuel sur la perception des lettres est différent pour les lecteurs anglais et hébreux. Plus précisément, les lecteurs de l'alphabet roman, dont la direction de lecture est de gauche à droite, sont plus précis lorsque la chaîne de caractères est présentée dans le CVD, en accord avec les hypothèses de dominance hémisphérique et de direction de lecture ; cependant, chez les lecteurs hébreux, dont la direction de lecture est de droite à gauche, cet avantage en faveur du CVD disparaît bien qu'il ne soit pas compensé par un avantage pour le CVG (voir Figure 9).

Figure 9. Scores de reconnaissance de lettres cibles contenues dans des séquences de 5 lettres romanes (à gauche) et hébreux (à droite) présentées brièvement et fixées sur la première lettre (présentation dans le champ visuel droit) ou sur la dernière lettre (présentation dans le champ visuel gauche). Les symboles pleins représentent les

performances dans le champ visuel droit et les symboles vides, les performances dans le champ visuel gauche (Nazir, et al., 2004, Expérience 3).

La disparition de l'asymétrie chez les lecteurs hébreux semble suggérer que la lisibilité des lettres dans les champs visuels est modulée par la combinaison de la dominance hémisphérique gauche pour le langage, et les habitudes liées à la direction de la lecture (voir Whitney, 2001, pour une proposition similaire).

3-4 Un modèle mathématique « visuel » (Nazir et al, 1991 ; 1998)

3-4-1 Probabilité de reconnaissance d'un mot

Afin de rendre compte de l'influence de l'asymétrie de la lisibilité des lettres sur la reconnaissance des mots, Nazir et collaborateurs (Nazir, et al., 1991 ; 1998) ont développé un modèle mathématique *probabiliste* basé sur des mesures empiriques de lisibilité des lettres (cf paragraphe 3-3, voir aussi McConkie et al., 1989 pour une proposition similaire). Ce modèle repose sur trois hypothèses. 1) La contribution d'une lettre à la reconnaissance d'un mot est proportionnelle à sa lisibilité. 2) Les lettres d'un mot sont reconnues indépendamment les unes des autres. 3) La probabilité de reconnaître un mot est fonction de la probabilité d'identifier toutes ses lettres. La probabilité de reconnaître un mot en fonction de sa longueur, et de la lettre fixée peut être estimée par l'équation (1) suivante:

où $P_{word}(f, a, l, bl, br)$ est la probabilité de reconnaître un mot en fonction de f , la position relative de la lettre fixée dans la séquence, a , la probabilité d'identifier la lettre directement fixée, bl et br , les taux de diminution de lisibilité des lettres en fonction de l'excentricité à gauche et à droite de la fixation respectivement et l , la longueur du mot (Nazir et al., 1998). La Table 1 présente la probabilité de reconnaître une séquence de 5 lettres, pour un taux de diminution à droite br arbitrairement fixé à .06.

La probabilité d'identifier la lettre directement fixée (a) prend la valeur de 1. Considérant la diminution linéaire de la lisibilité des lettres avec l'excentricité, la probabilité d'identifier les lettres chute d'une valeur constante à mesure que l'on s'écarte à droite (br) ou à gauche (bl) de la lettre fixée. Toutefois, le ratio de l'asymétrie bl/br a été estimé à 1.8 (cf paragraphe précédent, Nazir et al., 1991), autrement dit $bl = br * 1.8$. La probabilité de reconnaître la séquence entière pour une position du regard donnée est obtenue en multipliant les probabilités d'identification des lettres individuelles constitutives. Par exemple, la probabilité de reconnaître la séquence entière lorsque la 3ème lettre est fixée est égale à $0.78 * 0.89 * 1 * 0.94 * 0.88 = 0.57$. La dernière colonne de la Table 1 donne la probabilité de reconnaître une séquence de 5 lettres en fonction de

la lettre fixée (présentée dans la Figure 10).

Figure 10. Probabilité théorique (trait plein) et empirique (trait pointillé) de reconnaître un mot de 5 lettres en fonction de la position du regard dans le mot. La courbe empirique représente les données obtenues pour des mots de 5 lettres dans une tâche de décision lexicale (d'après Nazir et al., 1991).

Comme on peut le voir sur la Figure 10, la forme de la courbe de l'EPR prédite par le modèle de « multiplication des probabilités des lettres individuelles » (MPLI par la suite) est fortement similaire à la forme de la courbe de l'EPR obtenue empiriquement (Nazir et al., 1991). La probabilité la plus grande de reconnaître un mot est obtenue lorsque le regard se pose légèrement à la gauche du centre du mot et décroît progressivement à mesure que les yeux s'éloignent de cette position optimale.

Néanmoins, en dépit du fait que ce modèle rend compte de l'asymétrie de l'EPR dans les mots par l'asymétrie de la lisibilité des lettres dans les deux champs visuels, le MPLI rencontre plusieurs difficultés qui résident essentiellement dans l'absence de la prise en considération de l'influence des facteurs lexicaux sur la reconnaissance des mots.

3-4-2 Limites du modèle MPLI

La principale limite du modèle MPLI concerne l'ajustement de la *hauteur* et de la *forme* de la courbe théorique à la courbe obtenue empiriquement. Bien que la forme des courbes théorique et empirique de la fonction de l'EPR soit fortement similaire (cf Figure 10), la courbe calculée selon le MPLI se trouve largement en dessous de la courbe empirique. Compte tenu du fait que les estimations théoriques étaient entièrement basées sur des mesures empiriques de lisibilité des lettres, la différence de hauteur constatée entre les deux courbes a été interprétée comme la conséquence de l'intervention de « **facteurs lexicaux** » (Nazir et al., 1991). Plus précisément, Nazir et al. (1991, p.173) ont suggéré que le taux de diminution br/bl ne serait pas transposable directement aux mesures de reconnaissance des mots empiriques, pour lesquelles les « *conditions visuelles sont favorables* » (ces « conditions » n'ont pas été précisées par ces auteurs) ; la diminution de lisibilité des lettres serait ainsi moins « importante » lors de la reconnaissance des mots.

Dans le MPLI, la probabilité de reconnaître la lettre directement fixée ($a = 1$) et le ratio de l'asymétrie ($bl/br = 1.8$) sont des paramètres fixes ; le taux de diminution à droite (br) constitue par conséquent le seul paramètre libre du modèle permettant d'ajuster la courbe théorique aux données empiriques. La Figure 11 présente différentes courbes théoriques calculées pour différentes valeurs de br .

Figure 11. Probabilités théoriques d'identifier un mot de 5 lettres en fonction de la position du regard dans le mot pour différentes valeurs du taux de diminution de lisibilité des lettres br (le rapport d'asymétrie gauche/droite est maintenu constant à 1.8/1). Plus le taux de diminution (br) diminue, plus la courbe est élevée (augmentation de la hauteur). La courbe en ligne pleine représente les données empiriques (Nazir et al., 1991).

En conservant le rapport de l'asymétrie gauche/droite à 1.8/1 mais en diminuant

théoriquement la chute de lisibilité des lettres (via le paramètre br), la courbe théorique « gagne » en hauteur. Afin d'obtenir la courbe théorique qui ajuste le mieux les données empiriques, on peut calculer la « somme des carrés des erreurs » ou RMSD (« Root mean square deviation », voir par exemple Nazir et al., 1998 pour l'utilisation de cette même procédure et l'annexe 1) ; plus cette somme est petite, meilleur est l'ajustement. Une valeur minimale du RMSD a été obtenue pour $br = .024$ (RMSD = .05).

Toutefois, comme l'indique la Figure 11, à mesure que la courbe théorique approche la hauteur de la courbe empirique, la forme de la courbe ainsi prédite par le modèle MPLI « s'aplatit » (i.e., l'EPR est moins prononcé) et ne correspond plus alors à la forme de la courbe obtenue empiriquement, laquelle est par ailleurs mieux prédite pour $br = .06$.

En résumé, la prise en considération de la variation de la lisibilité des lettres en fonction de la position du regard a permis de montrer que l'EPR dans le mot ne peut se réduire à la seule contribution des facteurs liés à la contrainte lexicale des mots mais tirerait son origine de la combinaison de plusieurs facteurs. Un premier pas vers la quantification des effets intra-lexicaux de la contrainte lexicale et de la lisibilité des lettres sur la reconnaissance visuelle des mots a été réalisé grâce au développement de modèles dits « visuo-lexicaux ».

Outre le fait que ces modèles postulent tous l'existence d'une interaction entre les facteurs visuels et les facteurs lexicaux, l'entremise entre les deux diffère selon les auteurs (Clark & O'Regan, 1999 ; Kajii, 2000 ; Kajii & Osaka, 2000 ; Stevens & Grainger, 2003).

Chapitre 4 Effet de la position du regard et modèles mathématiques

4-1 Le modèle de Clark et O'Regan (1999)

Sur la base d'une hypothèse simplifiée concernant les lettres qui sont identifiées lorsque le regard se pose à différentes positions dans le mot, Clark et O'Regan (1999) ont défini une « **mesure d'ambiguïté** » qui fournit une estimation du nombre de mots compatibles avec l'information visuelle disponible à partir d'une fixation pour chaque position du regard dans le mot.

4-1-1 Estimation de l'information visuelle

En raison de la diminution de l'acuité visuelle, Clark et O'Regan (1999) supposent que seules les deux lettres les plus proches de la fixation du regard sont correctement identifiées pendant une seule fixation. Toutefois, étant donné que les deux lettres externes sont moins sujettes au phénomène de masquage latéral, elles sont également perçues parfaitement ; toutes les autres lettres ne peuvent être identifiées. Ainsi par exemple, seules les lettres « C », « H », « A » et « E » sont visibles lorsque le regard se

pose entre la deuxième et troisième lettre du mot « CHAISE ».

Ce schéma simplifié des lettres visibles va ainsi être utilisé afin de déterminer la nature des stimuli interférents avec la reconnaissance du mot cible.

4-1-2 Mesure d'ambiguïté du mot

Sur la base des quatre lettres les plus visibles, ces auteurs ont calculé statistiquement, à l'aide de dictionnaires, une mesure de l'ambiguïté lexicale d'un mot, définie comme le nombre de mots qui possèdent les mêmes 4 lettres visibles à ces positions. Par exemple, lorsque le mot anglais de 9 lettres « SCATTERED » est fixé entre la 3^e et la 4^e lettre (i.e., entre le « A » et le « T »), le patron visuel formé est la séquence suivante « S*AT**D », où « * » indique les lettres non reconnues. L'ambiguïté est estimée par la moyenne du nombre d'entrées compatibles avec les lettres identifiées. La longueur du mot est supposée connue (voir la Figure 12 pour une illustration).

Figure 12. Illustration théorique du nombre de mots compatibles avec l'information visuelle disponible lorsqu'un sujet fixe le mot anglais de 13 lettres « UNDERGRADUATE », A) vers le milieu du mot (i.e., entre la 5^e et la 6^e lettre) et, B) vers le début du mot (i.e., entre la 3^e et la 4^e lettre) (Figure tirée de Clark et O'Regan, 1999).

L'analyse statistique à partir de bases de données de mots anglais et français a ainsi révélé que la position optimale du regard (OVP) est proche de la position qui minimise l'ambiguïté du mot à partir de la reconnaissance incomplète du mot. Plus précisément, lorsque la fixation est légèrement à gauche du centre du mot, très peu de mots sont compatibles avec les 4 lettres les plus visibles ; cependant lorsque la fixation se situe vers les extrémités du mot, l'ambiguïté augmente car de plus en plus de mots partagent le même ensemble de lettres visibles (voir Figure 13). L'analyse théorique menée par Clark et O'Regan (1999) n'a cependant pas été testée empiriquement.

4-2 Le modèle de Stevens et Grainger (2003)

Stevens et Grainger (2003) ont combiné les données empiriques de lisibilité des lettres avec une mesure statistique de la redondance orthographique afin de rendre compte de l'EPR dans les mots.

4-2-1 Données empiriques de lisibilité des lettres

La lisibilité des lettres du mot a été définie par Stevens et Grainger (2003) à partir de mesures empiriques complètes de perception des lettres. Ces auteurs ont mesuré les variations de lisibilité des lettres contenues dans des séquences de « X » de 5 et 7 lettres à différentes positions sur la rétine. De plus, les séquences étaient fixées sur toutes les

positions de lettres et pas uniquement sur la première et la dernière lettre, comme précédemment décrit dans l'étude de Nazir et al. (1991 ; 1998, cf paragraphe 3-3).

La Figure 14 illustre les pourcentages d'identifications de lettres correctes pour les deux longueurs testées (5 et 7 lettres) en fonction de la position du regard dans la séquence. De manière surprenante, les courbes de position du regard ainsi obtenues adoptent la forme d'un « U » inversé parfaitement symétriques. L'implication de ces résultats sera abordée par la suite.

Figure 14. Pourcentages de reconnaissance des lettres contenues dans des séquences de 5 et 7 lettres en fonction de la position du regard relativement au centre de la séquence. Les données sont moyennées à travers toutes les positions des lettres cibles (données tirées de Stevens & Grainger, 2003, expérience 1).

Les scores d'identifications de lettres correctes pour chaque position du regard sont combinés à une mesure de la fréquence positionnelle des lettres composant les mots afin de prédire l'EPR dans les mots.

4-2-2 Fréquence positionnelle des lettres

Cette analyse a été motivée suite aux observations de Grainger et Jacobs (1993). Ces derniers ont observé que la fréquence positionnelle des lettres était un bon indicateur des effets d'amorçage orthographique. La fréquence positionnelle d'une lettre correspond à la somme des fréquences de tous les mots où cette lettre est apparue à une position donnée. Ainsi, moins la fréquence positionnelle des lettres de l'amorce est élevée, plus grands sont les effets de facilitation de l'amorçage. Par exemple, Grainger et Jacobs (1993) ont montré que les 4 lettres de l'amorce « TABL% » donnent lieu à des latences de décision lexicale plus courtes que les amorces « %ABLE » ou même les amorces « TA%LE », ayant une fréquence positionnelle des lettres relativement élevée.

Deux types de fréquence positionnelle des lettres ont été testés. (1) La fréquence positionnelle absolue fait référence au nombre de fois où une lettre apparaît à une position donnée dans une séquence de longueur fixe. (2) La fréquence positionnelle des lettres relative est calculée indépendamment de la longueur de la séquence, utilisant le schéma de codage suivant : première lettre, dernière lettre et lettres internes. Il s'agit de la somme des fréquences des mots de 3 à 10 lettres où la première lettre apparaît en première position; la dernière lettre en dernière position et les lettres internes en positions internes indépendamment de leurs positions absolues.

La probabilité de reconnaître un mot ($P(W)$) est estimée selon l'équation suivante :

où $W(freq)$ est la fréquence d'usage du mot, $Fpos(Lp)$ est la fréquence positionnelle de la lettre à la position p et, $Lisib(Lp)$ est la mesure de lisibilité empirique obtenue pour cette lettre. Le facteur lexical est estimé dans ce modèle par la probabilité qu'une lettre donnée appartienne au stimulus mot donné, (i.e., $[W(freq) / Fpos(Lp)]$).

Stevens et Grainger (2003) ont montré que la combinaison théorique des mesures

empiriques de lisibilité des lettres avec une mesure de la fréquence positionnelle des lettres utilisant un codage absolu de la position des lettres ne rend pas compte du patron de l'EPR dans les mots obtenu empiriquement ; en revanche, la mesure de la contrainte lexicale basée sur un codage de la position des lettres relatif génère des prédictions similaires aux données empiriques (voir Figure 15).

Figure 15. Probabilités théoriques (lignes pointillées) et empiriques (lignes pleines) de reconnaître un mot de 5 lettres (panel de gauche) et de 7 lettres (panel de droite) (données tirées de Stevens & Grainger, 2003).

4-3 Une proposition de Kajii (2000)

Selon une proposition de Kajii (2000) (voir aussi Kajii & Osaka, 2000), une extension du MPLI (Nazir et al., 1991 ; 1998) a été avancée afin de rendre compte de l'influence des facteurs visuels et des facteurs lexicaux sur la détermination de l'EPR dans les mots.

4-3-1 Étapes menant à la reconnaissance d'un mot.

Le modèle de Kajii (2000) a été construit pour simuler l'EPR dans les mots au cours d'une tâche d'identification perceptive où le participant doit identifier précisément un mot présenté très brièvement. L'information visuelle ainsi « limitée » favorise alors les erreurs d'identifications et les réponses basées sur le « *devinement* ». Parallèlement aux modèles classiques de la reconnaissance visuelle d'un mot (e.g., McClelland et Rumelhart, 1981 ; Grainger & Jacobs, 1996), Kajii (2000) suppose que l'information sensorielle issue de la présentation visuelle d'un mot implique trois étapes de traitement. La première étape conduit à la formation d'un patron visuel, qui code l'identité des lettres abstraite ainsi que la position des lettres, appelé « *CLIP* » (« codes for letter identity and position in word »). Dans une seconde étape, le CLIP est combiné à l'information lexicale afin de prédire le mot. Enfin dans la dernière étape, le mot est identifié.

Ce modèle que nous appellerons le *modèle CLIP* adopte l'hypothèse simplifiée que seule l'orthographe est nécessaire pour identifier un mot, bien que les processus phonologiques soient actifs dès les premières étapes de traitements (voir aussi Grainger & Jacobs, 1996).

4-3-2 Probabilité d'identifier un mot.

La probabilité d'identifier un mot dépend de deux conditions : (1) Conformément au MPLI décrit plus haut (Nazir et al., 1991 ; 1998), dans le modèle CLIP, l'identification d'un mot a lieu lorsque toutes les lettres du mot ont été identifiées avec succès. Cependant un mot peut aussi être reconnu à partir d'une information sensorielle partielle via l'aide des connaissances lexicales (voir par exemple Brysbaert et al., 1996 ; Clark & O'Regan,

1999 ; McClelland & Rumelhart, 1981 ; O'Regan et al, 1984 ; Paap et al., 1982). Par conséquent, le modèle CLIP suppose qu'un mot peut également être reconnu par une procédure d'*inférence* à partir du codage de certaines de ses lettres. (2) Par ailleurs, bien que l'identité d'une lettre peut être récupérée sans que sa position soit nécessairement codée, le modèle CLIP fait l'hypothèse simplifiée, dans un premier temps, que lorsqu'une lettre est identifiée, sa position est également codée (voir aussi McClelland & Rumelhart, 1981 pour une proposition similaire).

4-3-3 Codage de l'identité et la position des lettres.

*Figure 16. Illustration des différentes possibilités théoriques de CLIPs visuels pour mot de 5 lettres («L1L2L3L4L5»). Les chiffres indiquent le nombre de lettres correctement identifiées et le symbole « * » représente la lettre absente (Kajii, 2000).*

Comme on peut le voir sur la Figure 16, le CLIP « 5 » ou CLIP « total » est le cas où toutes les lettres du mot sont correctement identifiées («L1L2L3L4L5»). Le CLIP « 0 » ou CLIP « nul » (« ***** ») représente le cas où aucune lettre du mot n'a été identifiée. Les autres CLIPs, les CLIPs « partiels » possèdent au moins une lettre incorrecte (représentée par le symbole « * »), laquelle peut être différente de la lettre cible ou bien absente. Par simplification, on suppose ici que la lettre est manquante. Ainsi pour un stimulus de 5 lettres, on a 1 CLIP « total », 1 CLIP « nul » et 30 CLIPs « partiels ».

4-4 Le modèle CLIP

Kajii (2000) suppose que lorsqu'un stimulus mot est présenté brièvement dans un contexte d'identification perceptive, deux événements peuvent prédire l'identité du mot.

Événement 1. Le participant identifie toutes les lettres du mot (CLIP total) et la probabilité de prédire le mot à partir de ce CLIP est de 100% , on parle d' **identification perceptive totale** .

Événement 2. Le participant identifie un CLIP partiel, dans ce cas, il doit faire appel à ses connaissances lexicales afin de prédire le mot, on parle alors d' **inférence lexicale** .

De manière logique, le mot ne peut pas être prédit à partir du CLIP nul. Par ailleurs, il est hautement improbable de prédire un mot à partir du codage d'une seule lettre.

Ainsi, la probabilité de reconnaître un mot, $P(W)$ peut être estimée par l'équation (2) suivante:

où $P(C_i)$ est la probabilité d'occurrence du CLIP C_i , (parmi l'ensemble des n CLIPs possibles, un seul CLIP peut avoir lieu à la fois), $P(W|C_i)$ est la probabilité conditionnelle⁹ de prédire le mot complet lorsque le CLIP C_i a été codé.

L'équation (2) possède deux composantes: une composante *visuelle*, $P(C_i)$, i.e., la

probabilité d'occurrence d'un CLIP C_i , et une composante *lexicale*, $P(W|C_i)$, i.e., la probabilité de prédire le mot à partir du CLIP C_i .

4-4-1 Probabilité d'occurrence d'un CLIP (C_i)

Conformément au MPLI décrit précédemment (Nazir et al., 1991), la probabilité d'occurrence d'un CLIP est égale à la probabilité d'identifier les lettres (L_i) qui le composent. Par exemple, la probabilité de reconnaître toutes les lettres d'un mot de 5 lettres («L1L2L3L4L5»), c'est-à-dire le CLIP total, est égale à :

où $P(L_i)$ est la probabilité d'identifier la i ème lettre du mot (il s'agit précisément du MPLI proposé par Nazir et al., 1991 ; 1998).

La probabilité d'occurrence du CLIP lorsque la dernière lettre (L_5) n'a pas été identifiée («L1L2L3L4*») est égale à la probabilité d'identifier les 4 premières lettres et la probabilité de ne pas reconnaître la dernière lettre (L_5), soit :

La probabilité d'occurrence du CLIP lorsque la première lettre n'a pas été identifiée («*L2L3L4L5») est égale à :

4-4-2 Probabilité d'identifier une lettre (L_i)

Comme nous l'avons vu dans le paragraphe 3-3, la lisibilité d'une lettre dépend non seulement de sa position dans le champ visuel (gauche vs droit) mais aussi de sa position dans la séquence (lettres externes vs lettres internes). Ainsi, conformément à Nazir et al. (1991 ; 1998), Kajji (2000 ; voir aussi Kajji & Osaka, 2000) suppose que la lisibilité d'une lettre interne diminue plus vite dans le champ visuel gauche que dans le champ visuel droit (rapport d'asymétrie gauche droite maintenu à 1.8). Ainsi, la probabilité d'identifier une lettre interne donnée L_i est estimée par l'équation (4) suivante (cf paragraphe 4-1) :

La probabilité d'identifier une lettre est égale à a lorsque la lettre se trouve à l'emplacement du point de fixation f , et diminue selon un taux de br par lettre à droite de la lettre fixée et bl par lettre à gauche de la lettre fixée.

Par ailleurs, le modèle CLIP tient compte de la réduction du masquage latéral des lettres externes. Ainsi, la probabilité d'identifier une lettre externe est identique à la probabilité d'identifier la lettre immédiatement adjacente à gauche ou à droite (Kajji et Osaka, 2000). La probabilité d'occurrence d'un CLIP dépend par conséquent des seules contraintes visuelles liées à la lisibilité des lettres.

⁹ La probabilité conditionnelle de B étant donné A, notée $P(B|A)$ est définie comme la probabilité que l'événement B se produise étant donné que l'événement A s'est produit.

4-4-3 Contrainte lexicale et probabilité d'identifier un mot ($P(W|C_i)$)

La probabilité de prédire le mot, $P(W|C_i)$ lorsqu'un CLIP C_i a été identifié est fonction du nombre de candidats lexicaux activés, c'est-à-dire du nombre de mots partageant les mêmes lettres que le CLIP. Ainsi, plus le nombre de « candidats » est élevé, plus il sera difficile d'inférer le mot cible. La probabilité de prédire un mot à partir d'un CLIP donné (C_i) peut ainsi être estimée suivant l'équation (5):

où $P(W|C_i)$ est la probabilité de prédire le mot complet à partir du CLIP C_i , et $N(C_i)$ le nombre de mots qui partagent les mêmes lettres que le CLIP partiel (« voisins CLIP » par la suite). Le nombre de « voisins CLIP » doit être différent de 0. Lorsqu'il n'y a aucun « voisin CLIP », $N(C_i) = 1$, et la probabilité d'inférer le mot sera de 100%, c'est précisément le cas lorsque toutes les lettres du mot sont identifiées (CLIP total) ou lorsque le seul mot compatible avec le stimulus est le mot cible lui-même.

Le modèle CLIP est un modèle probabiliste construit pour rendre compte des réponses correctes des stimuli mots en fonction de la position du regard dans la tâche d'identification perceptive. Le MPLI décrit précédemment est un cas particulier du modèle CLIP, où le mot est identifié via l'identification de toutes ses lettres, c'est-à-dire le CLIP total. Ce processus de reconnaissance du mot que nous avons appelé « identification perceptive totale » est prédictible entièrement à partir de l'information visuelle. Le modèle CLIP considère par ailleurs les situations où le lecteur n'a pas identifié toutes les lettres du mot (CLIP partiel) mais peut quand même reconnaître le mot par un processus d'inférence lexicale. Ce processus, à l'inverse de l'identification perceptive vraie, dépend des contraintes lexicales.

Afin de démontrer la pertinence du modèle CLIP, nous allons dans la partie qui suit, modéliser les données empiriques obtenues par Brysbaert et al (1996) comme cela a été proposé par Kajii & Osaka (2000). L'étude de Brysbaert et al (1996) fournit des mesures empiriques de la probabilité de prédire un mot à partir de l'identification de quelques lettres et permet ainsi de tester directement le modèle CLIP.

4-5 Évaluation du modèle CLIP (Kajii & Osaka, 2000)

Brysbaert et collaborateurs (1996 ; Expérience 3) ont manipulé la position du regard dans des mots de 5 lettres à début informatif et à fin informative au cours d'une tâche d'identification perceptive (cf. chapitre 3). La variable mesurée était le pourcentage de mots correctement identifiés.

4-5-1 Probabilité d'occurrence des CLIPs

Table 2. Probabilités d'occurrences des CLIPs (moyenne supérieure à .01) pour un mot de 5 lettres en fonction de la position de la lettre fixée. La probabilité d'identifier la lettre directement fixée, a est égale à 1. Cette probabilité diminue linéairement d'une valeur arbitraire de .06 pour chaque lettre à droite de la lettre fixée (br) et de $.06 \times 1.8 = .11$ à gauche (bl).

4-5-2 Mesure de la contrainte lexicale

Dans l'expérience de Brysbaert et al. (1996), l'estimation de la contrainte lexicale a été obtenue par des mesures subjectives de prédiction des mots à l'issue d'une tâche de complétion de mots. 84% des mots à début informatif étaient prédictibles à partir du trigramme initial et 9% seulement à partir du trigramme final. À l'inverse, 8% des mots à fin informative étaient prédictibles à partir du trigramme initial et 71% à partir du trigramme final.

Autrement dit, la probabilité d'inférer les mots à début informatif lorsque le trigramme initial ou le trigramme final ont été codés respectivement est égale à :

De même, la probabilité d'inférer les mots à fin informative lorsque le trigramme initial ou le trigramme final ont été codés respectivement est égale à :

4-5-3 Reconnaissance du mot et position du regard

La probabilité théorique de reconnaître les mots à début informatif et les mots à fin informative a été calculée à partir des 7 CLIPs suivants: « L1L2L3L4L5 », « L1L2L3** », « **L3L4L5 », « *L2L3L4L5 », « L1*L3L4L5 », « L1L2L3L4* », « L1L2L3*L5 ». Le choix de ces CLIPs a été motivé par deux raisons. Premièrement, la probabilité d'occurrence de ces CLIPs est relativement élevée par rapport aux autres CLIPs, dont la contribution à la reconnaissance du mot a été considérée comme négligeable. Deuxièmement, les mesures empiriques de prédiction des mots disponibles dans cette expérience étaient limitées.

La probabilité d'inférer un mot lorsque toutes les lettres du mot ont été identifiées est de 100%. Par ailleurs, le nombre de mots pouvant être inféré à partir des CLIPs « L1L2L3** », « L1L2L3L4* » et « L1L2L3*L5 » a été considéré comme identique (Kajii, 2000). Par conséquent, les probabilités de prédire les mots à partir de ces CLIPs sont équivalentes, soit :

De même, les probabilités de prédire les mots à partir des CLIPs « **L3L4L5 », « *L2L3L4L5 », « L1*L3L4L5 » sont équivalentes, soit :

Les pourcentages de réponses correctes calculés pour les mots à début informatif et à fin informative en fonction de la position de la lettre fixée sont présentés dans la Figure 17. Afin d'obtenir la meilleure courbe théorique, le paramètre br a été ajusté afin d'obtenir un RMSD minimum (cf paragraphe 3-4). Ainsi, pour br = .05, les RMSD sont égales à .02 et .01 pour les mots à début informatif et à fin informative respectivement.

Figure 17. Probabilités de reconnaître un mot à début informatif (à gauche) et un mot à fin informative (à droite) en fonction de la position de la lettre fixée. Les lignes pointillées (symboles vides) représentent les données empiriques, les traits pleins (symboles pleins) représentent les données calculées d'après le modèle CLIP; enfin les lignes pointillées (symboles pleins) représentent la probabilité de coder le CLIP total (i.e., MPLI, Nazir et al., 1991).

Comme on peut le voir sur la figure, les courbes calculées théoriquement correspondent aux données obtenues empiriquement par Brysbaert et al. (1996). Ainsi, en ajoutant la contrainte lexicale à la composante visuelle (CLIP total), le modèle CLIP prédit de manière relativement précise les données empiriques. Autrement dit, les mesures de lisibilité des lettres utilisées en combinaison avec des estimations empiriques de la contrainte lexicale prédisent avec une grande justesse les données empiriques de Brysbaert et al. (1996).

En bref, le modèle CLIP (Kajii, 2000) de type visuo-lexical a été développé afin de rendre compte de l'EPR dans le mot dans la tâche d'identification perceptive. Il donne la probabilité de reconnaître un stimulus mot en fonction de la position du regard dans le mot. Selon ce modèle, la reconnaissance d'un mot a lieu en trois étapes. Dans une première étape, le participant identifie un CLIP visuel à partir de l'information sensorielle extraite du stimulus mot, ce CLIP peut être total (i.e., le participant a identifié toutes les lettres du stimulus) ou partiel (i.e., le participant n'a identifié que quelques lettres). Dans une seconde étape, le CLIP va « activer » un ensemble de candidats lexicaux (« voisins CLIP »), lesquels vont interférer avec la reconnaissance du stimulus mot. Selon le postulat que lorsqu'une lettre est identifiée, sa position est automatiquement codée (Kajii, 2000), ce modèle suppose qu'aucun mot ne vient interférer avec la reconnaissance du stimulus lorsque le CLIP identifié est total. Enfin, dans la dernière étape, le mot est identifié.

En résumé, les modèles de type « visuo-lexicaux » décrits dans le chapitre 4 sont des modèles de perception de lettres et de mots élaborés afin de rendre compte de l'EPR dans les mots dans le contexte de la tâche d'identification perceptive. Ces modèles sont inspirés des cadres théoriques classiques présentés plus haut (e.g., McClelland & Rumelhart, 1981 ; Paap et al., 1982) décrivant les processus mis en œuvre pour reconnaître un mot écrit.

En particulier, les mots sont reconnus sur la base de leurs lettres selon une identification totale (Stevens & Grainger, 2003), partielle (Clark & O'Regan, 1999) ou les deux (Kajii, 2000). Ces modèles prévoient explicitement un processus de compétition entre les mots lors de la reconnaissance visuelle des mots. Par ailleurs, dans les modèles de Clark et O'Regan (1999) et Kajii, (2000), la longueur du mot est codée et les lettres du mot le sont en fonction de leur position absolue à l'intérieur du mot. Seules les représentations des mots de même longueur compatibles avec les lettres identifiées s'activent lors de la présentation du stimulus visuel. En revanche, le modèle de Stevens et Grainger (2003) adopte un schéma de codage relatif de la position des lettres, indépendamment de la longueur des mots, lequel est par ailleurs davantage conforme aux données actuelles de la littérature. Ainsi, dans ce modèle, tous les mots ne

partageant même qu'une seule lettre avec le stimulus s'activent lors de la perception d'un mot.

Plusieurs études ont souligné une asymétrie de la lisibilité des lettres avec de meilleures performances de perception des lettres lorsque les séquences sont présentées dans le CVD que dans le CVG (Bouma, 1973 ; Kajii & Osaka, 2000; Nazir et al., 1991 ; 1992 ; 1998). Cette variable n'est néanmoins pas pris en compte dans les modèles de Stevens et Grainger (2003) et Clark et O'Regan (1999), considérant uniquement la diminution de l'acuité visuelle et la réduction du masquage latéral sur la perception des lettres externes.

Afin de mieux souligner l'influence de l'asymétrie de la lisibilité des lettres sur les processus de reconnaissance des mots, nous avons réalisé une simulation des données empiriques de Brysbaert et al. (1996) via le modèle CLIP en supprimant l'asymétrie droite/gauche du modèle. La démarche suivie était en tous points similaire à celle réalisée au paragraphe 4-5, à l'exception de la valeur du paramètre bl qui était cette fois égale à celui de br ($br = bl = .05$). Les données sont présentées dans la Figure 18.

Figure 18. Probabilités de reconnaître un mot à début informatif (à gauche) et à fin informative (à droite) en fonction de la position de la lettre fixée. Les lignes pointillées (symboles vides) représentent les données empiriques, les traits pleins (symboles pleins) représentent les données calculées d'après le modèle CLIP, enfin les lignes pointillées (symboles pleins) représentent la probabilité de coder le CLIP total lorsque $br = bl$.

Les données théoriques ainsi obtenues prédisent correctement de meilleures performances de reconnaissance des mots à début informatif lorsque les fixations sont sur le début des mots que lorsqu'elles se situent sur la fin des mots. En revanche, lorsque les mots sont à fin informative, l'absence d'asymétrie de la lisibilité des lettres prédit incorrectement un avantage lorsque le regard se pose sur la fin des mots que lorsqu'il se situe sur le début des mots, allant par conséquent à l'encontre des hypothèses des modèles de Stevens et Grainger (2003) et Clark et O'Regan (1999).

4-5-4 Objectifs de la contribution expérimentale

La partie expérimentale qui suit a pour objectif de tester les prédictions du modèle CLIP (Kajii, 2000) afin de mieux rendre compte du phénomène de l'EPR au cours de la reconnaissance visuelle des mots. Étant donné les limitations, précédemment mentionnées, des modèles de Stevens et Grainger (2003) et de Clark et O'Regan (1999), seul le modèle CLIP (Kajii, 2000) nous a guidé dans l'élaboration de nos hypothèses de travail et dans la mise en place de notre cadre expérimental.

Le modèle CLIP postule deux moyens par lesquels l'identification d'un mot peut avoir lieu : Par une identification perceptive totale de toutes les lettres constitutives du mot (CLIP total) ou par une inférence lexicale à partir de l'identification de quelques lettres du mot (CLIPs partiels).

Dans le cadre de ce présent travail, nous allons poursuivre l'évaluation de ce modèle. Dans ce contexte, nous allons tester la « réalité » des processus d'identification

perceptive totale et d'inférence lexicale, impliqués dans la reconnaissance visuelle des mots. Dans ce but, une tâche d'identification perceptive de mots et de pseudomots, combinée au paradigme de la position variable du regard, a été conduite. La logique de cette approche prévoit que l'identification correcte d'un pseudomot n'est possible que par la procédure d'identification totale, l'inférence lexicale mettant en jeu la récupération de la représentation lexicale du stimulus. Néanmoins, de la même manière que les mots, les CLIPs partiels identifiés à partir des pseudomots activent également des candidats lexicaux (e.g., Coltheart, Curtis, Atkins, & Haller, 1993 ; Coltheart, Rastle, Perry, Langdon, & Ziegler 2001 ; McCann & Besner, 1987) et tendraient à donner lieu à des erreurs de « lexicalisations ».

La comparaison des processus impliqués dans l'EPR, dans une tâche d'identification perceptive et dans une tâche de décision lexicale, permettra, en outre, d'avancer des hypothèses sur la spécificité des processus concernés.

Chapitre 5 EPR, lettres, position et lexique

Dans le but de tester l'hypothèse selon laquelle l'EPR dans le mot peut être considéré comme le résultat de l'identification perceptive totale et de l'inférence lexicale, des simulations des données obtenues expérimentalement, via le modèle CLIP, ont été réalisées. À cet égard, les données empiriques obtenues pour les pseudomots ont été modélisées via la procédure d'identification perceptive totale uniquement. En revanche, les deux procédures, identification perceptive totale et inférence lexicale, ont été utilisées pour simuler les données empiriques obtenues pour les mots. Par ailleurs, la mesure de la contrainte lexicale, initialement établit à partir de mesures empiriques basées sur une tâche de complétion de mots (Brysbaert et al., 1996), sera modélisée en tenant compte des données actuelles de la littérature.

5-1 Identification perceptive de mots et de pseudomots : Expérience 1

La tâche du participant était d'identifier précisément des mots et des pseudomots présentés brièvement selon le paradigme de la position variable du regard.

5-1-1 Méthode

Participants. Vingt participants droitiers, volontaires et de langue maternelle française ont participé à cette expérience. Tous avaient une vue normale ou corrigée. Les participants n'ont pas été rétribués pour leur participation.

Stimuli. Pour cette expérience, nous avons sélectionné ou construit 100 stimuli de 6 lettres: 50 mots français et 50 pseudomots prononçables et orthographiquement légaux. Les stimuli ont été sélectionnés à partir de la base de données lexicale informatisée "Brulex" (Content, Mousty & Radeau, 1990). Les mots ont une fréquence d'usage moyenne de 54.3 occurrences par million. La fréquence lexicale indiquée dans la base Brulex est reprise des tables publiées par le Centre de recherche pour un Trésor de la Langue Française (Imbs, 1971). Les pseudomots ont été construits en remplaçant, de façon pseudo aléatoire, une seule lettre dans des "mots de base", différents des « mots tests » utilisés pour l'expérience proprement dite. Les pseudomots respectaient les formes orthographiques légales de la langue française ; ils possédaient une fréquence des bigrammes moyenne similaire à celle des mots de base et des mots test (2.95, 2.92 et 2.97 occurrences par million respectivement pour les pseudomots, les mots de base et les mots test)¹⁰. Ces fréquences ont été calculées à partir des moyennes des logarithmes décimaux des fréquences textuelles en position initiale, interne et finale (établies sur la base d'un comptage pondéré par les fréquences d'usage des mots de tous les trigrammes constituant le mot (Content & Radeau, 1988). Tous les stimuli de l'expérience 1 sont présentés dans l'Annexe 3.

Dispositif expérimental. L'expérience était pilotée à l'aide d'un Power Macintosh 6500/250. Le script de l'expérience a été réalisé avec le logiciel Psyscope.1.2.2.PPC. Les stimuli étaient présentés à l'écran en lettres minuscules (taille 14, police Courier). À une distance de 57.3 cm, chaque lettre correspondait à 0.25 degrés d'angle visuel.

Plan expérimental. Pour le calcul des cinq positions de présentation sur l'écran, la longueur du stimulus a été divisée en cinq parties strictement égales. Le stimulus était présenté au niveau du centre de chacune de ces parties. Selon la technique de la position variable du regard précédemment décrite, le stimulus est déplacé latéralement par rapport au point de fixation (toujours au centre de l'écran) de sorte que chacune des 5 zones du stimulus puisse coïncider avec le point fixé par le participant (voir la Figure 19). Les participants ont été divisés en cinq groupes à raison de 4 participants par groupe, de manière à équilibrer la position du regard dans les stimuli. Ainsi, pour un stimulus donné, 4 participants fixaient la zone 1, 4 autres la zone 2 et ainsi de suite jusqu'à la zone 5 (carré latin). La présentation des stimuli ainsi que la succession des positions d'apparition étaient aléatoires.

Figure 19. Illustration de la condition des cinq positions du regard dans un stimulus de 6

¹⁰ La fréquence des bigrammes correspond à la fréquence d'occurrence des paires de lettres dans les mots du lexique français. Par exemple, le bigramme " ou " possède une fréquence d'occurrence plus élevée que le bigramme " oe ", c'est-à-dire que " ou " apparaît dans beaucoup plus de mots que " oe ".

lettres. Les chiffres à gauche indiquent la position du regard qui correspond au centre de chacune des 5 zones qui partagent le stimulus.

Procédure. Le participant était assis face à l'écran de l'ordinateur de façon à ce que ses yeux soient à une distance d'environ 50 cm du centre de l'écran. La consigne était présentée à l'écran. Le participant débutait par une phase d'entraînement comportant 20 essais (10 mots et 10 pseudomots) différents de ceux utilisés au cours de la session expérimentale proprement dite. Chaque essai commençait avec l'apparition d'un point de fixation au centre de l'écran (constitué par le signe "+") pendant 1500 ms. Il était ensuite remplacé par le stimulus présenté pendant une durée de 34 ms. Cette durée de présentation a été définie au cours d'un pré-test de façon à éviter « un niveau plafond » pour les mots. Un masque, constitué d'une suite de six dièses (#) succédait immédiatement à la présentation du stimulus ; il était présenté à la même position que le stimulus. Il restait à l'écran pendant 1000 ms. Les participants avaient pour consigne de focaliser leur attention au centre de l'écran (en fixant le point de fixation) et ensuite de taper le stimulus perçu à l'aide du clavier de l'ordinateur. La longueur des stimuli n'était pas précisée par l'expérimentateur. Les participants étaient encouragés à taper les quelques lettres perçues lorsqu'ils n'avaient pas identifié le stimulus dans sa totalité et ceci indépendamment de leurs positions dans la séquence. Les erreurs de frappe pouvaient être corrigées. Ils validaient leur réponse et déclenchaient l'essai suivant en appuyant sur la touche "entrée". L'expérience avait lieu en un seul bloc et durait environ 20 minutes. Une pause avait lieu après 50 essais. La *variable dépendante* mesurée était le pourcentage de réponses correctes. Une réponse était considérée comme correcte lorsque le participant rapportait toutes les lettres dans le bon ordre. Du fait de la difficulté de la tâche, le temps de réaction n'a pas été enregistré.

5-1-2 Résultats et discussion

La Figure 20 présente les pourcentages d'identifications correctes obtenus pour les mots et les pseudomots en fonction de la position du regard dans le stimulus. Bien que les mots soient, en général, mieux identifiés que les pseudomots, les courbes de l'EPR dans les mots et dans les pseudomots décrivent toutes deux une forme classique de « J » inversé, laquelle est moins prononcée dans le cas des mots.

Figure 20. Pourcentages d'identifications correctes et erreurs standard obtenus pour les mots et les pseudomots en fonction de la position du regard dans le stimulus (Expérience 1).

Les scores de réponses correctes ont été soumis à une analyse de variance (ANOVA) comportant la catégorie lexicale (mots vs pseudomots) et la position du regard dans le stimulus (position de 1 à 5) comme facteurs intra sujets. À l'issue de cette analyse, des effets significatifs de la catégorie lexicale [$F(1,19)=305.3$, $p<.0001$], de la position du regard dans le stimulus [$F(4,76)=30.1$, $p<.0001$] ainsi que de l'interaction entre ces deux facteurs [$F(4,76)=5.4$, $p<.001$] ont été mis en évidence. Ainsi les pourcentages d'identifications correctes étaient supérieurs pour les mots par rapport aux pseudomots.

Des analyses réalisées séparément pour les deux catégories lexicales avec la position du regard comme facteur ont mis en évidence une variation significative des performances en fonction de la position du regard pour les mots [$F(4,76)=8.4$, $p<.0001$] et, pour les pseudomots [$F(4,76)=37.1$, $p<.0001$]. Les performances d'identifications étaient meilleures lorsque le regard se posait sur le début de la séquence (positions 1 et 2) plutôt que sur la fin de la séquence (position 4 et 5 ; $F(1,19)=16.6$, $p=.0006$ et $F(1,19)=49.2$, $p<.0001$ respectivement pour les mots et les pseudomots).

Comme précédemment décrit dans la littérature, on retrouve un EPR classique dans les mots sur les pourcentages de réponses correctes au cours d'une tâche d'identification perceptive (e.g., Brysbaert et al., 1996 ; Kajii & Osaka, 2000 ; Montant et al., 1998 ; Nazir et al., 1992 ; Stevens & Grainger, 2003). La reconnaissance du mot est maximale lorsque les participants fixent légèrement à gauche du centre et diminue à mesure que l'œil s'écarte de cette position optimale avec une diminution plus importante lorsque les positions du regard sont situées à la fin du mot par rapport aux fixations sur le début du mot. Par ailleurs, on observe dans les pseudomots, un EPR de forme similaire à celui obtenu pour les mots. Toutefois, le niveau de performance pour les pseudomots reste inférieur à celui observé pour les mots, comme indiqué par la différence de hauteur des courbes.

Dans le but de rendre compte de la différence entre les réponses correctes pour les mots et pour les pseudomots, nous allons proposer une simulation des données empiriques de l'expérience 1, avec le modèle CLIP.

5-2 Simulations des données de l'expérience 1

Selon le modèle CLIP, un mot est reconnu (1) lorsque toutes ses lettres ont été identifiées (CLIP total), ou bien, (2) à partir de quelques lettres (CLIPs partiels), via une procédure d'inférence lexicale. Un pseudomot ne peut pas être « inféré » ou « récupéré » à partir d'une information partielle car il ne possède pas de représentation lexicale en tant que telle. Par conséquent, la probabilité théorique d'identifier un pseudomot ne peut être estimée que par la probabilité d'identifier toutes les lettres qui le composent (i.e., CLIP total).

Dans un premier temps, nous allons estimer les « facteurs visuels » du modèle à partir des données empiriques obtenues pour les pseudomots au cours de l'expérience 1. Plus précisément, nous allons calculer la probabilité d'occurrence du CLIP total, en ajustant le paramètre du taux de diminution de lisibilité des lettres (br) afin d'obtenir un écart minimum entre la forme des courbes théoriques et des courbes empiriques (procédure utilisée par Nazir et al., 1998, cf paragraphe 4-1). Puis, nous allons calculer la probabilité d'occurrence des CLIPs partiels, en conservant la valeur du paramètre visuel (br) obtenue pour la simulation des pseudomots. Ces CLIPs partiels seront alors combinés aux mesures de la contrainte lexicale afin de rendre compte de l'EPR dans les mots.

Deux types de mesures statistiques vont être testés afin d'estimer la contrainte lexicale (1) le nombre de « voisins CLIP » conformément à la proposition de Kajii (2000), c'est-à-dire le nombre de mots partageant les mêmes lettres que le CLIP partiel aux mêmes positions (codage absolu de la position des lettres) et, (2) le nombre de « voisins CLIP » utilisant un codage de la position des lettres relatif, parallèlement à Stevens et Grainger (2003).

5-2-1 Estimation du paramètre visuel

Afin d'estimer le taux de diminution de lisibilité des lettres (br), nous avons déterminé le meilleur ajustement théorique de la courbe de l'EPR empirique des pseudomots. La probabilité d'identifier un pseudomot de 6 lettres, en fonction de la position du regard dans le pseudomot $P(PW)$, est égale à la probabilité d'identifier toutes ses lettres (i.e., probabilité d'occurrence du CLIP total, Nazir et al., 1998) soit:

Étant donné que nous avons utilisé 5 positions du regard pour des stimuli de 6 lettres, les fixations ne correspondent pas exactement au centre des lettres comme dans l'équation (1) proposée par Nazir et al. (1998). L'équation (1) a, par conséquent, été légèrement modifiée afin de tenir compte de ces nouveaux paramètres (voir l'Annexe 2 pour une explication de ces modifications).

La probabilité d'identifier un mot a été calculée selon l'équation (6) suivante :

où $Pp_w(f, a, bl, br, n)$ est la probabilité de reconnaître une séquence de 6 lettres en fonction de w , la zone de fixation relative (de 1 à 5), a , la probabilité de reconnaître la lettre fixée ($a = 1$), b_l , b_r , le taux de diminution de lisibilité d'une lettre, en fonction de l'excentricité de cette lettre à gauche et à droite de la fixation respectivement, et n , la position sérielle de la lettre à identifier.

Une valeur minimale du RMSD (.028) a été obtenue pour br égale à .05 (avec $a = 1$ et $bl/br = 1.8$). La Figure 21 présente les résultats de la simulation ainsi que la courbe empirique de l'EPR dans les pseudomots de la Figure 20. Comme on peut le voir, la courbe théorique correspond aux données empiriques du point de vue de la forme et de la hauteur. Il faut souligner, par conséquent, que le modèle MPLI (Nazir et al., 1991 ; 1998) rend compte de l'identification de toutes les lettres sans influence lexicale avec une grande précision.

Figure 21. Probabilité théorique et empirique d'identifier un pseudomot de 6 lettres en fonction de la position du regard dans le stimulus. La probabilité théorique d'identifier un pseudomot est égale à la probabilité d'occurrence du CLIP total. Le taux de diminution à droite (br) est égal à .05 (RMSD = .028)

5-2-2 Probabilité d'occurrence des CLIPs

En conservant la même valeur du paramètre br (.05), nous avons ensuite calculé les probabilités d'occurrence des 62 CLIPs partiels pour un mot de 6 lettres, en fonction des cinq positions du regard. Le choix des CLIPs partiels, considérés comme contribuant principalement à la reconnaissance des mots, a été fondé en premier lieu sur la probabilité d'occurrence moyenne des CLIPs, dépendante du nombre de lettres identifiées. Les 6 « *CLIP-1* », (i.e., une lettre absente), possédant les plus grandes probabilités d'occurrence, ont ainsi été tous sélectionnés : « *L2L3L4L5L6 », « L1*L3L4L5L6 », « L1L2L3L4*L6 », « L1L2L3*L5L6 », « L1L2*L4L5L6 » et « L1L2L3L4L5* ». Parmi les « *CLIP-2* » (i.e., deux lettres absentes), les six CLIPs contenant les deux lettres externes ont également été sélectionnés : « L1**L4L5L6 », « L1*L3*L5L6 », « L1L2L3**L6 », « L1L2**L5L6 », « L1*L3L4*L6 » et « L1L2*L4*L6 » (cf., Table 3).

Table 3. Probabilités d'occurrence du CLIP total et des CLIPs partiels dont la moyenne est supérieure ou égale à .005 pour un mot de 6 lettres, en fonction de la position du regard (avec $a=1$, $br = .05$ et bl/br maintenu à 1.8/1).

5-2-3 Simulation 1 : Nombre de « voisins CLIP » de même longueur que le stimulus

Dans la simulation présentée au paragraphe 4-5 des données empiriques de Brysbaert et al., (1996), la contrainte lexicale, représentée par le nombre de mots « voisins CLIP » a été calculée à partir de probabilités empiriques d'inférence des mots, obtenues dans une tâche de complétion de trigrammes. Nous allons tester, dans un premier temps, si le nombre de voisins CLIP compté à partir d'une base de données lexicale peut rendre compte de l'EPR dans les mots. La nouvelle base de données « Lexique » (New, Pallier, Ferrand & Matos, 2001) a été utilisée pour estimer ces paramètres lexicaux.

Nous avons déterminé, pour chacun des CLIPs sélectionnés (cf., Table 3) et pour chaque stimulus mot de l'expérience 1, le nombre de « voisins CLIP » de 6 lettres $N(C_i)$, partageant strictement les mêmes lettres que le mot cible, aux mêmes positions, d'après la base « Lexique » (New et al., 2001). Par exemple, le CLIP partiel « BAN**E » identifié à partir du stimulus « BANQUE » active tous les mots de 6 lettres contenant ces 4 lettres aux mêmes positions comme « **BANQUE** », « **BANALE** », « **BANANE** », etc.

La probabilité de prédire un mot à partir d'un CLIP donné $P(W|C_i)$, a été estimée d'après l'équation (5) :

Enfin, la probabilité de reconnaître un mot a été calculée selon l'équation (2) :

Les résultats de ces calculs sont présentés dans la Figure 22 avec les données des pseudomots de la Figure 20 pour une comparaison.

Figure 22. Courbes théoriques et empiriques de l'EPR dans les mots et dans les

pseudomots. Les lignes continues représentent les données théoriques pour les mots (symboles ronds, pleins) et pour les pseudomots (symboles carrés, pleins) ; les lignes pointillées représentent les données empiriques de l'expérience 1 pour les mots (symboles ronds, vides) et pour les pseudomots (symboles carrés vides). La composante visuelle est identique entre les mots et les pseudomots. Le facteur lexical utilisé pour prédire un mot est le nombre moyen de « voisins CLIP » (simulation1).

Comme on peut le voir sur la figure, la courbe théorique de l'EPR dans les mots possède une forme classique en « J » inversé comme la courbe empirique ; cependant, la hauteur de la courbe est « sous-estimée », comme l'indique la valeur relativement élevée du RMSD (.11).

Les résultats, obtenus à l'issue des présentes simulations, indiquent que l'estimation des candidats lexicaux activés par le nombre de « voisins CLIP » de même longueur que le stimulus ne rend pas compte de l'effet exercé par la contrainte lexicale sur la reconnaissance des mots.

Dans ce qui suit, nous allons tenter d'améliorer l'ajustement de la courbe théorique aux données empiriques, en modifiant un certain nombre de paramètres.

(1) Le modèle CLIP repose sur un postulat en désaccord avec les données de la littérature (e.g., Humphreys et al., 1990 ; Peressoti & Grainger, 1995 ; 1999 ; Whitney, 2001). Dans ce modèle, l'unité de codage est la lettre à une position spécifique donnée ; autrement dit lorsqu'une lettre est identifiée, sa position dans la séquence est automatiquement codée (voir aussi McClelland & Rumelhart, 1981 ; Paap et al., 1982, pour une proposition similaire). Un codage de la position absolue des lettres dans les mots est cependant contestable étant donné que la perception de l'ordre des lettres dans une séquence sans signification est très instable, en particulier lorsque la durée de présentation est brève (Huey, 1908 ; Coltheart, 1984). Afin de tester davantage la relation entre l'identification des lettres et le codage de la position *sérielle* des lettres dans la séquence ¹¹, nous allons, dans le paragraphe qui suit, ré-analyser les données de l'expérience 1 au niveau de la lettre. Ainsi, au lieu de déterminer la proportion de stimuli complets correctement identifiés (cf. Figure 20), nous allons mesurer le nombre de lettres correctement identifiées pour chaque position du regard, indépendamment de leur position sérielle dans la séquence. Si, comme le postule le modèle CLIP, le codage de l'identité et de la position des lettres est automatique, le nombre de lettres correctement identifiées devrait être significativement plus élevé dans les mots que dans les pseudomots.

(2) De plus, au cours de la simulation précédente, tous les candidats lexicaux compatibles avec un CLIP donné étaient inclus dans l'estimation de la contrainte lexicale des mots. Néanmoins, comme le soulignent, par exemple, Grainger et collaborateurs (1989), ce n'est pas tant le nombre, mais plutôt, la fréquence des voisins qui interfère avec la reconnaissance du mot cible. La présence d'au moins un voisin de plus haute fréquence que le stimulus mot entraîne une interférence du traitement du stimulus (voir aussi Carreiras et al., 1997 ; Grainger & Jacobs, 1996 ; Grainger & Segui, 1990 ; Paap et

¹¹ Afin d'éviter toute confusion entre la position du regard et la position des lettres dans la séquence, nous emploierons le terme *position sérielle* pour faire référence à la position des lettres dans la séquence.

al., 1982). Afin de fournir une meilleure estimation de la nature des candidats lexicaux impliqués dans la reconnaissance des mots, nous allons conduire une analyse qualitative des erreurs commises par les participants au cours de l'expérience 1, lorsque ceux-ci n'avaient pas identifié le stimulus.

5-3 Analyses supplémentaires

5-3-1 Identification des lettres et codage de la position

Les réponses correctes et incorrectes ont été ré-analysées au niveau de la lettre. Une lettre présente dans le stimulus cible a été considérée comme étant correctement identifiée, indépendamment de sa position sérielle dans le stimulus réponse. Par exemple, les deux premières lettres du mot « BANQUE » étaient considérées comme correctes que le participant ait répondu « ba » ou « ab ». Dans ces deux cas, un score de « 1 » a été attribué aux deux premières lettres du mot et un score de « 0 » pour les autres lettres. Par ailleurs, lorsqu'une lettre était présente deux fois dans le stimulus cible (comme le « r » dans « BEURRE » par exemple) et que le participant ne la retranscrit qu'une seule fois, nous avons attribué un score de « 1 » pour ces deux lettres puisque nous ne pouvions pas décider laquelle des deux lettres avait été vue.

La Figure 23 présente les pourcentages de lettres individuelles correctement identifiées ainsi que les pourcentages d'identifications correctes des séquences de lettres (cf. Figure 20), en fonction de la position du regard dans le stimulus. Les résultats montrent clairement que le nombre de lettres individuelles correctement identifiées indépendamment de la position sérielle dans la séquence (panel de droite de la Figure 23) est beaucoup plus élevé que le nombre de séquences de lettres correctement rapportées (panel de gauche de la Figure 23). Par ailleurs, ces résultats mettent en avant le fait que la différence entre les mots et les pseudomots est beaucoup plus prononcée sur les proportions de séquences de lettres rapportées que sur les proportions de lettres individuelles correctement identifiées.

Figure 23. Pourcentages d'identification des séquences de lettres (à gauche) et pourcentages de lettres correctement identifiées, indépendamment de leur position dans la séquence (à droite), obtenus pour les mots (symboles pleins) et pour les pseudomots (symboles vides), en fonction de la position du regard dans le stimulus (expérience 1).

Une analyse de variance multivariée (MANOVA) a été réalisée sur les variables dépendantes « séquences de lettres » et « lettres individuelles correctes », comportant la catégorie lexicale (mots vs pseudomots) comme facteur. Cette analyse tient compte des éventuelles corrélations entre les variables dépendantes. L'analyse a révélé un effet significatif de la catégorie lexicale [$F(2,37) = 60.9$; $p < .0001$]. Les tests univariés ANOVA ont indiqué que la catégorie lexicale était significative sur les variables « séquences de

lettres » et « lettres individuelles correctes » [$F(1,38)= 63.5$; $p<.0001$; $F(1,38)= 13.7$; $p<.001$] respectivement.

Ainsi, alors que le nombre de séquences de lettres correctement rapportées diffère considérablement entre les mots et les pseudomots (79,5% et 48,5% réponses correctes en moyenne respectivement), le nombre total de lettres individuelles correctement identifiées lorsque la position sérielle n'est pas prise en considération est quasi identique dans les deux conditions (94.9% et 90,5% réponses correctes en moyenne respectivement), bien que cette différence soit significative.

Ces observations vont par conséquent à l'encontre de l'hypothèse du modèle CLIP selon laquelle l'identité et la position d'une lettre sont codées en même temps. Nos résultats montrent que lorsqu'une lettre est identifiée, sa position n'est pas systématiquement codée. L'analyse des lettres a en outre permis de mettre en évidence que l'avantage des mots sur les pseudomots n'était pas lié à une meilleure récupération de l'identité des lettres, mais plutôt à un meilleur codage de la position des lettres dans la séquence. Ces remarques suggèrent que la récupération de l'identité des lettres serait un processus « *visuo-perceptif* ». En revanche, le codage de la position serait un processus « *lexical* », faisant appel aux connaissances orthographiques (et phonologiques) de la langue.

5-3-2 Analyse qualitative des erreurs

a) Nature des erreurs

À l'issue de l'expérience 1, deux types d'erreurs ont été principalement recueillis : des erreurs dites « d'inférence lexicale » et des erreurs « d'omission ». Ces dernières étaient constituées par la production de quelques lettres du stimulus sans signification (généralement les 2-3 premières lettres du stimulus) et, dans une moindre mesure, de réponses « nulles », c'est-à-dire que le participant n'avait donné aucune réponse.

Les erreurs d'inférence lexicale étaient composées par des erreurs dites « *inférences incorrectes* » (i.e., production d'un mot à la place d'un stimulus mot) et, par des erreurs « *lexicalisations* », (i.e., production d'un mot à la place d'un stimulus pseudomot, comme « clavier » pour « claver » par exemple). La Table 4 résume les proportions d'erreurs d'inférence lexicale moyennées à travers les 5 positions du regard, par rapport au nombre total d'essais (proportions absolues) et au nombre total d'erreurs (proportions relatives).

Table 4. Nombres et proportions d'erreurs d'inférence lexicale moyennées à travers les 5 positions du regard, recueillis à partir des stimuli mots et pseudomots à l'issue de l'expérience 1.

La Figure 24 présente les proportions d'erreurs absolues de lexicalisation (panel de gauche) et d'inférence incorrecte (panel de droite) en fonction de la position du regard dans le stimulus.

Figure 24. Pourcentages d'erreurs d'inférence lexicale recueillies au cours de l'expérience

1 à partir des stimuli pseudomots (lexicalisations) et des stimuli mots (inférences incorrectes) en fonction de la position du regard dans le stimulus.

b) Nombre de lettres correctement identifiées

La Figure 25 indique le pourcentage de lettres correctement identifiées, lorsque les participants ont commis une erreur d'inférence lexicale ou une erreur d'omission, en fonction de la position sérielle des lettres dans les stimuli (procédure identique à celle du paragraphe 5-3-1). Par exemple, un score de « 1 » a été attribué pour la 1^o, 2^o, 3^o, 4^o, 5^o et 6^o lettre du stimulus mot « PERDRE » lorsque le participant avait mal identifié le mot de 7 lettres « PRENDRE ». Les scores ont été moyennés à travers les 5 positions du regard.

Figure 25. Pourcentages de lettres des stimuli mots et des stimuli pseudomots correctement identifiées en fonction de la position sérielle des lettres dans les stimuli (avec L1 correspondant à la 1^o lettre) pour les erreurs d'inférence lexicale (graphique de gauche) et les erreurs d'omission (graphique de droite).

En moyenne, 82,5% et 86,8% des lettres des stimuli ont été rapportés dans les erreurs d'inférence lexicale contre, 65,7% et 68,2% dans les erreurs d'omission, respectivement pour les stimuli mots et pseudomots. Comme le montre clairement la Figure 25, le patron général des résultats révèle une dissociation frappante de la forme générale des courbes en fonction du type d'erreurs produit. Plus précisément, pour les erreurs d'inférence lexicale, les courbes décrivent un « U » inversé, indiquant que les lettres externes des stimuli se retrouvent davantage dans les erreurs d'inférence lexicale que les lettres internes. Les courbes obtenues pour les erreurs d'omission sont, quant à elles, décroissantes avec un niveau de performance maximum pour la 1^o lettre qui chute ensuite quasi linéairement pour les lettres suivantes, suggérant que seules les premières lettres des stimuli sont rapportées indépendamment de la position du regard dans le stimuli. Ces résultats suggèrent par conséquent que les participants tendent à inférer un mot lorsqu'un maximum de lettres est lisible, en particulier lorsqu'ils ont identifié la première et la dernière lettre du stimulus.

c) Longueur des erreurs d'inférence lexicale

La Figure 26 montre que la majorité des erreurs d'inférence lexicale étaient des mots de 6 lettres (62,4% et 57,3% respectivement pour les stimuli mots et pseudomots). De plus, 19,3% et 19,2% étaient des mots de 7 lettres et, 15,6% et 18,7% des mots de 5 lettres, respectivement pour les stimuli mots et pseudomots. Une proportion négligeable de mots de 8 lettres et 4 lettres a également été enregistrée.

Figure 26. Pourcentages d'erreurs d'inférence lexicale recueillies pour les mots et pour les pseudomots en fonction de la longueur.

5-3-3 Discussion

L'ensemble des observations issues de l'analyse de l'identification des lettres permet donc d'affirmer que, contrairement aux hypothèses initiales du modèle CLIP (Kajii, 2000,

voir aussi McClelland et Rumelhart, 1981 ; Paap et al., 1982, pour une proposition similaire), le codage de l'identité et de la position d'une lettre n'est pas automatiquement lié. Nos résultats vont par conséquent dans le sens des études en faveur de l'hypothèse d'un codage de la position relative des lettres, laquelle a obtenu de plus en plus de crédit ces dernières années (e.g., Humphreys et al., 1990 ; Grainger & Jacobs, 1991 ; Peressotti & Grainger, 1995, 1999 ; Stevens et Grainger, 2003). L'analyse qualitative des erreurs a permis, par ailleurs, d'apporter des arguments supplémentaires à l'encontre d'un codage absolu de la position des lettres. Les erreurs d'inférence lexicale, commises par les participants, n'étaient pas strictement de la même longueur que le stimulus, suggérant que la longueur des stimuli n'est pas précisément codée.

Par ailleurs, il est important de souligner que 64% des inférences incorrectes recueillies, étaient des mots de plus haute fréquence que le stimulus. Cette remarque rejoint les précédentes observations rapportées par Grainger et Segui (1990) à la suite d'une tâche d'identification perceptive de mots. Les erreurs les plus fréquemment observées consistaient à donner un voisin orthographique plus fréquent que le stimulus. Des constatations similaires ont également été rapportées par Snodgrass & Mintzer (1993). Il a ainsi été suggéré que lorsque le participant était dans une situation de « devinement », le candidat lexical le plus fréquent s'impose alors naturellement (voir aussi Forster & Shen, 1996).

Ainsi, la prise en considération d'un certain nombre de facteurs, supposés influents pour la reconnaissance des mots, va nous permettre de faire de meilleures prédictions des données empiriques de l'expérience 1. En particulier, dans le but de tester davantage l'hypothèse d'un effet de la fréquence du voisinage, seuls les « voisins CLIP » dont la fréquence d'usage est supérieure aux stimuli mots seront introduits dans les simulations. De plus, afin de pallier aux limites d'un codage absolu de la position des lettres, deux « fourchettes » de longueur vont être testées : les mots de longueur comprise « entre 3 et 10 lettres », comme proposé par Stevens et Grainger (2003), et les mots « entre 5 et 7 lettres », conformément aux observations issues de l'analyse qualitative des erreurs. Ainsi, par exemple, les candidats lexicaux activés par le CLIP partiel « CHA**E », identifié à partir du stimulus « CHAISE » (possédant une fréquence d'usage de 48.45 occurrences par million), sont « CHANCE » (77.3), « CHAMBRE » (231.23) et « CHARGE » (68.23).

5-4 Simulations : Nombre de « voisins CLIP » de plus haute fréquence

5-4-1 Simulation 2 : « de 3 à 10 lettres »

Conformément à la démarche proposée par Stevens et Grainger (2003), nous avons inclus dans le calcul tous les « voisins CLIP » de 3 à 10 lettres possédant une fréquence d'usage supérieure au stimulus cible. La probabilité de reconnaître un mot a été calculée

selon la même procédure que le paragraphe précédent. Les résultats des simulations sont présentés dans la Figure 27.

Figure 27. Probabilité théorique de reconnaître un mot de 6 lettres (trait plein, symboles pleins) calculée selon la procédure de simulation 2. La courbe en trait pointillé, symboles pleins, représente la courbe théorique calculée selon la procédure de simulation 1 ; la courbe en trait pointillé, symboles vides, représente la courbe empirique des mots obtenue à l'issue de l'expérience 1.

Comme on peut le voir sur la Figure 27, les courbes théoriques obtenues à l'issue de la simulation 2 ($\text{RMSD} = .07$) s'approche davantage de la courbe empirique par comparaison avec la courbe théorique obtenue à l'issue de la simulation 1 ($\text{RMSD} = .11$). L'ajustement de la hauteur de la courbe avec les données empiriques reste cependant imprécis.

5-4-2 Simulation 3 : « de 5 à 7 lettres »

Conformément aux résultats obtenus à l'issue de l'analyse des erreurs, un comptage incluant uniquement les mots « voisins CLIP » de 5 à 7 lettres a été réalisé selon la même procédure que la simulation 2. Les résultats de la simulation 3 sont présentés dans la Figure 28.

Figure 28. Probabilité théorique de reconnaître un mot de 6 lettres (trait plein, symboles pleins) calculée selon la procédure de simulation 3. La courbe en trait pointillé, symboles vides, représente la courbe empirique des mots obtenue à l'issue de l'expérience 1. Les symboles carrés représentent les données obtenues pour les pseudomots empiriques (trait pointillé, symboles vides) et théoriques (trait plein, symboles pleins).

Comme l'indique très clairement la Figure 28, les courbes de l'EPR dans les mots théoriques et empiriques sont quasi confondues ($\text{RMSD} = .028$). Le meilleur ajustement des données empiriques est généré en utilisant un champ de voisinage contenant uniquement les mots de plus haute fréquence, ne possédant pas strictement le même nombre de lettres que le mot cible et qui préservent l'ordre absolu des lettres dans la séquence.

5-4-3 Simulation 4 : Erreurs d'inférence lexicale

Les prédictions pour les mots étant validées par nos simulations, nous pouvons alors clore ce chapitre par une simulation des erreurs d'inférence lexicale (lexicalisations et inférences incorrectes) commises à partir des stimuli pseudomots et mots (cf paragraphe 5-4).

Afin de simuler les erreurs de lexicalisations commises par les participants lors de la présentation d'un pseudomot, nous allons adopter une démarche identique à la simulation proposée pour les mots. Comme pour les mots, les CLIPs partiels, obtenus à partir d'un

stimulus pseudomot, activent également les représentations lexicales des mots (voir par exemple, Coltheart et al., 1977 ; Grainger & Jacobs, 1996 ; McClelland & Rumelhart, 1981 ; Paap et al, 1982). Cependant, dans le cadre de la simulation des performances d'identification d'un mot, nous avons prédit la probabilité qu'un mot spécifique donné soit correctement rapporté ; au contraire, dans le cadre de la simulation d'une erreur de lexicalisation, il s'agit de prédire la probabilité que n'importe quel mot soit rapporté. Il faut souligner que, contrairement aux mots, les CLIPs partiels identifiés à partir d'un stimuli pseudomot ne sont pas tous susceptibles de donner lieu à la perception d'un mot. Par exemple, le CLIP partiel « BLO**E » identifié à partir du stimulus pseudomot « BLOISE » peut potentiellement entraîner la production des mots « BLOUSE » ou « BLONDE » par exemple ; en revanche, aucun mot ne peut être inféré à partir du CLIP « REP**L » (stimulus « REPRIL »).

Ainsi, parallèlement à la probabilité d'inférer un mot, estimée par l'équation (2) :

La probabilité d'inférence lexicale à partir d'un CLIP **partiel** donné (C_{pi}), $P(Inf\ Lex)$ peut être estimée par l'équation (7) suivante :

où $P(C_{pi})$ est la probabilité d'occurrence du CLIP partiel C_{pi} et $P(Inf\ Lex / C_{pi})$ est la probabilité de produire un mot à partir d'un CLIP partiel donné.

La probabilité de produire un mot $P(Inf\ Lex \square C_{pi})$, lorsqu'un CLIP partiel C_{pi} a été identifié, est fonction du nombre de « CLIPs partiels susceptibles de donner lieu à un mot » ou « *CLIP lexical* ». La probabilité de produire un mot à partir d'un CLIP partiel donné (C_{pi}) peut ainsi être estimée par le rapport entre le nombre de « CLIPs lexicaux » et le nombre total de stimuli, soit l'équation (8) suivante :

Notez, que les CLIPs partiels identifiés à partir d'un stimulus mot étant tous « lexicaux », $P(Inf\ Lex \square C_{pi})$ prend toujours la valeur de 1 dans le cas des mots.

Ainsi, appliquée aux pseudomots, l'équation (7) prédit la probabilité de commettre une erreur de lexicalisation. Appliquée aux mots, elle prédit la probabilité de produire une inférence lexicale « totale », qu'elle soit « correcte » ou « incorrecte ».

La probabilité d'inférence correcte peut être estimée par la différence entre l'équation (2) et l'équation (3), soit, $(2) - (3)$.

La probabilité d'inférence incorrecte peut être estimée par la différence entre (7) et $((2) - (3))$.

Dans le but de déterminer si ces prédictions rendent compte des pourcentages d'erreurs d'inférence lexicale réellement commises par les participants (cf., paragraphe 5-3-2), nous avons calculé les courbes de l'EPR théoriques correspondantes selon une démarche strictement similaire à la simulation 3. Les résultats des simulations ainsi que les pourcentages d'erreurs d'inférence lexicale recueillies empiriquement sont présentés dans la Figure 29.

Figure 29. Pourcentages d'erreurs d'inférence lexicale théoriques (trait plein, symboles pleins) et empiriques (trait pointillé, symboles vides) en fonction de la position du regard dans le stimulus. À gauche, les probabilités d'erreurs de lexicalisation et à droite, les probabilités d'inférence incorrecte.

Ainsi, comme on peut le voir dans la Figure 29, les pourcentages d'erreurs d'inférence lexicale prédits par le modèle CLIP correspondent relativement bien aux données empiriques dans le cas des pseudomots (RMSD=.034) et des mots (RMSD=.035).

En résumé, l'expérience 1, réalisée dans le but d'explorer l'EPR dans les mots et dans les pseudomots via le modèle CLIP (Kajii, 2000), était basée sur la logique selon laquelle la différence observée entre les mots et les pseudomots serait due au processus d'inférence lexicale, qui tend à améliorer les performances d'identifications des mots. Les résultats obtenus à l'issue de ces simulations confortent cette hypothèse. L'ajout de la contrainte lexicale à la probabilité théorique d'identifier un pseudomot (estimée par la probabilité d'identifier le CLIP total) rend compte avec une remarquable précision de l'EPR dans les mots. Dans ce contexte, une mesure statistique de la contrainte lexicale des mots conforme aux données de la littérature a été établi.

Par ailleurs, une analyse qualitative des erreurs d'identification a permis de mettre en évidence trois processus sous-jacents à la reconnaissance d'un mot ou d'un pseudomot.

1. L'identification perceptive totale où les participants perçoivent toutes les lettres du stimulus et codent leurs positions dans la séquence (CLIP total). Ce processus est purement perceptif et s'applique aux mots et aux pseudomots de manière comparable. 2. L'inférence lexicale donne lieu à des inférences correctes et incorrectes dans le cas des stimuli mots et à des erreurs de lexicalisations dans le cas des pseudomots. 3. Enfin, lorsque aucun de ces processus n'a lieu, des erreurs d'omission se produisent, c'est le cas, en particulier, lorsque le nombre de lettres identifiées est trop faible, et/ou lorsque celles-ci ne sont pas compatibles avec un candidat lexical. L'apport principal de notre travail à l'issue des simulations des données de l'expérience 1 concerne la quantification théorique de ces trois processus via le modèle CLIP.

La probabilité d'occurrence de ces processus peut être résumée ainsi :

1. L'identification perceptive totale (CLIP total) est estimée par l'équation (3) :

2. La probabilité d'inférer un mot (correctement ou incorrectement) est estimée par l'équation (7) :

La probabilité de **rapporter** correctement un mot est estimée par l'équation (2) :

3. La probabilité de produire une erreur d'omission dans le cas des mots et des pseudomots est estimée par :

L'ensemble des observations, rapportées à l'issue de la première expérience, suggère que la mise en jeu de ces trois processus (identification perceptive totale, inférence lexicale, et omission) serait principalement dépendante des facteurs liés à la lisibilité des lettres et à la contrainte lexicale des mots.

Néanmoins, afin de s'assurer de la réalité de ces processus, tels que décrits dans le contexte de la tâche d'identification perceptive, il apparaît fondamental de déterminer leur conditions d'apparition en fonction du contexte expérimental. Autrement dit, si les processus cognitifs sont basés sur des principes généraux sous jacents à la reconnaissance visuelle des mots, leurs implications dans une autre tâche devraient être « détectables ».

Nous allons, par conséquent, au cours de la seconde partie expérimentale, comparer la tâche d'identification perceptive à une tâche de décision lexicale où le participant doit simplement décider si la séquence brièvement présentée est un mot (réponse « oui ») ou un nonmot (réponse « non »). Étant donné que les demandes propres à chacune de ces deux tâches sont différentes (dans un cas, il s'agit de rapporter précisément toutes les lettres du stimulus et dans l'autre cas, il faut simplement décider si la séquence est un mot ou un nonmot), les réponses comportementales mesurées (i.e., les réponses correctes et incorrectes) diffèrent entre les deux tâches. Ainsi, dans le contexte de la tâche de décision lexicale, les prédictions suivantes peuvent être énoncées. **1** L'identification perceptive totale devrait donner lieu à des réponses correctes, pour les mots (« Hits »), et pour les pseudomots (rejets corrects). **2** Le processus d'inférence lexicale devrait améliorer les performances pour les mots (Hits), et diminuer les performances pour les pseudomots par l'occurrence de fausses alarmes (i.e., dire « oui » pour un pseudomot). **3** À l'inverse, lorsque l'information visuelle est insuffisante, donnant lieu à des erreurs d'omissions dans la tâche d'identification perceptive, les performances pour les pseudomots tendraient à augmenter ; les performances pour les mots, au contraire, tendraient à diminuer, entraînant des « Misses » (i.e., dire « non » pour un mot). La Figure 30 résume l'ensemble de ces prédictions.

Figure 30. Modèle du comportement et des processus impliqués dans les tâches d'identification perceptive et de décision lexicale lors de la présentation brève d'un mot ou d'un pseudomot. Les chiffres entre parenthèses font références aux équations précédemment décrites.

Chapitre 6 EPR, stratégies et tâches

L'objectif poursuivi lors de la présente expérience consistait à comparer directement les performances dans une tâche d'identification perceptive et une tâche de décision lexicale.

6-1 Identification perceptive et Décision lexicale : Expérience 2

6-1-1 Méthode

Participants. Vingt participants droitiers, volontaires et de langue maternelle française ont participé à cette expérience. Tous avaient une vue normale ou corrigée. Les participants n'ont pas été rétribués pour leur participation.

Stimuli. Pour cette expérience, nous avons sélectionné ou construit deux cents stimuli de cinq lettres, 100 mots français et 100 pseudomots prononçables. Les stimuli ont été sélectionnés à partir de la base de données lexicale informatisée "Brulex" (Content et al., 1990). Les pseudomots ont été construits selon le même principe que l'expérience 1. La

fréquence lexicale moyenne des mots test et des mots de base est respectivement de 73.9 et 87 occurrences par million. Les pseudomots possèdent une fréquence de bigramme moyenne proche de celle des mots sources (2.7 et 2.55 occurrences par million respectivement). Tous les stimuli sont présentés dans l'Annexe 4.

Dispositif expérimental. Le dispositif expérimental de l'expérience 2 était identique à celui de l'expérience 1.

Plan expérimental. Le plan expérimental de l'expérience 2 était identique à celui de l'expérience 1, à l'exception des points suivants. Comme les stimuli avaient une longueur de 5 lettres, chacune des 5 zones fixées correspondait au centre de chacune des lettres. Par ailleurs, l'expérience consistait en deux blocs séparés. Dans un bloc, les stimuli étaient présentés pour la tâche d'identification perceptive, et dans l'autre bloc pour la tâche décision lexicale. Les 200 stimuli étaient divisés en deux listes homogènes contenant chacune 50 mots et 50 pseudomots. La moitié des participants voyait une liste en identification et l'autre liste en décision lexicale. On a inversé pour l'autre moitié des participants.

Procédure. La procédure de l'expérience 2 était identique à celle de l'expérience 1. Cependant les stimuli étaient présentés pendant une durée de 17 ms. Lors de la tâche d'identification, les participants avaient pour consigne de taper le stimulus perçu à l'aide du clavier. Les erreurs de frappe pouvaient être corrigées. Le participant validait sa réponse et déclenchait l'essai suivant en appuyant sur la touche " entrée ". Dans la tâche de décision lexicale, les participants devaient décider si le stimulus était un mot (réponse « oui ») ou un pseudomot (réponse « non »), en pressant une des deux touches correspondantes sur le clavier. Le fait de donner la réponse lançait l'essai suivant. Les participants étaient appelés à répondre avant tout le plus exactement possible. L'expérience était précédée de 20 essais d'entraînement différents de ceux utilisés au cours de l'expérience proprement dite (10 mots et 10 pseudomots). Les deux blocs expérimentaux duraient environ 30 minutes. Une pause avait lieu entre les deux tâches.

6-1-2 Résultats 1 : Pourcentages de bonnes réponses

La Figure 31 présente les pourcentages de réponses correctes, pour les mots et les pseudomots, en fonction de la position du regard dans le stimulus, dans la tâche d'identification perceptive (panel de gauche) et dans la tâche de décision lexicale (panel de droite).

Dans la tâche d'identification, les performances, pour les mots et les pseudomots, varient en fonction de la position du regard : on observe un maximum d'identifications correctes pour les positions du regard proches du centre du stimulus. En revanche, dans la tâche de décision lexicale, l'effet de la position du regard est de forme classique pour les mots mais pas pour les pseudomots. Les performances de rejets corrects pour les pseudomots sont identiques, quelle que soit la position du regard, autrement dit, la courbe est plate.

Figure 31. Pourcentages de réponses correctes et erreurs standard obtenues pour les

mots (symboles pleins) et les pseudomots (symboles vides) en fonction de la position du regard. Le graphique de gauche présente les résultats de la tâche d'identification perceptive et le graphique de droite présente les résultats de la tâche de décision lexicale.

Les performances obtenues, pour les mots et les pseudomots, ont d'abord été comparées à travers une analyse combinée de la tâche d'identification et de la tâche de décision lexicale sur les réponses correctes. Les résultats ont indiqué un effet significatif de la tâche [$F(1,19)=27.8$, $p<.0001$], de la catégorie lexicale [$F(1,19)=288.6$, $p<.0001$] et de la position du regard [$F(4,76)=29.7$, $p<.0001$]. Les interactions entre les facteurs tâche et catégorie lexicale [$F(1,19)=34.1$, $p<.0001$], tâche et position du regard [$F(4,76)=3.8$, $p<.005$] ainsi que tâche, catégorie lexicale et position du regard [$F(4,76)=4.6$, $p<.005$] se sont toutes avérées significatives.

Deux ANOVAs ont été réalisées sur les pourcentages de réponses correctes, séparément pour les deux tâches et comportant la catégorie lexicale et la position de regard dans le stimulus comme facteurs intra-sujets. Dans la tâche d'**identification**, des effets significatifs de la catégorie lexicale [$F(1,19)=220.8$, $p<.0001$] et de la position du regard [$F(4,76)=25.9$, $p<.0001$] ont été obtenus ; cependant l'interaction entre ces deux facteurs ne s'est pas révélée significative ($F<1$). Ainsi, bien que les performances étaient meilleures pour les mots, les résultats montrent que l'effet de la position du regard est similaire pour ces deux types de stimuli. Des analyses supplémentaires ont confirmé un effet significatif de la position du regard dans les mots [$F(4,76)=11.2$, $p<.0001$] et dans les pseudomots [$F(4,76)=12.7$, $p<.0001$]. Les courbes ont la forme d'un "J" inversé pour les mots et pour les pseudomots avec de meilleures performances d'identification lorsque les positions étaient situées au début que lorsqu'elles étaient situées à la fin des stimuli ($F(1,19)=18.04$, $p<.001$ et $F(1,19)=8.4$, $p<.01$; respectivement pour les mots et les pseudomots).

Dans la tâche de **décision lexicale**, l'analyse a révélé un effet significatif de la catégorie lexicale [$F(1,19)=24.3$, $p<.0001$] et de la position du regard [$F(4,76)=8.2$, $p<.0001$]. Cependant, contrairement à la tâche d'identification, une interaction significative entre ces deux facteurs a été mise en évidence [$F(4,76)=10.2$, $p<.0001$]. Des analyses supplémentaires ont révélé un effet significatif de la position du regard pour les mots [$F(4,19)=17.1$, $p<.01$], mais pas pour les pseudomots ($F<1$). Comme on peut le voir sur la Figure 31, les mots présentent une courbe classique de la position du regard en forme de « J » inversé avec un meilleur niveau de performance lorsque les participants fixaient le début du mot que lorsqu'ils fixaient la fin du mot [$F(1,19)=18.0$, $p<.0001$]. En revanche, pour les pseudomots, il n'y a pas de différence entre les positions situées au début et à la fin ($F<1$). Autrement dit, la courbe de position du regard dans les pseudomots est plate. Un test t (comparaison univariée à une norme) a permis de montrer que les performances globales pour les pseudomots (62%) différaient du niveau aléatoire de 50% ($t(199)=5.9$, $p<.0001$) à l'exception des performances de la première position ($t(19) = 1.5$, $p>1.0$).

Parallèlement aux données décrites précédemment, les résultats de cette expérience montrent un effet de la position du regard, dans les mots et dans les pseudomots, lorsque les participants devaient rapporter toutes les lettres du stimulus (tâche d'identification perceptive). Les pourcentages d'identification sont maximums lorsque le regard se pose vers le centre du stimulus et diminuent lorsqu'on s'en éloigne, avec une diminution plus

importante lorsque les participants fixent la fin de la séquence que lorsqu'ils fixent le début de la séquence. De façon surprenante, lorsque les participants doivent décider si un stimulus est un mot ou un nonmot (tâche de décision lexicale), l'effet de la position du regard persiste pour les mots, en revanche, il disparaît dans le cas des pseudomots. Les performances de rejets corrects ne dépendent plus alors de la position du regard et restent significativement supérieures au hasard.

6-1-3 Résultats 2 : Analyse qualitative des erreurs dans la tâche d'identification perceptive

a) Nature des erreurs

De manière générale, les mêmes observations que celles présentées dans l'expérience 1 ont été faites à l'issue de l'analyse des erreurs recueillies dans la tâche d'identification perceptive de l'expérience 2. En particulier, deux types d'erreurs ont été relevés, des erreurs d'inférence lexicale (des inférences incorrectes, dont 53.7% étaient constitués de mots de plus haute fréquence que le mot cible et des lexicalisations) et des erreurs d'omission (voir la Table 5 pour un résumé des proportions d'erreurs d'inférence lexicale moyennées à travers les 5 positions du regard).

Table 5. Nombres et proportions d'erreurs d'inférence lexicale moyennées à travers les 5 positions du regard, recueillis à partir des stimuli mots et pseudomots à l'issue de la tâche de l'identification perceptive (expérience 2).

La Figure 32 présente les proportions d'erreurs absolues des lexicalisation (à partir des stimuli pseudomots) et des inférences incorrectes (à partir des stimuli mots) en fonction de la position du regard dans le stimulus.

Figure 32. Pourcentages d'erreurs d'inférence lexicale recueillies à l'issue de la tâche d'identification perceptive à partir des stimuli pseudomots (lexicalisations) et des stimuli mots (inférences incorrectes) en fonction de la position du regard dans le stimulus.

b) Longueur des erreurs d'inférence lexicale

La Figure 33 montre que 92% des erreurs totales d'inférence lexicale étaient des mots de 5 lettres, 4.6% étaient des mots de 4 lettres et 3.5% des mots de 6 lettres.

Figure 33. Pourcentages d'erreurs d'inférence lexicale recueillies moyennées à travers les mots et les pseudomots en fonction de la longueur.

Dans ce qui suit, nous allons tester directement les prédictions énoncées dans la Figure 30. Pour ce faire, nous avons procédé à deux types de comparaison des données obtenues dans les deux tâches à l'issue de l'expérience 2. Une comparaison des données empiriques enregistrées dans les deux tâches, et une comparaison des données empiriques et théoriques via des simulations avec le modèle CLIP.

6-2 Comparaison Identification/Décision lexicale : données empiriques

6-2-1 Comparaison des données moyennées à travers les cinq positions du regard

La Table 6 présente les pourcentages de réponses correctes et incorrectes obtenues dans la tâche d'identification et dans la tâche de décision lexicale, moyennées à travers les 5 positions du regard.

Table 6. Pourcentages de réponses correctes et incorrectes obtenus pour les mots et pour les pseudomots dans la tâche d'identification et dans la tâche de décision lexicale. Les données sont moyennées à travers les 5 positions du regard.

Les proportions de Hits et de rejets corrects mesurées dans la tâche de décision lexicale ont été « calculées » à partir des réponses correctes et des erreurs obtenues dans la tâche d'identification perceptive.

les proportions de Hits mesurées dans la tâche de décision lexicale devraient ainsi être équivalentes à la somme des réponses correctes obtenues pour les mots (72.3) et pour les inférences incorrectes (11.5), mesurées dans la tâche d'identification perceptive ; soit un pourcentage moyen de Hits calculé, à partir de la tâche d'identification, s'élevant à 83.8% contre 76.5% de Hits effectivement observés.

De même, les proportions de rejets corrects dans la tâche de décision lexicale devraient être équivalentes à la somme des réponses correctes obtenues pour les pseudomots (29.2) et des erreurs d'omissions (33.2) enregistrées pour les pseudomots ; soit un pourcentage moyen de rejets corrects calculé, à partir de la tâche d'identification, s'élevant à 62.4%, lequel avoisine de très près les 61.4% de rejets corrects effectivement observés.

Dans ce qui suit, les mêmes données sont présentées en fonction des cinq positions du regard.

6-2-2 Comparaison des données en fonction de la position du regard dans le stimulus

Comme on peut le voir sur la Figure 34, exceptée la légère différence de hauteur des courbes de l'EPR dans les mots, les performances dans les deux tâches sont hautement similaires.

Figure 34. Pourcentages de hits (à gauche) et de rejets corrects (à droite) « calculés » à partir de l'identification (symboles pleins) et « observés » dans la tâche de décision lexicale (symboles vides) en fonction de la position du regard dans le stimulus.

Deux ANOVAs ont été réalisées séparément pour les mots et les pseudomots prenant chacune pour facteur, l'effet « observé vs calculé » et la position du regard dans le stimulus (positions de 1 à 5). Ces analyses n'ont pas mis en évidence de différence significative entre les données « observés » et les données « calculés », du moins dans le cas des pseudomots ($F < 1$). Une légère tendance significative a toutefois été obtenue pour les mots [$F(1,19) = 4.5$; $p = .05$]. Par ailleurs, une influence significative de la position du regard a été relevée [$F(4,76) = 31.2$, $p < .0001$; $F(4,76) = 4.2$, $p < .01$, respectivement pour les mots et les pseudomots] ; en revanche l'interaction entre les effets « observé vs calculé » et la position du regard ne s'est pas avérée significative ($F < 1$ dans les deux conditions).

Ces résultats montrent que le patron général des réponses correctes calculées empiriquement, à partir des réponses enregistrées dans la tâche d'identification perceptive, reproduit assez fidèlement celui observé empiriquement dans la tâche de décision lexicale.

6-3 Comparaison Identification/Décision lexicale : Simulations via le modèle CLIP

6-3-1 Estimation du paramètre visuel (br)

De manière générale, la même procédure de simulation que celle présentée pour l'expérience 1 a été suivie ici. Le paramètre visuel (br) a été calculé à partir des réponses correctes empiriques, obtenues pour les pseudomots, dans la tâche d'identification perceptive, en calculant la probabilité d'occurrence du CLIP total pour un RMSD minimum. Ainsi, pour $br = .12$ (avec $a = 1$ et $bl/br = 1.8$), la courbe théorique correspond aux données empiriques avec un RMSD minimum s'élevant à .06 (voir Figure 35).

Figure 35. Probabilités théorique et empirique d'identifier un pseudomot de 5 lettres en fonction de la position du regard dans le stimulus. La probabilité théorique d'identifier un pseudomot est égale à la probabilité d'occurrence du CLIP total ($br = .12$; $RMSD = .06$) ; la probabilité empirique d'identifier un pseudomot correspond aux données de l'expérience 2 (tâche d'identification perceptive).

6-3-2 Probabilités d'occurrence des CLIPs

La Table 7 présente les probabilités d'occurrence du CLIP total et des CLIPs partiels pour

un mot de 5 lettres, en fonction des cinq positions du regard calculées selon l'équation (3) pour $br = .12$. Dix CLIPs partiels ont été retenus suivant une procédure similaire à celle adoptée dans l'expérience 1.

Table 7. Probabilités d'occurrence du CLIP total et des CLIPs partiels dont la moyenne est supérieure ou égale à .01, pour une séquence de 5 lettres, en fonction de la position du regard (avec $a=1$, $br = .12$ et bl/br maintenu à 1.8/1).

6-3-3 Simulation 5 : Identifications correctes des mots et des pseudomots

Les données obtenues, dans la tâche d'identification perceptive, ont été simulées suivant une procédure identique à celle utilisée lors de l'expérience 1, pour les réponses correctes et les erreurs d'inférences lexicales.

Pour chacun des CLIPs sélectionnés (voir paragraphe précédent) et pour chaque stimulus mot de l'expérience 2, le nombre moyen de « voisins CLIP » a été déterminé d'après la base « Lexique » (New et al., 2001). Conformément à la démarche proposée dans les simulations 2 et 3 de l'expérience 1, nous avons inclus dans ce calcul tous les « voisins CLIP » de plus haute fréquence que le stimulus. Cependant, en raison du nombre élevé d'erreurs d'inférence lexicale de 5 lettres recueillies à l'issue de l'expérience 2 (92%), seuls les mots de 5 lettres ont été pris en compte. La probabilité de prédire correctement un mot, à partir d'un CLIP donné $P(W|Ci)$, a été estimée d'après l'équation (5) et la probabilité totale de reconnaître un mot a été calculée selon l'équation (2). Les résultats des simulations sont présentés dans la Figure 36.

Figure 36. Courbes théoriques et empiriques de l'EPR dans les mots et dans les pseudomots dans la tâche d'identification perceptive de l'expérience 2. Les lignes continues représentent les données théoriques et les lignes pointillées représentent les données empiriques.

Comme l'indique la Figure 36, les courbes de l'EPR dans les mots théoriques correspondent aux courbes empiriques avec un RMSD s'élevant à .06. Dans les deux cas, la forme de la courbe présente une asymétrie classique avec de meilleures performances pour les positions du regard se situant au début des mots que pour les positions situées sur la fin des mots. Il faut noter néanmoins que la valeur du RMSD reste relativement élevée en comparaison avec celle obtenue lors de la simulation 3 de l'expérience 1.

6-3-4 Simulation 6 : Erreurs d'inférence lexicale

Conformément à la démarche suivie dans la simulation 4 (paragraphe 5-4-3), la probabilité d'erreurs d'inférence incorrecte a été estimée par la formule suivante : $[(7)-(2)-(3))$; la probabilité de lexicalisation a été estimée par l'équation (7). Les

résultats des calculs sont présentés dans la Figure 37.

Figure 37. Pourcentages d'erreurs d'inférence lexicale théoriques (trait plein, symboles pleins) et empiriques (trait pointillé, symboles vides) en fonction de la position du regard dans le stimulus. À gauche, les probabilités d'erreurs de lexicalisation et à droite, les probabilités d'inférence incorrecte.

Ainsi, comme le montre la figure, bien que la forme des courbes calculées selon le modèle CLIP ne soit pas totalement conforme à celle obtenue empiriquement (RMSD = .06 pour les lexicalisations et pour les inférences incorrectes), les prédictions quantitatives des erreurs d'inférences lexicales (représentées par la différence de hauteur des courbes) restent relativement précises.

6-3-5 Simulation 7 : Décisions lexicales correctes pour les mots et les pseudomots

La même démarche que celle présentée dans le paragraphe précédent (paragraphe 6-3) a été suivie afin de prédire les pourcentages de Hits et de rejets corrects obtenus dans la tâche de décision lexicale via le modèle CLIP. Plus précisément, la probabilité de Hits a été estimée par la formule suivante : $[(2) + ((7) - (((2) - (3))))]$ et la probabilité de rejets corrects par : $[(3) + 1 - ((3) + (7))]$. Les résultats des calculs sont présentés dans la Figure 38.

Figure 38. Courbes théoriques et empiriques de l'EPR dans les mots et dans les pseudomots dans la tâche de décision lexicale (expérience 2). Les lignes continues représentent les données théoriques et les lignes pointillées, les données empiriques.

Les résultats montrent une fois encore que, bien que le chevauchement des courbes ne soit pas très précis (RMSD = .04 et .06, respectivement pour les mots et les rejets corrects), la hauteur des courbes théoriques correspond relativement bien à celle des courbes obtenues empiriquement.

Ainsi, sur l'ensemble des résultats obtenus à l'issue des simulations avec le modèle CLIP, lorsque les données sont présentées en fonction de la position du regard, l'ajustement des courbes empiriques et théoriques s'est avéré moins précis, en comparaison avec les données de l'expérience 1. Toutefois, si l'on présente ces mêmes données moyennées à travers les cinq positions du regard (voir la Figure 39), l'ajustement des données empiriques et théoriques est alors quasi parfait, comme indiqué par les valeurs plus faibles des RMSD (voir la Table 8). La précision des simulations est vraisemblablement liée aux choix des CLIPs partiels utilisés, lesquels diffèrent en fonction de la longueur des stimuli. Ce point sera abordé dans la discussion générale.

Figure 39. Comparaison des données obtenues empiriquement à l'issue de l'expérience 2 (histogrammes vides) et calculées théoriquement via le modèle CLIP (histogrammes pleins) dans la tâche d'identification perceptive (panel de gauche) et dans la tâche de

décision lexicale (panel de droite).

Table 8. Valeurs des RMSD minimum calculées pour les différents types de réponses entre les données empiriques et les données théoriques obtenues via le modèle CLIP (expérience 2), moyennées à travers les cinq positions du regard.

En résumé, sur l'ensemble des résultats présentés à l'issue de l'expérience 2, les performances de réponses correctes obtenues empiriquement dans la tâche de décision lexicale (Hits et rejets corrects) se sont avérées fortement prédictibles à partir des données obtenues dans la tâche d'identification perceptive, calculées empiriquement, et simulées via le modèle CLIP. Conformément à nos prédictions, ces observations indiquent, que bien que les réponses comportementales enregistrées dans la tâche d'identification et dans la tâche de décision lexicale soient différentes, les processus cognitifs sous-jacents à ces deux tâches semblent fortement similaires.

Ces processus cognitifs se sont révélés , en outre, fortement dépendants de la lisibilité des lettres et de la contrainte lexicale. Ainsi, le processus d'identification totale est optimisé lorsque la lisibilité des lettres (et la contrainte lexicale) sont maximales. L'inférence lexicale, à partir d'une information visuelle partielle (CLIP partiel), nécessite le codage d'une quantité minimale de lettres du stimulus et dépend de la contrainte lexicale du stimulus. Elle donne lieu à des inférences correctes et incorrectes dans le cas des stimuli mots et à des erreurs de lexicalisations dans le cas des pseudomots. En particulier, en fonction du type de candidats lexicaux activés, l'inférence lexicale est correcte ou incorrecte. Plus précisément, lorsque les contraintes lexicales sont fortes, l'inférence est correcte. C'est le cas, par exemple, lorsque le mot cible possède peu ou pas de voisins de plus haute fréquence ; autrement dit lorsque le voisinage est peu compétitif. Au contraire, lorsque le nombre de candidats lexicaux activés, de plus haute fréquence que le mot cible est très élevé, des erreurs d'inférence lexicale ont lieu, en particulier la production de mots de plus haute fréquence que le stimulus. Ce mécanisme donne également lieu à des erreurs de lexicalisations lorsque le stimulus est un pseudomot. Enfin, les omissions se produisent lorsque la lisibilité des lettres est minimale et/ou les contraintes lexicales insuffisantes pour inférer un mot.

Enfin, nous avons vu que les réponses comportementales enregistrées dépendent en partie de la tâche utilisée pour mesurer les processus cognitifs impliqués dans la reconnaissance des mots. Il est, par conséquent, crucial de développer outre une théorie de la lecture, une théorie de la tâche utilisée afin de mieux appréhender les mécanismes cognitifs impliqués (e.g., Balota & Chumblay, 1984 ; Grainger & Jacobs, 1994 ; Grainger & Jacobs, 1996).

Chapitre 7 EPR, asymétrie et familiarité visuelle

Les études, manipulant le format de présentation des mots, ont généralement été désignées pour tester si la reconnaissance du mot est basée sur l'information visuelle holistique (forme globale des mots) ou sur l'identification des lettres (e.g., Cattell, 1886 ; Henderson, 1982 pour des revues de questions). Toutefois, bien que les résultats vont dans le sens d'une absence de rôle de la forme globale des mots dans la lecture (Besner, 1989 ; Mayall & Humphreys, 1996 ; Mayall, Humphreys & Olson, 1997 ; Paap et al., 1984), la réduction des performances de lecture lorsque le format visuel d'un mot est perturbé indique que certains « aspects visuels » sont néanmoins stockés et utilisés pour la reconnaissance rapide et efficace d'un mot (e.g., Benboutayab, Michel, & Nazir, en révision). Adams (1979) et McClelland (1976) ont montré par exemple, que l'effet de supériorité du mot sur les pseudomots est entièrement préservé lorsque les mots sont présentés en cAsSe aLtErNéE (voir aussi Besner, 1983 ; Jordan, Redwood & Patching, 2003 ; McClelland, 1976) ou verticalement (Jordan & Patching, 2003).

Au cours de cette dernière partie expérimentale, nous allons tester la réalité des trois processus sous-jacents à la reconnaissance des mots, en manipulant le format de présentation des stimuli. Ainsi, si ces processus cognitifs sont réellement impliqués dans les mécanismes généraux de la reconnaissance visuelle des mots, la manipulation de la familiarité avec la configuration visuelle ne devrait pas affecter la mise en jeu de ces processus. Toutefois, étant donné, les différences « visuelles » entre ces deux conditions

de présentation, le niveau de performance pour les pseudomots va être ajusté entre les deux conditions afin d'obtenir des conditions visuelles comparables (autrement dit un CLIP total identique). Tous les stimuli étaient présentés dans la casse majuscule, supprimant ainsi tout effet éventuel lié à l'enveloppe visuelle (e.g., Bouma, 1971)¹².

7-1 Identification perceptive et format horizontal / vertical : Expérience 3

7-1-1 Méthode

Participants. Vingt participants droitiers, volontaires et de langue maternelle française ont participé à cette expérience. Tous avaient une vue normale ou corrigée. Les participants n'ont pas été rétribués pour leur participation.

Stimuli. Pour cette expérience, nous avons utilisé 100 mots français de six lettres et 100 pseudomots prononçables et orthographiquement légaux de six lettres. Les stimuli ont été sélectionnés à partir de la base de données lexicales "Brulex" (Content et al., 1990). Les pseudomots ont été construits selon le même principe que les expériences précédentes. La fréquence lexicale moyenne est respectivement de 75.9 et 76,7 occurrences par million pour les mots test et les mots de base. Tous les stimuli sont présentés dans l'annexe 5.

Dispositif expérimental. Le dispositif expérimental de l'expérience 3 était identique à celui de l'expérience 1 à l'exception des attributs des stimuli présentés. Les stimuli sont présentés à l'écran en lettres majuscules dans la taille 24 et la police « Courier ». À une distance approximative de 57 cm de l'écran, une lettre avait une étendue de 0.33 x 0.3 cm (hauteur x largeur). Ainsi, avec un espace entre les lettres de 0.05cm, une séquence horizontale avait une longueur de 2.05 cm (soit 2.05 degrés d'angle visuel). Une séquence verticale avait une longueur de 3.23 cm pour un espace entre les lettres de 0.25 cm.

Plan expérimental. Les 200 stimuli sont répartis en deux listes de 100 stimuli chacune (50 mots et 50 pseudomots) de manière à obtenir deux listes aussi semblables que possible, en particulier concernant la fréquence lexicale et la fréquence de bigramme. Chaque liste est associée à l'un des deux modes de présentations: « horizontal » ou « vertical » (voir Figure 40). Les deux conditions de présentations étaient présentées dans deux blocs séparés afin de pallier à d'éventuelles difficultés liées aux conditions de présentation (les participants pouvaient éprouver des difficultés à prédire l'orientation du prochain stimulus). La moitié des participants a commencé par la condition horizontale et l'autre moitié par la condition verticale. La présentation des stimuli ainsi que la succession

¹² Par ailleurs, afin de s'assurer que le succès des simulations obtenu au cours de l'expérience 1 n'est pas du à d'éventuels effets de la liste de stimuli utilisée, une seconde liste de stimuli de 6 lettres a été utilisée pour cette expérience.

des positions d'apparition étaient aléatoires.

Figure 40. Illustration de la condition des cinq positions du regard dans un stimulus de 6 lettres dans la condition de présentation horizontale (à gauche) et verticale (à droite). Les chiffres indiquent la position du regard qui correspond au centre de chacune des 5 zones qui partagent le stimulus.

Procédure. La procédure de l'expérience 3 était identique à celle de l'expérience 1 à l'exception des points suivants. Le temps de présentation des stimuli étaient de 17 ms pour la condition horizontale et 170 ms pour la condition verticale. Cette durée de présentation a été définie au cours d'un pré-test, de manière à obtenir un niveau de performances comparable entre les pseudomots présentés horizontalement et les pseudomots présentés verticalement. Les stimuli étaient déplacés latéralement (dans la condition horizontale) ou verticalement (dans la condition verticale) par rapport au point de fixation. L'expérience durait approximativement 30 minutes. Une pause avait lieu entre les deux blocs.

7-1-2 Résultats 1 : Identification des mots et des pseudomots

La Figure 41 résume les résultats obtenus. Le graphique de gauche représente les données obtenues dans la condition de présentation horizontale et le graphique de droite les données obtenues dans la condition de présentation verticale.

Figure 41. Pourcentages d'identification correcte obtenues pour les mots et les pseudomots en fonction du format de présentation horizontal (à gauche) et vertical (à droite) et de la position du regard dans le stimulus.

Les scores de réponses correctes ont été soumis à une ANOVA comportant la condition de présentation (horizontale vs verticale), la catégorie lexicale (mots vs pseudomots), et la position du regard dans le stimulus (position de 1 à 5) comme facteurs intra sujets. Les effets principaux de la condition de présentation [$F(1,19)=17.7$, $p=.0005$], de la catégorie lexicale [$F(1,19)=574.8$, $p<.0001$], et de la position du regard [$F(4,76)=78.3$, $p<.0001$] se sont tous révélés significatifs. En revanche, seule l'interaction entre condition de présentation et position du regard s'est avérée significative [$F(4,76)=3.8$, $p<.01$].

Une seconde analyse, réalisée séparément pour chacune des deux conditions de présentation, incluant la catégorie lexicale, et la position du regard dans le stimulus comme facteurs a été conduite. À l'issue de cette analyse, des effets significatifs de la catégorie lexicale [$F(1,19)=292.2$, $p<.0001$; $F(1,19)=237.9$, $p<.0001$] et de la position du regard [$F(4,76)=36.3$, $p<.0001$; $F(4,76)=40.1$, $p<.0001$] ont été observés dans la condition de présentation horizontale et verticale respectivement. L'interaction entre la catégorie lexicale et la position du regard s'est avérée significative, dans la condition de présentation verticale [$F(4,76)=5.8$, $p<.001$], mais pas dans la condition de présentation horizontale ($F<1$).

Comme l'indique la différence de hauteur des courbes de la Figure 41, les mots sont mieux identifiés que les pseudomots dans les deux conditions de présentation, horizontale et verticale. La forme des courbes s'apparente à un « J » inversé, lequel est plus prononcé pour les mots verticaux que pour les mots horizontaux. Il faut noter cependant, que les scores obtenus pour les pseudomots étaient relativement faibles, en particulier pour la position du regard 5. Des analyses supplémentaires ont indiqué un effet de la position du regard pour tous les stimuli [$F(4,76)=13.2$, $p<.0001$; $F(4,76)=12.0$, $p<.0001$; $F(4,76)=38.0$, $p<.0001$; $F(4,76)=10.2$, $p<.0001$ respectivement pour les mots horizontaux, les pseudomots horizontaux, les mots verticaux et les pseudomots verticaux]. Les performances d'identifications pour les mots étaient meilleures lorsque le regard était porté sur le début du stimulus plutôt que sur la fin [$F(1,19)=27.3$, $p<.0001$; $F(1,19)=43.8$, $p<.0001$, respectivement pour les mots horizontaux et verticaux].

Les résultats obtenus montrent un avantage significatif pour les mots par rapport aux pseudomots, dans la condition de présentation horizontale, et dans la condition de présentation verticale. Par ailleurs, l'absence d'interaction entre la condition de présentation et la catégorie lexicale ($F>1$) indique que cet effet est d'amplitude comparable entre les deux conditions (45% contre 40% en moyenne, respectivement dans la condition horizontale et verticale). La familiarité visuelle du format de présentation ne semble, par conséquent, pas contribuer à la supériorité des mots sur les pseudomots, du moins dans la tâche d'identification perceptive. Par ailleurs, les courbes de l'EPR dans les mots horizontaux et dans les mots verticaux présentent toutes deux une asymétrie, caractérisée par le fait que les performances d'identification sont généralement supérieures lorsque les participants fixent le début des mots que lorsqu'ils fixent la fin des mots, avec toutefois une asymétrie plus prononcée dans la condition verticale que dans la condition horizontale. Ainsi, l'EPR persiste dans les mots verticaux toutefois, il est nécessaire d'allonger considérablement la durée de présentation de ces stimuli (10 cycles) afin d'obtenir un niveau de performances comparable aux stimuli horizontaux.

7-1-3 Résultats 2 : Identification des lettres individuelles

Suivant le même principe que l'expérience 1, les réponses correctes et incorrectes ont été ré-analysées au niveau de la lettre. Les résultats sont présentés dans la Figure 42. Conformément aux résultats observés dans l'expérience 1, ces données montrent que la différence entre les mots et les pseudomots est due principalement au codage de la position des lettres et ce dans la condition de présentation horizontale et dans la condition de présentation verticale.

Figure 42. Pourcentages d'identifications correctes des lettres individuelles (à droite) et pourcentages d'identifications des séquences de lettres (à gauche) en fonction de la condition de présentation (horizontale vs verticale) de la catégorie lexicale (mots vs pseudomots) et de la position du regard dans le stimulus (expérience 2).

Les performances d'identification correcte « des lettres individuelles » et « des séquences de lettres » ont été analysées au moyen de deux MANOVAs réalisées

séparément pour la condition horizontale et verticale et comportant la catégorie lexicale (mots vs pseudomots) comme facteur. Ces analyses ont révélé un effet significatif de la catégorie lexicale dans la condition horizontale [$F(2,37)= 50.4$; $p<.0001$] et dans la condition verticale [$F(2,37)= 35.7$; $p<.0001$]. Les tests univariés ANOVA ont indiqué que la catégorie lexicale était significative sur les variables « séquences de lettres » [$F(1,38)= 91.1$; $p<.0001$; $F(1,38)= 59.8$; $p<.0001$] et « lettres individuelles correctes » [$F(1,38)= 15.4$; $p<.001$; $F(1,38)= 10.3$; $p<.005$] respectivement pour la condition horizontale et verticale.

Parallèlement aux données de l'expérience 1, les résultats issus de la présente analyse des lettres montrent, une fois encore, un plus faible effet de lexicalité sur la proportion de lettres individuelles correctement identifiées relativement au nombre de séquences de lettres correctement identifiées, suggérant que l'avantage de la reconnaissance des mots sur les pseudomots est principalement dû à la disponibilité de l'information sur la position sérielle des lettres et très peu à l'information sur l'identité des lettres. Les pourcentages de lettres correctement identifiées s'élevaient à 90,9% et 81% en moyenne pour les mots et les pseudomots dans la condition de présentation horizontale, et à 89% et 80% dans la condition de présentation verticale. Les pourcentages de séquences de lettres correctement rapportées s'élevaient à 69,9% et 24% en moyenne pour les mots et les pseudomots dans la condition de présentation horizontale, et à 55,7% et 17,2% dans la condition de présentation verticale.

7-1-4 Résultats 3 : Analyse qualitative des erreurs

De manière générale, l'analyse qualitative des erreurs a révélé le même patron de résultats que dans l'expérience 1 dans la condition de présentation horizontale et dans la condition de présentation verticale. Les erreurs recueillies à l'issue de l'expérience 3 étaient de deux types, des erreurs d'inférence (inférences incorrectes et lexicalisations) et des erreurs d'omission. Les proportions d'erreurs d'inférence lexicale mesurées dans les deux conditions de présentation, horizontale et verticale, sont présentées dans la Table 9. Les données sont moyennées à travers les 5 positions du regard.

Table 9. Nombre et pourcentages d'erreurs moyennées à travers les 5 positions du regard recueillies pour les stimuli mots et les stimuli pseudomots dans la condition de présentation horizontale et verticale.

Ainsi que présentées dans la Table 9, les proportions relatives d'erreurs d'inférence incorrecte se sont avérées similaires dans la condition horizontale et dans la condition verticale ($F>1$). De manière surprenante en revanche, les pourcentages d'erreurs de lexicalisation se sont révélés beaucoup plus élevés dans la condition de présentation horizontale que dans la condition de présentation verticale [$F(1,19) = 23.2$; $p<.0001$].

Une probable explication à cette incohérence provient vraisemblablement de la différence de performance entre les deux types de pseudomots. Malgré les précautions prises pour ajuster le niveau de performance des pseudomots dans les deux conditions de présentation, le nombre de pseudomots correctement identifiés était globalement

inférieur dans la condition verticale par rapport à la condition horizontale (18% en moyenne contre 25% ; $F(1,19)= 6.4$; $p=.02$). Ainsi, la moins grande lisibilité des lettres dans la condition verticale pourrait être à l'origine de ce patron de résultats et, un allongement du temps de présentation des stimuli verticaux pourrait, par conséquent, remédier à cette différence.

7-2 Lisibilité des lettres et format de présentation : Expérience 4

Afin de proposer une simulation des données empiriques obtenues dans l'expérience 3 avec le modèle CLIP, nous devons déterminer le ratio d'asymétrie entre les deux champs visuels du méridien vertical, le champ visuel supérieur (CVS) et le champ visuel inférieur (CVI). Kajii & Osaka (2000) ont étudié la lisibilité de caractères japonais dans le méridien horizontal et dans le méridien vertical. Alors que la lisibilité des caractères japonais est meilleure dans le CVD que dans le CVG dans le méridien horizontal comme pour le script roman (e.g., Nazir et al., 1991 ; 1998), ces auteurs ont montré que dans le méridien vertical, la lisibilité est similaire entre le CVI et le CVS. Toutefois, cette étude concernait le script japonais et ne peut donc pas être transposable directement au script roman (voir le chapitre 3 pour des différences de lisibilité des lettres entre le script roman et le script hébreu).

Dans le but d'estimer l'asymétrie du taux de diminution de lisibilité des lettres le long du méridien horizontal et vertical, nous avons mesuré la probabilité d'identifier une lettre cible parmi des « distracteurs » composés par la lettre majuscule « X » présentée à différentes excentricités selon une procédure identique à celle de Nazir et al. (1998). Afin de conserver autant que possible des conditions similaires à l'expérience 3, les séquences étaient présentées horizontalement ou verticalement en lettres majuscules. De plus, seules, deux positions du regard ont été utilisées correspondant à la position du regard 1 (la séquence était présentée dans le CVD ou le CVI) et la position du regard 5 (la séquence était présentée dans le CVG ou le CVS) de la Figure 40, respectivement dans la condition de présentation horizontale et dans la condition de présentation verticale. De plus, étant donné que les mots présentés horizontalement au cours de l'expérience 3 étaient en majuscules, nous allons délibérément utiliser le nouveau rapport d'asymétrie obtenu à l'issue de cette expérience 4 afin de simuler les données empiriques dans la condition horizontale et non le rapport d'asymétrie de 1.8, comme indiqué dans l'étude de lisibilité des lettres de Nazir et al. (1991) concernant des séquences de lettres minuscules.

7-2-1 Méthode

Participants. Six participants droitiers, volontaires et de langue maternelle française ont participé à cette expérience. Tous avaient une vue normale ou corrigée. Les participants n'ont pas été rétribués pour leur participation.

Stimuli. Pour cette expérience, 120 séquences de 6 lettres majuscules composées de la lettre « X » et d'une lettre cible ont été construites. Les séquences consistaient en une lettre *cible* et des *distracteurs* composés par la lettre « X ». Vingt lettres cibles ont été utilisées. Il s'agissait de toutes les lettres de l'alphabet excepté le « X » ainsi que les lettres « I, J, L, T et Q ». Ces dernières ont été évaluées comme hautement «discriminables » dans une séquence de « X » au cours d'une expérience pilote antérieure (Nazir et al., 1998). La même lettre cible apparaissait dans chacune des 6 positions de la séquence.

Dispositif expérimental. Le dispositif expérimental de l'expérience 4 était identique à celui de l'expérience 3.

Plan expérimental. Les séquences de lettres étaient présentées selon deux modes de présentation : « horizontal » ou « vertical » (voir le plan expérimental de l'expérience 3). La séquence de lettre horizontale était présentée soit dans le champ visuel droit « CVD », soit dans le champ visuel gauche « CVG ». La séquence de lettre verticale était présentée soit dans le champ visuel inférieur « CVI », soit dans le champ visuel supérieur « CVS ». Les deux conditions de présentation étaient présentées dans deux blocs séparés afin de pallier à d'éventuelles difficultés liées aux conditions de présentation. La moitié des participants a commencé par la condition horizontale et l'autre moitié par la condition verticale. L'expérience consistait en 480 essais pour chaque participant.

Procédure. La procédure de l'expérience 4 était identique à celle de l'expérience 3 à l'exception des points suivants. Le temps de présentation des stimuli étaient de 17 ou 34 ms pour les séquences horizontales et 68 ou 85 ms pour les séquences verticales. Ces temps de présentation ont été défini pour chaque participant au cours de la phase d'entraînement, de façon à éviter un niveau de performance trop bas ou trop élevé ainsi que pour avoir un niveau de performances comparables entre les deux conditions de présentation. Le participant avait pour consigne de taper la lettre cible perçue à l'aide du clavier. Les erreurs de frappe pouvaient être corrigées. Le participant validait sa réponse et déclenchait l'essai suivant en appuyant sur la touche " entrée ". L'expérience avait lieu en deux blocs et durait environ 1 heure 30. Le participant gérait lui-même les pauses.

7-2-2 Résultats et discussion

Afin d'évaluer la distance entre la lettre cible et la position du regard (i.e., l'excentricité de la lettre), la position sérielle des lettres a été référée en degrés d'angle visuel avec des valeurs positives à droite ou en dessous de la fixation (CVD et CVI respectivement) et des valeurs négatives pour les positions à gauche ou au dessus de la fixation (CVG et CVS respectivement). Conformément à la démarche préconisée par Nazir et al. (1991 ; 1998), les valeurs des lettres les plus excentrées dans les quatre champs visuels ont été exclues des analyses.

La Figure 43 présente les scores d'identification de lettres correctes avec les lignes de régression en fonction de la distance de la lettre cible par rapport à la fixation.

Figure 43. Pourcentages d'identification correcte de lettres cible et lignes de régression obtenues dans des séquences horizontales (panel de gauche) et verticales (panel de droite) en fonction de l'excentricité de la lettre en degré d'angle visuel et de la position dans le champ visuel (CVG, CVD, CVS et CVI).

Les scores de réponses correctes ont été soumis à une ANOVA réalisée séparément pour chacune des deux conditions de présentation (horizontale vs verticale) avec le champ visuel (droit/gauche et inférieur/supérieur) et la distance de la lettre cible par rapport à la fixation comme facteurs intra sujets.

À l'issue de ces analyses, un effet principal un effet de la distance de la lettre cible par rapport à la fixation s'est avéré significatif dans les deux conditions de présentation, horizontale et verticale [$F(4,20)= 18,6$; $p<.0001$ et $F(4,20)= 3,9$; $p<.01$ respectivement]. Ainsi, la probabilité d'identifier une lettre cible diminue à mesure que cette lettre dévie du point de fixation indiquant que l'acuité visuelle chute linéairement dans le méridien horizontal et dans le méridien vertical (Weymouth et al., 1928). Par ailleurs, conformément aux résultats antérieurs (e.g., Bouma, 1973 ; Kajii & Osaka, 2000 ; Nazir et al., 1991 ; 1998 ; 2004), une asymétrie de la lisibilité des lettres dans le méridien horizontal a été mise en évidence dans la condition de présentation horizontale [$F(1,5)= 12.4$; $p<.05$] mais pas dans la condition de présentation verticale ($F>1$). Enfin, l'interaction significative entre le champ visuel et la distance de la lettre cible par rapport à la fixation dans la condition de présentation horizontale et dans la condition de présentation verticale [$F(4,20)= 28,3$; $p<.0001$ et $F(4,20)= 160,8$; $p<.0001$ respectivement] suggère que la diminution de lisibilité des lettres est plus importante dans le CVG par rapport au CVD dans le méridien horizontal, et dans le CVS par rapport au CVI, dans le méridien vertical.

Le calcul des quatre lignes de régression linéaire a confirmé une diminution significative des scores d'identification de la lettre cible dans les quatre champs visuels ($p<.05$, voir Table 10). les rapports de l'asymétrie, gauche/droite s'élève à $.368 / .292 = 1.3$. et supérieur/inférieur s'élève à $.392 / .318 = 1.2$.

Table 10. Paramètres et coefficients de l'analyse de régression linéaire pour les champs visuel droit (CVD), gauche (CVG), inférieur (CVI) et supérieur (CVS) pour des séquences de 6 lettres. Les lettres externes ont été exclues des analyses (voir le texte).

Notez que les résultats obtenus à l'issue de cette expérience de lisibilité des lettres montrent que, même lorsque les facteurs liés à la dominance hémisphérique pour le langage et à la direction de lecture sont neutralisés par la présentation des séquences dans le méridien vertical, une asymétrie persiste. Il semblerait, comme le soulignent Schwartz et Kirsner (1982 ; Kirsner et Schwartz, 1986) que des biais attentionnels liés au « balayage horizontal » induisent également une asymétrie inférieure/supérieure avec de meilleures performances dans le « sens de la direction de lecture » (voir aussi Whitney, 2001). Ces résultats diffèrent en outre, de ceux obtenus avec des caractères japonais présentés dans des séquences verticales (Kajii & Osaka, 2000) où les taux de diminution de lisibilité des lettres dans le champ visuel inférieur et dans le champ visuel supérieur se sont avérés identiques. Une possible explication à cette divergence de résultats entre les deux scripts, roman et japonais, pourrait être liée au fait que les séquences de caractères utilisées dans l'expérience de lisibilité menée par Kajii et Osaka (2000) étaient différentes.

Il s'agissait de séquences composées d'un caractère japonais cible entouré de deux chiffres.

Les valeurs des rapports d'asymétrie dans le méridien horizontal (1.3) et dans le méridien vertical (1.2) obtenues à l'issue de cette expérience vont être utilisées pour prédire les probabilités empiriques de reconnaissance des mots de l'expérience 3.

7-2-3 Simulation des données de l'expérience 3

La même procédure de simulation que celle présentée dans l'expérience 1 a été suivie. Dans le cas des stimuli verticaux, nous avons remplacé les paramètres de diminution à droite et à gauche (br et bl) par les paramètres de diminution inférieur et supérieur (bi et bs) respectivement. Nous avons estimé la valeur du taux de diminution à droite, br, et la valeur du taux de diminution inférieur, bi, à partir des courbes empiriques de l'EPR dans les pseudomots pour la condition horizontale et verticale respectivement, obtenues à l'issue de l'expérience 3, pour une valeur du RMSD minimum (.07 et .017 dans la condition horizontale et verticale respectivement). Les rapports d'asymétrie, bl/br (1.3) et bs/bi (1.2), obtenus à l'issue de l'expérience 4, ont été utilisés. Les résultats sont présentés dans la Figure 44.

Figure 44. Courbes théorique et empirique de l'EPR dans les mots et dans les pseudomots dans la condition de présentation horizontale (à gauche) et verticale (à droite). Les lignes continues représentent les données théoriques pour les mots (symboles ronds, pleins) et pour les pseudomots (symboles carrés, pleins) ; les lignes pointillées représentent les données empiriques de l'expérience 3 pour les mots (symboles ronds, vides) et pour les pseudomots (symboles carrés vides ; avec br= 0.1 et bi= 0.13).

Comme on peut le voir sur la figure, combinés aux mesures de lisibilité des lettres de l'expérience 4, la contrainte lexicale rend compte, une fois encore, des données empiriques de l'expérience 3 avec une relative précision, dans la condition de présentation horizontale, et dans la condition de présentation verticale (RMSD = .07 et RMSD = .04 respectivement). Les courbes de l'EPR dans les mots horizontaux et verticaux présentent une forme classique de « J » inversé, indiquant que les meilleures performances sont obtenues lorsque le regard se pose légèrement à gauche du centre des mots et les scores diminuent de part et d'autre de cette position optimale.

Parallèlement aux expériences précédentes, les données obtenus à l'issue de l'expérience 3 vont être « utilisées » afin de prédire les performances dans une tâche de décision lexicale. Les proportions de Hits et de Rejets corrects « attendues », dans une tâche de décision lexicale, dans les mêmes conditions de présentation. Conformément à la démarche suivie dans le paragraphe 6-3, la probabilité de Hit a été estimée par la formule suivante : $[(2) + ((7) - ((2) - (3)))]$ et la probabilité de rejet correct par : $[(3) + 1 - ((3) - (7))]$. Les résultats des calculs sont présentés dans la Figure 45.

Figure 45. Pourcentages de Hits et de rejets corrects calculés à partir des données empiriques de la tâche d'identification perceptive (expérience 3).

Ainsi, il est attendu que la plus faible tendance à lexicaliser dans la condition de présentation verticale par rapport à la condition de présentation horizontale donne lieu à de meilleures performances pour les pseudomots que pour les mots dans une tâche de décision lexicale.

Afin de tester directement ces prédictions, les mêmes stimuli que l'expérience 3 ont été utilisées dans le cadre d'une tâche de décision lexicale (expérience 5), présentée dans des conditions expérimentales similaires à la tâche d'identification perceptive (expérience 3). Toutefois, par « maladresse », le temps de présentation des stimuli verticaux a été fixé à 68 ms (contre 170ms dans la tâche d'identification perceptive de l'expérience 3)¹³.

7-3 Décision lexicale et format horizontal / vertical : Expérience 5

7-3-1 Méthode

Participants. Vingt participants droitiers, volontaires et de langue maternelle française ont participé à cette expérience. Tous avaient une vue normale ou corrigée. Les participants n'ont pas été rétribués pour leur participation.

Stimuli, dispositif expérimental, plan expérimental de l'expérience 5 sont en tous points identiques à l'expérience 3.

Procédure. La procédure de l'expérience 5 était identique à celle de l'expérience 3 à l'exception des points suivants. Le temps de présentation des stimuli était de 17 ms pour la condition horizontale et de 68 ms pour la condition verticale. L'expérience durait approximativement 20 minutes.

7-3-2 Résultats

La Figure 46 résume les résultats obtenus à l'issue de l'expérience 5. Le graphique de gauche représente les données obtenues dans la condition de présentation horizontale ; le graphique de droite, les données obtenues dans la condition de présentation verticale.

Figure 46. Pourcentages de décision lexicale correcte obtenues pour les mots (symboles pleins) et les pseudomots (symboles vides) en fonction du format de présentation horizontal (à gauche) et vertical (à droite) et de la position du regard dans le stimulus.

¹³ Le temps de présentation des stimuli dans la tâche de décision lexicale a été ajusté afin d'obtenir un niveau de performance pour les pseudomots, comparable dans la condition de présentation de présentation horizontale et verticale (même procédure que dans l'expérience 3).

Les scores de réponses correctes ont été soumis à une analyse de variance réalisée séparément pour chacune des deux conditions de présentation (horizontale vs verticale) incluant la catégorie lexicale (mots vs pseudomots), et la position du regard dans le stimulus (position de 1 à 5) comme facteurs intra sujets.

Dans la **condition de présentation horizontale**, l'effet de la catégorie lexicale ne s'est pas avéré significatif ($F > 1$) ; en revanche des effets significatifs de la position du regard [$F(4,76)=12.5$, $p < .0001$] et de l'interaction entre la catégorie lexicale et la position du regard [$F(4,76)=10.8$, $p < .001$] ont été observés. Dans la **condition de présentation verticale**, des effets significatifs de la catégorie lexicale [$F(1,19)= 13.1$, $p = .002$] et de la position du regard [$F(4,76)=8.3$, $p < .0001$] ont été mis en évidence. L'interaction entre la catégorie lexicale et la position du regard s'est également avérée significative [$F(4,76)=7.1$, $p < .001$].

Comme on peut le voir sur la Figure 46 (graphique de gauche), les mots horizontaux sont mieux identifiés que les pseudomots horizontaux lorsque le regard se pose au début des stimuli ; en revanche, lorsque les participants fixent la fin des stimuli, les performances de décision lexicale sont similaires, bien que les pseudomots tendent à donner lieu à de meilleurs scores. En revanche, dans la condition verticale, les résultats sont quelque peu surprenants. En effet, les scores de réponses correctes, relevés pour les pseudomots, sont supérieurs à ceux relevés pour les mots, et cette différence est plus prononcée lorsque le regard se pose sur la fin des stimuli que lorsqu'il se pose sur le début du stimuli. Des analyses supplémentaires ont mis en évidence une variation significative des performances en fonction de la position du regard pour les mots horizontaux [$F(4,76)=19.9$, $p < .0001$] et pour les mots verticaux [$F(4,76)=12.5$, $p < .0001$]. Les hits se sont avérés plus élevés lorsque le regard se posait au début des mots que lorsque le regard était situé à la fin des mots [$F(1,19)=21.8$, $p < .001$; $F(1,19)=20.5$, $p < .001$, respectivement pour les mots horizontaux et verticaux]. Aucun effet de la position du regard n'a été obtenu pour les pseudomots ($F > 1$).

Les courbes de l'EPR dans les mots horizontaux et verticaux présentent une forme classique de « U » inversé avec de meilleures performances lorsque les participants fixent le début des stimuli que lorsqu'ils fixent la fin. En revanche, les courbes de l'EPR dans les pseudomots sont plates. Bien que les rejets corrects dans la condition de présentation verticale soient plus élevés pour les positions 4 et 5 que pour les positions 1 et 2, cette différence ne s'est pas avérée significative. Par ailleurs, bien que les mots soient généralement mieux reconnus que les pseudomots dans la condition de présentation horizontale, lorsque les stimulus sont disposés verticalement ce sont les pseudomots qui tendent à donner lieu à de meilleures performances de reconnaissance ; le niveau de performance des Hits est proche voire en dessous du niveau du hasard. Ces résultats vont par conséquent dans le sens de nos hypothèses prédisant des performances élevées pour les pseudomots verticaux du fait du plus faible nombre de fausses alarmes, au détriment des réponses correctes pour

les mots verticaux.

Les données empiriques obtenues à l'issue de l'expérience 5 (cf., Figure 46) ont été comparées directement aux données calculées à partir des réponses de la tâche

d'identification perceptive (expérience 3, cf., Figure 45). Ces comparaisons sont présentées dans la Figure 47.

Figure 47. Pourcentages de hits (symboles ronds) et de rejets corrects (symboles carrés), « calculés » à partir de l'identification (symboles pleins), et observés dans la tâche de décision lexicale (symboles vides) en fonction de la position du regard dans le stimulus et de la condition de présentation.

Étant donné que les temps de présentation des stimuli dans la **condition de présentation verticale** étaient différents dans la tâche d'identification perceptive (170ms) et dans la tâche de décision lexicale (68ms), les données « calculées » et « empiriques » sont présentées séparément pour cette condition. Ainsi, bien que le niveau de performance est différent entre les deux conditions « calculées » et « empiriques », le patron général des résultats « calculés » reproduit de manière fortement similaire celui observé empiriquement.

Dans la **condition de présentation horizontale** en revanche, les temps d'exposition des stimuli étaient identiques (17ms) entre les deux tâches, permettant d'établir une comparaison directe entre les deux. Ainsi, comme l'indique le graphique de gauche de la Figure 47, les pourcentages de hits et de rejets corrects « calculés » sont fortement similaires à ceux observés empiriquement non seulement d'un point de vue de la hauteur mais aussi de la forme des courbes (RMSD = .06 et .05, respectivement pour les hits et les rejets corrects).

En résumé, Le fait que le plus faible nombre d'erreurs de lexicalisation dans la condition de présentation verticale par rapport à la condition de présentation horizontale se répercute sur le nombre de fausses alarmes observé dans la tâche de décision lexicale apporte des arguments supplémentaires en faveur de notre hypothèse selon laquelle les processus cognitifs sont bien attribuables à des mécanismes généraux de la reconnaissance visuelle des mots.

Ces résultats indiquent, par ailleurs, que les meilleures performances observées pour les pseudomots verticaux par rapport aux mots verticaux dans la tâche de décision lexicale, ne sont pas dûs à une « meilleure identification » des pseudomots, mais au biais de réponse lié à la tâche. Des résultats comparables ont été rapportés par Driver et Baylis (1995) dans une tâche de décision lexicale. Dans cette étude, ces auteurs ont utilisé des mots et des pseudomots présentés verticalement ou inclinés de 90° dans le sens des aiguilles d'une montre. Ils observent, outre le fait que les stimuli inclinés de 90° entraînent de meilleures performances que les stimuli présentés verticalement, des décisions lexicales plus précises pour les pseudomots que pour les mots de basse fréquence dans les deux conditions de présentation. Cependant, les participants mettent plus de temps à rejeter les pseudomots qu'à accepter les mots. Une analyse des temps de réaction a été réalisée sur les réponses correctes mesurées à l'issue de l'expérience 5. Du fait que la consigne ne mettait pas l'accent sur la vitesse des réponses mais sur la précision, les données enregistrées sont relativement grossières. Un élagage a été conduit sur les données. Les temps de réponse non compris dans l'intervalle défini par la moyenne plus ou moins 2,5 fois l'écart-type n'ont pas été pris en compte. Les résultats obtenus indiquent que les temps de réponse moyens étaient généralement plus élevés pour les mots que

pour les pseudomots (867.8 ms et 982.6 ms respectivement) dans la condition de présentation horizontale et (1309.4 ms et 1335.5 ms respectivement) dans la condition de présentation verticale.

Chapitre 8 Discussion générale et Conclusion

8-1 Résumé

L'objectif de ce travail était de déterminer dans quelle mesure, un modèle mathématique relativement " simple ", considérant les facteurs visuels et les facteurs lexicaux, pouvait prédire la reconnaissance visuelle des mots chez les lecteurs experts. À partir d'un modèle, initialement proposé par Kajii (Kajii, 2000, Kajii & Osaka, 2000), nous avons élaboré un certain nombre d'hypothèses relatives aux aspects perceptifs et lexicaux de la reconnaissance de mots écrits. Dans ce contexte, nous avons modifié le modèle initial de manière à rendre compte, non seulement de la reconnaissance des mots (Kajii & Osaka, 2000), mais aussi de la reconnaissance des pseudomots, et de l'occurrence des différents types d'erreurs (lexicalisation des pseudomots, inférences incorrectes des mot et omissions), typiquement observées lors de présentations non optimales des stimuli orthographiques. Ce modèle est décrit dans l'équation suivante :

avec $P(C_i)$ représentant les aspects perceptifs et $P(W \square C_i)$ les aspects lexicaux. Contrairement à la version initiale du modèle, utilisant un codage absolu de la position des

lettres et une estimation empirique des facteurs lexicaux, nous avons choisi d'utiliser un codage de la position des lettres moins rigoureux et une estimation des facteurs lexicaux provenant d'une base de données lexicales établie (e.g. d'après la base " Lexique ", New et al., 2001). Ainsi, de manière cohérente avec les précédentes observations empiriques, mettant en évidence des effets d'amorçage variables en fonction du nombre et de la position des lettres partagées par l'amorce et la cible (e.g., Humphreys et al., 1990 ; Peressotti & Grainger, 1995, 1999), nous avons supposé qu'un ensemble de lettres identifiées, dans des conditions de présentation non optimale, pouvait activer des candidats lexicaux, légèrement plus courts ou plus longs que le mot cible. Ceci, à condition que le stimulus et les candidats lexicaux partagent les mêmes lettres initiale et finale ainsi que certaines lettres internes, dans le bon ordre mais pas nécessairement à la même position absolue. Par ailleurs, conformément aux études mesurant l'effet du voisinage orthographique sur la perception de mots cibles (voir Andrews, 1997 et Mathey, 2001 pour des revues récentes), nous avons fait l'hypothèse que seuls les candidats possédant une fréquence d'occurrence plus élevée que la cible pouvaient entrer en compétition durant le traitement du stimulus. L'application de ces différents postulats, a permis de mettre en évidence des ajustements remarquables entre les données théoriques et empiriques. Selon que $P(W \square C_i)$ soit défini comme la probabilité d'identifier un mot *spécifique* (équation (5)), ou la probabilité d'identifier n'importe quel mot (équation (8)), le modèle prédit les reconnaissances correctes et incorrectes du stimulus, indépendamment du fait qu'il s'agisse d'un mot ou d'un pseudomot. Enfin, la décomposition du processus sous-tendant la reconnaissance visuelle des mots, en différents sous-composants (identification perceptive totale, inférence lexicale et omission), nous a permis de prédire les performances dans deux tâches expérimentales différentes (identification perceptive et décision lexicale), donnant lieu à des profils de réponses qualitativement différents (e.g., Grainger & Jacobs, 1996). L'extension du modèle, tel que présenté à l'issue de ce travail, et initialement proposé par Kajii (2000), représente un outil « simple » mais pertinent pour l'étude de la reconnaissance visuelle de mots. Avant de discuter plus longuement de la validité du présent modèle, nous allons brièvement soulever quelques points, nécessitant d'être améliorés au cours de futures recherches.

8-2 Limites du modèle

8-2-1 Choix des CLIPs

L'ajustement entre les données théoriques et empiriques était particulièrement précis pour les stimuli de 6 lettres (expériences 1 et 3). Pour les stimuli de 5 lettres (Expérience 2), en revanche, nous avons observé quelques variations entre la théorie et l'observation. Les résultats plus bruités pour les mots de 5 lettres, pourraient vraisemblablement provenir du fait que nous n'avons pas encore clairement identifié les paramètres permettant une

sélection « réaliste » des CLIPs participants au processus de reconnaissance d'un mot. Pour les stimuli de 6 lettres, nous avons retenu les CLIPs comprenant les lettres initiale et finale de la cible (à l'exception des CLIP « *L2L3L4L5L6 » et « L1L2L3L4L5* ») ainsi que deux autres lettres. Pour les stimuli de 5 lettres en revanche, la même procédure de sélection conduit à un nombre de CLIPs trop restrictif. Le nombre de CLIPs a ainsi été augmenté, sans réel critère de sélection. Utiliser tout simplement la totalité des CLIPs ne constituait par ailleurs pas une solution puisque certains d'entre eux, comme le CLIP « **L3L4L5 », augmentaient considérablement la probabilité de lexicalisation. Toutes choses égales par ailleurs, inférer des mots connaissant leurs dernières lettres est généralement plus difficile qu'inférer des mots connaissant leurs premières lettres (e.g., « tab** » pour « table » vs. « **ire » pour « foire »). De futurs travaux sont par conséquent nécessaires afin d'améliorer ce point.

8-2-2 Champ du « voisin CLIP »

Alors que l'analyse des erreurs a révélé qu'un stimulus de 6 lettres peut donner lieu à des erreurs de différentes longueurs (mots de 5 à 7 lettres), la même analyse pour les stimuli de 5 lettres a révélé qu'une large majorité des erreurs commises étaient des mots de même longueur que le stimulus. Cette variation du champ du « voisinage CLIP » semble indiquer que la relation entre la longueur du stimulus et celles des candidats lexicaux, pouvant être activés, est relative et non pas absolue (e.g., Humphreys et al., 1990 ; Peressotti & Grainger, 1995, 1999 ; Stevens & Grainger, 2003 ; Whithney, 2001). On peut ainsi envisager que plus la longueur du mot cible est grande, plus la fourchette de longueur des candidats lexicaux activés est importante. Toutefois, n'ayant testé que deux longueurs de mots, cette remarque ne constitue qu'une hypothèse.

Par ailleurs, le codage de la position des lettres que nous avons utilisé peut encore être amélioré par la prise en compte des nouvelles conceptions théoriques du codage de la position des lettres. Une approche prometteuse, basée sur l'activation de « *bigrammes ouverts* » a été développée récemment par Whitney (2001 ; Whitney & Berndt, 1999; voir aussi, Grainger & Van Heuven, sous presse). Ainsi, par exemple, les quatre lettres "T", "A", "K" et "E" du mot cible « TAKE » entraînent l'activation des bigrammes ouverts, adjacents « TA », « AK », « KE » et, non adjacents « TK », « TE », « AE ». Le bigramme ouvert « AB » active les représentations des mots « TABLE », « TABAGISME » mais aussi « TANGLIBLE », « TRANSMISSIBLE », etc.

Enfin, de récentes observations relatives à l'effet d'*amorçage de transposition* (Perea & Lupker, 2003 ; Schoonbaert & Grainger, sous presse), indiquent la facilitation de la reconnaissance du mot cible lorsque l'amorce est formée par la transposition de deux lettres adjacentes (e.g., « humain-humian ») comparée à une amorce contrôle appropriée. Parallèlement à l'amorçage de transposition, on peut ainsi envisager que lorsque le participant a identifié toutes les lettres du mot « CRIER » (**CLIP total**), il code précisément les deux lettres externes ; en revanche il ne code pas l'ordre relatif des lettres internes et peut par conséquent faire des erreurs en produisant un mot « *anagramme de transposition* » comme « CIRER ». L'idée est que le mot anagramme « FRONDE » par exemple, est susceptible d'être activé par le CLIP total « FONDRE » ; en revanche,

l'anagramme « ALPINE » du stimulus mot « PLAINE » a très peu de chance d'interférer avec la reconnaissance du mot cible.

8-2-3 CLIP-total ou identification perceptive totale

Dans toutes nos simulations, nous avons estimé le CLIP-total à partir des performances d'identification obtenues pour les pseudomots, c'est-à-dire la probabilité de rapporter une séquence de lettres sans l'aide du lexique mental. Ce processus a été référencé comme étant une “ perception totale ”. Cependant, la reconnaissance correcte des pseudomots n'est pas basée sur des facteurs perceptifs « purs », étant donné que les performances d'identification de séquences de lettres sans signification varient énormément en fonction de la régularité orthographique (e.g., Baron & Thurston, 1973 ; Carr, 1986 ; Henderson, 1982 ; Massaro & Cohen, 1994). La Figure 48ci-dessous présente les probabilités d'identifier correctement des séquences de lettres sans signification, variant selon leur degré de régularité orthographique (pseudomots légaux (PML), pseudomots illégaux (PMI), et nonmots (NM), cf., Annexe 6 pour la méthode).

Figure 48. Pourcentages d'identification correctes obtenus en fonction du type de stimuli (mots, pseudomots légaux (PML), pseudomots illégaux (PMI), et nonmots (NM)) et de la position du regard dans le stimulus.

Ainsi que présenté dans la Figure 48, la familiarité orthographique a un impact conséquent sur la probabilité de reconnaître une séquence de lettres sans signification. La supériorité des pseudomots orthographiquement légaux sur les pseudomots orthographiquement illégaux et sur les nonmots indique clairement que des facteurs liés à l'apprentissage (autres que le lexique mental) contribuent à la perception des pseudomots. Il est intéressant de noter, par ailleurs, que si les performances d'identification pour les PMI avaient été utilisées lors de l'estimation des aspects perceptifs de la reconnaissance de mots, le modèle n'aurait pas permis de prédire correctement les performances des sujets.

Le fait que la familiarité orthographique provienne du développement d'unités “ sous-lexicales ”, améliorant la perception du groupement de lettres familier, comme le lexique mental, reste à être déterminé. Il est également important de rappeler que les résultats de l'expérience 3 montraient que les séquences de lettres présentées verticalement nécessitaient une augmentation importante du temps de présentation (10 cycles) pour atteindre le niveau de performance obtenu avec la présentation horizontale. Ainsi, la familiarité visuelle semble également contribuer à la perception des pseudomots. Ces différents points nécessiteraient d'être clarifiés afin de mieux contrôler les aspects “ perceptifs ” de notre modèle.

8-3 Le modèle CLIP et les autres modèles

visuo-lexicaux

Ainsi que mentionné dans le résumé, l'extension du modèle CLIP fournit un outil de prédiction des reconnaissances correctes et incorrectes d'un stimulus orthographique, que ce stimulus soit un mot ou un pseudomot, et ceci dans une tâche d'identification perceptive ou de décision lexicale. Aucun autre modèle, à notre connaissance, ne permet de faire de telles prédictions. Dans le modèle proposé par Clark & O'Regan (1999), les aspects lexicaux sont « irréalistes », puisqu'il est impossible pour un sujet de reconnaître correctement un pseudomot orthographiquement légal. Le changement d'une des lettres internes d'un mot (e.g. « chaussure » en « chaussere ») est indétectable par ce modèle, sauf si cette lettre est directement fixée. Le modèle proposé par Stevens et Grainger (2003) rencontre également des difficultés concernant la prédiction de la perception d'un pseudomot. Étant donné que les aspects lexicaux de ce modèle sont basés sur les fréquences positionnelles des lettres, le modèle ne permet pas de distinguer les mots des pseudomots. En d'autres termes, alors que le modèle de Clark et O'Regan (1999) ne serait pas suffisant pour identifier les pseudomots, celui proposé par Stevens et Grainger (2003) reconnaîtrait les pseudomots aussi aisément que les mots. Ces deux prédictions ne sont cependant pas compatibles avec les données empiriques.

8-4 Acquisition de la lecture

Outre l'avantage du modèle CLIP sur les deux autres modèles visuo-lexicaux, distinguer les aspects perceptifs des aspects lexicaux durant la reconnaissance de mots présente plusieurs autres avantages. Le fait que la contrainte lexicale, contribuant à la reconnaissance visuelle des mots, puisse être estimée à partir de bases de données lexicales comme « Lexique », permettra une utilisation « pratique » de notre modèle, afin de mieux appréhender les processus liés à l'acquisition de la lecture et les troubles du langage écrit. Par exemple, en adaptant la base de données lexicales, au niveau de lecture devant être normalement acquis, lors des premières années d'apprentissage, le modèle pourrait aisément prédire l'âge de lecture d'un enfant. Ainsi, au début de l'apprentissage, la reconnaissance des mots devrait fortement dépendre des aspects visuels, et à un moindre degré, des aspects lexicaux. La différence de performances entre les mots et les pseudomots devrait donc être faible (e.g., Grainger, Bouttevin, Truc, Bastien, & Ziegler, 2003), et les erreurs de lexicalisation relativement absentes. Au fur et à mesure que le lexique mental évolue avec l'expérience du lecteur, l'inférence lexicale devrait s'accroître, et les performances pour les mots devraient augmenter par rapport à celles pour les pseudomots. Cependant, il est intéressant de noter que l'inférence lexicale étant fonction du nombre de candidats lexicaux activés par le stimulus (ce nombre dépendant lui-même du nombre de mots déjà acquis) l'importance de l'inférence lexicale ne devrait pas augmenter progressivement jusqu'à atteindre le niveau des lecteurs

adultes, mais plutôt être transitoirement supérieure à celle des lecteurs experts. Quoi qu'il en soit, appliquer ce modèle à de nouveaux lecteurs pourrait certainement aider au dépistage précoce des problèmes de lecture.

8-5 Déficits de lecture acquis

Le modèle développé par Nazir et collaborateurs (Nazir et al., 1991 ; 1998), est un cas particulier du modèle CLIP qui, comme stipulé précédemment, a déjà permis de prédire avec succès, certaines fonctions défaillantes dans des cas de déficits de lecture acquis ou développementaux (Aghababian & Nazir, 2000; Montant et al. 1998; Nazir, Brustman, & Michel, en révision). L'élaboration de ce modèle CLIP offre un outil de diagnostic des déficits de lecture plus robuste, puisqu'il permet de distinguer clairement les déficits « perceptifs » de ceux liés à des facteurs lexicaux. Il est important de rappeler, étant donnés les points discutés en c)., que les déficits perceptifs ne doivent pas être réduits à des problèmes d'acuité visuelle. La patiente mentionnée au paragraphe 7 (Benboutayab, Michel, & Nazir, en révision), souffre, en l'occurrence, d'un déficit perceptif et non d'un déficit lexical, puisque les mots sont correctement identifiés dans le format horizontal. Le diagnostic précis d'un déficit devrait ainsi permettre l'élaboration de programmes de réhabilitation adéquats.

Références

- Adams, M.J. (1979). Models of word recognition. *Cognitive Psychology*, 11, 133-176.
- Aghababian, V., & Nazir, T.A., (2000). Developing normal reading skills: Aspects of the visual processes underlying word recognition. *Journal of Experimental Child Psychology*, 76, 123-150.
- Andrews, S. (1989). Frequency and neighborhood effects on lexical access: Activation or search? *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 15, 802-814.
- Andrews, S. (1992). Frequency and neighborhood effects on lexical access: Lexical similarity or orthographic redundancy? *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 18, 234-254.
- Andrews, S. (1997). The effect of orthographic similarity on lexical retrieval: Resolving neighborhood conflicts. *Psychonomic Bulletin & Review*, 4, 439-461.
- Allen, P.A., Wallace, B., & Weber, T.A. (1995). Influence of case type, word frequency and exposure duration in visual word recognition. *Journal of Experimental Psychology: Human Perception and Performance*, 21 (4), 914-934.
- Balota, D.A., & Chumbley, J. (1984). Are lexical decisions a good measure of lexical access? The role of word frequency in the neglected decision stage. *Journal of Experimental Psychology: Human Perception & Performance*, 10, 340-357.
- Baron, J., & Thurstone, I. (1973). An analysis of the word superiority effect. *Cognitive*

- Psychology*, 4, 207-228.
- Beauvillain, C. (1996). The integration of morphological and whole-word form information during eye fixations on prefixed and suffixed words. *Journal of Memory and Language*, 35, 801-820.
- Besner, D. (1983). Basic decoding components in reading: Two dissociable feature extraction processes. *Canadian Journal of Psychology*, 37, 429-438.
- Besner, D. (1989). On the role of outline shape and word-specific visual pattern in the identification of function words: NONE. *The Quarterly Journal of Experimental Psychology*, 41A(1), 91-105.
- Bouma, H. (1970). Interaction effects in parafoveal letter recognition. *Nature*, 226, 177-178.
- Bouma, H. (1971). Visual word recognition of isolated lower-case letters. *Vision Research*, 11, 459-474.
- Bouma, H. (1973). Visual interference in the parafoveal recognition of initial and final letters of words. *Vision Research*, 13, 767-782.
- Bowers, J.S. (2000). In defense of abstractionist theories of repetition priming and word identification. *Psychonomic Bulletin & Review*, 7 (1), 83-99.
- Bozon, F., & Carbonnel, S. (1996). Influence des voisins orthographiques sur l'identification de mots et de pseudomots. *Revue de neuropsychologie*, 2, 219-237.
- Bradshaw, J. L. & Nettleton, N. C. (1983). *Human cerebral asymmetry*. Englewood Cliffs, NJ: Prentice-Hall.
- Bryden, M. P. (1982). *Laterality: Functional asymmetry in the intact brain*. New York: Academic Press.
- Brysbaert, M. (1994). Interhemispheric transfer and the processing of foveally presented stimuli. *Behavioural Brain Research*, 64, 151-161.
- Brysbaert, M. (2004). The importance of interhemispheric transfer for foveal vision: A factor that has been overlooked in theories of visual word recognition and object perception. *Brain and Language*, 88, 259-267.
- Brysbaert, M., & d'Ydewalle, G. (1988). Callosal transmission in reading. In: G. Luer, U. Lass, & J. Shallo-Hoffman (Eds.), *Eye movement research: Physiological and psychological aspects*, (pp. 246-266). Göttingen: Hogrefe.
- Brysbaert, M., & Meyers, C. (1993). The optimal viewing position for children with normal and with poor reading abilities. *Facets of Dyslexia and its Remediation*, 107-123.
- Brysbaert, M., Vitu, F., & Schroyens, W. (1996). The right visual field advantage and the optimal viewing position effect: On the relation between foveal and parafoveal word recognition. *Neuropsychology*, 10 (3), 385-395.
- Carr, T.H. (1986). Perceiving visual language. In: K.R. Boff, L. Kaufman and J.P. Thomas (Eds.) *Handbook of Perception and Human Performance, Vol II*. Wiley, New York, pp.29:1- 29:92.
- Carreiras, M., Perea, M., & Grainger, J. (1997). Effects of orthographic neighborhood in visual word recognition: Cross-task comparisons. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 23, 857-871.

- Cattell, J.M. (1886). The time taken up by cerebral operations. *Mind*, 11, 220-242, 377-392, 524-538
- Chiarello, C., Liu, S., & Shears, C. (2001). Does global context modulate cerebral asymmetries? A review and new evidence on word imageability effects. *Brain and Language*, 79, 360-378.
- Clark, J.J. & O'Regan, J.K. (1999). Word ambiguity and the optimal viewing position in reading. *Vision Research*, 39 (4), 843-857 .
- Coltheart, M. (1984). Sensory memory: A tutorial review. In H. Bouma & D.B. Bouwhuis (Eds), *Attention and Performance X*, Hillsdale, NJ: Erlbaum.
- Coltheart, M., Curtis, B., Atkins, P., & Haller, M. (1993). Models of reading aloud: Dual-route and parallel-distributed-processing approaches. *Psychological Review*, 100, 589-608.
- Coltheart, M., Davelaar, E., Jonasson, J.T., Besner, D. (1977). Access to the internal lexicon. In S. Dornic (Ed.), *Attention and performance VI* (pp. 535-555). London: Academic Press.
- Coltheart, M., Rastle, K., Perry, C., Langdon, R. & Ziegler, J. (2001). DRC: A Dual Route Cascaded model of visual word recognition and reading aloud. *Psychological Review*, 108 (1), 204-256
- Content, A., Mousty, P., & Radeau, M. (1990). Brulex, une base de données lexicales informatisées pour le français écrit et parlé. *L'année Psychologique*, 90, 551-566.
- Content, A., & Radeau, M. (1988). Données statistiques sur la structure orthographique du français. *Cahiers de Psychologie Cognitive, European Bulletin of Cognitive Psychology*, N° hors-série.
- Deutsch, A., & Rayner, K. (1999). Initial fixation position in reading Hebrew words. *Language and Cognitive Processes*, 14, 393-421.
- Driver, J., & Baylis, G. (1995). Tilted letters and tilted words: A possible role for principal axes in visual word recognition. *Memory & Cognition*, 23, 560-568.
- Ducrot, S., Pynte, J. (2002). What determines the eyes' landing position in words? *Perception & Psychophysics*. 64(7):1130-44.
- Dunn-Rankin, P. (1978). The visual characteristics of words *Scientific American*, 238, 122-130.
- Erdmann, B. & Dodge, R. (1898). *Psychologische Untersuchungen über das Lesen auf experimenteller Grundlage*. Halle : Niemeyer.
- Eriksen, B.A., & Eriksen, C.W. (1974). The importance of being first: A tachistoscopic study of the contribution of each letter to the recognition of four-letter words. *Perception and Psychophysics*, 15, 66-72.
- Estes, W. K. (1972). Interactions of signal and background variables in visual processing. *Perception & Psychophysics*, 12, 278-286.
- Estes, W. K., Allmeyer, D. H., & Reder, S. M. (1976). Serial position functions for letter identification at brief and extended exposure durations. *Perception & Psychophysics*, 19, 1-15.
- Evet, L.J. & Humphreys, G.W. (1981). The use of abstract graphemic information in

- lexical access. *Quarterly Journal of Experimental Psychology*, 33A, 325-350.
- Farid, M., & Grainger, J. (1996). How initial fixation position influences visual word recognition: A comparison of French and Arabic. *Brain and Language*, 53, 351-368.
- Ferrand, L., & Grainger, J. (2001). Homophone interference effects. Manuscrit soumis à publication.
- Forster, K. I. (1976). Accessing the mental lexicon. In R.J. Wales & E.W. Walker (Eds.), *New approaches to language mechanisms* (pp.257-287). Amsterdam : North-Holland.
- Forster, K. I., & Shen, D. (1996). No enemies in the neighborhood: Absence of inhibitory effects in lexical decision and categorization. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 22, 696-713.
- Frost, R., Forster, K. I., & Deutsch, A. (1997). What can we learn from the morphology of Hebrew: A masked priming investigation of morphological representation. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 30, 370-384.
- Grainger, J., Bouttevin, S., Truc, C., Bastien, M., & Ziegler, J. (2003). Word superiority, pseudoword superiority, and learning to read: A comparison of dyslexic and normal readers. *Brain and Language* 87, 432-440.
- Grainger, J., & Jacobs, A.M. (1991). Masked constituent letter priming in an alphabetic decision task. *European Journal of Cognitive Psychology*, 3, 413-434.
- Grainger, J., & Jacobs, A. M. (1993). Masked partial-word priming in visual word recognition: Effects of positional letter frequency. *Journal of Experimental Psychology: Human Perception and Performance*, 19, 951-964.
- Grainger, J., & Jacobs, A. M. (1994). A dual read-out model of word context effects in letter perception: Further investigations of the word superiority effect. *Journal of Experimental Psychology: Human Perception and Performance*, 20, 1158-1176.
- Grainger, J. & Jacobs, A.M. (1996). Orthographic processing in visual word recognition: A multiple read-out model. *Psychological Review*, 103, 518-565.
- Grainger, J., O'Regan, J. K., Jacobs, A. M. & Segui, J. (1989). On the role of competing word units in visual word recognition: The neighborhood frequency effect. *Perception & Psychophysics*, 45, 189-195.
- Grainger, J., O'Regan, J. K., Jacobs, A. M., & Segui, J. (1992). Neighborhood frequency effects and letter visibility in visual word recognition. *Perception & Psychophysics*, 51, 49-56.
- Grainger, J. & Segui, J. (1990). Neighborhood frequency effects in visual word recognition: A comparison of lexical decision and masked identification latencies. *Perception & Psychophysics*, 47, 191-198
- Grainger, J., & Van Heuven, W.J.B. Modeling letter position coding in printed word perception. In *The Mental Lexicon* (Bonin, P., ed.), Nova Science, New York (sous presse).
- Greenberg, J.H. (1966). *Universals of Language*, Cambridge, MA : MIT Press.
- Healy, A.F., & Cunningham, T.F. (1992). A developmental evaluation of the role of word shape in word recognition. *Memory & Cognition*, 20 (2), 141-150.
- Henderson, L. (1982). Orthography and Word Recognition in *Reading*. Academic Press,

New York.

- Holmes, V.M., & O'Regan, J.K. (1987). Decomposing French words. In J.K. O'Regan & A. Levy-Schoen (Eds.), *Eye movements: From physiology to cognition* (pp. 459-466). Amsterdam: North-Holland.
- Huckauf, A., Heller, D., & Nazir, T. A. (1999). Lateral Masking: Limitations of the feature interaction account. *Perception & Psychophysics*, 61, 177-89.
- Huey, E.B. (1908). *The psychology and pedagogy of reading* (Reprinted 1968). Cambridge, MA: MIT Press.
- Humphreys, G.W., Evett, L. J., & Quinlan, P. T. (1990). Orthographic processing in visual word identification. *Cognitive Psychology*, 22, 517-560.
- Hyöna, J., & Pollatsek, A. (1998). Reading Finnish compound words: Eye fixations are affected by component morphemes. *Journal of Experimental Psychology: Human Perception and Performance*, 24, 1612-1627.
- Imbs, P. (1971). Etudes statistiques sur le vocabulaire français. Dictionnaire de fréquences. Vocabulaire littéraire des XIXe et Xxe siècles. *Centre de recherche pour un trésor de la langue française (CNRS)*, Nancy. Paris: Librairie Marcel Didier.
- Inhoff, A.W., Brihl, D., & Schwartz, J. (1996). Compound word effects differ in reading, on-line naming, and delayed naming tasks. *Memory & Cognition*, 24, 466-476.
- Inhoff, A.W., & Tousman, S. (1990). Lexical priming from partial-word previews. *Journal of Experimental Psychology: Learning, Memory & Cognition*, 16, 825-836.
- Jacobs, A. M., & Grainger, J. (1992). Testing a semistochastic variant of the interactive activation model in different word recognition experiments. *Journal of Experimental Psychology: Human Perception and Performance*, 18, 1174-1188.
- Jordan, T. R., & Patching, G. R. (2003). Assessing effects of stimulus orientation on perception of lateralized words and nonwords. *Neuropsychologia*, 41, 1693–1702.
- Jordan, T. R., Redwood, M., & Patching, G. R. (2003). Effects of format familiarity on perception of words pseudowords and nonwords in the two cerebral hemispheres. *Journal of Cognitive Neuroscience*, 15, 537– 548.
- Kajii, N. (2000). Unpublished dissertation, University of Kyoto.
- Kajii, N., & Osaka, N. (2000). Optimal viewing position in vertically and horizontally presented Japanese words. *Perception & Psychophysics*, 62, 1634-1644.
- Kirsner, K., & Schwartz, S. (1986). Words and hemifields: Do the hemispheres enjoy equal opportunity? *Brain and Cognition*, 5, 354–361.
- Knecht, S., Deppe, M., Dräger, B., Bobe, L., Lohmann, H., Ringlestein, E., Henningsen, H. (2000). Language lateralization in healthy right-handers. *Brain*, 123 (1), 74–81.
- Lima, S.D., & Pollatsek, A. (1983). Lexical access via an orthographic code: The basic orthographic syllabic structure (BOSS) reconsidered. *Journal of Verbal Learning and Verbal Behavior*, 22, 310-332.
- Levi, D.M., Klein, S.A., & Aitsebaomo, A.P. (1985). Vernier acuity, crowding and cortical magnification. *Vision Research*, 25, 963-977.
- Mason, M. (1982). Recognition time for letters and nonletters: Effects of serial position, array size, and processing order. *Journal of Experimental Psychology: Human*

- Perception & Performance*, 8, 724-738.
- Massaro, D. W., & Cohen, M. M. (1994). Visual, orthographic, and lexical influences in reading. *Journal of Experimental Psychology: Human Perception & Performance*, **20**, 1107-1128.
- Mathey, S. (2001). L'influence du voisinage orthographique lors de la reconnaissance des mots écrits. *Canadian Journal of Experimental Psychology*, 55 (1), 1-23.
- Mayall K, & Humphreys GW. (1996). Case mixing and the task-sensitive disruption of lexical processing. *Journal of Experimental Psychology Learning and Memory Cognition*, 22(2):278-294.
- Mayall K, Humphreys GW, Olson A. (1997). Disruption to word or letter processing? The origins of case-mixing effects. *Journal of Experimental Psychology Learning and Memory Cognition*, 23(5):1275-86.
- McCann, R.S., & Besner, D. (1987). Reading pseudohomophones: Implications for models of pronunciation assembly and the locus of word-frequency effects in naming. *Journal of Experimental Psychology : Human, Perception & Performance*, 13, 14-24.
- McClelland, J.L. (1976). Preliminary letter recognition in the perception of words and nonwords. *Journal of Experimental Psychology : Human, Perception & Performance*, 2, 80-91.
- McClelland, J.L., & Rumelhart, D.E. (1981). An interactive activation model of context effects in letter perception: Part I. An account of basic findings. *Psychological Review*, 88, 375-407.
- McConkie, G.W., Kerr, P.W., Reddix, M.D., & Zola, D. (1988). Eye movement control during reading: I. The location of initial eye fixations on words. *Vision Research*, 28, 1107-1118.
- McConkie, G.W., Kerr, P.W., Reddix, M.D., Zola, D. & Jacobs, A.M. (1989). Eye movement control during reading: II. Frequency of refixating a word. *Perception & Psychophysics*, 46, 245-253.
- Montant, M. Nazir, T.A., & Poncet, M. (1998). Pure alexia and the viewing position effect in printed words. *Cognitive Neuropsychology*, 15 (1/2), 93-140.
- Nazir, T.A. (1993). On the relation between the optimal and the preferred viewing position in words during reading. In G. Ydewalle & J. van Rensbergen (Eds.), *Perception & Cognition: Advances in eye movement research*, (pp. 349-361). Amsterdam: North-Holland.
- Nazir, T. A. (2000). Traces of print along the visual pathway. In: Kennedy, A., Radach, R., Heller, D. & Pynte, J. (Eds). *Reading as a perceptual process* . (pp. 3-23). North-Holland.
- Nazir, T . A. (2003) On hemispheric specialization and visual field effects in the perception of print. *Cognitive Neuropsychology*, 20, 73-80.
- Nazir, T. A., Benboutayab, N., Decoppet, N., Deutsch, A., & Frost, R. (2004). Reading habits, perceptual learning, and recognition of printed words. *Brain & Language*, 88, 244-311.
- Nazir, T. A., Decoppet, N., & Aghababian, V. (2003). On the origins of age-of-acquisition effects in the perception of written words. *Developmental Science*, 6 :2, 145-152.

- Nazir, T. A & Deprez, V. (2003). De l'oral à l'écrit, et réciproquement. *Science & Vie* Hors série.
- Nazir, T.A., Heller, D., & Sussmann C. (1992). Letter visibility and word recognition : The optimal viewing position in printed words. *Perception & Psychophysics*, 52, 315-328.
- Nazir, T.A., Jacobs, A.M, & O'Regan, J.K. (1998). Letter legibility and visual word recognition. *Memory & Cognition*, 26 (4), 810-821.
- Nazir, T.A., & Montant, M. (1996). Les étapes de traitement précoces dans la reconnaissance visuelle des mots. In S. Carbonnel, P. Gillet, M.-D. Martory & S.Valdois (Eds.), *Approche cognitive des troubles de la lecture et de l'écriture chez l'enfant et l'adulte*. (pp 183-193). Marseille: Solal.
- Nazir, T.A., O'Regan, J.K., Jacobs, A.M. (1991). On words and their letters. *Bulletin of the Psychonomic Society*, 29, 171-174.
- New B., Pallier C., Ferrand L., Matos R. (2001) Une base de données lexicales du français contemporain sur internet: LEXIQUE, *L'Année Psychologique*, 101, 447-462. <http://www.lexique.org>
- Olzak, L.A., & Thomas, J. P. (1986). Seeing spatial patterns. In: K. R. Boff, L. Kaufman, & J.P. Thomas (Eds.), *Handbook of perception and human performance*, (Vol. II, pp. 7:1-7:56). New york, Wiley.
- O'Regan, J.K. (1981). The convenient viewing position hypothesis. In D.F. Fisher, R.A. Monty, & J.W.Senders (Eds.), *Eye movements, cognition, and visual perception* (pp.289-298). Hillsdale, NJ: Erlbaum.
- O'Regan, J. K. (1990). Eye movements and reading. In E. Kowler (Ed.), *Eye movements and their role in visual and cognitive processes* (pp. 395-453). Amsterdam: Elsevier
- O'Regan, J.K., & Jacobs, A.M. (1992). Optimal viewing position effects in word recognition: A challenge to current theory. *Journal of Experimental Psychology; Human Perception and Performance*, 18, 185-197.
- O'Regan, J.K., Levy-Schoen, A., Pynte, J., & Brugailiere, B. (1984). Convenient fixation location within isolated words of different length and structure. *Journal of Experimental Psychology; Human Perception and Performance*, 10, 250-257.
- Paap, K. R., & Johansen, L. S. (1994). The case of the vanishing frequency effect: A retest of the verification model. *Journal of Experimental Psychology: Human Perception & Performance*, 20, 1129-1157.
- Paap, K. R., Newsome, S. L., McDonald, J. E., & Schvaneveldt, R. W. (1982). An activation-verification model for letter and word recognition: the word superiority effect. *Psychological Review*, 89, 573-594.
- Paap, K. R., Newsome, S. L., & Noel, R. W. (1984). Word shape's in poor shape for the race to the lexicon. *Journal of Experimental Psychology; Human Perception and Performance*, 10 (3), 413-428.
- Perea, M., & Lupker, S.J. (2003). Transposed-letter confusability effects in masked form priming. In *Masked Priming: State of the Art* (Kinoshita, S. and Lupker, S.J., eds), pp. 97–120, Psychology Press.

- Perea, M., & Pollatsek, A. (1998). The effects of neighborhood frequency in reading and lexical decision. *Journal of Experimental Psychology: Human Perception & Performance*, 24, 767-779.
- Peressotti, F., & Grainger, J. (1995). Letter position coding in random consonant arrays. *Perception & Psychophysics*, 57, 875-890.
- Peressotti, F. and Grainger, J. (1999) The role of letter identity and letter position in orthographic priming. *Perception & Psychophysics*, 61, 691–706.
- Pollatsek, A., Bolozky, S., Well, A. D., & Rayner, K; (1981). Asymetries in the perceptual span for Israeli readers. *Brain and Language*, 14, 174-180.
- Pollatsek, A., Perea, M., & Binder, K. (1999). The effects of neighborhood size in reading and lexical decision. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 1142-1158.
- Prinzmetal, W. (1992). The word-superiority effect does not require a T-scope. *Perception & Psychophysics*, 51, 473-484.
- Pujol, J., Deus, J., Losilla, J.M., Capdevila, A. (1999). Cerebral lateralization of language in normal left-handed people studied by functional MRI. *Neurology*, 52(5), 1038–43.
- Pynte J. (1996). Lexical control of within-word eye movements. *Journal of Experimental Psychology: Human Perception and Performance*, 22, 958-969.
- Pynte, J., Kennedy, A., & Murray, W. S. (1991). Within-word inspection strategies in continuous reading: time course of perceptual, lexical, and contextual processes. *Journal of Experimental Psychology : Human Perception and Performance*, 17, 458–570.
- Radach, R., & Kempe, V. (1993). An individual analysis of initial fixation positions in reading. In G. d'Ydewalle & J. Van Rensbergen (Eds.), *Perception and cognition: Advances in eye movement research (pp. 213-226)*. Amsterdam: North Holland.
- Rayner, K.(1979). Eye guidance in reading: fixation locations within words. *Perception*, 8, 21-30.
- Rayner, K., McConkie, G.W., & Zola, D. (1980). Integrating information across eye movements. *Cognitive Psychology*, 12, 206-226.
- Rayner, K., & Pollatsek, A. (1989). *The psychology of reading*. Englewood Cliffs, N. J.:Prentice Hall.
- Reicher, G. M. (1969). Perceptual recognition as a function of meaningfulness of stimulus material. *Journal of Experimental Psychology*, 81, 274-280.
- Rey-Debove, J. (1984). Le domaine de la morphologie lexicale. *Cahiers de Lexicologie*, 45, 3-19.
- Rumelhart, D.E., & Mc Clelland, J.L. (1982). An interactive activation model of context effects in letter perception: Part II. The contextual enhancement effect and some tests and extensions of the model. *Psychological Review*, 89, 60-94.
- Schoonbaert, S. and Grainger, J. (2004). Letter position coding in printed word perception: effects of repeated and transposed letters. *Language, Cognitive, Processes*, 19 (3), 333-367.

- Schwartz, S., & Kirsner, K. (1982). Laterality effects in visual information processing: Hemisphere specialization or the orientating of attention? *Quarterly Journal of Experimental Psychology*, 34A, 61–77.
- Sears, C. R., Hino, Y., & Lupker, S. J. (1995). Neighborhood frequency and neighborhood size effects in visual word recognition. *Journal of Experimental Psychology: Human Perception & Performance*, 21, 876-900.
- Segui, J. (1991). La reconnaissance visuelle des mots. In R. Kolinski, J. Morais & Segui (Eds), *La reconnaissance des mots dans les différentes modalités sensorielles: études de psycholinguistique cognitive* (pp. 99-117). Paris: PUF.
- Shillcock, R., Ellison, T.M. & Monaghan, P. (2000). Eye-fixation behavior, lexical storage, and visual word recognition in a split processing model. *Psychological Review*, 107(4), 824-851.
- Snodgrass, J. G., & Mintzer, M. (1993). Neighborhood effects in visual word recognition: Facilitatory or inhibitory? *Memory & Cognition*, 21, 247-266.
- Stevens, M., & Grainger, J. (2003). Letter visibility and the viewing position effect in visual word recognition. *Perception & Psychophysics*, 65 (1), 133-151.
- Townsend, J.T., Taylor, S.G., & Brown, D.R. (1971). Lateral masking for letters with unlimited viewing time. *Perception and Psychophysics*, 10, 375-378.
- Underwood, G., Clews, S., & Everatt, J. (1990). How do readers know where to look next? Local information distribution influence eye fixations. *Quarterly Journal of Experimental Psychology*, 42A, 39-65.
- Vitu, F. (1991). The influence of reading rhythm on the optimal landing position effect. *Perception and Psychophysics*, 50, 58-75.
- Vitu, F., O'Regan, J.K., & Mittau, M. (1990). Optimal landing position in reading isolated words and continuous texts. *Perception and Psychophysics*, 47, 583-600.
- Weymouth, F.W., Hines, D.C., Acres, L.H., Raaf, J.E. & Wheeler, M.C. (1928). Visual acuity within the area centralis and its relation to eye movements and fixation. *American Journal of Ophtalmology*, 11, 947-960.
- Whitney, C.S., & Berndt, R.S (1999). A new model of letter string encoding: Simulating right neglect dyslexia. *Progress in Brain Research*, 121, 143-163.
- Whitney, C.S. (2001). How the brain encodes the order of letters in a printed word: The SERIOL model and selective literature review. *Psychonomic Bulletin & Review*, 8: 221-243
- Ziegler, J.C., & Perry, C. (1998). No more problems in Coltheart's neighborhood: Resolving neighborhood conflicts in the lexical decision task. *Cognition*, 68, 53-62.
- Ziegler, J.C., Rey, A., & Jacobs, A.M. (1998). Simulating individual word identification thresholds and errors in the fragmentation task. *Memory & Cognition*, 26 (3), 490-501.
- Zola, D. (1984). Redundancy and word perception during reading. *Perception & Psychophysics*, 36, 2.

Annexes

Annexe 1

Calcul de l'efficacité de l'ajustement de valeurs théoriques à des valeurs empiriques.

La *méthode des moindres carrés* est basée sur la minimisation du « critère quadratique », c'est à dire la somme des carrés des écarts. Elle permet de trouver la meilleure courbe d'ajustement des valeurs théoriques aux valeurs empiriques.

Afin de mesurer l'efficacité d'un ajustement entre une courbe théorique et une courbe empirique, on effectue la somme des carrés des *erreurs* de l'ensemble des données ; plus cette somme est petite, meilleur est l'ajustement.

Autrement dit la valeur du RMSD (« Root mean square deviation ») doit être la plus petite possible. Le calcul du RMSD est donné par l'équation suivante :

graph4.gif

où O_i , représente la valeur empirique ; P_i , la valeur théorique et n , le nombre de valeurs.

Annexe 2 : Modification du modèle de Nazir et al. (1991 ; 1998)

Selon le modèle MPIL (Nazir et al., 1991), le nombre de positions du regard est égal au nombre de lettres du mot, par conséquent il est nécessaire de généraliser le modèle avec n'importe quel nombre de positions du regard et nombre de lettres.

Pour cela, on suppose que le ratio de diminution de la visibilité des lettres n'est pas discret mais continu quand les fixations ne sont pas exactement sur une lettre.

Calcul des distances entre le centre de chacune des lettres et la fixation du regard donnée.

Pour le calcul des cinq positions du regard, la longueur du stimulus (1 lettre est égale à 1 unité, soit 6 unités pour un mot de 6 lettres) a été divisée en cinq parties strictement égales ($6 / 5 = 1,2$). Le stimulus était présenté au niveau du centre de chacune de ces parties. Ainsi, la distance entre le début de la séquence et la fixation est égale à 0,6 ; 1,8 ; 3 ; 4,2 et 5,4 respectivement pour les fixations 1, 2, 3, 4 et 5 (lignes pointillées sur la figure).

La distance entre le centre de chacune des lettres et le début de la séquence est de 0,5 ; 1,5 ; 2,5 ; 3,5 ; 4,5 et 5,5 respectivement pour la 1°, 2°, 3°, 4°, 5° et 6° lettre.

Par exemple, la distance entre le centre de la 1° lettre et la fixation 1 est égale à $0,6 - 0,5 = 0,1$ (voir le tableau ci-dessous pour le calcul des autres distances).

Résolution d'équation polynomiale de degré 1 sous la forme:

$$distance = (fx + n + z),$$

avec f, la fixation du regard (de 1 à 5) et n, la lettre à identifier (de 1 à 6).

Pour ce faire, on pose un système à deux équations et deux inconnues (x et n).

Par exemple, lorsque les lettres sont à gauche de la fixation :

$$\text{Soit, } (3x - 2x) + (3 - 2) = (0,5 - 0,3), \text{ soit } x = 1,2.$$

Pour obtenir la valeur de z, on remplace x par 1,2 dans une des deux équations :

$$2 * 1,2 + 2 + z = 0,3 ; \text{ soit } z = -0,1.$$

Ainsi, la distance entre la fixation du regard et le centre d'une lettre à gauche de la fixation est : $1,2 * f - n - 0,1$.

3) Soit l'équation (6):

Annexe 3 : : stimuli de l'expérience 1

Annexe 4 : stimuli de l'expérience 2

Annexe 5 : stimuli de l'expérience 3

Annexe 6 : Expérience pilote x

Méthode

Participants. Dix participants droitiers, volontaires et de langue maternelle française ont participé à cette expérience. Tous avaient une vue normale ou corrigée. Les participants n'ont pas été rétribués pour leur participation.

Stimuli. Deux cents stimuli de cinq lettres ont été sélectionnés ou construits pour cette expérience: 50 mots, 50 pseudomots orthographiquement légaux, 50 pseudomots orthographiquement illégaux et enfin 50 non-mots imprononçables et orthographiquement illégaux. Les stimuli ont été sélectionnés à partir de la base de données lexicale "Brulex" (Content et al., 1990). La fréquence d'usage des mots est comprise entre 1 et 24 occurrences par million. Les pseudomots légaux ont été construits à partir de 50 mots de cinq et six lettres, en remplaçant ou en supprimant une lettre. Les mots de base ont une fréquence d'usage comprise entre 85 et 1684 occurrences par million. Les pseudomots illégaux sont prononçables cependant, la fréquence de leurs bigrammes est faible. Les non-mots sont imprononçables et la fréquence de leurs bigrammes est quasi nulle. Les stimuli proviennent essentiellement de la thèse Montant (1998, expérience 3).

Dispositif expérimental. Le dispositif expérimental de l'expérience pilote était identique à celui de l'expérience 2.

Plan expérimental. Le plan expérimental de l'expérience pilote était identique à celui de l'expérience 2.

Procédure. La procédure de l'expérience pilote était identique à celle de l'expérience 4. Le temps de présentation des stimuli était de 34 ms pour tous les stimuli. La durée de présentation des stimuli a été définie afin d'éviter un niveau de performances « plafond » pour les mots. L'expérience avait lieu en un seul bloc et durait environ 30 minutes. Une pause avait lieu après 100 essais.