



L'impact de la dégradation du signal de parole sur le langage, de sa représentation à sa compréhension

Marie Dekerle

► To cite this version:

Marie Dekerle. L'impact de la dégradation du signal de parole sur le langage, de sa représentation à sa compréhension. Neurosciences [q-bio.NC]. Université Claude Bernard - Lyon I, 2015. Français. NNT : 2015LYO10290 . tel-01315187

HAL Id: tel-01315187

<https://theses.hal.science/tel-01315187>

Submitted on 12 May 2016

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

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

ECOLE DOCTORALE Neurosciences et Cognition (NsCO)

THESE DE L'UNIVERSITE DE LYON

En vue de l'obtention du titre de
Docteur en Neurosciences

L'impact de la dégradation du signal de parole
sur le langage,
de sa représentation à sa compréhension

Soutenue publiquement le 14/12/2015
par :

Marie DEKERLE

Sous la direction de :
Fanny MEUNIER

JURY

Pr Séverine CASALIS	Rapporteur
Dr Sophie DONNADIEU	Examineur
Pr. Pierre FOURNERET	Examineur
Pr Frédéric LAVIGNE	Rapporteur
Dr Fanny MEUNIER	Directrice de Thèse

*« On ne fait jamais attention à ce qui a été fait ;
on ne voit que ce qui reste à faire »*
Marie Curie

RESUME

L'objectif de cette thèse est de s'intéresser à l'impact de la dégradation du signal de parole sur le traitement et les représentations du langage. Le signal peut être dégradé de façon transitoire (i.e., parole dans le bruit) ou permanente dans le cas d'un déficit au niveau des traitements auditifs centraux. Le travail expérimental s'est donc articulé autour de deux axes. L'Axe 1 étudie l'impact de la présence de bruit sur le traitement sémantique au niveau du mot isolé et dans un contexte phrastique. Deux études ont permis de mettre en évidence que lorsque le signal de parole est masqué, le traitement sémantique est moins efficace, voire disparaît. À la lumière de l'*Effortfulness Hypothesis* nous suggérons que le traitement sémantique n'est pas automatique mais dépendant de ressources cognitives. Lorsque le signal de parole est dégradé ces ressources cognitives sont allouées aux traitements linguistiques de bas niveau, les traitements de haut niveau comme le traitement sémantique sont donc moins efficaces.

L'Axe 2 évalue les liens entre la dégradation permanente du signal de parole du fait d'un manque de maturité des traitements auditifs centraux ou à un Trouble du Traitement Auditif (TTA) et les représentations langagières. Deux études se sont respectivement intéressées au développement des traitements auditifs centraux sur une population d'enfants de 6 à 11 ans et à la présence de TTA au sein d'une population d'adultes dyslexiques. Ces différentes populations ont été testées grâce à la Batterie d'Évaluation des Compétences Auditives Centrales (BECAC ; Donnadieu et al., 2014). Cette batterie ne contient aucun matériel verbal et permet ainsi d'étudier les liens entre les traitements auditifs centraux et les représentations langagières. Les résultats ont mis en évidence que les traitements auditifs centraux continuent leur maturation durant l'adolescence. Il semble que certaines compétences influencent la qualité des représentations du langage particulièrement autour de 8-9 ans. La dernière étude a permis de mettre en évidence chez la population dyslexique des difficultés dans les tests impliquant des traitements spectral et temporel. Ces compétences sont apparues liées au sein de cette population aux représentations phonologiques.

L'ensemble des résultats suggèrent que la dégradation du signal de parole a différents effets sur le langage selon sa nature. Ainsi, lorsqu'elle est transitoire elle impacte sa compréhension en dépit d'une intelligibilité préservée. Lorsque la dégradation est permanente, son impact sur les représentations du langage, semble évoluer au cours du développement pour disparaître à l'âge adulte sauf pour la population dyslexique.

Mots clés : Parole dans le bruit, traitement sémantique, *Effortfulness Hypothesis*, développement auditif, dyslexie, TTA

ABSTRACT

This thesis aims at investigating the effect of speech degradation on language processing and representation. Speech signal can be degraded temporary (i.e., speech in noise) or continually, when central auditory processes are deficient. Experimental work was therefore based on two main axes. The first one got interested in the effect of noise on semantic processing, despite preserved intelligibility. Two studies showed that semantic processing is less efficient when speech is presented in noisy condition. In light of the Effortfulness Hypothesis, we suggest that semantic processing relies on cognitive resources. When signal is degraded these resources are reallocated to low level processes and therefore few are left available to perform higher level processes.

Axe 2 aims at evaluating links between continuous speech degradation because of immaturity of central auditory processes or Auditory Processing Disorder (APD) and language representation. Two studies investigated the effect of auditory processing development on language on a children population (6-11 years old) and the effect of APD on language representation in an adult dyslexic population. Both populations were evaluated using the BECAC (Donnadieu et al. 2014). This battery aims at evaluating central auditory processes using non-verbal material so that auditory performances can be related to language competences. Results evidenced that central auditory processes mature until adulthood, and at some point in development (8-9 years old) are linked to language representation. Results of the last study showed that dyslexic adults are impaired in tests involving spectral and temporal processes; in addition these abilities are related with phonological awareness.

Altogether, these results indicate that speech degradation has distinct effects on language depending on its nature. Therefore when temporary, speech degradation impacts its comprehension despite intelligibility. When continuous, speech degradation's impact evolves during development and disappears in normal adults. However, it stays for dyslexics.

Keywords: Speech in noise, semantic processing, Effortfulness Hypothesis, auditory development, dyslexia, APD.

Remerciements

Je tiens tout d'abord à remercier la Professeure Séverine Casalis et le Professeur Frédéric Lavigne pour avoir accepté d'être mes rapporteurs, ainsi que le Professeur Pierre Fournier et le Docteur Sophie Donnadieu pour avoir accepté de faire partie de mon jury de thèse.

Ensuite, mes remerciements immenses vont à Fanny, pour m'avoir fait confiance et soutenue durant ces 6 dernières années. Il m'est souvent arrivé d'entrer dans votre bureau, découragée par la charge de travail et les difficultés, en ayant la certitude d'en ressortir une heure plus tard avec encore plus de travail, mais un nouvel enthousiasme, une grande motivation et hâte de tout recommencer ! Merci aussi pour cette question que vous m'avez posée et qui a marqué un tournant important dans ma thèse et dans ma vie personnelle : « Mais pourquoi tu n'essaierais pas de partir au Canada puisque tu as tant aimé ça ? ».

Je remercie à nouveau Sophie, d'avoir en premier lieu accepté de collaborer avec nous pour déposer ce projet de thèse qui m'a permis de continuer. Merci également de m'avoir confié tes données à analyser, répondu à mes questions et t'être autant investi dans mon travail.

Je remercie également Michel Hoen pour vos conseils, vos réponses aux questions, et pour votre enthousiasme et l'énergie que vous mettez à le communiquer.

Merci également à Brigitte Stemmer, qui m'a accueillie à Montréal, qui s'est battue contre l'administration pour que je puisse avoir un bureau (et que je quitte la cave !) et qui m'a donné accès à tout son matériel EEG.

Merci également à tous les membres du L2C2, à Tatjana et Anne, bien sûr pour nous avoir accueillis il y a maintenant 4 ans. Merci aux autres chercheurs et assimilés, spécialement Yves, pour leur accueil et la bienveillance envers les étudiants que j'ai pu ressentir durant ces 4 années.

D'un point de vue personnel, quoi que je ne sois pas sûre que la thèse soit le meilleur exercice pour séparer vie personnelle et professionnelle, j'ai eu l'immense chance de rencontrer des personnes exceptionnelles qui ont encore enrichie ces dernières années, et qui ont fait que les moments de détente étaient si intenses qu'ils permettaient de se remettre au travail sereinement. Alors merci à tous les étudiants, pour les pauses, les cigarettes, les cafés, les soirées, merci ! merci aux anciens pour leur accueil : Raph, Rawan, Pia, Marion A, ... Merci aux contemporains aussi : Quentin, pour les confidences, les intermèdes musicaux et les joutes verbales ! Marie-Laure (prends un carambar !) Manu (merci pour le chocolat, chocolat, chocolat !), Justine (pour les BBQ, et puis les pauses café), Marion (pour les cosplay têtards ;-)), Audrey, Thomas, Amandine, Auriane, Maël et les autres !

Et merci à ceux qui ont bravé avec moi l'hiver Québécois ! (-40°C quand même !) et sans qui je n'aurais pas survécu ! Mélo, pour les cigarettes, les trajets Montréal-Québec, et pour avoir représenté le petit morceau de Lyon dans tout cet inconnu. Merci à tous les habitants de Barbaland avec une mention très spéciale pour Barbara et pour Tanguy. Barbara, merci de m'avoir fait découvrir ta culture, j'espère que nos aventures continueront encore ici ou là-bas (ou même en Suisse ?). Tanguy merci pour ta présence qui a été pour moi un rayon de soleil ☺, et pour les sorties patins à glace. Merci à vous deux pour le festival de l'érable, le concert des Cowboys fringants et la piscine sur la terrasse ! Merci à Pierrick & Aurélie pour Montréal et Fougères ! Merci aux collègues du CRIUGM, pour les soirées, les cafés et la cabane à sucre ! Merci Bérengère, Ben, Stevan, Kristina, Matthieu, Simona, Christophe, Guillaume.

Et puis surtout, merci aux amis. En premier lieu, Aurèle & Damian, je vous ai déjà tout dit. Mathilde merci pour tes attentions et tes coups de pieds au cul, souvent mérités & avec Jéré, merci pour votre amitié et votre soutien aussi bien dans les délires que dans les coups durs. Merci à Emilie, pour ta bienveillance, ta perspicacité, et ton soutien envers et contre tout(-es les preuves de mes défaillances). Merci Gwen pour nos jeudis soirs, pour tes conseils, et ta capacité à me recadrer dans le réel (finalement c'est mieux que le monde de Marie). Merci à Léo pour ta présence, occasionnelle au bureau, inconditionnelle en dehors. Merci pour les discussions musiques, pour ce soir à Cardiff où je t'ai dit, comme dans la chanson de Vincent Delerm ? Merci pour le concert de Kent, pour la retranscription de Pierre Lapointe et la diffusion de Gilles Vigneaux, et pour avoir été l'exemple vivant qu'il n'y a pas que le labo dans la vie. Il y a l'accordéon et les Fleurs du Malt aussi. Merci aussi à Marjo, Grande Prêtresse de la dyslexie, pour tes conseils, nos délires, 32 fois merci ! ;-)

Merci à ma famille aussi, particulièrement Mamette & Grand-Père, et Jacques. Merci à mes sœurs, qui chacune à leur façon ont été pour moi des modèles. Merci à mon Gigot d'être devenu le frère que je n'avais pas. Et merci Héloïse, Robin, Paul et Gaëtan, simplement d'exister. Merci à mes parents bien sûr, qui m'ont soutenu tout au long de ces longues études et tout au long de tout. Je sais ce que je vous dois.

Enfin, à celui qui est entré dans ma vie, merci d'avoir su gérer mes torrents de larmes intempestifs et répétés, avec flegme et attention. Merci de rester confiant et serein devant mes montagnes d'angoisses sur la thèse, l'après thèse et la vie en général. En bref, merci d'être là, d'être toi et de prendre si bien soin de moi.

Table des matières

Table des matières.....	9
Table des figures.....	15
Partie Théorique.....	17
Introduction Générale.....	19
Chapitre I Le traitement de la parole.....	25
1. Le modèle des Logogènes (Morton, 1969).....	26
2. Le modèle Cohort (Marslen-Wilson & Welsh, 1978).....	27
3. Le modèle Cohort II (Marslen-Wilson, 1987).....	28
4. Le modèle Trace (McClelland & Elman, 1986).....	29
5. Résumé du Chapitre I.....	30
Chapitre II Le traitement sémantique.....	31
1. Modèles de mémoire sémantique.....	32
1.1. Propagation Automatique de l'Activation (ASA ; <i>Automatic Spreading Activation</i>).....	32
1.2. Le modèle stratégique de Becker (1980).....	37
1.3. Les théories post-lexicales.....	39
2. L'amorçage sémantique.....	40
2.1. L'automatisme du traitement sémantique.....	41
2.2. Amorçage sémantique pur ou amorçage associatif ?.....	48
3. La N400.....	49
3.1. Lexicale ou post-lexicale ?.....	50
3.2. Traitement sémantique automatique ou stratégique ?.....	52
4. Les mots ambigus.....	53
4.1. Le <i>context-sensitive model</i>	54
4.2. Le <i>re-ordered access model</i>	55
5. Résumé du Chapitre II.....	56
Chapitre III La parole dans le bruit.....	57
1. Les masquages énergétique et informationnel.....	58
1.1. Le masquage énergétique.....	58
1.2. Le masquage informationnel.....	58
2. Les indices démasquants.....	61
2.1. Les fluctuations temporelles du signal.....	61
2.2. La localisation.....	61
3. La ségrégation de flux.....	61
4. Parole dans le bruit et <i>Effortfulness Hypothesis</i>	62
5. Résumé du Chapitre III.....	64
Chapitre IV Le système auditif.....	67
1. Le système auditif.....	68
1.1. L'oreille externe.....	68

1.2.	L'oreille moyenne	69
1.3.	L'oreille interne.....	70
2.	Le système auditif central	71
2.1.	Le noyau cochléaire	72
2.2.	Le complexe olivaire supérieur	72
2.3.	Le lemnicus latéral.....	73
2.4.	Le colliculus inférieur	73
2.5.	Le corps genouillé médian du thalamus.....	73
2.6.	Le cortex auditif primaire.....	74
3.	Le développement des traitements auditifs centraux.....	75
3.1.	Aspects électrophysiologiques	76
3.2.	Développement comportemental	78
3.3.	Les facteurs non-sensoriels	83
4.	Les Troubles du Traitement Auditif (TTA)	84
4.1.	Définitions.....	84
4.2.	Diagnostic	85
4.3.	La Batterie d'Evaluation des Compétences Auditives Centrales	87
5.	Résumé du Chapitre IV	90
	Chapitre V La dyslexie.....	91
1.	Définition et diagnostic	92
2.	Classification des dyslexies.....	93
2.1.	Le modèle dual route.....	93
2.2.	Dyslexie phonologique et dyslexie de surface	94
3.	Le déficit phonologique.....	95
3.1.	La conscience phonologique	95
3.2.	La perception allophonique	96
4.	Les théories d'un déficit du traitement auditif	98
4.1.	La théorie d'un déficit des traitements auditifs rapides	99
4.2.	La théorie du déficit de la perception du rythme	101
4.3.	La théorie du déficit d'ancrage (<i>Anchor Theory</i>)	102
4.4.	La théorie de l'instabilité du traitement auditif.....	104
4.5.	Les TTA dans la dyslexie développementale	106
5.	Résumé du Chapitre V.....	108
	Synthèse et problématique générale	111
	Partie Expérimentale.....	113
	Chapitre VI Etude 1 : L'amorçage sémantique en situation de parole dans la parole ..	115
1.	Introduction.....	118
2.	Experiment 1.....	121
2.1.	Method	122
2.2.	Results	124
2.3.	Discussion	125
3.	Experiment 2	126

3.1.	Method	126
3.2.	Results	128
3.3.	Discussion	129
4.	Experiment 3	129
4.1.	Method	130
4.2.	Results	132
4.3.	Discussion	134
5.	Post-hoc analyses	134
6.	General Discussion	135
6.1.	Summary	135
6.2.	Lexical processing without semantic activation	136
6.3.	Cognitive load	137
7.	Conclusion	139
8.	Résumé du Chapitre VI	140
Chapitre VII Etude 2 : Le traitement des mots ambigus en situation de parole dans le bruit.....		141
1.	Introduction	144
1.1.	Speech in noise situation.....	144
1.2.	Effortfulness Hypothesis	145
1.3.	Ambiguous words	145
1.4.	Semantic processing in sentence context.....	146
1.5.	The present study	147
2.	Material & Methods	148
2.1.	Participants	148
2.2.	Stimuli	148
2.3.	EEG recordings	151
2.4.	Data Analysis	151
3.	Results	152
3.1.	Behavioral results.....	152
3.2.	EEG data.....	153
4.	Discussion	155
4.1.	Summary	155
4.2.	Behavioral results.....	156
4.3.	Congruency effect.....	156
4.4.	Ambiguity effect	157
4.5.	Phonological Mismatch Negativity.....	158
5.	Conclusion	159
6.	Résumé du Chapitre VII	160
De l'axe 1 à l'axe 2.....		161
Chapitre VIII Etude 3a. Les liens entre développement des traitements auditifs centraux et le langage		163
1.	Introduction.....	165

2.	Method	167
2.1.	Participants	167
2.2.	Verbal tests.....	167
2.3.	Auditory tests.....	168
2.4.	Procedure	169
3.	Results	169
3.1.	Auditory Tests.....	169
3.2.	Verbal tests.....	170
3.3.	Correlations between verbal skills and nonverbal tests	171
4.	Discussion	172
5.	Conclusion	174
6.	Résumé du Chapitre VIII.....	175
Chapitre IX Etude 3b. Le développement des traitements auditifs centraux		177
1.	Introduction	179
2.	Method	182
2.1.	Participants	182
2.2.	Auditory tests.....	183
2.3.	General procedure	187
3.	Results	187
3.1.	General remarks.....	188
3.2.	Lateralization	189
3.3.	Auditory Discrimination	190
3.4.	Central Masking.....	193
3.5.	Auditory identification and recognition.....	193
3.6.	Auditory Pattern Test (APT)	194
3.7.	Stream segregation	194
3.8.	Inter-individual differences.....	194
4.	Discussion	195
5.	Conclusion	199
6.	Résumé du chapitre IX	201
Chapitre X Etude 4. Les Troubles du Traitement Auditif dans la dyslexie		203
1.	Introduction	205
2.	Method	209
2.1.	Participants	209
2.2.	Auditory tests.....	209
2.3.	Language evaluation.....	212
2.4.	Procedure	212
3.	Tests Results.....	214
3.1.	Lateralization	214
3.2.	Discrimination tests	216
3.3.	Auditory Identification and auditory recognition	216

3.4.	Auditory Pattern Test (APT)	217
3.5.	Stream Segregation	217
3.6.	Prevalence of APD	217
4.	Correlations.....	219
4.1.	Frequency Discrimination.....	219
4.2.	Duration Discrimination	220
4.3.	Auditory pattern test. Frequency	221
4.4.	Auditory pattern test. Duration	222
5.	Discussion	223
6.	Conclusion	226
7.	Résumé du Chapitre X.....	227
	Discussion.....	229
	Chapitre XI Discussion de l'Axe 1	231
1.	Résumé des résultats	231
1.1.	Etude 1	231
1.2.	Etude 2.....	232
1.3.	Effet du masquage informationnel ou énergétique ?.....	234
2.	Traitement sémantique et <i>Effortfulness Hypothesis</i>	234
2.1.	<i>Effortfulness Hypothesis</i> et amorçage masqué	234
2.2.	Retour sur les modèles de traitement sémantique	235
2.3.	Automaticité et ressources cognitives	236
3.	Interactions entre processus <i>top-down</i> et <i>bottom-up</i>	237
3.1.	Au niveau comportemental.....	237
3.2.	Au niveau électrophysiologique.....	238
4.	Parole dans le bruit et rythme alpha	239
	Chapitre XII Discussion de l'Axe 2	243
1.	Résumé des résultats	243
1.1.	Etude 3.....	243
1.2.	Etude 4	244
2.	Les facteurs non-sensoriels	245
2.1.	Diminution des facteurs attentionnels et cognitifs	245
2.2.	Diminution de l'effet du paradigme	246
3.	Retour sur l' <i>Anchor Theory</i>	247
4.	Retour sur les TTA dans la dyslexie	248
5.	Critiques	249
6.	Liens entre compétences auditives centrales et langagières	251
	Chapitre XIII Discussion Générale	253
1.	Récapitulatif	253
2.	<i>Parole dans le bruit et dyslexie</i>	255
3.	Traitement sémantique et dyslexie.....	256
	Conclusion	259
	Références.....	261

Annexes.....	287
Curriculum Vitae	289
Liste de stimuli Etude 1.....	291
Liste de stimuli Etude 2	299

Table des figures

Figure 1	33
Figure 2	34
Figure 3	36
Figure 4	43
Figure 5	50
Figure 6	59
Figure 7	60
Figure 8	63
Figure 9	69
Figure 10	70
Figure 11	74
Figure 12	75
Figure 13	77
Figure 14	81
Figure 15	94
Figure 16	97
Figure 17	98
Figure 18	99
Figure 19	100
Figure 20	102
Figure 21	103
Figure 22	105
Figure 23	123
Figure 24	127
Figure 25	131
Figure 26	132
Figure 27	135
Figure 28	150
Figure 29	153
Figure 30	154
Figure 31	155
Figure 32	170
Figure 33	171
Figure 34	184
Figure 35	191
Figure 36	192

Figure 37.....	195
Figure 38	211
Figure 39	215
Figure 40	218
Figure 41.....	220
Figure 42	221
Figure 43	222
Figure 44.....	223
Figure 45	239
Figure 46.....	250
Figure 47	254

Partie Théorique

Introduction Générale

Le langage est traditionnellement étudié indépendamment des processus sensoriels. La linguistique s'intéresse à sa structure ; la psycholinguistique, à la façon dont l'homme le traite, le manipule, le comprend, l'apprend et le produit. Les processus les plus bas niveaux étudiés sont les traitements phonologique ou orthographique selon que l'on s'intéresse à la parole ou à l'écriture. Les compétences sensorielles jouent pourtant un rôle primordial dans notre capacité à apprendre le langage. Ainsi, les nourrissons sourds présentent souvent un retard de langage même après la mise en place d'un implant cochléaire et particulièrement si l'implantation a eu lieu après l'âge de deux ans. Il semble également que ces enfants aient moins d'intérêt et prêtent moins attention au signal parolier même après avoir été équipés d'un implant cochléaire. En effet, Houston, Pisoni, Kirk, Ying et Miyamoto (2003) grâce à un paradigme d'habituation visuelle, ont montré un plus grand intérêt de la part des nourrissons pour un stimulus visuel lorsque celui-ci est accompagné d'un son de parole. Dans une phase d'habituation les auteurs présentaient simultanément un stimulus visuel et un son (e.g., *hop hop hop*) ou uniquement un stimulus visuel. Une fois les nourrissons habitués à ces stimuli (i.e., diminution des temps de fixation), de nouveaux stimuli (e.g., *ah ah ah*) étaient présentés parmi les anciens. Les résultats montrent que les nourrissons regardent généralement davantage les nouveaux stimuli que les anciens. Ils regardent également davantage le stimulus visuel lorsque celui-ci est présenté en association avec le stimulus sonore. La différence est toutefois plus importante pour les nourrissons normo-entendants que pour les nourrissons ayant un implant cochléaire, même 6 mois après sa mise en place. Ces résultats suggèrent une attention moindre au signal de parole chez les nourrissons implantés malgré une nouvelle capacité à le percevoir. Ce manque d'intérêt pourrait alors avoir un effet délétère ajouté (i.e., en plus de la faible qualité du signal perçu). Un phénomène semblable a été décrit chez les nourrissons ayant souffert d'otites avec effusion. Cette pathologie otologique est très fréquente chez l'enfant et génère une hypoacousie temporaire. Polka et Rvachew (2005) ont montré que les nourrissons de 6 à 8 mois ayant une histoire d'otites à effusion présentent, même après recouvrement de l'audition normale, moins d'intérêt

pour la parole. Ces résultats illustrent ainsi comment la diminution des capacités perceptives impacte le traitement du langage et ce, même lorsque la perte auditive est relative et transitoire.

Toutefois, en dehors de ces cas pathologiques, les liens entre perception auditive et langage restent très peu étudiés. La conception du langage reste modulaire (Fodor, 1983), ce dernier étant considéré comme une entité indépendante et de haut niveau par opposition aux processus bas niveau (i.e., les processus perceptifs). Certains pans de la littérature développementale s'intéressent toutefois davantage aux liens entre processus perceptifs et cognitifs (Baltes & Lindenberger, 1997; Lindenberger, Scherer, & Baltes, 2001).

Cette thèse a pour objectif de contribuer à l'exploration des liens entre traitements perceptifs et cognitifs. Elle s'intéresse plus précisément à l'impact de la dégradation du signal de parole sur différents aspects du langage. Tout au long de ce travail, nous distinguons deux types de dégradation du signal de parole. Une première dégradation, transitoire, provenant du signal lui-même, dont l'exemple le plus parlant est la situation de parole dans le bruit (Cherry, 1953) dans laquelle le signal de parole est intégré dans un bruit (parolier ou non parolier). Un premier axe étudie ainsi l'impact de ces différents types de bruit sur le traitement de la parole. Particulièrement, nous distinguons les concepts d'intelligibilité et de compréhension du signal en nous intéressant à la modulation de la profondeur du traitement sémantique lorsque le signal est dégradé mais l'intelligibilité préservée. Le second axe concerne les dégradations permanentes provenant d'un mauvais traitement du signal. Dans ce cas, les traitements auditifs centraux déficients engendrent la perception d'un signal de mauvaise qualité. Nous étudierons ainsi l'impact de cette dégradation du signal sur les compétences langagières.

Ce travail de thèse est présenté en trois parties.

La première présente l'état de la littérature en 5 chapitres. Le chapitre premier s'intéresse particulièrement au traitement de la parole, via la présentation de 4 modèles de traitement du mot. Le second chapitre s'intéresse au traitement sémantique. Outre la présentation de différents modèles de traitement sémantique, il expose les mécanismes d'amorçage sémantique aussi bien au niveau comportemental qu'électrophysiologique avec l'étude du potentiel évoqué N400. Ce chapitre aborde également le cas spécifique des mots ambigus, et leur intérêt dans la question de l'accès au sens des mots. Dans un troisième chapitre nous nous intéressons à la situation de parole dans le bruit. Nous voyons tout d'abord les différents effets masquant présents en situation de parole dans le bruit, puis les indices bas niveau permettant d'améliorer l'intelligibilité. Enfin nous nous intéressons à l'*Effortfulness Hypothesis*, qui postule que la dégradation du signal de parole par le bruit a un impact

sur les traitements cognitifs de haut niveau, au-delà de la perte d'intelligibilité. Le chapitre IV a pour but de présenter les systèmes auditifs périphérique et central. Dans la seconde partie de ce chapitre IV, nous détaillons spécifiquement les traitements auditifs centraux, leur développement normal, leurs dysfonctionnements et les outils diagnostiques disponibles. Enfin, le chapitre V clôture cette première partie en présentant le cas spécifique de la dyslexie, en tant que conséquence possible des troubles du traitement auditif. Une proportion relativement importante de la population dyslexique (entre 40 et 60%) présentant des Troubles du Traitement Auditif (TTA ; Ahissar, 2007; Baldeweg, Richardson, Watkins, Foale, & Gruzelier, 1999; Goswami et al., 2002; Hornickel & Kraus, 2013; Ramus et al., 2003; Tallal & Gaab, 2006).

La seconde partie de ce manuscrit présente le travail expérimental effectué pendant la thèse en 4 études. Les deux premières s'inscrivent dans le premier axe de ce travail et s'intéressent à l'impact de la dégradation du signal de parole du fait de la présence de bruit sur le traitement sémantique au niveau du mot parlé. Une première étude, comprenant 3 expériences comportementales en situation de parole dans la parole a testé l'impact des interférences sémantiques sur le traitement sémantique d'un mot cible. Les participants devaient effectuer une tâche de décision lexicale sur un item cible inséré dans un bruit parolier composé de 1 à 4 voix. Les résultats obtenus mettent en évidence un effet du nombre de voix mais également de la saillance des amorces. Cette première étude a ainsi permis de mettre en évidence que le traitement sémantique semble reposer sur des processus moins automatiques que ce que la littérature pourrait suggérer et serait dépendant des ressources cognitives disponibles comme suggéré par l'*Effortfulness Hypothesis* (Dekerle, Boulenger, Hoen, & Meunier, 2014).

L'Etude 2 s'est intéressée à l'impact de la dégradation du signal sur le traitement sémantique d'un mot dans un contexte phrastique. Cette étude électrophysiologique a été mise en place en collaboration avec la Professeure Stemmer à l'Université de Montréal où j'ai passé 6 mois (bourse EXPLORA DOC, région Rhône-Alpes). L'objectif était de tester la modulation de l'activation des différents sens de mots ambigus (e.g., *avocat* le fruit ou le métier). Ils étaient intégrés dans des phrases favorisant leur sens dominant (i.e., le plus fréquent) et se terminant par un mot cible. Celui-ci pouvait être cohérent soit avec le sens dominant du mot ambigu soit avec le sens subordonné du mot ambigu ou encore être totalement incohérent. Les phrases étaient présentées dans 3 conditions sonores : bruit Parolier, bruit Stationnaire ou Silence. L'étude du potentiel évoqué N400 en réponse au mot cible a permis de mettre en évidence qu'en condition Silence, seul le sens dominant du mot ambigu était activé, en revanche, en condition Stationnaire, il semble que les deux sens du mot aient été activés. En condition bruit Parolier, comme dans l'Etude 1, le traitement sémantique n'a pu être mis en évidence, en dépit d'une intelligibilité préservée. Il semble alors que, comme le prédit

l'*Effortfulness Hypothesis*, la force d'activation du traitement sémantique dépende de l'intelligibilité. En condition Silence, le contexte peut être parfaitement traité sémantiquement, résultant que seul le sens pertinent du mot ambigu est activé (i.e., le sens dominant). En revanche, lorsque la phrase était présentée dans du bruit Stationnaire, le contexte était vraisemblablement traité moins efficacement, posant ainsi moins de contraintes sur le sens pertinent du mot ambigu. Cette deuxième étude a ainsi permis de mettre en évidence que la profondeur du traitement du contexte varie en fonction du bruit environnant jusqu'à ce que le traitement sémantique ne soit plus mesurable, lorsque la dégradation du signal de parole devient trop importante bien que l'intelligibilité soit préservée.

Le second axe de ce travail porte sur la dégradation permanente du signal du fait d'un déficit au niveau des traitements auditifs centraux et aux conséquences de ce déficit sur les représentations langagières. Les traitements auditifs centraux font référence à tous les traitements auditifs effectués au niveau central et non périphérique. Dans la troisième étude de ce manuscrit, nous avons étudié le développement normal de ces traitements grâce à la Batterie d'Evaluation des Compétences Auditives Centrales (BECAC ; Donnadieu, Gillet-Perret, Lassus-Sangosse, & Nguyen-Morel, 2014). J'ai ainsi analysé les données recueillies par Donnadieu et son équipe auprès des enfants au développement typique, répartis en 5 classes (du CP au CM2). Cette batterie constituée de 6 tests principaux permet d'évaluer les traitements auditifs centraux comme ils sont définis par l'American Speech-Language-Hearing Association (ASHA ; 2005) sans utiliser de matériel verbal. Cette étude a fait l'objet de deux articles. Le premier ne présentant qu'une partie des résultats recueillis grâce à la batterie et les liens entre les compétences auditives testées et les compétences langagières. Elle a permis de mettre en évidence que certaines compétences auditives testées étaient liées aux compétences phonologiques et au niveau de vocabulaire des enfants, uniquement sur les classes de CE2 et CM1. Une ré-analyse des données a donné lieu à l'écriture d'un second article, s'intéressant au développement des compétences de discrimination auditive, latéralisation, masquage central, reconnaissance de patterns auditifs, ségrégation de flux, identification et reconnaissance auditive. Les résultats ont mis en évidence un effet du développement sur tous les tests sauf la ségrégation de flux. De plus ce développement n'est ni homogène ni linéaire et se poursuit durant l'adolescence. En effet, lorsque les données sont comparées à celles d'un groupe d'adultes typiques, seuls deux des traitements évalués sont matures avant l'adolescence (la latéralisation, et la discrimination de durée).

Enfin, l'Etude 4 s'est intéressée à l'impact de la qualité des traitements auditifs centraux sur les compétences langagières au sein d'une population d'adultes sains et d'une population d'adultes dyslexiques. Ces TTA sont considérés par certains auteurs comme étant à l'origine du trouble phonologique générant les difficultés de lecture. J'ai

ainsi adapté la BECAC à une population adulte et l'ai administrée à un échantillon de 20 adultes dyslexiques et 20 adultes sains appariés en âge, en latéralité manuelle et en intelligence non-verbale. Les résultats ont permis de confirmer l'existence de déficits dans certains traitements auditifs centraux dans la population dyslexique, spécifiquement sur les tests impliquant des traitements de fréquence et de durée. Il est de plus apparu que ces compétences étaient liées à certaines compétences langagières (i.e., lecture et conscience phonologique), principalement au sein du groupe dyslexique. Il semble ainsi que la dégradation du signal de parole, même provenant d'un traitement auditif central déficitaire a des conséquences sur le langage. En particulier, il apparaît que le lien entre la qualité des traitements bas niveau et les représentations langagières reste très fort dans la population dyslexique ce qui ne semble pas être le cas pour les participants sains.

Enfin, la dernière partie du manuscrit permet de réaliser la synthèse des résultats de la partie expérimentale et de les discuter à la lumière de la littérature scientifique.

Chapitre I

Le traitement de la parole

Chez l'Homme, la perception de la parole débute dès la vie intra-utérine, les nourrissons sont capables de distinguer leur langue maternelle d'une autre dès 4 jours (Mehler et al., 1988). Jusqu'à environ 5-6 ans, l'essentiel voire toute l'information linguistique reçue par l'enfant est auditive. Même à l'âge adulte, la plupart des échanges se font oralement. Toutefois, cette asymétrie entre langage oral et écrit ne se reflète pas dans la littérature scientifique, elle est même inversée puisque la plupart des études s'intéressent au langage en modalité visuelle et donc à l'écrit. Il a ainsi longtemps été admis qu'un phénomène observable en modalité visuelle pouvait également l'être en modalité auditive. La lecture et le traitement de la parole étaient donc considérés comme sous-tendus par des processus similaires. Le traitement de la parole présente pourtant des spécificités par rapport à la lecture. En premier lieu simplement parce que la nature même du signal traité est très différente ; la parole est un signal qui ne persiste pas dans le temps. Son traitement doit alors être suffisamment rapide et efficace pour que le contenu du discours soit entendu, reconnu, compris, intégré alors (a) qu'il va disparaître et (b) que de nouvelles informations vont être présentées et devront être traitées. En effet, le rythme naturel moyen de production de parole est compris entre 200 et 300 mots par minutes. Une seconde contrainte spécifique à la parole provient du fait que celle-ci est un signal continu. Le flux de parole doit donc être segmenté en unités lexicales distinctes afin que celles-ci soient reconnues et comprises.

Dans ce premier chapitre nous présenterons 4 modèles du traitement du langage et plus particulièrement de reconnaissance des mots parlés (puisque c'est à cette échelle que se situent nos études expérimentales). Chacun de ces modèles part du même postulat : il existe différents niveaux de traitements. Tout d'abord le niveau phonétique, puis phonémique, puis lexical. Les modèles présentés ici sont des modèles d'activation parallèle, qui postulent que plusieurs candidats à la reconnaissance sont activés parallèlement et entrent en compétition jusqu'à la reconnaissance de l'input.

Nous nous intéressons particulièrement aux prédictions des modèles sur les interactions entre les processus *bottom-up* et *top-down*. En effet, ces interactions sont particulièrement pertinentes dans notre travail puisque notre partie expérimentale s'intéresse en partie aux liens entre les traitements bas (i.e., acoustique, phonologique) et haut niveau (i.e., sémantique). Le modèle des Logogènes de Morton, ne s'intéresse qu'aux informations *bottom-up*. Les deux versions du modèle Cohort donnent une importance aux informations de haut niveau, mais uniquement au moment de la sélection du candidat lexical. Enfin, le modèle TRACE lui, postule des interactions permanentes entre les différents niveaux de traitements, et donc des activations à la fois *bottom-up* et *top down*.

Les mécanismes d'accès au sens des mots traités, bien qu'ils soient inhérents au traitement de la parole, ne seront pas évoqués au sein de ce chapitre. En effet, la mémoire sémantique est décrite comme amodale, le traitement sémantique n'est donc pas spécifique à la compréhension de la parole. Les différents modèles de traitement sémantique seront donc présentés indépendamment, au sein du Chapitre II (cf. p32).

1. Le modèle des Logogènes (Morton, 1969)

L'unité de base du modèle de Morton est appelée le logogène (en Grec *logo* signifie mot et *gene* signifie naissance). Un logogène est un système, une unité, défini par un certain nombre de caractéristiques. Cette unité reçoit des informations à propos des caractéristiques sensorielles et contextuelles du stimulus. Les caractéristiques définissant un logogène sont divisées en trois sets [S], [V] et [A] correspondants aux caractéristiques Sémantique, Visuelle et Acoustique du mot représenté par le logogène. Lorsque suffisamment de caractéristiques sont partagées par le stimulus et le logogène, et qu'un certain seuil d'activation est dépassé, alors le mot est reconnu. Afin d'expliquer l'effet de fréquence, Morton suppose que le seuil d'activation est plus bas pour les mots ayant une haute fréquence d'occurrence que pour les mots ayant une basse fréquence d'occurrence. Un mot fréquent sera ainsi reconnu plus rapidement et avec moins d'informations qu'un mot peu fréquent. La décision d'activer le logogène se fait de façon intrinsèque au logogène, c'est-à-dire qu'il n'y a pas de compétitions directes (i.e., d'inhibition) entre les différents logogènes. Théoriquement plusieurs logogènes pourraient s'activer simultanément s'ils atteignaient chacun leurs seuils respectifs. Toutefois Morton contourne ce problème en stipulant que le buffer de sortie ne peut contenir qu'une seule réponse, seul le premier logogène à être activé sera ainsi sélectionné. De manière générale, Morton explique l'ensemble des effets de facilitations ou d'inhibitions qui peuvent être observés par un abaissement ou une augmentation du seuil de reconnaissance du logogène concerné.

2. Le modèle Cohort (Marslen-Wilson & Welsh, 1978)

Le modèle Cohort est le premier modèle conçu spécifiquement pour la modalité auditive. Il s'intéresse particulièrement à la question des interactions entre les processus *top-down* et *bottom-up*. C'est-à-dire, à la façon dont les informations acoustiques et linguistiques portées par le signal de parole interagissent avec les attentes du sujet (les contraintes sémantiques et syntaxiques portées par le contexte) afin de permettre la compréhension de la parole. D'après Marslen-Wilson et Welsh (1978), il est extrêmement important de considérer aussi bien les processus *top-down* que *bottom-up* puisque le signal de parole est très souvent dégradé (e.g., par du bruit, voir Chapitre III p58). Les auteurs suggèrent donc que l'interaction entre processus *top-down* et *bottom-up* se fait directement pendant la perception du signal de parole et ne provient pas simplement d'une ré-analyse du signal après que celui-ci ait été traité uniquement de façon *bottom-up*. Pour tester cette hypothèse, ils présentent à un groupe de participants des phrases dont un des mots contient un phonème déviant (sur un ou trois traits phonétiques) soit en première syllabe soit en dernière syllabe. De plus, ce mot peut être placé dans un contexte restrictif ou non restrictif. Les résultats ont mis en évidence que les participants, lorsqu'ils effectuaient une tâche de *shadowing* (i.e., les participants doivent répéter le signal de parole entendu au moment même où ils l'entendent) ne répétaient pas les déviations phonémiques alors qu'ils étaient capables de les détecter lorsque cela leur était demandé. Particulièrement, l'effet du contexte était significatif en *shadowing* et non en détection de phonèmes, suggérant ainsi que le contexte influence le percept du participant.

D'après Marslen-Wilson et Welsh (1978), 3 processus sont nécessaires à la reconnaissance d'un mot. Tout d'abord l'accès, puis la sélection et enfin l'intégration.

L'accès reposerait uniquement sur des processus *bottom-up*, c'est-à-dire qu'il serait totalement dépendant de l'input. Lorsque les deux ou trois premiers phonèmes d'un mot sont entendus, il y aurait une activation parallèle de toute une classe d'éléments de mémoire lexicale, i.e., la cohorte. La présentation des phonèmes /é/, /l/, /é/ va ainsi activer les éléments de mémoire lexicale /élément/, /éléphant/, /élection/... Lorsque le signal de parole continue, les éléments de mémoire lexicale ne correspondant plus au signal se désactivent les uns après les autres jusqu'à ce qu'il n'en reste plus qu'un : c'est le point de reconnaissance ou le point d'unicité. Le mot entendu est alors reconnu. Bien évidemment, le modèle propose que la correspondance entre l'input et la représentation du mot n'a pas besoin d'être parfaite ce qui explique les effets de non détection des changements de traits phonétiques. Toutefois, ce postulat ne permet pas d'expliquer les effets de contexte. Ici le modèle ne prend en compte que la reconnaissance du mot présenté en isolation et ne s'intéresse ainsi qu'aux effets *bottom-up*. Afin de prendre en compte l'effet du contexte, Marslen-Wilson et Welsh

(1978) suggèrent simplement que chaque élément de mémoire lexicale est informé des contraintes posées par le contexte. Ainsi, les mots incorrects sémantiquement ou syntaxiquement avec le contexte sont automatiquement désactivés. Aussi bien les processus *bottom-up* que *top-down* ont donc la possibilité de réduire la cohorte. Les contraintes du contexte permettent de diminuer la cohorte et de reconnaître plus rapidement le mot présenté. De plus, la présence du contexte confirme le choix du mot. Plus le contexte est contraignant, plus il renforce la validité du choix du mot. La force du contexte explique ainsi la non-détection des digressions phonétiques dans l'expérience présentée par Marslen-Wilson et Welsh (1978). Ce modèle reste toutefois avant tout *data-driven*, c'est-à-dire que les processus *top-down* ont un rôle de vérification, et de renforcement mais en aucun cas ne permettent de générer une attente sur le mot.

La sélection du mot (toujours parmi la cohorte de candidats) se situe au niveau de l'interaction entre les processus *bottom-up* et *top-down*. Les auteurs prennent pour évidence d'une telle interaction le phénomène d'*early selection* ou sélection précoce, lorsqu'un mot présenté dans un contexte est reconnu plus rapidement que lorsqu'il est présenté isolément. Cette accélération des processus de reconnaissance en présence d'un contexte suggère en effet que celui-ci joue un rôle dans la facilitation de la sélection d'un candidat au sein de la cohorte. En effet, Cohort considère, en l'absence de contexte que la reconnaissance d'un mot se fait au point de reconnaissance c'est-à-dire lorsqu'un mot est le seul qui reste de la cohorte initiale. Ce point de reconnaissance permet également de déterminer à partir de quel moment un item pourra être qualifié de pseudo-mot. En effet, lorsque le point de reconnaissance est passé et que plus aucun mot ne reste dans la cohorte alors seulement le système peut décider que l'item entendu est un pseudo-mot.

L'intégration en revanche serait complètement conduite par les processus *top-down*. Ce processus correspond à l'intégration du mot sélectionné dans le contexte.

3. Le modèle Cohort II (Marslen-Wilson, 1987)

L'un des problèmes de Cohort est de ne pas expliquer l'effet de fréquence. Pourtant, cet effet a clairement un impact sur le processus de sélection. En effet, après avoir présenté à des participants le début de mot CAPT (pouvant être *CAPTAIN* ou *CAPTIVE*), Zwitserlood (1989) a mis en évidence premièrement que cette présentation incomplète amorçait à la fois des mots sémantiquement liés à *CAPTAIN* (e.g., *ship*) et à *CAPTIVE* (e.g., *guard*). Ceci confirme donc le postulat effectué dans la première version de Cohort, que les caractéristiques sémantiques des candidats sont activées avant la sélection et intégrées à celle-ci. Deuxièmement, l'effet de facilitation était plus important pour *ship* qui est sémantiquement lié à *captain*, lui-même plus fréquent que

captive. En revanche, lorsque le mot lié est présenté tardivement, c'est-à-dire après la présentation complète du mot cible, cet effet disparaît. Seuls les mots liés au mot présenté sont activés. La nouvelle version de Cohort prend donc en compte cet effet en postulant que les mots fréquents auront une activation relative plus importante que les mots peu fréquents.

C'est dans cette modulation de l'activation que se joue la principale différence entre Cohort I et Cohort II. En effet, Cohort I prévoyait une activation en tout ou rien, où un mot était reconnu lorsqu'il était le seul candidat restant dans la cohorte. Toutefois, ce premier modèle ne laissait pas de possibilité d'erreurs de prononciation. Ainsi, si quelqu'un prononçait /pato/ au lieu de /bato/, le mot bateau ne pouvait pas être reconnu, puisque la perception de /p/ et /a/ aurait automatiquement engendré l'activation d'une cohorte dont bateau aurait été exclu. Dans Cohort II, il est possible que /bato/ soit activé et fasse partie de la cohorte initiale. Il sera simplement moins activé au début que d'autres mots comme /patin/. Puis, les interactions entre les processus *top-down* et *bottom-up* devraient permettre de pouvoir reconnaître /bato/ même s'il ne correspond pas parfaitement à l'input sensoriel. De même, un mot ne correspondant pas aux attentes syntaxiques et sémantiques ne pouvait pas être reconnu dans la première version de Cohort. Cette deuxième version suggère qu'un mot ne correspondant pas aux attentes sémantiques et syntaxiques pourrait tout à fait être reconnu dès lors que l'input est clair.

4. Le modèle Trace (McClelland & Elman, 1986)

Le modèle TRACE est un modèle spécifique à la compréhension de la parole. Il est basé sur le principe d'activation interactive. Le traitement de l'information se fait via des unités de traitement reliées entre elles par des liens excitateurs ou inhibiteurs. Chaque unité de traitement correspond à une hypothèse. Si deux unités de traitements soutiennent des hypothèses compatibles alors les liens entre elles seront excitateurs. Si deux unités de traitements soutiennent des hypothèses incompatibles alors les liens entre elles seront inhibiteurs. Ce principe d'activation et d'inhibition continue jusqu'à ce qu'un certain seuil soit atteint, qui va valider la dite hypothèse.

Les unités de traitement sont réparties en trois niveaux : les traits phonétiques, les phonèmes et les mots. La présentation du signal va activer tel ou tel trait phonétique, qui à son tour va activer les phonèmes comprenant ce trait phonétique. Les phonèmes vont ensuite activer les mots qui les contiennent. Afin de contourner le problème de difficulté de segmentation des mots dans la parole, les auteurs postulent que chaque phonème va activer l'ensemble des mots qui le contiennent, peu importe l'endroit où il se situe dans le mot. Par exemple la perception du phonème /a/ va générer l'activation de mots comme /ami/, /chat/ ou encore /tomate/. Les liens entre les différents niveaux

sont excitateurs, qu'ils soient *bottom-up* ou *top-down*. Par exemple, le niveau lexical peut exciter certaines unités au niveau phonémique. De ce fait, ce modèle postule une multiplication des lexiques mentaux puisque d'autres activations se mettent en place lorsque d'autres phonèmes seront entendus. Cette multiplication, implique que le système devrait gérer une très grande quantité d'informations ce qui rend le modèle peu probable.

5. Résumé du Chapitre I

Les 4 modèles de reconnaissance des mots parlés présentés dans ce chapitre émettent donc différentes hypothèses au regard des interactions entre les processus *bottom-up* et *top-down*. Le modèle des Logogènes postule que seuls des processus *bottom-up* sont impliqués dans la reconnaissance des mots. D'après Cohort I et II, les processus *top-down* ont un rôle uniquement pendant la phase de sélection. Enfin, le modèle TRACE pose l'hypothèse que les processus *top-down* interagissent avec les processus *bottom-up* à tous les niveaux. La notion d'interaction entre processus *bottom-up* et *top-down* est particulièrement importante dans l'étude de la parole dans le bruit. En effet, lorsque le signal de parole est dégradé, les informations *bottom-up* peuvent être insuffisantes pour un traitement efficace du signal. Des informations de plus haut niveau (e.g., contexte sémantique) pourront alors être utilisées afin de faciliter le traitement du mot. Ces hypothèses seront testées dans l'Axe 1, Etudes 1 et 2.

Chapitre II

Le traitement sémantique

Ce chapitre s'intéresse spécifiquement à la compréhension de la parole au sens propre, c'est-à-dire à l'accès au sens. En effet, l'un des objectifs de ce travail de thèse est d'étudier l'effet de la dégradation du signal de parole sur l'accès au sens. Nous verrons donc tout d'abord les différents modèles de traitement sémantique en situation de perception optimale. Ils se divisent en trois grandes catégories : les modèles automatiques, les modèles stratégiques et les modèles post-lexicaux. Les modèles automatiques postulent que l'accès au sens se fait de manière automatique, c'est-à-dire rapidement, sans que le sujet en ait conscience et indépendamment d'éventuelles autres tâches qui pourraient retenir l'attention du système cognitif. Les modèles stratégiques considèrent pour leur part que le traitement sémantique se fait de façon intentionnelle et dépendamment de l'attention du sujet. Ces deux types de modèles présentent toutefois un point commun : ils postulent que l'accès au sens se fait simultanément à l'accès lexical. Le troisième type de modèle présenté, les modèles post-lexicaux, postulent que l'accès au sens se fait après l'accès lexical. Une fois ces principaux modèles exposés, nous nous intéressons à la littérature sur l'amorçage sémantique, et plus particulièrement les deux principaux débats qui animent la communauté scientifique et reflètent les interrogations soulevées par les modèles : la question de l'automaticité du traitement sémantique (Marcel, 1983; Neely, 1991; Van den Bussche, Van den Noortgate, & Reynvoet, 2009) et la nature de l'organisation de la mémoire sémantique (Hutchison, 2003; Lucas, 2000). La nature automatique ou stratégique du traitement sémantique est particulièrement importante au sein de ce travail. Dans cette optique, les deux études présentées dans l'Axe 1 de la partie expérimentale (Etude 1 et Etude 2) portent sur l'effet de la dégradation du signal sur le traitement sémantique. Cette dégradation devrait avoir un effet différent selon que le traitement sémantique se fait automatiquement ou dépendamment des ressources cognitives des participants.

La troisième partie de ce chapitre s'intéresse aux corrélats électrophysiologiques du traitement sémantique, en présentant le potentiel évoqué N400. Celui-ci est utilisé comme indice du traitement sémantique durant l'Etude 2. Nous présentons donc les débats présents dans la littérature au regard de la nature des processus sous-tendant la N400, c'est-à-dire leur automaticité, et leur nature lexicale ou post-lexicale. Enfin, nous clôturons ce chapitre par l'étude du cas particulier des mots ambigus, étudiés dans l'Etude 2 afin d'évaluer la profondeur du traitement sémantique.

1. Modèles de mémoire sémantique

1.1. Propagation Automatique de l'Activation (ASA ; *Automatic Spreading Activation*)

1.1.1. Le modèle de Quillian (1967)

Le modèle de Quillian a été développé pour être modélisé sur ordinateur et ne peut donc répondre à toutes les exigences posées par les données expérimentales. Il est basé sur l'idée de propagation de l'activation depuis un ou plusieurs concepts jusqu'à ce qu'une intersection soit trouvée. Ce modèle postule que la mémoire à long terme est structurée en plans. Chaque plan est dédié à un concept appelé *type node*, c'est le *patriarche* de ce plan. Sur ce dernier, vont également être placés les *token nodes* qui vont servir à définir le *type node*. Ainsi, la Figure 1 montre 3 plans dédiés au mot anglais *plant* (un pour chaque sens du mot *plant*). Chacun de ces plans a pour *type node* le mot *plant*, et est composé de *token nodes* permettant de définir un des sens de ce mot. Le plan 1 est donc composé d'un *type node* faisant référence au mot *plant* comme un végétal, et de *token nodes* qui définissent ce concept. Il est également relié aux plans 2 et 3 faisant référence aux autres sens du mot *plant*. En cas de présentation du mot *plant*, l'activation se propage le long des *token nodes*.

Afin de réduire la taille du modèle, Quillian remplace tous les mots de liaison tels que *et*, *ou*, etc... par des liens spécifiques qui vont prendre en compte la relation entre deux *type node* ou deux *token node*. De même les pronoms sont représentés comme des références à un nom précédent. Quillian définit 5 types de liens (a) superordonné (b) modifiant (c) disjonctif (d) conjonctif (e) les liens résiduels dans lequel un lien est souvent lui-même un concept.

Toutefois, le modèle présente certaines limites : par exemple il va trouver des points communs entre des concepts très éloignés comme *homme* et *planter*. En effet, ces deux concepts vont se rejoindre au concept *personne* parce que *homme* est défini comme *est une personne* et *planter* est défini comme *quand une personne met quelque chose dans la terre*.

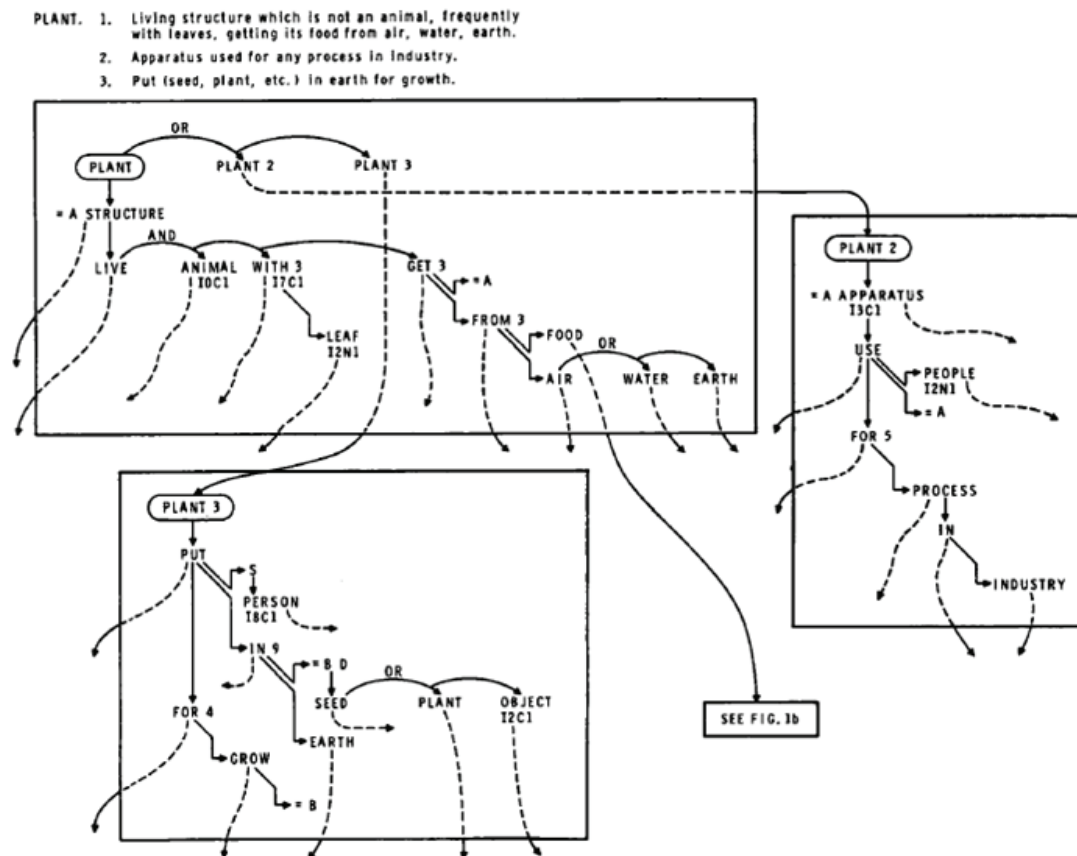


FIG. 1a. Three Planes Representing Three Meanings of "Plant."

Figure 1

Tiré de Quillian, 1967. Trois plans représentant les trois significations de *plant*. Les traits en pointillés indiquent le lien vers un nouveau plan dédié à un *type node* qui aura été le *token node* du plan précédent. Ainsi, *animal* est un *token node* du plan *PLANT* mais sera le *type node* du plan *ANIMAL*.

1.1.2. Le modèle de Collins et Loftus, 1975

Le modèle de Collins et Loftus (1975) reprend globalement l'idée de propagation de l'activation du modèle de Quillian (1967).

Les concepts sont également reliés entre eux par des liens. Plus les concepts partagent de caractéristiques plus ils sont proches au sein du réseau et plus l'activation de l'un engendre l'activation de l'autre. Il existe plusieurs types de liens qui spécifient les relations entre les concepts : les connections superordonnées, les propriétés, et les subordonnées. Ainsi un canard est un oiseau (superordonné) qui a les pattes palmées (propriété), et les colverts sont des canards (subordonnés).

Le modèle de Collins et Loftus précise 4 points principaux sur la propagation de l'activation :

- (a) Lorsqu'un concept est traité, l'activation se propage le long du réseau et diminue simultanément. C'est-à-dire que la diminution de l'activation est proportionnelle à la distance entre deux concepts (e.g., *nuage* sera moins activé que *soleil* suite à la présentation de *rouge* cf. Figure 2).
- (b) Plus un concept est traité continuellement et longtemps, plus la propagation de l'activation dure dans le temps. Pourtant, un seul concept peut être traité à la fois. L'activation se propage ainsi à partir d'un seul concept à la fois mais elle pourra ensuite se propager en parallèle à travers le réseau.
- (c) L'activation diminue avec le temps ou du fait d'une autre activité.

Les assumptions (b) et (c) supposent qu'il existe une quantité limitée d'activation parce que plus il y a de concepts amorcés moins ils vont être activés.

- (d) Puisque l'activation est une quantité variable, il est nécessaire d'avoir une notion de seuil à partir duquel un concept peut être activé.

En parallèle de ce réseau sémantique, il existerait un réseau lexical dans lequel les noms des concepts seraient organisés en fonction des similarités phonémiques et éventuellement orthographiques. L'existence d'un réseau lexical permettrait au système de contrôler si l'amorçage se fait de façon phonémique ou sémantique. L'accès lexical permettrait d'accéder au réseau sémantique.

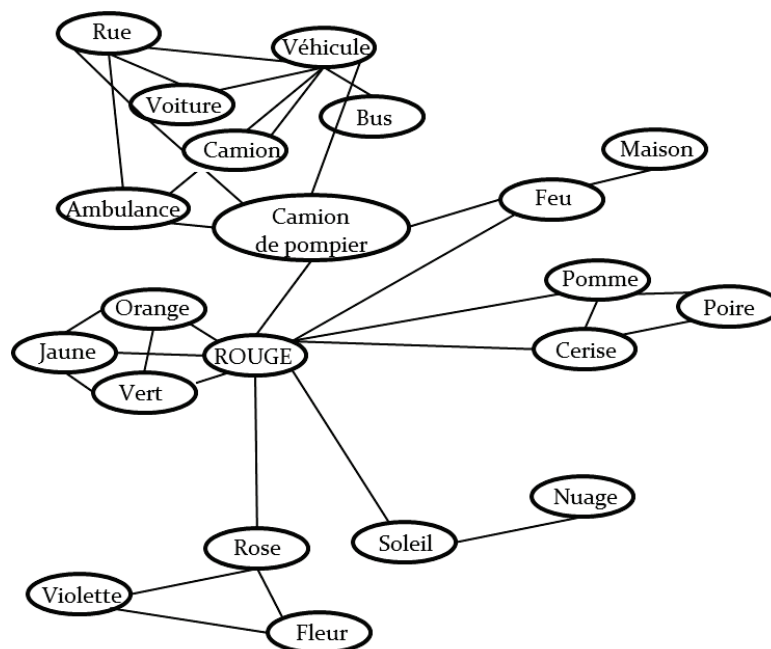


Figure 2

Exemple de réseau sémantique centré autour du concept ROUGE. La présentation de ROUGE engendrera une propagation de l'activation dans l'ensemble du réseau en suivant les assumptions présentées plus haut.

1.1.3. Théorie de Posner et Snyder (1975)

Posner et Snyder (1975) postulent qu'il existe deux mécanismes de récupération en mémoire à long terme : un mécanisme automatique et un mécanisme contrôlé.

Le mécanisme automatique qui est rapide, opère indépendamment de la conscience et de la volonté du sujet et n'inhibe pas les concepts non liés à ceux que l'on est en train de traiter et sur lesquels l'attention n'est pas portée (cf. Figure 3.A.).

Le mécanisme stratégique ou contrôlé qui est lent, ne peut pas opérer indépendamment de la conscience du sujet et inhibe les concepts non liés à celui que l'on traite et sur lesquels l'attention n'est pas portée (cf. Figure 3.B. et C). De plus, ce mécanisme dépend des ressources d'un processeur central, ses capacités sont limitées.

Cette théorie a principalement été testée et opérationnalisée par Neely (1977). En utilisant un paradigme d'amorçage sémantique (pour une définition de l'amorçage sémantique, cf. p40) Neely (1977) manipule les attentes des participants afin de tester l'existence des deux systèmes, automatique et stratégique.

Les participants doivent effectuer une tâche de décision lexicale (i.e., déterminer si l'item est un mot ou un pseudo-mot) sur un item cible présenté (entre 250 ms et 2000 ms) après une amorce. La catégorie sémantique de l'amorce permet aux participants d'émettre des prédictions sur la catégorie sémantique de la cible. L'amorce peut appartenir à l'une de ces 3 catégories sémantiques : un oiseau, une partie du corps, une partie de construction. Dans une dernière condition, neutre, l'amorce est une série de X. Les participants sont informés avant l'expérience que lorsque l'amorce est un oiseau, la cible, quand elle est un mot sera un oiseau dans 66% des cas. Lorsque l'amorce est une partie du corps, la cible quand elle est un mot sera une partie de construction dans 66% des cas et inversement. Enfin lorsque l'amorce est une série de X, la cible, lorsqu'elle est un mot pourra être un oiseau, une partie du corps ou une partie de construction.

Ce paradigme permet ainsi de tester d'une part l'amorçage automatique, résultat de la propagation de l'activation à travers le réseau. D'autre part, il teste l'effet d'amorçage stratégique et l'attente des participants. Neely (1977) prédit l'obtention d'une dissociation, entre ces deux types de traitement. Selon le modèle, l'amorçage sera uniquement automatique lorsque le SOA (*Stimulus Onset Asynchrony*) sera court (puisque les effets stratégiques sont lents) et uniquement stratégique lorsque le SOA sera long (puisque la propagation automatique diminue avec le temps).

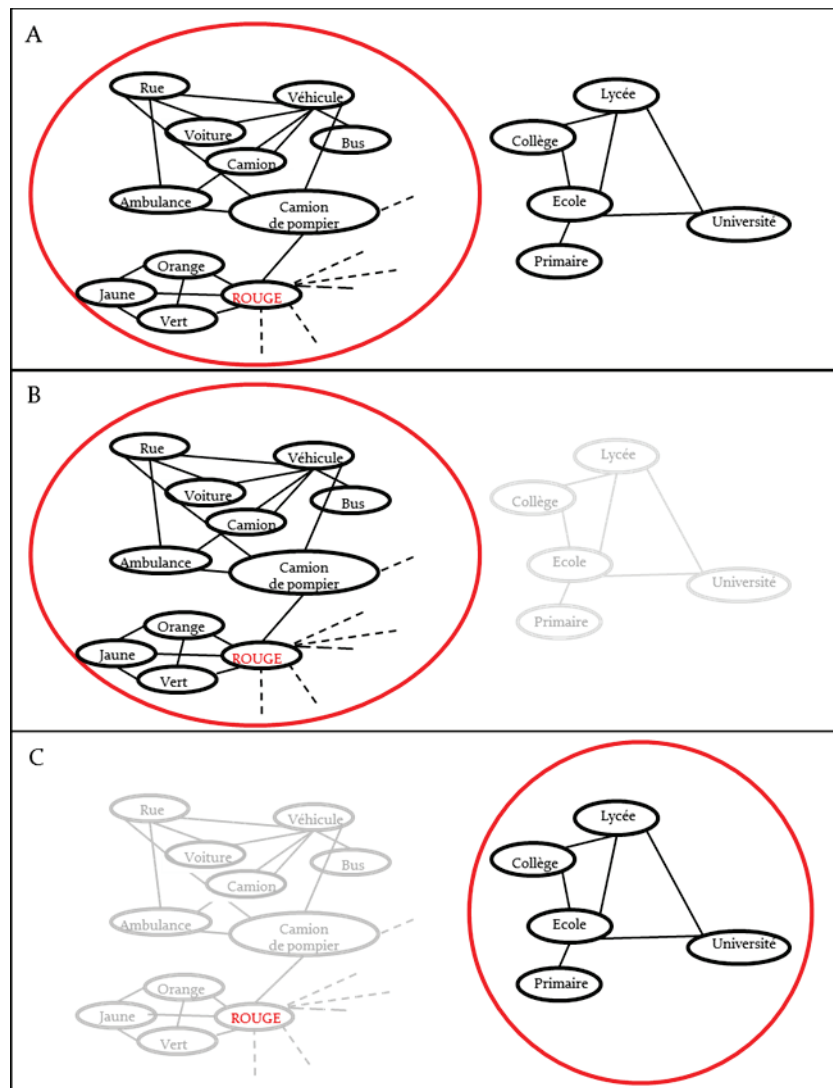


Figure 3

Représentation de deux réseaux sémantiques distincts.

A. Mécanisme Automatique. La présentation du mot ROUGE engendre l'activation automatique du réseau associé. Le réseau associé au cadre scolaire est au repos. B. Mécanisme Stratégique. La présentation du mot ROUGE engendre l'activation stratégique du réseau associé. Le réseau associé au cadre scolaire est inhibé. C. Mécanisme Stratégique. La présentation du mot ROUGE engendre l'activation stratégique du réseau associé au cadre scolaire. Le réseau associé au mot rouge est inhibé.

En s'appuyant sur la théorie de Posner & Snyder (1975), Neely (1977) émet les hypothèses suivantes :

Tout d'abord, il postule l'apparition d'un effet d'amorçage significatif lorsque l'amorce et la cible sont sémantiquement liées (oiseau – PIGEON), que le SOA soit court ou long. En effet, puisque l'amorce et la cible appartiennent à la même catégorie, des effets d'amorçage automatiques seront observés lorsque le SOA sera court. Lorsque le

SOA sera long, les effets d'amorçage stratégiques seront observables puisque les participants pourront prédire l'apparition de PIGEON dans 66% des cas.

Lorsque l'amorce et le mot cible ne sont pas liés mais que le mot cible fait partie de la catégorie attendue (i.e., partie de construction lorsque l'amorce est une partie du corps et inversement), un effet d'amorçage sera alors observé uniquement lorsque le SOA est long. En effet, puisque dans ces conditions, la propagation automatique de l'activation ne permettra pas l'activation de la cible, seuls les processus stratégiques seront à l'œuvre et ne se mettront donc pas en place suffisamment rapidement pour être observables lorsque le SOA est court.

Lorsque l'amorce et le mot cible sont liés mais que le mot cible ne fait pas partie de la catégorie attendue (i.e., partie du corps lorsque l'amorce est une partie du corps) alors un effet d'amorçage automatique sera observé lorsque le SOA est court. En revanche, lorsque le SOA est long alors un effet d'inhibition sera observé, puisque les mécanismes stratégiques génèrent une inhibition des réseaux sur lequel l'attention n'est pas portée (cf. Figure 3.C.).

Les résultats obtenus par Neely (1977) vérifient totalement les hypothèses posées, validant ainsi la théorie initialement proposée par Posner & Snyder, 1975. Cette étude a permis non seulement de tester la validité de la théorie de Posner & Snyder (1975), mais les résultats ont également montré que les processus stratégiques peuvent utiliser la présentation d'une amorce pour diriger l'attention sur un autre groupe de mots. Ces résultats suggèrent donc que les processus stratégiques peuvent avoir un autre rôle que l'équivalence simple de la propagation automatique de l'activation.

1.2. Le modèle stratégique de Becker (1980)

Contrairement aux modèles précédents, Becker propose un modèle basé uniquement sur des processus stratégiques. D'après ce modèle, la présentation d'un mot génère l'extraction de ses caractéristiques visuelles primitives. Celles-ci sont ensuite comparées à un set de mots, visuellement proches du stimulus, et classés en fonction de leur fréquence. C'est-à-dire que les mots visuellement proches du mot présenté et ayant une forte fréquence d'occurrence sont les premiers dans le set. Ainsi, la présentation de *CHIEN* engendre un set comprenant des mots comme *BIEN*, *RIEN*, *CHIEN*, *CHIOT*... Chacun de ces mots est ensuite comparé au percept jusqu'à ce que le mot présenté soit trouvé dans le lexique mental, c'est le processus de vérification. Les informations visuelles stockées en mémoire sur ce mot permettent de compléter les primitives qui en avaient été extraites et complètent ainsi la représentation du mot perçu. Cette représentation est ensuite comparée au percept en mémoire sensorielle et le mot est définitivement identifié.

Lorsqu'un contexte est présent, le mécanisme de reconnaissance des mots est modifié. En effet, le modèle considère que lorsque le participant perçoit la cible, il génère consciemment un set de mots liés sémantiquement à l'amorce. Ces mots ne sont pas classés en fonction de leur fréquence mais ordonnés de façon aléatoire au sein du set. Ils sont ensuite examinés un par un pour être comparés au percept, comme lors de la reconnaissance de mots classiques. Lorsque le participant repère que l'amorce et la cible sont généralement très proches sémantiquement, les sets sont plus petits. La recherche dans le set est plus rapide, et de ce fait, l'effet d'amorçage est plus important. En revanche lorsque l'amorce est sémantiquement moins proche de la cible, le set est plus grand. La recherche est donc plus longue et l'effet d'amorçage moins important. Lorsque la cible n'est pas liée à l'amorce, i.e., la cible ne fait pas partie du set présélectionné d'après les caractéristiques sémantiques, le système reprend la recherche en se basant sur les caractéristiques visuelles du mot, comme lorsqu'il est présenté sans contexte.

Ce modèle permet d'expliquer que l'effet d'amorçage résulte d'une facilitation lorsque dans la liste 1 la cible est systématiquement un antonyme de l'amorce (e.g., petit – GRAND) et d'une forte inhibition lorsque dans la liste 2 la cible est un associé ou un exemplaire de la cible (e.g., canard – OIE). Le set de recherche de cible dans la liste 1 est alors extrêmement petit, c'est le mécanisme de prédiction spécifique : l'effet d'amorçage est grand puisque la cible est facilement trouvée. Si la cible n'est pas trouvée dans le set, alors la recherche va reprendre sur la base des caractéristiques visuelles. Plus le set est petit, moins la perte de temps engendrée par la recherche au sein du set sera grande. L'effet d'inhibition sera donc faible.

En revanche, dans la liste 2 où la cible est moins liée à l'amorce (exemplaire de catégorie ou associé) alors le set de mot sélectionné sur la base des indices sémantiques est grand, c'est le mécanisme d'attente générale. Lorsque la cible et l'amorce sont effectivement liées alors la cible n'est généralement pas trouvée plus rapidement que si le système l'avait recherchée sur la base des caractéristiques visuelles. En revanche, lorsque l'amorce et la cible ne sont pas liées sémantiquement, la recherche dans le set est longue et une fois qu'elle a échoué, le système doit reprendre la recherche en utilisant les primitives sensorielles. Cette perte de temps est à l'origine de l'effet d'inhibition observé.

Ce modèle permettrait également d'expliquer une partie des résultats de Neely (1977). En effet, les participants ont appris à attendre un mot appartenant à la catégorie sémantique de la construction après présentation d'un mot faisant référence à une partie du corps. Toutefois, ce modèle n'explique pas l'effet d'amorçage observé par Neely lorsque le SOA est faible et que le mot cible n'appartient pas à la catégorie sémantique attendue mais est lié sémantiquement à l'amorce (e.g., poignet – MAIN). Si l'amorçage sémantique ne reposait que sur des processus stratégiques, aucun effet

d'amorçage n'apparaîtrait dans cette condition puisque le mot cible attendu n'était pas sémantiquement lié à l'amorce.

La particularité des modèles présentés ci-dessus est d'expliquer l'effet du contexte par une modification de l'accès lexical. Les modèles que nous allons à présent examiner ont pour particularité d'expliquer l'effet du contexte par des processus post-lexicaux.

1.3. Les théories post-lexicales

1.3.1. Le *context checking model* de Norris (1986)

Dans le modèle de vérification par le contexte de Norris (1986), l'impact du contexte sur le traitement du mot s'opère après que celui-ci ait été sélectionné ou au moins présélectionné. Dans un ordre d'idée similaire au modèle de Becker (1980), Norris postule que la perception d'un mot engendre la présélection d'une série de mots basée sur leur ressemblance visuelle avec le mot perçu. Le *poids* de ces mots présélectionnés est influencé par leur fréquence d'occurrence, c'est-à-dire qu'un mot plus fréquent a un *poids* plus important qu'un mot moins fréquent. Ensuite, une fois que le set de mots présélectionnés est suffisamment petit, chacun des mots qui le composent est activé, et le système vérifie la plausibilité de son apparition étant donné le contexte. Le seuil de reconnaissance pour un mot plausible est diminué, alors que le seuil de reconnaissance pour un mot peu ou pas plausible dans le contexte est augmenté.

1.3.2. *Compound cue model* de Ratcliff et McKoon

Un autre modèle à postuler un effet post-lexical du contexte est celui de Ratcliff et McKoon (McKoon & Ratcliff, 1992; Ratcliff & McKoon, 1988, 1994). Ce modèle a été créé pour s'implémenter dans des modèles de mémoire déjà existants, sans proposer d'architecture de la mémoire sémantique. Selon ce modèle, les mots présentés sont stockés en mémoire à court-terme. Ainsi, en situation d'amorçage, l'amorce et la cible vont toutes deux être présentes en mémoire à court-terme pour former un indice composé (i.e., *compound cue*). Cet indice composé est supposé avoir une certaine familiarité. Celle-ci est déterminée par la force des liens entre les représentations de ces items en mémoire à long-terme. Par exemple, les mots *docteur* et *infirmière* sont plus fortement liés que *docteur* et *pain*. La familiarité de *docteur* – *infirmière* est ainsi plus importante que la familiarité de *docteur* – *pain*. Cette plus grande familiarité va permettre de reconnaître la cible en tant que mot plus rapidement. Toutefois, une limite importante de ce modèle est qu'il ne permet pas d'expliquer les effets d'amorçage observés avec une tâche de naming (McNamara, 2005; Neely, 1991).

1.3.3. Le modèle hybride de Neely et Keefe (1989)

Enfin, Neely et Keefe, (1989) proposent un modèle impliquant 3 processus distincts et indépendants mis en place lors du traitement sémantique. Premièrement, la propagation automatique de l'activation (Collins & Loftus, 1975; Posner & Snyder, 1975) agit de manière automatique, rapide et inconsciente. Elle génère des effets de facilitation lors de la présentation de deux items sémantiquement liés mais n'explique pas les effets d'inhibition lors de la présentation d'items non liés sémantiquement (cf. Figure 3). Le second processus est un processus d'attentes stratégiques (Becker, 1980). Il est plus lent et agit donc principalement lorsque le SOA est long, il génère des effets facilitateurs lors de la présentation de deux items sémantiquement liés ainsi que des effets inhibiteurs lorsque deux items ne sont pas sémantiquement liés (cf Figure 3). Enfin, par rapport au modèle de Posner et Snyder, ce modèle hybride postule qu'il existe un phénomène de *semantic matching* post-lexical. Ce dernier processus explique les résultats obtenus lors de tâches de décision lexicale : s'il existe un lien entre l'amorce et la cible, alors la réponse est forcément *mot*, en revanche, s'il n'existe pas de liens entre l'amorce et la cible il y a de fortes chances pour que la réponse soit *pseudo-mot*. De ce fait, Neely et Keefe postulent qu'il se produit une vérification de l'existence de liens sémantiques entre l'amorce et la cible après l'accès lexical. De ce fait, ce modèle est hybride à la fois en terme de nature des processus (i.e., automatique vs stratégique) et de leur position dans la chaîne de traitement du langage (i.e., lexicale vs post-lexicale).

2. L'amorçage sémantique

L'amorçage sémantique se caractérise par l'amélioration des performances à une tâche (en terme de rapidité ou de justesse) effectuée sur un mot cible du fait de la présentation au préalable d'un autre mot sémantiquement lié à la cible, appelé amorce. Ce phénomène a été pour la première fois mis en évidence en modalité visuelle par Meyer et Schvaneveldt (1971). Les participants devaient déterminer si deux suites de lettres présentées simultanément étaient des mots ou des pseudo-mots (i.e., double tâche de décision lexicale). Les participants étaient plus rapides à classifier la paire de mots comme mots lorsque ceux-ci étaient sémantiquement liés que lorsqu'ils ne l'étaient pas (e.g., *infirmier-docteur* vs *pain-docteur*). De plus, plus le lien entre l'amorce et la cible était fort, plus l'effet d'amorçage était important. C'est à dire que lorsque les deux mots étaient identiques (e.g., *docteur-docteur*; amorçage de répétition) les participants étaient plus rapides qu'en condition d'amorçage sémantique (e.g., *infirmier-docteur*).

Depuis ce premier article, la littérature sur le sujet s'est énormément développée (pour une revue voir McNamara, 2005). L'amorçage sémantique a essentiellement été utilisé

pour comprendre et modéliser le fonctionnement de la mémoire sémantique et du traitement sémantique. Bien que la plupart des études s'effectuent en modalité visuelle, l'amorçage sémantique est également présent en modalité auditive (Donnenwerth-Nolan, Tanenhaus, & Seidenberg, 1981; Radeau, 1983; Radeau, Besson, Fonteneau, & Castro, 1998). Des études ont également mis en évidence des effets d'amorçage cross modal, confirmant l'amodalité de la mémoire sémantique (Kouider & Dehaene, 2009; Norris, Cutler, McQueen, & Butterfield, 2006).

Il existe deux principaux débats dans la littérature se rapportant à l'amorçage sémantique et donc au traitement sémantique. Tout d'abord, la question de son automaticité, qui reflète le débat se jouant sur les modèles (voir début du chapitre ; Neely 1991 pour une revue). Nous aborderons tout d'abord ce premier débat, ainsi que la façon dont il a évolué vers la question de l'existence de l'amorçage sémantique subliminal.

Le second débat porte davantage sur l'organisation de la mémoire sémantique, à savoir si les mots sont reliés entre eux du fait de la fréquence de leur association, ou via leurs caractéristiques sémantiques.

2.1. L'automaticité du traitement sémantique

2.1.1. Les méthodes classiques

Depuis 40 ans, la question de l'automaticité du traitement sémantique divise la littérature. Comme nous l'avons vu lors de la présentation des modèles de traitement sémantique certains auteurs considèrent que celui-ci se fait de façon automatique (Collins & Loftus, 1975 ; Posner & Snyder, 1975), alors que d'autres postulent que celui-ci est stratégique et dépend de la conscience du sujet (Becker, 1980). Plusieurs méthodes ont été utilisées pour tester l'automaticité du traitement sémantique.

2.1.1.1. Les variations de SOA

Tout d'abord, séparer l'amorce et la cible par un SOA très court, en effet comme nous l'avons vu au début de ce chapitre, les processus automatiques sous-tendant les effets d'amorçage sont supposés être très rapides, alors que les processus stratégiques sont plus longs à mettre en place. De ce fait, un effet d'amorçage lorsque l'amorce et la cible sont présentées très rapidement l'une après l'autre permettrait de mettre en évidence l'automaticité du traitement sémantique (Burke, White, & Diaz, 1987; Neely, 1977, 1991). Ainsi, Neely (1977) a rapporté un effet d'amorçage lorsque le SOA était de 250 ms alors même que les participants s'attendaient à voir une cible appartenant à une autre catégorie sémantique (cf. p35).

2.1.1.2. La *Relatedness Proportion*

Une autre façon de s'intéresser à l'automaticité du traitement sémantique est de manipuler la proportion de cibles sémantiquement liées à leur amorce (*Relatedness Proportion* ; *RP*). En effet, plus le nombre de cibles liées à l'amorce est important, plus les participants ont tendance à développer des stratégies. Cette manipulation repose sur le postulat suivant : si l'amorçage sémantique provient de processus stratégiques, alors augmenter la proportion de cible et d'amorces liées augmente la génération d'attentes de la part des participants (de Wit & Kinoshita, 2014, 2015b; pour une revue voir Hutchison, 2007). Ainsi, Neely, Keefe et Ross (1989) se sont intéressés à l'effet d'amorçage alors que la *RP* variait en fonction des groupes et pouvait être de .89, .67 ou .33. Comme attendu par les auteurs, l'effet d'amorçage était plus important lorsque la *RP* était plus élevée. Ces résultats suggèrent que des mécanismes aussi bien automatiques que stratégiques sont mis en œuvre lors du traitement sémantique. De plus, les auteurs ont remarqué que lors d'une tâche de décision lexicale, la probabilité de répondre *pseudo-mot* lorsque l'amorce et la cible ne sont pas liées varie avec la *RP*. Ainsi lorsque l'amorce et la cible n'étaient pas liées la probabilité de répondre *pseudo-mot* est plus importante lorsque la *RP* est haute. Ils ont également mis en évidence que le ratio de *pseudo-mot* jouait un rôle sur les temps de réponse, même lorsque la *RP* est maintenue constante.

La méthode la plus récente et maintenant la plus utilisée pour s'intéresser à l'automaticité de l'amorçage sémantique est de tester s'il peut apparaître lorsque l'amorce n'est pas perçue consciemment par le sujet.

2.1.2. L'amorçage sémantique subliminal

2.1.2.1. Masquage de l'amorce

Les premières études à s'intéresser à l'amorçage sémantique subliminal ont été effectuées dans les années 80 (Forster & Davis, 1984; Marcel, 1983) bien que l'intérêt pour la perception et le traitement subliminal soit bien plus ancien (e.g., Sidis, 1898). Le paradigme d'amorçage masqué le plus utilisé dans la littérature a été développé par Forster et Davis (1984) bien que d'autres méthodes existent (Greenwald, Klinger, & Liu, 1989; Luck, Vogel, & Shapiro, 1996). L'expérience de Forster et Davis (1984) s'intéressait en fait à l'amorçage de répétition, mais leur technique de masquage a ensuite été reprise presque systématiquement par les études s'intéressant à l'amorçage sémantique masqué (Abrams & Greenwald, 2000; Damian, 2001; Dehaene et al., 1998; Van den Bussche, Hughes, Humbeeck, & Reynvoet, 2010; Van den Bussche et al., 2009). L'objectif est de présenter une amorce au participant juste en dessous de son seuil de perception (i.e., subliminale). Si le participant ne perçoit pas consciemment l'amorce, il ne pourra alors pas mettre en place de stratégies pour améliorer ses performances. Le

principe de cette méthode consiste à présenter l'amorce pour un temps très bref, (quelques dizaines de ms) en la faisant précéder et suivre de masques (généralement une suite de caractères) afin d'empêcher tout phénomène de persistance rétinienne (cf. Figure 4). Les expériences sont habituellement composées de deux phases. L'objectif est de déterminer, dans une première phase, le temps de présentation de l'amorce suffisant pour que le participant soit capable de la traiter sans la percevoir. Pour cela, il est demandé aux participants de déterminer si l'amorce est présente ou non. Une fois le temps de présentation déterminé, la phase test peut avoir lieu.

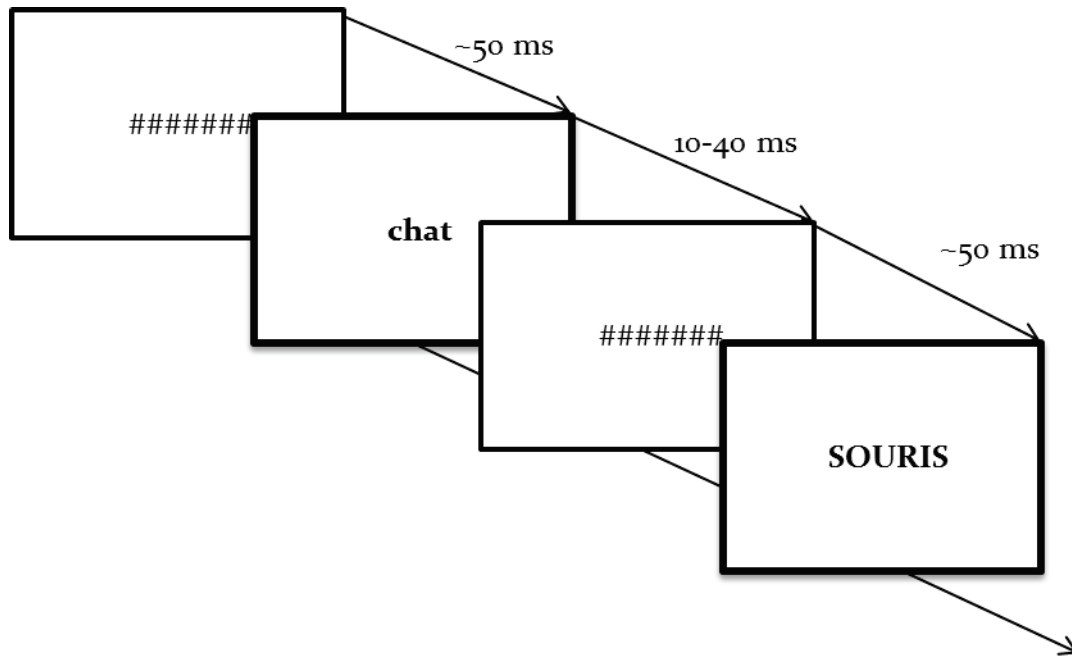


Figure 4

Représentation graphique du paradigme d'amorçage masqué (ici sémantique), comme développé par Forster et Davis (1984).

Ce paradigme a permis de mettre en évidence l'automaticité de traitements bas niveau comme les traitements orthographique ou lexical (Ferrand & Grainger, 1996; Grainger, Diependaele, Spinelli, Ferrand, & Farioli, 2003; Kouider & Dehaene, 2007). Pourtant les résultats concernant l'amorçage sémantique restent sujets à débat, ou tout au moins très limités et soulèvent de nombreuses critiques (Abrams & Greenwald, 2000; Abrams & Grinspan, 2007; Carr & Dagenbach, 1990; Damian, 2001; de Wit & Kinoshita, 2015a; Holender, 1986; Klinger, Burton, & Pitts, 2000; Kouider & Dupoux, 2004) même si certaines études ont obtenus des résultats intéressants (Dehaene & Naccache, 2001; Dehaene et al., 1998; Dell'Acqua & Grainger, 1999; Kunde, Kiesel, & Hoffmann, 2003).

2.1.2.2. De l'existence d'un traitement sémantique subliminal

Ainsi, Dehaene et al. (1998) ont observé un effet d'amorçage sémantique subliminal sur des chiffres. A chaque essai, un chiffre (présenté en lettres ou en caractère numérique ; e.g., quatre ou 4) était présenté de façon subliminale c'est-à-dire durant 43 ms et précédé et suivi d'un masque durant 71 ms. La cible était également un chiffre présenté en lettres ou en caractère numérique pendant 200 ms. Le participant devait alors indiquer si elle était inférieure ou supérieure à 5. Les résultats ont mis en évidence un effet d'amorçage significatif, c'est-à-dire que les participants étaient plus rapides à répondre à la cible lorsque celle-ci était précédée par une amorce qui était également inférieure ou supérieure à 5. Les auteurs postulent que cet effet provient d'un traitement sémantique de l'amorce : d'après eux, les participants appliqueraient inconsciemment la même tâche à l'amorce qu'à la cible. De ce fait, la présentation par exemple de l'amorce *SIX* engendrerait automatiquement la qualification *supérieur à 5* et donc la préparation motrice correspondant à la réponse. En conséquence lorsque la cible est présentée, la réponse est facilitée si elle est également supérieure à 5 puisqu'il y a déjà eu pré-activation motrice. Cette interprétation est confirmée par des données EEG et IRMf enregistrées simultanément, qui ont identifié une réponse motrice à l'amorce respectivement via une LRP (*Lateralized Readiness Potential*) et des activations au niveau du cortex moteur.

Toutefois, en réponse à ces conclusions, Damian (2001) argumente que le traitement de l'amorce n'est pas sémantique. En effet, puisque Dehaene et al. (1998) utilisent le même set de chiffres pour les amorces et les cibles, une fois que les participants ont effectué la tâche sur la cible, ils apprennent simplement la correspondance entre le percept (i.e., le chiffre ou la suite de lettre) et la réponse motrice. Damian (2001) appelle cet effet *stimulus-response mapping*. Dans son expérience, les sujets devaient effectuer une tâche de catégorisation, à savoir comparer la taille de l'objet désigné par le mot cible à un standard (e.g., cible : *ARAIGNEE* ; standard : *maison*). Encore une fois, une amorce masquée (précédée et suivie d'un masque pendant respectivement 54 ms et 29 ms) était présentée 43 ms. L'amorce et la cible pouvaient appartenir à la même catégorie (plus ou moins grandes que le standard) ou non, mais n'étaient jamais liées sémantiquement (e.g., l'amorce *voiture* et la cible *ARAIGNEE* sont plus petites que le standard *maison* mais ne sont pas liées sémantiquement). Les résultats ont montré un effet de congruence lorsque la cible et l'amorce devaient être catégorisées de la même manière (i.e., toutes deux plus petites ou plus grandes que le standard) sans que celui-ci ne reflète un amorçage sémantique. Damian a en outre mis en évidence dans une deuxième expérience une disparition de l'effet d'amorçage, lorsque les mots utilisés comme amorce n'étaient pas présentés en tant que cible au préalable. Ainsi, il suggère que dans la première expérience, l'amorce était simplement reconnue à un niveau orthographique et appelait la même réponse que lorsqu'elle avait été présentée

en tant que cible. Aucun traitement sémantique n'était donc effectué sur l'amorce et l'effet observé n'était que le reflet d'une association stimulus-réponse.

En réponse à cette critique toutefois, Naccache et Dehaene (2001) ont répliqué les résultats de Dehaene et al., (1998) alors que les amorces et cibles appartenaient à deux sets de chiffres différents.

2.1.2.3. Amorçage sémantique ou effet de congruence?

La nature des mécanismes sous-tendant ces effets reste toutefois sujette à débat, à savoir est ce que ces effets de congruence reflètent un amorçage sémantique ?

Dans cette optique, Reynvoet, Gevers et Caessens (2005) ont utilisé un paradigme d'amorçage masqué dont les stimuli pouvaient être des lettres ou des chiffres. Lorsqu'elle était un chiffre, il était demandé aux participants de déterminer si la cible était inférieure ou supérieure à 5. Lorsqu'elle était une lettre, ils devaient indiquer si elle était positionnée avant ou après O dans l'alphabet. Comme attendu, l'effet d'amorçage intra-catégorie est apparu significatif. La présentation d'une amorce inférieure à 5 engendrait une réponse plus rapide à une cible inférieure à 5 et inversement. De même lorsque l'amorce et la cible étaient des lettres présentes avant ou après O dans l'ordre alphabétique. Enfin, un effet d'amorçage était également observé en inter-catégorie. Ainsi, la présentation d'une lettre en amorce pouvait générer un effet d'amorçage sur une cible chiffre. Pour que cet effet inter-catégorie apparaisse il fallait que la réponse motrice évoquée par la présentation de l'amorce soit la même que celle évoquée par la présentation de la cible. Si la présentation d'un chiffre inférieur à 5 appelait une réponse de la main gauche et que la cible présentée était une lettre positionnée avant O, appelant également une réponse de la main gauche, alors un effet facilitateur inter-catégorie apparaissait. Ces résultats sont ainsi congruents avec ceux de Dehaene et al., (1998), mettant en évidence une préparation motrice en réponse à la présentation de l'amorce.

Dans une seconde expérience, Reynvoet et al. (2005) mettent d'ailleurs en évidence la nature motrice de cet effet de facilitation en demandant aux participants de répondre avec les mains aux cibles chiffres et avec les pieds aux cibles lettres. L'effet d'amorçage intra-catégorie est toujours observé, mais l'effet inter-catégorie disparaît, confirmant la nature motrice de l'effet d'amorçage de l'expérience précédente. Les auteurs interprètent leurs résultats comme une preuve que les stimuli sont traités sémantiquement mais que l'effet d'amorçage a une origine motrice. De plus, ils mettent en évidence que le traitement subliminal peut également se faire sur des lettres et pas uniquement sur des chiffres.

Bien que ces études mettent en évidence un traitement sémantique, celui-ci est largement dirigé par la tâche. En effet, pour qu'un effet d'amorçage apparaisse, l'amorce et la cible doivent être catégorisées de la même façon selon la tâche demandée au sujet (Eckstein & Perrig, 2007; Klinger et al., 2000; Spruyt, De Houwer, Everaert, & Hermans, 2012). Ainsi, Eckstein et Perrig (2007) ont demandé à deux groupes de participants d'effectuer soit une tâche de détermination de valence (i.e., positive vs. négative) soit une tâche de catégorisation animé vs. non-animé. Les mêmes stimuli étaient présentés à chaque groupe. Un effet d'amorçage a été mis en évidence pour chacun des groupes, mais dépendamment de la tâche. Ainsi, pour le groupe 1, *démon* amorçait *accident* (i.e., amorce et cible ayant une valence négative) mais pas *bébé* (i.e., amorce ayant une valence négative et cible une valence positive). Pour le groupe 2 en revanche, *démon* n'amorçait pas *accident* (i.e., amorce animée vs cible non-animée) mais *bébé* (i.e., amorce et cible animées).

Klinger et al. (2000) sont arrivés à des conclusions similaires en utilisant une tâche de décision lexicale. Une facilitation n'était observée que lorsque l'amorce et la cible partageaient la même nature lexicale (i.e., mots ou pseudo-mots). Lorsque l'amorce et la cible étaient des mots, qu'il existe un lien sémantique entre l'amorce et la cible n'influaient pas sur les temps de réponse. Ces résultats rappellent le constat de Neely (1977), qui mettaient en évidence que les attentes (i.e., les processus *top-down*) pouvaient générer des effets d'amorçages. Toutefois, dans l'article de Neely (1977) un effet d'amorçage est tout de même mis en évidence lorsque le SOA est si court qu'il est censé être automatique. Dans ces conditions, si l'effet de facilitation était celui observé par Neely, sous-tendu par la propagation automatique de l'activation, alors un effet d'amorçage plus important apparaîtrait lorsque amorce et cible sont deux mots sémantiquement liés.

L'allocation de l'attention semble également importante dans l'apparition d'un effet d'amorçage sémantique subliminal. Naccache, Blandin et Dehaene (2002) ont manipulé la régularité temporelle avec laquelle l'amorce masquée était présentée. Les résultats suggèrent que lorsque les participants ne peuvent pas prédire le moment d'apparition de l'amorce et ainsi allouer leur attention à un instant *t* à l'amorce, alors ils ne sont pas capables de la traiter sémantiquement.

Ces résultats suggèrent donc que l'effet d'amorçage est dépendant de la tâche et de l'allocation de l'attention du participant. D'ailleurs, l'effet de la tâche utilisée est très important en amorçage sémantique masqué, ainsi, la tâche de catégorisation sémantique génère des effets plus grands que la tâche de décision lexicale (Lucas, 2000), générant elle-même des effets plus grand que la tâche de *naming* (Van den Bussche et al., 2009). Il semble peu probable que le traitement sémantique effectué corresponde aux modèles présentés plus haut, comme la propagation automatique de l'activation. En effet, si tel était le cas, alors un effet d'amorçage apparaîtrait dès lors

que l'amorce et la cible partagent des caractéristiques sémantiques, indépendamment de la tâche. On observe ainsi davantage un effet de congruence sémantique qu'un effet d'amorçage sémantique *per se*. C'est-à-dire qu'un stimulus partageant des caractéristiques sémantiques avec un autre va pouvoir l'amorcer, si l'attention du sujet est explicitement portée sur ces caractéristiques. Cet effet ne correspond donc pas à un effet d'amorçage classique.

2.1.3. Le cas de la modalité auditive

Comme nous l'avons mentionné en introduction du Chapitre I (cf. p25), bien que la parole soit bien plus présente dans la vie d'un individu que l'écriture, la littérature scientifique s'est principalement intéressée à la modalité visuelle, en admettant que ce qui est observable à l'écrit le sera également à l'oral. Or, certaines différences entre l'oral et l'écrit génèrent des contraintes qui impliquent un traitement différent. Particulièrement, à l'oral, un mot n'est pas présenté dans son ensemble à un temps *t* mais il se déroule dans le temps. Toute l'information linguistique n'est donc pas disponible en même temps. Cette spécificité génère un problème majeur dans l'étude de l'amorçage sémantique subliminal en modalité auditive. En effet, alors qu'en modalité visuelle, le traitement conscient d'un mot est stoppé grâce à une présentation très brève et par l'ajout de masques, en modalité auditive, le mot est forcément présenté pendant suffisamment longtemps pour être traité consciemment.

De ce fait, différentes méthodes ont été proposées et testées pour étudier l'existence d'un traitement sémantique subliminal en modalité auditive.

Tout d'abord, l'écoute dichotique a été utilisée. Cette méthode consiste à faire porter un casque aux participants afin de présenter un signal différent dans chaque oreille. Il est ensuite demandé aux participants de prêter attention uniquement à un des deux signaux. Si Eich (1984) pensait avoir mis en évidence un effet d'amorçage sémantique sur la reconnaissance des mots présentés sur le signal à ignorer, la réplication de l'expérience par Wood, Stadler et Cowan (1997) a montré que ces résultats provenaient simplement de la capacité des participants à déplacer leur attention sur le signal à ignorer. Dans une autre expérience utilisant l'écoute dichotique Dupoux, Kouider et Mehler (2003) ont mis en évidence un effet de répétition sur une tâche de décision lexicale lorsque les participants écoutaient un des canaux et qu'un des mots cibles était également présenté dans l'oreille contralatérale. Toutefois, cet effet apparaissait uniquement lorsque le mot en question était suffisamment saillant pour que les participants aient conscience de l'avoir entendu.

Une autre technique pour étudier l'amorçage sémantique subliminal de la parole a été proposée par Daltrozzo, Signoret, Tillmann et Perrin (2011). L'amorce dans cette étude est présentée à une très faible intensité et est suivie d'un item cible, présenté à une

intensité confortable et sur lequel les participants doivent effectuer une tâche de décision lexicale. Après l'expérience, les auteurs se sont assurés que les participants n'étaient pas capable de catégoriser l'amorce comme mot ou pseudo-mot mais qu'ils étaient capables de détecter sa présence. Les participants ne répondant pas à ces critères ont été exclus. Les auteurs trouvent un effet d'amorçage significatif pour les participants les plus longs à répondre. Bien qu'ils déclarent que leurs résultats montrent qu'il est possible d'étudier le traitement sémantique inconscient en modalité auditive, quelques critiques peuvent être émises. Tout d'abord, l'amorçage sémantique n'est significatif que pour les participants les plus lents (~1100 ms), or ces temps de réponses paraissent extrêmement longs pour une cible présentée dans le silence. En effet, pour une tâche de décision lexicale sur un item cible présenté dans le silence, Radeau (1983) recueille des temps de réponse variant en moyenne autour de 500-600 ms. De plus, les longs temps de réponse augmente le risque de l'apparition de processus stratégiques (Neely, 1991) et certains auteurs soulignent que les effets d'amorçage sémantique subliminal sont de fait de très courte durée (Greenwald, Draine, & Abrams, 1996). Etant donné que les participants étaient conscients de la présence de l'amorce, il est également possible qu'un effet d'amorçage rétroactif soit à l'origine de ces résultats. Bien qu'ils soient intéressants, il reste donc des interrogations sur la validité des interprétations qui peuvent être retirées de ces résultats.

Contrairement à la modalité visuelle, aucune méthode ne permet donc à ce jour d'attester de l'existence d'un traitement sémantique subliminal en modalité auditive.

2.2. Amorçage sémantique pur ou amorçage associatif ?

Un second débat de la littérature sur le traitement sémantique s'intéresse à la nature des représentations sémantiques en mémoire à long-terme. C'est-à-dire, est-ce que les concepts se recoupent en termes de caractéristiques ou simplement parce qu'ils sont associés ? Ainsi, il est avéré que la présentation du mot *chien* va amorcer le mot *chat*. Toutefois, la question posée par certains scientifiques est de savoir si *chien* amorce *chat* parce que les deux mots sont associés (i.e., souvent présentés ensemble) ou parce qu'ils partagent des caractéristiques sémantiques communes (i.e., animaux de compagnie, mammifère, 4 pattes, fourrure, etc.). De plus, la définition même d'associatif pose problème puisque généralement, pour déterminer si un mot est associé à un autre, il va être demandé à un ensemble de participants d'associer librement un, deux ou trois mots au mot présenté (Alario & Ferrand, 1998; Ferrand, 2001). Ce type d'association ne prévient pourtant pas l'existence de liens sémantiques en plus des liens associatifs. Ainsi, selon les normes associatives pour le français, le mot le plus fortement associé à *diable* est *démon* (Alario & Ferrand, 1998) or ces deux mots partagent également beaucoup de caractéristiques sémantiques. Afin de contourner ce problème, et tester l'effet des liens associatifs, différents types d'associations ont été privilégiés. Tout d'abord les associations de mots fréquemment associés dans des

phrases mais ne partageant pas de caractéristiques sémantiques comme *diable* et *corps* (i.e., le diable au corps). Une autre solution est d'utiliser des associations créées pendant l'expérience ; ainsi, les expérimentateurs vont tenter de faire créer des liens aux participants entre deux items déjà présents en mémoire sémantique (Dagenbach, Horst, & Carr, 1990), ou alors de créer des liens entre deux pseudo-mots (Hayes & Bissett, 1998).

L'amorçage indirect peut également être utilisé pour tester l'effet des caractéristiques sémantiques. Par exemple, *lion* et *rayure* ne sont pas proches sémantiquement mais ils peuvent être liés via *tigre* et via *zèbre*. D'après le modèle de propagation automatique de l'activation (Collins & Loftus, 1975), la présentation de l'un devrait avoir un effet sur la reconnaissance de l'autre. Plusieurs études ont pu mettre en évidence des effets d'amorçage indirect significatifs (Bennett & McEvoy, 1999; Jones, 2012; McKoon & Ratcliff, 1992)

De nombreuses études ont ainsi visé à mettre en évidence un effet d'amorçage purement associatif et/ou purement sémantique (Frenck-Mestre & Bueno, 1999; pour une revue voir Hutchison, 2003). Dans une méta-analyse sur un total de 26 études, Lucas (2000) trouve des effets d'amorçage sémantique purs significatifs. De façon intéressante, ces effets semblent d'avantage automatiques que stratégiques puisqu'ils sont insensibles au SOA utilisé et à la RP. Il semble pourtant qu'il existe un effet de renforcement des relations associatives, en effet, la taille d'effet passe de $d = 0.25$ pour l'amorçage sémantique pur à $d = 0.50$ lorsque le lien entre l'amorce et la cible est à la fois sémantique et associatif. L'origine de ce renforcement associatif, proviendrait du fait que lorsque deux items partagent à la fois des liens associatifs et sémantiques, c'est qu'ils sont sémantiquement très proches (Lucas, 2000; Thompson-Schill, Kurtz, & Gabrieli, 1998). Une autre hypothèse serait que lorsqu'il existe une forte co-occurrence entre les mots associés, l'effet d'amorçage se fait alors via le lexique mental.

La recherche sur la structure de la mémoire sémantique et sur le traitement sémantique s'est, en premier lieu, faite grâce à des expériences comportementales. Toutefois presque 10 ans après la mise en évidence de l'effet d'amorçage sémantique en comportement, Kutas et Hillyard (1980) découvrait un indice électrophysiologique du traitement sémantique.

3. La N400

En 1980, en cherchant à étudier les variations de la P300 en réponse à un mot inattendu ou même incongruent dans son contexte (e.g., *Le papa tient son bébé dans ses BRAS* vs. *Le papa tient son bébé dans ses CHAUSSETTES*), Kutas et Hillyard (1980) ont mis en évidence l'existence d'un potentiel évoqué, distinct de la P300 et sensible à ces variations. Ce potentiel évoqué reflétant l'intégration d'un mot dans son contexte,

est caractérisé par une négativité relative apparaissant 400 ms après la présentation du stimulus et a donc été nommé N400. L'onde N400 est évoquée aussi bien en modalité auditive que visuelle (Anderson & Holcomb, 1995; Bentin, Kutas, & Hillyard, 1993, 1995). Elle est généralement localisée au niveau des électrodes centro-pariétales et légèrement latéralisée à droite en modalité visuelle (Holcomb & Anderson, 1993; Kutas, Van Petten, & Kluender, 1994) et à gauche en modalité auditive (Carey, Mercure, Pizzioli, & Aydelott, 2014; Hagoort & Brown, 2000; Holcomb & Neville, 1990). Elle n'est pas forcément négative dans l'absolu ; la présentation d'un mot inattendu ou incongruent dans son contexte va engendrer une réponse plus négative que lorsque le mot est attendu et congruent (cf. Figure 5). De plus, elle n'est pas sensible aux caractéristiques physiques du stimulus, ainsi, la présentation du stimulus *Il a rasé sa moustache et sa barbe* n'engendre pas une N400 de plus faible amplitude que la présentation de la phrase *Il a rasé sa moustache et sa BARBE* (Kutas & Federmeier, 2011).

C'est principalement son amplitude qui est mesurée comme indice de l'intégration d'un mot dans son contexte. Toutefois, la N400 n'est pas spécifique aux stimuli linguistiques, elle peut également être générée en réponse à des images (West & Holcomb, 2002) et reflèterait ainsi d'avantage une intégration des stimuli dans le but de faire sens plutôt qu'un traitement linguistique *per se*.

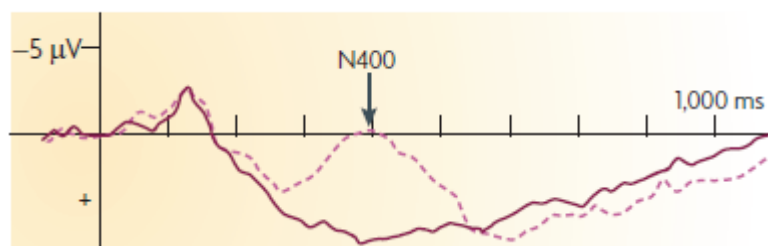


Figure 5

Tiré de Lau et al., 2008. Représentation graphique du potentiel évoqué N400. Les pointillés correspondent à la réponse électrophysiologique à un mot cible non attendu dans un contexte donné. La ligne continue correspond à la réponse électrophysiologique à un mot cible attendu dans un contexte donné.

3.1. Lexicale ou post-lexicale ?

La nature exacte des mécanismes sous-tendant l'effet N400 reste à ce jour sujette à débat à savoir, si la N400 est le reflet de processus lexicaux ou post-lexicaux (Lau, Phillips, & Poeppel, 2008). Si la N400 reflète un processus post-lexical, alors elle reflète l'intégration d'un stimulus (mot ou image) dans son contexte. Il a largement été montré que la N400 est corrélée ($r = .09$; Kutas & Federmeier, 2011) à la *cloze*

probability d'une phrase (Boulenger, Hoen, Jacquier, & Meunier, 2011; Connolly & Phillips, 1994; Federmeier & Kutas, 2001; Gunter, Friederici, & Schriefers, 2000; Oleser & Kotz, 2011). La *cloze probability* correspond à la probabilité qu'une phrase se termine par un mot donné. Ainsi, la N400 en réponse à un mot cible intégré dans une phrase à haute *cloze probability* sera moins importante qu'en réponse à un mot cible intégré dans une phrase à faible *cloze probability* (cf. Figure 5). Enfin, la latence assez longue de la N400 irait également dans le sens d'un processus post-lexical puisqu'elle serait trop tardive pour refléter un processus lexical (Hauk & Pulvermüller, 2004; Sereno, Rayner, & Posner, 1998). En effet, Sereno et al. (1998) ont mis en évidence un effet de la fréquence des mots présentés en modalité visuelle sur les composantes précoces comme la P1 et la N1 suggérant que l'accès lexical aurait plutôt lieu aux alentours de 100 – 150 ms. Cette estimation est congruente avec les données de Sereno et al. (1998), montrant grâce à une procédure d'*eye-tracking* qu'une saccade permettant de fixer un nouveau mot apparaît environ 300 ms après le début de la fixation précédente, suggérant ainsi qu'à cet intervalle de temps, l'accès lexical est terminé. Ces données seraient également compatibles avec la modalité auditive, où les tâches de *gating* (présentation du début de mot en fenêtres de plus en plus longues, jusqu'à ce que le participant soit capable de le reconnaître) ont permis d'évaluer qu'un mot est reconnu en moyenne dans les 200 ms après le début de sa présentation (Marslen-Wilson & Welsh, 1978).

Certains auteurs soutiennent toutefois que la N400 reflète davantage une difficulté de récupération d'un mot en mémoire, et donc plus spécifiquement une difficulté dans l'accès lexical que dans l'intégration. Cette hypothèse est en accord avec les études qui mettent en évidence que la N400 peut être induite par la présentation d'une paire de mots non liés sémantiquement (Barber, Vergara, & Carreiras, 2004; Bentin et al., 1995; Bentin, McCarthy, & Wood, 1985; Brown & Hagoort, 1993; Romei, Wambacq, Besing, Koehnke, & Jerger, 2011). En effet, en utilisant un simple paradigme d'amorçage, Bentin et collègues ont pu observer une modulation de l'amplitude de la N400 en réponse au mot cible suivant que celui-ci était lié à l'amorce qui le précédait ou qu'il ne l'était pas. Cette hypothèse implique ainsi que l'amplitude de la N400 serait modulée par n'importe quel paramètre permettant de faciliter ou de rendre plus difficile l'accès lexical. Plusieurs études ont permis de mettre en évidence que l'amplitude de la N400 était modulée par la fréquence des mots (i.e., un mot plus fréquent engendre une N400 de plus faible amplitude; Barber et al., 2004; Dufour, Brunellière, & Frauenfelder, 2013). La N400 est également sensible à l'effet de répétition ; la présentation répétée d'un mot engendre une diminution de son amplitude (Rugg, 1985). Enfin, Desroches, Newman et Joannis (2009), ont montré que la présentation d'une image d'objet suivie de la présentation d'un mot appartenant à la même cohorte (i.e., même chaîne de phonèmes sur la première syllabe ; p27) engendrait une diminution de l'amplitude de la N400. L'ensemble des résultats de ces travaux suggèrent donc que les variations de

l'amplitude de la N400 pourraient refléter plus spécifiquement l'accès lexical que l'intégration dans le contexte.

La nature lexicale ou post-lexicale de l'onde N400 s'inscrit également dans le débat concernant l'automatisme de l'amorçage sémantique. En effet, les études ayant mis en évidence un effet d'amorçage sémantique masqué sur la N400 vont dans le sens de sa nature lexicale (Deacon, Hewitt, Yang, & Nagata, 2000). De l'autre côté, les études ne l'ayant pas mis en évidence expliquent cette absence de modulation par la nature post-lexicale de la N400 (Brown & Hagoort, 1993; Holcomb, Reder, Misra, & Grainger, 2005). Nous allons nous intéresser à présent plus spécifiquement à cet aspect automatique vs stratégique du traitement sémantique.

3.2. Traitement sémantique automatique ou stratégique ?

Que l'on considère que l'effet N400 soit lexical ou post-lexical, il est évident qu'il est sensible au lien sémantique entre différents stimuli. Puisque la nature automatique ou stratégique du traitement sémantique est également sujette à discussion, la N400 a été utilisée comme indice de traitement sémantique dans la littérature sur ce sujet.

Comme pour les études comportementales, l'investigation de la nature automatique ou stratégique du traitement sémantique en électrophysiologie s'est principalement faite en modalité visuelle et en utilisant des paradigmes d'amorçage sémantique masqué. Toutefois, quelques études se sont intéressées à la modalité auditive (Bentin et al., 1993, 1995). Ainsi, Bentin et al., (1995) ont présenté en écoute dichotique (i.e., un signal différent pour chaque oreille) des listes de mots. Dans chacun des signaux, des paires de mots liés sémantiquement ou non étaient présentées simultanément. Lorsqu'ils étaient sémantiquement liés, ils étaient antonymes dans 50% des paires (e.g., *jour-nuit*) et étaient des exemplaires d'une même catégorie dans les 50% restant (e.g., *hiboux – buse*). Les participants avaient pour consigne d'écouter spécifiquement un des deux signaux. Il leur était demandé de porter particulièrement attention aux mots sous prétexte qu'ils auraient à effectuer une tâche mnésique dessus dans la suite de l'expérience. L'analyse des tracés EEG a permis de mettre en évidence un effet N400 en réponse aux mots présentés dans le signal écouté. En revanche, aucun effet n'était observé en réponse aux mots cibles présentés dans le signal ignoré. Ces résultats suggèrent donc que lorsque les participants ne prêtaient pas attention au signal, celui-ci n'était pas traité sémantiquement.

Dans une étude plus récente, Carey et al. (2014) ont cependant mis en évidence que la présentation d'un signal concurrent dans l'oreille controlatérale a un effet sur le traitement sémantique du signal cible. En effet, dans cette étude, les auteurs présentent aux participants deux phrases en condition dichotique, prononcées par la même voix de femme ou une seule voix en condition monaurale (i.e., le signal est

présenté dans une seule oreille). Dans ces deux conditions, le mot cible était prononcé par un homme et présenté binauralement (i.e., simultanément dans les deux oreilles). Les participants devaient prêter attention à l'une ou l'autre des phrases en fonction des essais, et le mot cible pouvait être congruent avec ce contexte ou pas. L'analyse des tracés EEG en réponse au mot cible a montré que l'amplitude de la N400 était plus faible, à la fois pour les mots cibles congruents et incongruents en condition dichotique, qu'en condition monaurale. Les auteurs interprètent ce résultat comme une preuve que la demande attentionnelle supplémentaire engendrée par la présence d'une voix concurrente perturbe le traitement sémantique de la phrase cible et donc l'intégration du mot cible. Ainsi, cette deuxième étude permet de compléter les apports de la première. En effet, il semblerait que bien que le système cognitif ne soit pas capable de traiter sémantiquement le signal présenté dans l'oreille ignorée, sa présence est perturbatrice pour le traitement sémantique du signal attendu. Ces deux études vont ainsi dans le sens d'un traitement sémantique non automatique, c'est-à-dire qui ne s'effectue pas en dehors de l'attention du sujet et qui est sensible au degré d'attention alloué au signal.

Ces résultats sont toutefois modulés par le fait que la modalité auditive ne soit certainement pas la plus aisée pour mettre en évidence un traitement automatique (cf. p47). En effet, les études s'étant intéressées à l'automaticité de l'effet N400 en modalité visuelle sont beaucoup moins consensuelles, et comme pour l'amorçage sémantique comportemental, sujettes à débat. En utilisant le paradigme d'amorçage sémantique masqué, certains auteurs ont pu mettre en évidence un effet N400 (Deacon et al., 2000; Holcomb et al., 2005; Kiefer, 2002; Kiefer & Brendel, 2006; Kiefer & Martens, 2010; Kiefer & Spitzer, 2000) et d'autres non (Brown & Hagoort, 1993).

Il est intéressant de noter que les questionnements sur la N400 font parfaitement écho aux questions posées par les études sur l'amorçage sémantique en comportement et par les modèles du traitement sémantique. Nous étudierons donc également ces phénomènes grâce à la N400 dans l'Etude 2.

4. Les mots ambigus

Les mots ambigus sont des mots partageant la même forme mais ayant des sens différents. Dans la littérature, le terme *mots ambigus* regroupe à la fois les homonymes (homophones et homographes) mais aussi les polysèmes. Alors que les homonymes (-phones ; -graphes) ont plusieurs significations bien distinctes, les polysèmes quant à eux désignent différents concepts très proches les uns des autres et ayant la même étymologie. Ainsi, *avocat* est considéré comme un homonyme, puisque ses deux sens (i.e., métier vs fruit) ne partagent pas la même étymologie. En revanche, *hôte* (i.e., celui qui reçoit vs celui qui est reçu) est un polysème car les deux sens du mot ont la

même origine. Ces deux types de mots ambigus sont traités et représentés différemment dans le lexique mental (Klepousniotou & Baum, 2007; Klepousniotou, Pike, Steinhauer, & Gracco, 2012). Il existe également des mots ambigus non seulement au niveau de leur sens, mais également en terme de catégorie grammaticale, par exemple les homophones *fête* et *faite* (Lee & Federmeier, 2009, 2012). Dans cette thèse nous nous sommes uniquement intéressés à l'aspect purement sémantique et non à l'aspect grammatical de l'ambiguïté, et de ce fait, nous ne présenterons que des études s'étant intéressées à des mots ambigus homonymes (-phones ; -graphes) partageant la même catégorie grammaticale (e.g., *la mère* ; *le/la maire* ; *la mer*). Ainsi, bien que le terme de *mots ambigus* recouvre de nombreuses réalités, nous l'utiliserons ici pour parler des homonymes (mots ayant une même prononciation –homophones- et/ou une même graphie –homographes- mais des sens différents), dont nous avons étudié le traitement dans la partie expérimentale.

Il est habituellement considéré que l'identification d'un mot dans le lexique mental entraîne l'activation de son sens. Dans le cas des mots ambigus, différents sens peuvent potentiellement être activés. Ces mots auraient une seule représentation dans le lexique mental, mais plusieurs représentations en mémoire sémantique (Klepousniotou, 2002). Lorsque le mot ambigu est présenté, ses différentes représentations entreraient alors en compétition, expliquant le fait que les mots ambigus soient reconnus plus lentement que les autres, et ce d'autant plus que le nombre de leur sens est important (Beretta, Fiorentino, & Poeppel, 2005; Rodd, 2004; Rodd, Gaskell, & Marslen-Wilson, 2002). Il est admis dans la littérature que lorsque le contexte de présentation est neutre, alors les différents sens du mot sont activés en fonction de leur fréquence. Le sens le plus fréquent (i.e., le sens dominant) sera davantage activé que le sens le moins fréquent (i.e., le sens subordonné). En revanche, lorsque le contexte apporte une indication sur le sens le plus pertinent, deux hypothèses peuvent être envisagées ; soit l'ensemble des sens sera activé, soit uniquement le sens pertinent dans le contexte donné sera activé.

4.1. Le *context-sensitive model*

Cette hypothèse postule que le contexte a un effet dominant sur la fréquence du sens des mots : seul le sens pertinent du mot ambigu dans un contexte donné est activé. Par exemple, la présentation de la phrase *Au tribunal, j'ai vu des AVOCATS* ne va engendrer l'activation que du sens juriste du mot *avocat* et non du sens fruit (Glucksberg, Kreuz, & Rho, 1986; Vu, Kellas, & Paul, 1998). Cette hypothèse serait donc en faveur des modèles interactifs de reconnaissance de la parole comme TRACE (McClelland & Elman, 1986). En effet, ce modèle considère qu'il existe des interactions entre les processus *top-down* et *bottom-up* à tous les niveaux de traitement de la parole.

4.2. Le *re-ordered access model*

Cette deuxième hypothèse, au contraire, postule un effet dominant de la fréquence du sens des mots sur le contexte, c'est-à-dire qu'elle considère que les différents sens du mot seront activés quel que soit le contexte (Duffy, Morris, & Rayner, 1988). Celui-ci influencera tout de même le degré d'activation du sens approprié sans pour autant inhiber l'autre : la présentation de la phrase *Au tribunal, j'ai vu des AVOCATS* devrait engendrer une forte activation du sens juriste du mot *avocat* et une faible activation du sens fruit du mot *avocat*. En revanche, la phrase *Au marché, j'ai vu des AVOCATS* devrait engendrer le pattern d'activation inverse (i.e., forte activation du sens fruit et faible activation du sens juriste). Cette seconde hypothèse se positionne davantage dans la lignée du modèle Cohort, qui considère que l'effet du contexte ne sera effectif qu'au moment de la sélection. Particulièrement, Cohort II postule une activation relative des différents candidats de la cohorte en fonction de leur fréquence d'occurrence. Il semblerait donc que le *re-ordered access model* soit compatible avec ce modèle de reconnaissance des mots parlés. En effet, la présentation d'un mot ambigu dans un contexte donné va entraîner l'activation relative de ses différents sens en fonction de leur fréquence et de leur pertinence dans le contexte. Les sens non sélectionnés ne sont pas pour autant inhibés.

Dans une étude en EEG, Swaab, Brown, et Hagoort (2003) ont mis en évidence que les différents sens d'un mot ambigu étaient activés, quel que soit le contexte de présentation (pour chaque mot ambigu, deux sens ont été testés). Des phrases faisant référence au sens dominant ou au sens subordonné d'un mot étaient présentées, suivies par une cible faisant elle-même référence au sens dominant ou au sens subordonné du mot ambigu. La cible et la phrase pouvaient être séparées par un ISI (*Intervalle Inter Stimuli*) court (100 ms) ou long (1250 ms). Les participants devaient écouter attentivement les phrases afin de pouvoir répondre à une éventuelle question de l'expérimentateur. L'étude de l'onde N400 en réponse au mot cible a révélé que lorsque l'ISI était court, l'intégration du sens favorisé par le contexte était facilitée, mais que le sens non favorisé par le contexte était tout de même activé. Ceci était reflété par une N400 de plus faible amplitude lorsque la phrase et la cible faisait référence à deux sens différents du mot ambigu que lorsqu'aucun lien n'était présent entre l'amorce et la cible. Lorsque l'ISI était long, la réponse EEG à la cible montrait une facilitation de l'intégration lorsque le sens de la cible et de la phrase était le même. L'intégration de la cible dominante était également facilitée lorsque la phrase faisait référence au sens subordonné, alors que lorsque la cible faisait référence au sens subordonné et la phrase au sens dominant, l'intégration du mot cible n'était plus facilitée par rapport à un mot complètement incohérent. Dans une expérience comportementale, Rodd, Lopez Cutrin, Kirsch, Millar, et Davis (2013) ont montré que

cette préférence pour le sens précédemment présenté aux participants pouvait perdurer au moins 20 minutes.

Cette hypothèse apporte une explication à l'effet de biais du subordonné (Rayner, Pacht, & Duffy, 1994). L'effet de biais du subordonné correspond à un ralentissement dans le traitement d'un mot ambigu (en termes de temps de fixation ou de temps de lecture) lorsque celui-ci est présenté dans un contexte favorisant son sens le moins fréquent. Selon le *re-ordered access model*, dans cette situation, le sens subordonné va être activé plus fortement et plus rapidement que dans un contexte neutre alors que le sens dominant sera lui toujours autant activé. C'est la compétition entre les deux sens qui va générer un temps de traitement plus long du mot ambigu dans un contexte favorisant le sens subordonné que dans un contexte favorisant le sens dominant.

5. Résumé du Chapitre II

L'objectif de ce chapitre était de présenter le traitement sémantique des mots. Différents modèles du traitement sémantique, différant par la nature des processus les sous-tendants ont été présentés. Certains modèles considèrent le traitement sémantique comme automatique, d'autres le considère comme stratégique. Enfin, selon une autre classe de modèle l'accès au sens se fait non pas simultanément mais après l'accès lexical. Ces deux distinctions sur les modèles de traitement sémantique (i.e., automatique vs. stratégique et lexical vs. post-lexical) influencent également la littérature sur l'amorçage sémantique, que celui-ci soit étudié de manière comportementale ou électrophysiologique. Enfin, une dernière partie de ce chapitre portait sur les mots ambigus et à l'impact du contexte sur l'activation de leurs différents sens.

Dans cette thèse, c'est principalement la question de l'automatisme du traitement sémantique qui va être étudiée. En effet, la dégradation du signal de parole par le bruit aura pour objectif de masquer le signal cible afin de perturber le traitement sémantique tout en préservant l'intelligibilité. Dans l'Étude 2 nous nous intéresserons également aux interactions entre les processus *bottom-up* et *top-down* en étudiant l'activation des différents sens d'un mot ambigu dans un contexte donné. Encore une fois, la présence de bruit nous permettra de moduler le traitement sémantique en dépit d'une bonne intelligibilité. Les mécanismes sous-tendant ce phénomène seront explicités au Chapitre III.

Chapitre III

La parole dans le bruit

Au quotidien, la parole est rarement perçue dans le silence, en effet, de nombreux bruits environnants comme la circulation, la musique ou encore des conversations alentours perturbent le signal cible (i.e., la parole de notre interlocuteur). Malgré ce constat évident, la recherche en psycholinguistique se fait habituellement dans un environnement extrêmement calme et silencieux. Ce n'est que très récemment que quelques équipes de recherche se sont intéressées au traitement de la parole en situation bruitée (Aydelott, Dick, & Mills, 2006; Boulenger, Hoen, Ferragne, Pellegrino, & Meunier, 2010; Gautreau, Hoen, & Meunier, 2013; Obleser & Kotz, 2010; Strauss, Kotz, & Obleser, 2013). En effet, la plupart des études s'intéressant à la compréhension de la parole dans le bruit se font d'un point de vue psychoacoustique. Ainsi, Cherry fut le premier en 1953 à s'intéresser à ce phénomène qu'il appellera *cocktail party phenomenon*. Cette situation est extrêmement demandante cognitivement parlant, et la capacité à y faire face évolue avec le développement. Les nourrissons par exemple sont capables de discriminer la voix d'un homme et d'une femme en situation de cocktail party (R. S. Newman & Jusczyk, 1996) mais également deux voix de femmes, lorsque l'une d'elle est la voix de leur mère et ce dès 7 mois (Barker & Newman, 2004). Ceci suggère donc qu'ils présentent déjà des capacités de ségrégation de flux, même si elles restent limitées et très dépendantes de la familiarité de la voix à distinguer. Les personnes âgées en revanche perdent rapidement leur capacité à percevoir la parole dans le bruit (Schoof & Rosen, 2014), et ceci représente une plainte importante (CHABA, 1988)

Dans cette thèse, nous allons nous intéresser à la situation de *cocktail party* ou parole dans le bruit d'un point de vue psycholinguistique. C'est-à-dire que la présentation des stimuli au sein d'un signal concurrent (i.e., bruit) est utilisée comme outils pour une meilleure appréhension des mécanismes sous-jacents au traitement de la parole en général. Les différents types de masquages présents en situation de parole dans le bruit sont présentés puis les indices utilisés pour diminuer les effets de ces masquages. Enfin

nous nous intéressons à l'*Effortfulness Hypothesis* qui postule l'existence de liens entre la facilité à effectuer les traitements bas niveau et l'efficacité des traitements de haut niveau.

1. Les masquages énergétique et informationnel

Que la présence d'un bruit puisse masquer le signal de parole est un savoir implicite, nous aurons d'ailleurs naturellement tendance à cesser de parler lorsqu'un bruit fort survient (e.g., une moto accélérant à côté de nous). Nous savons qu'il est inutile de poursuivre notre propos qui sera de toute façon inaudible pour notre interlocuteur. La conscience de l'effet de masquage au quotidien fait référence à l'idée selon laquelle le signal le plus fort sera le mieux entendu. En réalité deux types de masquages vont opérer en situation de parole dans le bruit : le masquage énergétique et le masquage informationnel (Brungart, 2001).

1.1. Le masquage énergétique

Le masquage énergétique fait référence à un masquage purement physique. Simplement, deux bruits partageant les mêmes bandes de fréquence vont stimuler les mêmes parties de la cochlée au même moment. L'un des deux bruits ne pourra ainsi pas être entendu. Le masquage va ainsi apparaître uniquement si le masqueur partage des caractéristiques spectro-temporelles avec le signal d'intérêt. En situation de parole dans la parole, l'effet de masquage énergétique va augmenter proportionnellement avec le nombre de voix dans le fond sonore (Simpson & Cooke, 2005). Le masquage énergétique est donc présent quelle que soit la nature du bruit masquant, du moment qu'il partage les mêmes bandes de fréquences que le signal cible. Le Rapport Signal/Bruit (SNR) c'est-à-dire l'intensité relative du masqueur par rapport au signal cible, va également faire varier l'importance du masquage énergétique. En effet, plus l'intensité du masqueur est importante comparativement au signal cible, plus ce dernier est masqué (Bronkhorst & Plomp, 1988).

1.2. Le masquage informationnel

La définition du masquage informationnel est plus complexe. En effet, il représente les effets de masquage ne pouvant pas être expliqués par le masquage énergétique (Durlach, 2006). Plus spécifiquement, le masquage informationnel fait référence aux interférences n'apparaissant pas au niveau périphérique. Il est particulièrement important en situation de parole dans la parole (Brungart, 2001; Brungart, Chang, Simpson, & Wang, 2006; Brungart, Simpson, Ericson, & Scott, 2001). A SNR (Rapport Signal/Bruit) égal, l'effet masquant du masquage informationnel est beaucoup plus important que celui du masquage énergétique. En effet, Festen et Plomp (1990) ont mis en évidence que le SNR nécessaire pour atteindre le Seuil de Perception de la Parole

(i.e., 50% de détection correcte d'une phrase présentée dans du bruit) était décalé de 6 à 8 dB en condition de parole dans la parole comparativement à une condition de parole dans du bruit fluctuant. La présence des deux types de masquage en situation de parole dans la parole diminue ainsi l'intelligibilité de la parole cible.

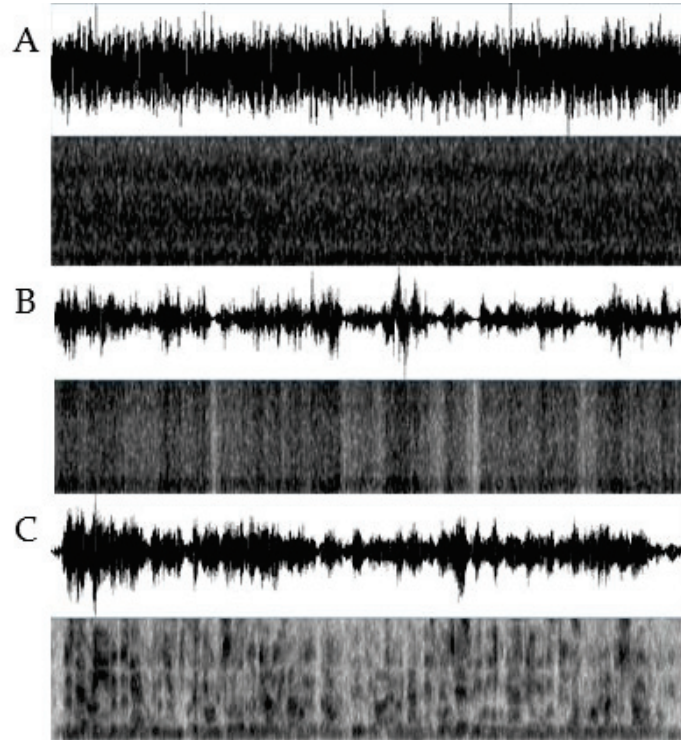


Figure 6

Forme et spectrogramme de différents types de bruits. A. Bruit Stationnaire. B. Bruit Fluctuant. C. Bruit parolier (4 voix).

En situation de parole dans la parole, le masquage informationnel provient principalement des interférences linguistiques entre le signal cible et le fond sonore. Ces interférences se situent aussi bien au niveau phonologique (Gautreau et al., 2013) lexical (Boulenger et al., 2010) que sémantique (Brouwer, Van Engen, Calandruccio, & Bradlow, 2012). Ainsi, Boulenger et al., (2010) ont mis en évidence un effet lexical du masquage informationnel. Des fonds sonores composés de 2, 4, 6, ou 8 voix prononçant des mots ayant une fréquence d'occurrence élevée ($F > 45$ par million) ou faible ($F < 1$ par million) étaient présentés aux participants. Un item cible était présenté 2.5 s après le début du stimulus. Les participants devaient effectuer une tâche de décision lexicale sur cet item. Les temps de réponses ont mis en évidence que la fréquence des mots composant les fonds sonores avait un impact sur le traitement des mots cibles. Ainsi, lorsque les mots du fond sonore avaient une fréquence importante, la reconnaissance des mots cibles était ralentie comparativement à lorsque les mots composant le fond sonore avaient une fréquence d'occurrence faible.

Le masquage informationnel peut également ne pas être en modalité auditive. Puisque sa définition implique simplement qu'il ne provient pas d'interférences au niveau périphérique Mattys et collègues (Mattys, Brooks, & Cooke, 2009; Mattys & Wiget, 2011) définissent le masquage informationnel comme n'importe quelle interférence au niveau central. Le masquage informationnel est donc amodal et peut être généré par l'exécution d'une double tâche. Dans une série d'expérience, Mattys et Wiget (2011) testent l'existence de l'effet Ganong (Ganong, 1980) en condition de charge cognitive. L'effet Ganong réfère à l'influence des connaissances lexicales sur l'identification des phonèmes, plus précisément à la lexicalisation d'un pseudo-mot contenant un phonème ambigu. Les participants doivent ainsi effectuer une tâche de recherche visuelle (i.e., indiquer la colonne et la ligne du carré rouge ; cf Figure 7) tout en écoutant des pseudo-mots (i.e., *gift/kift* ; *giss/kiss*). Ils doivent alors déterminer si le premier phonème était /g/ ou/k/.

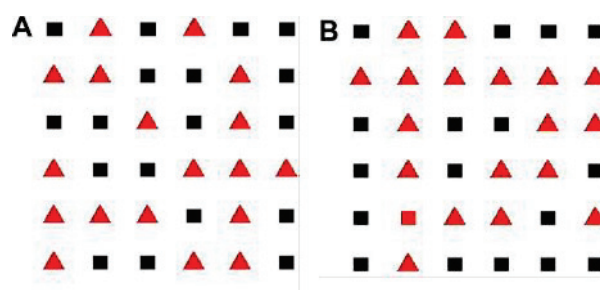


Figure 7

Tiré de Mattys et al., (2011). Exemple de tâche visuelle permettant d'augmenter la charge cognitive. Les participants doivent chercher le carré rouge parmi les carrés noirs et triangles rouges. A. Cible absente. B. Cible présente à la 2^e colonne, 5^e ligne.

Les résultats mettent en évidence une augmentation de l'effet de lexicalité en condition de charge cognitive, comme conséquence de l'appauvrissement du traitement des informations sensorielles. C'est-à-dire que les participants vont augmenter leur dépendance aux informations *top-down* pour compenser le manque de qualité des informations *bottom-up*.

Toutefois, la compréhension de la parole dans le bruit peut également être améliorée grâce à l'utilisation de certains indices bas niveau. Les indices démasquants diminuent l'effet de masquage sur le signal cible, les principaux sont les fluctuations temporelles du signal, et la différence de localisation entre le signal cible et le fond sonore. Les caractéristiques vocales de l'interlocuteur quant à elles facilitent la distinction du signal cible et du fond sonore. Nous présentons ici ces 3 effets, bien que seules les caractéristiques vocales de l'individu sont utilisées afin d'améliorer les performances des participants dans l'Etude 1.

2. Les indices démasquants

2.1. Les fluctuations temporelles du signal

Il existe différentes fluctuations temporelles au sein du signal de parole. L'enveloppe temporelle caractérisée par les variations lentes du signal (i.e., la prosodie) et la structure temporelle fine correspondant aux variations rapides (i.e., les syllabes, les phonèmes).

La présence de ces variations d'amplitude va générer la présence de creux dans le signal (cf. Figure 6C ; B). Au niveau de ces creux, le SNR est modifié, en effet, le signal cible devient plus fort que le masqueur, générant ainsi le démasquage de la cible (Buss, Whittle, Grose, & Hall, 2009; Howard-Jones & Rosen, 1993). De ce fait, un mot présenté dans du bruit fluctuant (Figure 6 B) sera mieux perçu qu'un mot présenté dans du bruit stationnaire (Figure 6 A).

2.2. La localisation

L'origine du son permet de distinguer le signal cible du bruit de fond. En effet, lorsque le signal cible et le fond sonore ne proviennent pas du même endroit de l'espace, l'intelligibilité du signal cible est améliorée, c'est le démasquage spatial (Bronkhorst, 2000; Drullman & Bronkhorst, 2000; Freyman, Helfer, McCall, & Clifton, 1999). Différents indices peuvent être modifiés pour générer un démasquage spatial comme la différence interaurale d'intensité ou la différence interaurale de temps (Bronkhorst & Plomp, 1988). Il est également possible de décaler les phases du signal entre les deux oreilles, donnant une impression de déplacement de la source du signal ayant subi cette modification.

3. La ségrégation de flux

Les indices propres à la voix de l'interlocuteur peuvent être utilisés afin de faciliter la ségrégation. De façon générale, lorsque les différentes voix (signal cible et fond sonore) ne sont pas du même sexe, l'intelligibilité est augmentée comparativement à une condition où les différentes voix appartiennent toutes au même sexe (Brungart, 2001). Plus finement, des différences de F_0 et de longueur de tractus vocal facilitent également la distinction entre deux voix même de même sexe (Darwin, Brungart, & Simpson, 2003). Récemment, un effet facilitateur de la familiarité de la voix a été mis en évidence (Johnsrude et al., 2013). Les auteurs ont présenté la procédure *Coordinate Response Measure (CRM)* à des participants. Cette procédure consiste à présenter des phrases constituées d'un mot d'appel (i.e., *Baron*) puis de deux cibles : un chiffre et une couleur (e.g., *Baron, go to SEVEN BLUE*). Il y a 6 mots d'appel différents, et les participants doivent uniquement écouter lorsque le mot d'appel est *Baron*. Ils doivent

ensuite indiquer, sur un écran le chiffre correspondant à ce qu'a dit la voix, ici le sept bleu. Deux voix étaient présentes dans chaque stimulus, une voix cible prononçant le mot d'appel *Baron* et une voix masqueur. Les participants étaient plus performants lorsque c'était leur époux (se) qui prononçait la phrase cible mais aussi lorsque c'était leur époux (se) le masqueur, suggérant que la familiarité aide à traquer la voix soit pour y prêter attention soit pour l'ignorer.

4. Parole dans le bruit et *Effortfulness Hypothesis*

Comme expliqué plus haut, la plupart des études concernant la parole dans le bruit, s'intéressent en fait à l'intelligibilité d'un signal cible lorsque celui-ci est présenté dans un fond sonore. Toutefois, certaines études se sont également intéressées à l'effet de la dégradation du signal sur un traitement spécifique effectué sur celui-ci. La première étude fut celle de Rabbitt (1968) dans laquelle il demandait à des participants de retenir des séries de 8 chiffres présentés dans du bruit. Les adultes normo-entendants étaient moins performants lorsque les séries de chiffres étaient présentées dans le bruit, alors même que ceux-ci étaient parfaitement intelligibles. Dans une seconde partie, les chiffres étaient présentés en deux séries de 4. Les participants étaient moins performants à rappeler la première série de chiffre qu'elle ait été présentée dans le silence ou dans le bruit lorsque la seconde série avait été présentée dans le bruit. A partir de ces résultats, Rabbitt développe l'*Effortfulness Hypothesis* selon laquelle l'effort engendré par la perception des chiffres dans le bruit laisse moins de ressources cognitives disponibles pour leur mémorisation en mémoire à court-terme.

Plus largement, l'*Effortfulness Hypothesis* postule que les traitements superficiels (i.e., acoustique, phonologique) se font de manière automatique et nécessitent peu de ressources cognitives. Les traitements plus profonds (i.e., mémoire de travail) en revanche, sont coûteux en termes de ressources cognitives. Lorsque l'intelligibilité est bonne, peu de ressources sont allouées aux traitements superficiels, et la majorité des ressources cognitives sont disponibles pour être allouées aux traitements profonds (cf. Figure 8.A.). En revanche, lorsque l'intelligibilité diminue, du fait de la présence d'un masqueur (cf. Figure 8.B.), les ressources cognitives seraient réallouées aux traitements superficiels, pour compenser le manque d'intelligibilité. De ce fait, moins de ressources sont disponibles pour les traitements profonds. En conséquence, on observe une diminution de l'efficacité de ces derniers.

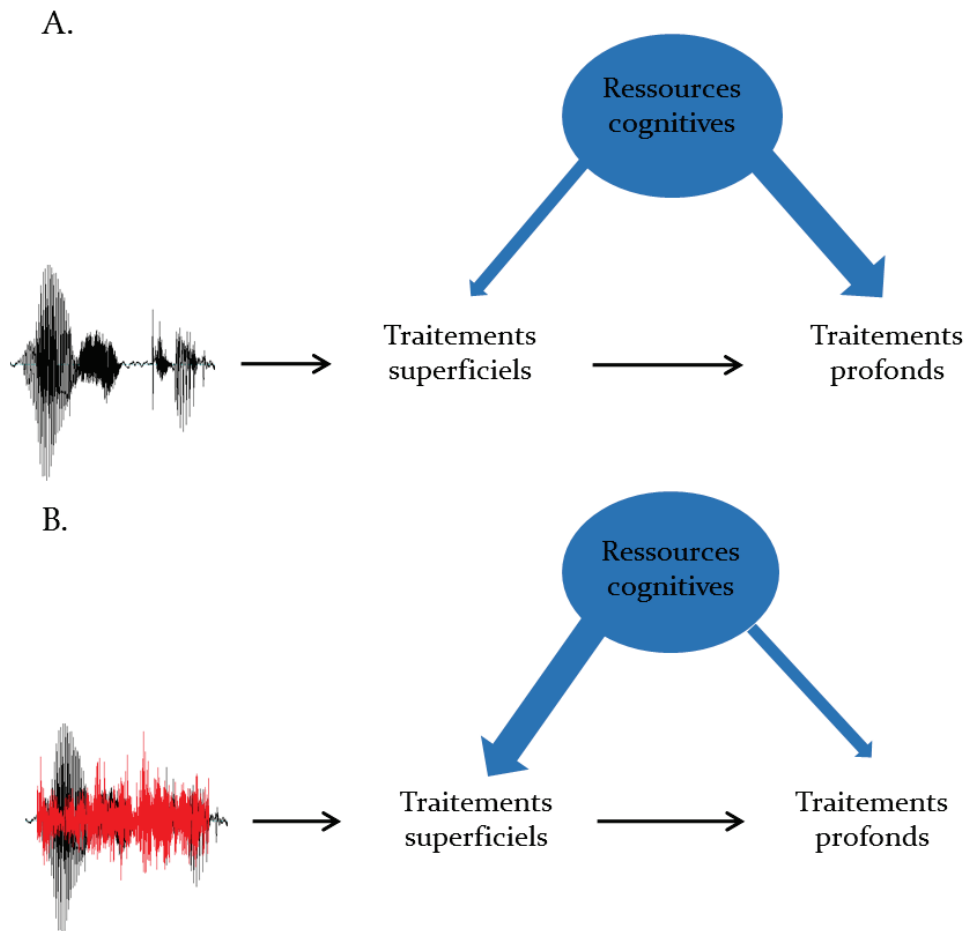


Figure 8

Représentation graphique de l'*Effortfulness Hypothesis* adaptée de Wingfield, Tun, et McCoy (2005) . A. En condition d'intelligibilité optimale, peu de ressources cognitives sont allouées aux traitements superficiels. Les ressources cognitives sont allouées aux traitements profonds. B. En condition d'intelligibilité diminuée, les ressources cognitives sont allouées aux traitements superficiels pour compenser le manque d'intelligibilité, peu de ressources sont donc disponibles pour le traitement profond.

Depuis ce premier papier, de nombreuses études ont répliqué ces résultats, notamment avec des personnes âgées (Murphy, Craik, Li, & Schneider, 2000; Pichora-Fuller, Schneider, & Daneman, 1995) chez qui l'effet est majoré du fait de la diminution des ressources cognitives associées au vieillissement normal (Salthouse, 1996).

L'effet de la saturation des ressources cognitives a aussi pu être observé lorsque l'effort ne provient pas de la qualité du signal mais de l'ajout d'une charge cognitive (e.g., double tâche). Tun, McCoy, et Wingfield (2009) se sont intéressés à la capacité de personnes âgées normo-entendantes et presbycousiques à retenir des listes de mots présentés dans le silence alors qu'elles effectuaient une seconde tâche ou pas. Les résultats ont mis en évidence que l'ajout d'une seconde tâche était plus délétère pour les participants presbycousiques que pour les participants normo-entendants alors

même que leurs performances étaient identiques en condition de tâche unique. Il semble donc que l'ajout d'une charge cognitive à la difficulté à traiter le signal du fait de la presbycousie sature les ressources cognitives ; en conséquence, les performances chutent. De plus, certains mots de la liste partageaient un lien sémantique. Pourtant, les résultats n'ont pas montré d'effet facilitateur significatif de ce lien sémantique suggérant ainsi que les participants ne sont pas capables d'utiliser le lien sémantique entre les mots pour améliorer la mémorisation, suggérant une difficulté dans ces conditions à mettre en place une stratégie de mémorisation. Encore une fois, le masquage informationnel était présenté sous forme de double tâche, confirmant que les ressources cognitives desquelles sont dépendantes les traitements de haut niveau sont amodales.

L'effet de la dégradation du signal a également été mis en évidence sur les traitements langagiers, en plus de la mémoire de travail. En effet, Gao et collaborateurs (Gao, Levinthal, & Stine-Morrow, 2012; Gao, Stine-Morrow, Noh, & Eskew, 2011) ont mis en évidence, en modalité visuelle, que la dégradation d'un texte (par une variation de brillance des pixels) engendrait des difficultés pour l'intégration sémantique des informations comparativement à une condition où aucun bruit visuel n'est présent. Si l'*Effortfulness Hypothesis* a été développée en modalité auditive, elle semble également se vérifier en modalité visuelle et reflèterait donc le lien général entre le traitement du langage et l'intégrité du percept.

5. Résumé du Chapitre III

Ce chapitre présente les processus spécifiques à la situation de parole dans le bruit. Nous avons tout d'abord vu les différents types de masquages existant : le masquage énergétique, dû à des interférences de bas niveau et le masquage informationnel dû à des interférences de haut niveau. Au sein de la partie expérimentale, l'impact du contenu sémantique du fond sonore (i.e., masquage informationnel) sur la reconnaissance d'un item cible est testé (Etude 1). L'augmentation du nombre de voix dans le fond sonore permet d'évaluer comment le traitement sémantique est impacté par l'augmentation du masquage énergétique et du masquage informationnel. Une deuxième étude vise quant à elle à différencier l'effet du masquage énergétique et de la combinaison des masquages énergétique et informationnel sur la profondeur de traitement sémantique dans un contexte phrastique.

Ce chapitre présente également les indices bas niveau permettant d'améliorer l'intelligibilité du signal cible. Particulièrement, les caractéristiques vocales de la voix cible ont été manipulées dans l'Etude 1 afin d'améliorer les performances des participants. Enfin, nous présentons l'*Effortfulness Hypothesis* selon laquelle la dégradation acoustique du signal va générer une diminution de l'efficacité des traitements de haut niveau en dépit d'une intelligibilité préservée. Cette hypothèse est

à la base du travail expérimental présenté, elle permet en effet d'expliquer les liens entre les processus bas niveau et haut niveau. L'objectif du travail expérimental sera donc entre autre d'approfondir cette hypothèse.

Chapitre IV

Le système auditif

Les premiers chapitres de la partie Théorique se sont intéressés au traitement de la parole, à son sens et enfin à la situation de parole dans le bruit. Pourtant, le traitement de la parole est nécessairement tributaire de l'intégrité des systèmes auditifs à la fois périphérique et central. Cette thèse vise à s'intéresser aux liens entre les processus bas niveau, ici le traitement auditif, et les processus haut niveau, ici les représentations langagières. L'Axe 1 s'intéresse aux dégradations acoustiques du signal lorsque la parole est dégradée par la présence de bruit. La dégradation est alors transitoire et n'affecte donc pas les représentations du langage. L'Axe 2 en revanche, s'intéresse aux dégradations permanentes du signal, du fait d'un traitement auditif déficitaire.

Nous laissons de côté les atteintes périphériques (i.e., la surdité) pour nous intéresser aux déficits au niveau des traitements auditifs centraux (i.e., non périphériques). Ce déficit du traitement auditif peut provenir d'un manque de maturité puisque les traitements auditifs sont matures très tardivement (Etude 3). Enfin, il peut provenir d'un Trouble du Traitement Auditif (TTA), souvent associé à des troubles du langage (Etude 4).

L'objectif de ce chapitre IV est de présenter les systèmes auditifs. Nous voyons tout d'abord brièvement le système auditif périphérique, la compréhension superficielle de son fonctionnement permet de mieux appréhender certains phénomènes comme le masquage énergétique (cf. p58). Puis nous nous intéressons à l'anatomie et au fonctionnement du système auditif central ainsi qu'à son développement, de la naissance à l'âge adulte. Ce chapitre présente ensuite le cas particulier des TTA, la définition de ce trouble et les outils diagnostiques disponibles pour les dépister. En dernier lieu, nous présentons la Batterie d'Evaluation des Compétences Auditives Centrales (BECAC ; Donnadieu et al., 2014) qui a été utilisée comme matériel expérimental dans les Etudes 3 et 4.

1. Le système auditif périphérique

Le système auditif périphérique est composé de trois parties, l'oreille externe, l'oreille moyenne et l'oreille interne.

1.1. L'oreille externe

L'oreille externe est composée du pavillon auriculaire, ainsi que du conduit auditif externe. La forme du pavillon auriculaire permet d'amplifier les différentes fréquences avant que l'onde sonore n'entre dans le conduit auditif externe. Lorsque le son est latéralisé, la tête et le torse vont modifier l'amplitude des différentes fréquences. L'amplitude des sons de basses fréquences (de 100 Hz à 10000 Hz) est ainsi amplifiée, alors que l'amplitude des sons de hautes fréquences (plus de 10000 Hz) sera diminuée. Ce phénomène s'explique par la petite taille des structures du pavillon auriculaire, de ce fait, seules les ondes à faible période (i.e., haute fréquence) vont être altérées.

En plus de ces modifications apportées par le torse, la tête et le pavillon auriculaire, l'oreille externe engendre également une amplification de l'amplitude des fréquences comprises entre 1500 Hz et 7000 Hz. Cette amplification provient majoritairement de la conque et du conduit auditif externe. En effet, du fait de sa taille, la conque a une fréquence de résonance de 5000 Hz alors que le conduit auditif externe a une fréquence de résonance d'environ 2500 Hz. Ces deux structures amplifient ainsi l'intensité des sons qui sont présentés à l'oreille et dont la fréquence est comprise entre 1500 Hz et 7000 Hz. Une fois l'onde sonore amplifiée sur ces différentes fréquences, elle se dirige vers l'oreille moyenne (cf. Figure 9).

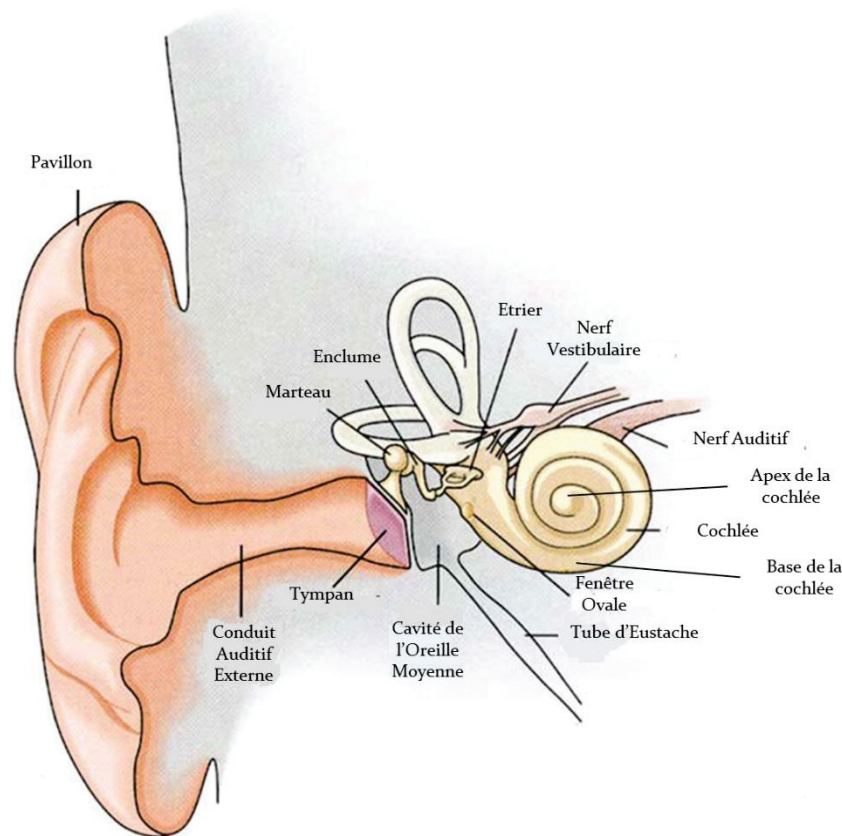


Figure 9

Tiré de Graven et Browne (2008). Représentation de l'oreille humaine

1.2. L'oreille moyenne

L'oreille moyenne commence au niveau de la membrane tympanique. Celle-ci va isoler le tympan du conduit auditif externe. Elle est attachée à la chaîne des osselets, composée du marteau, de l'enclume et de l'étrier eux même liés à la fenêtre ovale et donc à l'oreille interne. La principale fonction de l'oreille moyenne va être de transmettre l'onde sonore entre l'oreille externe et l'oreille interne. Pour ce faire, 3 solutions sont possibles : (a) le son peut être conduit par transmission osseuse (i.e., transmission par les os du crâne) ; (b) le son peut être conduit par l'air, c'est-à-dire que les vibrations de la membrane tympaniques font vibrer l'air contenu dans la cavité tympanique et la fenêtre ovale ; (c) le son peut être amplifié par les osselets. Les deux premières possibilités ne permettent pas l'amplification du son, or, comme le passage entre l'oreille moyenne et l'oreille interne correspond au passage d'un milieu aérien à un milieu liquide (beaucoup plus résistant), si le son n'est pas amplifié, les vibrations en milieu liquide ne sont pas assez importantes pour engendrer un percept. Pour exemple, la disparition de la chaîne des osselets chez un être humain engendrerait une perte auditive de 60 dB. Le son est donc principalement transmis de l'oreille moyenne à l'oreille interne via la chaîne des osselets. Lorsque la membrane tympanique vibre

sous l'effet des changements de pression de l'air du à l'onde sonore, elle met en mouvement la chaîne des osselets. Celle-ci agit comme un levier et répercute le mouvement de vibration sur la fenêtre ovale. Du fait de cet effet de levier d'un osselet sur l'autre et de la plus petite surface de la fenêtre ovale comparativement à la membrane tympanique, l'amplitude de l'onde sonore est amplifiée.

L'oreille moyenne a également pour fonction de protéger l'oreille interne en cas de bruit trop intense grâce au réflexe stapédien. En effet, en cas de trop grande pression sur la membrane tympanique, un muscle reliant la caisse du tympan et l'étrier se contracte afin de limiter la pression de l'étrier sur la fenêtre ovale, l'amplification de la transmission de l'onde sonore est ainsi diminuée.

1.3. L'oreille interne

L'oreille interne est incluse dans l'os temporal et est composée de 3 parties : les canaux semi-circulaires, le vestibule et la cochlée. Les canaux semi-circulaires et le vestibule, font partie du système vestibulaire responsable du maintien de l'équilibre et non de l'audition. Nous allons donc nous intéresser uniquement à la cochlée, au sein de laquelle se fait la transduction mécano-électrique.

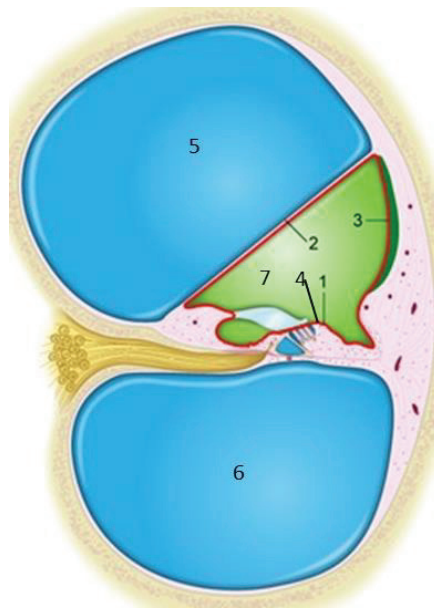


Figure 10

Schéma d'un tour de la cochlée. Adapté de <http://www.cochlea.eu/cochlea>. 1. Lamelle Réticulaire. 2. Membrane de Reissner. 3. Mur Latéral. 4. Organe de Corti. 5. Rampe Vestibulaire. 6. Rampe tympanique. 7. Canal Cochléaire.

1.3.1. Anatomie de la cochlée

La cochlée est un tube enroulé sur lui-même environ 2,6 fois dont le diamètre est décroissant de la base près de l'étrier vers l'apex (cf. Figure 9). Elle est composée de trois canaux remplis de liquide : (a) la rampe tympanique ; (b) la rampe vestibulaire ; toutes deux entourant (c) le canal cochléaire. La membrane basilaire sépare la rampe tympanique et le canal cochléaire alors que la membrane de Reissner sépare la rampe vestibulaire et le canal cochléaire (cf. Figure 10). Ce dernier est rempli d'endolymphe alors que les rampes vestibulaires et tympaniques sont remplies de périlymphe. La structure de la membrane basilaire donne lieu à une représentation tonotopique des fréquences. Les régions basales sont stimulées par les fréquences élevées alors que les régions apicales sont stimulées par les fréquences basses. L'organe de Corti présent le long de la membrane basilaire contient les cellules sensorielles (i.e., les cellules ciliées). Il en existe deux types, les cellules ciliées externes, organisées le long de la membrane basilaire en 3 à 5 rangées, et les cellules ciliées internes organisées en une seule rangée. Chaque cellule ciliée contient entre 50 et 150 stéréocils. Les stéréocils des cellules ciliées externes sont insérés dans la membrane tectoriale contrairement à ceux des cellules ciliées internes. L'organe de Corti est également composé de cellules de soutien : les cellules de Deier et les cellules de Henson.

1.3.2. Physiologie de la cochlée

Lorsqu'un son atteint la cochlée, la mise en mouvement de la périlymphe génère également un mouvement de la membrane basilaire. Du fait de son élargissement de la base vers l'apex, les mouvements de la périlymphe n'influencent pas la membrane basilaire de la même façon au même endroit. C'est ainsi que la fréquence du son perçu fait vibrer différemment la membrane, permettant un encodage différencié des fréquences. C'est la tonotopie passive. Les cellules ciliées externes présentes sur la membrane entrent également en mouvement. L'étirement de ces cellules du fait de leur insertion dans la membrane tectoriale engendre l'ouverture de canaux à cations. Le potassium contenu dans l'endolymphe entre ainsi dans les cellules ciliées externes et les dépolarise, créant ainsi un influx nerveux. Plus l'intensité du son perçu est importante, plus les cellules ciliées sont stimulées et ont une fréquence de décharge importante. En moyenne, les cellules adjacentes diffèrent par leur fréquence caractéristique de seulement 0.2%. L'influx nerveux ainsi créé se propage le long du nerf auditif travers le système auditif.

2. Le système auditif central

Le système auditif central est composé de différentes structures (cf. Figure 11) : (a) les noyaux cochléaires ; (b) les complexes olivaires supérieurs ; (c) les noyaux du lemniscus latéral ; (d) les colliculus inférieur ; (e) les corps géniculés médiaux et (f)

les aires auditives. À noter que toutes ces structures présentent une organisation tonotopique. De même que pour la cochlée, au sein de chaque structure il existe une correspondance entre la localisation topographique d'un neurone et sa fréquence de prédilection.

Les voies auditives ascendantes se projettent de façon bilatérale, bien qu'elles se projettent de façon prédominante vers l'hémisphère contralatéral (70-80%). À chaque étape, des commissures permettent ces projections vers le côté contralatéral, et la liaison de chaque structure avec son homologue droite ou gauche.

2.1. Le noyau cochléaire

Les neurones auditifs primaires provenant de la cochlée, se divisent en deux branches se projetant ipsilatéralement sur le noyau cochléaire (cf. Figure 11). La branche antérieure se termine alors dans le noyau cochléaire ventral antérieur (NCVa) alors que la branche postérieure se termine dans le noyau cochléaire ventral postérieur (NCVp) et le noyau cochléaire dorsal (NCD). C'est au niveau des noyaux cochléaires que s'effectue un premier décodage des caractéristiques de l'information sonore, en intensité, durée et fréquence.

2.2. Le complexe olivaire supérieur

Les neurones en provenance des noyaux cochléaires se projettent de façon bilatérale sur le complexe olivaire supérieur. L'information pour chacune de ces structures provient donc des deux oreilles. C'est à ce niveau que sont traités les différents indices permettant la localisation spatiale du son perçu. En effet, la réception des signaux provenant de chaque oreille permet de traiter les différences interaurales de temps (ITD ; *Interaural Time Difference*) et d'intensité (ILD ; *Interaural Level Difference*). Les ITD et ILD correspondent en fait aux différences de temps et d'intensité que présente le signal entre les deux oreilles lorsqu'il n'est pas présenté face à la tête. En effet, un son présenté à droite mettra moins de temps à atteindre l'oreille droite que la gauche, il perd également de son intensité entre les deux oreilles du fait de la présence de la tête.

Pour que les ITD puissent être traitées, il faut que le mécanisme de *phase locking* soit mis en œuvre. C'est-à-dire que la fréquence du son soit encodée par la fréquence de décharge des neurones. Les ITD ne sont donc utilisées que lorsque la fréquence du son est suffisamment basse (i.e., 3 kHz), elles sont traitées au sein de l'olive supérieure moyenne. Au-delà, c'est plutôt les ILD qui sont utilisées pour la latéralisation des sons, elles sont traitées au sein de l'olive supérieure latérale. De plus, les ILD sont très sensibles à la fréquence du son, les sons de haute fréquence produisent de plus grandes ILD que les sons de basse fréquence.

2.3. Le lemnicus latéral

Une partie des fibres nerveuses en provenance du noyau cochléaire se projette directement sur le lemnicus latéral contralatéral, sans passer par le complexe olivaire supérieur. Ce chemin particulier répond uniquement aux sons présentés monoralement. Cette structure encode le début et la durée de ce son, sans en encoder la fréquence. Toutefois le rôle précis de cette structure au sein des structures auditives centrales reste peu connu.

2.4. Le colliculus inférieur

Le colliculus inférieur reçoit des projections non seulement du lemnicus latéral, mais également du complexe olivaire supérieur et directement aussi du noyau cochléaire. Il a pour spécificité le traitement des informations temporelles complexes. Certains neurones du colliculus inférieur répondent uniquement aux sons présentant des modulations de fréquence ou encore à des sons d'une durée particulière.

2.5. Le corps genouillé médian du thalamus

Le dernier relai avant le cortex auditif primaire est le corps genouillé médian du thalamus, il reçoit ses projections du colliculus inférieur. Les informations spectrales et temporelles traitées au préalable y sont combinées. Ce relai est également considéré comme une zone de convergence multisensorielle.

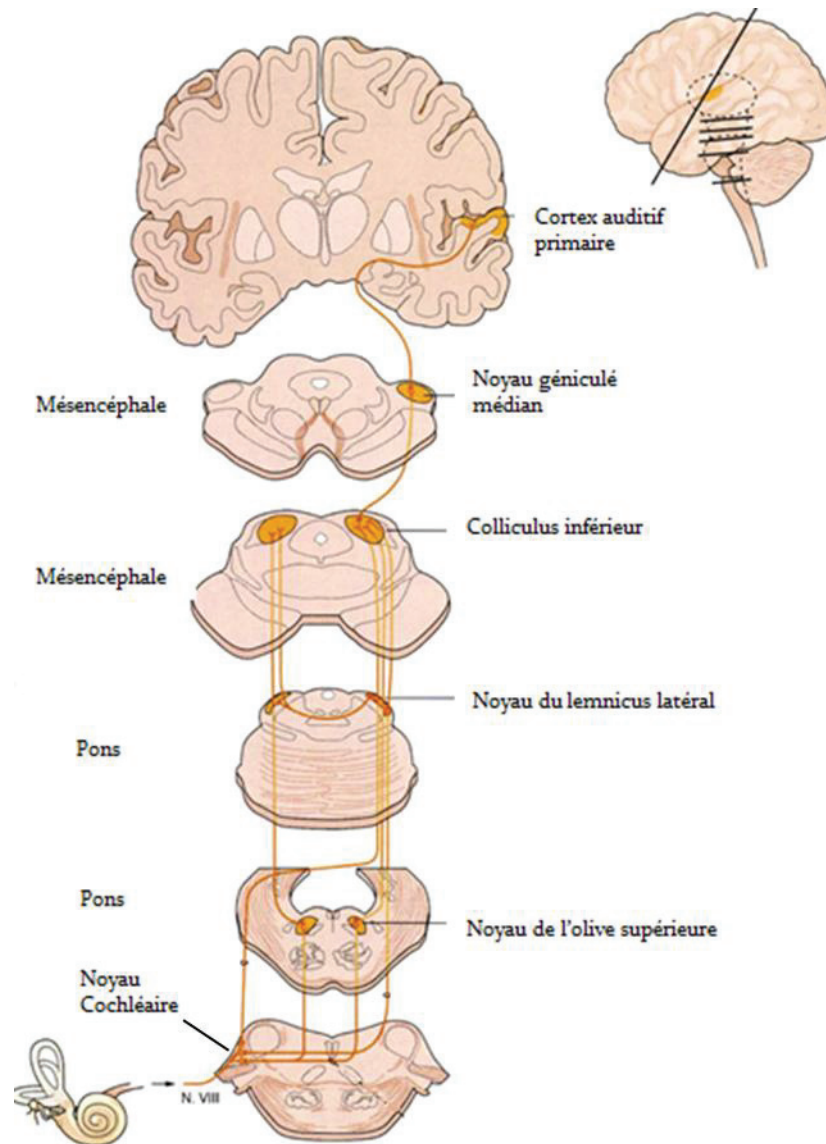


Figure 11

Adapté de Graven et Brown, (2008). Représentation des structures auditives centrales et des liens entre elles.

2.6. Le cortex auditif primaire

Le but ultime de l'information en provenance de la cochlée est bien sûr le cortex auditif primaire. Celui-ci est situé au niveau du gyrus de Heschl, à l'intérieur de la scissure sylvienne (cf. Figure 12.A). Il est entouré par le planum polaire au niveau antérieur et par le planum temporale au niveau postérieur. Comme les autres relais auditifs, le cortex auditif primaire présente une organisation tonotopique (cf. Figure 12.B). Le cortex auditif primaire est la première zone d'intégration du message auditif. Il va s'y mettre en place des mécanismes de décodage de la fréquence, de l'intensité et de la localisation du message.

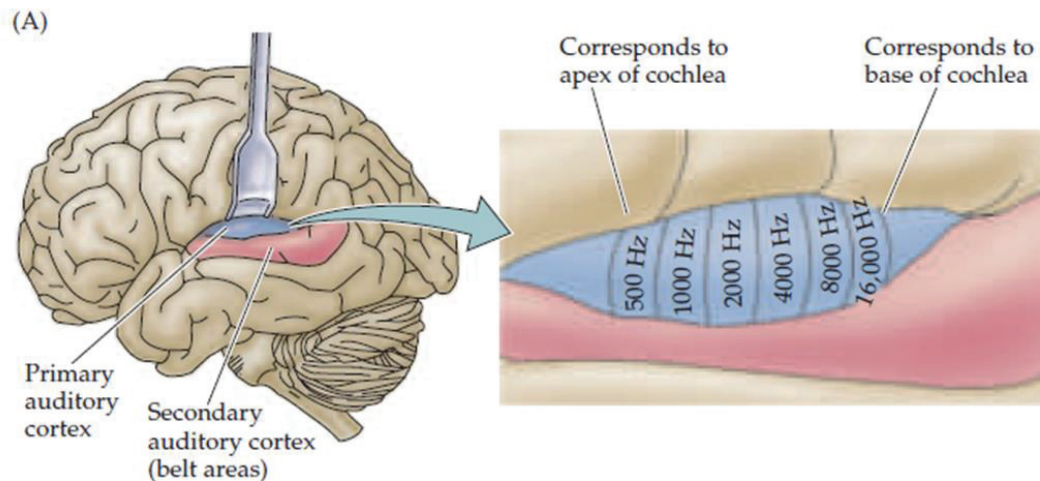


Figure 12

Tiré de Fitzpatrick (2004) Illustration de l'organisation tonotopique du cortex auditif primaire

La qualité du signal de parole perçu va ainsi dépendre du bon fonctionnement de ces structures. Celui-ci peut ne pas être optimal, soit du fait d'une immaturité s'inscrivant dans le développement normal (Etude 3), soit du fait de TTA (Etude 4). Nous nous intéressons donc en premier lieu au développement des compétences auditives centrales, puis à leur dysfonctionnement dans le cas des TTA. Bien que les fonctions des structures décrites soient connues, il est délicat de faire correspondre un traitement spécifique à une structure puisque les différents aspects du son sont traités à plusieurs reprises, plus ou moins finement tout au long du système auditif central. De ce fait, nous allons à présent nous intéresser au développement des traitements auditifs centraux, indépendamment du développement des structures présentées ci-dessus.

3. Le développement des traitements auditifs centraux

Le système auditif périphérique est mature très tôt dans le développement. La capacité à entendre débute dès la 25^e semaine in utéro (Graven & Browne, 2008), et les traitements périphériques sont matures dès la 1^{ère} année. De la 25^e semaine in utero à 5-6 mois, c'est la période critique. Durant cette période, le nerf auditif et les cellules ciliées de la cochlée vont se spécifier dans le traitement d'une fréquence donnée (cf. p71). Afin que cette spécification se fasse de façon adéquate, le nourrisson doit être exposé à différents sons et être en mesure de les discriminer. C'est-à-dire que le bruit de fond ne doit pas excéder 60 dB, en particulier dans les basses fréquences.

Les traitements auditifs centraux en revanche se développent beaucoup plus lentement, jusqu'à l'adolescence voir début de l'âge adulte (Ludwig et al., 2014; Moore, 2002).

Bien sûr, l'étude du développement implique de pouvoir déterminer la part respective des facteurs sensoriels et non-sensoriels des comportements observés (Moore, Cowan, Riley, Edmondson-Jones, & Ferguson, 2011). En effet, certains facteurs non-sensoriels comme l'attention se développent avec l'âge et ont bien évidemment un impact sur les performances obtenues à une tâche. Toutefois, de nombreuses études en électrophysiologie ont permis de mettre de côté l'hypothèse selon laquelle tous les effets développementaux seraient dus à la maturation des facteurs non-sensoriels (Moore, Ferguson, Halliday, & Riley, 2008), sans pour autant nier leur impact sur les résultats comportementaux (cf. p83). En effet, les potentiels évoqués générés par la présentation d'un stimulus sonore sont indépendants de l'attention des participants

3.1. Aspects électrophysiologiques

La maturation tardive du système auditif chez l'humain proviendrait, d'après des études morphologiques et physiologiques du développement tardif du système thalamocortical. Les structures du tronc cérébral quant à elles seraient matures dès la deuxième année de vie, comme montré par l'étude des réponses ABR (Auditory Brainstem Response ; Eggermont, 1988; Ponton, Eggermont, Coupland, & Winkelaar, 1992).

Au niveau cortical, l'étude de la latence et de l'amplitude des potentiels évoqués permet de mettre de côté les facteurs non-sensoriels (Draganova et al., 2005; Draganova, Eswaran, Murphy, Lowery, & Preissl, 2007; Eggermont & Ponton, 2003; He, Hotson, & Trainor, 2007). La présentation d'un son dans le silence, engendre systématiquement une réponse électrophysiologique composée de 2 complexes principaux P1/N1 et P2/N2. Ces deux complexes apparaissant dans les 200 ms suivant la présentation du stimulus sont censés refléter les traitements corticaux précoces des stimuli auditifs. Ainsi, P1 serait généré par le gyrus de Hesch et N1 par le cortex auditif temporal supérieur. Ces deux composantes restent immatures jusqu'à la fin de l'adolescence (i.e., 15-16 ans en ce qui concerne leur latence et leur amplitude ; Ponton, Eggermont, Kwong, & Don, 2000) (cf. Figure 13).

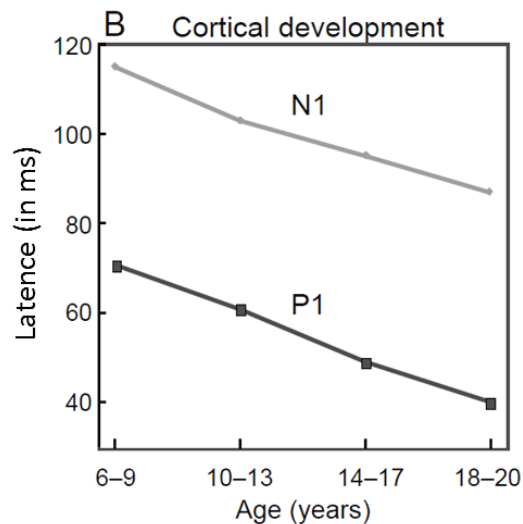


Figure 13

Tiré de Moore (2002) ; d'après les résultats de Ponton et al., (2000). Evolution de la latence (en ms) de P1 et N1 en fonction de l'âge.

P2 en revanche, ne proviendrait d'une activité neuronale au niveau du cortex temporal mais reflèterait au moins en partie l'activité du système réticulaire mésencéphalique. Ponton et al. (2000) n'observent pas de diminution de sa latence entre 6 et 18 ans, mais une diminution d'amplitude. Enfin, N2 serait influencée par l'attention et des variables du stimulus comme la probabilité ou l'intensité. Elle atteint les valeurs adultes à 17 ans en ce qui concerne son amplitude et 18 ans en ce qui concerne la latence.

De façon générale, Ponton et al., (2000) observent une évolution différente de la latence et de l'amplitude. La première diminuant de façon graduelle alors que la seconde évolue de façon abrupte. Cette différenciation confirme qu'il existe différents mécanismes sous-tendant ces deux caractéristiques. En effet, les modifications de latence des ERP vont dépendre principalement du degré de myélinisation des axones (Eggermont, 1988) alors que l'amplitude sera davantage dépendante du nombre de synapses contribuant au potentiel post-synaptique.

L'utilisation de l'électrophysiologie permet donc d'étudier le développement du traitement des stimuli auditifs mais elle est peu adaptée à la situation de diagnostic, aussi bien du fait du temps nécessaire à la mise en place du matériel, que du coût (en termes monétaire et temporel) de celui-ci. Puisqu'un des objectifs de cette thèse est de tester une batterie qui pourrait à terme être utilisée pour évaluer l'intégrité des traitements auditifs centraux, l'utilisation de paradigmes comportementaux semble plus adaptée.

3.2. Développement comportemental

Pour la plupart des compétences dont nous étudions le développement dans la partie Expérimentale, de nombreux tests ont été utilisés. Chaque test, même s'il a pour objectif de tester un traitement donné, implique obligatoirement d'autres processus. Par exemple, pour tester le traitement temporel d'un stimulus deux tests sont fréquemment utilisés : le test de *gap detection* (i.e., détection d'un silence au sein d'un bruit blanc) et la discrimination de durée (Phillips, Comeau, & Andrus, 2010; Werner, Fay, & Popper, 2011). Le test de *gap detection* permet d'évaluer le traitement temporel puisque les participants doivent détecter la fin d'un stimulus et le commencement d'un autre (i.e., avant et après le silence). Le test de discrimination de durée (type paradigme AX par exemple) permet aussi de tester le traitement temporel d'un stimulus auditif. Les tests de discrimination demandent également aux participants d'être capables de retenir et de comparer deux stimuli. Ces différences entre les tests peuvent donc engendrer des résultats difficilement comparables, c'est pourquoi seuls les tests utilisés dans la partie Expérimentale seront présentés dans cette partie. Pour 2 des tests de la batterie, aucunes données développementales n'ont été trouvées, ils ne seront donc pas mentionnés dans cette partie. Nous présentons également, les données développementales concernant la perception de la fréquence, la durée et l'intensité.

3.2.1. Traitement fréquentiel

3.2.1.1. Résolution fréquentielle

Bien que la cochlée soit mature dès la naissance, le développement de l'encodage spectral est plus long. L'étude de Schneider, Trehub, Morrongiello, et Thorpe (1989) est l'une des premières à s'être intéressée spécifiquement au développement de la résolution de fréquence. La résolution de fréquence fait référence à la capacité à traiter suffisamment finement les informations spectrales pour les distinguer au sein d'un son complexe. Des participants âgés de 6,5 mois à 20 ans ont ainsi été évalués sur leur capacité à détecter les différentes fréquences présentes dans un son complexe. Ils devaient ainsi déceler la présence d'un bruit cible d'une largeur fréquentielle d'une octave centrée à 0.4, 1, 2, 4, et 10 kHz au sein d'un bruit blanc. Les nourrissons étaient testés grâce à un paradigme de réflexe d'orientation conditionné. Ils sont conditionnés à tourner la tête en direction du stimulus lorsque celui-ci contenait le bruit cible au sein du bruit blanc. Le conditionnement se fait grâce à une récompense visuelle, i.e., un jouet qui bouge lorsque le bruit blanc contenait effectivement le bruit cible. Le même type de paradigme est utilisé pour l'ensemble des tests présentés sur les nourrissons. Les enfants dès 6 ans et les adultes devaient indiquer le côté où était présenté le signal en pressant le bouton correspondant. Comme attendu, l'analyse des résultats a mis en évidence une amélioration des performances avec l'âge, particulièrement après l'âge de 8 ans. Toutefois, la capacité à distinguer le bruit cible

au sein du masqueur était la même quelle que soit la fréquence du bruit cible, et ce, à tous les âges testés. Il semble donc que la résolution de fréquence soit mature dès l'âge de 3 mois. De plus, Spetner et Olsho (1990) ont mis en évidence que les nourrissons de 3 mois étaient autant capable de distinguer un son pur de 500 ou 1000 Hz au sein d'un masqueur de fréquence variable que les adultes. Lorsque le son pur était de 4000 Hz, seuls les nourrissons de 6 mois obtenaient des performances comparables à celles des adultes. Ces résultats semblent donc confirmer la maturation assez rapide de la résolution fréquentielle au cours du développement.

3.2.1.2. Discrimination de fréquence

La discrimination de fréquence semble se développer plus tardivement que la résolution fréquentielle. Concernant les hautes fréquences (i.e., 4000 Hz) les enfants semblent atteindre les performances de discrimination des adultes entre 6 mois et 12 mois (Olsho, Koch, & Halpin, 1987). Ce résultat est cohérent avec l'observation d'une maturité plus précoce pour le traitement des hautes fréquences des basses fréquences (Eggermont, 2014). Cette différence concernant l'âge de maturation entre les hautes et basses fréquences s'explique très vraisemblablement par le fait qu'elles soient encodées par différents mécanismes (i.e., *phase locking*). Lorsque la fréquence testée est inférieure à 4000 Hz, il semble que la maturation se fasse plus tardivement, vers la fin de l'enfance (i.e., de 8 à 11 ans selon les études ; Banai & Yuvak-Weiss, 2013; Moore et al., 2011; Moore et al., 2008) voir même durant l'adolescence (Halliday, Taylor, Edmondson-Jones, & Moore, 2008).

3.2.2. Traitement temporel

3.2.2.1. Perception de durée

Il existe différentes façons de s'intéresser au traitement temporel auditif. Une des techniques les plus populaires consiste à demander au participant d'effectuer une tâche de *gap detection*. C'est-à-dire que l'on présente au participant un bruit à large spectre. Ce bruit est interrompu pour une durée déterminée et sur une bande de fréquence déterminée (i.e., le *gap*). La durée de l'interruption est variée et le seuil, c'est-à-dire la plus petite interruption détectable par le participant est déterminé (Ahissar, Protopapas, Reid, & Merzenich, 2000; McAnally & Stein, 1996; Rosen, 2003). En utilisant ce paradigme, Buss, Hall, Porter, et Grose (2014) ont mis en évidence une maturation du traitement temporel variable et qui varierait entre 7 et 10 ans en fonction des caractéristiques du bruit. D'autres auteurs ont mis en évidence une maturation plus tardive, durant l'adolescence (Fischer & Hartnegg, 2004). D'autres ont en revanche montré une maturation autour de 6-7 ans (Trehub, Schneider, & Henderson, 1995) évoquant même la possibilité que cette différence ne soit que le reflet des différences de facteurs non-sensoriels.

3.2.2.2. Discrimination de durée

Les traitements temporels peuvent également être testés, comme le sont l'intensité et la fréquence avec des tâches de discrimination (Banai & Ahissar, 2006; Banai & Yuvak-Weiss, 2013; Jensen & Neff, 1993). Ainsi, Morrongiello et Trehub (1987) ont observé une amélioration des capacités de discrimination de durée entre un groupe de nourrissons (6 mois) et un groupe de jeunes enfants (5 ans et 6 mois) et entre les jeunes enfants et les adultes capables de détecter des changements d'au minimum 20, 15 et 10 ms respectivement pour un bruit blanc de 200 ms. Jensen et Neff (1993) obtiennent des résultats congruents avec ces données, puisque selon leur étude la capacité à discriminer la durée de deux stimuli continue de s'améliorer après l'âge de 6 ans.

3.2.3. Traitement de l'intensité

3.2.3.1. Perception de l'intensité

La perception de l'intensité d'un stimulus évolue tout au long de la vie (Werner et al., 2011; cf. Figure 14). La plupart des études s'étant intéressées au développement de la perception de l'intensité l'ont fait uniquement dans le but de s'intéresser au développement des traitements auditifs périphériques. Il semble que l'intensité nécessaire à la détection d'un son diminue au cours du développement jusqu'à l'âge adulte pour une fréquence donnée. Trehub, Schneider, et Endman (1980) ont pu tester la perception de l'intensité sur plusieurs bandes de fréquences centrées à 200 Hz, 400 Hz, 1000 Hz, 2000 Hz, 4000 Hz et 10000 Hz. Les nourrissons de 6, 12 et 18 mois présentent des seuils de détection plus élevés que les adultes sauf à 10000 Hz. Ainsi, les changements développementaux étaient plus importants pour les fréquences basses que pour les hautes fréquences. Les seuils augmentent ensuite à nouveau avec l'âge du fait du vieillissement du système auditif périphérique (i.e., presbycousie).

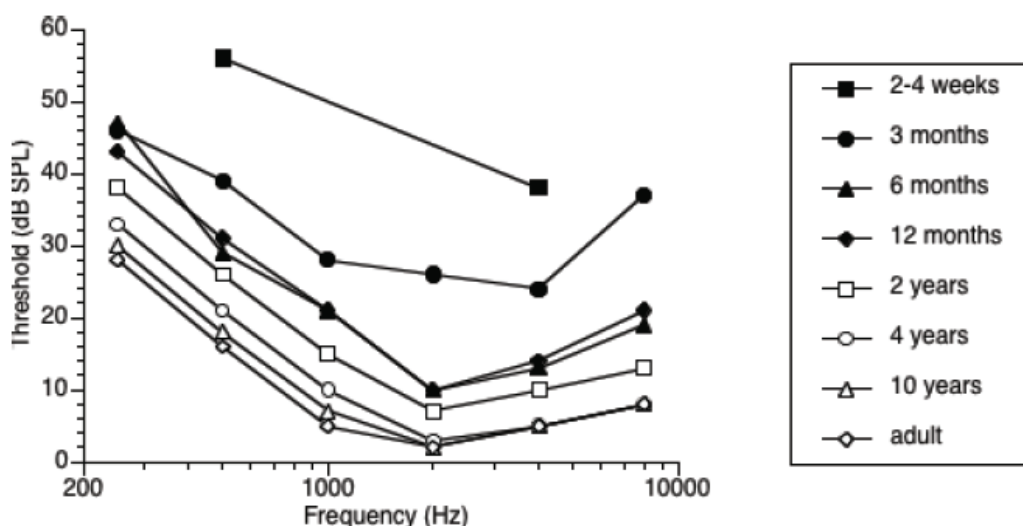


Figure 14

Tiré de Werner et al. (2011). Représentation graphique des seuils de perception en dB SPL en fonction de la fréquence au cours du développement.

3.2.3.2. Discrimination d'intensité

Les capacités de discrimination d'intensité peuvent être testées de différentes manières, soit par comparaison de deux stimuli, i.e., deux stimuli identiques en termes de fréquence et de durée sont présentés aux participants qui doivent indiquer lequel était présenté avec la plus grande intensité (Jensen & Neff, 1993). Une seconde technique consiste à augmenter légèrement l'intensité d'un stimulus présenté. Durant la présentation du stimulus en question, il est demandé au participant d'indiquer le moment où l'intensité augmente (Berg & Boswell, 2000). L'avantage de cette seconde tâche comparativement à la première est qu'elle peut être utilisée plus facilement sur des nourrissons et des jeunes enfants. Les résultats de ces différentes études montrent en général que la discrimination d'intensité est mature à partir de 5-6 ans.

3.2.4. Latéralisation et localisation

La capacité à situer un son dans l'espace apparaît très tôt dans le développement. Elle est d'ailleurs utilisée comme indice de réponse chez les nourrissons dès 3 mois. Toutefois, ces tâches demandent une capacité de latéralisation de 45°, alors qu'un être humain adulte est capable de latéraliser un son à 1° près (Fitzpatrick, 2004). La différence entre latéralisation et localisation provient du paradigme utilisé pour les tester. La localisation fait référence à une situation expérimentale où le participant est assis dans une pièce dans laquelle seront placés des haut-parleurs à différents angles. Il doit ensuite déterminer d'où provient le son perçu. La latéralisation en revanche, fait référence à une situation expérimentale où le participant est équipé d'un casque, et les

stimuli vont être manipulés pour donner l'impression que le son perçu provient d'un côté ou de l'autre. La latéralisation va principalement se faire grâce à deux indices, les différences interaurales de temps (ITD) et les différences interaurales d'intensité (ILD ; cf. p72). Habituellement, un seul de ces indices va être varié en condition expérimentale (Altmann et al., 2014). La localisation se fait également grâce à ces indices mais cette situation expérimentale ne permet pas de distinguer un indice de l'autre. Ce paradigme est également plus facilement utilisable lorsque l'on s'intéresse à une population de nouveau-nés et/ou nourrissons (Morrongiello 1987).

En utilisant un paradigme de localisation, Morrongiello, Fenwick, Hillier et Chance (1994), ont ainsi pu mettre en évidence que les nouveau-nés étaient capables non seulement de tourner la tête efficacement vers l'origine d'un son mais également le suivre s'il se déplace. L'angle minimum audible a été évalué à 30° pour les nouveaux nés, il atteindrait 14° à 7 mois et le niveau adulte vers 5 ans. Une étude s'intéressant au développement des compétences de latéralisation basée sur l'ITD chez des enfants âgés de 5 à 9 ans n'a pas mis en évidence d'effet du développement, suggérant que les enfants de 5 ans présentaient des compétences de latéralisation similaires à celles des enfants de 9 ans. Toutefois, une différence significative est apparue entre les enfants et le groupe adulte (Van Deun et al., 2009).

3.2.5. Ségrégation de Flux

Bregman et Campbell (1971) définissent le concept de flux auditif comme « une séquence d'évènements auditifs dont les éléments sont reliés perceptuellement les uns aux autres, ce flux étant perceptuellement distinct des autres évènements auditifs » [notre traduction]. La ségrégation de flux est donc la capacité à distinguer deux flux auditifs. Cette compétence est particulièrement importante en situation de parole dans le bruit (cf. p61). La ségrégation de flux se teste habituellement avec du matériel non langagier, par exemple, en présentant une série de sons purs, alternant en fréquence, et dont le SOA est plus ou moins important. Lallier et al. (2009) ont ainsi mis en évidence que les enfants bons lecteurs de 10 ans présentaient un seuil de fusion équivalent à celui des adultes. C'est-à-dire qu'ils avaient besoin du même SOA pour percevoir le stimulus comme contenant deux flux. Ces résultats suggèrent que la maturation de ce processus se fait avant l'âge de 10 ans. En variant la fréquence des sons purs, Sussman, Wong, Horbath, Winkler et Wang (2007) ont d'ailleurs mis en évidence une amélioration des compétences de ségrégation avec l'âge sur des enfants de 5 à 11 ans. Il serait donc possible que le développement de cette compétence se fasse durant l'enfance pour atteindre des performances équivalentes à celles des adultes vers 10-11 ans. Toutefois, alors que Lallier et al. (2009) varient le SOA pour faire varier le percept, Sussman et al. (2007) varient la fréquence des sons purs. Il est alors possible que la capacité à effectuer une ségrégation de flux soit sous-tendue par des processus

différents dans ces deux études, et n'aient donc pas la même trajectoire développementale.

3.2.6. Test de reconnaissance de pattern auditif

Dans les tests de reconnaissance de pattern auditif en termes de fréquence et de durée, les participants entendent des séries de trois stimuli différant entre eux en terme de fréquence ou de durée (les paradigmes sont détaillés p89). Lorsque les sons purs varient en fréquence, les résultats montrent une évolution assez linéaire entre 8 et 11 ans âge auquel les compétences sont équivalentes à celles des adultes à 75% de bonnes réponses (Musiek, 2002). Lorsque les sons purs varient en terme de durée, les résultats sont assez similaires, avec un âge de maturation aux alentours de 10 ans (Brooke Shinn, 2014). D'autres auteurs toutefois ne trouvent pas d'effet de développement significatif, sur aucune de ces deux dimensions (Neijenhuis, Snik, Priester, van Kordenoordt, & van den Broek, 2002).

Ces différences de résultats d'une étude à l'autre, pourtant censées mesurer les mêmes processus, mettent en évidence l'importance de l'impact des facteurs non-sensoriels, auxquels nous allons à présent nous intéresser.

3.3. Les facteurs non-sensoriels

La différenciation entre le développement des fonctions étudiées et celui des fonctions cognitives de haut niveau comme l'attention, est un problème inhérent à toute étude développementale (Moore et al., 2011). Chez les nourrissons, il a été estimé que le taux d'inattention devait avoisiner les 30 % (Bargones, Werner, & Marean, 1995) et varier entre 0 et 25% chez les enfants (P. Allen & Wightman, 1994). Bien sûr, l'inattention n'est pas la seule explication aux différences de performances entre enfants et adultes. En effet, comme nous l'avons vu dans ce chapitre (cf. p76) il existe de nombreuses évidences que le développement des compétences auditives centrales se poursuit durant l'enfance et l'adolescence indépendamment de l'attention.

Les capacités d'attention réduites chez les enfants comparativement à une population adulte posent aussi le problème d'une plus grande sensibilité aux contraintes données par la tâche. Ainsi, Sutcliffe et Bishop (2005) ont mis en évidence que les enfants de 6-7 ans bénéficiaient particulièrement des indices pouvant attirer leur attention lorsqu'ils devaient effectuer la tâche. En effet, ce groupe d'enfant obtenaient des seuils de discrimination plus bas lorsque le paradigme attirait leur attention sur le son pur cible. Par exemple, lorsque le son pur cible est toujours présenté après un autre son pur, les performances des participants s'améliorent. Les auteurs postulent que le premier son pur agit comme un indice pour indiquer à l'enfant de prêter attention au son pur suivant. Cet effet n'était pas présent chez les enfants plus âgés (8-9 ans) et les adultes. Ceci suggère donc que les problèmes d'inattention peuvent être minorés ou majorés

par le paradigme utilisé. Toutefois, d'autres études (cf. partie sur l'*Anchor Theory*, p102) ont mis en évidence que le paradigme utilisé influençait non seulement les performances des enfants mais aussi celles des adultes (Ahissar, 2007; Ahissar, Lubin, Putter-Katz, & Banai, 2006; Banai & Yifat, 2011; Banai & Yuvak-Weiss, 2013). Il semble ainsi que le paradigme utilisé impacte les performances des participants quel que soit leur âge.

Une façon de contrôler ce problème méthodologique sera d'utiliser le même paradigme sur toutes les populations, et à travers le plus de tests possibles, afin de les rendre comparables. Dans l'Axe 2, l'Etude 3 et l'Etude 4 utilisent donc des paradigmes indiquant la présence du stimulus cible par la présentation systématique d'un premier son pur (i.e., paradigme AX) sur plusieurs populations : enfants sains, adultes sains et adultes dyslexiques.

4. Les Troubles du Traitement Auditif (TTA)

Cette thèse s'intéresse à l'impact de la dégradation du signal de parole sur le langage. Le signal de parole peut être dégradé par la présence de bruit (cf. p57) ou par un mauvais traitement de celui-ci, comme dans le cas des Troubles du Traitement Auditif (TTA). Les TTA font référence à un déficit au niveau du traitement neuronal des stimuli auditifs. C'est-à-dire que ces troubles ne peuvent pas être dus à des dysfonctionnements au niveau des traitements auditifs périphériques. Ils ne doivent pas non plus résulter de dysfonctionnements de fonctions cognitives de plus haut niveau (mais voir Moore, Ferguson, Edmondson-Jones, Ratib, & Riley, 2010). La présence de TTA peut toutefois coïncider avec des troubles du langage écrit et / ou oral (Dawes, Bishop, Sirimanna, & Bamiau, 2008; Ferguson, Hall, Riley, & Moore, 2011; Moore, 2006; Moore, Rosen, Bamiau, Campbell, & Sirimanna, 2013; M. Sharma, Purdy, & Kelly, 2009). Les TTA sont également parfois considérés comme à l'origine des troubles langagiers (cf. p98). Dans l'Etude 4, la BECAC (Donnadieu et al., 2014) a été adaptée à une population adulte afin de s'intéresser à la prévalence des TTA dans un échantillon de dyslexiques adultes et aux liens entre ces troubles et la dyslexie. Cette dernière partie présente donc les définitions des TTA sur lesquelles s'est appuyé le développement de la BECAC, puis les outils diagnostics disponibles en France et enfin la BECAC.

4.1. Définitions

Les TTA ne correspondent pas réellement à un déficit défini mais plutôt à un ensemble de déficits fonctionnels. Il en existe plusieurs définitions et à ce jour, aucun *golden standard* n'a été trouvé pour les définir et les diagnostiquer (Cacace & McFarland, 2005). Rosen (2005), citant Churchill, estime d'ailleurs que les TTA peuvent être considérés comme « *a riddle wrapped in a mystery inside an enigma* ». Des définitions

proposées dans la littérature, nous présentons uniquement celles à partir desquelles la BECAC a été développée (mais voir : Cacace & McFarland, 2005; Moore, 2011; Rosen, 2005 pour d'autres définitions et critiques). Trois organismes principaux ont tenté de poser une définition des TTA : L'*American Speech-Language Hearing Association* (ASHA, 2005), L'*American Academy of Audiology* (AAA, 2010) et la *British Society of Audiology* (BSA, 2011). Nous allons voir les définitions de chacun de ces organismes, présentant des approches plus ou moins centrées sur les aspects audiolinguistiques (ASHA et AAA) ou sur les liens avec les troubles cognitifs (BSA). Tout d'abord, l'ASHA et l'AAA définissent les TTA comme « des difficultés dans le traitement de l'information auditive dans le système nerveux central, se reflétant par de mauvaises performances dans une ou plusieurs des habiletés suivantes : localisation et latéralisation des sons, discrimination auditive, reconnaissance de pattern auditif, aspects temporels de l'audition, performances auditives dans des signaux acoustiques compétitifs et les performances auditives avec des signaux acoustiques dégradés » [notre traduction]. D'après l'ASHA (2005) ces difficultés se traduiraient au niveau comportemental, par des difficultés à comprendre la parole dans des environnements bruyants, des réponses inappropriées, de fréquentes demandes de répétitions, des difficultés à prêter attention, de mauvaises compétences musicales.

La BSA (2011) quant à elle, définit les TTA comme des troubles caractérisés par une mauvaise perception à la fois des sons parolistes et non parolistes. Cette définition considère que ces troubles ont une origine neuronale, qu'ils ont un impact sur la vie quotidienne principalement du fait d'une faible capacité à écouter et donc à répondre pertinemment aux sons, enfin que les TTA sont un ensemble de symptômes qui vont fréquemment être co-occurents d'autres troubles développementaux.

La BECAC (Donnadieu et al., 2014) a donc été pensée pour tester l'ensemble de ces compétences chez les enfants. En effet, les TTA restent difficiles à diagnostiquer, particulièrement parce qu'ils sont, comme mentionnés dans la BSA souvent associés à d'autres troubles comme des troubles du spectre autistique, des troubles de l'attention ou encore, et c'est ce sera l'un des points d'intérêt de ce travail, des troubles du langage. Du fait de l'importante co-occurrence (voir même du lien de cause à effets) entre TTA et troubles du langage, la BECAC est une batterie entièrement non verbale.

4.2. Diagnostic

4.2.1. Critères diagnostics recommandés par l'ASHA (2005) et l'AAA (2010)

Parmi les 3 organisations s'étant intéressées à définir les TTA, seules l'ASHA et l'AAA ont proposé des critères diagnostics. Ainsi, l'ASHA suggère que l'on peut parler de TTA lorsqu'il a été démontré qu'il existe un déficit du traitement des stimuli auditifs, et que cette observation ne peut pas être due à un déficit de plus haut niveau (i.e., problème

attentionnel, mnésique, compréhension de consigne...). Le diagnostic peut également être posé si, à un ensemble de test évaluant les compétences présentées dans la définition le patient montre des scores inférieurs à 2 écart-types de la moyenne sur 2 des tests, ou inférieur à 3 écart-types de la moyenne sur 1 des tests, ou encore un score inférieur à 2 écart-types en dessous de la moyenne à 1 test, si c'est accompagné de difficultés comportementales manifestes.

Enfin, l'AAA considère que des TTA peuvent être diagnostiqués dès la présence d'un score inférieur à -2 écart-types de la moyenne, pour au moins une oreille sur au moins deux tests comportementaux portant sur les traitements auditifs.

4.2.2. Le BAC

En France, il n'existe pas de ligne directrice de critères diagnostics pour les TTA et à notre connaissance un seul test est mis à la disposition des praticiens pour leur diagnostic, le Bilan Auditif Central (BAC ; Demanez, Dony-Closon, Lhonneux-Ledoux, & Demanez, 2003). Cette batterie propose une série de 4 tests parmi les catégories les plus classiques, visant à évaluer différents aspects du traitement auditif central (Baran, 2014).

Une première partie de la batterie vise à tester la qualité du décodage phonémique lorsque la redondance du signal parolier (i.e., redondance extrinsèque) est dégradée par l'ajout de bruit grâce à une adaptation du test d'intégration auditive de Lafon (1964). Plusieurs séries de mots monosyllabiques sont présentées successivement dans le silence puis dans du bruit blanc avec un SNR de 0 dB. C'est la reconnaissance du phonème caractéristique du mot par le sujet qui va déterminer la réponse comme correcte ou incorrecte.

Une deuxième partie vise à tester grâce à l'écoute dichotique l'intégration binaurale et la séparation binaurale. Deux mots sont présentés simultanément (un dans chaque oreille). Les participants doivent soit répéter les deux mots (intégration binaurale) soit un des deux mots en fonction de l'oreille préalablement désignée par l'expérimentateur (séparation binaurale).

Ce bilan teste ensuite l'interaction binaurale, grâce au *Masking Level Difference* test (Licklider, 1948). Pour cela, des mots dissyllabiques sont présentés de façon binaurale dans du bruit blanc. Dans une seconde partie du test de démasquage, un déphasage de 180° est introduit entre les mots présentés par les écouteurs droit et gauche. Ce déphasage du signal cible crée un effet de démasquage spatial et réduit ainsi l'effet masquant du bruit blanc.

Enfin, le BAC propose un test dans lequel sont présentées des séquences de 3 tons pouvant varier en fréquence (Haut ou Bas) ou en durée (Long ou Court). Les participants doivent reproduire ces séries soit en fredonnant, soit en verbalisant (e.g., « Haut Haut Bas »). Ces deux tests, sont issues des *Frequency Pattern Test* et *Duration Pattern Test* développés par Musiek et Pinheiro (1987) et Musiek (1994) respectivement.

Le bilan proposé par Demanez et collègues (Demanez & Demanez, 2003; Demanez et al., 2003) teste donc la résistance à la faible redondance extrinsèque (test de Lafon, 1964), les capacités d'écoute dichotique, puis l'intégration binaurale et enfin la reconnaissance de patterns auditifs en fréquence et en durée. Les auteurs soulignent toutefois eux-mêmes que le BAC ne permet pas de vérifier l'ensemble des traitements auditifs centraux, mais le problème majeur de ce bilan provient de l'utilisation de signaux de paroles (dans 3 des 4 tests). En effet, le BAC teste un trouble non verbal en utilisant du matériel principalement verbal. Dès lors, de mauvais résultats aux tests peuvent ne pas provenir d'un déficit auditif central mais simplement d'un problème langagier. Etant donné que des TTA sont fréquemment observés chez les personnes souffrant de Troubles Spécifiques du Développement du Langage, (Bishop, Hardiman, & Barry, 2012; Bishop & McArthur, 2005; Ramus et al., 2003), cette confusion de troubles verbaux et non verbaux peuvent mener à un faux diagnostic. Ce bilan, associé à un contrôle orthophonique et neuropsychologique est à ce jour, le seul outil à disposition des praticiens francophones pour le diagnostic des TTA (Masquelier, 2011). C'est pour répondre à ce besoin que Donnadieu et al., (2014) ont développé la BECAC, permettant de tester les compétences auditives centrales en utilisant exclusivement du matériel non verbal. Cette batterie a été présentée à des enfants de 6 à 11 ans (Etude 3a et 3b) puis, après adaptation à une population adulte, à un groupe de jeunes adultes de 18 – 22 ans (Etude 3b) et à un groupe d'adultes dyslexique et de contrôles sains, appariés en âge, intelligence non verbale et latéralité (Etude 4).

4.3. La Batterie d'Evaluation des Compétences Auditives Centrales

La BECAC est composée de 6 tests principaux eux-mêmes comprenant un ou plusieurs sous-tests. Les sous-tests présentent l'avantage de tester différentes dimensions d'un même traitement auditif, en utilisant le même paradigme, afin de pouvoir limiter l'impact de celui-ci.

4.3.1. La latéralisation

Un son pur de 500 Hz et de 250 ms est présenté aux participants, avec une Différence Interaurale d'Intensité (*Interaural Level Difference* ; ILD) variable permettant de créer l'illusion d'une latéralisation. L'ILD pouvait être nulle (condition binaurale) ou varier de -5 dB à -25 dB par pas de 2 dB. C'est-à-dire que les participants peuvent entendre le son pur dans les deux oreilles, à gauche (si le ton est présenté moins fort dans l'oreille

droite que dans la gauche) ou à droite. La phase test est constituée de 66 items sans feed-back, incluant 22 essais dans lesquels le son était présenté à une égale intensité dans les 2 oreilles, 22 essais dans lesquels le son présenté dans l'oreille droite a une intensité moindre et 22 essais dans lesquels le son présenté dans l'oreille gauche a une intensité moindre. L'ordre de présentation des essais est aléatoire. Les participants doivent indiquer à chaque essai dans quelle oreille ils ont entendu le son (à gauche, à droite ou dans les deux).

4.3.2. La discrimination auditive

Le test de discrimination auditive est présenté sous forme d'un paradigme AX. Un premier son pur (A) de 75 ms et 520 Hz est présenté à une intensité de 70 dB. Un second son pur (X) est présenté 400 ms plus tard et peut différer de A en termes de fréquence, durée ou intensité selon le sous-test. Les participants doivent indiquer si les sons purs sont identiques ou différents sur 48 essais (50% sont identiques).

La fréquence. La fréquence de X varie entre 520 Hz (comme A) et 527 Hz, 535 Hz, 546 Hz, 562 Hz, 583 Hz ou 609 Hz.

La durée. La durée de X varie en fonction des essais, de 75 ms (même durée que A) à 27 ms par pas de 8 ms (67 ms, 59 ms, 51 ms, 43 ms, 35 ms, ou 27 ms).

L'intensité. L'intensité relative de X par rapport à A varie jusqu'à 15 dB moins fort, par pas de 2,5 dB (-2,5 dB, -5 dB, -7,5 dB, -10 dB, 12,5 dB, -15 dB).

4.3.3. L'identification et la reconnaissance auditive

Ces deux tests sont les seuls de la batterie à ne pas utiliser de sons purs, mais des bruits présents quotidiennement dans l'environnement. Chaque test comprend 28 essais. Ils visent à évaluer la capacité des participants à traiter un son représentant un objet et à l'apparier soit à une image (donc en passant obligatoirement par le traitement sémantique de l'objet) soit à une autre séquence sonore présentant un bruit caractéristique de l'objet (donc en ne passant pas obligatoirement pas le traitement sémantique de l'objet).

L'identification auditive. Un bruit est présenté (e.g., bruit d'aspirateur) puis le participant doit désigner l'image correspondant au bruit entendu parmi les trois présentées. Il s'agit donc d'apparier le son cible à l'image de l'objet sonore correspondant. L'image cible est présentée avec un distracteur sémantiquement lié (e.g., balai) et acoustiquement lié (e.g., avion).

La reconnaissance auditive. Dans ce test, l'expérimentateur présente un bruit au participant, puis 3 autres bruits. Il s'agit alors d'indiquer lequel des 3 bruits présentés est produit par le même objet ou être vivant que le bruit précédemment entendu alors même qu'ils diffèrent acoustiquement. Encore une fois, le bruit cible est présenté avec

un distracteur acoustique et un distracteur sémantique. Par exemple, si le son à comparer est un son de voiture, on présentera comme bruits de comparaison, un son de voiture différent, un son sémantiquement lié (e.g., un train) et un son acoustiquement lié (e.g., une tondeuse à gazon).

4.3.4. Pattern auditifs

Ce test vise à évaluer la capacité des participants à reconnaître, retenir et reproduire des séquences de 3 stimuli acoustiques, ils sont adaptés de l'*Auditory Pattern Test* (Pinheiro & Musiek, 1985). Dans un premier sous-test les stimuli varient en fréquence, et dans un second en durée. Les participants effectuaient un entraînement comprenant 12 essais (6 avec feedback, 6 sans feedback). Les participants doivent répéter le pattern entendu après chaque essai, en le chantant ou en le fredonnant.

Fréquence. Dans ce sous test, les sons purs de 500 ms et séparés par 300 ms varient en fréquence, Basse (B = 880 Hz) ou Haute (H = 1222 Hz). Six patterns distincts sont présentés, BBH, BHB, BHH, HBB, HBH, HHB. Les participants reproduisent le pattern entendu après chaque essai.

Durée. Trois sons purs de 1000 Hz sont présentés, séparés par un ISI de 300 ms. Ces sons purs peuvent être Longs (L = 500 ms) ou Courts (C = 250 ms). Les séquences contiennent deux tons de même durée, il y a donc 6 patterns différents : CCL, CLC, CLL, LCC, LCL, LLC.

4.3.5. Ségrégation des flux

Ce test a pour but d'évaluer la capacité des participants à percevoir des stimuli comme provenant d'une ou deux sources. Les séquences sonores de 5 sec sont composées de sons purs de 40 ms. Les tons présentés sont alternativement de haute fréquence, i.e., aigus (1000 Hz), et de basse fréquence, i.e., grave (400 Hz). Le SOA est manipulé en fonction de la réponse des participants. Les participants doivent indiquer s'ils ont perçus un ou deux flux auditifs en appuyant sur le bouton correspondant à leur réponse. Le premier SOA utilisé est de 260 ms afin que la réponse soit non ambiguë et corresponde à *un flux*. Le seuil de ségrégation individuel correspondant à le SOA à partir duquel le sujet perçoit deux flux distincts est établi selon une méthode adaptative *one-up, one-down*. Tant que la réponse du sujet est *un flux*, le SOA diminue, en revanche, il augmente lorsque la réponse est *deux flux*. Trente essais sont présentés. Pour les premiers essais, non ambigus, le SOA diminue de 40 ms après chaque réponse. Après le premier changement de réponse, le SOA augmente de 20 ms, puis de 10 ms lors du 2^{ème} changement et de 5 ms à partir du 3^{ème}. Le seuil de ségrégation est ainsi calculé en moyennant les SOA sur les 10 dernières séquences présentées (séquences 21 à 30). Cette procédure est identique à celle utilisée par Lallier et al. (2009).

4.3.6. Le masquage central

Ce test a pour objectif d'évaluer l'effet délétère d'un bruit blanc (i.e., le masqueur) sur la détection d'un son pur présenté dans l'oreille contralatérale. Trois séquences de bruit blanc de 200 ms (ISI = 200 ms) étaient présentées dans chaque essai dans une oreille, gauche ou droite. Dans 50% des essais, 3 sons purs de 500 Hz et 200 ms (ISI = 200 ms) étaient présentés simultanément dans l'oreille contralatérale. L'intensité des sons purs varie d'un essai à l'autre (entre -30 dB et -58 dB relativement à l'intensité du bruit blanc). Le test comprend au total 60 essais, dont 30 essais où les sons purs sont présentés (15 présentations dans chaque oreille).

Dans un second sous-test, les participants sont soumis à la même procédure, simplement, le bruit blanc n'est pas présenté. Seuls les sons purs sont présents, dans 50% des essais. L'effet du masque est évalué en comparant les deux sous-tests.

5. Résumé du Chapitre IV

L'objectif de ce chapitre était de s'intéresser au fonctionnement du système auditif. Tout d'abord, nous avons brièvement présenté le système auditif périphérique puis nous nous sommes intéressés de façon plus approfondie au système auditif central. Une revue de la littérature sur le développement des compétences auditives centrales a révélé que les données disponibles étaient relativement peu nombreuses. La plupart d'entre elles mettaient en évidence que le développement des compétences auditives centrales se poursuivait durant l'enfance jusqu'au début de l'adolescence. Peu d'études testant différentes compétences au sein d'une même population ont été trouvées. Ce chapitre s'intéresse ensuite au cas particulier des Troubles du Traitement Auditif. Après une brève présentation de leur définition et du Bilan Auditif Central (BAC ; Demanez et al., 2003) seul outil disponible en France pour les diagnostiquer nous présentons la Batterie d'Evaluation des Compétences Auditives Centrales (Donnadieu et al., 2014) qui permet d'évaluer les traitements auditifs centraux sans utiliser de matériel verbal, contrairement au BAC.

Chapitre V

La dyslexie

C'est en 1887 qu'apparaît le terme de dyslexie pour qualifier l'observation par Rudolf Berlin, un ophtalmologue allemand, d'un patient souffrant d'une incapacité acquise à déchiffrer les symboles écrits et

qui proviendrait de lésions cérébrales (Elliott & Grigorenko, 2014). Selon Berlin, une incapacité totale à lire résulterait en une alexie, alors qu'une capacité seulement minimale résulterait en une dyslexie. Presque une décennie plus tard, le premier cas de dyslexie développementale a été rapporté par Morgan (1896). Il décrit le cas d'un jeune garçon de 14 ans, présentant une intelligence normale, de bonnes habiletés en arithmétique et ayant reçu une éducation depuis l'âge de 7 ans et pourtant incapable de déchiffrer une phrase complète. Le nombre de cas rapportés a ensuite beaucoup augmenté au cours du XXe siècle avec la généralisation de l'accès à l'éducation. Malgré un intérêt grandissant aussi bien scientifique que sociétal pour la dyslexie, cette pathologie reste un sujet de débat à tous les niveaux.

Au sein de ce chapitre, nous présentons en premier lieu les principales définitions proposées et les questions qu'elles soulèvent. Nous nous intéressons ensuite aux différents types de dyslexie. La deuxième moitié du chapitre porte sur les théories explicatives de la dyslexie proposées par la littérature. Étant donné leur nombre important (hypothèse attentionnelle : Bosse, Tainturier, & Valdois, 2007; hypothèse cérébelleuse : Fawcett & Nicolson, 1999; hypothèse magnocellulaire : Stein & Walsh, 1997), nous ne présentons que le déficit phonologique et les théories tentant de l'expliquer par un déficit du traitement auditif puisque c'est celles-ci auxquelles le travail expérimental s'est intéressé (pour des revues sur les théories explicatives de la dyslexie et leurs limites voir Demonet, Taylor, & Chaix, 2004; Peterson & Pennington, 2012; Ramus, 2001, 2003).

1. Définition et diagnostic

La dyslexie est un trouble développemental qui touche environ 7% de la population (Peterson & Pennington, 2012; Snowling, 2000) et davantage les garçons que les filles (Rutter et al., 2004). Sa définition a évolué avec le temps (Fletcher, 2009) et sa définition la plus usitée est qu'elle correspond à « (...) un trouble spécifique des apprentissage dont l'origine est neurobiologique. Elle est caractérisée par des difficultés à reconnaître rapidement et précisément les mots et par des compétences déficitaires en déchiffrement et orthographe. Ces difficultés résultent typiquement d'un déficit phonologique qui est souvent inattendu comparativement aux autres compétences cognitives. Elle peut avoir comme conséquences secondaires des problèmes de compréhension de lecture et la réduction de l'expérience en lecture peut impacter de manière négative le vocabulaire et les connaissances générales » [notre traduction] Lyon, Shaywitz, et Shaywitz (2003).

Cette définition ne permet pas de définir de critères diagnostics. La difficulté principale soulevée par la pose du diagnostic de dyslexie est de mettre en évidence que le retard de lecture provient d'un dysfonctionnement neurobiologique spécifique à l'apprentissage de la lecture. C'est-à-dire qu'il ne doit pas être la conséquence d'un autre trouble plus général touchant d'autres fonctions cognitives. Afin de répondre à ces exigences les critères diagnostics habituellement employés sont les suivants : un retard dans l'apprentissage de la lecture (en France 18 mois) en dépit d'une intelligence non-verbale normale, d'un milieu socio-culturel suffisamment stimulant, d'un accès normal à l'éducation et de l'absence de troubles psychiatriques et perceptifs majeurs (visuel ou auditif). Cette définition par exclusion présente l'avantage de distinguer une population très spécifique de dyslexiques, dont les difficultés sont sous-tendues par des causes neurobiologiques, d'une population communément appelée de mauvais lecteurs représentant simplement un extrême au sein de la population saine. Si la littérature scientifique accepte majoritairement cette définition puisqu'elle permet d'obtenir un groupe *homogène* qui ne se distingue de la population contrôle que par le trouble de lecture, sa validité, particulièrement dans les domaines plus appliqués (i.e., enseignement, rééducation) est très controversée (Elliott & Grigorenko, 2014).

En effet, si l'appartenance d'un enfant à un milieu socio-culturel considéré comme défavorisé et peu stimulant empêche de certifier que le trouble de lecture a une origine neurobiologique, en aucun cas il ne permet de certifier que ce n'est pas le cas. D'autant plus qu'il a récemment été mis en évidence, sur la population française qu'au même âge, des enfants issus de milieux défavorisés ont presque 10 fois plus de risques de présenter un retard de lecture (Fluss et al., 2009). De même, la définition posée par Lyon et al., (2003) puisqu'elle exclut les enfants présentant d'autres troubles cognitifs, ne tient pas compte de la forte comorbidité entre notamment, dyslexie et TDAH

(Troubles De l'Attention/Hyperactivité ; McGrath et al., 2011) ou encore dyslexie et dysphasie (Bishop & Snowling, 2004; Catts, Adlof, Hogan, & Weismer, 2005).

La pertinence de la définition et du diagnostic de la dyslexie dépend donc de l'objectif du diagnostic. Dans un objectif de recherche, la définition par exclusion nous paraît plus appropriée puisqu'elle permet de sélectionner une population dont la seule distinction d'avec la population contrôle est de présenter des troubles d'apprentissage de la lecture. Dans un objectif de rééducation ou d'éducation en revanche, cette définition pose le problème d'une trop grande exclusivité, et risque de retirer à des enfants en difficulté le droit à une éducation spécifique plus adaptée sur la base par exemple de leur milieu socio-culturel considéré comme pas assez stimulant.

2. Classification des dyslexies

2.1. Le modèle dual route

Trois types de dyslexies sont théoriquement distingués, la dyslexie de surface, la dyslexie phonologique et enfin la dyslexie mixte. Elles se différencient par les patterns d'erreurs commises par les patients et expliqués par le modèle *Dual route* proposé par Castles et Coltheart (1993). Ce modèle suggère que deux voies peuvent être empruntées pour la reconnaissance d'un mot écrit (cf. Figure 15). Tout d'abord, la procédure phonologique (ou voie d'assemblage), qui correspond à une conversion graphème-phonème systématique. Chaque graphème (ou groupe de graphème) est reconnu et apparié au phonème correspondant, pour former la forme orale du mot. Cette voie est obligatoirement empruntée pour déchiffrer les nouveaux mots ainsi que les pseudo-mots.

L'importance de la procédure phonologique dans l'apprentissage de la lecture dépend de l'opacité de l'orthographe de la langue considérée. Ainsi, l'anglais est considéré comme ayant l'orthographe la plus opaque des langues utilisant un système alphabétique. En effet, 40 phonèmes peuvent être écrits par 1120 graphèmes. Ainsi, un même graphème peut correspondre à plusieurs phonèmes (e.g., *love* vs. *move*). L'espagnol et l'italien en revanche, ont des orthographes transparentes, en italien par exemple, seuls 33 graphèmes sont nécessaires à la retranscription de 25 phonèmes. Le français est considéré comme moyennement opaque, 170 graphèmes sont nécessaires pour retranscrire 36 phonèmes. Ainsi, d'après le modèle *dual route* la procédure phonologique sera plus utilisée dans l'apprentissage de la lecture en italien ou en espagnol qu'en anglais.

La procédure phonologique à elle seule ne peut toutefois pas expliquer certains effets des plus robustes observés durant la reconnaissance de mots : les effets de lexicalité (i.e., reconnaissance plus rapide des mots que des pseudo-mots) et de fréquence (i.e.,

reconnaissance plus rapide des mots fréquents que rares). Elle ne permet pas non plus d'expliquer la reconnaissance des mots irréguliers comme *agenda*. En effet, la lecture de mots irréguliers en utilisant la procédure phonologique engendrerait des prononciations erronées.

Le modèle postule donc qu'il existe une seconde procédure permettant de reconnaître les mots écrits, la procédure lexicale (ou voie d'adressage). La forme globale du mot écrit serait reconnue dans le lexique orthographique et appariée directement à la fois à sa forme phonologique au sein du lexique phonologique et à son sens au sein du système sémantique. L'existence de cette procédure lexicale permet d'expliquer l'effet de lexicalité puisque les pseudo-mots n'ont par définition jamais été rencontrés au préalable et leur forme globale visuelle ne peut pas être reconnue. Les mots sont donc reconnus via la procédure lexicale, plus rapide que la procédure phonologique. Un déficit dans l'une ou l'autre voie permet ainsi de distinguer différents types de dyslexie.

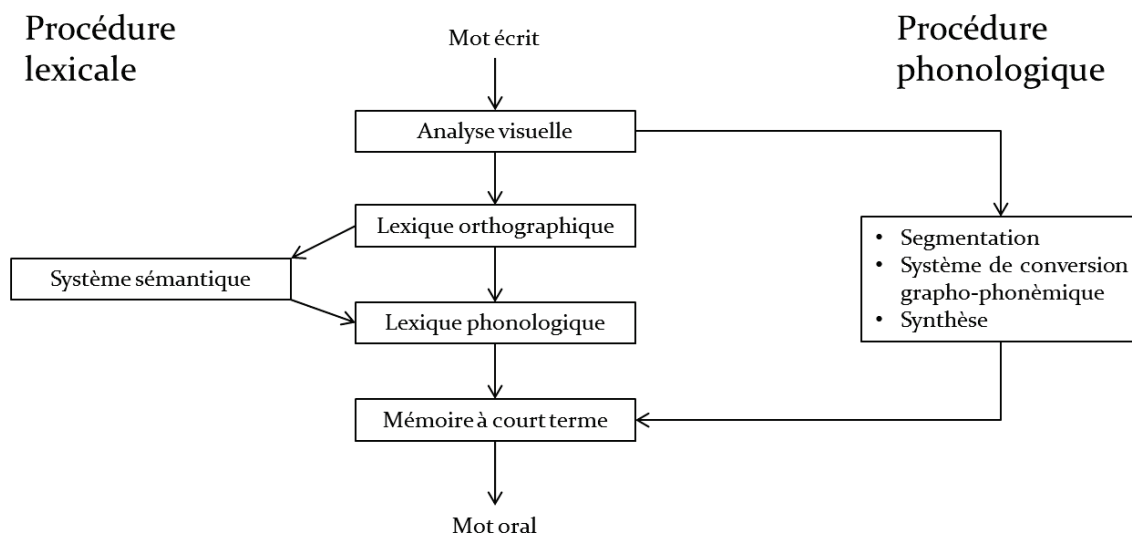


Figure 15

Tiré de Jacquier-Roux, Valdois, Zorman, Lequette, et Pouget (2005). Représentation graphique du modèle double voie.

2.2. Dyslexie phonologique et dyslexie de surface

La dyslexie phonologique se caractérise par un déficit dans la procédure phonologique, ou voie d'assemblage. Les personnes présentant une dyslexie phonologique vont donc présenter une difficulté accrue à déchiffrer les nouveaux mots et les pseudo-mots (Castles & Coltheart, 1993; Joanisse, Manis, Keating, & Seidenberg, 2000; Sprenger-Charolles, Cole, Lacert, & Serniclaes, 2000; Stanovich, Siegel, & Gottardo, 1997).

La dyslexie de surface correspond davantage à un délai dans l'apprentissage qu'à une déviance (Joanisse et al., 2000; Stanovich et al., 1997). Elle est caractérisée par une difficulté dans la procédure lexicale, se manifestant par des performances inférieures à celles des sujets contrôles de même âge en lecture de mots irréguliers (e.g., *femme* va être lu /fem/ et non /fam/). En revanche, lorsque leurs performances sont comparées à des contrôles appariés en termes de niveau de lecture (i.e., plus jeunes) les performances du groupe dyslexique en lecture de mots irréguliers ne diffèrent pas de celles du groupe contrôle (Joanisse et al., 2000).

Le modèle *dual route* permet de répondre à une demande de rééducation et de mieux appréhender les troubles des patients dyslexiques. En termes de recherche, particulièrement chez l'adulte, il est moins pertinent. En effet, les adultes dyslexiques ayant pour la plupart bénéficié de rééducation, ces différences si elles étaient observables durant l'enfance ne le sont plus. De plus, si la dyslexie de surface correspond à un délai d'apprentissage, elle ne devrait plus être observable chez les sujets adultes. Les dyslexiques recrutés dans l'Etude 4 ne feront donc pas l'objet d'une telle classification.

3. Le déficit phonologique

La théorie explicative de la dyslexie la plus prégnante ces dernières décennies est la théorie reposant sur l'idée d'un déficit phonologique. Ce déficit phonologique se traduirait par des représentations phonologiques instables, rendant difficile l'association graphème-phonème et donc l'apprentissage de la lecture. Nous présentons les manifestations comportementales de ce trouble, qui se retrouvent chez tous les dyslexiques (Ramus, 2003; White et al., 2006).

3.1. La conscience phonologique

Il existe un consensus selon lequel les personnes dyslexiques présentent des difficultés dans les tâches qui ont trait au traitement phonologique. Nous avons vu dans le chapitre 1, que le flux de parole se compose de plusieurs unités. Les unités de sens : les phrases, mots et morphèmes et les unités de forme : syllabes et phonèmes, ces dernières constituant la plus petite unité du signal de parole. Lors de l'apprentissage de la lecture, les enfants apprennent d'abord à distinguer les syllabes dans le flux de parole puis les phonèmes. La distinction des phonèmes est indispensable puisqu'elle va permettre de former des représentations phonologiques qui seront ensuite appariées à leur correspondance orthographique, permettant la conversion graphème/phonème. C'est cette capacité à apparier un signe graphique au son de parole lui correspondant qui permet l'apprentissage de la lecture (Castles & Coltheart, 1993). L'implication des compétences phonologiques dans la lecture a largement été démontrée et elles sont

apparues comme un des meilleurs prédicteurs des compétences en lecture (Duff, Hayiou-Thomas, & Hulme, 2012; Stanovich, 1988).

Le déficit phonologique a été mis en évidence sur l'ensemble de la population dyslexique (Ramus et al., 2003; Sprenger-Charolles et al., 2000; mais voir Bosse, Tainturier & Valdois, 2007 pour une hypothèse attentionnelle). Il se manifeste traditionnellement (a) par un déficit en conscience phonologique, qui se manifeste par une difficulté à accéder consciemment aux représentations phonémiques et à les manipuler, (b) par un déficit en mémoire à court-terme verbale et (c) par une difficulté à récupérer le code phonologique faisant référence à un objet en mémoire à long-terme (mais voir Wolf & Bowers, 2000 pour l'hypothèse du double déficit). Ces compétences vont être testées respectivement avec des tests comme (a) la suppression de phonème ou les contrepèteries, (b) la mémoire de travail (testée avec un empan digital ou la répétition de pseudo-mots et enfin (c) la dénomination rapide (RAN ; *Rapid Automated Naming* ; Denckla & Rudel, 1976). Ces difficultés vont persister à l'âge adulte, même après rééducation (J. Martin et al., 2010) car l'exécution de ces tâches nécessite l'intégrité des représentations phonologiques. L'observation des difficultés des dyslexiques dans ces tâches suggère ainsi que leurs représentations phonologiques sont dégradées ou instables (mais voir aussi Boets et al., 2013; Ramus & Szenkovits, 2008 pour une hypothèse alternative expliquée dans la partie sur l'Anchor Theory p100). La conscience phonologique est toutefois également dépendante du niveau de lecture, puisque l'apprentissage de la conversion graphème-phonème va faciliter la manipulation de ces derniers. Pourtant, le faible niveau en lecture des dyslexiques ne permet pas d'expliquer leurs mauvaises performances dans les tâches impliquant la conscience phonologique. En effet, même comparés à des contrôles appariés en niveau de lecture, les participants dyslexiques présentent des scores déficitaires à ces tâches. La conscience phonologique est également un prédicteur de l'apprentissage de la lecture chez les enfants pré-lecteurs (Sprenger-Charolles et al., 2000).

3.2. La perception allophonique

Un autre argument en faveur d'un déficit de représentations phonologiques des dyslexiques provient des études s'intéressant à la perception allophonique vs catégorielle. Les différents phonèmes sont composés de traits phonétiques comme le lieu d'articulation, le voisement, et le mode d'articulation. Ces différents traits vont permettre de distinguer deux phonèmes. Par exemple le phonème /b/ et /p/ en français ne diffèrent que par le VOT (*Voice Onset Time*) c'est-à-dire le délai entre l'ouverture du conduit vocal et la mise en vibration des cordes vocales. Lorsque le VOT est négatif, les orateurs français vont percevoir un /b/, alors que lorsque le VOT est positif, ils vont percevoir un /p/. Le positionnement de cette frontière dépend de la langue (cf. Figure 16). La perception des phonèmes évolue en fonction de la langue dans laquelle le nourrisson est élevé. A la naissance, celui-ci a une perception

allophonique, et sera capable de distinguer deux stimuli positionnés du même côté de la frontière catégorielle. Cette perception allophonique évolue vers une perception catégorielle à force d'exposition à une langue spécifique. Il semble pourtant que les dyslexiques n'accèdent pas ou peu à cette perception catégorielle et conservent une perception allophonique (Bogliotti, Serniclaes, Messaoud-Galusi, & Sprenger-Charolles, 2008; Serniclaes, Van Heghe, Mousty, Carre, & Sprenger-Charolles, 2004).

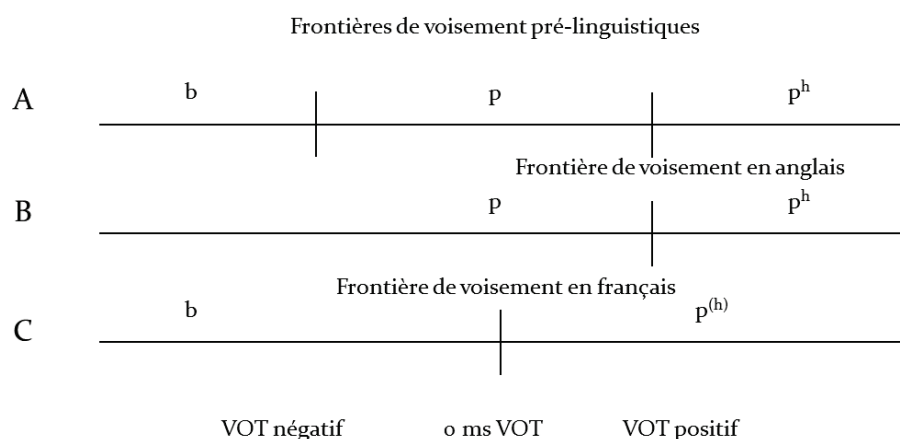


Figure 16

Adapté de Serniclaes et al. (2004). A. Illustration des frontières perceptuelles pour les enfants pré-linguistiques. Illustration (B) de la perception catégorielle en anglais et (C) en français.

Les études s'intéressant à la perception allophonique vs catégorielle étudient la perception des participants en fonction des caractéristiques acoustiques des stimuli présentés. C'est-à-dire qu'une série de stimuli est créée à partir d'une syllabe, par exemple /to/. Le VOT est ensuite manipulé, afin de créer une série de stimuli créant un continuum de /to/ vers /do/. Les participants doivent ensuite déterminer si tel ou tel stimulus est un /to/ ou un /do/, ou encore, deux stimuli sont présentés à chaque essai et les participants doivent déterminer s'ils sont identiques ou différents. Classiquement, la différence de durée de VOT entre les deux stimuli présentés n'a pas d'influence sur la perception d'une différence entre les deux stimuli. C'est le placement de l'un ou l'autre des stimuli d'un côté ou de l'autre de la frontière catégorielle qui va permettre la distinction entre les deux phonèmes. Ainsi une différence de 20 ms va permettre de distinguer les deux stimuli si chacun est d'un côté de la frontière alors qu'une différence de 20 ms ne permettra pas de les distinguer s'ils sont du même côté de la frontière (cf. Figure 17). Ce déficit des dyslexiques en perception catégorielle peut expliquer en partie la difficulté à apprendre les correspondances graphème-phonèmes. En effet, si les dyslexiques conservent les frontières pré-linguistiques (i.e., perception allophonique ; cf. Figure 16 et Figure 17), alors ils vont percevoir par exemple deux sortes de p alors que les lecteurs français n'en percevront qu'un. Il sera alors difficile d'apparier deux perceptions différentes à un même graphème.

Le déficit phonologique peut aussi être expliqué par la présence d'un trouble de bas niveau comme un déficit du traitement auditif.

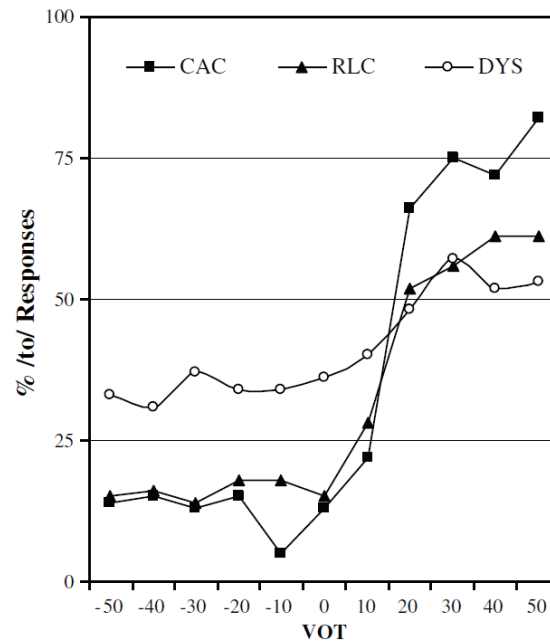


Figure 17

Tiré de Bogliotti et al., (2008). Pourcentage de réponses /to/ (vs /do/) en fonction de la durée du VOT et du groupe de participants. DYS = Dyslexique ; RLC = Contrôle âge lecture ; CAC = Contrôle âge chronologique.

4. Les théories d'un déficit du traitement auditif

L'ensemble des théories explicatives de la dyslexie postulant l'existence d'un trouble du traitement auditif de bas niveau sont compatibles avec la théorie phonologique. Elles postulent que le trouble phonologique est une conséquence d'un déficit du traitement auditif (cf. Figure 18).

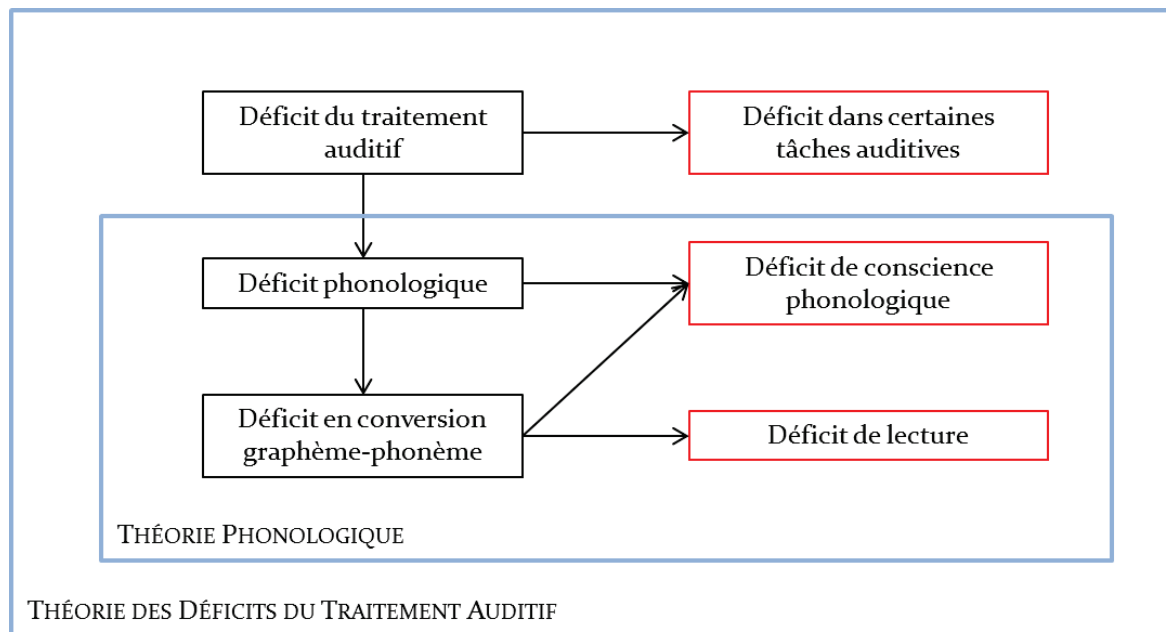


Figure 18

Adapté de Ramus (2001). Schéma représentant les processus cognitifs déficitaires d'après la théorie phonologique (encadrés en noir) et les théories d'un déficit du traitement auditif et leurs conséquences comportementales (encadrées en rouge).

4.1. La théorie d'un déficit des traitements auditifs rapides

La théorie d'un déficit des traitements auditifs rapides a été initialement proposée par Tallal (1980). Elle a en effet mis en évidence que les enfants dyslexiques présentaient un déficit aussi bien en discrimination de fréquence que dans une tâche de jugement d'ordre temporel (*Temporal Order Judgment ; TOJ*) lorsque la présentation des stimuli était rapide. Dans cet article fondateur, on présente aux enfants dyslexiques deux sons purs de 100 Hz et 305 Hz et d'une durée de 75 ms. Après une période d'entraînement, les participants devaient appuyer sur le bouton correspondant (un pour chaque fréquence) dans l'ordre dans lequel les stimuli avaient été présentés. Il est apparu que les participants dyslexiques obtenaient de moins bons scores que les participants contrôles lorsque l'ISI était court (i.e., 8 à 305 ms) réduisant le temps de traitement, alors que les performances des deux groupes étaient identiques lorsque l'ISI était long (i.e., 428 ms). La même observation a été faite pour la tâche de discrimination de fréquence. De plus, Tallal a mis en évidence que le déficit lorsque l'ISI est court est équivalent dans les deux tâches, ce qui suggérerait que ce déficit traduit une incapacité des participants dyslexiques à traiter les stimuli auditifs rapides. Ceci entraînerait alors une impossibilité de traiter les transitions rapides de formants permettant de discriminer les phonèmes (cf. Figure 19).

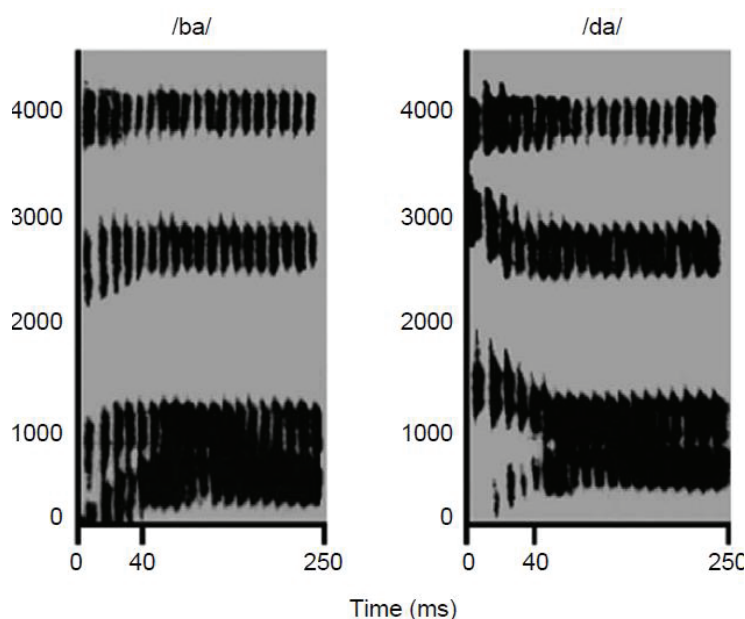


Figure 19

Tiré de Tallal et Gaab (2006) . Deux spectrogrames représentant les syllabes /ba/ et /da/. La différenciation entre les deux phonèmes /b/ et /d/, caractérisée par des transitions formantiques différentes se fait dans les premières 40 ms. La distinction entre les deux stimuli dépend donc de la capacité à traiter des changements de fréquence rapide.

Toutefois, il est apparu que les déficits en termes de traitements auditifs n'étaient pas spécifiques aux stimuli présentés de façon rapides. En effet, les dyslexiques présentent des scores moins élevés que les participants contrôles dans les tâches de discrimination de fréquence, quel que soit l'ISI utilisé (ISI = 800 ms ; Ahissar et al., 2000; ISI = 1000 ms ; Banai & Ahissar, 2006; ISI = 400 ms ; France et al., 2002; ISI = 300 ms ; McAnally & Stein, 1996). Certains auteurs trouvent même une différence entre contrôles et dyslexiques uniquement lorsque l'ISI est long (Share, Jorm, Maclean, & Matthews, 2002). Une autre critique apportée à cette théorie est que tous les dyslexiques ne présentent pas de déficit, en effet, seuls 30 à 60% des dyslexiques testés ont montré un déficit en traitement auditifs rapides (Adlard & Hazan, 1998; Ramus et al., 2003; Rosen & Manganari, 2001).

Enfin, certaines études n'ont pas retrouvé ce déficit en traitement de transition de formants (Mody, Studdert-Kennedy, & Brady, 1997; Nitttrouer, 1999; Rosen & Manganari, 2001). Toutefois ces études présentent des biais méthodologiques qui diminuent l'impact de leur absence de résultat (cf. p106 et p205).

Bien que la théorie portée par Tallal et ses collègues soit critiquée et puisse être en partie remise en cause, l'idée d'un déficit auditif à l'origine du trouble phonologique est toujours d'actualité, et d'autres théories explicatives ont été développées.

4.2. La théorie du déficit de la perception du rythme

Goswami et son équipe ont ainsi proposé une théorie s'intéressant aux différentes variations du rythme de parole. Tout d'abord, les variations rapides (entre 30 et 50 Hz) correspondant à la structure fine du signal et aux caractéristiques phonémiques, comme les transitions formantiques. Puis les variations apparaissant à un rythme de 4 à 7 Hz, correspondant au rythme syllabique, et enfin, les variations plus lentes, correspondant à l'intonation ou la prosodie avec un rythme de 1-2 Hz. La bonne perception de ces rythmes est nécessaire à la distinction des syllabes au sein du signal continu qu'est le signal de parole. En effet, la perception de la prosodie et du stress lexical est liée à la conscience phonologique, même chez les enfants pré-lecteurs (Beattie & Manis, 2014; Goodman, Libenson, & Wade-Wooley, 2010). Particulièrement, c'est la bonne perception des syllabes qui va permettre ainsi de développer la conscience phonologique. L'hypothèse soutenue par Goswami et ses collègues (Goswami et al., 2002; Leong & Goswami, 2014; Richardson, Thomson, Scott, & Goswami, 2004) est que les personnes dyslexiques présentent des difficultés dans la perception des rythmes syllabique et prosodique. Plus spécifiquement, les dyslexiques percevraient moins les modulations d'amplitude que les normo-lecteurs. Dans l'article original, Goswami et al., 2002 ont présenté une sinusoïdale de 500 Hz avec une profondeur de modulation d'amplitude de 50%, à une fréquence de 0.7 Hz. Le temps de hausse d'amplitude, variait de 300 ms (i.e., percept d'un son continu dont l'intensité varie) à 15 ms (i.e., percept correspondant à un rythme ; cf. Figure 20). Les participants, dyslexiques et contrôles devaient indiquer leur percept en fonction du temps de hausse d'amplitude. Les résultats ont révélé que les participants dyslexiques présentaient une courbe de réponse plus plate que les sujets contrôles, c'est-à-dire qu'ils étaient moins sensibles au changement de temps de hausse d'amplitude. De plus, les résultats à cette tâche étaient corrélés aux tâches de conscience phonologique, RAN, mémoire phonologique et lecture.

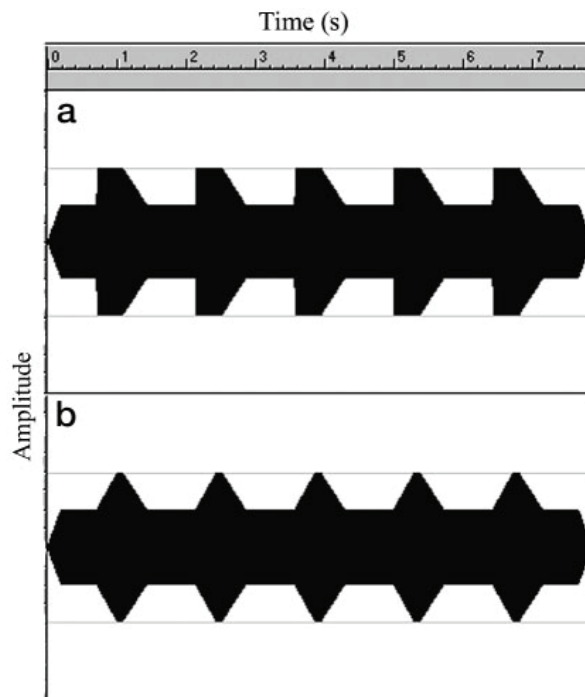


Figure 20

Tiré de Goswami et al., (2002). Exemple de stimulus, lorsque que le temps de hausse d'amplitude était (a) de 15 ms et (b) de 300 ms.

Goswami et al., interprètent ces résultats comme une difficulté à détecter les *perceptual centers* ou *P centers*. Les *P centers* correspondent, dans le signal de parole au moment de perception et sont associés avec une augmentation rapide de l'énergie sur le milieu de la bande spectrale. Leur détection permet ainsi de mieux différencier les syllabes et donc devraient influencer les performances en conscience phonologique. De façon intéressante, cette sensibilité au rythme est corrélée avec les compétences phonologiques dans plusieurs langues, notamment en chinois, n'utilisant pourtant pas le système alphabétique (Goswami et al., 2011). Plus généralement, les compétences rythmiques sont liées aux compétences linguistiques (Goswami, Huss, Mead, Fosker, & Verney, 2013; Huss, Verney, Fosker, Mead, & Goswami, 2011). Cette théorie ne permet toutefois pas d'expliquer les autres troubles auditifs observés au sein de la population dyslexique. Nous allons donc nous intéresser à la théorie du déficit d'ancrage qui vise à expliquer différents résultats de la littérature observés sur les tâches auditives.

4.3. La théorie du déficit d'ancrage (*Anchor Theory*)

Ahissar et collègues (Ahissar, 2007; Ahissar et al., 2006) ont également développé une théorie intéressante puisqu'elle permet à la fois d'expliquer le trouble phonologique et le manque de consistance des données de la littérature sur les traitements auditifs centraux dans la dyslexie. D'après cette théorie, les dyslexiques auraient des difficultés à former une ancre perceptuelle. C'est à dire qu'ils ne parviendraient pas à conserver en

mémoire à court-terme une représentation du stimulus, de ce fait dans les tâches demandant une comparaison entre deux stimuli, typiquement un standard et une cible (i.e., paradigme AX). Dans une tâche de discrimination de fréquence, Ahissar et al., (2006) mettent en évidence que les participants dyslexiques obtiennent des scores inférieurs à ceux des contrôles lorsque le paradigme utilise la répétition d'un même ton de comparaison. Dans ce cas, les dyslexiques ne bénéficient pas de la répétition de A. Alors que les participants contrôles comparent simplement X à leur représentation de A, les dyslexiques, eux, effectuent une comparaison à chaque essai entre les deux stimuli. Une telle comparaison est plus coûteuse cognitivement, la tâche est donc plus difficile et les scores sont moins bons. En effet, de façon intéressante, lorsque les deux sons purs sont changés à chaque essai (i.e., condition sans référence) les deux groupes obtiennent des performances similaires (cf. Figure 21).

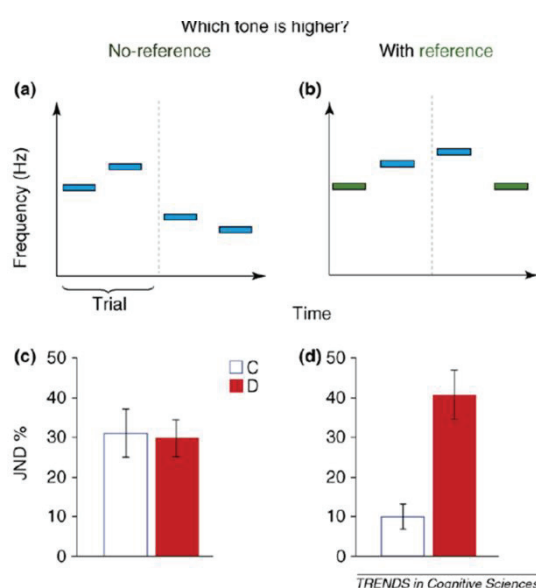


Figure 21

Tiré de Ahissar (2007). *Just Noticeable Difference (JND ; i.e., plus petite différence perçue)* en fonction du type de paradigme utilisé : sans ton de comparaison à gauche et avec un ton de comparaison à droite.

C = Contrôles et D = Dyslexiques.

La même observation est faite en modalité visuelle, les dyslexiques sont moins bons à discriminer deux fréquences spatiales uniquement si elles sont présentées l'une après l'autre (Ben-Yehudah & Ahissar, 2004). Quand les deux fréquences spatiales sont présentées simultanément sur le même écran alors les participants dyslexiques sont tout aussi performants à les discriminer que les participants contrôles.

Cette théorie présente l'avantage de proposer une seule explication pour de nombreux domaines connus pour être déficitaires dans la population dyslexique comme la mémoire de travail, les difficultés auditives et visuelles (Amitay, Ben-Yehudah, Banai, & Ahissar, 2002; Ramus et al., 2003; White et al., 2006) et la plus grande sensibilité au bruit (Dole, Hoen, & Meunier, 2012; Sperling, Lu, Manis, & Seidenberg, 2005; Ziegler,

Pech-Georgel, George, & Lorenzi, 2009). Elle explique également les troubles de la conscience phonologique. En effet, certains auteurs suggèrent que les mauvaises performances des dyslexiques aux tâches de conscience phonologique seraient dues à un déficit d'accès aux représentations phonologiques plutôt qu'à des représentations mal spécifiées (Boets et al., 2013; Ramus & Szenkovits, 2008). L'*Anchor Theory* postule également que c'est la nature des tâches qui génère les mauvaises performances. Ainsi, la tâche de RAN demande de faire référence un grand nombre de fois à un petit set de stimuli. Les participants contrôles vont donc maintenir en mémoire l'image correspondant au mot et y accéderont rapidement alors que les participants dyslexiques auront plus de difficulté parce qu'ils ne créeront pas cet ancrage. Ahissar (2007) postule donc que les résultats des dyslexiques à cette tâche seraient équivalents à ceux des sujets contrôles si le nombre de stimuli était plus important.

A l'origine, cette théorie présente tout de même une limite importante, en effet, si elle permet d'expliquer les déficits associés à la dyslexie, elle ne permet de créer un lien direct avec les compétences en lecture (Ahissar, 2007). Toutefois, Ramus et Ahissar (2012) postulent que l'incapacité à effectuer une ancre perceptuelle peut mener à terme à une mauvaise représentation des régularités de la parole (comme les transitions formantiques) et donc à de mauvaises représentations phonologiques.

Une autre limite de cette théorie serait de ne reposer que sur des données comportementales et ne permet pas, pour l'instant, d'expliquer les données électrophysiologiques.

4.4. La théorie de l'instabilité du traitement auditif

L'équipe de Kraus et collègues (Hornickel, Anderson, Skoe, Yi, & Kraus, 2012; Hornickel, Chandrasekaran, Zecker, & Kraus, 2011; Hornickel, Knowles, & Kraus, 2012; Hornickel & Kraus, 2013; Kraus et al., 1996; White-Schwoch et al., 2015) s'est intéressée à la stabilité des réponses du système nerveux aux sons de paroles. Le raisonnement étant que chez les enfants dyslexiques la réponse du système nerveux aux stimuli auditifs est instable, de ce fait les représentations phonologiques sont instables. Les auteurs postulent alors que cette instabilité rend plus difficile l'accès à la conscience phonologique durant la période préscolaire. En retour, le mauvais développement de la conscience phonologique diminue la stabilité des réponses du système nerveux aux sons de paroles.

Ainsi, Hornickel et Kraus (2013) ont étudié la réponse du tronc cérébral à la présentation de syllabes (i.e., *speech Auditory Brainstem Response* ; ABR). Cent enfants âgés de 6 à 13 ans ont été testés, ils étaient divisés en trois groupes mauvais, bons ou moyens lecteurs. L'analyse des résultats a mis en évidence que les réponses ABR des bons lecteurs étaient beaucoup plus stables dans le temps et présentaient moins de variabilités inter-essais que pour les dyslexiques (cf. Figure 22)

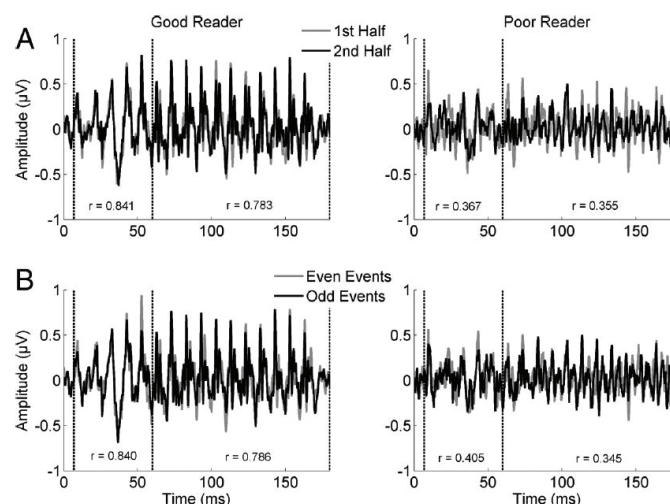


Figure 22

Tiré de Hornickel et Kraus (2013). Représentation graphique des réponses du tronc cérébral aux syllabes d'un bon lecteur à gauche et d'un mauvais lecteur à droite. A. les essais sont triés en fonction du moment de leur présentation, 1ère ou 2ème moitié de l'expérience pour illustrer la fatigue neuronale. B. Les essais sont triés en fonction de leur ordre de présentation (paire ou impaire) pour illustrer la variabilité des réponses neuronales.

De façon intéressante, cette théorie rend également compte de la difficulté des personnes (adultes et enfants) dyslexiques (Bradlow, Kraus, & Hayes, 2003; Dole et al., 2012; Ziegler et al., 2009) ou même des enfants pré-lecteurs issus de famille à risques (Boets, Ghesquiere, van Wieringen, & Wouters, 2007) à percevoir la parole dans le bruit. En effet, la présence de bruit va rendre encore plus difficile la perception du signal de parole, augmentant son instabilité (i.e., le bruit masque les informations redondantes du signal de parole) et augmentant la fatigue neuronale et donc la variabilité des réponses du système nerveux. De plus, bien que la situation de parole dans le bruit soit transitoire, elle est fréquente, et les difficultés de traitements de la parole récurrentes pourraient à terme entraver le développement du langage. Ainsi, Boets et al. (2011) ont mis en évidence dans une étude longitudinale que la capacité des enfants de *kindergarten* et en première année d'école primaire à percevoir la parole dans le bruit ainsi que la sensibilité à la modulation de fréquence permettait de prédire les compétences en lecture. Ces résultats ont été répliqués dans une étude électrophysiologique, mettant en évidence que la précision avec laquelle la parole dans le bruit est traitée par le système auditif chez les enfants de 3 ans prédit les compétences en pré-lecture l'année suivante (White-Schwoch et al., 2015).

De façon générale, les dyslexiques présentent des compétences auditives instables et peu répliquables, c'est d'ailleurs la critique majeure apportée à la littérature s'intéressant aux TTA dans la dyslexie.

4.5. Les TTA dans la dyslexie développementale

La littérature s'intéressant aux liens entre les TTA et la dyslexie développementale est très importante mais peu consistante. En effet, de nombreuses études trouvent des résultats contradictoires dus vraisemblablement à des différences méthodologiques, que ce soit au niveau du choix des tests psychoacoustiques, ou des critères de sélection de la population dyslexique. Cependant certains tests ont mis en évidence de façon consistante des différences significatives entre la population dyslexique et la population contrôle. Nous présentons ici les tests les plus utilisés ainsi que ceux constituant la BECAC ayant déjà été testés sur une population dyslexique. Ainsi, il a été démontré que les participants dyslexiques sont généralement significativement déficients dans beaucoup de tests impliquant un traitement spectral des stimuli. En effet, les études testant les compétences en discrimination de fréquence montrent généralement un déficit chez les dyslexiques (Ahissar et al., 2006; Ahissar et al., 2000; Amitay et al., 2002; Baldeweg et al., 1999; Banai & Ahissar, 2006; Cohen-Mimran & Sapir, 2007; France et al., 2002; Hari, Saaskilahti, Helenius, & Uutela, 1999; McAnally & Stein, 1996). Toutefois, certaines études ne mettent pas en évidence de différence significative (N. I. Hill, Bailey, Griffiths, & Snowling, 1999). Certaines de ces études ont de plus mis en évidence des corrélations significatives entre les performances en discrimination de fréquence et les compétences en lecture (sur l'ensemble de la population : Baldeweg et al., 1999) ou en conscience phonologique (uniquement sur le groupe dyslexique : Banai & Ahissar, 2006) et même en mémoire de travail (sur l'ensemble de la population et uniquement sur le groupe dyslexique : Ahissar et al., 2006).

La discrimination de durée a été moins étudiée que la discrimination de fréquence, et les résultats sont peu clairs. Chez les adultes il semble que la discrimination de durée ne soit pas déficitaire chez les dyslexiques (Baldeweg et al., 1999) bien qu'une corrélation significative avec les scores aux tâches de lecture apparaisse sur l'ensemble de la population. Sur une population d'enfants les résultats sont contradictoires, il semble que les enfants dyslexiques soient moins performants que les enfants contrôles lorsque les sons purs utilisés sont relativement courts (i.e., 100 ms) mais pas lorsqu'ils sont plus longs (i.e., 400 ms ; Banai & Ahissar, 2006). Enfin en utilisant du matériel verbal, Richardson et al. (2004) ont montré que les enfants dyslexiques étaient moins performants que les enfants contrôles à discriminer la durée de deux exemplaires du même pseudo-mot (i.e., [ata]). Cette étude a également mis en évidence des corrélations significatives entre les performances à ce test et les scores en lecture, orthographe et lecture de pseudo-mot sur l'ensemble de la population.

La discrimination d'intensité a rarement été testée, mais deux études ont mis en évidence que les dyslexiques adultes (Amitay et al., 2002) et enfants (McArthur & Hogben, 2001) étaient moins performants à cette tâche que les contrôles. En revanche,

Ahissar et al. (2000) n'ont pas montré de différence entre les performances des adultes contrôles et des adultes avec une histoire de dyslexie à la tâche de discrimination d'intensité. Sur une population d'enfants Richardson et al. (2004) n'ont pas non plus mis en évidence de différences entre les performances des enfants contrôles et les enfants dyslexiques. Les résultats ne corrôlaient pas non plus avec les compétences en lecture dans aucun des deux groupes. En revanche, ces scores étaient significativement corrôlés avec la mémoire phonologique à court-terme à la fois dans l'étude de Richardson et al. (2004) et dans l'étude de Amitay et al., (2002).

Ensuite, dans des tests de *backward masking* ou masquage rétroactif, il semble que les participants dyslexiques (enfants et adultes) ne soient pas déficitaires comparativement aux contrôles (Hill et al. 1999 ; Ramus et al. 2003) et chez les adultes les performances à ce test ne sont pas corrôlées aux compétences en lecture (Ahissar et al., 2000). En revanche, Rosen et Manganari (2001) trouvent un déficit en masquage rétroactif lorsque les stimuli utilisés sont des stimuli de parole.

Les tests de *temporal order judgment* (TOJ) ou jugement d'ordre temporel génèrent également des résultats assez confus. Ainsi, sur une population d'adultes Ramus et al. (2003) trouvent un déficit chez les sujets dyslexiques que la durée des stimuli soit longue (115 ms) ou courte (30 ms). Les compétences de jugement d'ordre temporel évaluées avec des stimuli variant soit en latéralisation soit en fréquence sont également apparues déficitaires dans la population dyslexique adulte (Ben-Artzi, Fostick, & Babkoff, 2005). Toutefois, seules les compétences en jugement d'ordre temporel fréquence étaient corrôlées aux compétences en lecture et uniquement au niveau de l'ensemble de la population. Dans l'étude longitudinale de Share et al. (2002) les jeunes enfants se révélant dyslexiques en grandissant ne sont déficitaires à cette tâche que lorsque l'ISI est long (i.e., 400 ms) mais pas lorsqu'il est court (8 ; 15 ; 30 ; 60 ; 150 ms) contredisant ainsi les résultats de Tallal (1980).

Les études s'intéressant à la capacité à percevoir les modulations de fréquence présentent des résultats peu consistants. Ainsi, en 1998, Witton et al., trouvent cette compétence déficitaire chez les adultes dyslexiques lorsque cette modulation est relativement lente (2 Hz et 40 Hz) mais pas lorsqu'elle est rapide (240 Hz). Ces résultats sont partiellement répliqués par Ramus et al. (2003) qui trouvent une différence entre le groupe contrôle et le groupe dyslexique lorsque la modulation est de 2 Hz mais pas 240 Hz (condition 40 Hz non testée). Finalement, Hill et al., (1999) en estimant la plus petite modulation de fréquence détectée sur un son pur de 1 kHz et sur un son pur de 6 kHz ne trouvent pas de différence significative en détection de modulation de fréquence entre les dyslexiques et les contrôles.

Sur les tests de démasquage binaural, Hill et al., (1999) ne trouvent pas de différences entre les groupes dyslexiques et contrôles, alors que McAnnally et Stein (1996) mettent

en évidence une moins bonne détection du signal par les dyslexiques que par les contrôles lorsqu'il est présenté en antiphasse entre les deux oreilles.

Enfin, concernant la compréhension de parole dans le bruit, les dyslexiques adultes sont apparus déficitaires en reconnaissance de mot isolés (Dole et al. 2012). Les enfants sont également apparus déficitaires dans la plupart des études (Bradlow et al., 2003; Calcus, Colin, Deltenre, & Kolinsky, 2015; Ziegler et al., 2009). En revanche, Messaoud-Galusi, Hazan et Rosen (2011) n'ont pas mis en évidence de tel déficit sur une population d'enfants, que ce soit en reconnaissance de mots isolés ou de mots dans un contexte.

Comme cette partie le démontre, de nombreuses études ont mis en évidence la présence de TTA chez les dyslexiques, et également nombre d'entre elles ont mis en évidence des corrélations significatives entre les performances auditives et les compétences en lecture ou phonologique. Ramus et al. (2003) par exemple en calculant un score composite pour l'ensemble des tests auditifs trouvent une corrélation significative à la fois avec les compétences en lecture et la conscience phonologique sur l'ensemble de la population.

Deux principales critiques peuvent être émises à l'égard de ces études. Tout d'abord, certaines d'entre elles testent un très petit nombre de participants (e.g., 10 enfants dyslexiques : Calcus et al., 2015; 13 adultes dyslexiques : Helenius, Uutela, & Hari, 1999; 12 adultes dyslexiques : N. I. Hill et al., 1999; 12 adultes dyslexiques : Lallier et al., 2009; 8 enfants dyslexiques : Rosen & Manganiari, 2001). La présence de TTA est attendue uniquement sur une sous-partie de la population dyslexique (entre 40% et 60% des dyslexiques ; Ramus, 2003). De ce fait, tester un trop petit échantillon ne permettra pas de mettre en évidence un trouble éventuel. Ensuite, beaucoup de ces études mettent en évidence des corrélations entre les compétences auditives et langagières mais uniquement sur l'ensemble des participants (groupe contrôle + groupe dyslexique. Dès lors, ces corrélations ne peuvent réellement être interprétées comme représentant un lien mais simplement un effet de groupe. L'objectif de l'Etude 4 de l'Axe 2 sera donc de clarifier les liens entre les compétences auditives centrales et la dyslexie en testant 20 participants dans chaque population sur un nombre plus grand de tests.

5. Résumé du Chapitre V

L'objectif de ce chapitre était de présenter la dyslexie. Nous avons vu que la dyslexie est un trouble développemental, caractérisé par un retard d'au moins 18 mois dans l'apprentissage de la lecture en dépit d'une intelligence non verbale normale, d'un environnement suffisamment stimulant et d'aucun trouble affectif, psychiatrique ou neurologique. Plusieurs types de dyslexies ont été décrites dont la classification est particulièrement pertinente pour la rééducation. Enfin nous nous sommes intéressés à deux grands pans de la littérature sur les origines de la dyslexie : la théorie

phonologique et la théorie d'un déficit du traitement auditif. L'existence d'un déficit phonologique au sein de la quasi-totalité de la population fait l'objet d'un consensus. Le débat en revanche reste ouvert quant au fait que ce déficit soit à l'origine de la dyslexie (i.e., théorie phonologique) ou soit une conséquence d'un autre trouble, comme un trouble auditif. La théorie du traitement du déficit auditif est composée de plusieurs théories. En effet, si de nombreux auteurs admettent qu'il existe un trouble du traitement auditif chez les dyslexiques la nature exacte de ce déficit est très discutée. Nous présentons 4 principales théories tentant d'expliquer la dyslexie par un trouble de l'audition. Enfin, nous passons en revue les autres troubles du traitement auditif observés régulièrement au sein de la population dyslexique sans qu'ils ne fassent l'objet d'une théorie ou d'une hypothèse spécifique. Un des objectifs de l'Etude 4 sera de clarifier (a) l'existence de ces troubles du traitement auditif dans la dyslexie et (b) leurs liens avec les compétences langagières.

Synthèse et problématique générale

Cette partie Théorique est composée de 5 chapitres présentant les principaux thèmes abordés dans ce travail. Le premier chapitre présente 4 modèles de reconnaissance du mot parlé. De façon générale, ces modèles considèrent que le signal de parole va activer les représentations des phonèmes entendus. Celles-ci vont à leur tour activer les mots contenant ces phonèmes (i.e., processus *bottom-up*). Certains modèles prévoient un effet du contexte (i.e., processus *top-down*) au moment de la reconnaissance du mot ou avant.

Le second chapitre s'intéresse au traitement sémantique, et souligne le débat de la littérature quant à l'automatisme de l'accès au sens des mots. Au regard de la littérature, il semble que les effets d'amorçage obtenus avec des paradigmes d'amorçage sémantique masqué soient davantage des effets de congruence de réponse que des effets d'amorçage sémantique *per se*. En modalité auditive, aucun résultat concluant n'a été rapporté que ce soit en utilisant des méthodes comportementales ou électrophysiologiques. Ceci suggère alors que le traitement sémantique ne serait pas entièrement automatique. Il pourrait d'ailleurs être modulé par le contexte, comme suggéré par la littérature sur les mots ambigus qui postule que le contexte dans lequel le mot ambigu est présenté influence la force d'activation de ses différentes significations.

Le troisième chapitre présente la situation de parole dans le bruit. Il distingue deux types de masquages du signal cible par le fond sonore : le masquage énergétique, résultant d'interférences au niveau spectro-temporel et le masquage informationnel, résultant d'interférences à un plus haut niveau (e.g., linguistique). Ces deux types de masquages apparaissent principalement en situation de parole dans le bruit et dans le cas spécifique de la parole dans le bruit parolier respectivement. Ils sont à l'origine de la diminution de l'intelligibilité du signal cible. D'après l'*Effortfulness Hypothesis* l'effet du bruit peut être plus important qu'une simple perte d'intelligibilité. En effet, la présence de bruit concurrent au signal cible peut perturber les traitements de haut niveau comme la mémoire de travail.

Le chapitre 4 s'intéresse au système auditif, dont le traitement de la parole est nécessairement tributaire. Il présente en premier lieu l'anatomie du système auditif périphérique puis du système auditif central. Leurs développements respectifs sont ensuite étudiés. Si le premier est mature dès les premiers mois de vie, le développement du second se poursuit jusqu'à l'âge adulte. Pourtant, le développement de l'ensemble des compétences auditives centrales reste peu étudié, la plupart des études ne testant qu'une ou deux d'entre elles. La revue de la littérature montre que le développement observé est très dépendant du type de test utilisé, même lorsque le même processus sous-jacent est testé. Ensuite, le chapitre se poursuit par la présentation des Troubles du Traitement Auditif (TTA), et du peu de moyens diagnostics à disposition des praticiens en France. Enfin, la Batterie d'Évaluation des Compétences Auditives Centrales (BECAC) est présentée.

Enfin, le dernier chapitre s'intéresse à la dyslexie dont les TTA pourraient être à l'origine. La définition et les critères diagnostics de la dyslexie peuvent être sujets à critiques, pourtant, il existe un consensus sur l'existence d'un trouble phonologique à l'origine de la difficulté à apprendre à lire. Certaines théories postulent que le trouble phonologique provient d'un déficit de perception du langage oral du fait de TTA. Seules les théories postulant que la dyslexie provient d'un déficit auditif de bas niveau sont présentées au sein du chapitre puisque c'est cet aspect de la dyslexie qui est étudié dans la partie Expérimentale. Quatre théories sont présentées, postulant l'existence d'un trouble auditif bas niveau spécifique, permettant d'expliquer les déficits en conscience phonologique et en lecture.

L'objectif global de cette thèse est donc de s'intéresser aux liens entre la dégradation du signal de parole et le langage.

Deux axes sont développés, le premier vise à déterminer (1) si et comment la dégradation du signal du à l'ajout de bruit impacte le traitement sémantique malgré une intelligibilité préservée (Etude 1 et Etude 2). Cette question principale permet également de s'intéresser (2) à l'automatisme du traitement sémantique (Etude 1 et 2) et (3) aux interactions entre processus *bottom-up* et *top-down* dans l'accès au sens des mots (Etude 2).

Le second axe a pour objectif d'étudier (4) comment la dégradation du signal du fait d'un mauvais traitement auditif va impacter les représentations du langage sur une population d'enfants au développement normal et d'adultes sains et dyslexiques (Etude 3 et Etude 4). Cette question principale permet également de (5) s'intéresser au développement normal des traitements auditifs centraux et à l'impact des facteurs non-sensoriels dans ce développement (Etude 3).

Partie Expérimentale

Chapitre VI

Etude 1 : L'amorçage sémantique en situation de parole dans la parole

Cette première étude s'intéresse à l'effet de la dégradation du signal sur le traitement sémantique d'un mot inséré au sein d'un bruit parolier.

Comme nous l'avons vu au sein de la partie Théorique, l'automatisme du traitement sémantique est habituellement testée en modalité visuelle grâce à un paradigme d'amorçage masqué. A ce jour aucune alternative à ce paradigme n'a été trouvée pour la modalité auditive. Nous proposons ici d'utiliser la situation de parole dans la parole pour tester la résistance de l'amorçage sémantique à la dégradation du signal de parole. Cette étude a fait l'objet d'une publication dans la revue *Frontiers in Human Neuroscience*.

Cette première étude vise principalement à déterminer si et comment le traitement sémantique de l'amorce résiste au masquage de celle-ci. Par ailleurs, il a été montré que des interférences entre la cible et le masqueur peuvent apparaître au niveau phonologique, lexical, et sémantique dans le cadre d'une phrase. Cette étude permet alors d'évaluer si des interférences sémantiques peuvent apparaître au niveau du mot entre l'amorce et les masqueurs.

Dans 3 expériences, 72 participants (i.e., 3 x 24 participants) devaient effectuer une tâche de décision lexicale sur un item cible intégré au sein de fond sonores composés de bruit parolier. Dans l'Expérience 1, le fond sonore était composé d'une ou deux voix Sémantiquement Consistantes (SC) c'est-à-dire prononçant des mots sémantiquement liés entre eux (e.g., *château, cabane, villa, manoir, appartement, habiter* ; cf. p291) et liés au mot cible dans 50% des essais *mot*. Ces voix SC jouent donc le rôle d'amorce dans notre paradigme. Dans les Expériences 2 et 3 respectivement 1 et 2 voix sémantiquement inconsistantes (SI) étaient ajoutées à chaque fond sonore. Ces voix prononçaient des mots sémantiquement indépendant les uns des autres, et non liés au mot cible (e.g., *étagère, policier, tasse, peinture, dormir*). Les voix SC jouent donc le rôle de masqueurs. Le ratio de voix SC par rapport au nombre de voix SI est ainsi varié en

fonction des conditions et des expériences (Exp 1 : 1SC/oSI et 2SC/oSI ; Exp 2 : 1SC/1SI et 2SC/1SI ; Exp 3 : 1SC/2SI et 2SC/2SI). L'objectif de cette variation est d'évaluer d'une part le besoin d'intelligibilité pour le traitement sémantique de l'amorce : en variant le nombre de voix total dans le fond sonore, i.e., de 1 à 4 voix. D'autre part, elle permet d'évaluer le besoin de saillance de l'amorce : le nombre de voix SC peut être inférieur, égal ou supérieur au nombre de voix SI. Les temps de réponses et les pourcentages de bonnes réponses ont été recueillis. Nous posons l'hypothèse que le masquage informationnel peut apparaître au niveau du traitement sémantique du mot. Un effet d'amorçage sémantique significatif devrait donc apparaître au moins dans l'Expérience 1 où seules les voix SC constituent le fond sonore. Nous postulons que si le traitement sémantique est automatique comme postulé par les modèles de Quillians (1967) et Collins et Loftus (1975), alors il ne sera pas influencé par l'augmentation du nombre de voix dans les fonds sonores ni par la saillance de l'amorce. En revanche, si le traitement sémantique n'est pas automatique, alors celui-ci devrait disparaître/diminuer avec l'augmentation du nombre de voix en fond sonore et/ou avec la diminution du ratio de nombre de voix SC/SI.

Multi-talker background and semantic priming

Marie Dekerle^{a,b}, Véronique Boulenger^{b,c}, Michel Hoen^{b,d}, Fanny Meunier^{a,b}

(a) CNRS, UMR 5304, Laboratoire sur le Langage, le Cerveau et la Cognition, Lyon, F-69000, France

(b) University of Lyon, F-69000, France

(c) CNRS, UMR5596, Laboratoire Dynamique Du Langage, Lyon, F-69000, France

(d) CNRS, UMR5292, INSERM, U1028, Lyon Neuroscience Research Center, Brain Dynamics and Cognition Team, Lyon, F-69000, France

Abstract

The reported studies have aimed to investigate whether informational masking in a multi-talker background relies on semantic interference between the background and target using an adapted semantic priming paradigm. In 3 experiments, participants were required to perform a lexical decision task on a target item embedded in backgrounds composed of 1 to 4 voices. These voices were Semantically Consistent (SC) voices (i.e., pronouncing words sharing semantic features with the target) or Semantically Inconsistent (SI) voices (i.e., pronouncing words semantically unrelated to each other and to the target). In the first experiment, backgrounds consisted of 1 or 2 SC voices. One and 2 SI voices were added in Experiments 2 and 3 respectively. The results showed a semantic priming effect only in the conditions where the number of SC voices was greater than the number of SI voices, suggesting that semantic priming depended on prime intelligibility and strategic processes. However, even if backgrounds were composed of 3 or 4 voices, reducing intelligibility, participants were able to recognize words from these backgrounds, although no semantic priming effect on the targets was observed. Overall this finding suggests that informational masking can occur at a semantic level if intelligibility is sufficient. Based on the Effortfulness Hypothesis, we also suggest that when there is an increased difficulty in extracting target signals (caused by a relatively high number of voices in the background), more cognitive resources were allocated to formal processes (i.e., acoustic and phonological), leading to a decrease in available resources for deeper semantic processing of background words, therefore preventing semantic priming from occurring.

Key words: semantic priming, informational masking, cocktail party, cognitive load, Effortfulness Hypothesis.

1. Introduction

In daily life, speech is rarely perceived in silence, but with interference from wind, music or other people's conversation. Although often used to study psychoacoustic topics (Brungart, 2001; Brungart et al., 2001; McDermott, 2009), speech-in-noise and cocktail party situations, (i.e., speech-in-speech; Cherry, 1953)) also appear to be interesting paradigms to tackle linguistic processes and competition occurring between backgrounds and targets (Boulenger et al., 2010; Hoen et al., 2007). Our study aimed to investigate the extent to which multi-talker background is processed semantically when listening to speech-in-speech and therefore how the cocktail party situation can be used to study the automaticity of word semantic activation.

The cocktail party situation is described as involving two types of masking effects: energetic and informational masking (Brungart, 2001). Energetic masking relies on the spectro-temporal features of sounds and results from different sounds stimulating the same part of the cochlea at the same time so that one of them cannot be heard (i.e., as two signals increasingly share spectro-temporal characteristics, energetic masking becomes more efficient). In multi-talker background situations, the magnitude of energetic masking is proportional to the number of voices that comprise the background (Simpson & Cooke, 2005). Informational masking, however, usually refers to masking effects that cannot be attributed to energetic masking. Specifically, it is related to the overlap of information carried by the different signals at a higher level (e.g., lexical level and working memory; see Cooke, Garcia Lecumberri, & Barker, 2008; Durlach, 2006; Mattys, Brooks, & Cooke, 2009; Mattys & Wiget, 2011). Whereas background noise mainly elicits energetic masking, a speech background produces both energetic and informational masking (Brungart et al., 2006). Despite the masking, it is still possible to detect and recognize a word or linguistic token embedded in a babble. Of course, as more voices are present in the babble, participants become less accurate (Freyman, Balakrishnan, & Helfer, 2004). However, it is interesting to note that Simpson and Cooke (2005) showed, using a -6 dB SNR, that intelligibility decreases as a monotonic function of the number of speakers in babbles of up to 8 voices. Specifically, participants' accuracy to detect the target token decreases as the number of voices increases up to 8 voices. Further increasing the number of voices does not lead to a decrease in accuracy. These results suggest that if energetic masking is too high, informational masking decreases with the diminution of the available linguistic cues. For example, with more than 8 talkers, phonetic cues are not or less available and therefore cannot be attributed incorrectly to the target.

The first aim of this paper is to test whether semantic features are involved in informational masking. It has been established that informational masking is not monolithic and occurs at many linguistic levels. Indeed, a multi-talker background will create less interference on a target word, if it is pronounced in a different language (Brouwer et al., 2012; Van Engen & Bradlow, 2007) and different languages will not have the same masking power (Gautreau et al., 2013). By manipulating the number of talkers in the background, Boulenger et al. (2010) revealed lexical competitions between a 2-talker background and target speech using a lexical decision task. Increasing the number of voices in the background, however, led to the disappearance of lexical interference because of increased energetic masking (i.e., words from the background became less intelligible and therefore competed less with target processing). However, using the same paradigm but with an intelligibility task, Hoen et al. (2007) showed that lexical processing of a background could be performed with up to 4 concurrent voices; beyond that threshold, masking was too high and seemed to prevent linguistic processes. Although it has been shown that phonological and lexical information contribute to the informational masking effect, our experiments tested the role of semantic information.

Processing of the background's semantic content has already been highlighted with 2 talkers, pronouncing either semantically correct sentences (i.e., 'rice is often served in round bowls') or incorrect sentences (i.e., 'the great car met the milk', Brouwer et al., 2012). Semantic incoherence in the background impacts the recognition of the target sentence. This result suggests that the background signal with 2 talkers is processed semantically. Our experiments aimed to identify how many talkers are allowed in this semantic processing using words and how semantic information from the backgrounds interferes with the identification of target words.

The ability to semantically process auditory words presented outside of the attentional focus is traditionally studied using dichotic listening. This paradigm allows to study pure informational masking as no energetic masking occurs in dichotic listening. However, discrepant results have been reported (Cherry, 1953; Dupoux et al., 2003; Eich, 1984; Lewis, 1970; Wood et al., 1997). In 1984, Eich showed a semantic effect on the recognition of words presented in the unattended channel. However, this effect resulted, at least partially, from an attentional shift toward the to-be-ignored channel (as suggested by Wood et al., 1997). As the speaker rate was very slow in Eich's experiment (85 words per minute), it allowed participants to listen to the supposedly unattended channel without disturbing the primary task (in the case of Eich's study, a shadowing task). In replicating Eich's study using the same speech rate, Wood and collaborators (1997) observed the same semantic effect; however, it disappeared if the speaker's rate was increased to 170 words per minute, corresponding to a more ecologically valid rate. The authors concluded that as this faster speech rate demanded

more cognitive resources, participants could no longer shift attention to the unattended channel while performing the primary task, suggesting that at least in dichotic listening, informational masking does not involve semantic information. The issue raised by this paradigm is that the spatial separation of auditory signals creates a masking release compared to a binaural condition and therefore facilitates stream segregation that could prevent competition between the to-be-ignored and target speech (Drullman & Bronkhorst, 2000; Hawley, Litovsky, & Culling, 2004).

Concerning semantic activation and according to traditional theoretical models, semantic memory is organized into networks. The recognition of one word leads to its activation in semantic memory, and this activation is supposed to spread automatically to other related concepts (Collins & Loftus, 1975; Collins & Quillian, 1969). This supposition is derived from semantic priming paradigms, shown in auditory and visual modalities, in which the presentation of prime word leads to faster recognition of a semantically related target word (Donnenwerth-Nolan et al., 1981; Meyer & Schvaneveldt, 1971; Radeau, 1983; Radeau et al., 1998; Schacter & Church, 1992; Spruyt et al., 2012). For example, the presentation of the prime ‘nurse’ before the target ‘doctor’ facilitates the recognition of the target word ‘doctor’ compared to a condition in which the prime is unrelated to the target (Meyer & Schvaneveldt, 1971). Adapting this paradigm to the cocktail party situation will allow us to investigate if the semantic content of the background is processed and interferes with the target word, despite decreased intelligibility. Some background words will therefore act as primes.

In the current study, we used the rationale of a priming paradigm by manipulating the association between words pronounced in the background and target words. Additionally, we varied the amount of masking to evaluate how it modulates semantic priming effects. Participants were required to perform a lexical decision task on a target item (i.e., decide whether the target item is a word or a pseudo-word) embedded in backgrounds composed of 1 to 4 voices depending on the experiment. These voices could pronounce words that were semantically related to each other and that were related or unrelated to the target. They acted as primes and were called Semantically Consistent (SC) voices. Additional voices pronounced words that were always unrelated to each other and unrelated to the target, acting as maskers. They were called Semantically Inconsistent (SI) voices.

Across experiments, we manipulated the ratio between SC and SI voices. The aim was to test the preservation of the semantic processing of SC voices despite increased masking (i.e., more SI voices). In Experiment 1, backgrounds were composed of 1 or 2 SC voices. In Experiments 2 and 3, respectively, 1 and 2 SI voices were added to each background to increase masking and therefore, decrease the intelligibility of the SC voices. Consequently, in Experiment 2, backgrounds in one condition consisted of 1 SC voice and 1 SI voice and in a second condition of 2 SC voices and 1 SI voice. In

Experiment 3 they comprised 1 SC voice and 2 SI voices in one condition and 2 SC voices and 2 SI voices in the other condition.

Overall increasing the number of voices allowed us to examine if and how semantic priming can be impacted by the increase in the number of talkers in the background. Additionally, the variation in the number of SC voices compared to the number of SI voices allowed us to study the effect of prime saliency on semantic processing and therefore its participation in informational masking. Indeed, across experiments, backgrounds can consist of the same number of voices whereas the number of SC voices compared to the number of SI voices could differ (e.g., 3 voices in the background: either 2SC/1SI in Exp 2 or 1SC/2SI in Exp 3).

If semantic processing can occur automatically, semantic priming should be observed at least as long as background words are intelligible and should not be disturbed by increased masking and decreased prime saliency. Indeed, automaticity is defined as a strategy free processing that occurs without using the resources of a limited capacity central processor (Neely, 1977). Therefore, if semantic processing is strategy free, it should occur even if participants are not aware that a given word is presented to them (as is done in visual modality in classical masked priming paradigms, see Forster and Davis, 1984).

2. Experiment 1

The aim of this experiment was to first establish set up and test our paradigm and experimental materials. Backgrounds were composed of 1 or 2 SC voices that pronounced words sharing semantic features with each other. In the related condition, target words belonged to the same semantic field as the prime, but they did not in the unrelated condition. We therefore expected to observe a semantic priming effect: participants should more quickly and accurately identify target words in the related compared to the unrelated condition. The second aim of this first experiment was to test if the presence of 2 voices in the background would affect participants' performance as suggested by the psychoacoustic literature (Brungart, 2001; Brungart et al., 2001). We therefore hypothesized that target words would be answered to more slowly and less accurately in the 2SC condition compared to the 1SC condition. Finally, we examined whether the semantic priming effect was modulated by increased energetic and informational masking caused by the augmentation of the number of voices in the background.

2.1. Method

2.1.1. Participants

Twenty-seven participants (20 females) volunteered for this experiment. All were right-handed, French native speakers and reported no known hearing or language disorder. Subjects' ages ranged from 18 to 25 years old. All participants gave written informed consent and were not aware of the experiment's purpose. They were compensated for their participation. The protocol that was used in this experiment was approved by the local ethics committee (CPP Sud-Est IV, Lyon; ID RCB: 2008-A00708-47).

2.1.2. Stimuli

Forty-eight disyllabic target words ($M_{\text{lexical frequency}} = 21.94$ per million, $SD = 18.75$ according to the French database Lexique 3, New et al., 2005) were selected, and each word belonged to a specific semantic field (e.g., *CAROTTE* 'carrot'; *MÉTRO* 'subway'). Each target word was matched to 10 words belonging to the same semantic field (e.g., *CAROTTE* 'carrot' was associated with *légume*, *chou*, *céleri*, *salade*, *tomate* 'vegetable, cabbage, celery, lettuce, tomato'). As participants had to perform a lexical decision task, 48 pseudo-words respecting French phonotactic rules were created (e.g., *PLARO*, *HUMEL*). Ten words sharing semantic features with each other were arbitrarily associated with each pseudo-word target, resulting in a total of 96 groups of 10 words ($M_{\text{lexical frequency}} = 21.86$, $SD = 18.20$). As each background comprised 1 or 2 SC voices (related or not to the target), each group was divided into two subgroups of 5 words one of the subgroups was spoken by a first speaker (S_1), and the other by a second speaker (S_2).

Target words were presented with a semantically related (related condition) or semantically unrelated background (unrelated condition). In the unrelated condition, SC voices pronounced words that were semantically related to each other but not to the target (see Figure 23). Backgrounds comprised 1 SC voice (1SC condition) or 2 SC voices (2SC condition). The 48 target words were divided into 4 groups of 12 words, the mean frequency did not differ significantly between the groups ($F < 1$), nor did the number of phonemes ($M = 6.97$, $SD = 5.65$; $F(3,44) = 1.1$, *n.s.*) and phonological neighbors ($M = 4.75$, $SD = 0.81$; $F(3,44) = 2.2$, *n.s.*). Each group of twelve target words was assigned to a condition (1SC related, 1SC unrelated, 2SC related, 2SC unrelated) depending on the experimental list. The same was true for pseudo-words. Four experimental lists of 96 stimuli (i.e., 48 target words and 48 target pseudo-words) were created so that each target word was presented in each condition, but only once in a list (each participant was presented with one list only).

Targets and SC voices were recorded by 3 different French native female speakers (age: 21-22) in a sound-proof room (22 kHz, mono, 16 bits). Auditory sequences of 5 words from Speakers 1 and 2 (S1 and S2) were segmented into 3 sec periods. The periods were then normalized at an intensity of 60 dB-A and mixed together to create backgrounds. All audio files were synchronized at the beginning, so all voices started to speak at the same time. However, as all voices pronounced words of different lengths, they soon became desynchronized, and there was always one speaker talking in the background. Targets recorded by Target Speaker (TS; also normalized at an intensity of 60 dB-A) were inserted 2 sec after the start of the backgrounds (so that each participant always had the same exposure to the background before the target speech was presented), with a 0 dB SNR (Signal/Noise Ratio; see Figure 23). Because the backgrounds, which comprised 1 or 2 voices, generated different amounts of energy, the intensity of all stimuli was varied over a +/- 3 dB range in 1 dB steps to prevent participants from predicting condition depending on individual stimuli intensity.

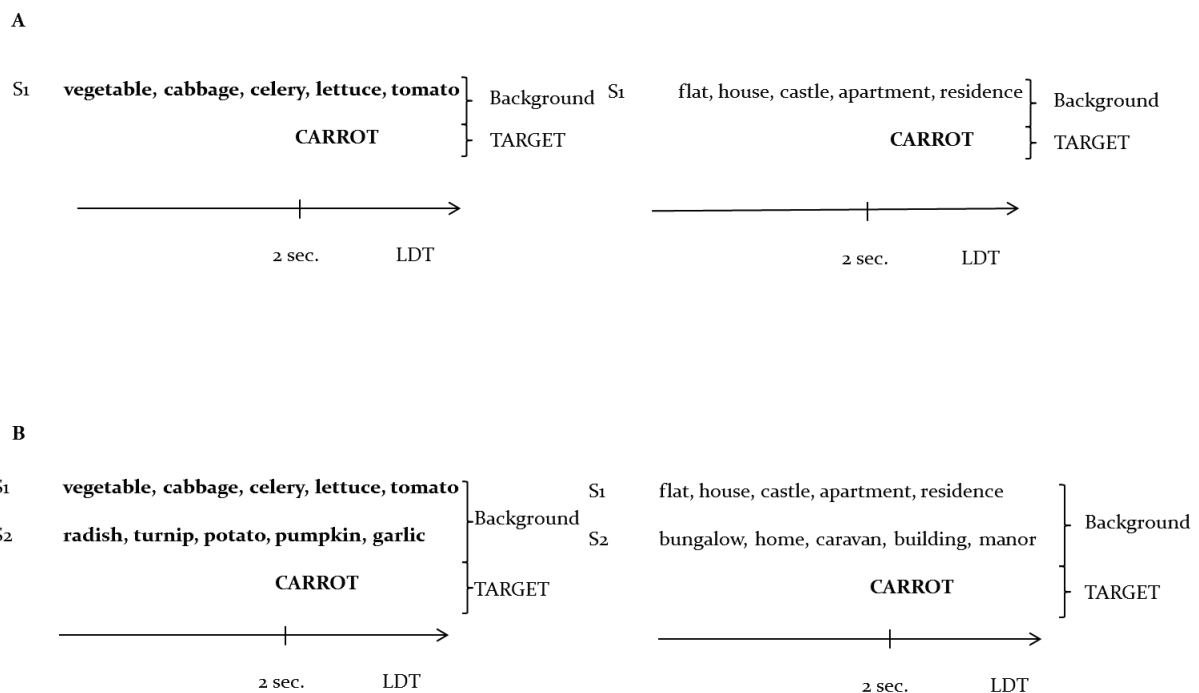


Figure 23

(a) Example of two backgrounds in the 1SC condition of Experiment 1, presented with a semantically related target word (left; related condition) or not (right; unrelated condition). S1 = speaker 1, LDT = Lexical Decision Task. (b) Example of a background in the V2 condition of Experiment 1, presented with a semantically related target word (left; related condition) or not (right; unrelated condition). S2 = speaker 2.

2.1.3. Procedure

Participants sat in front of a computer screen and heard the stimuli binaurally through headphones at a comfortable level (mean level 65 dB-A, ranging from 62 dB-A to 68 dB-A, normalized using an artificial ear). A fixation cross was presented on the screen at the beginning of each trial and remained on the screen during stimulus presentation. Participants were asked to listen to the stimuli to decide as quickly and accurately as possible whether the target was a word or a pseudo-word, by pressing one of two pre-specified keys. After a response was given, a string of hash marks indicated that the trial was over; participants could then press a key to start the next trial. Half of the participants gave the response to “word” with their left hand and to “pseudo-word” with their right hand. As all participants were right-handed, they might answer faster with their right hand than their left hand. To avoid this confounding effect, the other half were given the opposite instruction. A training session composed of twelve trials (different from the experimental stimuli) preceded the test session so that participants could acclimate to the stimuli and the task.

2.2. Results

Two two-way repeated measures analyses of variance (ANOVAs) by participants (F_1) and by items (F_2) were conducted, with Response Times (RTs, in ms) and Error Rates (ERs) for target word identification as dependent variables. We included Number of Voices in the background (1 Voice, 1SC or 2 Voices, 2SC) and Semantic Link between prime and target (related or unrelated) as within-subjects factors. Three participants were excluded from analyses because of very high Error Rates (ERs; more than 40%). Four target words error rates greater than 50% were also excluded from Item analyses (*POIGNET*, *RIDEAU*, *RATON* and *RACINE* ‘wrist, curtain, baby rat, root’). Trials with RTs below or above 2.5 standard deviations from the individual means (4.5%) and trials in which participants made mistakes (19.5%) were not included in RTs analysis. Means and SDs (Standard Deviations) of RTs and ERs are summarized in Table 1.

		1SC		2SC	
		Related	Unrelated	Related	Unrelated
RT (ms)	Mean	980	1035	1014	1070
	SD	170	142	169	150
ER (%)	Mean	7.08	16.33	19.5	20.77
	SD	6.82	12.34	19.04	14.05

Table 1 Means and Standard Deviations (SDs) of Response Times (RTs) and Error Rates (ERs) depending on the Number of Voices in the background and the Semantic Link between prime and target in Experiment 1. 1SC = 1 SC voice condition; 2SC = 2 SC voices condition; related = semantic link between the prime and the target; unrelated = no semantic link between the prime and the target.

The ANOVA by participants first revealed a significant main effect of the Number of Voices: participants were faster ($F_{1(1,23)} = 4.25$, $p = .05$) and more accurate ($F_{1(1,23)} = 9.12$, $p < .01$) to identify targets in the 1SC condition ($M_{RT} = 1008$ ms, $SD = 157$; $M_{ER} = 11.7\%$, $SD = 10.9$) than in the 2SC condition ($M_{RT} = 1042$ ms, $SD = 161$; $M_{ER} = 20.1\%$, $SD = 16.5$). The Item analysis, however, did not highlight an effect of the Number of Voices on RT ($F_{2(1,43)} = 1.9$, $p = .1$), although target words were better categorized as words in the 1SC condition ($F_{2(1,43)} = 13.43$, $p < .001$).

The main effect of Semantic Link also appeared to be significant on RTs ($F_{1(1,23)} = 14.24$, $p < .001$; $F_{2(1,43)} = 4.63$, $p < .05$), participants responded faster if targets shared semantic features with the prime ($M = 997$ ms, $SD = 169$) than if they did not ($M = 1053$ ms, $SD = 145$); this resulted in a 56 ms priming effect. This effect was also significant for ERs in the participant analysis ($F_{1(1,23)} = 3.93$, $p = .05$), and there was only a trend in the item analysis ($F_{2(1,43)} = 3.07$, $p < .10$). Participants tended to be more accurate in the related condition ($M = 15.4\%$, $SD = 15.4$) than in the unrelated condition ($M = 18.55\%$, $SD = 13.2$). There was no significant interaction between the two factors for RTs ($F_{1<1}$; $F_{2<1}$) and ERs ($F_{1(1,23)} = 2.34$, *n.s.*; $F_{2(1,43)} = 1.97$, *n.s.*), suggesting that the semantic priming effect was not modulated by the Number of Voices (one or two) in the background.

2.3. Discussion

These results first highlight that participants were slowed by the increase in the number of voices in the background. This effect is certainly attributable to enhanced target masking in the two-voice condition (Brungart, 2001; Brungart et al., 2001). Interestingly, participants' performance was improved by the semantic relationship between the prime and target, and this effect was independent of the number of voices, suggesting that the increase in masking from one to two background voices, was not sufficient to prevent semantic processing. However, in this first experiment, prime was salient in both conditions (1SC voice and no SI voice or 2SC voices and no SI

voice). To test whether participants could still take advantage of the semantic relationship between target and prime if the intelligibility of the SC voices was further decreased, we conducted a second experiment in which a SI voice was added to each background.

3. Experiment 2

This second experiment aimed to investigate whether the semantic priming effect would resist increased masking. An SI voice was therefore added to each background. This voice pronounced words sharing no semantic features with each other or with the target word, whatever the condition. The purpose was to use the same material and procedure as in Experiment 1 with the addition of mask on the SC voices. In Experiment 2, backgrounds were composed of two voices (1 SC voice + 1 SI voice) or 3 voices (2 SC voices + 1 SI voice). A deleterious effect of the number of voices on participants' performance was predicted, and we expected that the presence of SI voice would not affect the semantic priming effect if this latter effect is automatic.

Another change was made in Experiment 2 regarding target items. In Experiment 1, target items were pronounced by a female speaker, and were consequently, difficult to detect among the other female speakers (S₁ and S₂). These difficulties might partly explain the low accuracy and long response times to target words inserted in babbles that were only composed of one or two voices. To avoid flux segregation difficulties (Brungart et al., 2001; Festen & Plomp, 1990), target items were therefore pronounced by a male speaker (Target Speaker 2; TS₂) in the two following experiments.

3.1. Method

3.1.1. Participants

Twenty-four right-handed French native speakers (18 females), aged 18-30 years, participated in this second experiment. They had no known auditory or language disorders. All participants gave their written informed consent and were compensated for their participation. None of the participants had been tested in Experiment 1 and they were not aware of the aim of the study before testing.

3.1.2. Stimuli

To add an SI voice to each background used in Experiment 1, 96 groups of 5 words ($M_{\text{lexical frequency}} = 18.15$, $SD = 9.75$), not semantically related to each other, were generated. Average lexical frequency did not differ between SC voices (from Experiment 1) and the SI voice as shown by an ANOVA ($F(2,190) = 1.16$, *n.s.*). Each group was selected to mask a specific prime (composed of 1 or 2 SC voices), which

shared no semantic link (e.g., the prime *légume, chou, céleri, salade, tomate* ‘vegetable, cabbage, celery, lettuce, tomato’ was always presented with the S1 voice pronouncing *policier, intéressant, cour, affiche, étagère* ‘policeman, interesting, yard, poster, shelf’). This S1 voice was recorded by another French native female speaker (S3, age = 20) using the same method as in Experiment 1.

Backgrounds composed of 1 SC voice (S1) in Experiment 1 were now composed of 1 SC voice and 1 SI voice (S1+S3), corresponding to the 1SC/1SI condition, and backgrounds composed of 2 SC voices (S1+S2) in Experiment 1 were now composed of 2 SC voices and 1 SI voice (S1+S2+S3), corresponding to the 2SC/1SI condition. The 4 groups of 12 target words and pseudo-words created for Experiment 1 were used. Target words were presented in each of the 4 conditions: 1SC/1SI related, 1SC/1SI unrelated, 2SC/1SI related and 2SC/1SI unrelated. The corresponding number of pseudo-words was also presented. Four experimental lists were created so that each target word was seen in each condition but only once in a list.

Recordings of S1 and S2, used in the previous experiment, were mixed with S3 following the previously established experimental lists. Targets were recorded by a French native male speaker (Target Speaker 2; age = 20) and were inserted into backgrounds 2 sec after the start of the sequence (see Figure 24).

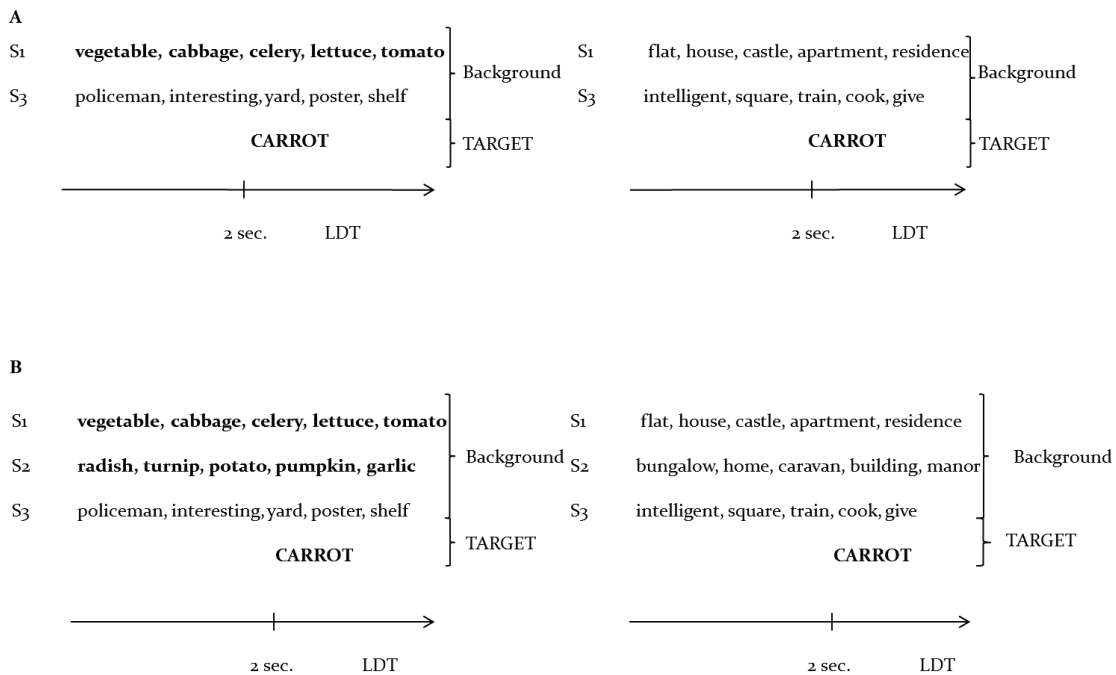


Figure 24

(a) Example of a background in the 1SC/1SI condition of Experiment 2, presented with a semantically related target word (left; related condition) or not (right; unrelated condition). S3 = speaker 3 (see the legend of Figure 23 for other abbreviations). (b) Example of a background in the 2SC/1SI condition of

Experiment 2, presented with a semantically related target word (left; related condition) or not (right; unrelated condition).

3.1.3. Procedure

The procedure was the same as in Experiment 1.

3.2. Results

Similar analyses as in Experiment 1 were performed by subjects (F_1) and by items (F_2), with Response Times (RTs, in ms) and Error Rates (ERs) for target word identification as the dependent variables. We included Number of Voices in the background (two voices, 1SC/1SI or three voices, 2SC/1SI) and Semantic Link between prime and target (related or unrelated) as within-subjects factors. As in Experiment 1, target words for which less than 50% of participants answered correctly were not analyzed. The target words *BILLET* and *CHIGNON* ('ticket' and 'bun') were therefore not included. Moreover, 15% of data were excluded (13% of errors and 2% of extremes values) from RT analysis. Mean RTs and ERs with SDs are summarized in Table 2.

		1SC/1SI		2SC/1SI	
		Related	Unrelated	Related	Unrelated
RT (ms)	Mean	933	958	945	1000
	SD	165	153	162	187
ER (%)	Mean	6.47	10.7	10.7	14.33
	SD	7.29	9.14	9.32	8.78

Table 2

Means and SDs of RTs and ERs depending of the Number of Voices in the background and the Semantic Link between prime and target in Experiment 2. 1SC/1SI = 1 SC voice and 1 SI voice condition; 2SC/1SI = 2 SC voices and 1 SI voice condition.

The two-way repeated measures ANOVAs showed a significant main effect of the Number of Voices in the backgrounds ($F_{1(1,23)} = 6.08$, $p < .05$; $F_{2(1,45)} = 4.34$, $p < .05$). Participants responded faster to target words if backgrounds comprised 2 voices ($M = 946$ ms, $SD = 157$) compared to 3 voices ($M = 973$ ms, $SD = 175$). This main effect was also significant for ERs ($F_{1(1,23)} = 7.18$, $p < .05$) but only in the participants' analysis ($F_{2(1,45)} = 3.72$, $p < .10$). Responses were more accurate in the 1SC/1SI condition ($M = 8.6\%$, $SD = 8.2$) than in the 2SC/1SI condition ($M = 12.5\%$, $SD = 9.1$).

The main effect of Semantic Link was also significant ($F_{1(1,23)} = 5.32$, $p < .05$; $F_{2(1,45)} = 5.38$, $p < .05$), responses were 40 ms faster if the prime and target shared semantic features ($M = 939$ ms, $SD = 162$) compared to if they did not share features ($M = 979$

ms, SD = 170). This effect also reached significance for ERs in the participants' analysis ($F_{1(1,23)} = 5.33$, $p < .05$; $F_{2(1,45)} = 1.98$, $p = .10$). ERs decreased by 3.9% if the prime and target were semantically related ($M = 8.61\%$, $SD = 8.5$ in the related condition and $M = 12.53\%$, $SD = 8.8$ in the unrelated condition).

There was no significant interaction between the Number of Voices in the backgrounds and the Semantic Link between prime and target for RTs ($F_{1(1,23)} = 1.92$, $n.s.$; $F_{2 < 1}$) and ERs ($F_{1 < 1}$; $F_{2 < 1}$), indicating that the priming effect was not affected by the increase in the number of voices in the backgrounds.

3.3. Discussion

As in Experiment 1, performances improved if the prime and target were semantically related, although participants were slower in condition 2SC/1SI than 1SC/1SI because of increased masking (Bronkhorst, 2000; Brungart, 2001; Brungart et al., 2001). Interestingly, there was again no indication that increased masking reduced the semantic priming effect. To further test this resistance of the semantic priming effect to prime intelligibility loss, we conducted a third experiment in which a second SI voice was added to each background to further decrease prime saliency. However, the intelligibility of the target item may decrease with the addition of a second SI voice in the background. However, as the target item is pronounced by a male voice and embedded in a female voice background, it is still quite salient compared to background voices (Brungart et al., 2001). Indeed, Brungart and collaborators showed that participants more easily recognize a target sentence if it was embedded in a background composed of voices of a different sex than if all the sentences (background and target) were pronounced by speakers of the same sex. Additionally, in our experiment, target items were presented at the same time (2 sec after the beginning of the stimulus), and this regular timing helps participants to detect the target, as they know when to listen to it. Consequently, in our paradigm, the masking effect of the SI voice on the target item was quite low compared to its effect on SC voices.

4. Experiment 3

This last experiment was the same as Experiment 2, except that a second SI voice was added to each background (i.e., backgrounds comprised 3 or 4 voices). The same method was used as in Experiment 2. The aim was to further increase the masking of SC voices to test the resistance and automaticity of semantic processing. As the number of voices in the backgrounds increased (i.e., up to 4 voices), we wanted to make sure whether participants could still identify words from the background. Therefore, after the main experiment, we asked participants to perform a recognition test. This post-test was designed to examine whether, in agreement with Hoen and collaborators (2007), words in the background were still intelligible. It aimed to clarify,

in the case of significant semantic priming, if this effect resulted from preserved intelligibility of the prime words by testing lexical access or from automatic processing. In case of preserved intelligibility we cannot prove that the semantic effect results from automatic processing, it may be either automatic or strategic. However, if a priming effect is found without preserved intelligibility, semantic processing is automatic (as shown with masked priming paradigm in visual modality, see Dehaene et al., 1998; Naccache & Dehaene, 2001; Spruyt et al., 2012). Proof of non-intelligibility was consequently needed in the case of a significant semantic priming effect to provide evidence for an automatic process. Otherwise, we would not be able to detect whether automatic components are present in semantic processes. As in our previous experiments, we expected to observe a significant effect of the number of voices in the background; increasing the number of voices has shown to decrease intelligibility for up to 8 voices (Simpson & Cooke, 2005). No interaction between the number of voices and the semantic association between prime and target was expected, at least if the words composing the background were still intelligible.

4.1. Method

4.1.1. Participants

Twenty-four participants (19 females) were recruited for this experiment (age: 18-34). All were right-handed French native speakers and reported no hearing or speech disorder. They gave written informed consent and were compensated for their participation. None of them had participated in Experiments 1 or 2, and they were unaware of the experiment's purpose prior to testing.

4.1.2. Stimuli

To create an extra SI voice, pronounced by a fourth speaker (S4), 96 groups of 5 words, that were semantically unrelated to each other, were generated ($M_{\text{lexical frequency}} = 20.88$, $SD = 1.22$). The word mean frequency did not differ between the different voices (SC and SI voices; $F < 1$). As in Experiment 2, each group of words was matched to a prime (composed of 1 or 2 SC voices) with which it did not share any semantic features and was systematically presented with (e.g., the prime *légume, chou, céleri, salade, tomate* 'vegetable, cabbage, celery, lettuce, tomato' was always presented with the masker *étui, liberté, drôle, global, sympathie* 'case, freedom, funny, global, sympathy'). Therefore, with the addition of this second SI voice, backgrounds composed of 2 voices (S1+S3) from Experiment 2 became 3-voice backgrounds (S1+S3+S4; 1 SC voice and 2 SI voices; 1SC/2SI condition) and 3-voice backgrounds (S1+S2+S3) from Experiment 2 became 4-voice backgrounds (S1+S2+S3+S4; 2 SC voices and 2 SI voices; 2SC/2SI condition). The 4 groups of 12 target words and pseudo-words created in Experiment 1 were used and presented in the following conditions: 1SC/2SI related, 1SC/2SI unrelated, 2SC/2SI

related and 2SC/2SI unrelated. Four experimental lists were created so that each target word was presented in each condition but only once in a list.

The second SI voice (S4) was recorded by a French native female speaker (age = 23), using the same procedure as in previous experiments. S1, S2 and S3's auditory sequences were mixed with S4's. Targets used in Experiment 2 (recorded by TS2) were embedded in the backgrounds 2 sec after their beginning start (see Figure 25 and Figure 26)

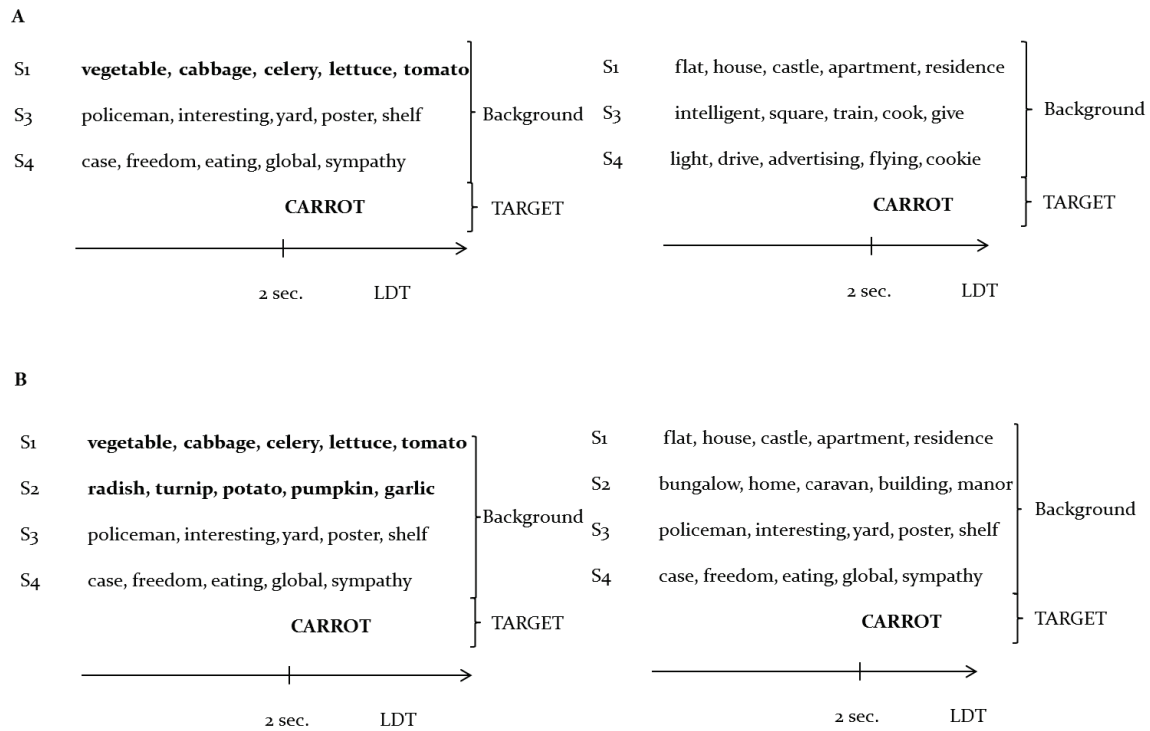


Figure 25

(a) Example of a background in the 1SC/2SI condition of Experiment 3, presented with a semantically related target word (left; related condition) or not (right; unrelated condition). S4 = speaker 4 (see the legend of Figure 23 and Figure 24 for other abbreviations). (b) Example of a background in the 2SC/2SI condition of Experiment 3, presented with a semantically related target word (left; related condition) or not (right; unrelated condition).

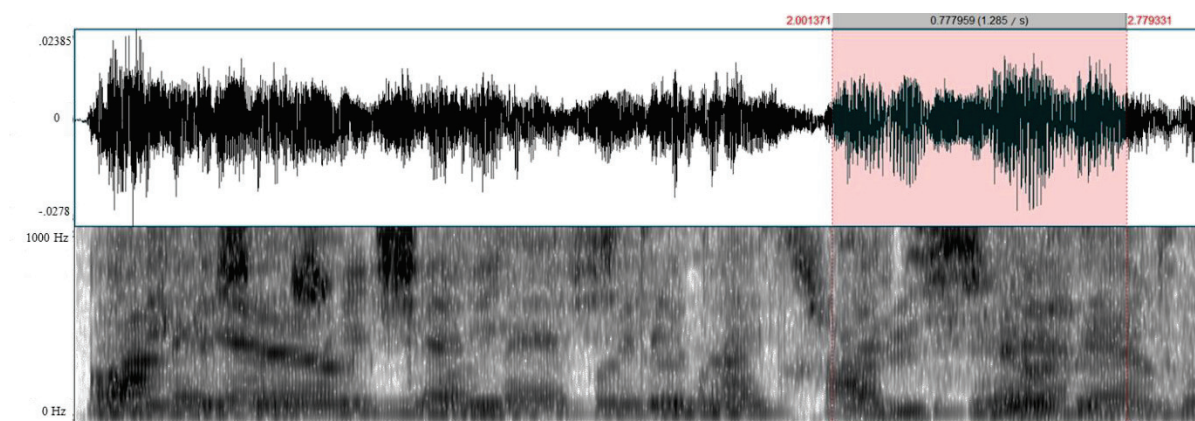


Figure 26

Example of a spectrogram and a wave form of a 2SC/2SI condition stimulus. Red part corresponds to the occurrence of the target word.

4.1.3. Recognition post-test

A recognition test was devised to test whether participants could recognize words previously heard during the experiment. Fifty words were presented to participants after the main experiment on a sheet of paper, 20 had been previously presented in the backgrounds, whereas 30 were new words not used as stimuli in the experiment. A list of new words was generated ($M_{\text{lexical frequency}} = 23.15$, $SD = 48.29$), and their lexical frequency did not significantly differ from that of previously heard words ($M_{\text{lexical frequency}} = 18.47$, $SD = 20.61$; $F < 1$).

4.1.4. Procedure

The same procedure was used as in Experiments 1 and 2. At the end of the experiment, a post-test was given to participants, instructing them to decide (i.e., write down) whether they had heard the given word during the experiment. They were asked to do the best they could but not to think about it too much, and to simply answer what they thought was right.

4.2. Results

4.2.1. Test

The same statistical analyses as in Experiments 1 and 2 were performed by participants (F_1) and by items (F_2), with Response Times (RTs, in ms) and Error Rates (ERs) for target word identification as dependent variables. We included Number of Voices in the background (3 voices, 1SC/2SI condition or 4 voices, 2SC/2SI condition) and Semantic Link between prime and target (related or unrelated) as within-subjects

factors. Target words answered correctly by less than half of participants were excluded from analyses (*CHIGNON*, *EPAULE*, and *HIBOU*, ‘bun, shoulder, owl’). In total, 16.8% of data were excluded from RTs analysis because of errors (15%) or extreme values (1.8%). Means and SDs for RTs and ERs are summarized in Table 3.

		1SC/2SI		2SC/2SI	
		Related	Unrelated	Related	Unrelated
RT (ms)	Mean	924	928	928	953
	SD	157	158	156	157
ER (%)	Mean	7.03	11.1	13.00	15.27
	SD	7.30	9.64	11.69	7.88

Table 3

Means and SDs of RTs and ERs depending on the Number of Voices in the background and the Semantic Link between prime and target in Experiment 3. 2SC/1SI = 1 SC voice and 2 SI voices condition; 2SC/2SI = 2 SC voices and 2 SI voices condition.

The analysis by participants revealed a significant main effect of the Number of Voices on RTs ($F_{(1,23)} = 4.01$, $p = .05$; $F_{2<1}$). Participants more quickly identified target words if backgrounds were composed of 3 voices (1SC/2SI condition; $M = 925$ ms, $SD = 155$) than if they were composed of 4 voices (2SC/2SI condition; $M = 941$ ms, $SD = 155$). The ERs analysis also showed that participants responded significantly more accurately ($F_{(1,23)} = 6.19$, $p < .05$; $F_{2(1,44)} = 5.80$, $p < .05$) in the 1SC/2SI condition ($M = 10.02\%$, $SD = 10.10$) than in the 2SC/2SI condition ($M = 13.21\%$, $SD = 8.95$). The main effect of Semantic Link was not significant for either RTs ($F_{(1,23)} = 1.51$, $n.s.$; $F_{2(1,44)} = 3.21$, $n.s.$) or ERs ($F_{(1,23)} = 3.19$, $n.s.$; $F_{2(1,44)} = 2.92$, $n.s.$). No significant interaction emerged between the 2 factors for RTs ($F_{(1,23)} = 1.39$, $n.s.$; $F_{2(1,44)} = 1.10$, $n.s.$) and ERs ($F_{1<1}$; $F_{2<1}$). These last results indicate that participants performed better in the lexical decision task, both in terms of speed and accuracy, if the backgrounds were composed of 3 voices rather than 4 voices. This finding highlights the impact of masking on target recognition. However, no significant semantic priming effect was observed in these conditions, suggesting that informational masking from 1SC/2SI and 2SC/2SI backgrounds was so efficient that it prevented semantic processing of the prime, which could therefore not affect target word identification. Additionally, in the 2SC/2SI condition, energetic masking was more important; this factor might also explain the important masking of prime and explain the lack of semantic priming.

4.2.2. Post-test

Analysis of the participants' answers showed a mean ER of 44.5% ($SD = 7$) and $d' = 0.26$ ($SD = 0.5$). Although these results are close to chance, both were significant results, as shown by a one-sample t test (ER $p < .01$; $d' p < .05$). No significant correlation was found

between d' and priming effect ($r = -0.39$, *n.s.*). This finding implies that participants heard some background words, but those words were not used to improve performances. Additionally, a repeated measure ANOVA was performed with ERs as the dependent variable and number of Voices (3 or 4) in the background and Semantic Link between presented word and target as a within subjects factor. Neither the effect of the Number of Voices nor the effect of Semantic Link were significant ($F_{\text{Number of Voices}} < 1$; $F_{\text{Semantic Link}} < 1$).

4.3. Discussion

In this third experiment, target word identification was again disturbed by the increase in the number of voices in the backgrounds, confirming that masking is more efficient in the 4-voice than in the 3-voice condition. Disappearance of the semantic priming effect also suggests that the number of SC voices compared to the number of SI voices was too small. To compare the data of the 3 experiments, we considered the ratio of SC voices over the total number of voices. Therefore, the ratio of SC/total voices is 1 in 1SC and 2SC conditions; ratio 2/3 in 2SC/1SI; ratio 1/2 in 1SC/2SI and 2SC/4SI; and ratio 1/3 in 1SC/2SI. An ANCOVA on all data using the number of voices as the independent variable and the ratio of SC voices/total voices as covariate on priming effect confirmed this hypothesis (effect of ratio: $p < .05$): when the ratio was too low, SC voices were not salient enough for participants to perform semantic processing. In a post-test, participants were asked to perform a recognition task immediately after the experiment. Participants scored significantly better than chance at the recognition test, showing that at least some words in the background were identified. This finding is consistent with the results by Hoen et al. (2007) who showed that in a transcription task, with up to 4 voices in the background, participants gave words from the background as responses instead of target words. Overall Experiment 3 showed that participants were unable to process the prime at a semantic level (i.e., no priming effect was observed) although the post-test results suggest that they hear it sufficiently to recognize it in a recognition post-test. This finding suggests that a word can be heard and implicitly encoded without being sufficiently deeply processed to elicit semantic priming.

5. Post-hoc analyses

To better analyze the impact of the SC/total voices ratio, we conducted post-hoc analyses. A HSD Tukey test showed that up to 4 voices in the background if the ratio of SC/total voices was inferior to 1/2, there was no significant semantic priming effect (Exp 2. 1SC/1SI; Exp 3. 1SC/2SI and 2SC/2SI). However, if the ratio was superior to 1/2, semantic priming was significant (Exp 1: 1SC, $p < .05$ Cohen's $d = 0.31$; 2SC, $p < .05$, Cohen's $d = 0.35$; Exp 2. 2SC/1SI, $p < .05$, Cohen's $d = 0.31$; cf. Figure 27).

Ratio	1		2/3	1/2		1/3
Condition	1SC (EXP ₁)	2SC (EXP ₁)	2SC/1SI (EXP ₂)	1SC/1SI (EXP ₂)	2SC/2SI (EXP ₃)	1SC/2SI (EXP ₃)
Mean	55	56	55	25	26	4
SD	70	119	110	88	83	65
p	0.03*	0.03*	0.007*	0.88	0.21	0.99

Table 4

Mean priming effect (in ms), SD, and p value for each ratio. * indicates significant priming. The ratio is expressed in terms of the number of SC voices over the total number of voices in the background: 1 = 1 SC voice; 2/3 = 2 SC voices, 1 SI voice; 1/2 = 1 SC voice, 1 SI voice; 1/3 = 1 SC voice, 2 SI voices.

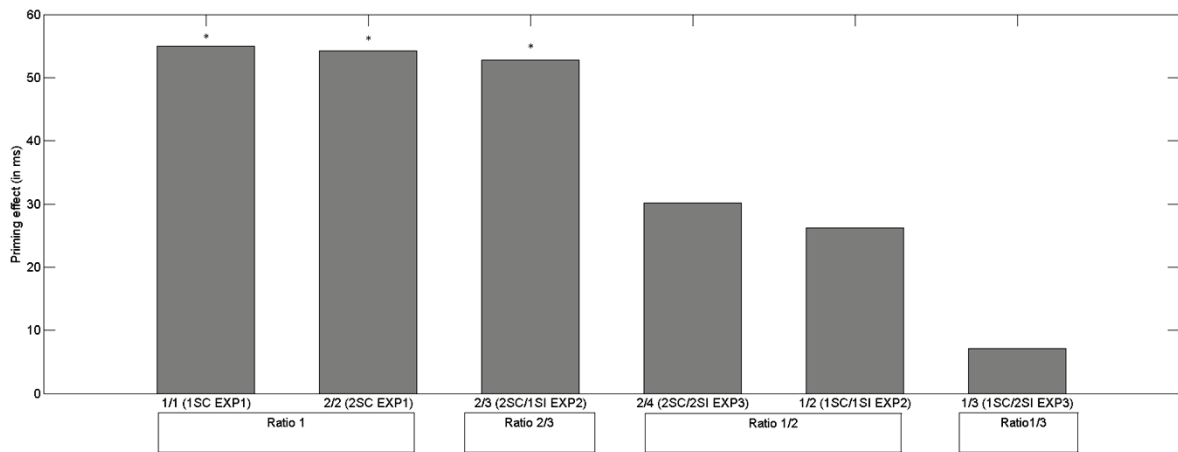


Figure 27

Semantic priming effect (i.e., difference between RTs for unrelated and related targets; in ms) depending on the ratio of SC voices in the background over the total number of voices. 1SC EXP₁ = 1SC condition in Experiment 1 (1 SC voice; no SI voice); 2SC EXP₁ = 2SC condition in Experiment 1 (2 SC voices; no SI voice); 2SC/1SI EXP₂ = 2SC/1SI condition in Experiment 2 (2 SC voices; 1 SI voice); 1SC/1SI EXP₂ = 1SC/1SI condition in Experiment 2 (1 SC voice; 1 SI voice); and 2SC/2SI EXP₃ = 2SC/2SI condition in Experiment 3 (2 SC voices; 2 SI voices); 1SC/2SI EXP₃ = 1SC/2SI condition in Experiment 3 (1 SC voice; 2 SI voices).

6. General Discussion

6.1. Summary

The aim of this study was to investigate to what extent the semantic content of the multi-talker babble could interfere with the processing of targets in a cocktail party situation. In the three reported experiments, participants were required to perform a lexical decision task on a target item embedded in a multi-talker background. The backgrounds were composed of 1 or 2 SC voices pronouncing words that shared semantic features with each other and could be semantically related to the target or

not. The ratio of SC voices over the total number of voices in the background was varied across experiments. In Experiment 1, ratios were 1/1 and 2/2 (i.e., only SC voices were presented in the background). In Experiments 2 and 3, 1 and 2 SI voices, respectively, that pronounced semantically unrelated words, were added to backgrounds to decrease the intelligibility of the SC voices, which acted as the prime. The ratio of SC/total voices in Experiment 2 was therefore 1/2 and 2/3 and in Experiment 3 it was 1/3 and 2/4.

The main effect of the Number of Voices was significant in each experiment; participants responded faster and more accurately to target words if a smaller number of voices composed the backgrounds. This delayed response time with increasing number of voices is a well-known phenomenon (Boulenger et al., 2010; Brungart et al., 2001). As the number of voices increases, energetic masking is enhanced so that the signal is saturated, and target items are consequently more difficult to process. In addition to physical masking, more information is perceived and must be processed (i.e., informational masking), leading to slower response times.

Overall the main effect of Semantic Link between prime and target seems to be an all-or-none phenomenon. The semantic priming effect could have decreased with an increasing number of voices in the background, but our results suggest this effect does not occur, as no interaction between the Number of Voices and the Semantic Link was observed in any of the three experiments. The main effect of Semantic Link depends on the ratio of SC voices over the total number of voices in the background. A semantic priming effect emerged in our experiments spanning from 1 to 4 voices in the background but only if the number of SC voices was higher than the number of SI voices. This ratio of SC voices is interesting because it highlights the necessity of prime saliency for semantic processing to occur; it also gives an objective measure of this prime saliency across experiments. This finding suggests that informational masking does occur at the semantic level but only if the prime is sufficiently salient. If informational and energetic masking had the same role, a significant semantic priming effect would appear in conditions containing 3 voices (i.e., Exp 2 2SC/1SI and Exp 3 1SC/2SI). However, this effect did not occur; we therefore conclude that informational masking can be semantic only if it is sufficiently salient to be used to increase performance.

6.2. Lexical processing without semantic activation

Our results suggest that semantic processing in cocktail party situation is not automatic. In fact, it seems that in challenging listening situations, one can hear and activate a mental representation of a word without deeply processing it at a semantic level. Background words can be semantically processed with up to 3 voices (2SC/1SI, Experiment 2); however no priming effect emerged in Experiment 3, despite the

recognition post-test's results suggesting that primes were heard and recognized. This finding is consistent with the way the language system has been modeled; most models of word processing, both in the auditory and the visual modalities, suggest a distinction between lexical and semantic levels of processing (Grainger & Holcomb, 2009; Marslen-Wilson & Welsh, 1978; McClelland & Elman, 1986). If one considers that these processes are independent stages of word recognition, it seems reasonable that one of these stages can be reached (i.e., lexical) without strongly activating the deepest stage (i.e., semantic).

Many studies have highlighted the automaticity of semantic processing using masked semantic priming (Klauer, Eder, Greenwald, & Abrams, 2007; Kouider & Dehaene, 2009; Naccache & Dehaene, 2001; Spruyt et al., 2012). However, masked semantic priming is only found in very specific conditions in the visual modality and with semantic categorization tasks (see Van den Bussche, Van den Noortgate, & Reynvoet, 2009, for a review). One can therefore assume that semantic activation remains superficial in masked priming paradigm, and only activates very close concepts such as superordinates, which is enough to create a priming effect in semantic categorization tasks. In our experiments however, a deeper semantic processing was necessary. For example, in the semantic field of birds, the SC voice pronounced: *corbeau*, *rossignol*, *cage*, *voler*, *nid* ('craw, nightingale, cage, fly, nest') and *PIGEON* ('pigeon') was the target word. In a condition of high intelligibility, these words primed the target word; however, with decreased intelligibility, if the participants only heard the word 'cage', this word may have activated its superordinate (e.g., 'object') but not its associates such as 'bird'. Consequently, it may not have primed 'pigeon'. The absence of a semantic priming effect in addition to the decrease in intelligibility seems to show that greater difficulty in processing auditory signals at a superficial level, causes decreased processing at a (higher) linguistic level.

Consistent with our results, and using an auditory masked priming paradigm, Kouider and Dupoux (2005) demonstrated lexical access without semantic priming. In their experiments, prime was a time-compressed word embedded in a masker composed of time reversed and compressed words. They manipulated the compression rate of the prime to vary its intelligibility. If the prime was not intelligible, they found a significant repetition priming effect on words but not non-words suggesting that this effect 'involved a lexical activation of abstract word form'. In the same condition of prime compression, they did not find any semantic priming effect. This finding would therefore suggest a dissociation between lexical and semantic processing.

6.3. Cognitive load

Overall our results are compatible with the idea of cognitive load, in which simultaneously processed information and interactions can either under-load or

overload the finite amount of processing capacities. As shown by our results, it is more difficult to process items with more voices in the background. This finding might be partly due to the high perceptual load and because the target item was embedded in backgrounds that could not be completely ignored (Lavie, 1995, 2005). Processing words in the background might have been particularly cognitively effortful and, semantic processing of a word can be delayed and even prevented if participants have to perform an additional task (i.e., high cognitive load, Hohlfeld, Sangals & Sommer, 2004; Hohlfeld & Sommer, 2005; Van Petten, 2014). In Experiment 3 we argue that some background words were heard but not deeply semantically processed because it was both very demanding and irrelevant to perform the task. In Experiment 2, 2SC/1SI background words were also very difficult to hear and process (as in Experiment 3 1SC/2SI); however, they were semantically processed as revealed by the semantic priming effect observed. As SC voices were more salient in Experiment 2 2SC/1SI than in Experiment 3 1SC/2SI, participants might have heard more related words and one could argue that semantic priming in our study in fact relied on the chance for participants to hear a SC word. Although this hypothesis is interesting, it does not seem sufficient to explain our results. Indeed, according to this hypothesis, a priming effect should have appeared also in Experiment 3 where participants, in line with previous results from the literature (Brungart, 2001; Hoen et al., 2007), recognized SC words presented in the post-test. Given the overall results of our experiments, we argue that if the ratio of SC/total voices was $<1/2$, participants heard some SC words, but, because of high cognitive load (i.e., intelligibility was low and they focused on the target item) these words were not processed sufficiently deeply to lead to semantic priming (a similar dissociation between lexical and semantic processing was also found by Kouider and Dupoux, 2005). However, if the ratio was $>1/2$, SC words were salient, and participants thus, processed them sufficiently to improve their performance. As this interpretation relies on the post-test effect which is quite small, although significant, more experiments should be performed. For example, using only SC voices and degrading intelligibility by adding noise or filtering the signal instead of adding SI voice would be a good way to test this hypothesis in future studies.

Our results suggest that the increased cognitive load necessary to reconstruct the degraded signal reduced available resources for higher level processes. This claim is consistent with the Effortfulness Hypothesis (Rabbitt, 1968) that states degraded signals require allocating many cognitive resources to formal processes (i.e., orthographic or phonological), leaving less available cognitive resources to perform higher level processes (e.g., lexical). It has been shown that hearing-impaired participants are less accurate at recalling previously heard final sentence words than their control peers (Pichora-Fuller et al., 1995). The underlying assumption is that for hearing-impaired participants, the auditory signal is highly degraded and therefore demands more cognitive resources to be formally (i.e., phonologically) processed. As

studies usually use recall tasks of lists of unrelated words or digits (Murphy et al., 2000; Surprenant, 1999; Wingfield et al., 2005), verbal working memory only relies on the phonological loop (Baddeley, 2000; Baddeley & Hitch, 1974) and only involves formal processes. Our results showed that higher levels of processing such as semantic activation may be specifically impacted by signal degradation.

Recently, semantic and syntactic integration difficulties if a signal is degraded have been reported in the visual modality (Gao et al., 2012; Gao et al., 2011). Experiments conducted by Gao et al. (2011) using visual noise (i.e., pixel's brightness variation) showed that if participants allocate more resources to formal processes, semantic integration is affected. Indeed, after reading an entire text, participants were worse at recalling the main proposition in the noisy condition. Altogether, these findings suggest that the availability of cognitive resources is involved at various levels during language processing. Whereas previous studies have shown that noise impairs memory systems, our study provides evidence that semantic activation is linked to cognitive resources, independently of memory.

7. Conclusion

This study explored the semantic nature of informational masking in a cocktail party situation. The results of three behavioral studies reveal that the emergence of semantic priming effects relies on prime intelligibility and saliency. These findings question the assumption that signal degradation has no effect on speech processing if target signals can be recognized. The results reveal that high level processes, such as semantic processing, might not be as automatic as previously thought but are subjected to the limits of cognitive resources. Our study also demonstrates how the cocktail party situation can be used to study the automaticity of linguistic processes.

Acknowledgements

The first author is funded by a PhD grant from Rhône-Alpes region, France. This research was supported by a European Research Council grant to the SpiN project (no. 209234).

8. Résumé du Chapitre VI

L'objectif de cette première étude était d'évaluer l'impact de la dégradation du signal de parole sur le traitement sémantique.

Nous avons pu mettre en évidence un effet d'amorçage sémantique significatif lorsque le nombre de voix SC était inférieur au nombre de voix SI, au moins jusqu'à 4 voix présentes dans le fond sonore. C'est-à-dire que les participants étaient plus rapides à catégoriser l'item cible comme mot lorsque celui-ci partageait un lien sémantique avec les mots prononcés par les voix SC. Toutefois, cet effet n'apparaissait que lorsque celles-ci étaient en nombre plus important que les voix SI. Il est apparu que l'effet d'amorçage sémantique ne dépendait pas uniquement de l'intelligibilité des mots prononcés par les voix SC. En effet, dans l'Expérience 3 aucun effet d'amorçage n'apparaissait alors que les participants étaient capables de reconnaître les mots composant le fond sonore parmi des distracteurs. Ceci implique donc que les participants étaient encore capables d'entendre et de reconnaître les mots composant les fonds sonores même lorsque ceux-ci étaient composés de 3 et 4 voix. Ces résultats suggèrent ainsi, comme postulé par l'*Effortfulness Hypothesis*, que la dégradation du signal de parole a un impact sur les traitements haut niveau, notamment le traitement sémantique. Si celui-ci est effectivement dépendant de la disponibilité des ressources cognitives, alors il semble qu'il corresponde davantage aux processus stratégiques décrits par Posner et Snyder (1975) qu'à un mécanisme purement automatique de propagation automatique de l'activation (Collins et Loftus, 1975).

Chapitre VII

Etude 2 : Le traitement des mots ambigus en situation de parole dans le bruit

Cette deuxième étude a été effectuée en collaboration avec la Professeure Stemmer au Centre de Recherche de l'Institut Universitaire Gériatrique de Montréal, où j'ai passé 6 mois (séjour financé par une bourse Explora Doc de la Région Rhône Alpes).

L'Etude 1 a permis de mettre en évidence que le traitement sémantique est modulé par la présence de bruit. Notamment, les résultats suggèrent qu'un traitement lexical est effectué sur les mots composant le fond sonore mais que le traitement sémantique est trop superficiel dans certaines conditions pour qu'apparaisse un effet d'amorçage. Dans cette deuxième étude, nous voulons affiner l'étude de la modulation du traitement sémantique en utilisant les mots ambigus. Nous posons l'hypothèse que l'activation des différents sens d'un mot ambigu peut être modulée par la présence de bruit. Cette étude vise également à évaluer l'interaction relative entre les traitements *bottom-up* et *top-down* en fonction de la dégradation du signal.

Dans cet objectif, des mots ambigus ont été présentés au sein de phrases composées de deux propositions indépendantes. La première proposition présentait le mot ambigu dans son sens dominant (i.e., le plus fréquent). La seconde proposition contenait le mot cible en dernière position. Celui-ci pouvait faire référence au sens dominant du mot ambigu (i.e., condition Congruent), au sens subordonné du mot ambigu (i.e., condition Ambigüe) ou à aucun des sens du mot ambigu (i.e., condition Incongruent). Comme présenté dans la partie Théorique, la présentation d'un mot ambigu est supposée générer l'activation relative de ses différents sens. Particulièrement, s'il est présenté au sein d'un contexte favorisant son sens dominant, son (ses) sens subordonné(s) seront eux aussi activés (cf. p53). Enfin, afin de varier la dégradation du signal, la phrase pouvait être présentée dans le silence (e.g., condition Silence), dans

du bruit stationnaire (e.g., condition Bruit Stationnaire) ou dans du bruit parolier (e.g., condition Bruit Parolier). Le bruit stationnaire ne génère que du masquage énergétique alors que le bruit parolier génère principalement du masquage informationnel. Nous verrons donc l'effet relatif de ces différents types de masquages sur le traitement sémantique. Dans toutes les conditions, le mot cible était présenté dans le silence afin de s'assurer que celui-ci soit traité de façon optimale. Le but de cette étude est (1) d'évaluer la force d'activation du traitement sémantique du mot ambigu en fonction du bruit et (2) d'étudier la modulation des interactions entre les processus *bottom-up* et *top-down* en fonction du bruit.

Semantic processing of ambiguous word in speech in noise situation

Marie Dekerle ^{a, b}, Brigitte Stemmer ^{c, d}, Fanny Meunier ^{a, b}

(a) Laboratory on Brain, Language and Cognition (L2C2) CNRS, UMR 5304, Lyon France

(b) Université de Lyon

(c) Centre de Recherche de l'Institut Universitaire Gériatrique de Montréal, Montréal, Québec, Canada

(d) Université de Montréal

ABSTRACT

This paper aims at investigating the modulation of semantic processing of a word in its context in speech in noise situation. To this purpose, ambiguous words were presented in sentences. The first part of the sentence referred to the dominant (i.e., the more frequent) meaning of the ambiguous word and contained it. The second part ended by a target word, referring whether to the dominant meaning of the ambiguous word (i.e., Congruent condition), to its subordinate meaning (i.e., Ambiguous condition) or to none meaning (i.e., Incongruent condition). In addition, sentences were embedded in different kind of noise (i.e., Silence, Stationary noise, Babble noise). According to the Effortfulness Hypothesis, this should vary the amount of cognitive resources allocated to the semantic processing of speech. The electrophysiological response to target word was recorded. It was hypothesized that N400 in response to Incongruent target word would be more important than in response to Congruent target word in all noise condition (i.e., congruency effect). N400 amplitude in response to ambiguous target word was supposed to be medium, reflecting the activation of both meaning of ambiguous word when semantic processing strength was sufficient (i.e., ambiguity effect). Results did not evidence any N400 effect in Babble condition suggesting that too many resources were allocated to low level processes to perform an efficient semantic processing despite preserved intelligibility. In Silence condition, congruency and ambiguity effect were both significant but did not differ from each other. In Silence condition, semantic processing of context might have been sufficiently strong that only the pertinent meaning of ambiguous word was activated. Finally, the expected effect was observed in Stationary condition only, with ambiguity effect being significantly stronger than zero and weaker than congruency effect. It therefore seems that in Stationary condition, context is not sufficiently processed to elicit strong predictions, therefore both meaning of ambiguous word are activated.

1. Introduction

In daily life speech is rarely perceived in silence, most of the time it is masked by other sounds like wind, music, or parallel conversations, this is called the cocktail party phenomenon (Cherry, 1953). Target speech must therefore be distinguished from background. Since Cherry's (1953) seminal paper, most of the literature got interested to this phenomenon on a psychoacoustic point of view, studying the gain and loss of intelligibility (Brungart, 2001; Brungart et al., 2001; Simpson & Cooke, 2005). However, recently, cocktail party phenomenon has also been used as a tool to tackle linguistic processing (Boulenger et al., 2010; Dekerle et al., 2014; Hoen et al., 2007; Obleser & Kotz, 2011; Strauss et al., 2013). The aim of this study was twofold: first, evaluating whether the addition of noise can influence the strength of activation of different meaning of ambiguous words and second, investigating the interaction between top-down and bottom-up processes.

1.1. Speech in noise situation

Speech in noise situation generates two main kind of masking: energetic masking and informational masking (Brungart, 2001; Brungart et al., 2006; Brungart et al., 2001; Cherry, 1953). Energetic masking refers to the overlap of sounds sharing the same spectrotemporal characteristics. They will consequently stimulate the same part of the cochlea at the same time, and therefore only one of them will be heard. This masking is purely acoustic, and in speech in speech situation, is proportional to the number of voices presented in the background, until the signal is saturated and the addition of another voice does not enhance energetic masking anymore (i.e., more than 8 voices; Simpson & Cooke, 2005).

Informational masking on the other hand results from higher level interferences (Durlach, 2006; Mattys et al., 2009; Mattys & Wiget, 2011; Simpson & Cooke, 2005). It is more important than energetic masking in speech in speech situation (Brungart, 2001). Indeed, interferences can occur at different level, whether phonological (Gautreau et al., 2013) lexical (Boulenger et al., 2010) and semantic (Brouwer et al., 2012; Dekerle et al., 2014). It is also more deleterious than energetic masking, indeed, Festen and Plomp (1990) evidenced that the needed SNR (Signal / Noise Ratio) to attain SRT (Speech Reception Threshold, corresponding to a 50% performance at recognition of speech in noise) was 6 dB to 8 dB higher in speech in speech situation than in speech in non-speech noise. In our study, stimuli are embedded in both noise and speech. This will first decrease the intelligibility of speech and allow investigating the consequences of the restrain intelligibility on semantic processing and second, allow to investigate the relative effects of different maskers

1.2. Effortfulness Hypothesis

As mentioned earlier, the majority of the literature on speech in noise gets interested to the impact of energetic and informational masking on target speech intelligibility. However, some studies also got interested in the impact of noise on high level processing despite intelligibility. The Effortfulness Hypothesis, (Rabbitt, 1968; Wingfield et al., 2005) suggests that when signal's quality is good, low level processing (i.e., acoustic and phonologic) are performed automatically, whereas high level processing (e.g., working memory, semantic processing) are subjected to cognitive resources. In adverse condition, when target signal is degraded, cognitive resources are reallocated to low level processing in order to compensate for signal degradation and fewer resources are therefore left to perform higher level processing. This reallocation of resources is reflected by a decrease in performances in noisy situation. Interestingly, some studies evidenced that a signal degradation because of presbycusis majored this phenomenon (Pichora-Fuller, 2003; Pichora-Fuller et al., 1995) highlighting the concern about ageing and speech in noise. Most studies interested in the Effortfulness Hypothesis focused on working memory performances and evidenced a decrease of them in noisy situation (i.e., remember a list of number or words; Murphy et al., 2000; Obleser, Wostmann, Hellbernd, Wilsch, & Maess, 2012; Surprenant, 1999). However, some studies also got interested on the impact of noise on language processing per se. For example, Gao et al. (2011) evidenced, in visual modality the impact of noisy signal (i.e., degradation of visibility by pixel's bright variation) on syntactic and semantic processes. Our precedent study also suggested that in speech in speech situation, a word can be heard and recognized without being sufficiently deeply processed at the semantic level to lead to a significant semantic priming (Dekerle et al., 2014). This paper will therefore go further in the investigation of semantic processing modulation.

1.3. Ambiguous words

This study aims at evaluating the impact of different type of masking on semantic processing, using ambiguous words, that is, words that have the same pronunciation but different meanings (e.g., bank-money vs bank-river vs bank-snow). Ambiguous words (i.e., homonyms, homophones, homographs) are commonly used in the literature to study the access of meaning (Foraker & Murphy, 2012; C. Martin, Vu, Kellas, & Metcalf, 1999; Rayner et al., 1994; Van Petten & Kutas, 1987; Vu et al., 1998; Wiley & Rayner, 2000). The main interest of the literature is about the activation of different meanings of an ambiguous word in a given context. In a neutral context the presentation of an ambiguous word generates the activation of its meanings depending on their frequency. Therefore, the dominant meaning will be more activated than subordinate meanings (e.g., *BANK* will generate a strong activation of *money* meaning and a weaker activation of *river* meaning and an even weaker activation of *snow* meaning).

However, when the context is not neutral, the way it influences the activation of the ambiguous word's different meanings is subject to two hypotheses.

(1) the Context-sensitive model (Vu et al., 1998) postulates that the effect of context is more important than the frequency of the meaning. This means that only the appropriate meaning will be activated in a given context, whatever its frequency.

(2) the Re-ordered access model (Duffy et al., 1988) suggests that all meanings of a word will be activated whatever the context. However, it suggests that the context will re-order the activation depending on the congruency of a meaning in the context, but will not inhibit the other meaning. Therefore, when presented in dominant context, the dominant meaning of the ambiguous word is more activated than the subordinate meaning. When presented in subordinate context however, activation of subordinate increases whereas the activation of dominant meaning remains the same.

In addition, this 2nd model gives an explanation to the subordinate bias (C. Martin et al., 1999; Rayner et al., 1994). It refers to the increased processing time of an ambiguous word as measured by the increased of reading time when it is presented in a context favoring the subordinate meaning. Indeed, when presented in a subordinate context both meanings are strongly activated. The competition between the two meanings is thus responsible for the increased time processing in subordinate context.

Recently, a link between context strength and subordinate bias has been evidenced (Colbert-Getz & Cook, 2013). Subordinate bias was found in weak context (i.e., one sentence preceding the ambiguous word) however, when context was strong (i.e., four sentences preceding the ambiguous word), the subordinate bias disappeared. This implies that when context is weak, both meanings are activated as suggested by the Re-ordered access model (Duffy et al., 1988). However, when context is strong, only the meaning congruent with the context is activated. Thus, when context is strong, meaning activation seems to be more Context-sensitive (Vu et al., 1998).

Therefore it suggests that the impact of context on relative meanings activation depends on its strength. These results explain and combine the two already presented models.

Following these findings, our study will use ambiguous words presented in a weak dominant context (i.e., in our study, a preposition), assuring that (1) both meanings of the ambiguous word are activated and (2) one meaning is more activated than the other (i.e., dominant meaning is more activated than subordinate meaning).

1.4. Semantic processing in sentence context

One main cue used to study semantic processing in context is the N400 Event Related Potential (ERP; Kutas & Federmeier, 2011; Kutas & Hillyard, 1980). The N400 is a well-

known and studied centro-parietal negative ERP, occurring 400 ms after the presentation of a target word. It is interpreted as reflecting the semantic integration of this word in its context. The more a word is attended in its context, the smaller the amplitude of the N400. It therefore depends on the restrictivity of the sentence. That is, when the last word is very attended (e.g., I will take my coffee with sugar and CREAM) the presentation of another word than the expected one will generate a N400 of higher amplitude. Interestingly it has also been evidenced that the presentation of an incongruent word semantically linked to the attended word generates a N400 of medium amplitude (i.e., higher than in response to the attended word but smaller than in response to an incongruent word non semantically linked to the attended word; Kutas & Hillyard, 1984). For example, the sentence *He liked lemon and sugar in his TEA* generated a N400 of small amplitude. The N400 amplitude was medium when the last word was *COFFEE* (i.e., semantically linked to the expected target *TEA*). However, if the target word was completely unrelated to the expected target (e.g., *DOG*), N400 amplitude was more important. This is explained by the fact that the expectation of the word *TEA* leads to its activation, which spreads to its related concepts (Collins & Loftus, 1975). The integration of a related context will thus be eased and this will be reflected by the relatively small amplitude of the N400. According to these results, in our study, the presentation of a target word linked to the subordinate meaning of the ambiguous word should lead to smaller N400 amplitude than an incongruent word if both meanings of the ambiguous word are activated.

N400 is also sensitive to the quality of signal, indeed, although some studies found N400 effect in subliminal priming paradigm (Deacon et al., 2000; Kiefer, 2002; Kiefer & Brendel, 2006), in auditory modality, N400 disappears/diminishes when intelligibility decreases (Aydelott et al., 2006; Boulenger et al., 2011; Strauss et al., 2013).

1.5. The present study

The aim of this study was twofold: first, investigate the relative activation of both meanings of an ambiguous word depending on the degradation of signal. Second, study the interactions of top-down and bottom-up processes when the last ones are impaired by noise. To this purpose, ambiguous words were embedded in sentences whom first part referred to the dominant meaning of the ambiguous word (e.g., *The businessman called from the bank...*). The second part of the sentence contained the target word (in response to which EEG signal were recorded). This target word could refer to the dominant meaning of the ambiguous word (e.g., *The businessman called from the bank, he needed MONEY*), or the subordinate meaning of the ambiguous word (e.g., *The businessman called from the bank, he needed BOATS*) or to neither meaning (e.g., *The businessman called from the bank, he needed TABLES*). As we are interested in the way different kind of masking influences semantic processing, sentences but not target word were embedded in stationary noise (i.e., energetic masking) or babble

noise (i.e., energetic masking + informational masking). They could also be presented in silence as control condition. Following the rationale of the Effortfulness Hypothesis and our previous study (Dekerle et al., 2014) we postulate that N400 amplitude in response to ambiguous target word will be smaller than in response to incongruent target but higher than in response to congruent target. These different responses should be attenuated in Stationary and Babble condition compared to Silence condition.

2. Material & Methods

2.1. Participants

Twenty participants were recruited for this experiment, all of them were native French-Canadian speaker, right handed as assessed by Edinburgh score ($M = 85.78$; $SD = 15.38$; Oldfield, 1971) and had no history of neurological or psychiatric disorder. As population with reading disorders are known to show speech in noise comprehension deficits (Bradlow et al., 2003; Dole et al., 2012; Dole, Meunier, & Hoen, 2013; Ziegler et al., 2009) participants were tested with the Alouette test (Levafrais, 1967) the number of Correctly Read Word per Minute (CRWM) was calculated ($M = 156.1$; $SD = 26$). Participants did not report any known hearing disorders. They gave written informed consent and were compensated for their participation. The protocol was approved by local ethics committee (CER-IUGM 13-14-035).

2.2. Stimuli

2.2.1. Ambiguous words

Two-hundred and seventy ambiguous words were selected from different sources (Ferrand, 1999; Larousse, 2013). When they were not homographs the frequency of both meanings was found in Lexique 3 (New, Pallier, Ferrand, & Matos, 2001) and the more frequent one was defined as the dominant meaning (e.g., *la mer* : (the sea) frequency = 106.6 vs. *la mère* (the mother) frequency = 686; therefore, mother meaning was considered as dominant meaning). However, for 106 ambiguous words both meanings had exactly the same orthography, same gender and same grammatical category (e.g., *un avocat* can mean a lawyer or an avocado). These 106 words were recorded by a French-Canadian native woman (age = 23) and then presented through the internet to 20 native French-Canadian participants. They had to write down for each of these words, the first word coming to their mind: in the case of *avocat* it could be for example *guacamole* or *court*. The number of words referring to each meaning was then calculated for each word, and the dominant meaning was the one which was the most referred to.

2.2.2. Sentences

After each ambiguous word's dominant meaning has been identified, 270 sentences were created. All of them were made of two independent propositions, the first one ending by the ambiguous word and referring to its dominant meaning. The second proposition ended systematically with the target word, which could refer to the dominant meaning of the ambiguous word (i.e., congruent condition; *Le joueur est resté dans le golf, en fait il a perdu son BATON* (The player stayed in the golf (subordinate meaning: golf), in fact he lost his CLUB)). It could refer to the subordinate meaning (i.e., ambiguous condition; *Près du tribunal j'ai vu des avocats, mais je n'aime pas le GUACAMOLE*; Near court I saw lawyers (subordinate meaning: avocados), but I do not like GUACAMOLE). Finally target word could be incongruent, that is, refer neither to the dominant nor the subordinate meaning (i.e., incongruent condition; *Au magasin j'ai acheté un duvet, mais je suis mieux avec une DIFFERENCE*; At the shop, I bought a sleeping bag (subordinate meaning: down), but I feel better with a DIFFERENCE).

To control for (a) cloze probability of each sentence, and (b) probability of appearance of target word, fifteen French-Canadian native speakers were recruited through the internet and were presented to the sentences except for final word (i.e., *Au magasin j'ai acheté un duvet mais je suis mieux avec une ...* ; At the shop, I bought a sleeping bag, but I feel better with a ...). The number of participants answering the same word for each sentence was registered. When the sentence was meant to be presented in the congruent condition, the most cited final word was chosen. The mean cloze probability of sentences was controlled across condition as shown by a one way ANOVA ($F < 1$).

2.2.3. Target words

Target words were all nouns. They were controlled across conditions for frequency ($M = 29.77$; $SD = 47.61$; $F < 1$), number of phonemes ($M = 5.1$; $SD = 1.91$; $F < 1$) and number of phonological neighbors ($M = 7.17$; $SD = 8.43$; $F < 1$) following Lexique 3 (New et al., 2001).

2.2.4. Recordings and Noises

2.2.4.1. Babble noise

Four French-Canadian native speakers were recorded (2 females; mean age = 27.25; $SD = 8.53$) in a quiet room of the laboratory. They read selected chapters of "The Little Prince" by Antoine de Saint Exupéry. This book was chosen as its vocabulary is neither childish nor refined. Recordings were then treated: silences and pauses exceeding 500 ms were removed as were sentences containing mispronunciations. Intensity was then calibrated at 70 dB-A. Recordings were then cut in 5 sec chunks using Praat software. Chunks were then mixed using Matlab software, 90 babble noises made of 4

voices were thus created. Previous studies (Gautreau et al., 2013; Hoen et al., 2007) suggested that in 4 voices babbles, informational masking is still efficient as phonological and lexical information are still processed although it might not be the case for semantic information (Dekerle et al., 2014).

2.2.4.2. Stationary noise

Stationary noise was derived from the previously created babbles. Energy root mean square (rms) of the original sample was computed, and temporal envelope was extracted by applying a 60 Hz low pass filter. A Fast Fourier Transform (FFT) was then used to extract the power spectrum and phase distribution of the original signals. Original phase transformation was randomized. Finally, an inverse FFT generated new signals that were normalized to the rms power of the original sample.

2.2.4.3. Target sentences

Target sentences were recorded by a female French-Canadian (age = 31) in a quiet room of the laboratory. Recordings were cut so that there was an audio file for each sentence. Sentences were then normalized at 70 dB-A. Finally, final word (i.e., target word) was cut off each sentence so that two different audio files were created.

- A. *Near court, I saw lawyers, but I don't like GUACAMOLE*
- B. *Near court, I saw lawyers, but I don't like GUACAMOLE*
- C. *All grown-up now, I saw lawyers, but I don't like GUACAMOLE*

Figure 28

Visual representation of auditory stimuli. A. Silence condition; B. Stationary condition and C. Babble condition. In all conditions target word is presented in silence to ensure its optimal semantic processing.

2.2.5. Stimuli creation

Sentences were then inserted 1 sec after the beginning in a background (i.e., stationary or babble noise or silence) with a 0 dB SNR. Those stimuli were then cut at the end of the sentence, so the background stopped at the same time as the sentence, before the target word (cf. Figure 28.B and C.). Target words were thus presented alone in separate files.

2.3. EEG recordings

The EEG signal was recorded with Biosemi system, from 64 electrodes (Electro-Cap International INC) accorded to the extended 10-20 standard system. Six additional electrodes were placed: 2 VOG and 2 HOG and 2 on mastoids. Signal was digitized on-line with a sampling rate of 512 Hz. Participants were installed in an electrically shielded and sound-proof EEG booth. They were asked to sit comfortably and listen to sentences. At the end of each sentence, they were asked to indicate whether the previously heard sentence was congruent or not, they could also indicate that they did not hear the sentence by pressing the corresponding keys. They were asked to blink and move if needed during this period. After they answered the fixation cross appeared again and experimenter could wait until the EEG signal was stable to launch the following trial. There were 270 test trials and 6 training trials so participants could get used to the stimuli. A break was proposed after 135 trials. Trials order was pseudo-randomized, so that the same condition could not be presented twice in a row. Triggers were inserted at the presentation of the target word (i.e., Target triggers). A second set of triggers were used to mark the answer of the participant at the coherence question (i.e., yes/no/not intelligible; Answer trigger).

At the end of the experimental session, participants were administered 3 tests: the Edinburgh test (Oldfield, 1971), the Alouette (Levafrais, 1967) and a phoneme suppression test (from ECLA 16+). The overall session lasted about 1h45.

2.4. Data Analysis

Data analyses were performed using the toolbox EEGLAB from Matlab. Data were first down-sampled at 256 Hz and re-referenced to linked mastoids. Epochs of 7.2 sec (200 ms pre-stimulus baseline) around Target trigger were extracted. The aim was to include in each epoch the Answer trigger to be able to sort trials depending on the participant's answer. As each condition only contained 30 trials, data were cleaned by performing an ICA (Independent Component Analysis). ICA in these conditions is interesting as it allows cleaning data from artifact without excluding trials. Artifact components (i.e., eye movements, muscular and cardiac artifact) were rejected. Epochs were then filtered with a low-pass filter of 30 Hz and high-pass filter of 1 Hz. Epochs were then sorted depending on whether participants heard the sentence or not. Trials where participants answered "not intelligible" were therefore excluded. Data of 2 participants were excluded due to technical issues during recordings, and 3 participants were excluded due to too weak intelligibility rate in the babble condition (<50%). Shorter epoch of 1.4 sec (200 pre-stimulus baseline) were then extracted.

Congruency effect was then calculated, by subtracting the congruent and incongruent condition for each background condition (i.e., Silence; Stationary; Babble). Ambiguity effect was also calculated, by subtracting the Congruent condition from the

Ambiguous condition for the 3 background conditions. Grand Averages were then calculated for each of this condition: Congruency effect Silence; Congruency effect Stationary; Congruency effect Babble; Ambiguity effect Silence; Ambiguity effect Stationary; Ambiguity effect Babble. The scalp midline was defined as the ROI (Fz, FCz, Cz, CPz, Pz, POz; Strauss et al., 2013). The signal was therefore averaged across these electrodes by taking mean amplitude of each electrode in 50 ms time laps from 250 ms to 600 ms (i.e., 7 periods). We calculated 7 t-tests against zero to cover the period from 250 ms to 600 ms and we corrected for multiple comparisons using the False Discovery Rate (FDR).

3. Results

3.1. Behavioral results

Intelligibility rates were analyzed with a repeated measure ANOVA taking Noise (Silence vs Stationary vs Babble) and Semantic Link (Congruent vs Incongruent vs Ambiguous) as within subject factors. The main effect of Noise appeared to be significant ($F(2,28) = 64.93$, $p < .001$), participants being less accurate at hearing the target sentence when embedded in Babble condition ($M = 71.97\%$; $SD = 13.48$) than in Stationary ($M = 88.97\%$; $SD = 8.80$) and in Silence ($M = 99.06\%$; $SD = 3.21$). The effect of Semantic Link on the other hand was not significant, therefore the congruency of the sentence did not have an impact on its intelligibility ($F(2,28) = 1.43$, $p > .10$). Participants were not better at hearing the target sentence when target word was congruent ($M = 87.9\%$; $SD = 14.0$) than when it was incongruent ($M = 85.44\%$; $SD = 16.06$) or ambiguous ($M = 86.67\%$; $SD = 13.9$). However the interaction between Noise and Semantic Link was significant ($F(4,56) = 2.80$, $p < .05$). A post-hoc (HSD Tukey) revealed that participants were impaired by the presence of Babble noise vs Stationary noise in all condition and by Stationary noise compared to Silence condition only in ambiguous and incongruent conditions (cf. Figure 29).

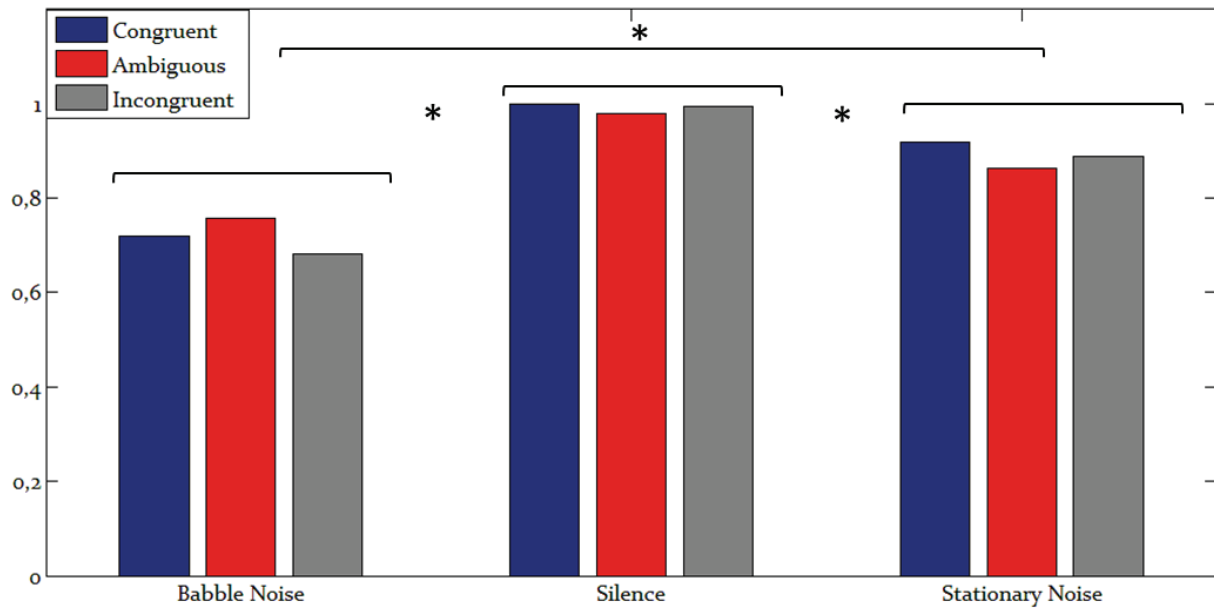


Figure 29

Behavioral results. Intelligibility rate depending on Noise condition and Link between target word and ambiguous word. * indicates significance.

3.2. EEG data

Our first hypothesis was that Incongruent and Ambiguous target would generate a N400 of higher amplitude than Congruent target. Incongruent condition was therefore subtracted to Congruent condition for all Noise conditions to obtain the congruency effect. Ambiguity effect was calculated by subtracting Ambiguous condition to Congruent condition. Mean amplitude for each time laps of 50 ms between 250 ms and 600 ms was averaged and compared to zero. After applying FDR, results revealed significant congruency and ambiguity effect only for Silence and Stationary conditions (cf. Figure 30 and Figure 31).

In Silence condition, congruency effect appeared to be significant from 250 ms to 550 ms showing that N400 in response to target word was significantly more negative in Incongruent Silence condition than in Congruent Silence condition. The ambiguity effect was significant from 250 ms to 600 ms after presentation of target word with the N400 being more negative in the Ambiguous condition than in the Congruent condition. These results reflect the difficulty of integration of incongruent and ambiguous target word in context. When sentences were embedded in a Stationary noise background, N400 in response to incongruent target word was more negative than in response to congruent target word from 250 ms to 600 ms after target word presentation. Ambiguity effect was also significant, first at 250-300 ms and then from 400 ms to 600 ms. This means therefore that incongruent and ambiguous target word were more difficult to integrate in context than congruent target word. When target sentence was embedded in babble background, no significant congruency effect was

found, suggesting that N400 in response to incongruent target word was not more negative than in response to congruent target word. No significant effect was found either for ambiguity effect.

The second hypothesis we raised in this paper is that N400 in response to ambiguous target word would be smaller than in response to incongruent target word. Data recorded in Babble condition were discarded from analysis as no N400 effect appeared neither in Incongruent nor in Ambiguous condition.

For Silence condition, a repeated measure ANOVA was performed on mean amplitude it included Effect (Congruency vs Ambiguity), Electrodes (POz vs Pz vs CPz vs Cz vs FCz vs Fz) and Time period (250 ms vs 300 ms vs 350 ms vs 400 ms vs 450 ms vs 500 ms) as within subject factor. Results revealed that the N400 effect was not different in Incongruent condition than in Ambiguity condition ($F(1,14) = 1.18, n.s.$). The effect of Electrode was not significant ($F < 1$), the effect of Time period was significant ($F(5,70) = 4.08, p < .01$); mean amplitude being more negative during the 400-450 ms time period. Only the interaction between Time period and Electrode appeared to be significant ($F(25,350) = 1.97, p < .01$), POz seeming to vary less with time.

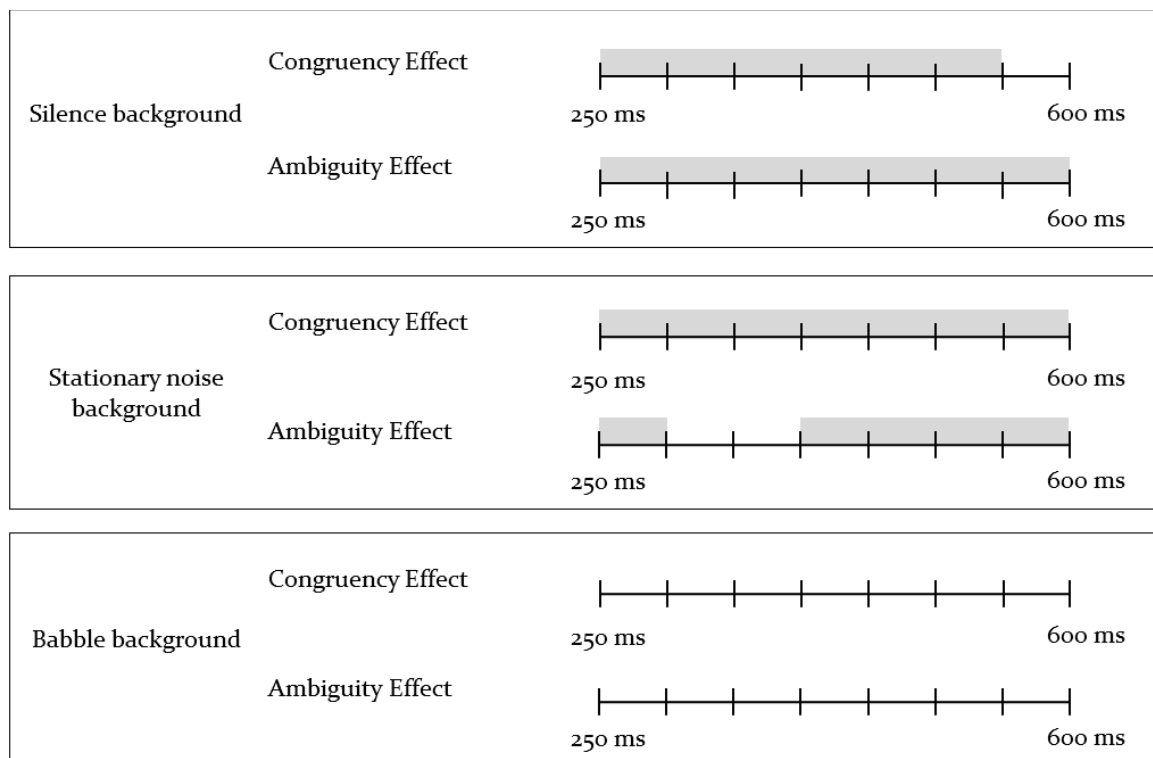


Figure 30

Results of t-test against zero after FDR depending on Noise condition (Silence vs. Stationary vs. Babble) and Effect (Congruency vs. Ambiguity). Grey time laps indicates significance.

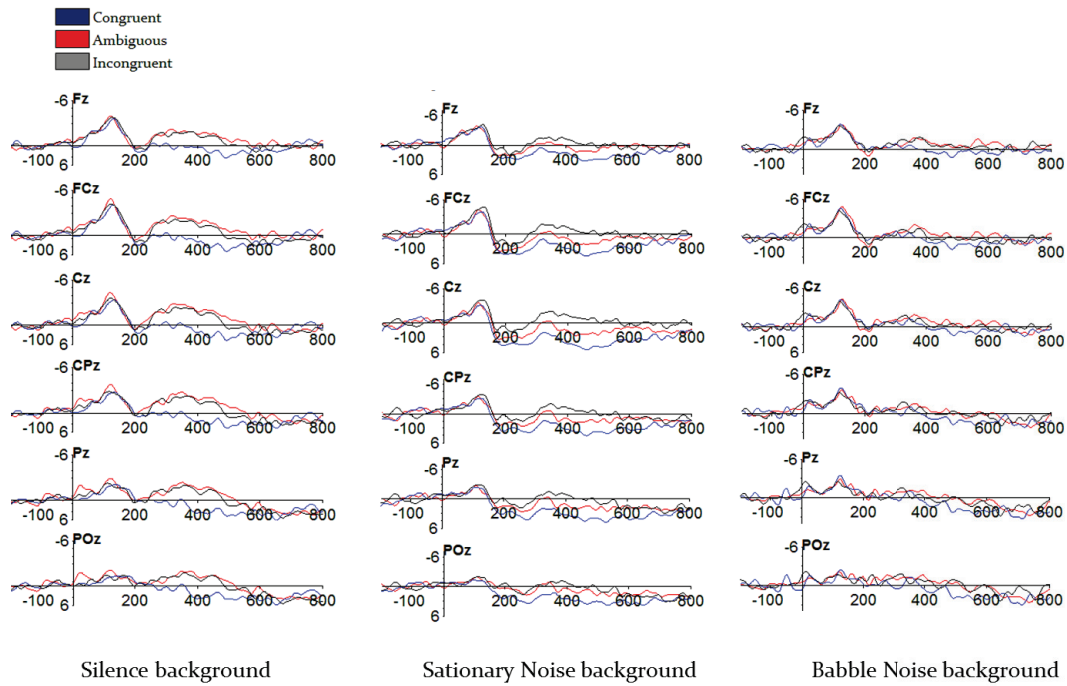


Figure 31

EEG signals in the ROI for Babble Noise, Silence and Stationary Noise conditions. For significant differences refer to Figure 30

Data collected in Stationary noise were analyzed with a repeated measures ANOVA including Effect (Congruency vs Ambiguity), Time period (400 ms vs 450 ms vs 500 ms vs 550 ms), and Electrode (POz vs Pz vs CPz vs Cz vs FCz vs Fz) as within subject factors. ANOVA revealed a significant effect of Electrodes ($F(5,70) = 2.74, p < .05$). Negativity seemed to be more important at Cz site. Nor the type of Effect ($F(1,14) = 2.60, n.s.$) neither the effect of Time period ($F(3,42) = 1.74, n.s.$) were significant, however the interaction between these factor was significant ($F(3,42) = 2.97, p < .05$). A post-hoc test (HSD Tukey) revealed that the Type of Effect was significant only for the time period 400-500 ms, where N400 in response to incongruent target word was more negative than N400 in response to ambiguous target word. No other interaction was significant.

Peak latencies in response to target word were also analyzed but no significant effect was found.

4. Discussion

4.1. Summary

The aim of this study was to investigate the modulation of the strength of activation of the different meaning of an ambiguous word in its context depending on the degradation of the signal. In order to investigate this, we presented ambiguous words embedded in sentences. The first part of each sentence referred to the dominant

meaning of the ambiguous word and the second part could refer to its dominant meaning (i.e., Congruent condition), to its subordinate meaning (i.e., Ambiguous condition) or to neither of its meaning (i.e., Incongruent condition). In order to vary the amount of masking, sentences could be presented in Silence, in Stationary noise or in Babble noise. Therefore, 9 conditions were presented to participants (Congruent Silence; Congruent Stationary; Congruent Babble; Ambiguous Silence; Ambiguous Stationary; Ambiguous Babble; Incongruent Silence; Incongruent Stationary; Incongruent Babble).

4.2. Behavioral results

Intelligibility scores revealed that participants had better intelligibility rates when sentences were presented in silence than in stationary noise and in both of them, than in babble noise. These intelligibility rates replicate what has been found by many other studies, that is, informational masking is more deleterious to intelligibility than energetic masking (Brungart, 2001; Brungart et al., 2001). However, there was no effect of the congruency, that is, the congruency of target word did not impact intelligibility whatever the background noise.

4.3. Congruency effect

A significant congruency effect has been found in condition Silence and Stationary Noise. That is, N400 was more negative in response to Incongruent target word than in response to Congruent target word in these two background conditions. This confirms that participants had more difficulty integrating incongruent target word in its context than congruent word. Interestingly, no congruency effect was found in the Babble Noise condition. That is, despite preserved intelligibility (only trials that were rated as intelligible were analyzed), no difference in integration processes was found between Congruent and Incongruent condition. This is congruent with previous studies investigating N400 effect in highly degraded signal (Strauss, 2013) and with previous behavioral results suggesting that a word could be heard and recognized without being sufficiently semantically processed to lead to a semantic priming effect (Dekerle et al., 2014). This is in line with the *Effortfulness Hypothesis*, postulating that despite preserved intelligibility, listening to the target sentence in Babble noise is so effortful, that it leaves less cognitive resources available to semantically process speech. In Silence and Stationary Noise conditions on the other hand, few cognitive resources are allocated to low level processes because speech signal degradation is very low, it can thus be processed semantically in a sufficiently efficient way that a N400 effect appears.

4.4. Ambiguity effect

Results evidenced a significant difference in N400 response to ambiguous target compared to congruent target word in Silence and Stationary Noise condition. Once again, no significant difference was found for Babble condition.

Following the rationale of *Effortfulness Hypothesis*, we postulated that when signal quality was good enough (i.e., Silence and Stationary Noise condition) both meaning of ambiguous word would be activated; therefore integration of ambiguous target word would be easier than integration of incongruent target. This relative facilitation would be reflected by a significantly smaller ambiguity effect compared to congruency effect at least in Silence condition. Surprisingly, however, ambiguity effect was significantly smaller than congruency effect only in Stationary Noise condition. These results do not support our hypotheses. On the contrary, they seem to imply that in Silent condition, dominant context generates the activation of dominant meaning only, whereas both meanings are activated in Stationary condition.

These results may reflect the fact that in Silence condition participants' context processing is sharper. Therefore, dominant meaning of ambiguous word is strongly primed and is the only meaning activated. In this condition, activating both meaning of ambiguous word would be useless. Our results therefore suggest that participants rely more on top-down processes in Silence condition than in Stationary condition although both meanings are supposed to be activated in weak context (Colbert-Getz & Cook, 2013). It seems that context in our study, although short, could be considered as strong as it was directly linked to the dominant meaning of ambiguous word.

Our results are not consistent with what is found by others (Strauss et al., 2013). Indeed, our study suggests that degraded speech signal leads to broader expectancies than does clear speech signals whereas Strauss et al., (2013) results suggest that decreased intelligibility leads to narrow expectancies. Indeed, they suggest that signal degradation causes higher reliance on top-down processes. They presented sentences that could end with a congruent and typical target word (e.g., She often peels POTATOES) or a congruent but non-typical target word (e.g., She often peels BANANAS) or an incongruent target word (e.g., She often peels GUM). Typicality was determined by the probability of co-occurrence of the verb and the target word (i.e., "to peel" is more often encountered with "patato" than with "banana"). Sentences could be presented in clear speech, or in vocoded speech (4-band or 8-band). In the 4-band condition, they did not evidence any N400 effect because of too much degradation. In clear speech condition, congruent but non-typical target word led to a N400 of medium amplitude whereas non congruent target word led to a N400 of high amplitude. In 8-band condition however, results showed no difference between congruent non-typical target and incongruent target, however, in both condition, N400 amplitude was higher than in response to congruent typical target word. Authors

therefore conclude that when speech signal is degraded, participants rely more on top-down processes to compensate for the uncertainty of the bottom-up processes.

However, these interesting results imply that the context can be efficiently processed even when its intelligibility is decreased so that strong predictions are made about the end of the sentence. It is therefore not in accordance with the *Effortfulness Hypothesis*. On the other hand, our results suggest that slightly decreased intelligibility leads to weaker predictions because bottom-up processes cannot be relied on.

One possible explanation to these differences is that target words in Strauss et al., (2013) experiment were degraded, whereas they were not in ours. The aim in our study, as mentioned in the introduction, was to be sure that the target word was efficiently processed. It is very probable, in Strauss et al.'s study that the semantic processing of target word was less efficient in 8-band condition than in clear speech. According to the *Effortfulness Hypothesis* it implies less semantic processing and integration might rely on lower level features (e.g., lexical features). Strauss et al. (2013) found that participants integrated more easily target words that were frequently presented with the verb. At the lexical level, typical target were more associated to the verb. It is therefore possible that the easier integration of typical target word relies on a lexical facilitatory effect (i.e., frequent simultaneous presentation) rather than a semantic effect.

4.5. Phonological Mismatch Negativity

As can be seen in Figure 30 the observed N400 effects are very long lasted. Particularly they start very early around 250 ms. One can therefore assume that the observed effect may be confounded with the Phonological Mismatch Negativity (PMN; Connolly, Phillips, Stewart, & Brake, 1992; Connolly, Stewart, & Phillips, 1990). The PMN or N250 is a negative ERP reflecting the violation of the expectation at a phonological level. This ERP often co-occurs with the N400 as usually the violation of a semantic expectation is linked to violation of phonological expectation and vice-versa. However, they have been proved to be independent (Connolly & Phillips, 1994; Desroches et al., 2009; R. L. Newman & Connolly, 2009). This component can explain that ambiguity effect is significantly smaller than congruence effect in Stationary Noise, only between 400 and 500 ms. Indeed, in both condition the target word phonologically violated participants' expectations. Therefore, both ambiguous and incongruent target word presentation elicited a PMN and consequently no significant difference is observed between these conditions at the corresponding time period. These results thus suggest that participants had expectations regarding the target word in Stationary and Silence condition.

5. Conclusion

This study investigated semantic processing of ambiguous word in adverse listening conditions. The purpose was to test the modulation of the strength of semantic processing in sentence and the interaction of top-down and bottom-up processes. Our results are congruent with the Effortfulness Hypothesis and confirm that semantic processing depends on the quality of bottom-up processes, even when intelligibility is preserved. Top-down predictions were consequently stronger in Silence than in Stationary Noise although phonological expectations did not differ in those conditions.

Acknowledgments

The first author is funded by a PhD grant from Région Rhône Alpes and EXPLORA DOC grant from Région Rhône Alpes.

6. Résumé du Chapitre VII

Le chapitre VII présente le travail expérimental effectué à Montréal, sur la force d'activation des différents sens d'un mot ambigu. Les résultats de cette étude en EEG ont permis de mettre en évidence que la présentation d'un mot ambigu dans un contexte favorisant son sens dominant lorsque toutes les ressources cognitives sont disponibles (i.e., dans notre cas, dans le Silence) génère de fortes prédictions (i.e., processus *top-down*). De ce fait, seul le sens dominant du mot ambigu est activé, puisqu'il est le seul pertinent. En revanche, lorsque la qualité du signal de parole (i.e., processus *bottom-up*) diminue légèrement en présence de masquage énergétique le traitement sémantique du contexte est moins efficace. De ce fait, la sélection du sens approprié du mot ambigu est moins précise, et les deux sens de celui-ci sont activés. Malgré l'activation de ces deux sens, il semble que des prédictions soient effectuées sur le mot cible au niveau phonologique. C'est-à-dire que les participants s'attendent à entendre le mot cible faisant référence au sens dominant du mot ambigu. Enfin, lorsque le signal est perturbé par du masquage informationnel (i.e. bruit parolier), le traitement sémantique n'est pas mesurable. Ces derniers résultats étant concordants avec ceux mis en évidence dans l'Etude 1 au Chapitre VI.

Cette étude a donc permis de mettre en évidence que la présence de bruit impacte le traitement sémantique au niveau de la phrase même lorsque l'intelligibilité est préservée. Une fois encore le masquage informationnel est apparu comme plus délétère à l'intelligibilité que le masquage énergétique. Il semble également qu'il soit plus délétère au traitement sémantique puisque aucun effet N400 n'apparaît significatif lorsqu'il est prédominant (i.e., en condition Bruit Parolier).

De l'Axe 1 à l'Axe 2

Les deux premières études de la partie Expérimentale s'intéressent à la modulation du traitement sémantique dans la situation de parole dans le bruit.

L'Etude 1 s'intéresse au traitement sémantique au niveau du mot, alors que l'Etude 2 s'intéresse au traitement sémantique du mot présenté dans un contexte phrastique. Les résultats de ces deux études ont permis de montrer que le traitement sémantique du mot est impacté par la présence de bruit, même lorsque l'intelligibilité est préservée. Il semble que le traitement sémantique observé dans ces études, dépende de processus stratégiques, puisqu'il nécessite une certaine saillance du signal (Etude 1) et est soumis à la disponibilité des ressources cognitives (Etude 1 et Etude 2).

Les résultats de ces deux premières études ont permis de mettre en évidence que la présence de bruit a un impact délétère sur la compréhension du langage même lorsque l'intelligibilité est préservée. Toutefois, le signal de parole peut également être dégradé de manière permanente du fait de traitements auditifs centraux déficients.

Dans le second axe de cette partie Expérimentale, nous nous intéressons à l'effet de la mauvaise qualité du signal dû à des traitements auditifs centraux soit immatures (Etude 3) soit déficients du fait de la présence de Troubles du Traitement Auditif (Etude 4). Ce second axe permet de s'intéresser aux conséquences que peut avoir un signal de parole en permanence légèrement dégradé sur les représentations du langage.

Chapitre VIII

Etude 3a. Les liens entre développement des traitements auditifs centraux et le langage

L'Etude 3 s'est intéressée au développement des traitements auditifs centraux sur une population d'enfants au développement typique, âgés de 6 à 11 ans. La BECAC (Donnadieu et al., 2014) leur a été administrée. Elle est constituée de 6 principaux tests, eux même composés de un ou plusieurs sous-tests entièrement non verbaux. L'objectif est d'évaluer les compétences auditives centrales comme elles sont décrites par l'ASHA (2005). De plus, comme recommandé par la BSA (2011), la BECAC ne contient pas de matériel verbal afin de limiter l'impact du développement du langage ou des troubles langagiers. Les passations ont été effectuées par Donnadieu et son équipe, et j'ai analysé les résultats.

Une partie des résultats a fait l'objet d'une communication orale avec actes édités en août 2013 à la 14^e édition de la conférence Interspeech. Nous présentons ici la publication correspondante (Etude 3a) ainsi qu'un second article, contenant toutes les données sur les tests auditifs (Etude 3b). Nous présentons ces deux publications car les analyses effectuées sur les données sont différentes, et chacune présente un intérêt.

Dans l'Etude 3a nous présentons les données des enfants à 3 des 6 tests de la BECAC. Spécifiquement, l'article présente des données développementales au test de latéralisation, discrimination auditive (en termes de fréquence et durée) et masquage central. Afin d'étudier les liens entre le développement de ces compétences auditives et le développement des représentations langagières, la conscience phonologique, le vocabulaire ainsi que la compréhension orale ont été testés.

Comme mentionné dans la partie Théorique, la distinction entre le développement des facteurs sensoriels et non-sensoriels est un enjeu majeur dans l'étude de la population infantile (cf. p83). De ce fait, dans ce premier article, l'analyse des données recueillies a été effectuée en utilisant les Hit – FA pour le test de discrimination et le test de masquage central.

L'objectif de cette première publication était de s'intéresser aux liens spécifiques entre le développement de certaines compétences auditives centrales et le développement des représentations langagières. Il est évident, à la lumière des études sur les populations présentant des déficits auditifs périphériques qu'un traitement auditif de qualité est nécessaire au bon développement des représentations langagières. Toutefois, à notre connaissance peu d'études se sont intéressées aux liens entre les compétences auditives centrales et le développement du langage (Chonchaiya et al., 2013) en dehors des cas spécifiques de troubles de l'apprentissage (Bishop & McArthur, 2005; Fischer & Hartnegg, 2004; Mengler, Hogben, Michie, & Bishop, 2005).

Development of Central Auditory Processes and their links with language skills in typically developing children

Marie Dekerle ¹, Fanny Meunier ¹, Marie-Ange N'Guyen ², Estelle Gillet-Perret ², Delphine Lassus-Sangosse ², Sophie Donnadieu ^{3, 4}

¹ Laboratoire L2C2, CNRS - UMR 5304, Auditory Language Processing (ALP) research group, Lyon, France

² CHU de Grenoble, Centre du Langage, Grenoble, France

³ Université de Savoie, Chambéry, France

⁴ Laboratoire de Psychologie et NeuroCognition, CNRS - UMR 5105, Grenoble, France

Abstract

This study aimed at investigating the development of central auditory processes and their links with language skills. Seventy nine typically developing children divided up in five levels groups were recruited among primary schools. The development of central auditory processes was assessed with three main tasks. A lateralization task, a discrimination task and a central masking task were presented. These tasks were selected as each of them may correspond to important auditory processes underlying different speech abilities, and thus playing a role on its development. Verbal skills were evaluated on three levels: comprehension, vocabulary, and phonological awareness. Results confirmed a developmental effect both on auditory and verbal skills. In addition, vocabulary and phonological awareness performances correlated with auditory skills, highlighting links between central auditory processing and language development.

Index Terms: Central Auditory Processes, language development

1. Introduction

Central Auditory Processing (CAP) and their disorders (Central Auditory Processing Disorders; CAPD) is a growing field of interest, as suggested by the wide literature on the topic (Bishop, 2007; Demanez et al., 2003; Moore, 2012; Rosen, 2005) CAP refers to the processing of acoustic stimuli by the central nervous system. Many anatomical structures are involved in these processes and are in charge of acoustic processing like sound spatialization, detection of frequency change, pitch discrimination, temporal order judgment... Children with CAPD meet difficulties in understanding speech-in-noise and rapid speech, in localizing source, and focusing attention, they also tend to ask for repetition, and confound sounds (ASHA, 2005).

Some authors suggest that CAPD may have a causal relation with language developmental disorder (i.e., Specific Reading Impairment; SRI and Specific Language Impairment; SLI). Many studies indeed reported auditory processing deficit in children and adults with language disorder (Goswami et al., 2002; McArthur, Ellis, Atkinson, & Coltheart, 2008; Ramus et al., 2003; Tallal, 1980). According to them, the deficit in auditory processing could lead to language developmental impairment. Other studies however argue that only a minority of people with language disorder are affected by CAPD and that this might reflect a simple co-occurrence rather than a causal relation (Rosen, 2003)].

The first aim of this study was to look at the development of CAP in typically developing children. Although development of language skills is well known, the evolution of CAP is still little documented while it appears that whilst the human cochlea has completed its development by birth, the brain's auditory pathways and centers develop slower and progressively, from the brain stem to the auditory cortex, for at least the first decade of life (Moore, 2002). How is this physiologic development reflected in auditory skills?

The second interest of this study was to find out if development of CAP was related to language development in typically developing children. Indeed, if SLI and SRI are linked, at least for some subgroup to auditory processing deficit, CAP should be related to language abilities, in typically developing children as well as in children suffering from SLI or SRI. In this paper, we therefore chose to test three main CAP abilities which are supposed to be important in language development. These data are part of a larger study, aiming at normalizing a battery evaluating CAP in children. As this battery aims at testing children suffering from SLI and SRI, it does not involve verbal material.

The first task was a lateralization task, aiming at evaluating the development of the capacity for typically developing children to use Interaural Level Difference (ILD) to lateralize sounds. It is widely known that spatial cues are very helpful for speech comprehension particularly in noise (K. Allen, Carlile, & Alais, 2008; Bronkhorst & Plomp, 1988; Dole et al., 2012). The improvement of this ability should thus lead to a better understanding of speech in noisy condition (e.g., classroom) and ease language development.

As detection of frequency and duration changes are known to be involved in accurate phoneme discrimination, we proposed to test, in a second task, frequency and duration discrimination acuities. Indeed, if there is a direct link between slow frequency discrimination and SLI and SRI, as proposed by the auditory deficit theory [9], then these discrimination abilities should develop with age and correlate with participants' phonological scores.

Finally, participants had to perform a central masking task. As mentioned above speech is mostly perceived in noise in daily life, different kind of masking can thus occur and render speech unintelligible. One can argue then, that the ability with which participants can still hear and process accurately the target sound is crucial for language comprehension. Three main kind of masking have been conceptualized: one of them, energetic masking occurs at the peripheral level of the auditory system. It corresponds to the overlap of two sounds in the same spectral band; cochlea will therefore not be able to encode both of them. The second one is named informational masking, and corresponds to higher level interferences (i.e., linguistic level). Finally, the third one, on which was our interest is the central masking. Central masking occurs when two sounds are presented dichotically. This masking is highlighted by presenting a pure tone to one ear and a white noise (i.e., wide band masker containing all frequencies) to the other, with a Signal/ Noise Ratio varying across trials. Although no energetic masking occurs, participants are less accurate (i.e., higher SNR) at detecting the pure tone when the white noise is presented because of integration difficulty at the central level (McQueen & Terhune, 2011).

In this paper we presented the results showing the development of three CAP abilities and their correlational links with verbal development. Children from 5 age groups have been tested, going from 6 years old to 12.

2. Method

2.1. Participants

Seventy-nine children from 74 to 140 month ($M = 103.26$; $SD = 17.35$) were tested. They were selected by teachers as having no particular disorder. Their non-verbal intelligence was controlled with Raven's progressive matrices ($M = 9.96$; $SD = 3.13$). The interest of the study being the development of central auditory processing and its relation with language skills development, participants were spread in 5 groups (corresponding to their school grade; $M_{1st\ grade} = 80.65$; $SD_{1st\ grade} = 4.38$; $M_{2nd\ grade} = 91.72$; $SD_{2nd\ grade} = 4.36$; $M_{3rd\ grade} = 102.42$; $SD_{3rd\ grade} = 3.38$; $M_{4th\ grade} = 114.66$; $SD_{4th\ grade} = 5.36$; $M_{5th\ grade} = 127.47$; $SD_{5th\ grade} = 5.56$). Participants' age was significantly different across groups ($F(4,84) = 285.45$, $p < .0001$). Nonverbal intelligence was not different across groups ($F = 1.04$, *n.s.*).

2.2. Verbal tests

Three main language abilities were tested for each participant. Lexical level was evaluated using EVIP (Dunn, Theriault-Whalen, & Dunn, 1993). Phonological awareness was assessed with word, pseudo-word and logatome repetition took from the BALE (Jacquier-Roux, Lequette, Pouget, Valdois, & Zorman, 2010). Finally, oral comprehension was assessed by the ECOSSE (Lecoq, 1996).

2.3. Auditory tests

Auditory central processing was assessed using 3 main tasks. First a lateralization task was presented, then a frequency and a duration discrimination tasks and finally a central masking task.

2.3.1. Lateralization task

This first test aimed at evaluating participants' ability to detect the origin of a pure tone based on an Interaural Level Difference (ILD). A 500 Hz pure tone was presented for 250 ms in both ears. ILD could be of 0 dB (binaural condition) or ranged from -5 dB to -25 dB with 2 dB steps. Participants could thus hear the pure tone as coming from right, left or front. Binaural condition was presented on 22 trials. Each ILD was presented twice leading to 22 trials with the pure tone coming from the right and 22 trials with the pure tone coming from the left. Trials were presented randomly. Children were asked to point the direction from where the pure tone was coming.

2.3.2. Discrimination task

Participants had to compare a 500 Hz, 75 ms standard pure tone preceding a comparison tone by 400 ms and to indicate if the two tones were identical or different. Standard and comparison tones differed in half trials only on one dimension: duration or frequency.

In the *duration discrimination task* the stimulus could be of the same length or shorter than the standard. Its duration varied from 27 ms to 75 ms by 8 ms steps.

The *frequency discrimination task* presented stimuli of the same or higher frequency than the standard. It varied from 500 Hz to 609 Hz by 7 Hz steps.

This task is an adaptation to behavioral measurement of a MMN study (Pakarinen, Takegata, Rinne, Huotilainen, & Naatanen, 2007).

2.3.3. Central masking task

This task, constructed using the same paradigm as Breier, Gray, Klaas, Fletcher, and Foorman (2002) aims at measuring the detection threshold of a 500 Hz pure tone presented in two conditions. In one of them, a masking white noise was presented in the contralateral ear, whereas not in the other. The pure tone intensity varied from -30 dB to -58 dB by 2 dB steps compared to the masking noise. Half of the 60 trials did not contain pure tone. Participants had to indicate whether they had heard the pure tone or not by pressing two pre-specified key.

2.4. Procedure

Participants were seated in front of a computer and heard stimuli through earphones (Sennheiser HD 230 pro). Intensity was controlled and fixed at 65 dB-A with a sonometer (Votcraft SL-50).

3. Results

3.1. Auditory Tests

Mean accuracy for auditory tests are summarized in Figure 32. All the nonverbal tests were analyzed with a repeated measure ANOVA with the Group (1st grade; 2nd grade; 3rd grade; 4th grade; 5th grade) as between subject factor and participants' scores as the dependent variable. Within subject factor will be specified for each task. For the discrimination and the central masking tasks, data were analyzed using Hit minus False Alarm.

- *Lateralization task.* Lateralization (i.e., ILD : -5 dB; -7 dB; -9 dB; -11 dB; -13 dB; -15 dB; -17 dB; -19 dB; -21 dB; -23 dB; -25 dB) was taken as within subject factor. Analyses revealed a significant effect of Lateralization ($F(10,740) = 52.72, p < .01$). Participants were more accurate with increased ILD. The effect of Group approached significance ($F(4,74) = 2.33, p < .10$), older participants tended to be better than younger. Finally the interaction Group*Lateralization ($F < 1$) was not significant, thus suggesting that the pattern of results was the same along development.
- *Duration discrimination.* As expected, the effect of the within-subject factor Duration differences (27 ms; 35 ms; 43 ms; 51 ms; 59 ms; 67 ms) was significant ($F(5,365) = 123.53, p < .001$), the greater is the difference between the standard and comparison tones, the more participants were accurate at detecting differences. Group effect was fully significant, ($F(4,73) = 3.74, p < .01$) accuracy increased with age. In addition, the interaction Duration*Group was also significant ($F(20,365) = 2.11, p < .01$) suggesting that the pattern of results was different across group.
- *Frequency discrimination.* The within-subject factor Frequency (527 Hz; 532 Hz; 546 Hz; 562 Hz; 583 Hz; 609 Hz) appeared significant ($F(5,370) = 111.52, p < .01$). Predictably, participants were more able to detect a difference between the standard and the stimulus with increased frequency difference. Older participants were more accurate than younger as revealed by the significant effect of Group ($F(4,74) = 4.20, p < .01$). The interaction Frequency*Group was not significant ($F(20,370) = 1.13, n.s.$).
- *Central masking.* The within-subject factor Conditions (with masking noise; without masking noise) was significant ($F(1,74) = 26.32, p < .0001$), as participants

were more accurate at detecting a pure tone in silence than with a masking noise in the contralateral ear. The effect of Group was also significant, with an increased accuracy with age ($F(4,74) = 2.63, p < .05$). The interaction Group*Mask however, was not significant ($F < 1$). To further understand our results, post-hoc analyses were performed. Although the interaction did not reach significance, it appeared that the effect of mask was significant only for children in 1st ($p < .01$), 2nd ($p < .05$), and 3rd ($p < .05$).

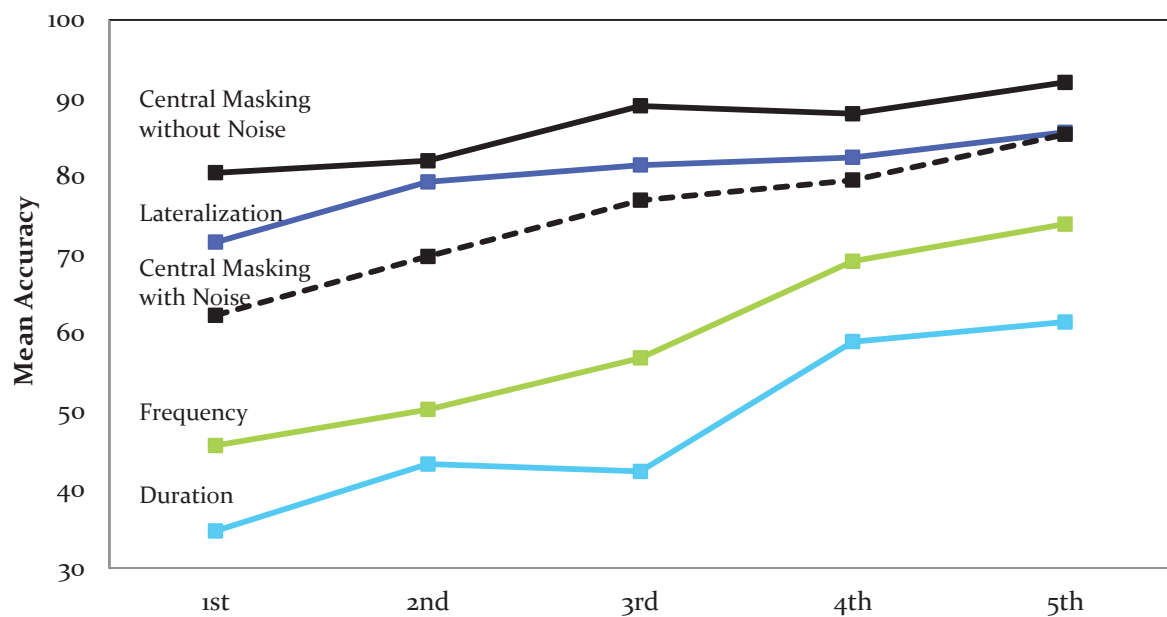


Figure 32

Mean accuracy by groups for each auditory task.

3.2. Verbal tests

Scores at verbal tests are presented in Figure 33.

EVIP Test. Participants' scores were normalized, accounting for their age. Those scores were not significantly different across groups ($F < 1$). However, when turned into an equivalent age score, Group effect was significant ($F(4,74) = 13.30, p < .001$), highlighting that participants vocabulary improved with age.

Phonological awareness. Older participants were more accurate to repeat word ($F(4,74) = 6.95, p < .0001$); pseudo word ($F(4,74) = 4.14, p < .01$) and logatomes ($F(4,79) = 3.75, p < .01$).

ECOSSE Test. The factor Group was significant in this comprehension task ($F(4,79) = 10.05, p < .001$). Older participants performed better than younger.

To sum up, verbal tests are sensitive to participants' group. However, it is interesting to specify that when scores are transformed into equivalent age score, no difference appears, thus suggesting that each group has the same level compared to its aged match.

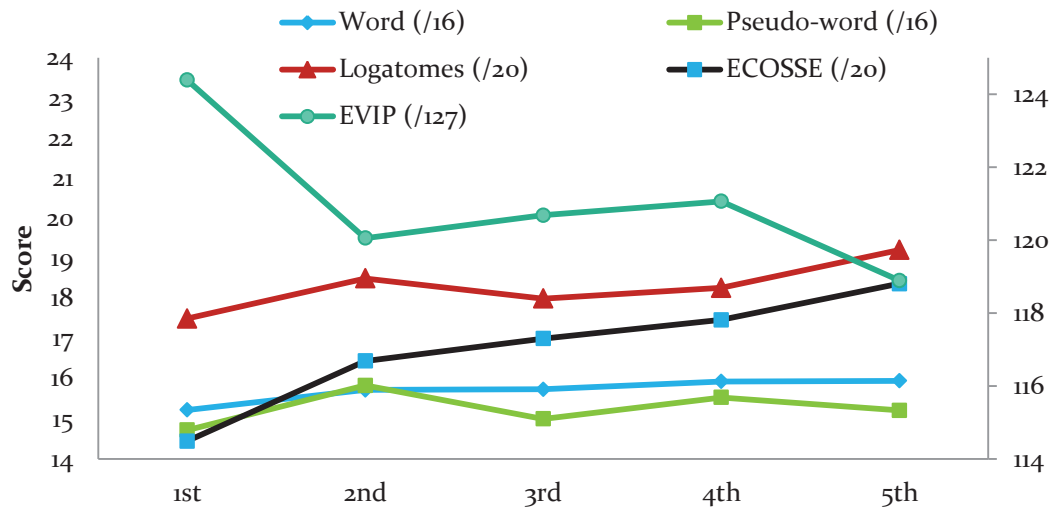


Figure 33

Score at verbal tests depending on groups.

3.3. Correlations between verbal skills and nonverbal tests

To investigate the link between language and auditory non-verbal skills we performed correlations matrices on our data. As mentioned in the results the effect of Group was very important, correlations matrices were consequently performed for each group so that correlation cannot be influenced by a simple developmental effect.

- *Phonological awareness.* A significant correlation was found between the performance at the logatome repetition task and the ability to discriminate frequency, in 4th grade children ($r = 0.65$; $p < .05$). High frequency discrimination accuracy was associated with high score to logatome repetition task. In addition, it appeared that participants who were less sensitive to central masking had better score in pseudo-word repetition in 3rd grade ($r = -0.55$; $p < .05$).
- *Vocabulary.* Scores to the EVIP test improved with accurate duration discrimination for 3rd grade children ($r = 0.60$; $p < .01$) and with frequency discrimination for 4th grade children ($r = 0.54$; $p < .05$).

4. Discussion

The first aim of the study was to investigate the behavioral consequences of the central auditory processes maturation across 5 groups of children, ranging from 1st to 5th grade of primary school with comparable non-verbal intelligence accounting for age. Participants performed 3 tests selected as they evaluate important abilities for speech development.

First, a lateralization task was presented, participants were asked to indicate from where the pure tone was presented (i.e., left, right or center). As expected, accuracy significantly increased with the ILD, that is, the more the sound is lateralized, the better children are to point its direction. It seemed however that this ability is already well developed by the age of 6 (1st grade) as the Group effect misses significance. The ability to lateralize sounds is crucial in speech-in-noise comprehension; as it seems to be operational early in development, it might indicate that young children can efficiently use spatial cues to release masking in adverse situation.

A discrimination task was then presented to children. They were asked to evaluate if the two presented pure tones (a standard followed by a stimulus) were the same or not. In the first subtest, the stimulus duration could be the same or shorter than the standard's. It appeared that participants increased their performance when the difference between the standard's and the stimulus' duration was more important. In addition, development seemed to improve their abilities to process stimulus duration, as revealed by the significant effect of Group. The second subtest consisted in comparing two pure tones which this time could differ in frequency (i.e., the stimulus being of higher frequency than the standard in half trials). Once again, accuracy increased with the difference of frequency between standard and stimulus, and older participants performed better than younger, suggesting that maturation of central auditory processes improves frequency discrimination acuity.

Finally, participants were asked to perform a central masking task. This task might reflect the sensibility of participants to masking in daily life. It was sensitive to participant's age, as older performed better than younger. In addition, post-hoc results revealed that the deleterious effect of mask was significant only for young children (from 1st to 3rd grade), older children however were not significantly disturbed by the masking noise. It seems, thus that the maturation of central auditory processes leads to a better resistance to central masking, and consequently, to a better ability to understand speech-in-noise.

Despite the probably high impact of such central auditory abilities on the development of language skills, few studies have experimentally investigated their development in parallel with language development in children. These results clearly highlight that central auditory processes improve with age and this improvement have direct

repercussion on behavior. Indeed if there is truly a relationship between CAPD and developmental language and its disorders, we should find correlation between the children performances in auditory and language tasks. Interestingly significant correlation between auditory and language skills appeared for children in 3rd and 4th grade. Our results highlighted that phonological awareness was positively related to frequency discrimination skills. This is consistent with Tallal's theory (Tallal, 1980; Tallal & Gaab, 2006) of a temporal perception deficit. According to her, children suffering from SRI and SLI fail to detect rapid frequency change which leads to a difficulty to discriminate accurately phonemes. Third grade results confirm clearly the link between the ability to discriminate frequency (even without time constraint) and phonological awareness. Sensitivity to central masking also influenced phonological awareness as evidenced by the improved accuracy to repeat logatomes with the decrease of the difference between masked and unmasked condition. This suggests that a better ability to resist masking provides a better phonological awareness. It seems quite intuitive, as speech is often perceived in noise, that children performing well in noise create more robust phonological representation. This result is also consistent with data on SRI. Their phonological representations indeed are known to be weak and they perform poorly in noise (Bradlow et al., 2003; Ziegler et al., 2009).

Other correlation between language and auditory skills appeared to be significant. Vocabulary is linked to discrimination task for both 3rd grade (duration) and 4th grade (frequency). The ability to process low acoustic features influencing higher level of language than phonological awareness is very interesting. Indeed, it suggest that other SLI symptoms than phonological representation weakness (e.g., naming deficit, poor vocabulary) could be linked to a CAPD. As speech comprehension models suggest that the oral input is processed in cascade (Marslen-Wilson & Welsh, 1978) it seems coherent that a poor phonological resolution leads to a poor comprehension. The imprecise phonological input could be less easily associated to a meaning. This would imply then that the maturation of CAP does have consequences on language development, whether in typically developing children or SRI or SLI children.

An interesting way to pursue this study would be to test SRI and SLI children, to evaluate if the central auditory processes we are testing here are specifically impaired in those children. In addition, if participants encountering difficulties in vocabulary tasks show more auditory deficit than others this would confirm that low level lack of precision have consequences on high linguistic processes and/or representations. This would also mean that low level signal degradation (whether because of poor processes or external noise) influences language development, and that noisy environment could generate other problems than simple lack of intelligibility.

5. Conclusion

This study showed that central auditory processes mature among primary school (from 1st to 5th grade). This maturation affects principally the ability to discriminate frequency and duration and the sensitivity to central masking. It seems however that the ability to lateralize sound is already efficient in 1st grade children. In addition, some language skills (i.e., phonological awareness and lexical level) correlate with auditory abilities, thus suggesting that low level auditory processes are involved in language development at different level.

Acknowledgements

This research was financially supported by the “SFR Santé et Société”, Université Pierre Mendès-France BSHM, BP 47, 38040 Grenoble Cedex 9, www.sfr-sante-societe.net. The first two authors are also supported by the SpiN project - European Research Council grant (no. 209234). The first author is supported by a PhD fund from Région Rhône Alpes.

We thank A. Risso, A. Didelot and C. Martin for their valuable help in the recruitment and the testing of the children. We thank all the children, their family and their teachers for their participation in this research.

6. Résumé du Chapitre VIII

Cette première publication s'intéresse uniquement à 3 des 6 tests compris dans la BECAC et aux liens entre les compétences auditives ainsi testées et les compétences langagières. Ces résultats mettent en évidence que le développement des compétences auditives a des conséquences sur le développement langagier chez les enfants au développement typique.

Particulièrement ces liens apparaissent chez les enfants de CE₂ et CM₁, donc aux alentours de 8-9 ans, âge auquel le langage atteint le niveau des adultes pour ce qui est de la forme (i.e., prononciation). Or, il est vraisemblable que les traitements auditifs centraux ne soient réellement liés qu'à la forme du langage, c'est-à-dire au traitement phonologique. Les corrélations obtenues avec les compétences langagières de plus haut niveau ne reflètent vraisemblablement que les conséquences de difficultés de traitements de plus bas niveau. Ces résultats suggèrent que chez les enfants, des différences de maturation dans les traitements de bas niveau sont associées à des différences en termes de compétences langagières. En revanche, il semble que ces liens évoluent avec l'âge, en effet, ils ne sont pas significatifs sur les petites classes (i.e., CP et CE₁) et ne le sont plus non plus en CM₂. Ce changement dans l'impact de la qualité des traitements sensoriels sur le langage peut être lié avec la littérature s'intéressant aux liens entre l'intégrité des traitements sensoriels et les compétences cognitives (Baltes & Lindenberger, 1997; Lindenberger et al., 2001). Il est ainsi possible que certaines périodes du développement montrent une plus grande vulnérabilité des fonctions cognitives de haut niveau aux dysfonctionnements sensoriels.

Toutefois, l'important développement du langage durant la période étudiée fait que les tests langagiers pouvaient entraîner un effet plancher pour les groupes plus jeunes et un effet plafond pour les groupes plus âgés. Le manque de variabilité qui en découle peut également expliquer l'absence de corrélations significatives pour les groupes CP, CE₁ et CM₂. Cette explication ne permet toutefois pas de justifier l'absence de corrélations sur ces tranches d'âges pour l'EVIP puisque ce test est conçu pour tester les enfants de 2 à 16 ans.

Chapitre IX

Etude 3b. Le développement des traitements auditifs centraux

Cette première analyse des données développementales a ainsi permis de mettre en évidence les liens entre le développement des traitements auditifs centraux et des compétences langagières sur une population d'enfants âgés de 6 à 11 ans. Toutefois, cette première analyse posait deux problèmes méthodologiques. D'une part, le langage se développe de façon importante entre 6 et 11 ans. De ce fait, les enfants les plus âgés (9-11 ans) présentent des scores plafonds à certains des tests employés (i.e., répétition de mots ; répétition de logatomes ; ECOSSE). D'autre part, la plupart des études s'intéressant aux compétences auditives centrales rapportent des seuils de plus petite différence détectée. Afin que nos résultats puissent être comparés à la littérature, tout en utilisant des tests plus rapides que les autres études, nous avons calculé la courbe psychométrique la plus proche des données recueillies pour chaque participant. De ce fait, pour les tests de latéralisation et de discrimination (3 sous-tests) nous avons calculé la différence nécessaire entre le stimulus et le standard pour que les participants soient capables de les distinguer à 75%. Le seuil de 75% permet de s'assurer que l'enfant est capable de discriminer le standard du stimulus et que ces réponses ne résultent pas du hasard (par opposition à un seuil de 50%).

Enfin, comme présenté au sein de la partie Théorique, les études s'intéressant au développement des traitements auditifs centraux s'intéressent souvent à un ou deux traitement auditifs (mais voir Ludwig et al., 2014; Moore, 2012). Notre étude en revanche teste la même population sur plusieurs aspects du traitement auditif. De ce fait, les courbes développementales des différents processus testés peuvent être comparées et leur maturation relatives évaluées.

Central auditory processing development in primary school children

Marie Dekerle^{a, b}, Marie-Ange N'Guyen^c, Estelle Gillet-Perret^c, Delphine Lassus-Sangosse^c, Fanny Meunier^{a, b}, Sophie Donnadieu^{d, e}

(a) Laboratory on Brain, Language and Cognition (L2C2) CNRS, UMR 5304, Lyon France

(b) University of Lyon

(c) CHU de Grenoble, Centre du Langage, Grenoble, France

(d) University of Savoie, Chambéry, France

(e) Laboratoire de Psychologie et NeuroCognition, CNRS – UMR 5105, Grenoble, France

ABSTRACT

The aim of this study was to investigate the development of central auditory processes in primary school children. Eighty-nine typically developing children from 1st to 5th grade and 23 young adults were tested on several auditory tasks, testing central auditory processes as defined by ASHA (2005). Indeed, the here proposed battery aims at testing a wide range of central auditory processes: lateralization, discrimination (i.e., frequency, duration and intensity), auditory recognition and identification, auditory pattern recognition, central masking and stream segregation. All tests use non-verbal materials to (a) minimize the effect of language development in the present study and (b) be used with population presenting dyslexia and/or SLI. Results evidenced developmental effect in all but one proposed task: stream segregation. Interestingly, developmental trajectories differed depending on the task, presenting whether developmental steps and maturation before adulthood or linear development leading to maturation during adolescence.

Keywords: Central auditory processes, auditory development, typical development

1. Introduction

Central auditory processing are defined as auditory processes taking place at the central level, that is, excluding the external, middle and internal ear. According to the American Speech-Language-Hearing Association (ASHA, 2005), these processes are responsible for the following behaviors: sound localization, auditory discrimination,

auditory pattern recognition, temporal aspects of audition, auditory performance in competing acoustic signals and auditory performance with degraded acoustic signal. Central auditory processes therefore cover a large area of auditory scene analysis. Although it is well known that central auditory processing are not mature at birth but continue to develop during childhood (e.g., Dawes & Bishop, 2008; Rickard, Smales, & Rickard, 2013; Sussman et al., 2007) to our knowledge very few studies got interested in the time course of this development in more than one or two processes (but see Ludwig et al., 2014; Moore et al., 2011). In fact, central auditory processing are more usually studied in relation to Auditory Processing Disorders (APD) and their implication in Specific Language Impairment (SLI) and dyslexia (Bishop & McArthur, 2005; Boets et al., 2011; Hamalainen, Salminen, & Leppanen, 2013; Leong & Goswami, 2014; Rosen, 2003). This fact rises 2 issues (a) central auditory processes which are not found to be impaired in population presenting SLI or dyslexia are under-represented in the literature compared to processes impaired in such population. For example intensity discrimination does not seem to be impaired and thus, is not much studied both in children and adult population (but see Ahissar et al., 2000; Buss, Hall, & Grose, 2006; P. R. Hill, Hogben, & Bishop, 2005; Jensen & Neff, 1993; Mengler et al., 2005). In contrast, frequency discrimination, which is suspected to be impaired in these populations is over-represented (Ahissar et al., 2006; Ahissar et al., 2000; Banai & Ahissar, 2006; Banai & Yifat, 2011; Banai & Yuvak-Weiss, 2013; Bishop & McArthur, 2005; France et al., 2002; Halliday et al., 2008; P. R. Hill et al., 2005; McAnally & Stein, 1996; Moore et al., 2008; Pakarinen et al., 2007). The second issue is that (b) being able to comprehend the normal development of central auditory processing is a prerequisite to the study of their dysfunction particularly as some authors claim that APD in developmental disorders may come from auditory processing immaturity (Bishop & McArthur, 2005; Wright & Zecker, 2004). The aim of the present study is therefore to provide a view of the development of central auditory processes in typically developing children between 6 and 11 years old. The results of these children will also be compared to a group of normal young adults to light up the adult outcome of the developmental process. Moreover, the here presented tests are entirely nonverbal as British Society of Audiology (BSA, 2011) recommends, results are therefore less influenced by the literacy skills of participants. Indeed, it is of particular interest to present non-verbal stimuli to children for the following reasons. First, as mentioned earlier, many children suffering from a language development disorder show APD and using verbal material to evaluate them can lead to an over-estimation of the co-morbidity between these two disorders. Second, language develops importantly during childhood, and even in typically developing children, at different rates; testing APD using verbal material could therefore enhance the confound between verbal and auditory abilities.

Although peripheral auditory processes are already efficient during pre-natal life and are mature within 30 days after birth, central auditory processes on the other hand are

longer to develop. Auditory cortices mature gradually until adulthood whereas midbrain and brainstem mature during the first years (Eggermont, 2014). Sensory encoding of sounds attributes is thus mature only in young adulthood, as reflected by the evolution of the Auditory Evoked Potential (AEP) P₁ and N₁. Both electrophysiological components occur around 100 ms after the presentation of the stimulus; with the positive one, P₁, reflecting the synaptic delay between peripheral and central auditory processes and the negative one, N₁, being evoked in the primary and associative auditory cortices in response to an unexpected stimulus. Although P₁ and N₁ are very robust in adults, they are not in children and infants. P₁ is already present in 3 months infants' responses to vowel (Shafer, Yu, & Wagner, 2014) but its amplitude increases in early childhood (up to 4 years old) and then decreases until adulthood (Shafer et al., 2014; Sussman, Steinschneider, Gumenyuk, Grushko, & Lawson, 2008; Wunderlich, Cone-Wesson, & Shepherd, 2006). N₁ on the other hand is present on its final form only at the end of adolescence (A. Sharma, Kraus, McGee, & Nicol, 1997; but see Wunderlich et al., 2006 for observation on some younger children). For both components, latency decreases with age due to myelination and augmentation of synaptic efficiency occurring during childhood and adolescence (Eggermont, 2014). Although these results indicate that auditory stimuli encoding is still immature in infancy, childhood and adolescence, the ability to discriminate sounds arises very early in development (Kral & Pallas, 2011). Indeed, immature MMN were found in response to a deviant pure tone (Draganova et al., 2005; Draganova et al., 2007) and piano tone (He et al., 2007) using an oddball paradigm in fetuses, newborns and 4 months infants. A mature MMN was found in 6 months infants using gap detection as the contrast (Trainor et al., 2003).

These findings imply that auditory processing matures with age and most importantly, as the here presented AEP are pre-attentive, these observations are independent of higher cognitive skills. However, testing several auditory processes with EEG would be extremely long, and therefore not possible for example in diagnostic situations. Settings used to behaviorally investigate auditory development are much shorter and easier to adapt for diagnostic purpose, for instance in case of APD. A secondary aim of our study is thus to develop a battery investigating a large part of central auditory processes and able to provide a reliable cue on auditory processing development using tests adapted to diagnostic situation (i.e., fast, sensitive and easily interpretable). This study is thus a first step in this direction testing the battery before it can be normalized on a larger population.

In addition, investigating many auditory processes on the same population makes the development of different processes more comparable. In addition, different dimension of the same process are tested with the same paradigm in the case of auditory discrimination and auditory pattern test. Indeed, looking at the literature on auditory processes development it appears that it is difficult to compare studies as children's

performances are very sensitive to the paradigm used. Sutcliffe and Bishop (2005) for example evidenced that children's ability to discriminate frequency varies depending on the paradigm presented and for instance on whether it is cued when to attend to the target tone or not. Children are also able to use the repetition of the comparison tone to remember it (i.e., form a perceptual anchor; Banai, Sabin, & Wright, 2011; Banai & Yifat, 2011) and therefore compare it to the target tone more easily. Thus, children's performances at discrimination tests improve when the comparison tone is constant (Banai & Yifat, 2011).

Another issue regarding the comparison between the different studies is the choice of threshold, as mentioned earlier, development of frequency discrimination for example is very well documented (Banai & Ahissar, 2006; Bishop & McArthur, 2005; Halliday et al., 2008; P. R. Hill et al., 2005; Moore et al., 2011; Sutcliffe & Bishop, 2005), however, none of the studies use the same paradigm and/or report the same threshold (i.e., 70,7%; 75% or 79%). This study will therefore enable to compare children of different age's performances and developmental variation across different dimensions. For example, discrimination ability will be tested on three dimensions (frequency, duration and intensity) using the same paradigm so their developmental trajectories would be comparable.

The aim of this study was twofold: (a) report developmental course of central auditory processing in school age children (b) using test adapted to a diagnostic use. On this purpose, participants were presented to 6 main tests composed of one or more subtest in order to evaluate central auditory processing, as they are defined by ASHA (2005). Participants aged from 6 to 11 years old therefore went through lateralization test, auditory discrimination test (involving frequency, duration or intensity), central masking test, auditory identification and recognition test, auditory pattern recognition test (involving frequency and duration) and stream segregation test. More precisely, this study will allow us to investigate whether all auditory processes develop in the same way (in term of speed, developmental step, age of maturity) as this would be a cue to underlying processes (i.e., different trajectories would suggest different processing; Dawes & Bishop, 2008). A young adult population was also tested for control and as evaluation of the outcome of the developmental process, it will thus be possible to investigate whether the here tested ability are mature by the end of childhood or continue to develop during adolescence.

2. Method

2.1. Participants

Eighty-nine children (37 boys) were administered the tests. They were classified following their school class from 1st to 5th grade corresponding to chronological age (see Table 1). They were recruited in two primary schools and were selected by teachers

for presenting no developmental disorder. They were all French native speakers and reported no hearing disorder. All of them had normal nonverbal IQ as shown by matrices subtest of the WISC (Wechsler Intelligence Scale for Children; Wechsler, 2005; see Table 1). Non-verbal IQ did not differ across Group, as evidenced by a one way ANOVA ($F(4.84) = 1.04, p > .10$).

A group of 23 young adults (18-22 years old; $M = 20.2$; $SD = 1.0$) has been tested. All adults were French native speaker with audiometric pure-tone thresholds < 20 dB on a frequency range from 250 Hz to 8000 Hz. All of them had a normal non-verbal IQ assessed by Raven standard progressive matrices ($M = 47.9$; $SD = 6.5$; Raven, 1938).

Participants reported no history of psychiatric or neurological disorder. Children and their representatives, and adults signed written informed consent. Children were given a diploma of good willingness and adults were compensated for their participation.

	1st grade			2nd grade			3rd grade			4th grade			5th grade		
	n	M	SD	n	M	SD	n	M	SD	n	M	SD	n	M	SD
AGE (in months)	18	80.66	4.38	18	91.72	4.36	19	102.42	3.39	15	114.67	5.37	19	127.47	5.56
Matrices score	18	10.28	2.59	18	9.67	3.05	19	9.68	3.65	15	11.27	3.69	19	9.21	2.62

Table 5

Mean age and matrices score for children depending on class group.

2.2. Auditory tests

Six main tests were presented to participants. Schematic representation of all tests is shown in Figure 34.

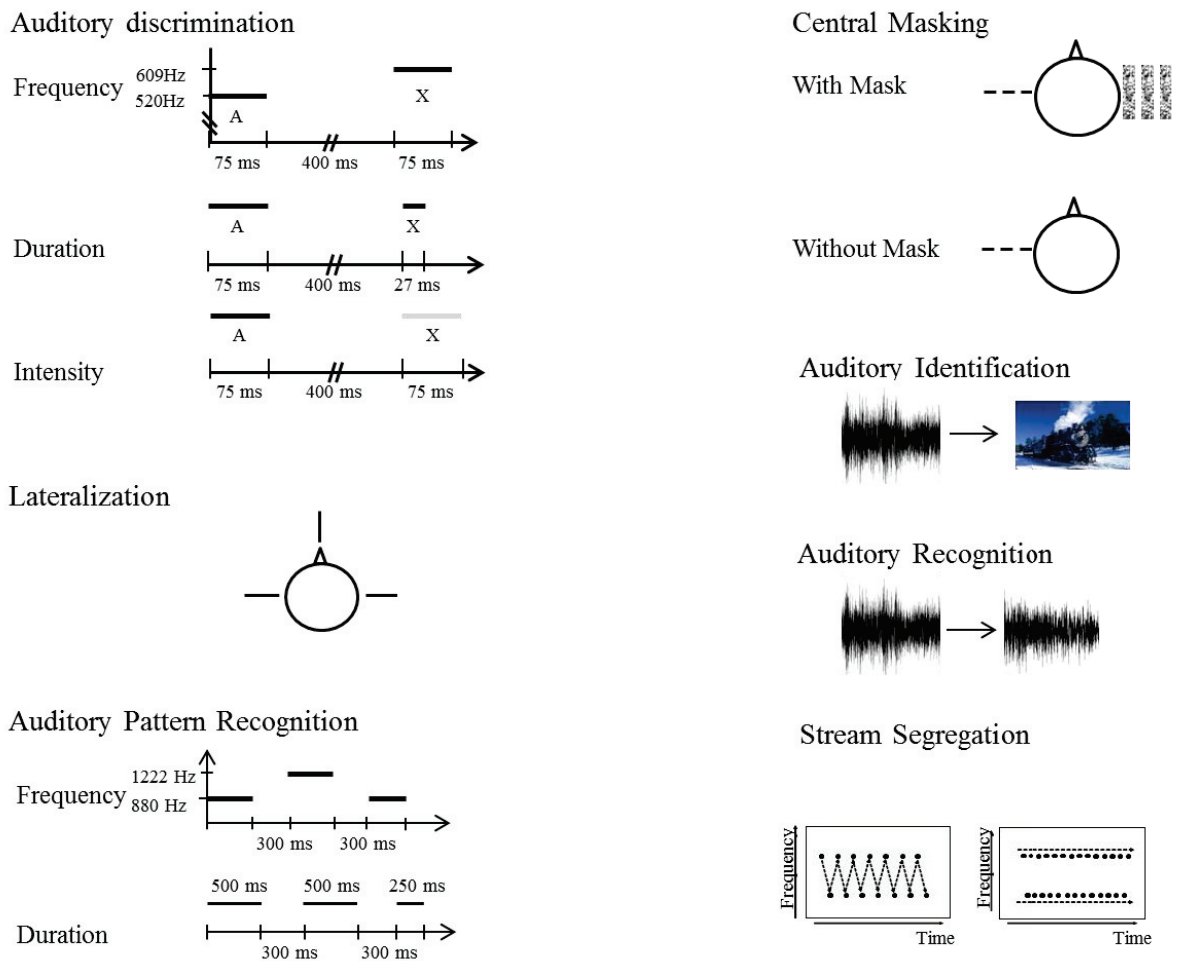


Figure 34

Schematic representation of all presented tests.

Lateralization. Lateralization skills were assessed by presenting binaurally a pure tone (250 ms; 500 Hz; rise-fall ramps 5 ms) with an Interaural Level Difference (ILD) varying from -5 dB to -25 dB by 2 dB steps (i.e., 11 ILD). Each ILD was presented 4 times (twice with the left ear intensity lower than the right ear's and twice with the opposite configuration). In addition, 22 of the stimuli were presented with a 0 dB ILD (i.e., the same level in both ear). This test therefore contained 66 stimuli. A training session providing a feedback and including 6 trials (2*-25 dB ILD; 2*-5 dB ILD and 2* 0 dB ILD condition) for adults and 12 (2*-25 dB ILD; 2*-21 dB ILD; 2*-15 dB ILD; 2*-11 dB ILD; 2*-5 dB ILD; 2*0 dB ILD) for children was included. Participants were asked to indicate the origin of the pure tone (i.e., left ear, right ear or the two ears), by pointing the ear where the sound was heard for children and by pressing the corresponding key for adults. No feedback was provided during the test session.

Discrimination. The discrimination test was composed of 3 subtests testing respectively frequency, duration and intensity discrimination abilities. An AX paradigm was presented where participants heard a first pure tone A (75 ms; 520 Hz;

rise-fall ramps 5 ms) following by a second pure tone X after a 400 ms ISI. The second tone could be the same as the first one or could vary in frequency, duration or intensity depending on the subtest:

For the *frequency discrimination test*, X could have a frequency of 527 Hz; 535 Hz; 546 Hz; 562 Hz; 583 Hz; 609 Hz.

For the *duration discrimination test*, X duration varied from 27 ms to 67 ms with 8 ms steps (i.e., 27 ms; 35 ms; 43 ms; 51 ms; 59 ms; 67 ms).

For the *intensity discrimination test*, X intensity compared to A's varied from -2.5 dB to -15 dB with 2.5 dB steps (i.e., -2.5 dB; -5 dB; -7.5 dB; -10 dB; -12.5 dB; 15 dB).

Each value of X was presented 4 times (i.e., 6 value of X * 4 trials = 24 trials) and X was the same as A in half trials (i.e., 24 trials). Subtests were thus made of 48 stimuli pairs randomly presented. A training session including a feedback was presented before testing. During this phase 12 trials were presented to children, in half of them X was different from A (i.e., one for each value of X). Feedback consisted of a happy or sad smiley face. Adults were presented to 6 training trials, half of them included an X different from A (i.e., the 1st, the 3rd and the last easier to discriminate from A). Feedback was provided by the computer displaying *correct* or *incorrect*. Participants had to indicate whether the two pure tones were identical or different by pressing a pre-specified key.

Central masking. The central masking test was composed of two subtests, a Without mask subtest and a With mask subtest.

In both subtests, 3 pure tones (200 ms; 500 Hz; rise-fall ramps 10 ms) separated by a 200 ms ISI were presented to one ear in half trials. Nothing was presented to the other ear in the Without mask subtest whereas 3 bursts of white noise (i.e., the mask) were simultaneously presented to the other ear in the With mask subtest. The intensity (in dB) of the pure tones ranged from 30 dB to 2 dB with 2 dB steps, each condition was presented twice for each ear. The mask intensity stayed at 60 dB, therefore in With mask condition the relative intensity of pure tones were -30 dB to -58 dB with 2 dB steps.

The other half trials were silent in the Without mask subtest, and only contained the white noise in one ear in the With mask subtest. Each subtest contained therefore 120 stimuli presented randomly, given the high number of stimuli and the monotony of the task, a break was inserted after the first 60 trials. Participants were asked to indicate whether or not they heard the 3 pure tones by pressing the corresponding key. Six training trials were proposed to adults and 10 to children. Feedback was provided by smileys for children and by displaying *correct* or *incorrect* to adults.

Auditory Identification and Recognition. These two subtests are the only one not involving pure tones. Both subtest consisted in presenting a target auditory sequence corresponding to a daily encountered object (e.g., a fire noise). Twenty-eight environmental sounds were presented to participants they lasted on average 12 sec (SD = 15). One practice item was presented in each subtest.

Auditory Identification was tested by presenting the subject with sound samples. For each sound samples, three pictures were simultaneously presented to participants. One of them representing the object producing the target sound samples (e.g., a fire), another one representing a semantically (but not acoustically) related object (e.g., fireman truck) and the last one consisted of an acoustically (but not semantically) related object (e.g., crumpled paper). Children participants had to point out the correct drawing and adults had to press the corresponding key.

Auditory Recognition was tested by presenting a target sound sample (e.g., a starting car) to the subject followed by three sound samples produced by daily encountered objects. One of them was produced by the same object that produced the target auditory sequence (but was not the same sequence; e.g., a driven car), another one was produced by semantically associated object (i.e., a horn) and the last one was produced by an object acoustically related to the target (e.g., a mower). The answer was given in the same way as for the Auditory Identification subtest.

Auditory pattern recognition. This test was composed of two subtests, one of them assessing the frequency and the other one the duration. Both consisted in the presentation of 3 pure tones separated by a 300 ms ISI. They included a training session containing 10 trials (a feedback was provided by the experimenter for the 5 first items). After this phase, 20 trials were presented binaurally. Children had to sing or hum the heard sequence, adults had to reproduce it on a piano keyboard, they did not have to produce the exact note, only the relative frequency or duration of the notes were taken into account.

For the *auditory pattern recognition on frequency*, the 3 presented pure tones had the same length (500 ms; rise-fall ramps 5 ms) but could vary in term of frequency. They could be 'High' (H; i.e., 1122 Hz) or 'Low' (L; i.e., 880 Hz). Six possible patterns were presented to participants: HHL; HLH; HLL; LHH; LHL; LLH.

For the *auditory pattern recognition on duration*, the pure tones frequency was held constant (1000 Hz; rise-fall ramps 5 ms) but their length could be 'Short' (S; i.e., 250 ms) or 'Long' (L; i.e., 500 ms). The six following possible patterns were presented to participants: SSL; SLS; SLL; LSS; LSL; LLS.

Stream segregation. The stream segregation test was the same as the one used by Lallier et al. (2009). Tone sequences of alternating pure tones of 1000 Hz and 400 Hz were presented to participants for 5 s. Each pure tone lasted for 40 ms (rise-fall ramps

10 ms). Participants had to report whether they perceived one stream or two streams according to a forced choice paradigm. A training phase was presented to ensure that participants understood the correspondence between the percept and the concept. For this purpose, two practice trials including an unambiguous 'one stream' percept (SOA = 400 ms) and an unambiguous 'two streams' percept (SOA = 50 ms) were presented and associated with corresponding drawings. After each sequence, children answered by pointing at the drawing corresponding to the pattern they had perceived; adults responded by pressing the corresponding key. The SOA was varied using a one-up, one-down adaptive procedure, to estimate the fission threshold (i.e., 50% chance of hearing one or the other percept; Levitt, 1971). As long as the answer was 'one stream', SOA was decreased. On the contrary, as long as the answer was 'two streams' the SOA was increased automatically. Each session included 30 sequences. The initial SOA (300 ms) was chosen to be perceived as 'one stream'. In the first trials, the SOA was decreased by steps of 40 ms. After the first answer reversal, the step was set to 20 ms, then to 10 ms after the second reversal and 5 ms after the third reversal. The stream segregation thresholds were then computed by averaging the SOAs of the last 10 trials (sequences 21 to 30).

2.3. General procedure

Children. Children were assessed the tests during two sessions in a quiet room at school. They could ask for a break when needed. Tests were presented using Presentation software (www.neurobs.com) except for the auditory pattern recognition tasks, and auditory identification and recognition tests which were presented with power point software (2010). Finally the stream segregation test was presented using E-prime2 (Psychology Software Tools, Pittsburgh, PA). Stimuli were presented through headphones (Sennheiser HD 280 pro) at an intensity of 65 dB-A.

Adults. Adults were tested in a quiet room at the Laboratory on Language, Brain and Cognition (L2C2). All tests were presented using E-prime2 (Psychology Software Tools, Pittsburgh, PA). Stimuli were presented through headphones (Sennheiser HD 448) at an intensity of 65 dB-A.

RT were not recorded, the use of different software between populations may therefore not have consequences on results.

3. Results

For each test and subtest an ANOVA evaluates the effect of Development on participants' performances (i.e., 1st grade vs 2nd grade vs 3rd grade vs 4th grade vs 5th grade vs Adults), post-hoc (HSD Tukey) were then calculated in order to investigate more precisely developmental trajectories. Data are presented in Table 2.

3.1. General remarks

3.1.1. Psychometric function

For the lateralization and the discrimination tests, the psychometric function fitting best each participant data was computed. We then computed for each participant the threshold where they were able to lateralize a pure tone or discriminate two pure tones correctly at 75%. This threshold was chosen because, as mentioned in the introduction, it is very difficult when working with children to distinguish perceptual developmental effect and more general cognitive developmental effect. Therefore, it is possible that some children would answer randomly without really performing the task for example because they were not strongly involved in the task. A 50% threshold would thus not really evidence a perceptual threshold. The aim of computing the psychometric function was (1) to compare our results with developmental literature, and (2) diminish the impact of non-sensory factors. Indeed, analyzing the threshold derived from the function decreases the impact of an unattended trial. Participants whom 75% threshold did not fall into the extreme value of the test stimuli were excluded from analysis. For each test and subtest, a χ^2 test was performed to ensure that the number of excluded participants was not affected by the developmental group (1st grade; 2nd grade; 3rd grade; 4th grade; 5th grade and Adults). The aim was to evaluate whether younger children were more likely to obtain thresholds not included within the extreme values of the test stimuli. The χ^2 test for the frequency discrimination subtest appeared to be significant; less adults were excluded from analysis based on their performances. Two adults (the best and the worst performers) were therefore excluded from the analysis of the frequency discrimination test so that the number of excluded participants was not different across developmental group (χ^2 $p > .05$).

3.1.2. Stream segregation test

The stream segregation test was built with an adaptive task; participants were therefore supposed to approach their fission threshold within the last ten trials. Participants who responded more than 5 times “one” or “two” streams in a row during the last 10 trials were considered as not understanding/performing the task and were thus excluded from the analysis.

Development of Central Auditory Processes

	1st grade		2nd grade		3rd grade		4th grade		5th grade		Adults	
	M	SD	M	SD	M	SD	M	SD	M	SD	M	SD
Lateralization ILD in dB	13.42	4.93	11.76	3.71	10.66	4.28	9.25	2.85	9.49	3.42	9.70	3.81
Discrimination												
Frequency (in Hz)	46.38	20.23	36.97	19.17	46.77	18.93	34.82	20.45	36.41	17.23	17.06	7.34
Duration (in ms)	31.22	10.63	29.60	9.55	26.79	10.46	19.84	5.36	23.92	7.40	18.50	5.78
Intensity (in dB)	7.59	3.47	7.18	3.20	7.19	1.96	6.49	2.96	4.75	1.91	3.53	1.51
Central Masking												
With Mask (Hit – FA)												
10 dB	0.55	0.54	0.75	0.30	0.80	0.21	0.86	0.28	0.80	0.27		
8 dB	0.56	0.35	0.60	0.29	0.68	0.27	0.74	0.34	0.79	0.26		
6 dB	0.48	0.35	0.56	0.40	0.56	0.37	0.70	0.32	0.76	0.33		
4 dB	0.44	0.37	0.50	0.27	0.68	0.29	0.49	0.47	0.62	0.37		
2 dB	0.30	0.30	0.40	0.41	0.61	0.26	0.50	0.38	0.55	0.42		
Without Mask (Hit – FA)												
10 dB	0.76	0.24	0.90	0.14	0.90	0.17	0.90	0.17	0.94	0.10		
8 dB	0.81	0.25	0.75	0.29	0.90	0.25	0.90	0.15	0.93	0.14		
6 dB	0.69	0.32	0.77	0.25	0.87	0.21	0.85	0.19	0.91	0.16		
4 dB	0.70	0.28	0.55	0.42	0.85	0.24	0.78	0.27	0.84	0.22		
2 dB	0.67	0.39	0.62	0.40	0.90	0.17	0.65	0.42	0.84	0.30		
Stream Segregation (in ms)	140.25	67.12	136.39	86.39	86.62	59.72	150.25	68.40	155.18	86.40	127.51	69.89
Auditory Identification												
Acoustic Errors	3.33	1.64	2.44	1.25	3.21	1.87	1.73	1.28	1.53	1.07	1.39	0.94
Semantic Errors	0.50	0.71	0.83	0.92	0.42	0.69	0.53	0.64	0.47	0.84	0.61	0.99
Auditory Recognition												
Acoustic Errors	7.06	2.65	4.50	1.58	5.37	1.54	4.53	1.92	3.68	2.00	2.70	1.18
Semantic Errors	1.06	1.06	1.78	0.81	0.79	0.92	0.87	1.19	0.32	0.48	0.91	1.08
Auditory pattern recognition												
Frequency	15.44	4.58	17.44	2.75	17.42	2.69	18.13	2.36	18.84	1.50	19.48	1.12
Duration	10.61	4.33	13.39	2.97	13.58	0.90	15.20	3.43	17.21	2.12	18.35	1.80

Table 6

Mean (M) and Standard Deviation (SD) for each test and developmental group.

3.2. Lateralization

Following exclusion rules (see section 3.1.a.), the analyses were performed on 96 participants. A one way ANOVA on the 75% threshold comprised between -2.5 dB and -25 dB revealed a significant effect of Development ($F(5,90) = 2.53$, $p < .05$; $\eta^2 = 0.12$), younger participants needed a larger ILD to lateralize a pure tone at 75% than older participants. Examining the trajectory of the data, it seems that thresholds tend to decrease from 1st to 4th grade with a slope of -1.36 between the two means of

performances ($p = .09$; $y = -1.36x + 14.76$) and then stabilizes, as shown by a flat developmental trajectory between 4th grade and Adults ($p = 0.99$; $y = 0.23x + 9.03$; cf. Figure 35.A). This would therefore suggest that maturation of lateralization processing occurs in late childhood, up to 4th grade.

3.3. Auditory Discrimination

3.3.1. Frequency

Ninety-seven participants had a threshold comprised between 7 Hz and 89 Hz and were therefore included in the analysis. A one way ANOVA on discrimination thresholds revealed a significant effect of Development ($F(5,91) = 7.35$, $p < .001$; $\eta^2 = 0.28$). Younger participants needed significantly higher differences between A and X to discriminate them at 75%. Although the comparison between the 1st and the 5th grades shows no significant differences ($p = .60$). When examining the data it appears that overall thresholds decrease slightly ($y = -2.21x + 46.89$) between 1st and 5th grades. This decrease is more important ($y = -19.34x + 55.75$) and it is significant between 5th grade and adults participants ($p < .01$). Fifth grade participants needed a significantly higher difference between A and X to discriminate them at 75% compared to Adults (cf. Figure 35.B). Therefore it appears that frequency discrimination capacities mainly mature between 5th grade and adulthood.

3.3.2. Duration

Thresholds of 100 participants comprised within the extreme values of X (i.e., 8 ms to 67 ms) were analyzed. A one way ANOVA revealed a significant effect of Development ($F(5,94) = 6.36$, $p < .001$; $\eta^2 = 0.25$) indicating that younger participants needed a larger difference between A and X to be able to discriminate them at 75%. Data show a decrease in the 75% thresholds between 1st and 4th grade ($p < .01$; $y = -3.7x + 36.1$) and a stabilization between 4th grade and Adulthood ($p = .99$; $y = -0.67x + 22.08$; cf. Figure 35.C). Overall these results suggest a maturation of duration discrimination ending in 4th grade.

3.3.3. Intensity

Intensity thresholds of 102 participants, ranging from -2.5 dB to -15 dB were analyzed. The ANOVA revealed a significant effect of Development ($F(5,96) = 7.54$, $p < .001$; $\eta^2 = 0.28$) suggesting that younger participants needed a greater difference between A and X to discriminate them than older. The developmental trajectory does not seem to be linear, indeed, thresholds seems to be stable from 1st to 3rd grade ($p = .99$; $y = -0.20x + 7.72$) and decreases between 3rd grade and Adulthood ($p < .001$; $y = -1.27x + 8.67$). It seems thus that the ability to discriminate intensity develops quite lately, since 3rd grade (cf. Figure 35.D).

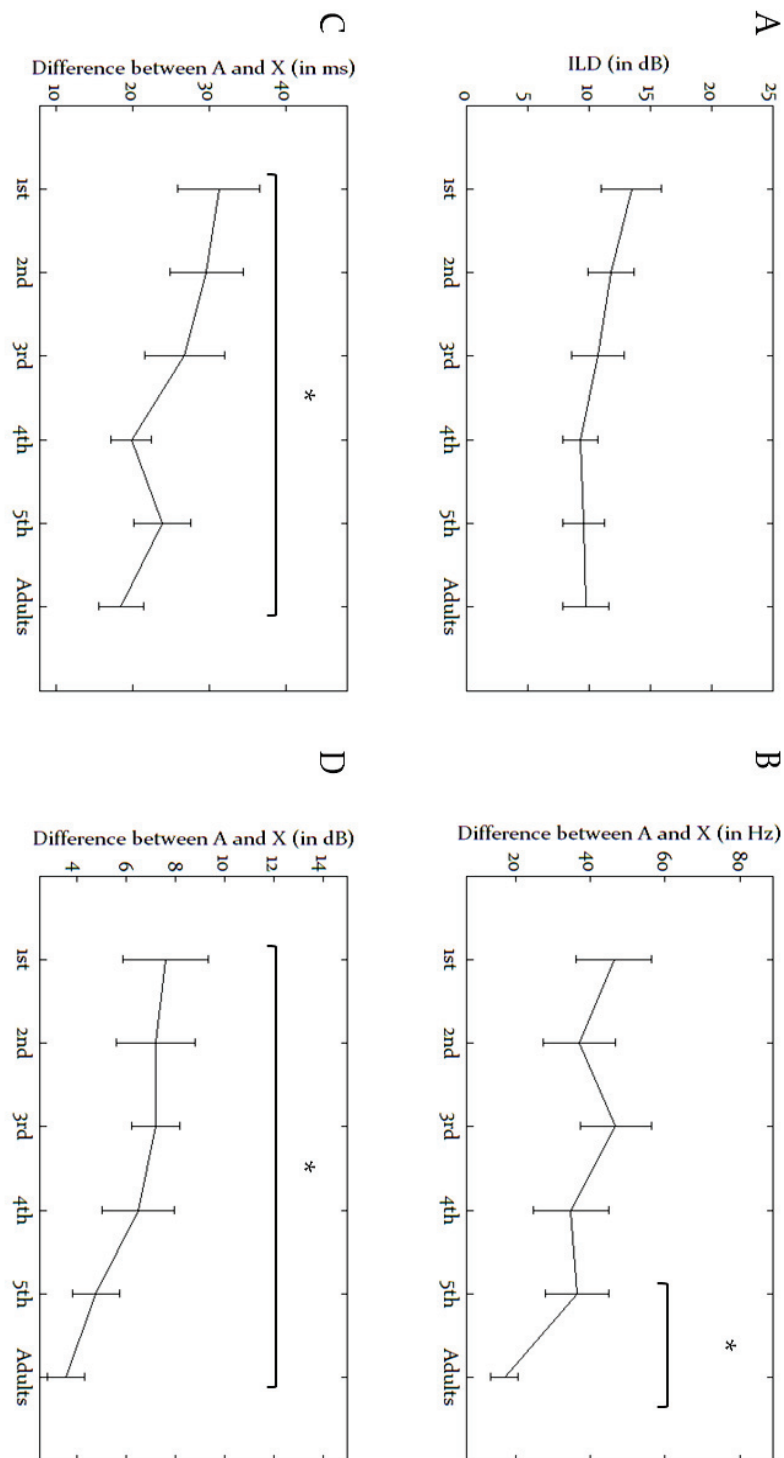


Figure 35

A. Interaural Level Difference (ILD; in dB) needed for participants to lateralize pure tone at 75% depending on the population. B. Frequency difference between A and X (in Hz) needed for participants to discriminate them accurately at 75% depending on the population. C. Duration difference between A and X (in ms) needed for participants to discriminate them accurately at 75% depending on the population. D. Intensity difference between A and X (in dB) needed for participants to discriminate them accurately at 75% depending on the population.

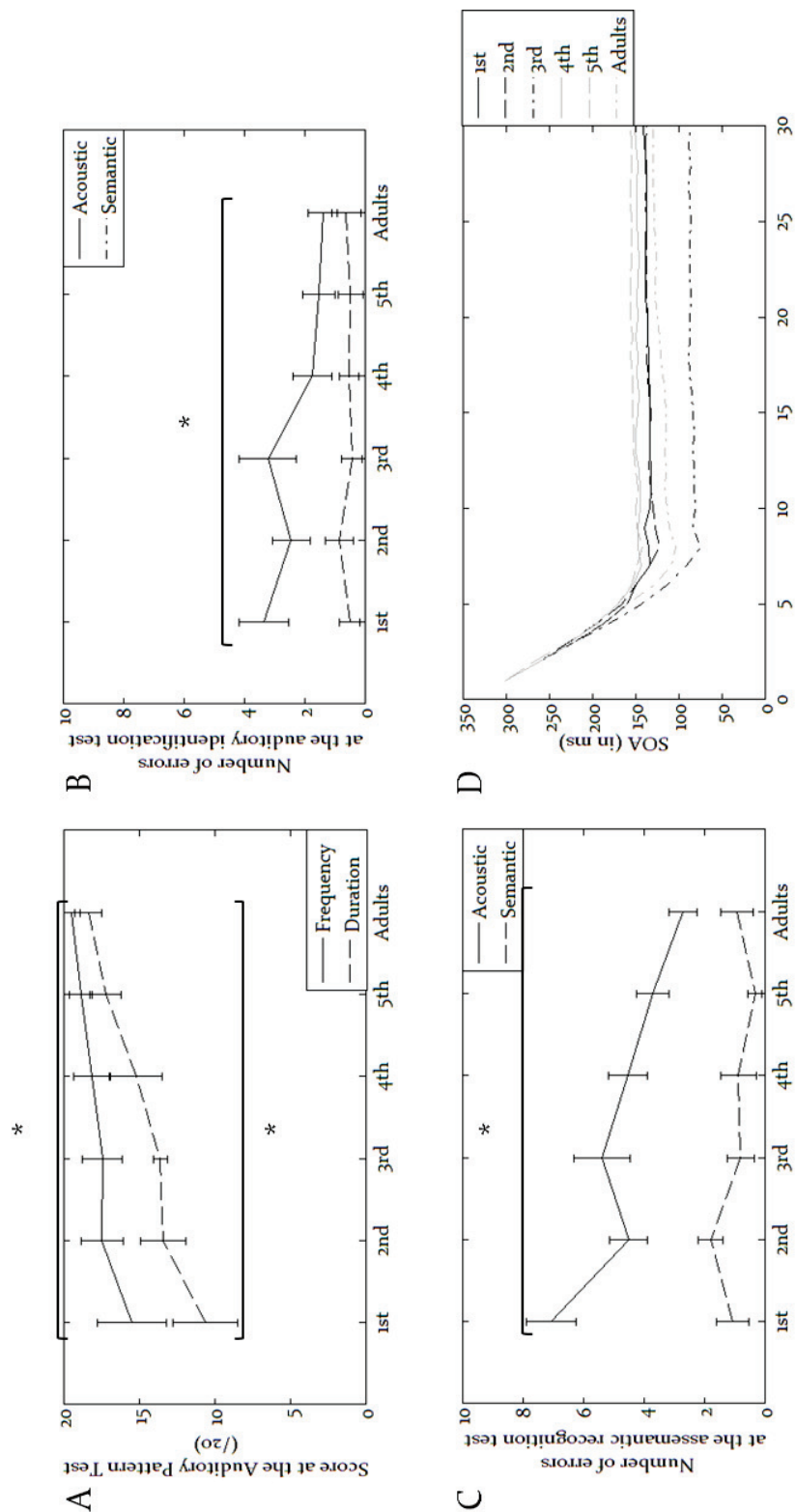


Figure 36

A. Auditory Identification test. Number of Acoustic and Semantic Errors depending on the population.
B. Assemantic Recognition test. Number of Acoustic and Semantic Errors depending on the population.
C. Auditory pattern recognition test. Score at Auditory pattern recognition test depending on the population.
D. Stream segregation test. Evolution of the SOA (in ms) depending on the trial number and population.

3.4. Central Masking

Due to technical issue, data of adult participants could not be analyzed. This test therefore only gets interested in children performances. Hit rates (proportion of “present” answers when the Tone was present) and false alarm rates (proportion of “present” answers when the tone was absent) were computed across all conditions for each participant.

A repeated measure ANOVA was performed on Hit – FA including the between subject factor Development and the within subject factor Mask (With vs Without) and Intensity of the pure tone to detect (10 dB vs 8 dB vs 6 dB vs 4 dB vs 2 dB). Development factor was significant ($F(4,74) = 2.66$, $p < .05$; $\eta^2 = 0.12$), older participants had better score than younger. As expected, Mask and Intensity factors were also significant. Participants being less accurate in the With Mask subtest than in the Without Mask subtest ($F(1,74) = 48.26$, $p < .001$; $\eta^2 = 0.39$) and being more accurate at detecting pure tones when they were presented louder ($F(4,296) = 34.82$, $p < .001$; $\eta^2 = 0.37$).

Only the factors Development and Intensity had a significant interaction ($F(16,296) = 2.57$, $p < .001$; $\eta^2 = 0.03$). Intensity effects are significant only once in 1st grade (2 dB vs 10 dB, $p < .01$), twice in 2nd grade (2 dB vs 10 dB, $p < .01$; 4 dB vs 8 dB, $p < .01$) and then mostly in 4th grade (4 dB vs 10 dB, $p < .01$; 4 dB vs 8 dB, $p < .01$; -58 dB vs -50 dB, $p < .01$; 2 dB vs 8 dB, $p < .01$ and 2 dB vs 6 dB, $p < .01$) and 5th grade (2 dB vs 10 dB, $p < .01$; 2 dB vs 8 dB, $p < .01$).

3.5. Auditory identification and recognition

3.5.1. Auditory identification

Number of error were analyzed with a repeated measure ANOVA, including Development as between subject factor and Type of Errors (Acoustic vs Semantic) as within subject factor. The effect of Development was significant, younger participants committed more errors than older ($F(5,106) = 5.79$, $p < .001$; $\eta^2 = 0.21$). Participants made significantly more Acoustic than Semantic errors ($F(1,106) = 116.76$, $p < .001$; $\eta^2 = 0.46$). Interestingly, the interaction between the two factors was also significant ($F(5,106) = 5.79$, $p < .001$; $\eta^2 = 0.11$). Indeed, as shown by the data, the number of Acoustic Errors is quite stable between 1st and 3rd grade ($p = 1$; $y = -0.06x + 3.12$), and decreases significantly between 3rd grade and Adulthood ($p < .001$; $y = -0.57x + 3.38$). The number of Semantic Errors on the other hand stays stable and very low from 1st grade to Adulthood ($p = 1$; $y = -0.01x + 0.60$; cf. Figure 36.A.).

3.5.2. Auditory recognition

A repeated measure ANOVA was performed on the number of error. It included Development as between subject factor and Type of Error as within subject factor (Acoustic vs Semantic). The effect of Development was significant ($F(5,106) = 23.40$, $p < .001$; $\eta^2 = 0.52$), younger participants committed more mistakes than older. All participants made significantly more Acoustic errors than Semantic errors ($F(1,106) = 312.11$, $p < .001$; $\eta^2 = 0.63$). The interaction between the two factors Development and Type of Error was also significant ($F(5,106) = 14.39$, $p < .001$; $\eta^2 = 0.15$). When looking at the data, it appears that the number of Acoustic Errors decreases with no obvious step from 1st grade to Adulthood ($p < .001$; $y = -0.14x + 1.46$) whereas the number of Semantic Errors stays stable and low during all development ($p = .90$; $y = -0.23x + 1.68$; cf. Figure 36.B).

3.6. Auditory Pattern Test (APT)

A repeated measure ANOVA including Development as between subject factor and Dimension as within subject factor was performed to be able to compare the developmental trajectories. Indeed, it revealed a significant effect of Development ($F(5,106) = 16.36$, $p < .001$; $\eta^2 = 0.43$) and Dimension ($F(1,5) = 64.61$, $p < .001$; $\eta^2 = 0.35$). In Frequency subtest, younger participants were less accurate than older, as can be seen on the Figure 36.C, the developmental trajectory does not seem to reveal developmental step from 1st grade to Adulthood ($p < .001$; $y = 0.72x + 15.28$). The same seems to be true for the Duration subtest, the developmental trajectory does not show any obvious developmental step from 1st grade to Adulthood ($p < .001$; $y = 1.48x + 9.55$; cf. Figure 36.C).

In addition participants were less accurate at performing the Duration subtest than the Frequency subtest. The interaction between these two factors was significant ($F(5,106) = 2.6$, $p < .05$; $\eta^2 = 0.07$). A post-hoc test (HSD Tukey) showed that only children from 1st ($p < .001$), 2nd ($p < .01$) and 3rd ($p < .01$) had lower scores at the Duration subtest than at the Frequency subtest.

3.7. Stream segregation

A one way ANOVA was performed on the mean of the last 10 trials for 100 participants. The effect of Development was not significant ($F(5,94) = 1.86$, $p > .05$; $\eta^2 = 0.09$), suggesting that all participants had the same ability to perform stream segregation based on ISI variation (cf. Figure 36.D).

3.8. Inter-individual differences

Data of each child was also compared to adult population. Children's results falling within adults performances range (i.e., $> \text{Adults mean} - 2 * \text{Adults SD}$) were looked for.

Interestingly, in all but two tests (i.e., frequency and intensity discrimination) more than 50% of children performed within adults' performances range. When excluding the stream segregation test (showing no significant developmental effect) and investigating individual performances, only two children do not reach adults performance at any test (one 1st and one 3rd grade). Four children in 5th grade and 1 in 4th grade showed adult-like performances in all tests. In addition the mean number of tests where children had adult-like performances increased significantly during primary school as shown by an ANOVA ($F(4,84) = 11.26$, $p < .001$; $\eta^2 = 0.42$; cf. Figure 37).

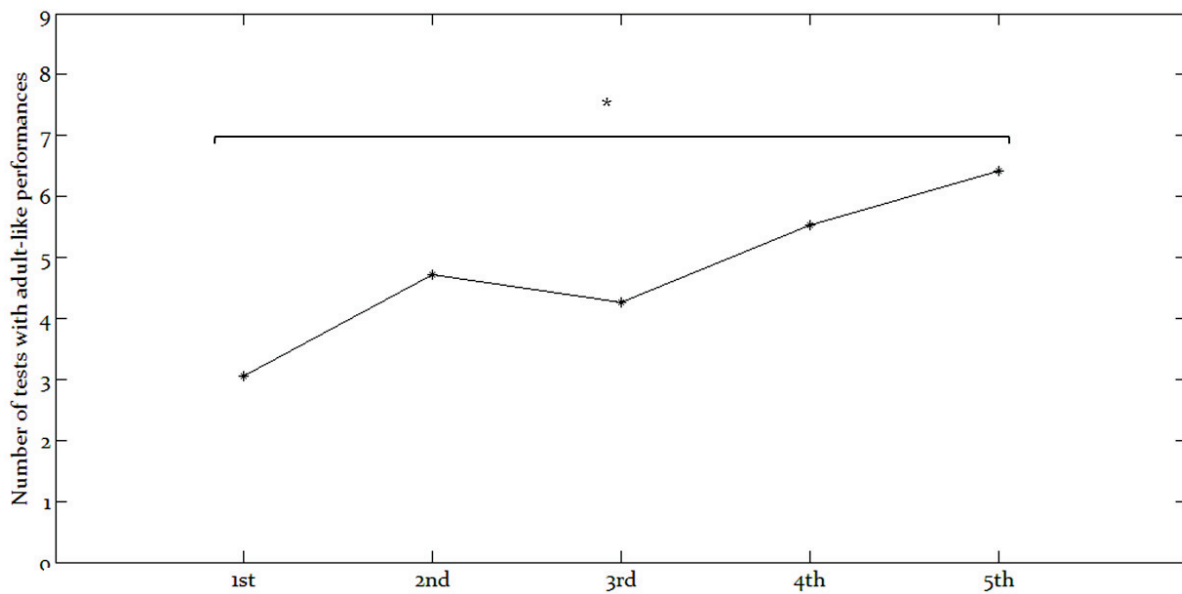


Figure 37

Graphical representation of the mean number of tests where participants obtained adult-like performances depending on the group.

4. Discussion

The aim of this study was to investigate behaviorally the development of Central Auditory Processes in typically developing children in primary school using paradigms adapted to a diagnostic situation. A group of young adults was also tested to evaluate the outcome of this development. Six main tests were presented to each participant; they were developed in order to evaluate efficiently and rapidly central auditory processes. In addition none of these tests included language; they thus can be used on population with language development disorders. We will discuss our results and compare them to literature when possible, to evaluate the validity of our simplified paradigms.

First of all, the lateralization test showed that the ability to use ILD evolves during childhood, and appears to be mature by 4th grade. Indeed, the needed ILD to accurately lateralize a pure tone decreases from 1st (13.42 dB) to 4th grade and then stabilizes around 9 dB. No literature measuring lateralization thresholds using ILD in a developmental view was found. Indeed, although Van Deun et al. (2009) got interested in the development of lateralization in children from 5 to 9 years old, they tested the ability to use Interaural Time Difference (ITD), however they did not evidence any developmental effect (except between children and adults). These results combined to ours are very interesting as they might suggest that the ability to use ILD mature earlier than the ability to use ITD, and thus that cues used to lateralize sounds might evolve during development. More recently, Moossavi and colleagues (Moossavi, Mehrkian, Lotfi, Faghihzadeh, & Sajedi, 2014) conducted a study on lateralization abilities in 9 years old children using ILD. Unfortunately, they did not computed any threshold and only report their results in terms of percent correct lateralization. However, a very recent study conducted by Altmann et al. (2014), evidenced that adult listeners needed a mean ILD of 6.2 dB (5.4 dB on one side and 7 dB on the other side) to accurately lateralize a 500 Hz pure tone at 50%. As our threshold was set at 75%, their results seem to be congruent with ours.

As mentioned in the introduction, literature on discrimination is much more complete, particularly for frequency discrimination. As it is known that frequency discrimination develops differently depending on the pure tone frequency (Eggermont, 2014), and as studies use pure tones of different frequency, threshold values will be expressed in percent of the standard (e.g., a 30 Hz needed difference between A and X, when A is set at 500 Hz corresponds to a 5.7% needed difference). Our data show that frequency discrimination between 520 Hz and 609 Hz does not significantly improve during childhood. Thresholds however significantly decrease between 5th grade and adulthood (from 7% to 3.3%). This suggests thus that frequency discrimination matures during adolescence. Our results are congruent with the literature where the needed difference between two pure tones is comprised between 5% and 10% when threshold is set at 79% (due to the use of an adaptive 3-down/1-up staircase procedure; Banai & Yifat, 2011; Banai & Yuvak-Weiss, 2013; Halliday et al., 2008; Moore et al., 2008). Although our results are comprised within the literature range, one must note that no significant developmental effect is found between children in our study, although a difference is usually found in the literature. However, as frequency processing in higher spectral range seems to be mature earlier than lower ranges (Werner & Marean, 1996), it would not be surprising that the ability to discriminate pure tones in the tested spectral range (i.e., 520 Hz to 609 Hz) does not evolve during childhood but during adolescence, as shown by the significant difference observed between 5th grade and adults.

Concerning the duration discrimination subtest, the choice was made to use shorter pure tone than other studies (75 ms vs 400 ms for Jensen and Neff, 1993 or 200 ms for Banai et al., 2013). The aim was to tackle durations that were closer to temporal variation in speech (e.g., Voice Onset Time; VOT). The ability to discriminate duration appears to improve between 1st (31.2 ms; 41.6%) and 4th grade (19.40 ms; 26.4%). It then stabilizes until adulthood (18.45 ms; 24.6%). As for previous studies, our data show a significant developmental effect for duration discrimination, confirming that it develops along childhood (Werner & Marean, 1996). Overall this pattern is coherent between studies, even though the duration of the standard used is very different (from 75 ms to 400 ms). However, looking in more details, it appears that while in all studies adults thresholds seems to converge around 20 ms, the results observed for children are more variable. This suggests that children's ability to discriminate duration depends on the duration of the stimulus, indeed, a 20 ms difference on a pure tone lasting for 75 ms is much more noticeable than on a 400 ms pure tone. However, it seems that maturation leads to a threshold more independent from the stimulus length in adulthood.

Our intensity discrimination subtest shows no developmental effect between 1st (7.59 dB) and 3rd grade (7.19 dB) and then, a decrease between 3rd grade and adulthood (3.53 dB). As mentioned in the introduction, very few data were found on development of intensity discrimination. However, in line with our results, a study by Buss et al. (2006) revealed that children from 5 to 10 years old need, on average 6.7 dB, to accurately discriminate 500 Hz pure tones at 70.7%, whereas adults needed 1.5 dB only.

The central masking test aimed at investigating whether the processing of a pure tone would be disturbed by the presence of a simultaneous mask in the contralateral ear. Results evidenced an increase of accurately detected pure tones in development. However, although all participants were less performant when presented to the With Mask subtest, it does not seem that the effect of Mask evolves with Development. The lack of data from adults prevents us to state whether sensitivity to a contralateral masker is mature before 1st grade or whether it matures during adolescence. Further investigations are needed to answer this question.

The Auditory Identification scores showed that the ability to pair a daily encountered sound to a picture without explicitly using language evolves during childhood. Interestingly, the number of acoustic errors is stable from 1st (3.33) to 3rd grade (3.21) and then decreases until adulthood (1.39) whereas the number of semantic errors does not evolve during childhood and during adolescence. This means that the improvement relies on auditory processing efficiency *per se* and not on the knowledge of concepts. Indeed, participants need to know and be able to recognize the presented concept as the association between an object sound and its picture implies that the object is recognized. The Auditory Recognition test highlights that participants were

able to match two different sounds produced by the same object or living being without visual support. This ability evolves during primary school, while once again, semantic errors did not decrease during development and stay constant. This suggests that participants' scores increase due to a better auditory recognition process.

The Auditory pattern recognition tests aimed at evaluating the ability of children to perceive, remember and reproduce sequences of three pure tones differing on one dimension (i.e., frequency or duration) without oral verbalization.

In the frequency subtest, performances increased linearly, with no obvious developmental step, to reach a quasi-perfect score in adulthood. Results obtain at our frequency subtest are around 20% better than those reported by Musiek (2002) with the Frequency Pattern Test developed by Musiek and Pinheiro (1987). However there were several differences in the paradigm used that could explain these differences. First of all, the duration of stimuli is longer in our frequency subtest than in the FPT. The rhythm of presentation is therefore slowed down and might easier the task. Second, the number of trials is smaller, one can assume that participants are less tired; more focused and thus are more performant in our test. Finally, the last difference is the fact that participants in previous paradigm are required to answer verbally to the test, whereas they were asked in our subtest to hum or sing the sequence for children and to reproduce it on a piano keyboard for adults. The intervention of verbal ability in this test might help explaining lower performances in FPT, because children could easily have difficulties remembering which word corresponds to which percept (particularly as 8 years old children might not be very familiar to 'high frequency' and 'low frequency' concepts). However, although the observed performances are better than those reported on the FPT, the two developmental trajectories are the same, with an increase of performances during primary school and maturation occurring around 10 years old. It seems thus, that after 10 years old the ability to judge, retain and recall a frequency pattern is mature.

The duration subtest was developed in accordance with the Duration Pattern Test (Musiek & Pinheiro, 1987). Children's performances increase regularly with age from 66% at 7 to 92% in adulthood. Again performances in our study are slightly better than those reported with the DPT (Bellis, 2003, cited in Brooke Shinn, 2014). The difference observed between the two set of data is very large around 7 years old (25% correct answers compared to 66% in our study) and much smaller around 10 (70% vs. 86% in our study). Adults on the other hand, seem to perform equally in our test and in the DPT (Musiek, 1994). The number of trials, the ear of presentation (both ear in our test vs one after the other in FPT and DPT) and the type of response differed between the two paradigms and all of them can certainly explain the differences in performances. However, developmental trajectories for both tests seem to be equivalent although our results are slightly better.

The stream segregation task did not reveal any significant developmental effect. Although it is known that stream segregation occurs very early in development (Smith & Trainor, 2011). It seems that segregation competences do not evolve during childhood and adolescence, at least when using ISI as cue for segregation. Lallier et al., (2009; using exactly the same task as us) already evidenced that the ISI corresponding to the change of percept (1 stream vs 2 streams) does not evolve between 10 years old good reader children and normal adults. In the light of our results it therefore seems that stream segregation based on ISI variation is mature as soon as 1st grade. As for lateralization abilities, it is possible that cues used for segregating streams efficiently change during childhood, for example Sussman et al. (2007) found a developmental effect on stream segregation task using frequency as cue for segregation.

Overall, our results show that development of central auditory processes in children is neither homogenous nor linear. They also confirm that its development is long lasting. Indeed, only 3 of the tested processing are mature before the end of childhood (i.e., lateralization, duration discrimination and stream segregation). Other tested processes reach adult efficiency during adolescence. While these trajectories are of mere interest, it has to be underlay that development of central auditory processes is also very variable across children. For example, 6 children (2 in 4th grade and 4 in 5th grade) reached adult level for almost all tests.

The high variability of children's performances also highlights the question of the impact of non-sensory factors (Moore et al., 2008). Indeed, an issue raised when working on children development is whether their results mainly rely on sensory or non-sensory factors. Children indeed present lower cognitive abilities than adults; a lack of attention and working memory could explain the observed differences when comparing children to adults. When studying development one main question is therefore whether developmental effect originates in part from attentional and cognitive differences between younger and older children but mostly how much of the results can be explained only by non-sensory factors (Halliday et al., 2008; Moore et al., 2011; Moore et al., 2008; Sutcliffe & Bishop, 2005). Interestingly, in our study, developmental trajectories are different despite different dimensions of the same process were tested using the same paradigm (i.e., discrimination test; auditory pattern recognition test). It therefore suggests that the observed developmental effects do not only rely on cognitive development (but see Moore et al., 2008), but on auditory developments *per se*.

5. Conclusion

The first aim of our study was to get interested in the development of central auditory processes as they are defined by ASHA (2005) in primary school children, and to compare it with the outcome of development represented by adulthood. The second aim was to use an entirely non-verbal battery sufficiently brief to be used in a

diagnostic situation. An entirely non-verbal battery was thus developed, using only pure tones and daily encountered noises. Results show a developmental effect for all tasks except stream segregation. Development appears to be non-linear with developmental step differing across different tests and subtests. Except stream segregation test which did not reveal any development effect, only two tested abilities appeared to be fully mature during childhood: lateralization based on ILD and duration discrimination. The other tested abilities appeared to pursue their development during adolescence confirming that auditory development is a very long lasting process.

Acknowledgment

The first author is funded by a PhD grant from Région Rhône Alpes.

6. Résumé du chapitre IX

Cet article présente l'intégralité des données développementales recueillies par Donnadieu et son équipe grâce à la BECAC. A ces données sont ajoutées celles d'un groupe de jeunes adultes, afin de pouvoir déterminer la finalité du développement des traitements auditifs centraux.

Les données recueillies sur ces enfants ont permis de mettre en évidence tout d'abord que le développement des traitements auditifs centraux est long et se poursuit durant l'adolescence. L'ensemble des tests ont mis en évidence un effet significatif du développement, à l'exception de la ségrégation de flux, qui d'après nos données semble être mature avant l'école primaire. Le développement des autres compétences testées ne parait pas linéaire, mais présente des étapes de développement. Enfin, seules trois des compétences testées sont matures avant l'adolescence, la latéralisation, la discrimination de durée et la ségrégation de flux.

L'Etude 3 a ainsi permis de s'intéresser au développement des traitements auditifs centraux et à leurs liens avec le développement du langage chez les enfants sains. Afin d'explorer plus avant l'impact entre ces traitements auditifs et les compétences langagières, une dernière étude a été mise en place, s'intéressant au TTA au sein de la population dyslexique adulte.

Chapitre X

Etude 4. Les Troubles du Traitement Auditif dans la dyslexie

Comme présenté au sein de la partie Théorique, certains auteurs postulent que la dyslexie pourrait être causée par la présence de TTA. Il a en effet été mis en évidence qu'environ 40 à 60% de la population dyslexique présentait un déficit du traitement auditif comparativement à la population contrôle. Ces résultats sont toutefois controversés, en effet, certaines études ne retrouvent pas de déficit auditif (Mody et al., 1997; Nitttrouer, 1999). D'autres suggèrent que la nature des liens entre les TTA et la dyslexie n'est pas causale (Rosen & Manganari, 2001).

De plus, la littérature n'est pas homogène, en effet, des différences de méthodologie génèrent des résultats différents, c'est-à-dire que l'utilisation d'un paradigme entraîne l'apparition d'une différence significative entre les deux populations alors qu'un autre paradigme testant le même traitement ne met pas en évidence de différences même au sein d'une seule population (Ahissar et al., 2006; Amitay et al., 2002).

L'objectif de l'Etude 4 est donc de tester une population adulte dyslexique sur la BECAC dans le but d'évaluer les traitements auditifs centraux comme ils sont définis par l'ASHA (2005). Cette population dyslexique est comparée à une population contrôle appariée en âge, intelligence non verbale et latéralité manuelle. Comme nous l'avons vu dans la partie Théorique et dans l'Etude 3 (a et b) la BECAC permet de tester, sur certains traitements, différentes dimensions avec le même paradigme (i.e., discrimination auditive et test de pattern auditif). L'utilisation sur ces tests d'un seul paradigme pour différentes dimensions du son permettra de mieux distinguer la nature du déficit s'il en est un.

De plus nous testons dans cette étude différents niveaux de compétences langagières : la conscience phonologique, la lecture et le vocabulaire ; nous évaluons également la mémoire à court-terme et la mémoire de travail. L'investigation des liens entre les traitements auditifs centraux et les différents niveaux de représentations langagières permettra ainsi d'évaluer si un déficit de bas niveau peut impacter des compétences

relativement haut niveau, comme nous avons pu le mettre en évidence dans l'Axe 1 dans le cadre d'une dégradation transitoire de la qualité du signal de parole.

Enfin nous tentons, dans l'Etude 4 de distinguer des profils de déficits spécifiques qui pourraient donner lieu à différents types de dyslexies, ou différents degré de dyslexie.

Auditory Processing Disorders in dyslexic adults.

Marie Dekerle ^{1,2}, Sophie Donnadieu ^{3,4}, Fanny Meunier ^{1,2}

- (1) Laboratoire sur le Langage, le Cerveau et la Cognition, L2C2, UMR 5304, Lyon, France
- (2) Université de Lyon, Lyon, France
- (3) Université de Savoie, Chambéry, France
- (4) Laboratoire de Psychologie et NeuroCognition, CNRS – UMR 5105, Grenoble, France

ABSTRACT

This paper aims at investigating Auditory Processing Disorders (APD) in a population of dyslexic adults. APD were screened using a new battery especially developed to test Central Auditory Processes as they are defined by ASHA (2005). This battery therefore evaluates lateralization; auditory discrimination; central masking; auditory pattern recognition; auditory identification and asemantic recognition; and stream segregation. In addition, as recommended by BSA (2011) it is entirely non-verbal, therefore the impact of speech deficit is minimized. Participants' verbal abilities were also evaluated at different level including phonological awareness, reading and vocabulary. Short-term and working verbal memories were also tested. Performances were compared to those of a control group, matched in terms of age, non-verbal intelligence and laterality to dyslexics. Results evidenced deficits in tests involving spectral and temporal processes (i.e., frequency discrimination, duration discrimination, Frequency Pattern Test and Duration Pattern Test). Interestingly, only those tests showed significant correlations with verbal abilities, both at the population level (dyslexics + controls) and at the dyslexic group level. Our data show that 9 out 20 dyslexics were impaired in at least 2 tests whereas 2 controls reached this threshold. Our results suggest that spectral and temporal processes are deficient in dyslexic population and that these abilities are linked to phonological awareness and reading efficiency. Interestingly, none of the tested auditory abilities were linked to higher level of representation (i.e., vocabulary).

1. Introduction

Dyslexia is a developmental disorder defined by a reading delay of more than 18 month in normally intelligent children, with no neurological or psychiatric disorder and with a regular access to education. Its prevalence is evaluated at 7% (Fluss et al., 2009).

As dyslexia is only defined in a behavioral way, its definition gives no cue on the underlying mechanisms of the reading deficit. However, it is now commonly agreed

that dyslexia is underlied by a phonological deficit (but see Bosse et al., 2007; and Hari & Renvall, 2001 for attentional hypothesis). Dyslexic people would show a deficit whether in representation, storage or retrieval of phonological information (Ramus, 2003; Ramus & Szenkovits, 2008; Sprenger-Charolles et al., 2000). As learning to read in an alphabetic system implies to be able to convert a grapheme into a phoneme, the phonological deficit would be at the origin of the inability for dyslexic children to acquire reading, despite normal intelligence. Indeed, mapping a grapheme to a phoneme when its representation or the access to its representation is impaired might be problematic. One of the proposed hypotheses to explain this phonological deficit is the presence of subtle low level auditory deficit (Ahissar, 2007; Baldeweg et al., 1999; Goswami et al., 2002; Hornickel & Kraus, 2013; Ramus et al., 2003; Tallal & Gaab, 2006). This auditory deficit is supposed to prevent efficient processing of sounds, including speech. It would lead to poor phonemic representation weakening grapheme-phoneme conversion and therefore preventing efficient reading/reading acquisition.

Historically, Tallal and colleagues were the first to highlight this issue on dyslexics (Tallal, 1980) after investigating SLI (Tallal & Piercy, 1973). Dyslexic children showed deficits in both frequency discrimination and a frequency sequencing test when ISI were short. According to the temporal deficit hypothesis, the inability to process accurately rapid auditory stimuli would lead to difficulties in processing rapid formant transient, which are indispensable cues to perform accurate phoneme discrimination (Tallal & Gaab, 2006). Since then, the defect in frequency discrimination in dyslexics has largely been emphasized both in children and adults even when using long ISI (800 ms, Ahissar et al., 2000; 1000 ms, Banai & Ahissar, 2006; 400 ms, France et al., 2002; 300 ms, McAnally & Stein, 1997). Thus, if the hypothesis of a low level auditory deficit seems to pertain, it might not be specific to rapid temporal stimuli. Other low level auditory processes were found to be impaired in dyslexics using tasks like duration discrimination (Banai & Ahissar, 2006), intensity discrimination (McArthur & Hogben, 2001), Temporal Order Judgment (Tallal, 1980), Frequency Pattern test (Billiet & Bellis, 2011) dichotic listening (Billiet & Bellis, 2011), stream segregation (Hari & Renvall, 2001; Helenius, Uutela, et al., 1999; Lallier et al., 2009), amplitude modulation (Goswami et al., 2002), rhythm detection (Leong & Goswami, 2014), binaural unmasking (McAnally & Stein, 1996) and backward masking (Rosen & Manganari, 2001).

Despite the large amount of evidences, the theory of an underlying low level auditory processing is controverted for two main reasons. First, some studies failed to find a deficit in rapid formant transient which is supposed to be at the origin of the phonological deficit (Mody et al., 1997; Nitttrouer, 1999; Rosen & Manganari, 2001). However, it might be important to state that these studies also present some methodological weaknesses that could at least in part explain their inability to evidence a deficit. For example, Nitttrouer's (1999) choice of participant might be questionable; indeed, author chose to exclude children performing worst than 30th

percentile at the sounds in words subtest of the Goldman-Fristoe Test of Articulation (Goldman & Fristoe, 1986). Although it is justified to exclude children presenting language production deficit, this criterion might be too strict and might exclude particularly phonologically impaired children particularly as some children with dyslexia present deficit in speech too (Bishop & Snowling, 2004; Kamhi & Catts, 1986). Additionally, the criterion used to discriminate poor phonological processing group was to be one standard deviation under the mean at the reading subtest of the wide range achievement test revised (WRAT-R; Jastak & Wilkinson, 1984). Once again, this criterion is questionable as 1 SD is not supposed to be deviant. It is therefore possible that their test group did not include only dyslexic children, (as suggested by the percentage of their population included in the Poor Phonological Processing group, 16% which is way more than the 7% occurrence of dyslexia). Therefore, the fact that poor readers do not present deficit in rapid formant transient is interesting but does not allow strong conclusion on the dyslexic population. Rosen and Manganari (2001) study might also be subjected to criticism, as their sample is very small; only 8 dyslexic teenagers were tested (see Tallal & Graab, 2006 for a critical view of Mody et al., 1997). Overall, it seems that although not all dyslexics present low auditory processing (Ramus et al., 2003; White et al., 2006), 40% to 60% of them depending on the study present difficulties in different low level auditory processes.

The second issue weakening the hypothesis of a causal role between APD and dyslexia is that no typical profile has been found. Indeed, most of studies test only one or two processes, or did not evidenced any regular pattern of auditory deficit (Ahissar et al., 2000; Ramus, 2003; White et al., 2006).

One aim of our study is to test dyslexic participants on a wide range of central auditory processes, as they are defined by ASHA (2005). This study indeed aims at providing a wider view of potential APD in dyslexics in order to evaluate different processing and potentially to distinguish different profiles. A new developed battery already tested on typically developing children and normal adults (Dekerle et al., submitted) will be used. This battery provides 6 main tests made of one or more subtest evaluating: lateralization, discrimination, central masking, auditory pattern test, auditory recognition and identification and stream segregation. This battery presents 2 main advantages (1) it is entirely non-verbal, minimizing the interferences with verbal deficits (2) two of its main tests contains subtests allowing to test several dimension of them process with one paradigm so they can be compared.

Testing different aspect of auditory processing using the same paradigm is very interesting as it has been largely emphasized that the occurrence of APD in dyslexic population highly relied on the nature of the task. For example, when changing the paradigm, Ahissar et al. (2006), found or not an impairment in frequency discrimination in dyslexics adults. More precisely, when presented to an AX paradigm

(i.e., first pure tone A is always the same frequency whereas the discrimination has to be made on the second pure tone X) dyslexic participants perform poorer than controls. However, when frequency of both pure tones change, dyslexics are not less accurate than controls at determining whether they are same or different. Authors thus suggested that more than having low auditory deficit, dyslexics present a difficulty to form a perceptual anchor. The wide range of paradigms used to evaluate dyslexics' auditory abilities across studies renders them difficult to compare. Our battery uses the same paradigm to test discrimination of 3 different dimensions of sound frequency, duration and intensity discrimination. Similarly the auditory pattern test tested two dimensions using the same paradigm: frequency and duration. This regularity across dimension presents the advantages of (a) make them comparable and (b) investigating the impact of the used paradigm. Indeed, if all three discrimination abilities are found to be impaired in the same way in dyslexics, this could mean that dyslexics are less performant on this specific paradigm. However if not all dimensions are impaired, this would rid away the hypothesis of a simple inability of dyslexics to perform the task.

Finally, verbal abilities of participants will also be evaluated so their links to central auditory processes will be investigated. Indeed, at the developmental literature, it appears that depending on population cognitive performances and sensory acuity are more or less linked. For example, it has been evidenced that older people's cognitive performances are linked to their sensory acuity more than younger people's (Baltes & Lindenberger, 1997; Lindenberger et al., 2001). The effect of sensory acuity is also well known during infancy, for instance, deaf children present speech development delay even after being implanted, they also show less interest to speech (Houston et al., 2003). In a less radical way, infants presenting a history of otitis with effusion show speech perception deficit like phoneme discrimination compared to control infants, even when auditory acuity is recovered (Polka & Rvachew, 2005). In young age, it appears that sensory abilities are much linked to cognitive abilities, and it decreases during development while cognitive abilities automatized. Following this rationale, it could be that in dyslexic population, low level auditory processes are still strongly linked to verbal performances even in young adults. In our study we will compare the link between participant's abilities in central auditory processes and language processes for dyslexics and controls.

Using a non-verbal battery will enable us (1) to test participants presenting language disorders. This is interesting for two main reasons, (a) one of the literature debate is that dyslexics auditory processing deficit is specific to speech (Mody et al., 1997; Schulte-Korne, Deimel, Bartling, & Remschmidt, 1998), therefore, testing their auditory processes with speech stimuli will be helpless to answer this question and (b) some evidences suggests that dyslexics often show oral language impairment, whatever might be the link between dyslexia and SLI (Bishop & Snowling, 2004; Catts et al.,

2005; McArthur, Hogben, Edwards, Heath, & Mengler, 2000); therefore asking for verbal answers might lead to poor performances from dyslexics participants independently of their auditory abilities. (2) We will evaluate the different links between different auditory processes and different language abilities. Indeed, more than evaluating the links between APD and phonological awareness and reading ability, vocabulary level of participants was also evaluated to test whether a deficit in low level auditory process could impact higher level of language abilities (i.e., Matthew effect; Stanovich, 1986b). In addition, we want to investigate whether correlations between verbal performances and central auditory processes are found in both population or particularly in dyslexic population. This battery also provides the advantage of being designed to be shortly administered and thus could easily be adapted to a diagnostic situation.

2. Method

2.1. Participants

Twenty-four dyslexic adults (12 females) were tested (18-48 years old; SD = 7.96). Three dyslexic participants were excluded from the study for showing a hearing loss greater than 20 dB in one of the tested frequency (125 Hz, 250 Hz, 500 Hz, 750 Hz, 1000 Hz, 1500 Hz, 2000 Hz, 3000 Hz, 4000 Hz, 6000 Hz, 8000 Hz). One participant was also excluded due to poor score at Raven's matrices (Raven, 1941: Score = 27; M = 50.15; SD = 6.36). The remaining 20 participants were matched to 20 normal readers for age, non-verbal IQ and handedness (see Table 1). All of them were French native speaker.

2.2. Auditory tests

Auditory tests are visually represented in Figure 38.

Lateralization test. Lateralization skills were assessed by presenting binaurally a pure tone (250 ms; 500 Hz; rise-fall ramps 5 ms) with an Interaural Level Difference (ILD) varying from -5 dB to -25 dB by 2 dB steps (i.e., 11 ILD). One third of the pure tones were lateralized on the right; one third on the left; one third was presented in a binaural condition (ILD = 0 dB). Sixty-six trials were presented randomly and participants had to press the pre-specified key corresponding to the ear where they heard the sound (i.e., left, right or both). Six training trials including feedback were included.

Auditory discrimination test. All three discrimination tests were built following an AX paradigm. A first pure tone A (75 ms; 520 Hz; rise-fall ramps 5 ms) was presented, followed by X after a 400 ms ISI. In half trials A and X were identical; the remaining 24 were divided in 6 different conditions specified below. Participants had to indicate whether A and X were identical or different by pressing the corresponding key. Frequency, duration and intensity discrimination were tested following this procedure.

Values of the different pure tones' dimension were chosen following Pakarinen et al. (2007). Six training trials with feedback were presented prior testing.

Frequency. X frequency could be 520 Hz (i.e., the same as A) or 527 Hz; 535 Hz; 546 Hz; 562 Hz; 583 Hz or 609 Hz.

Duration. X length varied from 75 ms (i.e., the same as A) to 67 ms; 59 ms; 51 ms; 43 ms; 35 ms or 27 ms.

Intensity. X could be presented at the same intensity as A or with an intensity difference of -2.5 dB; -5 dB; -7.5 dB; -10 dB; -12.5 dB; -15 dB.

Auditory pattern test. This test was made of two subtests, one of them testing the frequency and the other one the duration. They are adaptation of the Frequency Pattern Test and the Duration Pattern Test (Musiek, 1994; Musiek & Pinheiro, 1987). Three pure tones separated by a 300 ms ISI were presented to participants in each trial. Participants had to go through two versions of this test in which stimuli could vary in terms of frequency or duration. Participants had to reproduce the heard pattern on a musical keyboard, albeit participants did not have to reproduce the exact note, and only the relative frequency was taken into account. Response pattern were written by the experimenter. Each test consisted of 12 training trials (including feedback from the experimenter for the first 6). Twenty trials were then presented.

Frequency. Three pure tones (500 ms; rise-fall ramps 5 ms) were presented as stimulus. Each pure tone could be 'High' (i.e., 1222 Hz) or 'Low' (i.e., 880 Hz). Six different patterns were presented: HHL; HLL; HLH; LHH; LLH; LHL.

Duration. Three pure tones of constant frequency (1000 Hz; rise-fall ramps 5 ms) were presented. They could be whether 'Short' (i.e., 250 ms) or 'Long' (i.e., 500 ms). These 6 different patterns were presented to participants: SSL; SLS; SLL; LSS; LSL; LLS.

Auditory Identification. Participants hear a noise referring to a daily encountered object (e.g., a fire) 3 pictures are simultaneously displayed. Participants had to indicate by pressing the corresponding key which picture corresponds to the noise they heard. These 3 pictures included the target picture (e.g., a fire); an acoustically related distractor (e.g., a crumpled paper) and a semantically related distractor (e.g., a fireman truck). Twenty trials were presented to participants. Once they answered, the next trial started automatically.

Auditory Recognition. In this test, participants were presented, as in the Auditory Identification test, to a target noise referring to a daily object. Once they heard the target noise, participants had to press the space key to hear 3 other noises, one after the other. They then had to choose, amongst the 3 heard noises, which one was

produced by the same object that produced the target noise. After they answered, the next trial began automatically.

Stream Segregation. The stream segregation test was the same as the one used by Lallier et al. (2009). Tone sequences of alternating pure tones of 1000 Hz and 400 Hz were presented to participants for 5 s. Each pure tone lasted for 49 ms (rise-fall ramps 10 ms). Participants were asked to report whether they perceived one or two streams. A training phase was thus presented to ensure that participants understood the correspondence between the percept and the concept. The SOA was varied using a one-up, one-down adaptive procedure, in order to estimate the fission threshold (i.e., 50% chance of hearing one or two streams; Levitt, 1971). The initial SOA (300 ms) was chosen to be perceived as ‘one stream’. SOA then decreased by 40 ms steps. After the first reversal, the step was set at 20 ms, after the second reversal, it was set at 10 ms and at 5 ms after the third reversal. Thresholds were then computed by averaging the last 10 trials.

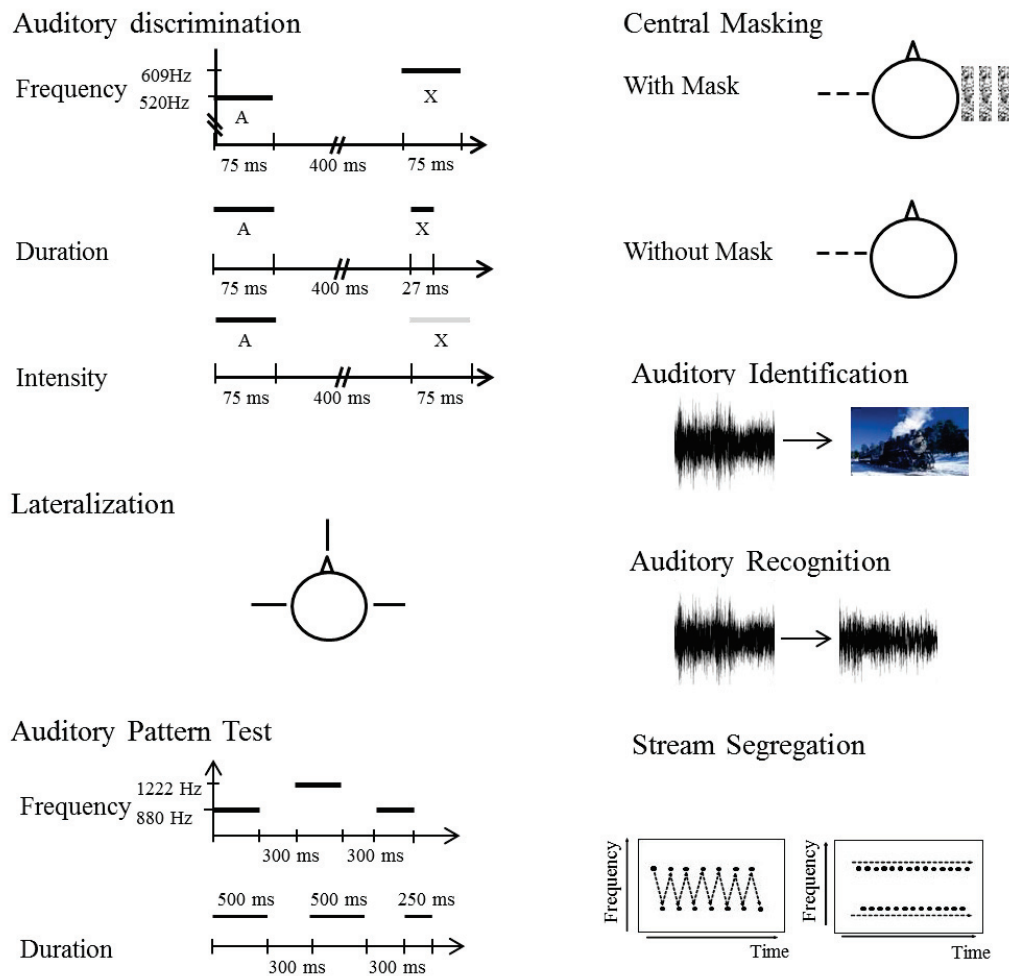


Figure 38

Visual representation of auditory tests.

2.3. Language evaluation

Participants' language skills were tested using the ECLA 16+ (Evaluation des Compétences de Lecture chez l'Adulte de 16 ans et plus; Cola-Asmussen, Lequette, Pouget, Rouyet, & Zorman, 2011).

Verbal Short-term and working memory

Participants' verbal short term memory and working memory were assessed using ascending and descending digit span.

Reading

Participants were submitted to several reading test. First, they underwent the Alouette test (Levafrais, 1967), for which was calculated the number of Word Correctly Read per Minute (WCRM; i.e., the speed of reading) and number of errors. They also performed a Regular Word, Irregular Word and Pseudo-Word reading test. Participants' time and number of error were recorded.

Phonological Awareness

Two main tests from the ECLA 16+ were used to evaluate participant's phonological awareness: phoneme suppression and spoonerism. In both test, time and number of errors were recorded. The time of two participants (one dyslexic and one control) are not available for spoonerism test due to technical issue.

Results and statistics are summarized in Table 1.

2.4. Procedure

Participants were tested in 2 sessions lasting approximately 1 hour and a half each. Each session alternated language and auditory tests to avoid monotony. There was no scheduled break but participants were told at the beginning of each session they can rest whenever they needed it. They sat in front of a computer when performing the auditory tests and computer was discarded when performing verbal tests. They were compensated for their participation.

	Dyslexics			Controls			Statistics		
	n	M	SD	n	M	SD	dl	t	p
Age	20	25.6	7.96	20	26.00	8.29	38	-0.24	n.s.
Handedness	20	83.5	21.09	20	83.5	17.55	38	-0.16	n.s.
Non Verbal IQ (Raven's Matrice) Score/60	20	50.15	6.36	20	50.15	6.1	38	0.00	n.s.
EVIP	20	158.45	10.05	20	161.90	9.46	38	-1.11	n.s.
Verbal Short-term Memory									
Ascending Span	20	6.25	0.91	20	6.9	1.11	38	-2.01	n.s.
Descending Span	20	4.65	1.03	20	5.2	1.32	38	-1.46	.05
Reading									
Alouette - CRWM	20	123.4	38.09	20	159.55	27.21	38	-3.45	<.001
Alouette - Number of errors	20	8.85	4.6	20	5.45	4.01	38	2.48	<.02
Regular Words Reading - Score (/20)	20	18.95	0.82	20	19.2	2.44	38	-0.43	n.s.
Regular Words Reading - Time (s)	20	20.7	9.85	20	12.05	3.51	38	3.69	<.001
Irregular Words Reading - Score (/20)	20	17.00	1.65	20	18.5	2.46	38	-2.26	<.03
Irregular Words Reading - Time (s)	20	19.2	9.35	20	11.8	3.03	38	3.36	<.01
Pseudo-word Reading - Score (/20)	20	16.4	3.97	20	18.25	2.91	38	-1.67	.10
Pseudo-word Reading - Time (s)	20	33.3	14.44	20	19.7	6.17	38	3.87	<.001
Phonological Awareness									
Phoneme Suppression - Score (/10)	20	7.5	2.25	20	8.9	1.29	38	-2.32	<.02
Phoneme Suppression - Time (s)	20	48.2	20.67	20	30.05	9.03	38	3.6	<.001
Spoonerism - Score (/20)	20	13.35	6.17	20	18.55	1.95	38	-3.59	<.001
Spoonerism - Time (s)	19	157.89	92.94	19	76.35	25.7	37	3.77	<.001

Table 7

Mean scores and Standard Deviations depending on Groups and Statistics.

3. Tests Results

Results for each participant at lateralization test were computed to find the most fitting psychometric function using Matlab software. This function allowed us to calculate the Interaural Level Difference needed for each participant to lateralize correctly the pure tone at 75%. The same was done for the discrimination tests, psychometric function were calculated to fit the results of each participants. Allowing thus to calculate the need of difference between A and X for participant to discriminate them at 75%. One way ANOVAs including Group (Dyslexics vs Controls) as between subject factor were run on data unless specified.

3.1. Lateralization

Four participants (1 control and 3 dyslexics participants) were excluded from the analysis as their needed ILD to lateralize a pure tone correctly at 75% was not included between tested ILD (-5 dB to -25 dB). No significant difference between the two groups was observed on ILD to lateralize a pure tone (Dyslexics: $M = 8.12$; $SD = 2.03$; Controls: $M = 8.50$; $SD = 3.33$; $F < 1$; cf Figure 39).

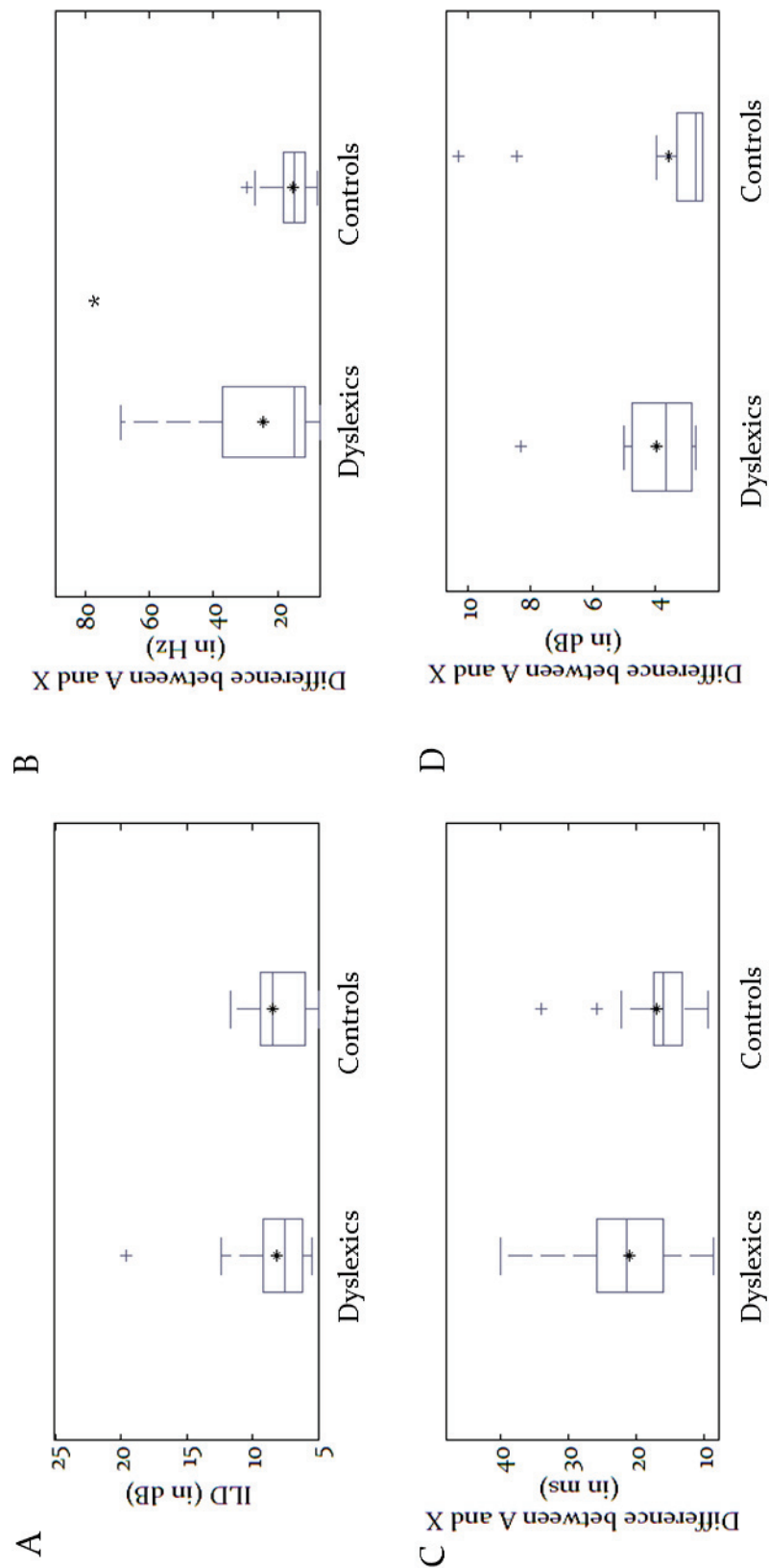


Figure 39

Boxplots and Means for Dyslexic and Control populations on A. Lateralization test. B. Frequency Discrimination test. C. Duration Discrimination test. D. Intensity discrimination test. + represents outliers, and * represents significance ($p < .05$).

3.2. Discrimination tests

3.2.1. Frequency discrimination

One control participant was excluded from this analysis because his/her need of difference between A and X (in Hz) was not in the range of tested values. Dyslexics participants significantly needed a bigger difference between A and X to discriminate them ($M = 24.34$ Hz; $SD = 17.48$) than did controls ($M = 15.41$ Hz; $SD = 6.34$; $F(1,37) = 4.45$, $p < .05$). Dyslexic participants are less efficient than controls in frequency discrimination (cf. Figure 39).

3.2.2. Duration discrimination

Two dyslexic participants whose need of difference between A and X to discriminate them was not in the tested range were not analyzed. On average, Dyslexics participants ($M = 21$ ms; $SD = 8$) needed 4 ms more than did controls ($M = 17$ ms; $SD = 6$) to discriminate A and X. This difference was marginally significant ($F(36,1) = 3.68$, $p = .06$; cf. Figure 39).

3.2.3. Intensity discrimination

Data of two control participants and 4 dyslexics were not included in the analysis because their need of difference between A and X to discriminate them was not part of the tested values. The two Groups need of intensity difference was not significantly different ($M_{\text{Dyslexics}} = 3.96$ dB; $SD_{\text{Dyslexics}} = 1.46$; $M_{\text{Controls}} = 3.59$ dB; $SD_{\text{Controls}} = 2.17$; cf. Figure 39).

3.3. Auditory Identification and auditory recognition

3.3.1. Auditory Identification

An ANOVA was performed on participants' number of error. It took the factor Group (Control vs Dyslexic) as a between subjects factor and the Type of Error (Semantic vs Acoustic) as a within subject factor. The two Groups performed exactly the same ($M = 2$; $SD = 1.07$), the factor Group was thus, not significant ($p = 1$). The Type of Error factor was significant ($F(1,38) = 12.75$, $p < .001$), revealing that participants committed more Acoustic Errors ($M = 1.25$; $SD = 0.93$) than Semantic Errors ($M = 0.75$; $SD = 0.74$). The interaction between the two factors was not significant as none of the post-hoc analyses (HSD Tukey) therefore suggesting that Dyslexics committed the same amount of Acoustic Errors and Semantic Errors than Controls ($p = .88$ and $p = .76$, respectively; cf. Figure 40).

3.3.2. Auditory Recognition

Due to technical difficulties, the results of a Dyslexic participant could not be recorded. Analyses were thus performed on 39 participants. An ANOVA including the

Group (Control vs Group) as a between subjects factor and including the Type of Error (Acoustic vs Semantic) as a within subjects factor was conducted. The Group factor was not significant (Dyslexics: $M = 2.69$; $SD = 1.85$; Controls: $M = 3.45$; $SD = 1.50$; $F(1,37) = 2.01$, $p > .1$). However, as for identification, the factor Type of Error is significant, participants committed more Acoustic Errors ($M = 1.99$; $SD = 1.26$) than Semantic ($M = 0.91$; $SD = 1.00$). Again the interaction between Group and Type of Error was not significant, neither the post-hoc tests (HSD Tukey) revealing no significant differences between Dyslexics and Controls either in the number of Acoustic Error ($p = 0.99$) nor in the number of Semantic Error ($p = .70$; cf. Figure 40).

3.4. Auditory Pattern Test (APT)

Data from the APT were analyzed with an ANOVA including the between participants factor Group (Control vs Dyslexics) and the within participants factor Dimension (Frequency vs Duration). In both Dimensions, Dyslexic participants ($M = 16.42$; $ET = 3.64$) were significantly less accurate than Controls ($M = 19.05$; $SD = 1.49$; $F(1,38) = 7.12$, $p < .05$). The Dimension factor also had a significant impact on participants performances, both Groups were more accurate when stimuli differed in term of Frequency ($M = 18.40$; $SD = 2.22$) than in Duration ($M = 17.05$; $SD = 2.91$). The interaction between the two factors was not significant ($F < 1$). As the aim of this study is to compare Dyslexic and Control participants performances, we performed a post-hoc test (HSD Tukey) which confirmed that Dyslexics ($M_{\text{Frequency}} = 17.25$; $SD_{\text{Frequency}} = 3.21$; $M_{\text{Duration}} = 15.60$; $SD_{\text{Duration}} = 4.07$) were less accurate than Controls ($M_{\text{Frequency}} = 19.55$; $SD_{\text{Frequency}} = 1.23$; $M_{\text{Duration}} = 18.55$; $SD_{\text{Duration}} = 1.76$) in both Dimension ($p_{\text{Frequency}} = .05$; $p_{\text{Duration}} < .01$; cf. Figure 40).

3.5. Stream Segregation

Participants who responded more than 5 times consecutively the same thing (one or two stream) on the last ten trials were discarded from the analyses (1 Control and 5 Dyslexics). Three other Control participants were excluded due to material deficiency. A one way ANOVA including Group factor (Control vs Dyslexic) was performed on the remaining 31 participants (16 Controls and 15 Dyslexics). No significant difference is observed (Dyslexics: $M = 172$ ms; $SD = 84$ to perceive 2 streams and Controls: $M = 135$ ms; $SD = 69$; $F(1,30) = 2.67$, $p = .11$ cf. Figure 40).

3.6. Prevalence of APD

Following the diagnostic criteria proposed by ASHA (2005), we calculated for each test the mean of control participants $- 2$ SD. Nine dyslexic participants reached the criterion for APD (scores inferior at control mean $- 2SD$ for at least 2 tests), whereas only two control participants reached this criterion.

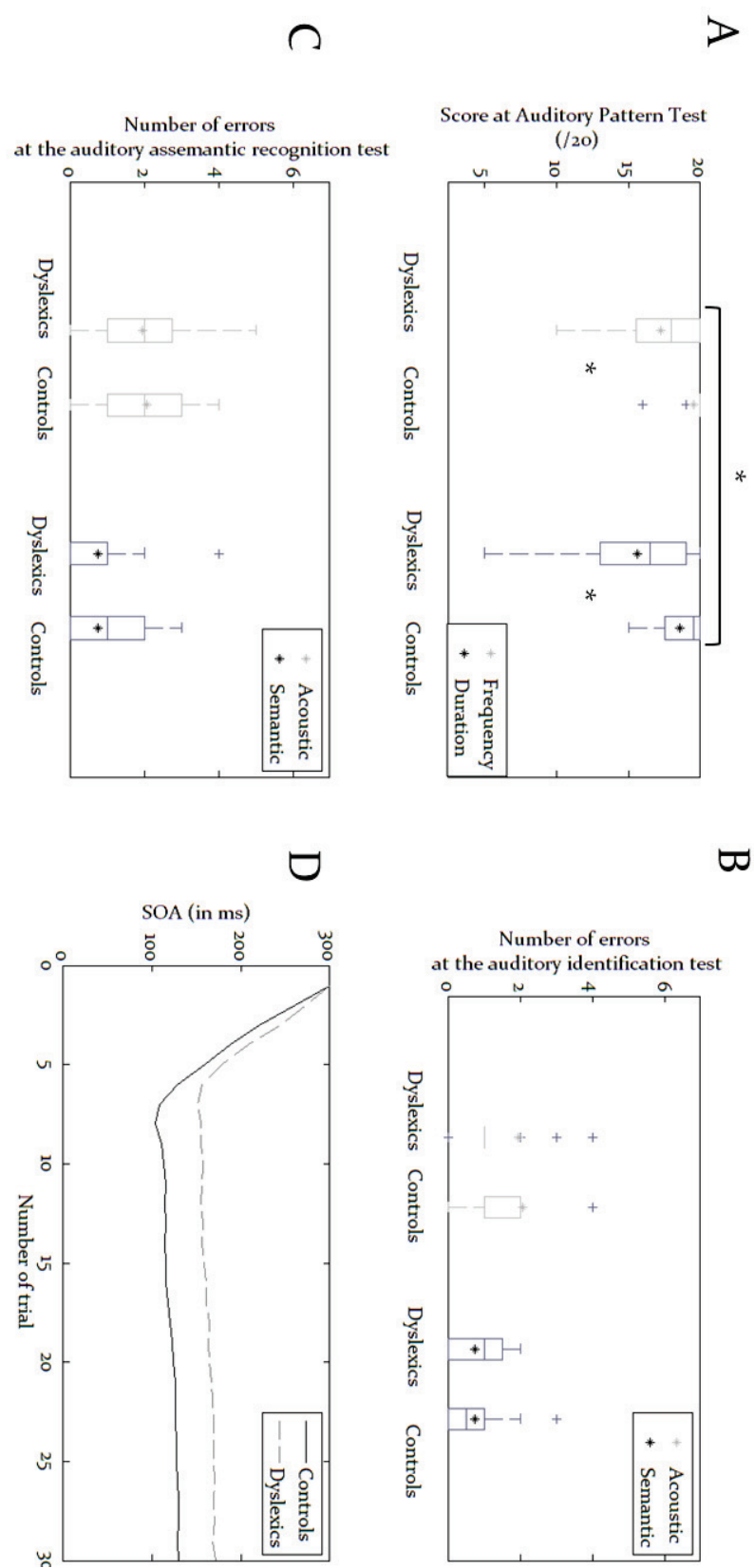


Figure 40

Boxplots and Means for Dyslexic and Control populations on A. Auditory Pattern Test. B. Auditory Identification test. C. Auditory Assemantic recognition test. D. Stream Segregation test. + indicates outliers and * indicates significant differences (p < .05).

4. Correlations

The second aim of this paper was to investigate the different links between language skills and central auditory processes. Pearson coefficients were thus calculated between auditory and language tests for the entire population. When a significant correlation appeared between the results of an auditory and a language test, it was recalculated for each group separately. This aimed at disentangle correlation and simple group effect.

4.1. Frequency Discrimination

Total Population. Results at the frequency discrimination task correlated significantly with the number of Word Correctly Read per minute (WCRM) at the Alouette ($r = -0.34$; $p < .05$). Participants who needed a bigger difference between A and X (in Hz) to differentiate them were also longer at this reading test. They were also longer in reading regular words ($r = 0.52$; $p < .05$), irregular words ($r = 0.4$; $p < .05$) and pseudo-words ($r = 0.59$; $p < .05$).

The ability to discriminate frequency was also correlated to phonological awareness, both for accuracy (phoneme suppression: $r = -0.35$; $p < .05$; spoonerism: $r = -0.45$; $p < .05$) and speed (phoneme suppression: $r = 0.38$; $p < .05$; spoonerism: $r = -0.45$; $p < .05$). These correlations suggest that participants who needed less difference between A and X (in Hz) to differentiate them at 75% have a better phonological awareness.

Control Group. No correlation was found between the results at the frequency discrimination test and any language skill in Control group.

Dyslexic Group. The difference between A and X (in Hz) necessary for their discrimination was correlated to the speed at the pseudo-word reading test ($r = 0.59$; $p < .05$). Dyslexic participants who performed better at discriminating frequency were also faster at reading pseudo-words. They were also faster at the spoonerism test ($r = 0.6$; $p < .05$; cf. Figure 41).

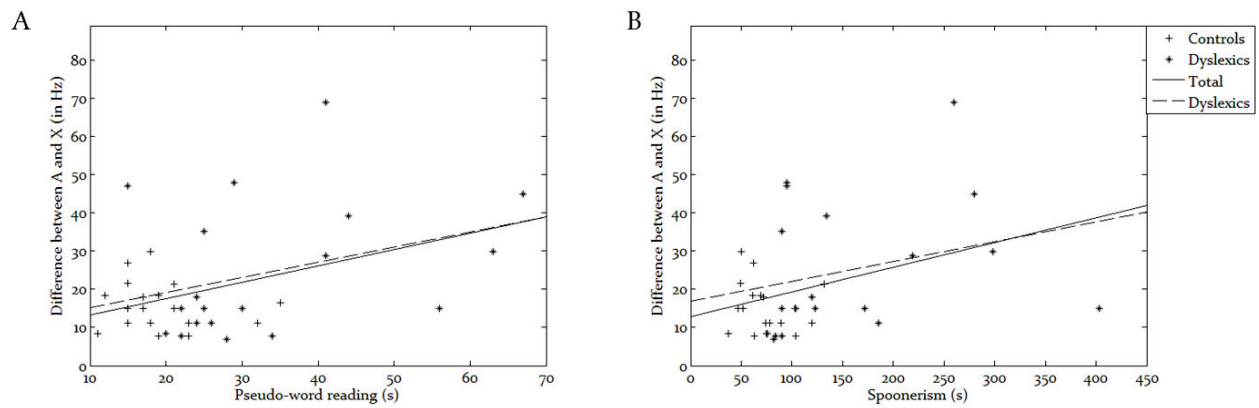


Figure 41

Correlations for Dyslexic and Total populations between threshold at frequency discrimination test and A. Pseudo word reading (in s) and B. Spoonerism (in s).

4.2. Duration Discrimination

Total Population. Participants who needed a smaller difference between A and X (in ms) to differentiate them were also able to read faster at the Alouette test ($r = -0.4$; $p < .05$). They were also faster in regular ($r = 0.50$; $p < .05$) and irregular ($r = 0.38$; $p < .05$) word and pseudo-word reading ($r = 0.54$; $p < .05$).

Results at the duration discrimination also correlated with speed at the phoneme suppression test ($r = 0.46$; $p < .05$) and spoonerism ($r = 0.56$; $p < .05$). They were also linked to the scores at spoonerism ($r = 0.56$; $p < .05$). These correlations imply that participants who were better at duration discrimination were faster at the phoneme suppression test and were both more accurate and faster at the spoonerism test.

Control Group. There were no correlation between results at the Duration discrimination test and language tests.

Dyslexic Group. Dyslexics who were more accurate at the Duration discrimination test were faster in pseudo-word reading ($r = 0.67$; $p < .05$; cf. Figure 42).

Results at the Duration discrimination were also correlated to the speed at spoonerism ($r = 0.63$; $p < .05$). Dyslexic participants needing a smaller difference between A and X to discriminate them were also faster at spoonerism (cf. Figure 42).

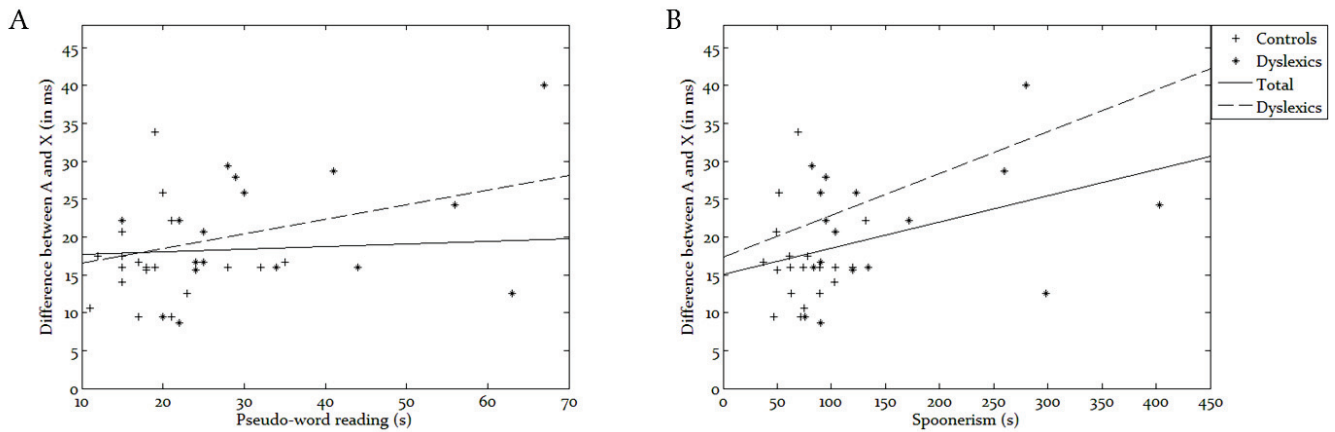


Figure 42

Correlations between thresholds at the duration discrimination test A. Pseudo-word reading test (in s).
B. Spoonerism test (in s).

4.3. Auditory pattern test. Frequency

Total Population. Performances at APT correlated significantly with all reading tests. These correlations suggested that the better participants were at the APT test the better they were at reading both in term of accuracy and speed. Indeed, at the Alouette test, participants who had better score at the APT were also the faster ($r = 0.34$; $p < .05$) and were more accurate ($r = -0.69$; $p < .05$). The same observation was made on regular ($r_{\text{score}} = 0.53$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.47$; $p_{\text{time}} < .05$), irregular ($r_{\text{score}} = 0.49$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.43$; $p_{\text{time}} < .05$) and pseudo-word reading tests ($r_{\text{score}} = 0.71$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.51$; $p_{\text{time}} < .05$).

Phonological awareness was also correlated with the ability to perform a auditory pattern test on frequency. Participants who were more accurate at APT were also better at phoneme suppression both for accuracy ($r = 0.43$; $p < .05$) and speed ($r = -0.36$; $p < .05$). Similar correlation appeared to be significant for the spoonerism test, speaking of accuracy ($r = 0.45$; $p < .05$) and speed ($r = -0.61$; $p < .05$).

Control Group. Scores at the APT were significantly correlated with some of the reading test. Participants who were better at reproducing the sequence of pure tone committed less error at the Alouette test ($r = -0.59$; $p < .05$). They also were more accurate at reading regular words ($r = 0.55$; $p < .05$) and pseudo-word ($r = 0.79$; $p < .05$).

Dyslexic Group. Participants scores at the APT were also correlated to the number of errors at the Alouette test ($r = -0.7$; $p < .05$), and with pseudo-words reading score ($r = 0.65$; $p < .05$). These correlations suggest that participants who were more accurate at the APT were also better readers. They also were faster at the spoonerism test ($r = -0.57$; $p < .05$; cf. Figure 43).

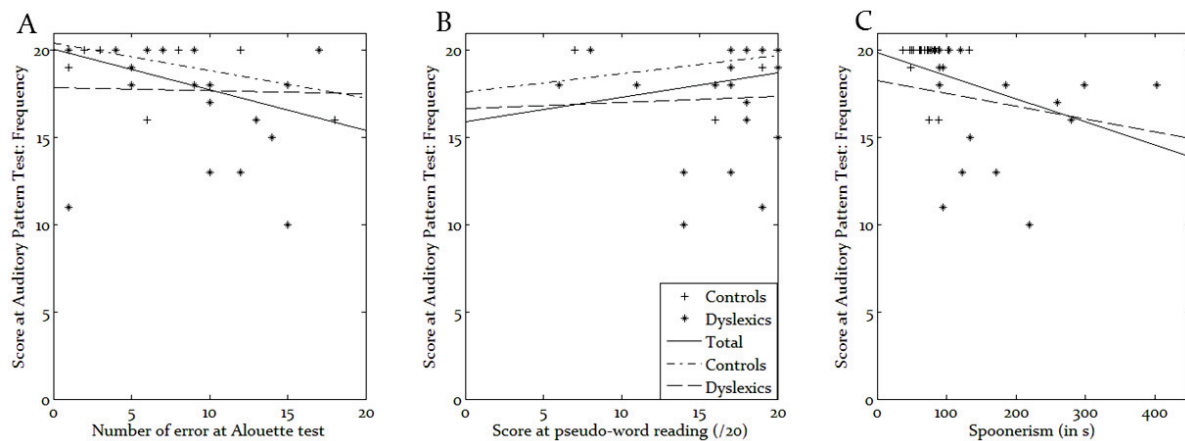


Figure 43

Correlations between Score at the Auditory Pattern Test on frequency and A. Number of error at Alouette test. B. Score at pseudo-word reading. C. Spoonerisms (in s).

4.4. Auditory pattern test. Duration

Total Population. Results at the APT were correlated with all reading test except for the WCRM at the Alouette test. Participants performing better at the APT were more accurate at reading the Alouette ($r = -0.58$; $p < .05$). They also were better and faster at reading regular words ($r_{\text{score}} = 0.5$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.53$; $p_{\text{time}} < .05$), irregular words ($r_{\text{score}} = 0.38$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.52$; $p_{\text{time}} < .05$) and pseudo-words ($r_{\text{score}} = 0.35$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.62$; $p_{\text{time}} < .05$).

Phonological awareness also seems to be related with the ability to perform the APT. Participants who were better at reproducing the sequence of stimuli were also better at the phonological awareness tests. They were more accurate and faster both at the phoneme suppression test ($r_{\text{score}} = 0.50$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.55$; $p_{\text{time}} < .05$) and at the spoonerism test ($r_{\text{score}} = 0.47$; $p_{\text{score}} < .05$; $r_{\text{time}} = -0.58$; $p_{\text{time}} < .05$).

Control Group. Results at the APT for control Group did not correlate with any language test.

Dyslexic Group. Participants who were better at the APT also committed less error at the Alouette test ($r = -0.62$; $p < .05$). They were also better at reading regular words ($r = 0.51$; $p < .05$) and faster at reading pseudo-words ($r = -0.6$; $p < .05$).

The ability to perform the APT was also linked to phonological awareness. Participants who had better score at the APT were more accurate ($r = 0.57$; $p < .05$) and faster ($r = -0.56$; $p < .05$) at the phoneme suppression test. Results at the APT also correlated with the speed to perform the spoonerism test ($r = -0.55$; $p < .05$; cf. Figure 44).

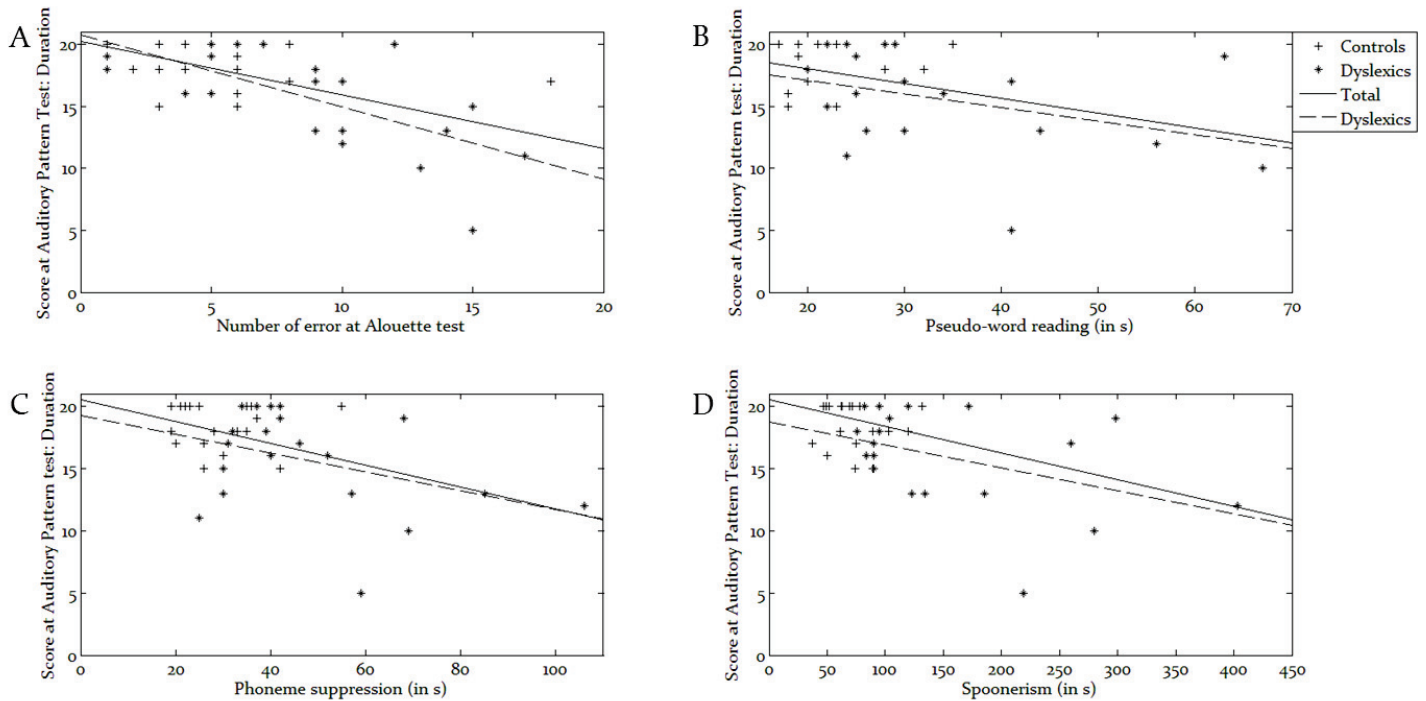


Figure 44

Correlations between Score at Auditory Pattern Test: Duration and A. Number of errors at Alouette test. B. Pseudo-word reading (in s). C. Phoneme suppression (in s). D. Spoonerism (in s).

5. Discussion

The aim of this study was to investigate Auditory Processing Disorders in dyslexics and their links with verbal skills. We used a previously developed battery in order to evaluate Central Auditory Processes as they are defined by ASHA, and to investigate whether performances at the central auditory processes tests were linked to verbal performances in one or both population. The second aim of the study was to try to evidence different profiles of dyslexics, according to the impaired processes.

Results do not evidence any deficit in lateralization abilities in dyslexics, no previous data about sound lateralization performances and dyslexia were found. However, our results seem to be congruent with previous findings of preserved lateralization abilities suggested by preserved spatial unmasking in speech in noise (Dole et al., 2012) or for pure tone (N. I. Hill et al., 1999). Geiger et al. (2008) evidenced that dyslexics perform less accurately than controls in speech in noise situation when noise is lateralized, however, no condition presented noise in the same direction as target speech, therefore, it's not clear if dyslexic are more impaired by lateralized noise than controls as it might be that they just are more impaired in speech in noise than controls.

Overall, it seems that dyslexics' ability to lateralize sounds using ILD is the same as for controls.

Our study evidenced that dyslexics are less performant than controls at discrimination test when the tested dimension is frequency. These results are congruent with the literature as it has been widely emphasized that dyslexics show impairments in frequency processing both in children (Banai & Ahissar, 2006) and adults (Ahissar et al., 2006; Ahissar et al., 2000; France et al., 2002). However, as mentioned in the introduction, links between frequency processing and language skills remain unclear, indeed, some studies claim that there is no evidence for a link between low level auditory deficit and phonological awareness and others claim that the deficit in frequency processing comes from a deficit to perform the paradigm more than an inability to discriminate frequency (Ahissar, 2007; Banai & Ahissar, 2010). Although the Anchor Theory is very interesting, our results suggest that it might be a part of the explanation of dyslexics' poor performances at frequency discrimination test but not the only one. Indeed, one of the interest of our study is that it tests discrimination of several sound dimensions (frequency, duration and intensity) using the same paradigm. Therefore, it allows testing the impact of the paradigm on the recorded performances. Indeed, if dyslexic participants are lower performers than controls in one or two of the tested dimension but not all, it therefore means that the ability to perform the task does not only depends on the paradigm. Interestingly, it is the case in our study; dyslexic participants are found to be impaired in frequency and tend to be in duration discrimination however they are not impaired in intensity discrimination. Therefore, our paradigm which is supposed to favor control participants compared to dyslexic might not underlay this effect. At least, it would mean that the anchoring deficit is specific to frequency and maybe duration and thus results from a low level auditory deficit. Very few studies reported results of duration discrimination, in fact only one study getting interested in duration discrimination differences between dyslexics and controls was found (Banai & Ahissar, 2006). It evidenced that 7th grade children with reading difficulties needed a greater difference (in ms) between two pure tones than controls to determine whether they were different or not. Our results suggested therefore that the difficulties to discriminate duration pertains during adulthood.

Results at both asemantic recognition and auditory identification do not differ between dyslexic and control participants. More, in both tasks, pattern of results are equivalent, that is, participants commit few mistakes and principally they commit acoustic errors and one or no semantic errors. Concerning the auditory identification test, the fact that dyslexic participants are not found to be impaired is coherent with the absence of difference of vocabulary level in both groups. Indeed, this task needs participants to access the amodal concept when hearing the sound to be able to pair it to the corresponding picture. Once again, it suggests that dyslexia has not impacted

higher level of language representation than phonological awareness in our dyslexic population; however it has to be taken into account that our dyslexic population is very qualified and well compensated and may not be representative on that specific issue of a more general population of dyslexics. Our dyslexic participants were not impaired either in the auditory identification test; they performed like controls with the same pattern of errors suggesting that the same processes were used to perform the task.

Dyslexic participants were also found to be impaired in both subtests of the auditory pattern test. These results are congruent with the literature who usually find deficit in dyslexics in FPT and DPT (Billiet & Bellis, 2011; King, Lombardino, Crandell, & Leonard, 2003). It is therefore interesting to note that our simplified test allow us to obtain a similar pattern in less time. Interestingly, both populations follow the same pattern of results, that is, their performances are worst in the duration subtest than in the frequency subtest.

A second aim of the study was to investigate the correlations between the here tested abilities and the language processes. However, it must be noted that correlations might not always reflect a causal link; in order to rule out this possibility we only report correlations which were found to be significant both at the population level (controls + dyslexics) and at the dyslexic level. Interestingly, only performances which were previously found to be impaired in dyslexics were also found to be linked to language processing. It must also be noted that no here tested central auditory ability was found to be correlated with vocabulary level, suggesting low level auditory processes has no impact on higher level language representation (once again this claim must be moderated, as our dyslexic population might not be very representative of global dyslexic population on this domain). It seems, given our results that spectral and temporal processing are both impaired in dyslexics and that these impairments are linked with reading and phonological abilities, even in well compensated dyslexics.

No correlation was found between central auditory processes and verbal abilities for control group, except between reading abilities and FPT. It seems that with language acquisition, phonological representation become independent from sensory processing in normal population but not in dyslexic population. It therefore appears that during development, language representation do not tear off from low level processing in dyslexic population. This is particularly interesting as our dyslexic sample is very well compensated and received reeducation. Their phonological awareness and reading abilities have thus certainly evolve since childhood but are still dependent of low level processes.

The third aim of this paper was to try to evidence one or more profile of auditory low level deficit which could eventually be associated to specific language processing. Our results showed that dyslexic participants were impaired in some but not all the tested

processes. They particularly showed a deficit in tasks involving frequency processing (as frequency discrimination), duration processing (as duration discrimination) and auditory pattern test both in term of frequency and duration. These 4 tests showed correlation with reading and phonological skills not only at the population level but also within the dyslexic group. This is very interesting as it suggests that the found correlation does not only rely on group effect. Particularly, it seems that in reading, the Alouette test and the pseudo-word reading test are more sensible than the other reading tests. It can easily be explained by the fact that both of these tests might not be targeted by active reeducation. (1) the Alouette test is the main test to diagnose dyslexia in French speaking countries, all dyslexic participants had therefore already encounter this text (2) our dyslexic population is mainly made of university students, they have therefore benefit from efficient reeducation and can be considered as very well compensated dyslexics. They might have been specifically trained to be able to read quite efficiently regular and irregular words leading to less variable performances than for pseudo-words (cf. Table 1). The same observation can be made regarding the performances at the phonological awareness tests, inter individuality variance is more important in speed than in scores because our dyslexic participants are very well compensated; they were therefore trained to be able to perform that kind of tasks correctly.

Finally, looking at the number of abilities impaired and congruent with literature, 9 out of our 20 dyslexics were impaired in 2 or more of the tested abilities although only 2 of the control participants reached this threshold. Impaired was defined, as suggested by ASHA, as score < 2 SD of control group. Particularly 9 of the dyslexic participants were impaired at the Auditory Pattern Test, in terms of duration and 8 of them were impaired at the same task in the frequency subtest. This is congruent with literature evidencing that dyslexics are particularly impaired in temporal and spectral processing, however, no specific profile was highlighted.

6. Conclusion

The aim of this study was to investigate the integrity of central auditory processing in a dyslexic population, as they are defined by ASHA, using entirely non-verbal stimuli. Results evidenced deficit in 40 % of our dyslexic population particularly on tasks involving temporal and spectral processes. Particularly, 15 out of 20 dyslexics show a deficit in at least one test vs 6 out of 20 control participants; however no obvious links were found between central auditory processing disorders and language abilities. Nevertheless, results suggest that dyslexics' language representation stay dependent of sensory processes, despite reeducation whereas those processes seem to disentangle in control population.

7. Résumé du Chapitre X

Cette dernière étude a permis de mettre en évidence des déficits aux tests impliquant un traitement spectral et temporel chez les adultes dyslexiques. De plus, les performances à ces mêmes tests sont corrélées aux performances en conscience phonologique et lecture, aussi bien pour l'ensemble de la population testée (groupe contrôle + groupe dyslexique) qu'au niveau du groupe dyslexique uniquement. Ceci suggère donc que les corrélations observées ne reflètent pas uniquement un effet de groupe mais bien l'existence d'un lien entre les performances auditives et langagières. Il est intéressant de noter que le niveau de vocabulaire n'est pas perturbé par la présence TTA. Ceci suggérerait que chez les adultes le manque de qualité du signal de parole ne perturbe pas les représentations sémantiques, alors que, comme suggéré par les résultats obtenus dans l'Axe 1, la dégradation transitoire du signal le traitement sémantique. Toutefois, la population dyslexique que nous avons testée est essentiellement constituée d'étudiants de niveau universitaire, ils ont donc bien compensé leur dyslexie et il est possible que sur ce point ils ne soient pas représentatifs de la population dyslexique.

Discussion

Chapitre XI

Discussion de l'Axe 1

L'objectif de ce travail était de s'intéresser aux liens entre la qualité de la perception du signal de parole et les traitements cognitifs effectués sur celui-ci. Plus spécifiquement nous nous sommes intéressés à l'impact de la dégradation transitoire ou permanente du signal de parole sur le traitement du langage et sur ses représentations. L'Axe 1 visait à évaluer l'impact de la dégradation transitoire du signal de parole sur l'accès au sens. Dans un premier temps, la discussion de l'Axe 1 présente un résumé des Etudes 1 et 2 puis reprend les principaux questionnements de la littérature exposés au cours de la partie Théorique à la lumière de nos résultats.

1. Résumé des résultats

1.1. Etude 1

Dans l'Axe 1, le travail expérimental a porté sur les effets de la dégradation transitoire du signal de parole provenant de la présence de bruit (i.e., la situation de *cocktail party* ; Cherry, 1953) sur le traitement du sens des mots. L'Etude 1 consistait en l'adaptation du paradigme d'amorçage sémantique à la situation de *cocktail party*. Les participants devaient effectuer une tâche de décision lexicale sur un item cible inséré au sein d'un fond sonore composé de bruit parolier. L'objectif était d'évaluer si le contenu sémantique de ce dernier (i.e., jouant le rôle d'amorce) pouvait influencer sur le traitement du mot cible (sémantiquement lié ou non à l'amorce). Dans la première expérience, le fond sonore était composé d'une ou deux voix Sémantiquement Consistantes (i.e., SC ; prononçant des mots sémantiquement liés entre eux et sémantiquement liés ou non au mot cible). Dans les 2^e et 3^e expériences, respectivement 1 et 2 voix Sémantiquement Inconsistantes (i.e., SI ; prononçant des mots sémantiquement indépendants les uns des autres et jamais liés au mot cible) ont été ajoutées à chaque fond sonore afin de masquer les voix SC.

Les résultats ont mis en évidence que le traitement sémantique des voix SC n'avait d'influence sur le traitement du mot cible que lorsqu'elles étaient suffisamment saillantes (i.e., le nombre de voix SC était supérieur au nombre de voix SI ; jusqu'à 4

voix au total). Enfin les résultats suggèrent que lorsque les fonds sonores étaient composés de 3 et 4 voix, les participants étaient capables d'entendre les mots les composants et de les reconnaître dans un post-test, bien qu'aucun effet d'amorçage sémantique ne soit apparu. Ces données suggèrent alors que le traitement sémantique ne s'effectue pas de manière automatique et qu'il est dépendant de la saillance de l'amorce ainsi que de son intelligibilité. Ces observations sont cohérentes avec l'*Effortfulness Hypothesis* (Gao et al., 2011; Rabbitt, 1968; Wingfield et al., 2005). Elle postule que lorsque le signal est perçu de façon optimale, les traitements bas niveau se font de façon automatique, alors que les traitements haut niveau sont dépendants de la disponibilité des ressources cognitives. Lorsque le signal est dégradé, les ressources cognitives sont allouées aux traitements bas niveau pour compenser la dégradation du signal. De ce fait, moins de ressources sont allouées aux traitements de haut niveau, ceux-ci sont donc moins efficaces. Dans l'Etude 1, l'effet d'amorçage disparaît quand le nombre de voix SI par rapport au nombre de voix SC augmente, bien que les participants soient toujours capables de percevoir et reconnaître les mots du fond sonore. Nous postulons, en suivant l'*Effortfulness Hypothesis* que la dégradation du signal de parole a engendré une réallocation des ressources cognitives sur les traitements bas niveau au détriment du traitement sémantique. Celui-ci n'était plus suffisant pour engendrer l'apparition d'un effet d'amorçage.

Afin de poursuivre nos investigations sur la modulation du traitement sémantique par la dégradation du signal, l'Etude 2 a été mise en place. En effet, les résultats de l'Etude 1 suggèrent que la profondeur du traitement sémantique est soumise à la disponibilité des ressources cognitives, qui peut être modulée par la présence de bruit. Toutefois le paradigme utilisé ne permet pas de déterminer si le traitement sémantique des mots du cocktail est superficiel ou totalement absent. L'Etude 2 a ainsi permis de nous intéresser plus spécifiquement à la modulation du traitement sémantique grâce à l'utilisation des mots ambigus. Cette étude visait également à mettre en évidence les interactions entre les processus *top-down* et *bottom-up*, c'est-à-dire à identifier comment le contexte et sa dégradation influencent le traitement du mot cible. Cette étude s'est déroulée en collaboration avec le Professeur Stemmer de l'Université à Montréal.

1.2. Etude 2

Les mots ambigus étaient présentés au sein d'une phrase, composée de deux parties. La première, contenant le mot ambigu, faisait obligatoirement référence à son sens dominant (e.g., *Au tribunal, j'ai vu des avocats*). Le mot cible était le dernier mot de la phrase. En condition congruente, il faisait référence au sens dominant du mot ambigu (i.e., *Le légiste s'occupe des morts, il cherche aussi la cause du DECES*). En condition ambiguë, il faisait référence au sens subordonné du mot ambigu (i.e., *Tu lui as fait une bonne blague, espérons qu'il ne perdra plus son TABAC*). Enfin, en condition

incongruente, il ne faisait référence à aucun des sens du mot ambigu (i.e., *Le paysan a un grain, il va pouvoir faire une USINE*). De plus, ces phrases pouvaient être présentées au sein d'un bruit parolier composé de 4 voix (i.e., condition Bruit Parolier), au sein d'un bruit stationnaire partageant les mêmes caractéristiques spectrales que le bruit parolier (i.e., condition Bruit Stationnaire) ou dans le silence (i.e., condition Silence). Les variations d'amplitude de la N400 en réponse au mot cible ont été analysées.

Les résultats ont mis en évidence une N400 de plus grande amplitude en réponse au mot cible lorsque celui-ci était incongruent que lorsqu'il était congruent (i.e., effet de congruence), et ce, uniquement lorsque la phrase était présentée dans le silence ou dans du bruit stationnaire. Dans ces conditions la réponse N400 avait également une amplitude plus importante lorsque le mot cible faisait référence au sens subordonné du mot ambigu que lorsqu'il faisait référence à son sens dominant (i.e., effet d'ambiguïté). En revanche, l'effet d'ambiguïté était significativement plus petit que l'effet de congruence uniquement lorsque la phrase était présentée dans du bruit stationnaire.

Nous avons émis l'hypothèse selon laquelle l'effet d'ambiguïté serait significativement moins important que l'effet de congruence aussi bien en condition Silence qu'en condition Bruit Stationnaire. En effet, nous pensions que le peu de ressources allouées aux traitements superficiels dans ces conditions laisserait beaucoup de ressources disponibles pour les traitements de plus haut niveau. De ce fait, nous pensions que l'activation sémantique du mot ambigu serait suffisante pour activer ses différents sens. Contrairement à l'hypothèse posée, il semble que les deux sens du mot ambigu ne soient pas activés lorsque celui-ci est présenté dans un contexte favorisant son sens dominant en condition Silence. Dans cette condition, du fait de la bonne qualité du signal de parole, les ressources pouvaient être allouées au traitement sémantique. Le contexte était vraisemblablement traité profondément et de ce fait, seul le sens pertinent dans le contexte (ici le sens dominant) a été activé et a mené à une prédiction forte sur le mot cible. Finalement, l'effet d'ambiguïté était significativement plus faible que l'effet de congruence lorsque la phrase était présentée en bruit stationnaire. En effet, la présentation du mot cible faisant référence au sens subordonné du mot ambigu engendre une N400 de plus faible amplitude que la présentation d'un mot incongruent. Cette différence d'effet suggère alors que le sens subordonné du mot ambigu avait été activé. Dans cette condition, le contexte était vraisemblablement traité de façon moins efficace du fait de la légère dégradation du signal. Moins de prédictions étaient effectuées sur le sens du mot ambigu, et donc ses différents sens pouvaient être activés. Comme suggéré par l'étude de Colbert-Getz et Cook (2013), c'est la force du contexte qui a déterminé l'activation des différents sens du mot ambigu. La force du contexte ici, est modulée par la présence de bruit alors que dans l'étude citée, elle est modulée par la longueur du contexte (i.e., 1 phrase ou 4 phrases).

1.3. Effet du masquage informationnel ou énergétique ?

Les résultats de ces deux premières études ont permis de montrer que la dégradation du signal de parole a un effet sur le traitement sémantique. Nous n'avons pas pu mettre en évidence de traitement sémantique lorsque le signal cible était intégré à un bruit parolier composé de 4 interlocuteurs, et ce alors que dans les deux études les participants sont capables d'entendre les mots présentés. Il semble donc qu'à partir d'un certain seuil, le traitement sémantique ne soit plus mesurable, ni comportementalement ni électrophysiologiquement. Toutefois, la cause de cette diminution importante de la profondeur du traitement sémantique n'est pas connue. En effet, dans les deux études, les bruits de fonds sont des bruits paroliers à 4 voix, plusieurs facteurs peuvent donc expliquer le phénomène observé : le masquage énergétique, les interférences dans les traitements au niveau central ou une combinaison des deux. Si seule la dégradation physique du signal de parole est responsable de l'absence de traitement sémantique mesurable, alors il ne devrait pas être possible non plus de mesurer un traitement sémantique lorsque le taux d'intelligibilité est le même mais que le masquage est uniquement énergétique (i.e., utilisation de bruit stationnaire ou fluctuant, à un SNR inférieur à 0). A l'inverse, utiliser une double tâche permettrait d'évaluer si la capacité à effectuer un traitement sémantique est modulée du fait d'une appropriation des ressources cognitives par les traitements superficiels ou d'une surcharge directe de ces ressources par une trop forte demande. Les participants pourraient par exemple écouter des phrases dans le silence (i.e., pas de demande de bas niveau) tout en effectuant une tâche visuelle comme dans l'étude de Mattys et al. (2009 ; i.e., forte demande en ressources cognitives). La comparaison de la réponse N400 suite à la présentation d'un mot cible congruent ou non congruent selon si la phrase a été présentée dans du bruit stationnaire (i.e., masquage énergétique) ou dans le silence accompagné d'une double tâche (i.e., masquage informationnel) permettrait d'évaluer si l'un ou l'autre ou les deux types de masquage sont responsables de la modulation du traitement sémantique.

Bien que l'*Effortfulness Hypothesis* explique nos résultats, de nombreuses études, ont suggéré que le traitement sémantique est automatique. Nous allons donc reprendre les principaux résultats de la littérature sur le sujet afin de déterminer s'ils sont compatibles avec l'*Effortfulness Hypothesis*.

2. Traitement sémantique et *Effortfulness Hypothesis*

2.1. *Effortfulness Hypothesis* et amorçage masqué

Comme nous l'avons vu dans la partie théorique, le débat quant à l'automaticité du traitement sémantique reste d'actualité. Il est maintenant admis qu'il est possible en modalité visuelle de générer un traitement sémantique subliminal donnant lieu à un effet de congruence (Van den Bussche et al., 2009). La modalité auditive rend toutefois

difficile le masquage du signal de parole. En effet, puisque celui-ci dure dans le temps, son traitement se fait au fur et à mesure de sa présentation et ne peut donc pas être stoppé avant de devenir conscient. Plusieurs méthodes ont été proposées pour tester l'automatisme du traitement sémantique en modalité auditive (Daltrozzo et al., 2011; Eich, 1984; Kouider & Dupoux, 2005) mais aucune ne s'est révélée satisfaisante. Dans cette thèse, nous avons pu tester l'automatisme du traitement sémantique sans pour autant rendre le signal inaudible. En effet, dans deux études distinctes nous avons observé une disparition du traitement sémantique malgré une intelligibilité préservée. Nos résultats s'inscrivent dans l'*Effortfulness Hypothesis*, suggérant que le traitement sémantique n'est pas automatique mais soumis à la disponibilité des ressources cognitives. Cette interprétation semble au premier abord contredire la littérature sur l'amorçage masqué en modalité visuelle, qui est censé mettre en évidence l'automatisme du traitement sémantique. En réalité, il paraît possible de réinterpréter les résultats des études d'amorçage visuel masqué en termes d'allocation de ressources.

Comme nous l'avons vu dans la partie Théorique (cf. p42), les effets observés grâce au paradigme d'amorçage masqué sont très limités. Un effet d'amorçage n'apparaît que lorsque le lien entre l'amorce et la cible est pertinent en regard de la tâche demandée au participant (i.e., tâche de catégorisation : être vivant vs objet manufacturé ou encore valence positive vs valence négative). Il se pourrait ainsi que le masquage de l'amorce génère une réallocation des ressources cognitives sur les traitements superficiels. De ce fait, très peu de ressources cognitives sont disponibles pour effectuer le traitement sémantique. Il est alors vraisemblable que le système alloue le peu de ressources disponibles à l'activation des sens pertinents dans le contexte. Ceci expliquerait l'impact important de la tâche sur le traitement sémantique reflété par l'effet de congruence (Eckstein & Perrig, 2007; Klinger et al., 2000; Spruyt et al., 2012).

2.2. Retour sur les modèles de traitement sémantique

Certains modèles de traitement sémantique présentés en partie Théorique sont compatibles avec l'idée selon laquelle le traitement sémantique serait dépendant de ressources cognitives.

Ainsi, le modèle de Collins et Loftus (1975), suppose que la propagation de l'activation le long du réseau diminue avec le temps et s'arrête lorsque l'attention est dirigée vers une autre tâche, suggérant que le traitement sémantique est coûteux en ressources. Toutefois ce modèle explique difficilement les résultats de la littérature sur l'amorçage subliminal. En effet, comme nous l'avons vu en partie théorique, la propagation de l'activation ne permet pas d'expliquer les effets de congruence observés, par exemple, l'amorçage de *diable* par *bébé* quand la tâche de catégorisation se fait sur les concepts *être vivant* et *objet manufacturé* mais pas quand elle se fait sur les concepts *valence positive* et *valence négative* (Eckstein & Perrig, 2007).

L'idée de l'existence d'un administrateur central du modèle proposé par Posner et Snyder (1975) est cohérente avec l'*Effortfulness Hypothesis*. Toutefois, ce modèle prévoit également l'existence de processus purement automatiques et indépendants de l'attention du sujet ou de toute forme de fonction cognitive de plus haut niveau. De plus, la principale caractéristique proposée par Posner et Snyder quant à la différenciation entre les processus stratégiques et automatiques est l'inhibition des concepts non-pertinents. Nos données ne permettent pas d'émettre d'hypothèse quant à cet aspect du modèle.

Les modèles purement stratégiques comme le modèle de Becker (1980), le modèle de Norris (1986) et le modèle de Ratcliff et McKoon (1988) ne permettent pas d'expliquer les données de la littérature sur l'amorçage subliminal. En effet, le modèle de Becker suggère l'existence de processus uniquement stratégiques, au cours desquels le système sélectionnerait un ensemble de mots vraisemblables. Toutefois, ces processus stratégiques ne permettent pas d'expliquer les effets d'amorçages obtenus aux tâches de catégorisation impliquant de larges sets de stimuli. Ainsi, certaines tâches de catégorisation demandent aux participants de déterminer si un concept a une valence positive ou une valence négative. Même si les participants, après avoir perçu l'amorce masquée et l'avoir traitée, produisaient un set de mots attendus, celui-ci serait si large qu'il est peu probable qu'il entraîne une reconnaissance plus rapide du mot cible. De plus, après la présentation d'une amorce ayant une valence positive, la probabilité pour que la cible ait également une valence positive est de 50%. Il n'y a donc aucune raison pour que les participants génèrent une attente sur une cible ayant une valence positive puisque générer une attente sur une cible ayant une valence négative serait potentiellement tout aussi bénéfique. De même pour le modèle de Ratcliff et McKoon, les tâches de catégorisation utilisées dans les paradigmes d'amorçage subliminal génèrent des sets si larges qu'il rend le principe de facilitation par familiarité peu probable. Par exemple, selon les normes françaises d'associations de mots abstraits, *diable* (valence négative) est le plus souvent associé à *enfer* (valence négative) et *ange* (valence positive). Selon ce modèle, la présentation de *diable* avec le mot *ange* devrait entraîner un effet d'amorçage plus grand que *diable* et *accident* qui ont tous deux une valence négative.

Enfin le modèle hybride de Neely et Keefe (1989) est soumis aux mêmes limites que les modèles stratégiques et automatiques présentés dont il est une compilation.

2.3. Automaticité et ressources cognitives

Comme nous l'avons vu dans la partie Théorique et la présentation des modèles de traitement sémantique (cf. p32) il existe une distinction systématique entre la notion d'automaticité et de coût cognitif. Les données recueillies ici tendent pourtant à mettre en évidence que le traitement sémantique est dépendant de la disponibilité des ressources cognitives. En effet, nous avons mis en évidence dans deux études

distinctes, utilisant des indices différents, que le traitement sémantique pouvait ne pas être mesurable lorsque le signal de parole était dégradé, en dépit d'une intelligibilité préservée. Ces résultats tendraient donc à suggérer que le traitement sémantique est purement stratégique, mais cette hypothèse n'est pas congruente avec la littérature sur l'amorçage sémantique masqué.

A la lumière de la littérature et de nos résultats, nous proposons un modèle adapté de celui de Posner et Snyder (1975). Nous postulons comme eux, qu'il existe deux processus responsables du traitement sémantique. L'un est stratégique et permet d'expliquer les données recueillies par Neely (1977 ; cf. p35). Il résulte en fait simplement d'une mise en place de stratégies conscientes dans le but d'améliorer les performances à une tâche et n'est vraisemblablement observable qu'en conditions expérimentales. L'autre processus serait semi-automatique. Nous employons le terme semi-automatique ici pour signifier que le traitement sémantique serait dépendant des ressources cognitives sans pour autant être stratégique *per se*. C'est-à-dire que le système met en place des « stratégies » pour améliorer les performances à la tâche sans pour autant que le sujet en soit conscient. Ceci expliquerait ainsi les effets attentionnels (Naccache et al., 2002) et de la nature de la tâche (Eckstein & Perrig, 2007; Spruyt et al., 2012) observés dans la littérature sur l'amorçage subliminal.

Ce second processus reposerait sur une propagation automatique de l'activation. L'activation serait dépendante des ressources disponibles, lorsque peu de ressources sont disponibles, l'activation se propage peu. Comme postulé par Collins et Loftus (1975) l'activation doit dépasser un certain seuil pour qu'un concept soit activé. La nature de la tâche pourrait moduler ce seuil. En d'autres termes, un concept pertinent pour la tâche demandée a un seuil d'activation plus bas qu'un concept non pertinent. Cette variation des seuils permet de compenser le manque d'efficacité du traitement dû au manque de ressources cognitives en condition d'amorçage masqué.

Le contexte pourrait également moduler la propagation de l'activation. Ainsi, comme nous l'avons observé dans l'Etude 2, un contexte fort va empêcher la propagation de l'activation vers des concepts jugés non pertinents.

3. Interactions entre processus *top-down* et *bottom-up*

3.1. Au niveau comportemental

Comme nous l'avons vu dans la partie Théorique, les interactions entre processus *top-down* et *bottom-up* sont au cœur de nombreux débats, aussi bien au niveau de l'accès lexical que du traitement sémantique. La situation de parole dans le bruit, comme les autres dégradations de la parole (e.g., *noise vocoded speech*) permet de s'intéresser à ces interactions en diminuant la qualité des informations *bottom-up* (Aydelott et al., 2006; Obleser & Kotz, 2010; Strauss et al., 2013). Ainsi, lorsque la *cloze probability* d'une phrase est importante (e.g., *Elle boit une BIERE*), alors son intelligibilité est meilleure

lorsque le signal est dégradé comparativement à une phrase dont la *cloze probability* est faible (e.g., *Elle voit une BIERE*). En d'autres termes, l'augmentation des informations *top-down* disponibles diminue généralement l'effet de la dégradation des informations *bottom-up* (Obleser & Kotz, 2010, 2011; Strauss et al., 2013).

Dans l'Etude 2, nous ne retrouvons aucun effet significatif du contexte sur l'intelligibilité des phrases. Plusieurs facteurs pourraient expliquer cette absence d'effet. D'une part, la dégradation du signal pouvait être à la fois insuffisante dans le bruit stationnaire et trop importante dans le bruit parolier. Dans la première condition, les participants n'auraient pas eu besoin du contexte pour reconnaître les mots de la phrase. Dans la seconde, la dégradation du signal est telle que la présence du contexte est inutile.

D'autre part, la *cloze probability* des phrases au sein de nos différentes conditions étaient contrôlées. C'est-à-dire que les phrases présentées en condition Congruente généraient autant d'attentes sur le mot cible que les phrases présentées en condition Ambigüe ou Incongruente. Simplement, en condition Congruente le mot cible était le mot attendu. Toutefois, comme le mot cible était présenté dans le silence, le contexte, qu'il soit cohérent ou non n'avait aucun impact sur son intelligibilité.

Dans cette étude donc les résultats comportementaux ne permettent pas de déterminer si les informations *top-down* ont facilité le traitement des informations *bottom-up*.

3.2. Au niveau électrophysiologique

Nous proposons, au vu de nos résultats que lorsque le contexte est traité efficacement (i.e., en silence), les prédictions *top-down* sont importantes et précises. En revanche, lorsque le signal est dégradé, le contexte est traité de manière plus superficielle, et les prédictions sont moins efficaces (cf. Figure 45). De ce fait, l'ensemble des sens du mot ambigu sont vraisemblablement activés, avec une préférence pour le sens correspondant au contexte. Ceci est reflété par une N400 d'amplitude plus petite en réponse à un mot cible faisant référence au sens subordonné du mot ambigu qu'en réponse à un mot cible incongruent.

La prédiction serait donc moins forte au niveau sémantique lorsque le contexte est légèrement dégradé. Toutefois, il est intéressant de noter qu'une PMN est observée aussi bien lorsque le mot cible est incongruent que lorsqu'il fait référence au sens subordonné du mot cible. Il semble ainsi qu'en dépit de l'activation des deux sens du mot ambigu, les participants prédisent la présentation d'un mot congruent avec le sens favorisé par le contexte. De ce fait, la cohorte pré-sélectionnée correspond à un mot congruent avec ce sens, et le changement de cohorte engendré par la présentation d'un mot cible incongruent ou ambigu, génère une PMN.

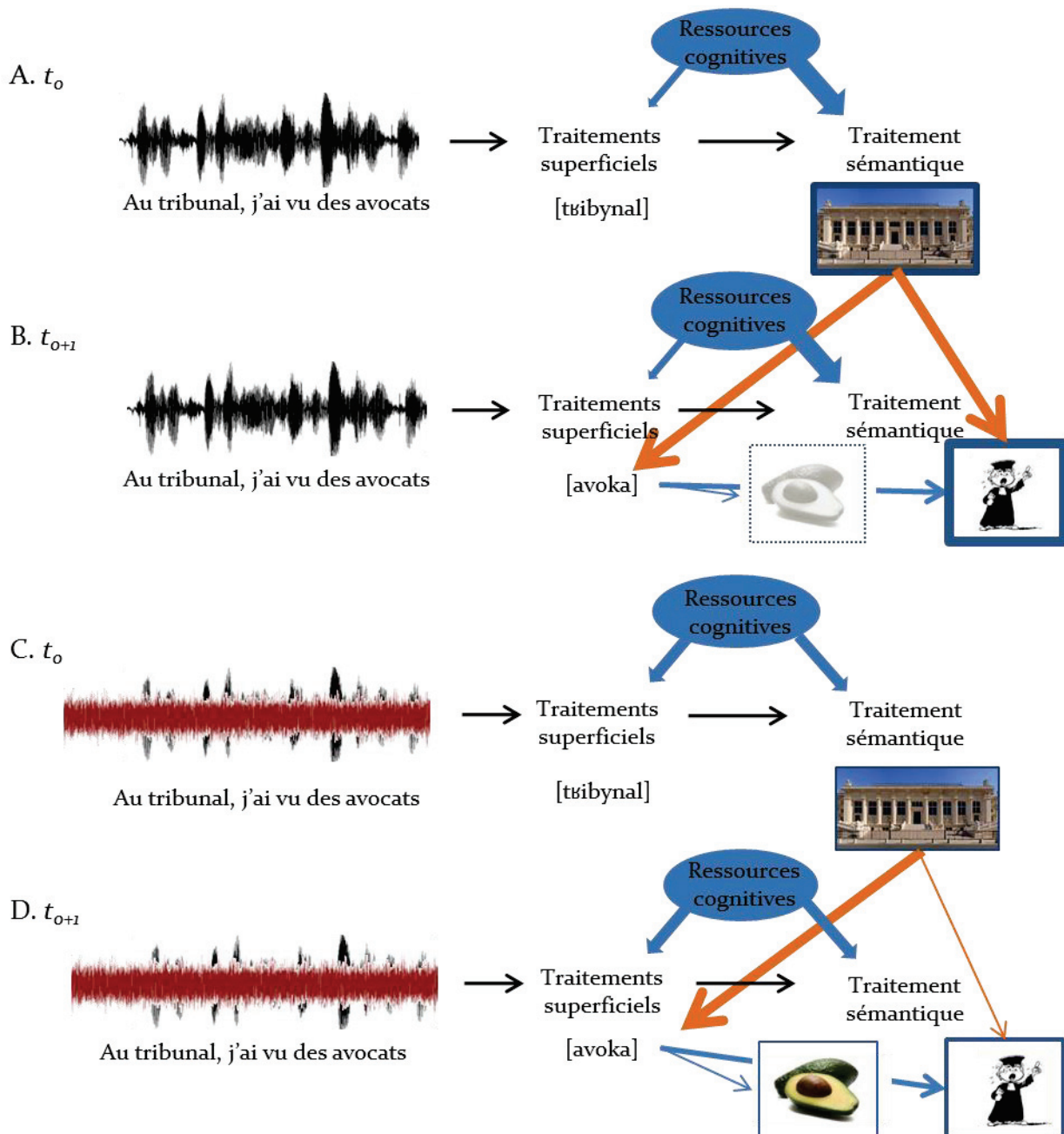


Figure 45

Représentation graphique des interactions *bottom-up* et *top-down* dans l'Etude 2. En condition (A) et (B) silence et (C) et (D) stationnaire. Traitement du mot *tribunal* à t_0 et traitement du mot *avocat* à $t+1$. Les flèches orange représentent l'amorçage. L'épaisseur des flèches représente la force de l'activation.

Il apparaît alors que dans le silence, le contexte est plus important que les informations *bottom-up*, alors que dans le bruit stationnaire, le traitement du contexte est moins sûr, les prédictions sont donc plus larges.

4. Parole dans le bruit et rythme α

Les résultats que nous avons obtenus s'inscrivent parfaitement au sein de l'*Effortfulness Hypothesis*. Ils mettent en avant l'importance de l'intégrité des

informations de bas niveau sur les traitements de haut niveau. Toutefois, si ce travail nous a permis de mettre en évidence l'impact de la dégradation du signal sur l'efficacité des traitements sémantiques, les substrats neuronaux, responsables de ce phénomène n'ont pas été étudiés. Il serait pertinent de poursuivre ce travail en s'intéressant à l'origine de cet effet.

Bien que l'*Effortfulness Hypothesis* soit exprimée en termes d'activations, des données récentes suggèrent qu'elle pourrait en fait reposer sur des perturbations de l'inhibition. En effet, il a été montré que les oscillations *alpha* observées en situation de parole dégradée reflétaient l'inhibition des informations non-pertinentes (e.g., bruit ; Obleser & Weisz, 2012; Strauss, Wostmann, & Obleser, 2014). Ainsi, des différences de puissance d'*alpha* entre les hémisphères au niveau pariétal ont été mises en évidence en condition d'écoute dichotique en fonction du canal à écouter (Kerlin, Shahin, & Miller, 2010). Dans une condition d'écoute dichotique, les participants devaient indiquer si la phrase présentée dans le canal droit ou gauche en fonction des essais était cohérente ou non. L'analyse en temps fréquence des données EEG recueillies ont mis en évidence une suppression des oscillations *alpha* au niveau pariétal contralatéral accompagnée d'une augmentation de la puissance *alpha* ipsilatérale lorsque les participants parvenaient à écouter la phrase cible. Le rôle inhibiteur des oscillations *alpha* a également été souligné dans une tâche de mémoire de travail. Dans une étude en MEG, Obleser et ses collègues (Obleser et al., 2012) ont mis en évidence que conserver les informations en mémoire de travail (i.e., inhiber la relâche en mémoire) génère des oscillations *alpha* provenant du même réseau que celles générées par la présence de bruit. Au vu de ces résultats, les auteurs proposent que l'effet d'*Effortfulness* résulte d'une surcharge de ce réseau, qui ne peut pas inhiber à la fois le bruit et la relâche de l'information simultanément.

En s'appuyant sur ces résultats nous posons l'hypothèse selon laquelle un phénomène similaire pourrait être à l'origine des perturbations du traitement sémantique en situation de parole dans le bruit. En effet les oscillations *alpha* semblent également être impliquées dans le traitement sémantique (Rohm, Klimesch, Haider, & Doppelmayr, 2001; Willems, Oostenveld, & Hagoort, 2008). Ainsi, Rohm et al., 2001, ont mis en évidence une désynchronisation des oscillations *alpha* pendant le traitement sémantique. Des morceaux de phrases étaient présentés visuellement aux participants. Ceux-ci devaient soit lire les phrases soit générer pour le troisième groupe de mots présentés au sein de chaque phrase, le superordonné (i.e., le représentant de la catégorie). Par exemple pour la phrase *Un lapin est dans la BOITE, en train de se cacher*, les participants devaient générer le mot *récepteur* comme superordonné de *BOITE*. Puisque les oscillations *alpha* sont généralement associées à l'activité d'inter-neurones inhibiteurs, la suppression du rythme *alpha* lors du traitement sémantique pourrait refléter la levée d'inhibition et donc la propagation de l'activation au sein du réseau sémantique.

De ce fait, en présence de bruit, il est possible que l'augmentation de puissance des oscillations *alpha* reflétant l'inhibition du bruit empêche la suppression *alpha*, et donc la levée d'inhibition et la propagation de l'activation sémantique.

L'Axe 1 a ainsi permis de s'intéresser à l'effet de la dégradation transitoire du signal de parole sur la compréhension de celle-ci. Il a permis de remettre en question l'automaticité du traitement sémantique qui serait plutôt dépendant de la disponibilité de ressources cognitives. De plus, il semble que lorsque la dégradation du signal augmente, le système repose moins sur les informations *top-down* puisque le contexte est traité de façon moins efficace.

Chapitre XII

Discussion de l'Axe 2

Le deuxième axe du travail expérimental était constitué de deux études s'intéressant aux traitements auditifs centraux, à leur développement au sein d'une population d'enfants présentant un développement typique, à leurs troubles au sein d'une population d'adultes dyslexique et à leurs liens avec les compétences langagières. Ces deux études ont utilisé la Batterie d'Evaluation des Compétences Auditives Centrales (BECAC ; Donnadieu et al., 2014)

1. Résumé des résultats

1.1. Etude 3

L'Etude 3 a donné lieu à deux publications distinctes. La première s'est intéressée au développement des compétences auditives centrales testées par 3 des 6 tests de la BECAC et à leurs liens avec le développement du langage. Les tests analysés dans cette étude étaient le test de latéralisation, le test de discrimination auditive (fréquence et durée) et le test de masquage central. Les résultats ont été analysés en calculant les Hit- FA pour chaque participant afin de minimiser l'impact des facteurs non-sensoriels, particulièrement important au sein de la population infantine (Moore, 2012). Une seconde analyse des données recueillies sur la BECAC par Donnadieu et son équipe a permis d'obtenir des résultats comparables à ceux rapportés par la littérature déjà existante sur le développement auditif. Les données langagières présentant des effets planchers et plafonds importants, elles n'ont pas été incluses dans la nouvelle analyse. Dans cette nouvelle analyse, a également été inclus un groupe de jeunes adultes sains, permettant de représenter l'issue du développement observé.

Afin d'extraire le seuil auquel chaque participant était capable de latéraliser un son pur ou de discriminer deux sons purs à 75%, la fonction psychométrique se rapprochant le plus possible des performances, a été calculée. En plus de permettre de comparer nos données avec celles de la littérature, l'estimation de la fonction psychométrique a permis de minimiser l'impact des erreurs d'inattention chez les participants.

La première analyse n'a pas mis en évidence d'effet développemental sur les compétences de latéralisation d'un son, alors que cet effet est apparu significatif lors de la deuxième analyse. Toutefois, dans les deux études, les résultats suivent la même trajectoire, avec une amélioration entre le CP et le CE2 puis une stabilisation entre le CM1 et le CM2. Les résultats des autres tests sont cohérents entre les deux analyses. L'analyse des résultats a révélé un effet significatif du développement sur l'ensemble des tests administrés à l'exception du test de ségrégation de flux. Parmi les autres compétences, seules deux étaient matures avant la fin de l'école primaire : la latéralisation et la discrimination de durée. Les autres compétences évaluées dans cette étude semblent avoir un développement plus lent, se terminant vraisemblablement durant l'adolescence, ce qui semble cohérent avec d'autres données concernant le développement auditif (Eggermont, 1988, 2014; Moore, 2002). Il semble donc y avoir une grande variabilité dans le développement des compétences auditives testées. Cette variabilité se retrouve aussi au sein de la population. En effet certains enfants obtiennent des performances similaires à celles des adultes dans plusieurs tests voire tous les tests (4 enfants) alors que d'autres n'atteignent ces performances à aucun des tests présentés (2 enfants).

Les compétences auditives centrales évaluées étaient également liées aux performances à certains tests verbaux. Ainsi, la capacité à discriminer la fréquence de deux sons purs était liée chez les enfants de CM1 à la conscience phonologique et au vocabulaire. Les performances au test de masquage central ainsi qu'au test de discrimination de durée étaient respectivement liées à la conscience phonologique et au niveau de vocabulaire chez les enfants de CE2.

Seules les compétences en latéralisation ne présentaient pas de liens avec les compétences langagières. Ces résultats semblent cohérents puisque les performances en latéralisation semblent atteindre la maturité tôt durant le développement, au plus tard autour de 9-10 ans. Elles sont vraisemblablement stables d'un enfant à l'autre et n'ont donc pas (ou tout au moins plus) d'influence sur la perception du langage et donc sur ses représentations.

L'Etude 3 a ainsi permis de s'intéresser au développement des compétences auditives centrales chez les enfants d'école primaire. L'utilisation de la BECAC n'utilisant pas de matériel verbal a permis de minimiser l'impact du développement langagier sur ces performances.

1.2. Etude 4

Dans la dernière étude de la partie Expérimentale, nous nous sommes intéressés à la présence des TTA dans une population de dyslexiques adultes. L'objectif était d'évaluer l'ensemble des compétences auditives centrales comme définies par l'ASHA (2005) afin de déterminer celles qui étaient déficitaires et leurs liens avec les compétences

langagières (conscience phonologique, lecture et vocabulaire). La BECAC a donc été adaptée à l'adulte, et présentée à deux types de populations : un groupe de dyslexiques et un groupe contrôle, appariés en âge, intelligence non-verbale, et latéralité manuelle. Comme pour la deuxième analyse de l'Etude 3, les fonctions psychométriques se rapprochant le plus des performances de chaque participant ont été calculées pour les tests de latéralisation et de discrimination auditive. La comparaison entre les groupes a mis en évidence que les participants dyslexiques sont significativement moins performants que les contrôles dans les tests impliquant un traitement spectral (i.e., test de discrimination de fréquence et test de reconnaissance de pattern auditif fréquentiel) ou temporel (i.e., test de discrimination de durée et test de reconnaissance de pattern auditif temporel). De plus, seules les performances à ces tests sont apparues comme ayant un lien avec les compétences langagières des participants. Particulièrement, des corrélations entre les performances à ces tests et les performances aux tests de conscience phonologique et de lecture, aussi bien au niveau de la population totale qu'au niveau du groupe dyslexique uniquement se sont révélées significatives. Suggérant ainsi un véritable lien de cause à effet et non un simple effet de groupe. Enfin les participants avaient effectué un test de vocabulaire afin d'évaluer si la dégradation du signal impactait également les représentations langagières de haut niveau. Aucune différence significative n'est apparue entre les deux groupes et les résultats au test de vocabulaire ne corrélaient avec aucun des tests auditifs. Il semble ainsi que contrairement à ce qui est observé dans le cas de la parole dans le bruit, la dégradation du signal par un déficit au niveau des traitements auditifs centraux n'a pas d'effets sur les représentations langagières de haut niveau. Notons toutefois que notre propos se doit d'être modulé par le fait que notre groupe dyslexique était principalement composé d'étudiants, ayant donc bénéficié de rééducation et ayant réussi à compenser au moins partiellement leur déficit. Ils ne sont donc peut-être pas représentatifs de la population dyslexique de ce point de vue.

2. Les facteurs non-sensoriels

2.1. Diminution des facteurs attentionnels et cognitifs

Comme nous l'avons mentionné au sein de la partie théorique, il est très difficile de distinguer les facteurs sensoriels des facteurs non-sensoriels particulièrement dans l'étude du développement. En effet, au sein d'une population infantine, les facteurs non-sensoriels sont susceptibles d'être plus importants qu'au sein d'une population adulte, et donc d'influer sur les résultats (Dawes & Bishop, 2008; Moore et al., 2011; Moore et al., 2008). Toutefois, ces effets ont été minorés au sein de notre étude, d'une part par le choix de l'analyse et d'autre part par le choix des paradigmes.

En effet, tout d'abord, l'analyse par la méthode des Hit – FA dans l'Etude 3a permet d'estimer et de prendre en compte les Fausses Alarmes, donc les erreurs d'inattention des participants afin que celles-ci influencent moins les résultats finaux (Abdi, 2007).

Toutefois cette analyse ne permettait pas de comparer les résultats obtenus à ceux de la littérature qui rapporte le plus souvent les seuils à partir desquels les participants sont capables de discriminer deux stimuli. Une deuxième analyse a donc été effectuée, visant toujours à minimiser l'impact des facteurs non-sensoriels. En effet, l'estimation des courbes psychométriques correspondant le mieux aux données de l'individu a permis de diminuer l'impact de certaines réponses aberrantes. Ainsi, un participant pouvait différencier A et X à 75% pour une différence donnée (a) et ne les différencier qu'à 25% lorsque la différence était un peu plus grande (b). Dans cette situation, notre analyse permettait d'estimer la fonction psychométrique qui allait minimiser la différenciation en situation (a) et la maximiser en situation (b) afin d'estimer un seuil plausible au regard des résultats de l'individu.

Enfin, un test de χ^2 a été effectué pour les tests pour lesquels avait été estimée une fonction psychométrique. L'objectif de ce test était d'évaluer si la proportion de participants exclus dans chaque groupe était équivalente. Ce test est apparu significatif uniquement sur les performances au test de discrimination de fréquence. Ceci suggère donc que le choix d'extraire les seuils de discrimination à 75% permet de limiter l'impact des facteurs non sensoriels sur les résultats. En effet, dans les autres tests (i.e., latéralisation et discrimination de durée et d'intensité), la proportion de participants ayant des scores considérés comme aberrants étaient équivalente au sein des différents groupes.

2.2. Diminution de l'effet du paradigme

Tester les compétences en discrimination auditive avec le même paradigme sur plusieurs dimensions permet d'une part de comparer ces différentes dimensions et d'autre part, d'écarter l'hypothèse d'un effet provenant uniquement du paradigme utilisé. Ainsi, dans l'Etude 3b, un effet significatif apparaît entre la classe de CM2 et l'âge adulte pour la discrimination de fréquence, entre le CP et le CM1 pour la discrimination de durée et du CE2 à l'âge adulte pour la discrimination d'intensité. L'observation de trajectoires développementales différentes permet de s'assurer que les effets observés ne reflètent pas uniquement l'amélioration de la capacité à exécuter la tâche demandée. De même, dans l'Etude 4, les dyslexiques obtiennent des seuils significativement plus élevés que ceux des sujets contrôles en discrimination de fréquence, alors que la différence de seuil est seulement tendancielle en discrimination de durée et n'est pas significative en discrimination d'intensité. Le paradigme choisi ne semble donc pas avoir d'influence sur les différences minimales entre A et X nécessaire à leur discrimination. Nous reviendrons sur le cas spécifique du paradigme AX dans la partie sur l'*Anchor Theory* (cf. p247).

Concernant le test de reconnaissance de pattern auditif, pour l'Etude 3 il est possible également de voir que si le paradigme a certainement un effet sur les résultats des participants, un effet du développement des compétences auditives centrales est

également présent. En effet, deux trajectoires développementales différentes sont observées entre le sous-test fréquence et le sous-test durée, suggérant ainsi que deux processus différents sont à l'origine de ces effets.

3. Retour sur l'Anchor Theory

Comme nous l'avons vu dans la partie Théorique, l'*Anchor Theory* (Ahissar, 2007; Ahissar et al., 2006) postule que les dyslexiques présentent des difficultés à former un ancrage perceptuel. C'est-à-dire qu'il leur est plus difficile de former et de conserver une représentation d'un stimulus perçu à plusieurs reprises afin d'effectuer des comparaisons avec les autres stimuli présentés.

Les résultats obtenus dans l'Etude 4 ne sont pas totalement en accord avec cette théorie. En effet, selon l'*Anchor Theory*, le paradigme qui a été choisi pour le test de discrimination auditive (i.e., paradigme AX) devrait entraîner des résultats déficitaires pour l'ensemble des sous-tests dans la population dyslexique. En effet, selon cette théorie, le déficit habituellement observé en discrimination de fréquence au sein de la population dyslexique provient de leur incapacité à bénéficier de la répétition du standard comparativement aux contrôles. Les performances de la population dyslexique devraient donc se trouver impactées sur l'ensemble des dimensions testées (i.e., fréquence, durée et intensité). Or, les résultats mettent en évidence que les participants dyslexiques sont moins performants que les contrôles uniquement lorsque les sons purs diffèrent en fréquence. Ils ont tendance à être moins performants que les contrôles en discrimination de durée, mais obtiennent les mêmes scores en discrimination d'intensité.

Il est toutefois intéressant de noter que les participants dyslexiques obtiennent de moins bonnes performances que les contrôles au test de reconnaissance de pattern auditif, aussi bien en fréquence qu'en durée. Dans chacun de ces sous-tests, deux sons purs sont répétés dans différentes configurations : sons purs Aigu et Grave pour le sous-test de fréquence et Court et Long pour la durée. Utiliser une ancre perceptuelle pourrait être très bénéfique dans ces tests. En effet, avoir une bonne représentation interne des sons purs utilisés rend le test plus facile puisqu'il suffit alors de se rappeler l'ordre de présentation pour effectuer la tâche (i.e., Long – Court). En revanche si les participants dyslexiques doivent à chaque essais traiter à nouveau les sons purs et les comparer entre eux (e.g., le premier son pur est plus long que le second) alors la tâche est beaucoup plus coûteuse, pouvant expliquer les moins bonnes performances des dyslexiques.

Comme expliqué dans la partie Théorique, Ahissar et collègues ont observé que les participants contrôles bénéficient habituellement des paradigmes répétant le même ton de comparaison essai après essai. Les dyslexiques en revanche ne bénéficient pas de cette répétition. Dans notre l'absence de différence significative entre les seuils de

discrimination de l'intensité de deux stimuli des participants dyslexiques et contrôles provient peut être de la facilité de la tâche. En effet, si la discrimination d'intensité est plus facile que la discrimination de fréquence alors il est possible que l'utilisation du même ton de comparaison ne soit pas bénéfique (i.e., les résultats plafonnent). De ce fait, les participants dyslexiques peuvent obtenir les mêmes scores que les contrôles, sans bénéficier de la répétition du ton de comparaison. Toutefois cette hypothèse est à moduler : les résultats obtenus dans l'Etude 3b, ne semblent pas indiquer que la discrimination d'intensité soit particulièrement facile puisqu'elle ne devient pas mature avant l'âge adulte.

Enfin une dernière possibilité serait que les participants dyslexiques soient particulièrement déficitaires à former une ancre perceptuelle d'informations spectrales ou temporelles. Ainsi, le son pur représenté en mémoire iconique serait peu précis en termes de fréquence et de durée mais pas en termes d'intensité. Cette hypothèse suggérerait alors que l'*Anchor Theory* est spécifique à certaines dimensions, contrairement à ce qui avait été pensé par Ahissar et collègues (Ahissar et al., 2006 ; Ahissar, 2007) initialement. Encore une fois, il semble que le déficit auditif observé chez les dyslexiques soit variable d'une population à l'autre.

4. Retour sur les TTA dans la dyslexie

La partie Théorique nous a permis de souligner l'inconsistance des résultats de la littérature sur les liens entre les TTA et la dyslexie. Certains des tests que nous avons utilisés sont souvent présentés dans la littérature et nous allons donc voir comment nos résultats peuvent s'y intégrer. Concernant la discrimination de fréquence, les résultats obtenus sont cohérents avec la grande majorité de la littérature, mettant en évidence un déficit chez les dyslexiques, comparativement aux contrôles.

En discrimination de durée, les données de la littérature montraient une différence significative entre les dyslexiques et les contrôles selon la durée des stimuli. Lorsque le son pur de référence durait 400 ms celui-ci était aussi bien discriminé par les deux populations alors que les enfants contrôles étaient plus performants lorsque le ton de référence durait 100 ms (Banai & Ahissar, 2006). En revanche, les adultes dyslexiques n'étaient pas déficitaires lorsque le son pur de référence durait 200 ms (Baldeweg et al., 1999). Nos résultats montrent une tendance des dyslexiques à avoir besoin d'une plus grande différence entre A et X pour les discriminer. Etant donné la durée du son pur de référence (75 ms) il semble possible que les dyslexiques présentent un déficit à discriminer les durées des stimuli courts lorsqu'ils sont enfants et cette difficulté pourrait diminuer mais persister à l'âge adulte, expliquant nos résultats. Comme nous le soulignons dans la discussion de l'Etude 3b, il semble que la durée du stimulus ait un impact sur la différence de temps nécessaire pour discriminer deux stimuli et également sur l'âge de maturation de cette compétence suggérant peut-être l'intervention de deux processus distincts.

Enfin, les résultats aux tests de FPT et DPT sont également cohérents avec la littérature suggérant que les participants dyslexiques présentent un déficit dans la reconnaissance de pattern auditif (King et al., 2003). Il est également intéressant de noter, dans notre étude que les participants dyslexiques et contrôles sont plus déficitaires lorsque la durée des sons purs est variée que lorsque c'est la fréquence. Il pourrait être argumenté alors que les dyslexiques ont simplement plus de difficulté à exécuter la tâche, à cause du paradigme et non à cause de TTA. Toutefois, le fait que l'on retrouve des performances déficitaires à la fois en discrimination de fréquence et en discrimination de durée, suggère qu'il existe au sein de la population testée un déficit spécifique sur ces traitements. De plus, nous avons observé dans l'Etude 3 des trajectoires développementales différentes entre les tests de FPT et DPT. Ceci suggère alors que les processus mis en jeu par ces deux sous-tests sont distincts, et que les déficits observés au sein de la population dyslexique ne peuvent pas être uniquement attribués à des difficultés à effectuer le paradigme.

5. Critiques

Malgré les nombreuses évidences de la présence de TTA dans la population dyslexique, certains auteurs remettent en cause la présence d'un lien de cause à effet entre TTA et dyslexie. Un des arguments utilisé est que tous les dyslexiques ne présentent pas de TTA. Entre 40 et 60% des dyslexiques en moyenne présentent un déficit auditif central, nous retrouvons également cette proportion dans notre échantillon. Neuf dyslexiques sur 20 présentent des scores déficitaires c'est-à-dire distants de plus de 2 écart-types de la moyenne des contrôles à au moins 2 des tests proposés (critère diagnostic proposé par l'ASHA cf. p85). Certains auteurs considèrent donc que puisque seule une minorité des dyslexiques présentent un déficit au niveau des TTA, ces troubles ne peuvent être considérés comme étant à l'origine de la dyslexie (Ramus & Ahissar, 2012; Rosen, 2003). Toutefois, bien qu'une majorité de dyslexiques ne présentent pas de TTA, c'est le trouble le plus fréquent après le trouble phonologique (cf. Figure 46). De plus, la définition de la dyslexie repose sur la description de ses symptômes, elle n'exclut donc pas la possibilité que les mêmes symptômes soient causés par différents déficits. En effet, il est maintenant accepté qu'il existe une sous-population spécifique de dyslexiques, ne présentant pas de trouble phonologique mais un déficit attentionnel, générant les mêmes difficultés à apprendre à lire (Bosse et al., 2007; Lallier et al., 2009). Il serait possible que le trouble phonologique soit également généré par différents déficits.

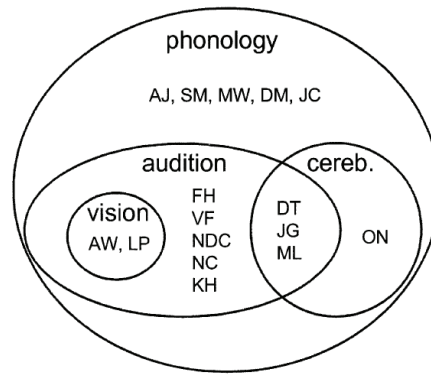


Figure 46

Tiré de Ramus et al. (2003). Répartition des différents troubles associés à la dyslexie (i.e., trouble phonologique ; trouble cérébelleux ; trouble auditif ; trouble visuel). Les initiales correspondent à chaque participant de l'étude.

Une autre critique apportée à l'hypothèse d'un lien causal entre TTA et dyslexie, est que certaines études ne trouvent pas de déficit au sein de leur groupe dyslexique. Nous avons déjà évoqué dans la partie Théorique la petitesse des échantillons testés dans certaines études (e.g., Rosen & Mangarini, 2001). Toutefois d'autres problèmes peuvent être soulevés sur le choix de la population. Certains auteurs vont ainsi choisir des participants avec une histoire de dyslexie durant l'enfance (Ahissar et al., 2000). Ainsi il est délicat de savoir si ces participants présentent toujours des symptômes de dyslexie ou si leur trouble est maintenant complètement compensé. D'autant plus qu'un diagnostic fait durant l'enfance répond à un autre objectif qu'un diagnostic effectué pour une recherche et pourrait donc s'avérer moins restrictif (cf. p92). De ce fait, il est délicat, sinon abusif de tirer des conclusions sur l'origine de la dyslexie sans évaluer les compétences langagières de la population évaluée au moment du test.

Le contrôle du QI non verbal peut également être questionnable. Souvent, les auteurs s'assurent simplement qu'il est dans la norme (i.e., >80 ; e.g., Hornickel et al., 2012). Il est alors impossible de distinguer les mauvais lecteurs et les dyslexiques. En effet, si les participants ont simplement un QI non verbal et des scores en lectures plutôt bas, ils seront davantage considérés comme mauvais lecteurs que comme dyslexiques. Le diagnostic de dyslexie demande en effet un retard en lecture comparativement à ce qui pourrait être attendu au regard du QI non verbal. Avec un QI non verbal à 80, un score en lecture peut être déficitaire sans pour autant être très différent de ce qui est attendu au regard du QI, déjà à la limite inférieure de la norme. Le terme de mauvais lecteur serait donc plus approprié, il ferait référence à un déficit en lecture relativement modéré au regard du QI non verbal. Dans un objectif de recherche, il semble alors abusif de les considérer comme dyslexiques, particulièrement si l'on considère la dyslexie comme un trouble à part entière résultant de mécanismes différents que ceux sous-tendants de simples difficultés en lecture. De plus, les tâches présentées aux

participants sont coûteuses en attention et demandent la compréhension de concepts assez abstraits (e.g., modulation de fréquence). Rosen (2003) montre d'ailleurs l'importance du QI non verbal par exemple dans les compétences en lecture de pseudo-mot dans l'étude de Ahissar et al., (45 % de la variance à cette tâche est expliquée par le QI non verbal).

L'ensemble de ces observations pourraient en partie expliquer le manque de consistance des résultats de la littérature, ou tout au moins, empêche de tirer des conclusions quant à l'absence de résultats de certaines études.

Notre étude avait ainsi l'avantage de tester un plus grand échantillon de dyslexiques, présentant de très bonnes compétences non-verbales et appariés aux sujets contrôles un à un. De ce fait, les différences de performances aux tests présentés ne peuvent pas provenir de facteurs non-sensoriels.

6. Liens entre compétences auditives centrales et langagières

Nous mentionnons dans la partie Théorique que certaines études s'étant intéressées aux liens entre les TTA et les compétences langagières ne mettent en évidence de corrélations qu'au niveau de la population totale. Ces corrélations sont alors susceptibles de refléter uniquement un effet de groupe plutôt qu'un lien réel entre le TTA et la compétence langagière étudiée. En effet, si les participants contrôles sont à la fois meilleurs aux tests auditifs et langagiers alors les corrélations seront inévitablement significatives entre ces différents tests. Elles ne seront pourtant que le reflet de l'effet de groupe et non d'un lien entre les compétences auditives et langagières.

Notre étude a permis de mettre en évidence au sein de la population dyslexique spécifiquement l'existence de liens entre les traitements spectral et temporel et les performances aux tests de conscience phonologique et de lecture.

Il est intéressant de noter que des liens significatifs entre traitements auditifs centraux et compétences langagières sont révélés uniquement sur la population dyslexique, en effet, il semble que les compétences auditives centrales aient moins d'impact sur le langage chez les participants contrôles. Lors du développement du langage, il est évident que les représentations phonologiques vont dépendre de la bonne qualité du percept. Nous postulons que chez le sujet contrôle, lorsque les compétences langagières sont matures, elles deviennent indépendantes de la qualité du percept. De ce fait, les plus ou moins bonnes compétences auditives centrales n'ont pas d'impact sur la qualité des représentations phonologiques. En revanche, il semble que dans la population dyslexique, les représentations phonologiques n'atteignent pas cette indépendance, expliquant les corrélations significatives.

Cette variabilité dans l'impact des compétences auditives centrales sur le langage peut être rapprochée des corrélations obtenues chez les enfants dans l'Etude 3a et variant en fonction du groupe. Il semble ainsi que les liens entre compétences auditives centrales et langagières ne soient pas figés et évoluent avec le temps. Spécifiquement, si les compétences auditives centrales semblent avoir un lien avec le niveau de vocabulaire au cours du développement, ce n'est plus le cas à l'âge adulte, même chez les participants dyslexiques (voir aussi Brosnan et al., 2002). Ceci pourrait ainsi refléter une certaine abstraction des concepts sémantiques (reflétés par le niveau de vocabulaire) qui deviennent indépendants.

Des variations dans la force des liens entre perception et cognition seraient en accord avec les études sur la compréhension de la parole et le vieillissement. Ainsi, il a été mis en évidence, chez des personnes âgées présentant un déficit auditif, une modification des circuits neuronaux responsables du traitement de la parole (Peelle, Troiani, Grossman, & Wingfield, 2011). Plus précisément, les auteurs ont pu mettre en évidence que l'acuité auditive permettait de prédire non seulement la réponse neuronale à la présentation de stimuli de parole mais également la quantité de matière grise dans les régions auditives primaires.

Afin d'affiner cette hypothèse, il serait également intéressant de tester une population d'enfants dyslexiques. Etant donnée l'immaturité des traitements auditifs centraux, et le problème de l'impact des facteurs non sensoriels, nous avons préféré dans un premier temps tester une population d'adultes dyslexiques. Ceci nous permet d'isoler spécifiquement le lien entre dyslexie et TTA en limitant l'impact d'autres facteurs, comme l'immaturité des traitements auditifs centraux. Toutefois, de futurs travaux testant des enfants dyslexiques permettraient de déterminer les courbes développementales de cette population. Ainsi, il serait possible d'évaluer si les déficits observés chez les adultes proviennent d'une immaturité, d'une déviance dès le début du développement ou d'une déviance au cours du développement. L'hypothèse d'un déficit dès le début du développement semble privilégiée puisque les études s'étant penchées sur la question suggèrent qu'il existe des déficits dans les traitements auditifs très tôt chez les sujets à risque pour la dyslexie (Boets et al., 2007; Hornickel et al., 2011). Le test d'une population d'adolescents permettrait également de compléter les trajectoires développementales observées.

Chapitre XIII

Discussion Générale

Cette thèse s'intéresse à l'impact de la dégradation du signal de parole sur le traitement et les représentations du langage, plus généralement elle vise à étudier les liens entre la perception et la cognition.

Le travail expérimental, articulé autour de deux axes a donné lieu à 4 études. Deux études se sont intéressées à la profondeur du traitement sémantique en situation de parole dans le bruit, et les deux autres ont visé à s'intéresser aux compétences auditives centrales, et à leurs liens avec les représentations langagières. L'idée générale sous-tendant ce travail est qu'il existe un lien direct entre la qualité de la perception du signal parolier et le langage, soit en termes de compréhension de celui-ci soit en termes de ses représentations, selon la nature (transitoire ou continue) et le degré de la dégradation.

1. Récapitulatif

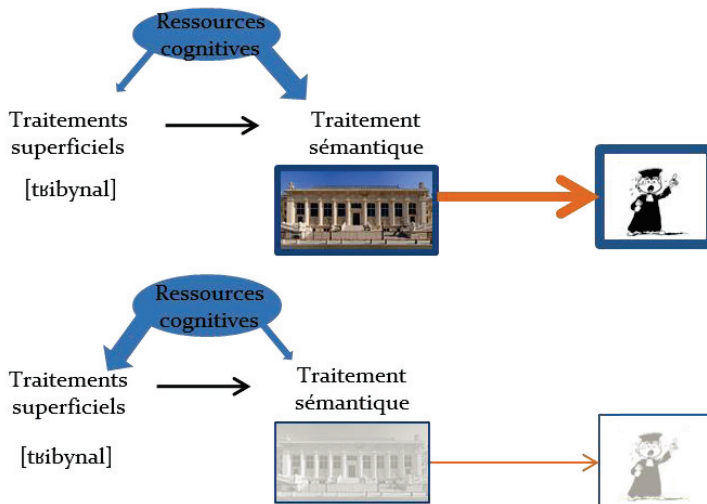
Le premier axe a permis de mettre en évidence que le traitement sémantique au niveau du mot et de la phrase était impacté par la présence de bruit en dépit d'une intelligibilité préservée. Alors que dans le silence, le traitement du contexte était suffisamment profond pour mener à des prédictions précises sur la fin de la phrase, la présence de bruit diminuant légèrement la qualité du signal de parole, gêne le traitement sémantique du contexte et entraîne des prédictions plus larges. Enfin, lorsque le fond sonore était constitué d'un bruit parolier composé de 4 voix, le traitement sémantique n'était plus mesurable aussi bien au niveau du mot que de la phrase et ce, malgré une intelligibilité préservée.

La seconde partie du travail expérimental s'est intéressée à la dégradation permanente du signal de parole du fait de mauvais traitements auditifs centraux, et à l'impact de ces traitements de mauvaise qualité sur le langage. Nous avons montré que les

compétences auditives centrales évoluaient au cours du développement, de l'école primaire à l'âge adulte. Il est apparu que les liens entre ces compétences et les représentations langagières n'étaient pas constants au cours du développement et étaient très limités à l'âge adulte. En revanche, au sein de la population dyslexique même adulte, les représentations phonologiques semblent encore très dépendantes de la qualité des traitements auditifs centraux (cf. Figure 47)

A. Adultes normo-lecteurs

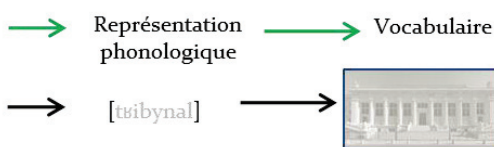
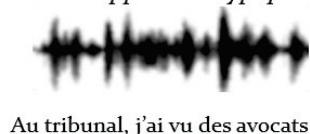
Silence



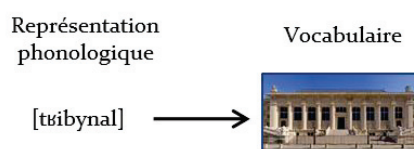
Bruit Stationnaire



B. Enfants développement typique



C. Adultes normo-lecteurs



D. Adultes dyslexiques

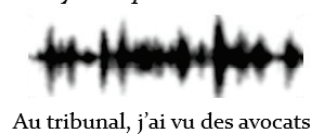


Figure 47

Représentation graphique des résultats. A. Illustration de l'effet de la dégradation du signal de parole par un fond sonore sur le traitement sémantique au sein d'une population d'adultes normo-lecteurs. B. Illustration de l'effet de la dégradation permanente du signal de parole du fait de l'immaturité des traitements auditifs centraux sur les représentations langagières chez des enfants au développement typique. C. Illustration du traitement du signal de parole non dégradé chez l'adulte normo-lecteur. D. Illustration de l'effet de la dégradation permanente du signal de parole du fait de TTA sur les représentations du langage au sein d'une population de dyslexiques adultes. Les flèches oranges représentent les effets d'amorçage, les flèches noires représentent l'enchaînement des différents niveaux

de traitement et les flèches vertes représentent les influences de la qualité du signal sur les différents niveaux de représentation langagière.

Cette thèse permet donc de montrer différents types de liens entre la dégradation du signal de parole et les traitements et représentations du langage. Lorsque la dégradation du signal est transitoire, comme dans la parole dans le bruit, celle-ci provoque une réallocation des ressources cognitives sur les traitements de bas niveau. De ce fait, peu de ressources sont disponibles pour les traitements de plus haut niveau notamment le traitement sémantique. Lorsque la dégradation est permanente, elle impacte les bas et hauts niveaux de représentation langagière chez les enfants normo-lecteurs. Chez les adultes dyslexiques en revanche, elle n'influe que sur la qualité des représentations bas niveau, les représentations de haut niveau semblant être indépendantes de la qualité du percept.

2. Parole dans le bruit et dyslexie

Nous avons pu voir dans ce travail l'effet de différents types de dégradation du signal sur le langage. Une prochaine étape serait de s'intéresser à l'impact de la dégradation du signal sur le traitement sémantique au sein d'une population dyslexique. En effet, si nous avons vu que la présence de TTA ne semble pas influencer sur le niveau de vocabulaire des dyslexiques, il serait intéressant d'évaluer si la diminution de qualité du signal génère une difficulté supplémentaire au traitement du sens du signal. Il a largement été montré que les enfants dyslexiques (Bradlow et al., 2003; Calcus et al., 2015; Ziegler et al., 2009) ainsi que les adultes (Dole et al., 2012) présentent une moins bonne intelligibilité en situation de parole dans le bruit.

Pourtant, l'impact de cette intelligibilité diminuée sur la qualité des traitements linguistiques de haut niveau et particulièrement sémantique n'a pas été investiguée. Calcus et al. (2015) ont étudié l'effet du bruit sur la reconnaissance de syllabe [CV] en fonction des caractéristiques phonétiques des consonnes. Ils ont mis en évidence, sur trois types de bruit (i.e., bruit parolier, bruit fluctuant ou bruit stationnaire) que les enfants dyslexiques âgés de 11 ans étaient moins performants que les enfants contrôles de même âge mais pas que les enfants contrôles de même niveau de lecture. Aucune différence significative n'a été mise en évidence pour aucun des groupes entre bruit fluctuant et bruit stationnaire, suggérant que les enfants testés ne bénéficiaient pas des variations temporelles du signal (cf. p61). Les résultats ont montré que quel que soit le bruit dans lequel la syllabe était présentée le lieu d'articulation était mieux identifié que le voisement et mode d'articulation. Pourtant, il ne semble pas, que les dyslexiques aient présenté un pattern différent que les contrôles (aussi bien en âge qu'en niveau de lecture).

Ziegler et collègues (2009) en revanche, en étudiant l'identification de pseudo-mots [VCV] présentés dans du bruit fluctuant et dans du bruit stationnaire, n'ont pas observés les mêmes résultats. Les enfants dyslexiques testés obtenaient des moins bons

scores que les deux groupes contrôles (appariés en âge ou en niveau de lecture). Sur la moyenne des différents bruits, les participants dyslexiques étaient significativement moins performants à identifier la place d'articulation que les contrôles en âge et que les contrôles en niveau de lecture. Ils étaient également moins performants à reconnaître le mode d'articulation que les contrôles en âge. Encore une fois, les données de la littérature sur les déficits de perception auditive au sein de la population dyslexique sont variables. Nous pouvons toutefois noter encore une fois la petitesse de l'échantillon testé par Calcus et al., (2015 ; i.e., 10 enfants par groupe). Bien que la nature exacte du déficit ne soit donc pas connue pour l'instant, il semble que les dyslexiques présentent des difficultés à traiter la parole bruitée. De plus, selon les résultats de Ziegler et al. (2009), les performances de reconnaissance de phonèmes sont impactées. Ceci devrait alors majorer les difficultés à effectuer les traitements linguistiques de haut niveau dans le bruit.

3. Traitement sémantique et dyslexie

A l'heure actuelle, peu d'études se sont intéressées à l'efficacité des traitements linguistiques, à l'exception du traitement phonologique, dans la population dyslexique. En effet, puisque les dyslexiques présentent une bonne intelligibilité de la parole dans le silence, il a été considéré que les traitements effectués sur celle-ci étaient efficaces. Certaines équipes de recherches ont toutefois étudié la question du traitement sémantique chez les dyslexiques mais toujours dans des conditions de perception optimales.

Ainsi, dans une étude étudiant le potentiel évoqué N400, Jednorog, Marchewka, Tacikowski et Grabowska (2010) se sont intéressés au traitement sémantique sur une population d'enfants dyslexiques (âgés de 10 ans). Les participants devaient écouter une série de 6 amorces, sémantiquement liées les unes aux autres, et suivies par un mot cible pouvant être ou non lié aux mots amorces. La réponse électrophysiologique à ce mot cible était enregistrée. En comparant les tracés EEG à la fois pour la condition Congruent et la condition Incongruent, les auteurs visaient à déterminer si un éventuel déficit au niveau des dyslexiques provenait d'une difficulté de traitement d'un stimulus congruent ou incongruent. Les résultats ont mis en évidence une N400 plus tardive en réponse au mot cible en condition Incongruent chez les dyslexiques comparativement aux contrôles. Ces résultats suggèrent donc que les dyslexiques seraient plus longs à intégrer un mot sémantiquement incongruent dans son contexte que les contrôles. Les auteurs suggèrent que ce délai serait dû à une difficulté à détecter la déviance. Il se pourrait que ce délai provienne de concepts sémantiques moins bien spécifiés, les participants dyslexiques pourraient bénéficier des effets d'amorçage mais être plus lent à déterminer qu'un concept n'appartient pas au même réseau sémantique que le contexte.

Toutefois, d'autres auteurs n'ont pas mis en évidence de différence ni dans l'amplitude ni dans le délai de la N400 générée par un mot incohérent au sein d'une phrase, chez des enfants dyslexiques par rapport à leurs contrôles (Sabisch, Hahne, Glass, von Suchodoletz, & Friederici, 2006).

D'autres études se sont intéressées au traitement sémantique en modalité visuelle. Ainsi, un délai dans l'apparition de la N400 a également été mis en évidence chez des adultes dyslexiques en MEG, lorsque les stimuli étaient des phrases pouvant se terminer par un mot attendu, par un mot non attendu mais cohérent ou par un mot non attendu et non cohérent (Helenius, Salmelin, Service, & Connolly, 1999). Ce délai a également été retrouvé avec un paradigme similaire, par Robichon, Besson et Habib, (2002) en plus d'une N400 d'amplitude plus importante aussi bien lorsque le mot cible était congruent qu'incongruent dans la phrase. Dans une tâche de jugement sémantique toutefois, Russeler et collègues (Russeler, Becker, Johannes, & Munte, 2007) ne mettent pas en évidence de différence entre l'effet N400 observé chez les sujets contrôles que chez les dyslexiques. Toutefois ces études se déroulent en modalité visuelle il apparaît ainsi délicat de distinguer l'impact des difficultés de lecture des participants de difficultés purement sémantiques. En effet, il semble évident que les dyslexiques seront plus lents à traiter le langage écrit que les contrôles, particulièrement lorsque leur prédiction n'est pas vérifiée, en situation d'incongruence sémantique.

Malheureusement, aucune des études ci-dessus, ne mesure le niveau de vocabulaire des participants. Il se pourrait pourtant que le déficit observé dans certaines d'entre elles provienne seulement de moins bonnes représentations de mémoire sémantique (e.g., moins précises, moins nombreuses). Les participants dyslexiques auront de ce fait besoin de plus longtemps pour se rendre compte que le mot n'est pas approprié dans le contexte. Bien que nous ne mettions pas en évidence dans l'Etude 4 de tel déficit en vocabulaire, il semble important de tester ces compétences. En effet, les déficits en conscience phonologique pourrait facilement engendrer des déficits en vocabulaire, particulièrement chez les enfants (i.e., Matthew Effect ; Stanovich, 1986a).

Evaluer sur une population dyslexique l'impact de la dégradation du signal sonore par du bruit sur la profondeur du traitement sémantique permettrait de s'intéresser à l'effet cumulé de la dégradation transitoire et continue du signal de parole. De plus, d'un point de vue plus appliqué, elle permettrait de mettre en évidence, si elles existent d'autres difficultés de la population dyslexique. En effet, selon l'*Effortfulness Hypothesis*, les difficultés dans le traitement sémantique chez les dyslexiques devraient être majorées comparativement aux contrôles, puisque avec le même bruit, l'intelligibilité en condition de parole dans le bruit est moins bonne chez les dyslexiques, vraisemblablement du fait de la présence de TTA.

Conclusion

Cette thèse s'intéressait à l'impact de la dégradation du signal de parole sur le traitement et les représentations du langage. D'une part, nous nous sommes intéressés à l'effet de la dégradation transitoire du signal de parole sur la profondeur du traitement sémantique, c'est-à-dire sur sa compréhension *per se*. Une première Etude comportementale, adaptant un paradigme d'amorçage sémantique à la situation de parole dans le bruit a permis de mettre en évidence que le traitement sémantique est dépendant de la saillance de l'amorce. De plus, lorsque 3 et 4 voix étaient présentes dans le fond sonore, le traitement sémantique n'était plus mesurable malgré une intelligibilité préservée. Ces résultats ont été interprétés à la lumière de l'*Effortfulness Hypothesis*, suggérant que le traitement sémantique était dépendant des ressources cognitives et non totalement automatique. Une seconde Etude en EEG, s'intéressant au traitement sémantique des mots ambigus présentés dans un contexte phrastique a confirmé que l'intelligibilité était nécessaire mais pas suffisante au traitement sémantique. Cette seconde Etude a également montré que lorsque la phrase est présentée dans le silence, le traitement sémantique de celle-ci est suffisamment profond pour mener à de fortes prédictions. En revanche, lorsque le signal de parole est dégradé par la présence de bruit stationnaire, le traitement sémantique du contexte est moins profond, les prédictions sont donc moins précises.

D'autre part, cette thèse avait pour objectif de s'intéresser aux liens entre la dégradation du signal du fait de déficits des traitements auditifs centraux et les représentations langagières. Deux types de déficits ont été considérés ; tout d'abord l'immaturité. De ce fait l'Etude 3a s'est intéressée aux liens entre le développement de certains de ces traitements auditifs et le langage chez des enfants d'école primaire. Une ré-analyse des données, a permis d'étudier le développement de l'ensemble des traitements auditifs centraux, comme ils sont définis par l'ASHA (2005) et de les comparer à la littérature sur le sujet (Etude 3b). L'Etude 3 a ainsi permis de montrer que les compétences auditives centrales évoluent avec l'âge et se développent jusqu'à l'âge adulte. De plus, les liens entre compétences auditives centrales et représentations

langagières ne semblent pas non plus être stables dans le temps et évoluent avec le développement. Une dernière étude a visé à s'intéresser aux liens entre compétences auditives centrales et dyslexie, chez l'adulte. Elle a montré que les participants dyslexiques étaient, au niveau du groupe, déficitaires aux tâches impliquant des traitements temporel et spectral. De plus, des corrélations entre ces compétences auditives centrales et les représentations langagières sont apparues significatives uniquement au sein de cette population. Ces résultats suggèrent donc que les corrélations observées ne reflètent pas simplement un effet de groupe mais des liens de cause à effet. De façon générale les résultats de l'Axe 2 suggèrent que les liens entre compétences auditives centrales et représentations langagières évoluent au cours du développement pour disparaître chez l'adulte normo-lecteur. En revanche, au sein de la population dyslexique il semble que les compétences phonologiques et le niveau de lecture soient toujours dépendants des compétences de traitement spectral et temporel.

Cette thèse a ainsi permis de s'intéresser aux liens entre l'intégrité du signal acoustique et le traitement et les représentations du langage. Bien que ces liens semblent évidents, ils restent très peu étudiés en particulier sur une population non pathologique. Ce travail permet donc d'affirmer que la dégradation du signal quelle que soit sa nature, impacte différents aspects du langage comme le traitement de son sens et ses représentations et présente donc un intérêt important pour la recherche appliquée. D'un point de vue plus fondamental, il permet de souligner l'intérêt d'utiliser la dégradation du signal pour étudier le fonctionnement normal du traitement du langage.

Références

- AAA (Producer). (2010). Diagnosis, treatment and management of children and adults with central auditory processing disorder. Retrieved from http://audiology-web.s3.amazonaws.com/migrated/CAPD%20Guidelines%208-2010.pdf_539952af956c79.73897613.pdf
- Abdi, H. (2007). Signal detection theory (SDT). *Encyclopedia of measurement and statistics*, 886-889.
- Abrams, R. L., & Greenwald, A. G. (2000). Parts outweigh the whole (word) in unconscious analysis of meaning. *Psychol Sci*, 11(2), 118-124.
- Abrams, R. L., & Grinspan, J. (2007). Unconscious semantic priming in the absence of partial awareness. *Conscious Cogn*, 16(4), 942-953; . doi: 10.1016/j.concog.2006.07.003
- Adlard, A., & Hazan, V. (1998). Speech perception in children with specific reading difficulties (dyslexia). *Q J Exp Psychol A*, 51(1), 153-177. doi: 10.1080/713755750
- Ahissar, M. (2007). Dyslexia and the anchoring-deficit hypothesis. *Trends Cogn Sci*, 11(11), 458-465. doi: 10.1016/j.tics.2007.08.015
- Ahissar, M., Lubin, Y., Putter-Katz, H., & Banai, K. (2006). Dyslexia and the failure to form a perceptual anchor. *Nat Neurosci*, 9(12), 1558-1564. doi: 10.1038/nm1800
- Ahissar, M., Protopapas, A., Reid, M., & Merzenich, M. M. (2000). Auditory processing parallels reading abilities in adults. *Proc Natl Acad Sci U S A*, 97(12), 6832-6837.
- Alario, F.-X., & Ferrand, L. (1998). Normes d'associations verbales pour 366 noms d'objets concrets. *L'année Psychologique*, 98(4), 659-709.
- Allen, K., Carlile, S., & Alais, D. (2008). Contributions of talker characteristics and spatial location to auditory streaming. [Research Support, Non-U.S. Gov't]. *J Acoust Soc Am*, 123(3), 1562-1570. doi: 10.1121/1.2831774
- Allen, P., & Wightman, F. (1994). Psychometric functions for children's detection of tones in noise. *J Speech Hear Res*, 37(1), 205-215.
- Altmann, C. F., Terada, S., Kashino, M., Goto, K., Mima, T., Fukuyama, H., & Furukawa, S. (2014). Independent or integrated processing of interaural time and level differences in human auditory cortex? *Hear Res*, 312, 121-127. doi: 10.1016/j.heares.2014.03.009
- Amitay, S., Ben-Yehudah, G., Banai, K., & Ahissar, M. (2002). Disabled readers suffer from visual and auditory impairments but not from a specific magnocellular deficit. *Brain*, 125(Pt 10), 2272-2285.
- Anderson, J. E., & Holcomb, P. J. (1995). Auditory and visual semantic priming using different stimulus onset asynchronies: an event-related brain potential study. *Psychophysiology*, 32(2), 177-190.
- ASHA (Producer). (2005). (Central) Auditory Processing Disorders - The role of the Audiologist [Position statement]. Retrieved from <http://www.asha.org/members/deskref-journals/deskref/default>
- Aydelott, J., Dick, F., & Mills, D. L. (2006). Effects of acoustic distortion and semantic context on event-related potentials to spoken words. *Psychophysiology*, 43(5), 454-464. doi: 10.1111/j.1469-8986.2006.00448.x
- Baddeley, A. (2000). The episodic buffer: a new component of working memory? *Trends Cogn Sci*, 4(11), 417-423.
- Baddeley, A., & Hitch, G. (1974). Working Memory. In G. H. Bower (Ed.), *The psychology of learning and motivation: Advances in research and theory* (Vol. 8, pp. 47-89). New York: Academic Press.

- Baldeweg, T., Richardson, A., Watkins, S., Foale, C., & Gruzelier, J. (1999). Impaired auditory frequency discrimination in dyslexia detected with mismatch evoked potentials. *Ann Neurol*, 45(4), 495-503.
- Baltes, P. B., & Lindenberger, U. (1997). Emergence of a powerful connection between sensory and cognitive functions across the adult life span: a new window to the study of cognitive aging? *Psychol Aging*, 12(1), 12-21.
- Banai, K., & Ahissar, M. (2006). Auditory processing deficits in dyslexia: task or stimulus related? *Cereb Cortex*, 16(12), 1718-1728. doi: 10.1093/cercor/bhj107
- Banai, K., & Ahissar, M. (2010). On the importance of anchoring and the consequences of its impairment in dyslexia. *Dyslexia*, 16(3), 240-257. doi: 10.1002/dys.407
- Banai, K., Sabin, A. T., & Wright, B. A. (2011). Separable developmental trajectories for the abilities to detect auditory amplitude and frequency modulation. *Hear Res*, 280(1-2), 219-227. doi: 10.1016/j.heares.2011.05.019
- Banai, K., & Yifat, R. (2011). Perceptual anchoring in preschool children: not adultlike, but there. *PLoS One*, 6(5), e19769. doi: 10.1371/journal.pone.0019769
- Banai, K., & Yuvak-Weiss, N. (2013). Prolonged development of auditory skills: A role for perceptual anchoring? *Cognitive Development*, 28, 300-311.
- Baran, J. A. (2014). Test battery principles and considerations *Handbook of central auditory processing disorders* (Vol. 1). San Diego: Plural Publishing.
- Barber, H., Vergara, M., & Carreiras, M. (2004). Syllable-frequency effects in visual word recognition: evidence from ERPs. *Neuroreport*, 15(3), 545-548.
- Bargones, J. Y., Werner, L. A., & Marean, G. C. (1995). Infant psychometric functions for detection: mechanisms of immature sensitivity. *J Acoust Soc Am*, 98(1), 99-111.
- Barker, B. A., & Newman, R. S. (2004). Listen to your mother! The role of talker familiarity in infant streaming. *Cognition*, 94(2), B45-53. doi: 10.1016/j.cognition.2004.06.001
- Beattie, R. L., & Manis, F. R. (2014). The relationship between prosodic perception, phonological awareness and vocabulary in emergent literacy. *Journal of Research in Reading*, 37(2), 119-137.
- Becker, C. A. (1980). Semantic context effects in visual word recognition: an analysis of semantic strategies. *Mem Cognit*, 8(6), 493-512.
- Bellis, T. J. (2003). *Assessment and management of central auditory processing disorders in the educational setting: From science to practice*. New York: Thomson Delmar Learning.
- Ben-Artzi, E., Fostick, L., & Babkoff, H. (2005). Deficits in temporal-order judgments in dyslexia: evidence from diotic stimuli differing spectrally and from dichotic stimuli differing only by perceived location. *Neuropsychologia*, 43(5), 714-723. doi: 10.1016/j.neuropsychologia.2004.08.004
- Ben-Yehudah, G., & Ahissar, M. (2004). Sequential spatial frequency discrimination is consistently impaired among adult dyslexics. *Vision Res*, 44(10), 1047-1063. doi: 10.1016/j.visres.2003.12.001
- Bennett, D. J., & McEvoy, C. L. (1999). Mediated priming in younger and older adults. *Exp Aging Res*, 25(2), 141-159. doi: 10.1080/036107399244066
- Bentin, S., Kutas, M., & Hillyard, S. A. (1993). Electrophysiological evidence for task effects on semantic priming in auditory word processing. *Psychophysiology*, 30(2), 161-169.

- Bentin, S., Kutas, M., & Hillyard, S. A. (1995). Semantic processing and memory for attended and unattended words in dichotic listening: behavioral and electrophysiological evidence. *J Exp Psychol Hum Percept Perform*, 21(1), 54-67.
- Bentin, S., McCarthy, G., & Wood, C. C. (1985). Event-related potentials, lexical decision and semantic priming. *Electroencephalogr Clin Neurophysiol*, 60(4), 343-355.
- Beretta, A., Fiorentino, R., & Poeppel, D. (2005). The effects of homonymy and polysemy on lexical access: an MEG study. *Brain Res Cogn Brain Res*, 24(1), 57-65. doi: 10.1016/j.cogbrainres.2004.12.006
- Berg, K. M., & Boswell, A. E. (2000). Noise increment detection in children 1 to 3 years of age. *Percept Psychophys*, 62(4), 868-873.
- Billiet, C. R., & Bellis, T. J. (2011). The relationship between brainstem temporal processing and performance on tests of central auditory function in children with reading disorders. *J Speech Lang Hear Res*, 54(1), 228-242. doi: 10.1044/1092-4388(2010/09-0239)
- Bishop, D. V. (2007). Using mismatch negativity to study central auditory processing in developmental language and literacy impairments: where are we, and where should we be going? *Psychol Bull*, 133(4), 651-672. doi: 10.1037/0033-2909.133.4.651
- Bishop, D. V., Hardiman, M. J., & Barry, J. G. (2012). Auditory deficit as a consequence rather than endophenotype of specific language impairment: electrophysiological evidence. *PLoS One*, 7(5), e35851. doi: 10.1371/journal.pone.0035851
- Bishop, D. V., & McArthur, G. M. (2005). Individual differences in auditory processing in specific language impairment: a follow-up study using event-related potentials and behavioural thresholds. *Cortex*, 41(3), 327-341.
- Bishop, D. V., & Snowling, M. J. (2004). Developmental dyslexia and specific language impairment: same or different? *Psychol Bull*, 130(6), 858-886. doi: 10.1037/0033-2909.130.6.858
- Boets, B., Ghesquiere, P., van Wieringen, A., & Wouters, J. (2007). Speech perception in preschoolers at family risk for dyslexia: relations with low-level auditory processing and phonological ability. *Brain Lang*, 101(1), 19-30. doi: 10.1016/j.bandl.2006.06.009
- Boets, B., Op de Beeck, H. P., Vandermosten, M., Scott, S. K., Gillebert, C. R., Mantini, D., . . . Ghesquiere, P. (2013). Intact but less accessible phonetic representations in adults with dyslexia. *Science*, 342(6163), 1251-1254. doi: 10.1126/science.1244333
- Boets, B., Vandermosten, M., Poelmans, H., Luts, H., Wouters, J., & Ghesquiere, P. (2011). Preschool impairments in auditory processing and speech perception uniquely predict future reading problems. *Res Dev Disabil*, 32(2), 560-570. doi: 10.1016/j.ridd.2010.12.020
- Bogliotti, C., Serniclaes, W., Messaoud-Galusi, S., & Sprenger-Charolles, L. (2008). Discrimination of speech sounds by children with dyslexia: comparisons with chronological age and reading level controls. *J Exp Child Psychol*, 101(2), 137-155. doi: 10.1016/j.jecp.2008.03.006
- Bosse, M. L., Tainturier, M. J., & Valdois, S. (2007). Developmental dyslexia: the visual attention span deficit hypothesis. *Cognition*, 104(2), 198-230. doi: 10.1016/j.cognition.2006.05.009

- Boulenger, V., Hoen, M., Ferragne, E., Pellegrino, F., & Meunier, F. (2010). Real-time lexical competitions during speech-in-speech comprehension. *Speech Communication*, 52(3), 246-253.
- Boulenger, V., Hoen, M., Jacquier, C., & Meunier, F. (2011). Interplay between acoustic/phonetic and semantic processes during spoken sentence comprehension: an ERP study. *Brain Lang*, 116(2), 51-63. doi: 10.1016/j.bandl.2010.09.011
- Bradlow, A. R., Kraus, N., & Hayes, E. (2003). Speaking clearly for children with learning disabilities: sentence perception in noise. *J Speech Lang Hear Res*, 46(1), 80-97.
- Bregman, A. S., & Campbell, J. (1971). Primary auditory stream segregation and perception of order in rapid sequences of tones. *J Exp Psychol*, 89(2), 244-249.
- Breier, J. I., Gray, L. C., Klaas, P., Fletcher, J. M., & Foorman, B. (2002). Dissociation of sensitivity and response bias in children with attention deficit/hyperactivity disorder during central auditory masking. *Neuropsychology*, 16(1), 28-34.
- Bronkhorst, A. W. (2000). The cocktail party phenomenon: a review of research on speech intelligibility in multiple-talker conditions. *Acustica*, 86, 117-128.
- Bronkhorst, A. W., & Plomp, R. (1988). The effect of head-induced interaural time and level differences on speech intelligibility in noise. *Journal of the Acoustical Society of America*, 83(4), 1508-1516.
- Brooke Shinn, J. (2014). Temporal Processing Tests. In F. E. Musiek & G. D. Chermak (Eds.), *Handbook of central auditory processing disorder* (Vol. 1, pp. 405-434). San Diego: Plural Publishing.
- Brosnan, M., Demetre, J., Hamill, S., Robson, K., Shepherd, H., & Cody, G. (2002). Executive functioning in adults and children with developmental dyslexia. *Neuropsychologia*, 40(12), 2144-2155.
- Brouwer, S., Van Engen, K., Calandruccio, L., & Bradlow, A. (2012). Linguistic contributions to speech-on-speech masking for native and non-native listeners: Language familiarity and semantic content. *Journal of the Acoustical Society of America*, 131(2), 1449-1463.
- Brown, C. R., & Hagoort, P. (1993). The processing nature of the N400: Evidence from Masked Priming. *Journal of Cognitive Neuroscience*, 5(1), 34-44.
- Brungart, D. S. (2001). Informational and energetic masking effects in the perception of two simultaneous talkers. *J Acoust Soc Am*, 109(3), 1101-1109.
- Brungart, D. S., Chang, P. S., Simpson, B. D., & Wang, D. (2006). Isolating the energetic component of speech-on-speech masking with ideal time-frequency segregation. *J Acoust Soc Am*, 120(6), 4007-4018.
- Brungart, D. S., Simpson, B. D., Ericson, M. A., & Scott, K. R. (2001). Informational and energetic masking effects in the perception of multiple simultaneous talkers. *J Acoust Soc Am*, 110(5 Pt 1), 2527-2538.
- BSA (Producer). (2011). An overview of current management of auditory processing disorder (APD). Retrieved from <http://www.thebsa.org.uk/wp-content/uploads/2014/04/BSA-APD-Management-1Aug11-FINAL-amended17Oct11.pdf>
- Burke, D. M., White, H., & Diaz, D. L. (1987). Semantic priming in young and older adults: evidence for age constancy in automatic and attentional processes. *J Exp Psychol Hum Percept Perform*, 13(1), 79-88.

- Buss, E., Hall, J. W., 3rd, & Grose, J. H. (2006). Development and the role of internal noise in detection and discrimination thresholds with narrow band stimuli. *J Acoust Soc Am*, 120(5 Pt 1), 2777-2788.
- Buss, E., Hall, J. W., 3rd, Porter, H., & Grose, J. H. (2014). Gap detection in school-age children and adults: effects of inherent envelope modulation and the availability of cues across frequency. *J Speech Lang Hear Res*, 57(3), 1098-1107. doi: 10.1044/2014_JSLHR-H-13-0132
- Buss, E., Whittle, L. N., Grose, J. H., & Hall, J. W., 3rd. (2009). Masking release for words in amplitude-modulated noise as a function of modulation rate and task. *J Acoust Soc Am*, 126(1), 269-280. doi: 10.1121/1.3129506
- Cacace, A. T., & McFarland, D. J. (2005). The importance of modality specificity in diagnosing central auditory processing disorder. *Am J Audiol*, 14(2), 112-123. doi: 10.1044/1059-0889(2005/012)
- Calcus, A., Colin, C., Deltenre, P., & Kolinsky, R. (2015). Informational masking of speech in dyslexic children. *J Acoust Soc Am*, 137(6), EL496-502. doi: 10.1121/1.4922012
- Carey, D., Mercure, E., Pizzioli, F., & Aydelott, J. (2014). Auditory semantic processing in dichotic listening: effects of competing speech, ear of presentation, and sentential bias on N400s to spoken words in context. *Neuropsychologia*, 65, 102-112. doi: 10.1016/j.neuropsychologia.2014.10.016
- Carr, T. H., & Dagenbach, D. (1990). Semantic priming and repetition priming from masked words: evidence for a center-surround attentional mechanism in perceptual recognition. *J Exp Psychol Learn Mem Cogn*, 16(2), 341-350.
- Castles, A., & Coltheart, M. (1993). Varieties of developmental dyslexia. *Cognition*, 47(2), 149-180.
- Catts, H. W., Adlof, S. M., Hogan, T. P., & Weismer, S. E. (2005). Are specific language impairment and dyslexia distinct disorders? *J Speech Lang Hear Res*, 48(6), 1378-1396. doi: 10.1044/1092-4388(2005/096)
- CHABA. (1988). Speech understanding and aging. *Journal of the Acoustical Society of America*, 83(3), 859-895.
- Cherry, C. (1953). Some experiments on the recognition of speech, with one and with two ears. *The Journal of the Acoustical Society of America*, 25(5), 975-979.
- Chonchaiya, W., Tardif, T., Mai, X., Xu, L., Li, M., Kaciroti, N., . . . Lozoff, B. (2013). Developmental trends in auditory processing can provide early predictions of language acquisition in young infants. *Dev Sci*, 16(2), 159-172. doi: 10.1111/desc.12012
- Cohen-Mimran, R., & Sapir, S. (2007). Deficits in working memory in young adults with reading disabilities. *J Commun Disord*, 40(2), 168-183. doi: 10.1016/j.jcomdis.2006.06.006
- Cola-Asmussen, C., Lequette, C., Pouget, G., Rouyet, C., & Zorman, M. (2011). ECLA 16+. http://www.cognisciences.com/article.php3?id_article=86.
- Colbert-Getz, J., & Cook, A. E. (2013). Revisiting effects of contextual strength on the subordinate bias effect: evidence from eye movements. *Mem Cognit*, 41(8), 1172-1184. doi: 10.3758/s13421-013-0328-3
- Collins, A., & Loftus, E. F. (1975). A spreading-activation theory of semantic processing. *Psychological Review*, 82(6), 407-428.

- Collins, A., & Quillian, M. R. (1969). Retrieval time from semantic memory. *Journal of verbal learning and verbal behavior*, 9(4), 240-247.
- Connolly, J. F., & Phillips, N. A. (1994). Event-related potential components reflect phonological and semantic processing of the terminal word of spoken sentences. *Journal of Cognitive Neuroscience*, 6(3), 256-266.
- Connolly, J. F., Phillips, N. A., Stewart, S. H., & Brake, W. G. (1992). Event-related potential sensitivity to acoustic and semantic properties of terminal words in sentences. *Brain Lang*, 43(1), 1-18.
- Connolly, J. F., Stewart, S. H., & Phillips, N. A. (1990). The effects of processing requirements on neurophysiological responses to spoken sentences. *Brain Lang*, 39(2), 302-318.
- Cooke, M., Garcia Lecumberri, M. L., & Barker, J. (2008). The foreign language cocktail party problem: Energetic and informational masking effects in non-native speech perception. *J Acoust Soc Am*, 123(1), 414-427. doi: 10.1121/1.2804952
- Dagenbach, D., Horst, S., & Carr, T. H. (1990). Adding new information to semantic memory: how much learning is enough to produce automatic priming? *J Exp Psychol Learn Mem Cogn*, 16(4), 581-591.
- Daltrozzo, J., Signoret, C., Tillmann, B., & Perrin, F. (2011). Subliminal semantic priming in speech. *PLoS One*, 6(5), e20273. doi: 10.1371/journal.pone.0020273
- Damian, M. F. (2001). Congruity effects evoked by subliminally presented primes: automaticity rather than semantic processing. *J Exp Psychol Hum Percept Perform*, 27(1), 154-165.
- Darwin, C. J., Brungart, D. S., & Simpson, B. D. (2003). Effects of fundamental frequency and vocal-tract length changes on attention to one of two simultaneous talkers. *J Acoust Soc Am*, 114(5), 2913-2922.
- Dawes, P., & Bishop, D. V. (2008). Maturation of visual and auditory temporal processing in school-aged children. *J Speech Lang Hear Res*, 51(4), 1002-1015. doi: 10.1044/1092-4388(2008/073)
- Dawes, P., Bishop, D. V., Sirimanna, T., & Bamiou, D. E. (2008). Profile and aetiology of children diagnosed with auditory processing disorder (APD). [Research Support, Non-U.S. Gov't]. *Int J Pediatr Otorhinolaryngol*, 72(4), 483-489. doi: 10.1016/j.ijporl.2007.12.007
- de Wit, B., & Kinoshita, S. (2014). Relatedness proportion effects in semantic categorization: reconsidering the automatic spreading activation process. *J Exp Psychol Learn Mem Cogn*, 40(6), 1733-1744. doi: 10.1037/xlm0000004
- de Wit, B., & Kinoshita, S. (2015a). The masked semantic priming effect is task dependent: Reconsidering the automatic spreading activation process. *J Exp Psychol Learn Mem Cogn*, 41(4), 1062-1075. doi: 10.1037/xlm0000074
- de Wit, B., & Kinoshita, S. (2015b). An RT distribution analysis of relatedness proportion effects in lexical decision and semantic categorization reveals different mechanisms. *Mem Cognit*, 43(1), 99-110. doi: 10.3758/s13421-014-0446-6
- Deacon, D., Hewitt, S., Yang, C., & Nagata, M. (2000). Event-related potential indices of semantic priming using masked and unmasked words: evidence that the N400 does not reflect a post-lexical process. *Brain Res Cogn Brain Res*, 9(2), 137-146.
- Dehaene, S., & Naccache, L. (2001). Towards a cognitive neuroscience of consciousness: basic evidence and a workspace framework. *Cognition*, 79(1-2), 1-37.

- Dehaene, S., Naccache, L., Le Clec, H. G., Koechlin, E., Mueller, M., Dehaene-Lambertz, G., . . . Le Bihan, D. (1998). Imaging unconscious semantic priming. *Nature*, 395(6702), 597-600. doi: 10.1038/26967
- Dekerle, M., Boulenger, V., Hoen, M., & Meunier, F. (2014). Multi-talker background and semantic priming effect. *Front Hum Neurosci*, 8, 878. doi: 10.3389/fnhum.2014.00878
- Dell'Acqua, R., & Grainger, J. (1999). Unconscious semantic priming from pictures. *Cognition*, 73(1), B1-B15.
- Demanez, L., & Demanez, J. P. (2003). Central auditory processing assessment. *Acta Otorhinolaryngol Belg*, 57(4), 243-252.
- Demanez, L., Dony-Closon, B., Lhonneux-Ledoux, E., & Demanez, J. P. (2003). Central auditory processing assessment: a French-speaking battery. *Acta Otorhinolaryngol Belg*, 57(4), 275-290.
- Demonet, J. F., Taylor, M. J., & Chaix, Y. (2004). Developmental dyslexia. *Lancet*, 363(9419), 1451-1460. doi: 10.1016/S0140-6736(04)16106-0
- Denckla, M. B., & Rudel, R. G. (1976). Rapid "automatized" naming (R.A.N): dyslexia differentiated from other learning disabilities. *Neuropsychologia*, 14(4), 471-479.
- Desroches, A. S., Newman, R. L., & Joanisse, M. F. (2009). Investigating the time course of spoken word recognition: electrophysiological evidence for the influences of phonological similarity. *J Cogn Neurosci*, 21(10), 1893-1906. doi: 10.1162/jocn.2008.21142
- Dole, M., Hoen, M., & Meunier, F. (2012). Speech-in-noise perception deficit in adults with dyslexia: effects of background type and listening configuration. *Neuropsychologia*, 50(7), 1543-1552. doi: 10.1016/j.neuropsychologia.2012.03.007
- Dole, M., Meunier, F., & Hoen, M. (2013). Gray and white matter distribution in dyslexia: a VBM study of superior temporal gyrus asymmetry. *PLoS One*, 8(10), e76823. doi: 10.1371/journal.pone.0076823
- Donnadieu, S., Gillet-Perret, E., Lassus-Sangosse, D., & Nguyen-Morel, M.-A. (2014). Evaluation et implications des Troubles Auditifs Centraux (TAC) dans la dysphasie : Etude de cas. *Parole*.
- Donnenwerth-Nolan, S., Tanenhaus, M. K., & Seidenberg, M. S. (1981). Multiple code activation in word recognition: evidence from rhyme monitoring. *J Exp Psychol Hum Learn*, 7(3), 170-180.
- Draganova, R., Eswaran, H., Murphy, P., Huotilainen, M., Lowery, C., & Preissl, H. (2005). Sound frequency change detection in fetuses and newborns, a magnetoencephalographic study. *Neuroimage*, 28(2), 354-361. doi: 10.1016/j.neuroimage.2005.06.011
- Draganova, R., Eswaran, H., Murphy, P., Lowery, C., & Preissl, H. (2007). Serial magnetoencephalographic study of fetal and newborn auditory discriminative evoked responses. *Early Hum Dev*, 83(3), 199-207. doi: 10.1016/j.earlhumdev.2006.05.018
- Drullman, R., & Bronkhorst, A. W. (2000). Multichannel speech intelligibility and talker recognition using monaural, binaural, and three-dimensional auditory presentation. *J Acoust Soc Am*, 107(4), 2224-2235.
- Duff, F. J., Hayiou-Thomas, M. E., & Hulme, C. (2012). Evaluating the effectiveness of a phonologically based reading intervention for struggling readers with varying language profiles. *Reading and Writing*, 621-640.

- Duffy, S. A., Morris, R. K., & Rayner, K. (1988). Lexical ambiguity and fixation times in reading. *Journal of Memory and Language*, 27, 429-446.
- Dufour, S., Brunellière, A., & Frauenfelder, U. H. (2013). Tracking the time course of word-frequency effects in auditory word recognition with event-related potentials. *Cognitive Science*, 34, 489-507.
- Dunn, L. M., Theriault-Whalen, C. M., & Dunn, L. M. (1993). *Echelle de vocabulaire en images Peabody. Adaptation française du Peabody Picture Vocabulary test-revised*.
- Dupoux, E., Kouider, S., & Mehler, J. (2003). Lexical access without attention? Explorations using dichotic priming. *J Exp Psychol Hum Percept Perform*, 29(1), 172-184.
- Durlach, N. (2006). Auditory masking: need for improved conceptual structure. *J Acoust Soc Am*, 120(4), 1787-1790.
- Eckstein, D., & Perrig, W. J. (2007). The influence of intention on masked priming: a study with semantic classification of words. *Cognition*, 104(2), 345-376. doi: 10.1016/j.cognition.2006.07.005
- Eggermont, J. J. (1988). On the rate of maturation of sensory evoked potentials. *Electroencephalogr Clin Neurophysiol*, 70(4), 293-305.
- Eggermont, J. J. (2014). Development of the central auditory nervous system. In F. E. Musiek & G. D. Chermak (Eds.), *Handbook of central auditory processing disorder* (Vol. 1, pp. 59-88). San Diego: Plural Publishing.
- Eggermont, J. J., & Ponton, C. W. (2003). Auditory-evoked potential studies of cortical maturation in normal hearing and implanted children: correlations with changes in structure and speech perception. *Acta Otolaryngol*, 123(2), 249-252.
- Eich, E. (1984). Memory for unattended events: remembering with and without awareness. *Mem Cognit*, 12(2), 105-111.
- Elliott, J. G., & Grigorenko, E. L. (2014). *The dyslexia debate*. New York: Cambridge University Press.
- Fawcett, A. J., & Nicolson, R. I. (1999). Performance of Dyslexic Children on Cerebellar and Cognitive Tests. *J Mot Behav*, 31(1), 68-78. doi: 10.1080/00222899909601892
- Federmeier, K. D., & Kutas, M. (2001). Meaning and modality: influences of context, semantic memory organization, and perceptual predictability on picture processing. *J Exp Psychol Learn Mem Cogn*, 27(1), 202-224.
- Ferguson, M. A., Hall, R. L., Riley, A., & Moore, D. R. (2011). Communication, listening, cognitive and speech perception skills in children with auditory processing disorder (APD) or Specific Language Impairment (SLI). *J Speech Lang Hear Res*, 54(1), 211-227. doi: 10.1044/1092-4388(2010/09-0167)
- Ferrand, L. (1999). 640 homophones et leurs caractéristiques. *L'année Psychologique*, 99(4), 687-708.
- Ferrand, L. (2001). Normes d'associations verbales pour 260 mots "abstraits". *L'année Psychologique*, 101(4), 683-721.
- Ferrand, L., & Grainger, J. (1996). List context effects on masked phonological priming in the lexical decision task. *Psychon Bull Rev*, 3(4), 515-519. doi: 10.3758/BF03214557
- Festen, J. M., & Plomp, R. (1990). Effects of fluctuating noise and interfering speech on the speech-reception threshold for impaired and normal hearing. *J Acoust Soc Am*, 88(4), 1725-1736.

- Fischer, B., & Hartnegg, K. (2004). On the development of low-level auditory discrimination and deficits in dyslexia. *Dyslexia*, 10(2), 105-118. doi: 10.1002/dys.268
- Fitzpatrick, D. (2004). The Auditory System. In D. Purves (Ed.), *Neuroscience*. Sunderland, MA: Sinauer Associates.
- Fletcher, J. M. (2009). Dyslexia: The evolution of a scientific concept. *J Int Neuropsychol Soc*, 15(4), 501-508. doi: 10.1017/S1355617709090900
- Fluss, J., Ziegler, J. C., Warszawski, J., Ducot, B., Richard, G., & Billard, C. (2009). Poor reading in French elementary school: the interplay of cognitive, behavioral, and socioeconomic factors. *J Dev Behav Pediatr*, 30(3), 206-216. doi: 10.1097/DBP.0b013e3181a7ed6c
- Fodor, J. (1983). *The modularity of mind*. Cambridge, MA: MIT Press.
- Foraker, S., & Murphy, G. L. (2012). Polysemy in Sentence Comprehension: Effects of Meaning Dominance. *J Mem Lang*, 67(4), 407-425. doi: 10.1016/j.jml.2012.07.010
- Forster, K. I., & Davis, C. (1984). Repetition Priming and Frequency Attenuation in Lexical Access. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 10(4), 680-698.
- France, S. J., Rosner, B. S., Hansen, P. C., Calvin, C., Talcott, J. B., Richardson, A. J., & Stein, J. F. (2002). Auditory frequency discrimination in adult developmental dyslexics. *Percept Psychophys*, 64(2), 169-179.
- Frenck-Mestre, C., & Bueno, S. (1999). Semantic features and semantic categories: differences in rapid activation of the lexicon. *Brain Lang*, 68(1-2), 199-204. doi: 10.1006/brln.1999.2079
- Freyman, R. L., Balakrishnan, U., & Helfer, K. S. (2004). Effect of number of masking talkers and auditory priming on informational masking in speech recognition. *J Acoust Soc Am*, 115(5 Pt 1), 2246-2256.
- Freyman, R. L., Helfer, K. S., McCall, D. D., & Clifton, R. K. (1999). The role of perceived spatial separation in the unmasking of speech. *J Acoust Soc Am*, 106(6), 3578-3588.
- Ganong, W. F., 3rd. (1980). Phonetic categorization in auditory word perception. *J Exp Psychol Hum Percept Perform*, 6(1), 110-125.
- Gao, X., Levinthal, B. R., & Stine-Morrow, E. A. (2012). The effects of ageing and visual noise on conceptual integration during sentence reading. *Q J Exp Psychol (Hove)*, 65(9), 1833-1847. doi: 10.1080/17470218.2012.674146
- Gao, X., Stine-Morrow, E. A., Noh, S. R., & Eskew, R. T., Jr. (2011). Visual noise disrupts conceptual integration in reading. *Psychon Bull Rev*, 18(1), 83-88. doi: 10.3758/s13423-010-0014-4
- Gautreau, A., Hoen, M., & Meunier, F. (2013). Let's all speak together! Exploring the masking effects of various languages on spoken word identification in multi-linguistic babble. *PLoS One*, 8(6), e65668. doi: 10.1371/journal.pone.0065668
- Geiger, G., Cattaneo, C., Galli, R., Pozzoli, U., Lorusso, M. L., Facoetti, A., & Molteni, M. (2008). Wide and diffuse perceptual modes characterize dyslexics in vision and audition. *Perception*, 37(11), 1745-1764.
- Glucksberg, S., Kreuz, R. J., & Rho, S. H. (1986). Context can constrain lexical access: Implications for models of language comprehension. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 12(3), 323-335.

- Goldman, R., & Fristoe, M. (1986). *Goldman, Fristoe test of articulation*. MN: American Guidance Service.
- Goodman, I., Libenson, A., & Wade-Wooley, L. (2010). Sensitivity to linguistic stress, phonological awareness and early reading ability in preschoolers. *Journal of Research in Reading*, 33(2).
- Goswami, U., Huss, M., Mead, N., Fosker, T., & Verney, J. P. (2013). Perception of patterns of musical beat distribution in phonological developmental dyslexia: significant longitudinal relations with word reading and reading comprehension. *Cortex*, 49(5), 1363-1376. doi: 10.1016/j.cortex.2012.05.005
- Goswami, U., Thomson, J., Richardson, U., Stainthorp, R., Hughes, D., Rosen, S., & Scott, S. K. (2002). Amplitude envelope onsets and developmental dyslexia: A new hypothesis. *Proc Natl Acad Sci U S A*, 99(16), 10911-10916. doi: 10.1073/pnas.122368599
- Goswami, U., Wang, H. L., Cruz, A., Fosker, T., Mead, N., & Huss, M. (2011). Language-universal sensory deficits in developmental dyslexia: English, Spanish, and Chinese. *J Cogn Neurosci*, 23(2), 325-337. doi: 10.1162/jocn.2010.21453
- Grainger, J., Diependaele, K., Spinelli, E., Ferrand, L., & Farioli, F. (2003). Masked repetition and phonological priming within and across modalities. *J Exp Psychol Learn Mem Cogn*, 29(6), 1256-1269. doi: 10.1037/0278-7393.29.6.1256
- Grainger, J., & Holcomb, P. J. (2009). Watching the Word Go by: On the Time-course of Component Processes in Visual Word Recognition. *Lang Linguist Compass*, 3(1), 128-156. doi: 10.1111/j.1749-818X.2008.00121.x
- Graven, S. N., & Browne, J. V. (2008). Auditory development in the fetus and infant. *Newborn & Infant Nursing Reviews*, 8(4), 187-193.
- Greenwald, A. G., Draine, S. C., & Abrams, R. L. (1996). Three cognitive markers of unconscious semantic activation. *Science*, 273(5282), 1699-1702.
- Greenwald, A. G., Klinger, M. R., & Liu, T. J. (1989). Unconscious processing of dichoptically masked words. *Mem Cognit*, 17(1), 35-47.
- Gunter, T. C., Friederici, A. D., & Schriefers, H. (2000). Syntactic gender and semantic expectancy: ERPs reveal early autonomy and late interaction. *J Cogn Neurosci*, 12(4), 556-568.
- Hagoort, P., & Brown, C. M. (2000). ERP effects of listening to speech: semantic ERP effects. *Neuropsychologia*, 38(11), 1518-1530.
- Halliday, L. F., Taylor, J. L., Edmondson-Jones, A. M., & Moore, D. R. (2008). Frequency discrimination learning in children. *J Acoust Soc Am*, 123(6), 4393-4402. doi: 10.1121/1.2890749
- Hamalainen, J. A., Salminen, H. K., & Leppanen, P. H. (2013). Basic auditory processing deficits in dyslexia: systematic review of the behavioral and event-related potential/ field evidence. *J Learn Disabil*, 46(5), 413-427. doi: 10.1177/0022219411436213
- Hari, R., & Renvall, H. (2001). Impaired processing of rapid stimulus sequences in dyslexia. *Trends Cogn Sci*, 5(12), 525-532.
- Hari, R., Saaskilahti, A., Helenius, P., & Uutela, K. (1999). Non-impaired auditory phase locking in dyslexic adults. *Neuroreport*, 10(11), 2347-2348.
- Hauk, O., & Pulvermüller, F. (2004). Effects of word length and frequency on the human event-related potential. *Clin Neurophysiol*, 115(5), 1090-1103. doi: 10.1016/j.clinph.2003.12.020

- Hawley, M. L., Litovsky, R. Y., & Culling, J. F. (2004). The benefit of binaural hearing in a cocktail party: effect of location and type of interferer. *J Acoust Soc Am*, 115(2), 833-843.
- Hayes, S. C., & Bissett, R. T. (1998). Derived stimulus relations produce mediated and episodic priming. *The psychological record*, 48, 617-630.
- He, C., Hotson, L., & Trainor, L. J. (2007). Mismatch responses to pitch changes in early infancy. *J Cogn Neurosci*, 19(5), 878-892. doi: 10.1162/jocn.2007.19.5.878
- Helenius, P., Salmelin, R., Service, E., & Connolly, J. F. (1999). Semantic cortical activation in dyslexic readers. *J Cogn Neurosci*, 11(5), 535-550.
- Helenius, P., Uutela, K., & Hari, R. (1999). Auditory stream segregation in dyslexic adults. *Brain*, 122 (Pt 5), 907-913.
- Hill, N. I., Bailey, P. J., Griffiths, Y. M., & Snowling, M. J. (1999). Frequency acuity and binaural masking release in dyslexic listeners. *J Acoust Soc Am*, 106(6), L53-58.
- Hill, P. R., Hogben, J. H., & Bishop, D. M. (2005). Auditory frequency discrimination in children with specific language impairment: a longitudinal study. *J Speech Lang Hear Res*, 48(5), 1136-1146. doi: 10.1044/1092-4388(2005/080)
- Hoen, M., Meunier, F., Grataloup, C.-L., Pellegrino, F., Grimault, N., Perrin, F., . . . Collet, L. (2007). Phonetic and lexical interferences in informational masking during speech-in-speech comprehension. *Speech Communication*, 49(12), 905-916.
- Hohlfeld, A., Sangals, J., & Sommer, W. (2004). Effects of additional tasks on language perception: an event-related brain potential investigation. *J Exp Psychol Learn Mem Cogn*, 30(5), 1012-1025. doi: 10.1037/0278-7393.30.5.1012
- Hohlfeld, A., & Sommer, W. (2005). Semantic processing of unattended meaning is modulated by additional task load: evidence from electrophysiology. *Brain Res Cogn Brain Res*, 24(3), 500-512. doi: 10.1016/j.cogbrainres.2005.03.001
- Holcomb, P. J., & Anderson, J. (1993). Cross-modal semantic priming: a time course analysis using event-related brain potentials. *Language and cognitive processes*, 8(4), 379-411.
- Holcomb, P. J., & Neville, H. J. (1990). Auditory and visual semantic priming in lexical decision: a comparison using event-related brain potentials. *Language and cognitive processes*, 5(4), 281-312.
- Holcomb, P. J., Reder, L., Misra, M., & Grainger, J. (2005). The effects of prime visibility on ERP measures of masked priming. *Brain Res Cogn Brain Res*, 24(1), 155-172. doi: 10.1016/j.cogbrainres.2005.01.003
- Holender, D. (1986). Semantic activation without conscious identification in dichotic listening, parafoveal vision and visual masking: a survey and appraisal. *Behavioural and Brain Sciences*, 9(1), 1-23.
- Hornickel, J., Anderson, S., Skoe, E., Yi, H. G., & Kraus, N. (2012). Subcortical representation of speech fine structure relates to reading ability. *Neuroreport*, 23(1), 6-9. doi: 10.1097/WNR.0b013e32834d2ffd
- Hornickel, J., Chandrasekaran, B., Zecker, S., & Kraus, N. (2011). Auditory brainstem measures predict reading and speech-in-noise perception in school-aged children. *Behav Brain Res*, 216(2), 597-605. doi: 10.1016/j.bbr.2010.08.051
- Hornickel, J., Knowles, E., & Kraus, N. (2012). Test-retest consistency of speech-evoked auditory brainstem responses in typically-developing children. *Hear Res*, 284(1-2), 52-58. doi: 10.1016/j.heares.2011.12.005

- Hornickel, J., & Kraus, N. (2013). Unstable representation of sound: a biological marker of dyslexia. *J Neurosci*, 33(8), 3500-3504. doi: 10.1523/JNEUROSCI.4205-12.2013
- Houston, D. M., Pisoni, D. B., Kirk, K. I., Ying, E. A., & Miyamoto, R. T. (2003). Speech perception skills of deaf infants following cochlear implantation: a first report. *Int J Pediatr Otorhinolaryngol*, 67(5), 479-495.
- Howard-Jones, P. A., & Rosen, S. (1993). Uncomodulated glimpsing in "checkerboard" noise. *J Acoust Soc Am*, 93(5), 2915-2922.
- Huss, M., Verney, J. P., Fosker, T., Mead, N., & Goswami, U. (2011). Music, rhythm, rise time perception and developmental dyslexia: perception of musical meter predicts reading and phonology. *Cortex*, 47(6), 674-689. doi: 10.1016/j.cortex.2010.07.010
- Hutchison, K. A. (2003). Is semantic priming due to association strength or feature overlap? A microanalytic review. *Psychon Bull Rev*, 10(4), 785-813.
- Hutchison, K. A. (2007). Attentional control and the relatedness proportion effect in semantic priming. *J Exp Psychol Learn Mem Cogn*, 33(4), 645-662. doi: 10.1037/0278-7393.33.4.645
- Jacquier-Roux, M., Lequette, C., Pouget, G., Valdois, S., & Zorman, M. (2010). *Batterie Analytique du Langage Ecrit, Edition 2010*.
- Jacquier-Roux, M., Valdois, S., Zorman, M., Lequette, C., & Pouget, G. (2005). *ODEDYS : un outil de dépistage des dyslexies version 2*. Grenoble: Laboratoire cognisciences.
- Jastak, S., & Wilkinson, G. S. (1984). *The wide range achievement test-revised*. DE: Jastak associates.
- Jensen, J. K., & Neff, D. L. (1993). Development of basic auditory discrimination in preschool children. *Psychological Science*, 4(2), 104-107.
- Joanisse, M. F., Manis, F. R., Keating, P., & Seidenberg, M. S. (2000). Language deficits in dyslexic children: speech perception, phonology, and morphology. *J Exp Child Psychol*, 77(1), 30-60. doi: 10.1006/jecp.1999.2553
- Johnsrude, I. S., Mackey, A., Hakyemez, H., Alexander, E., Trang, H. P., & Carlyon, R. P. (2013). Swinging at a cocktail party: voice familiarity aids speech perception in the presence of a competing voice. *Psychol Sci*, 24(10), 1995-2004. doi: 10.1177/0956797613482467
- Jones, L. L. (2012). Prospective and retrospective processing in associative mediated priming. *Journal of Memory and Language*, 66, 52-67.
- Kamhi, A. G., & Catts, H. W. (1986). Toward an understanding of developmental language and reading disorders. *J Speech Hear Disord*, 51(4), 337-347.
- Kerlin, J. R., Shahin, A. J., & Miller, L. M. (2010). Attentional gain control of ongoing cortical speech representations in a "cocktail party". *J Neurosci*, 30(2), 620-628. doi: 10.1523/JNEUROSCI.3631-09.2010
- Kiefer, M. (2002). The N400 is modulated by unconsciously perceived masked words: further evidence for an automatic spreading activation account of N400 priming effects. *Brain Res Cogn Brain Res*, 13(1), 27-39.
- Kiefer, M., & Brendel, D. (2006). Attentional modulation of unconscious "automatic" processes: evidence from event-related potentials in a masked priming paradigm. *J Cogn Neurosci*, 18(2), 184-198. doi: 10.1162/089892906775783688

- Kiefer, M., & Martens, U. (2010). Attentional sensitization of unconscious cognition: task sets modulate subsequent masked semantic priming. *J Exp Psychol Gen*, 139(3), 464-489. doi: 10.1037/a0019561
- Kiefer, M., & Spitzer, M. (2000). Time course of conscious and unconscious semantic brain activation. *Neuroreport*, 11(11), 2401-2407.
- King, W. M., Lombardino, L. J., Crandell, C. C., & Leonard, C. M. (2003). Comorbid auditory processing disorder in developmental dyslexia. *Ear Hear*, 24(5), 448-456. doi: 10.1097/01.AUD.0000090437.10978.1A
- Klauer, K. C., Eder, A. B., Greenwald, A. G., & Abrams, R. L. (2007). Priming of semantic classifications by novel subliminal prime words. *Conscious Cogn*, 16(1), 63-83. doi: 10.1016/j.concog.2005.12.002
- Klepousniotou, E. (2002). The processing of lexical ambiguity: homonymy and polysemy in the mental lexicon. *Brain Lang*, 81(1-3), 205-223.
- Klepousniotou, E., & Baum, S. R. (2007). Disambiguating the ambiguity advantage effect in word recognition: An advantage for polysemous but not homonymous word. *Journal of Neurolinguistics*, 20, 1-24.
- Klepousniotou, E., Pike, G. B., Steinhauer, K., & Gracco, V. (2012). Not all ambiguous words are created equal: an EEG investigation of homonymy and polysemy. *Brain Lang*, 123(1), 11-21. doi: 10.1016/j.bandl.2012.06.007
- Klinger, M. R., Burton, P. C., & Pitts, G. S. (2000). Mechanisms of unconscious priming: I. Response competition, not spreading activation. *J Exp Psychol Learn Mem Cogn*, 26(2), 441-455.
- Kouider, S., & Dehaene, S. (2007). Levels of processing during non-conscious perception: a critical review of visual masking. *Philos Trans R Soc Lond B Biol Sci*, 362(1481), 857-875. doi: 10.1098/rstb.2007.2093
- Kouider, S., & Dehaene, S. (2009). Subliminal number priming within and across the visual and auditory modalities. *Exp Psychol*, 56(6), 418-433. doi: 10.1027/1618-3169.56.6.418
- Kouider, S., & Dupoux, E. (2004). Partial awareness creates the "illusion" of subliminal semantic priming. *Psychol Sci*, 15(2), 75-81.
- Kouider, S., & Dupoux, E. (2005). Subliminal speech priming. *Psychol Sci*, 16(8), 617-625. doi: 10.1111/j.1467-9280.2005.01584.x
- Kral, A., & Pallas, S. L. (2011). Development of auditory cortex. In J. Winer & C. E. Schreiner (Eds.), *The Auditory Cortex* (pp. 443-464). New York: Springer.
- Kraus, N., McGee, T. J., Carrell, T. D., Zecker, S. G., Nicol, T. G., & Koch, D. B. (1996). Auditory neurophysiologic responses and discrimination deficits in children with learning problems. *Science*, 273(5277), 971-973.
- Kunde, W., Kiesel, A., & Hoffmann, J. (2003). Conscious control over the content of unconscious cognition. *Cognition*, 88(2), 223-242.
- Kutas, M., & Federmeier, K. D. (2011). Thirty years and counting: finding meaning in the N400 component of the event-related brain potential (ERP). *Annu Rev Psychol*, 62, 621-647. doi: 10.1146/annurev.psych.093008.131123
- Kutas, M., & Hillyard, S. A. (1980). Reading senseless sentences: brain potentials reflect semantic incongruity. *Science*, 207(4427), 203-205.
- Kutas, M., & Hillyard, S. A. (1984). Brain potentials during reading reflect word expectancy and semantic association. *Nature*, 307(5947), 161-163.

- Kutas, M., Van Petten, C., & Kluender, R. (1994). Psycholinguistics Electrified. In M. A. Gernsbacher & M. Traxler (Eds.), *Handbook of Psycholinguistics*. New York: Elsevier Press.
- Lafon, J. C. (1964). *Le test phonétique et la mesure de l'audition*. Eindhoven: Editions Centrex.
- Lallier, M., Thierry, G., Tainturier, M. J., Donnadieu, S., Peyrin, C., Billard, C., & Valdois, S. (2009). Auditory and visual stream segregation in children and adults: an assessment of the amodality assumption of the 'sluggish attentional shifting' theory of dyslexia. *Brain Res*, 1302, 132-147. doi: 10.1016/j.brainres.2009.07.037
- Larousse. (2013). *Dictionnaire des Homonymes*.
- Lau, E. F., Phillips, C., & Poeppel, D. (2008). A cortical network for semantics: (de)constructing the N400. *Nat Rev Neurosci*, 9(12), 920-933. doi: 10.1038/nrn2532
- Lavie, N. (1995). Perceptual load as a necessary condition for selective attention. *J Exp Psychol Hum Percept Perform*, 21(3), 451-468.
- Lavie, N. (2005). Distracted and confused?: selective attention under load. *Trends Cogn Sci*, 9(2), 75-82. doi: 10.1016/j.tics.2004.12.004
- Lecoq, P. (1996). *l'Epreuve de Compréhension Syntaxico-Sémantique*.
- Lee, C. L., & Federmeier, K. D. (2009). Wave-ering: An ERP study of syntactic and semantic context effects on ambiguity resolution for noun/verb homographs. *J Mem Lang*, 61(4), 538-555. doi: 10.1016/j.jml.2009.08.003
- Lee, C. L., & Federmeier, K. D. (2012). Ambiguity's aftermath: how age differences in resolving lexical ambiguity affect subsequent comprehension. *Neuropsychologia*, 50(5), 869-879. doi: 10.1016/j.neuropsychologia.2012.01.027
- Leong, V., & Goswami, U. (2014). Impaired extraction of speech rhythm from temporal modulation patterns in speech in developmental dyslexia. *Front Hum Neurosci*, 8, 96. doi: 10.3389/fnhum.2014.00096
- Levafrais, P. (1967). *Test de l'Alouette* (2nd Edition ed.). Paris.
- Levitt, H. (1971). Transformed up-down methods in psychoacoustics. *J Acoust Soc Am*, 49(2), Suppl 2:467+.
- Lewis, J. L. (1970). Semantic processing of unattended messages using dichotic listening. *J Exp Psychol*, 85(2), 225-228.
- Licklider, J. C. (1948). The influence of interaural phase relations upon the masking of speech by white noise. *Journal of the Acoustical Society of America*, 20(2), 150-159.
- Lindenberger, U., Scherer, H., & Baltes, P. B. (2001). The strong connection between sensory and cognitive performance in old age: not due to sensory acuity reductions operating during cognitive assessment. *Psychol Aging*, 16(2), 196-205.
- Lucas, M. (2000). Semantic priming without association: a meta-analytic review. [Meta-Analysis]. *Psychon Bull Rev*, 7(4), 618-630.
- Luck, S. J., Vogel, E. K., & Shapiro, K. L. (1996). Word meanings can be accessed but not reported during the attentional blink. *Nature*, 383(6601), 616-618. doi: 10.1038/383616a0
- Ludwig, A. A., Fuchs, M., Kruse, E., Uhlig, B., Kotz, S. A., & Rubsamen, R. (2014). Auditory processing disorders with and without central auditory discrimination deficits. *J Assoc Res Otolaryngol*, 15(3), 441-464. doi: 10.1007/s10162-014-0450-3

- Marcel, A. J. (1983). Conscious and unconscious perception: experiments on visual masking and word recognition. *Cogn Psychol*, 15(2), 197-237.
- Marslen-Wilson, W. (1987). Functional parallelism in spoken word-recognition. *Cognition*, 25(1-2), 71-102.
- Marslen-Wilson, W., & Welsh, A. (1978). Processing interactions and lexical access during word recognition in continuous speech. *Cognitive Psychology*, 10, 29-63.
- Martin, C., Vu, H., Kellas, G., & Metcalfe, K. (1999). Strength of discourse context as a determinant of the subordinate bias effect. *Q J Exp Psychol A*, 52(4), 813-839. doi: 10.1080/713755861
- Martin, J., Cole, P., Leuwers, C., Casalis, S., Zorman, M., & Sprenger-Charolles, L. (2010). Reading in French-speaking adults with dyslexia. *Ann Dyslexia*, 60(2), 238-264. doi: 10.1007/s11881-010-0043-8
- Masquelier, M. P. (2011). Rémediation des troubles auditifs centraux chez les enfants. *Les Cahiers de l'Audition*, 24(1), 37-45.
- Mattys, S. L., Brooks, J., & Cooke, M. (2009). Recognizing speech under a processing load: Dissociating energetic from informational factors. *Cognitive psychology*, 59, 203-243.
- Mattys, S. L., & Wiget, L. (2011). Effects of cognitive load on speech recognition. *Journal of Memory and Language*, 65, 145-160.
- McAnally, K. I., & Stein, J. F. (1996). Auditory temporal coding in dyslexia. *Proc Biol Sci*, 263(1373), 961-965. doi: 10.1098/rspb.1996.0142
- McAnally, K. I., & Stein, J. F. (1997). Scalp potentials evoked by amplitude-modulated tones in dyslexia. *J Speech Lang Hear Res*, 40(4), 939-945.
- McArthur, G. M., Ellis, D., Atkinson, C. M., & Coltheart, M. (2008). Auditory processing deficits in children with reading and language impairments: can they (and should they) be treated? *Cognition*, 107(3), 946-977. doi: 10.1016/j.cognition.2007.12.005
- McArthur, G. M., & Hogben, J. H. (2001). Auditory backward recognition masking in children with a specific language impairment and children with a specific reading disability. *J Acoust Soc Am*, 109(3), 1092-1100.
- McArthur, G. M., Hogben, J. H., Edwards, V. T., Heath, S. M., & Mengler, E. D. (2000). On the "specifics" of specific reading disability and specific language impairment. *J Child Psychol Psychiatry*, 41(7), 869-874.
- McClelland, J. L., & Elman, J. L. (1986). The TRACE model of speech perception. *Cogn Psychol*, 18(1), 1-86.
- McDermott, J. H. (2009). The cocktail party problem. *Curr Biol*, 19(22), R1024-1027. doi: 10.1016/j.cub.2009.09.005
- McKoon, G., & Ratcliff, R. (1992). Spreading activation versus compound cue accounts of priming: mediated priming revisited. *J Exp Psychol Learn Mem Cogn*, 18(6), 1155-1172.
- McNamara, T. P. (2005). *Semantic priming perspectives from memory and word recognition*. New York: Psychology Press.
- McQueen, A., & Terhune, J. G. (2011). Central masking: Fact or Artifact. *The New School Psychology Bulletin*, 9(1), 34-39.
- Mehler, J., Jusczyk, P., Lambertz, G., Halsted, N., Bertoncini, J., & Amiel-Tison, C. (1988). A precursor of language acquisition in young infants. *Cognition*, 29(2), 143-178.

- Mengler, E. D., Hogben, J. H., Michie, P., & Bishop, D. V. (2005). Poor frequency discrimination is related to oral language disorder in children: a psychoacoustic study. *Dyslexia*, 11(3), 155-173.
- Messaoud-Galusi, S., Hazan, V., & Rosen, S. (2011). Investigating speech perception in children with dyslexia: is there evidence of a consistent deficit in individuals? *J Speech Lang Hear Res*, 54(6), 1682-1701. doi: 10.1044/1092-4388(2011/09-0261)
- Meyer, D. E., & Schvaneveldt, R. W. (1971). Facilitation in recognizing pairs of words: evidence of a dependence between retrieval operations. *J Exp Psychol*, 90(2), 227-234.
- Mody, M., Studdert-Kennedy, M., & Brady, S. (1997). Speech perception deficits in poor readers: auditory processing or phonological coding? *J Exp Child Psychol*, 64(2), 199-231. doi: 10.1006/jecp.1996.2343
- Moore, D. R. (2002). Auditory development and the role of experience. *Br Med Bull*, 63, 171-181.
- Moore, D. R. (2006). Auditory processing disorder (APD): Definition, diagnosis, neural basis and interaction. *Audiological Medicine*, 4, 4-11.
- Moore, D. R. (2011). The diagnosis and management of auditory processing disorder. *Lang Speech Hear Serv Sch*, 42(3), 303-308. doi: 10.1044/0161-1461(2011/10-0032)
- Moore, D. R. (2012). Listening difficulties in children: bottom-up and top-down contributions. *J Commun Disord*, 45(6), 411-418. doi: 10.1016/j.jcomdis.2012.06.006
- Moore, D. R., Cowan, J. A., Riley, A., Edmondson-Jones, A. M., & Ferguson, M. A. (2011). Development of auditory processing in 6- to 11-yr-old children. *Ear Hear*, 32(3), 269-285. doi: 10.1097/AUD.0b013e318201c468
- Moore, D. R., Ferguson, M. A., Edmondson-Jones, A. M., Ratib, S., & Riley, A. (2010). Nature of auditory processing disorder in children. *Pediatrics*, 126(2), e382-390. doi: 10.1542/peds.2009-2826
- Moore, D. R., Ferguson, M. A., Halliday, L. F., & Riley, A. (2008). Frequency discrimination in children: perception, learning and attention. *Hear Res*, 238(1-2), 147-154. doi: 10.1016/j.heares.2007.11.013
- Moore, D. R., Rosen, S., Bamiou, D. E., Campbell, N. G., & Sirimanna, T. (2013). Evolving concepts of developmental auditory processing disorder (APD): a British Society of Audiology APD special interest group 'white paper'. *Int J Audiol*, 52(1), 3-13. doi: 10.3109/14992027.2012.723143
- Moossavi, A., Mehrkian, S., Lotfi, Y., Faghihzadeh, S., & Sajedi, H. (2014). The relation between working memory capacity and auditory lateralization in children with auditory processing disorders. *Int J Pediatr Otorhinolaryngol*, 78(11), 1981-1986. doi: 10.1016/j.ijporl.2014.09.003
- Morgan, W. P. (1896). A Case of Congenital Word Blindness. *Br Med J*, 2(1871), 1378.
- Morrongiello, B. A., Fenwick, K. D., Hillier, L., & Chance, G. (1994). Sound localization in newborn human infants. *Dev Psychobiol*, 27(8), 519-538. doi: 10.1002/dev.420270805
- Morrongiello, B. A., & Trehub, S. E. (1987). Age-related changes in auditory temporal perception. *J Exp Child Psychol*, 44(3), 413-426.
- Morton, J. (1969). Interaction of information in word recognition. *Psychological Review*, 76(2), 165-178.

- Murphy, D. R., Craik, F. I., Li, K. Z., & Schneider, B. A. (2000). Comparing the effects of aging and background noise on short-term memory performance. *Psychol Aging*, 15(2), 323-334.
- Musiek, F. E. (1994). Frequency (pitch) and duration pattern tests. *J Am Acad Audiol*, 5(4), 265-268.
- Musiek, F. E. (2002). The frequency pattern test: A guide. *Hearing Journal*, 55(6), 58.
- Musiek, F. E., & Pinheiro, M. L. (1987). Frequency patterns in cochlear, brainstem, and cerebral lesions. *Audiology*, 26(2), 79-88.
- Naccache, L., Blandin, E., & Dehaene, S. (2002). Unconscious masked priming depends on temporal attention. *Psychol Sci*, 13(5), 416-424.
- Naccache, L., & Dehaene, S. (2001). Unconscious semantic priming extends to novel unseen stimuli. *Cognition*, 80(3), 215-229.
- Neely, J. H. (1977). Semantic priming and retrieval from lexical memory: Roles of Inhibitionless spreading activation limited-capacity attention. *Journal of Experimental Psychology: General*, 106(3), 226-254.
- Neely, J. H. (1991). Semantic priming effects in visual word recognition: a selective review of current findings and theories. In D. Besner & G. W. Humphreys (Eds.), *Basic processes in reading: Visual word recognition*. Hillsdale, NJ: Lawrence Erlbaum Associates, Inc.
- Neely, J. H., & Keefe, D. E. (1989). Semantic context effects in visual word processing: A hybrid prospective-retrospective processing theory. In G. H. Bower (Ed.), *The psychology of learning and motivation: Advances in research and theory*. New York: Academic Press.
- Neely, J. H., Keefe, D. E., & Ross, K. L. (1989). Semantic priming in the lexical decision task: roles of prospective prime-generated expectancies and retrospective semantic matching. *J Exp Psychol Learn Mem Cogn*, 15(6), 1003-1019.
- Neijenhuis, K., Snik, A., Priester, G., van Kordenoordt, S., & van den Broek, P. (2002). Age effects and normative data on a Dutch test battery for auditory processing disorders. *Int J Audiol*, 41(6), 334-346.
- 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.
- Newman, R. L., & Connolly, J. F. (2009). Electrophysiological markers of pre-lexical speech processing: evidence for bottom-up and top-down effects on spoken word processing. *Biol Psychol*, 80(1), 114-121. doi: 10.1016/j.biopsycho.2008.04.008
- Newman, R. S., & Jusczyk, P. W. (1996). The cocktail party effect in infants. *Percept Psychophys*, 58(8), 1145-1156.
- Nittrouer, S. (1999). Do temporal processing deficits cause phonological processing problems? *J Speech Lang Hear Res*, 42(4), 925-942.
- Norris, D. (1986). Word recognition: context effects without priming. *Cognition*, 22(2), 93-136.
- Norris, D., Cutler, A., McQueen, J. M., & Butterfield, S. (2006). Phonological and conceptual activation in speech comprehension. *Cogn Psychol*, 53(2), 146-193. doi: 10.1016/j.cogpsych.2006.03.001
- Obleser, J., & Kotz, S. A. (2010). Expectancy constraints in degraded speech modulate the language comprehension network. *Cereb Cortex*, 20(3), 633-640. doi: 10.1093/cercor/bhp128

- Obleser, J., & Kotz, S. A. (2011). Multiple brain signatures of integration in the comprehension of degraded speech. *Neuroimage*, 55(2), 713-723. doi: 10.1016/j.neuroimage.2010.12.020
- Obleser, J., & Weisz, N. (2012). Suppressed alpha oscillations predict intelligibility of speech and its acoustic details. *Cereb Cortex*, 22(11), 2466-2477. doi: 10.1093/cercor/bhr325
- Obleser, J., Wostmann, M., Hellbernd, N., Wilsch, A., & Maess, B. (2012). Adverse listening conditions and memory load drive a common alpha oscillatory network. *J Neurosci*, 32(36), 12376-12383. doi: 10.1523/JNEUROSCI.4908-11.2012
- Oldfield, R. C. (1971). The assessment and analysis of handedness: the Edinburgh inventory. *Neuropsychologia*, 9(1), 97-113.
- Olsho, L. W., Koch, E. G., & Halpin, C. F. (1987). Level and age effects in infant frequency discrimination. *J Acoust Soc Am*, 82(2), 454-464.
- Pakarinen, S., Takegata, R., Rinne, T., Huotilainen, M., & Naatanen, R. (2007). Measurement of extensive auditory discrimination profiles using the mismatch negativity (MMN) of the auditory event-related potential (ERP). *Clin Neurophysiol*, 118(1), 177-185. doi: 10.1016/j.clinph.2006.09.001
- Peelle, J. E., Troiani, V., Grossman, M., & Wingfield, A. (2011). Hearing loss in older adults affects neural systems supporting speech comprehension. *J Neurosci*, 31(35), 12638-12643. doi: 10.1523/JNEUROSCI.2559-11.2011
- Peterson, R. L., & Pennington, B. F. (2012). Developmental dyslexia. *Lancet*, 379(9830), 1997-2007. doi: 10.1016/S0140-6736(12)60198-6
- Phillips, D. P., Comeau, M., & Andrus, J. N. (2010). Auditory temporal gap detection in children with and without auditory processing disorder. *J Am Acad Audiol*, 21(6), 404-408. doi: 10.3766/jaaa.21.6.5
- Pichora-Fuller, M. K. (2003). Cognitive aging and auditory information processing. *Int J Audiol*, 42 Suppl 2, 2S26-32.
- Pichora-Fuller, M. K., Schneider, B. A., & Daneman, M. (1995). How young and old adults listen to and remember speech in noise. *J Acoust Soc Am*, 97(1), 593-608.
- Pinheiro, M. L., & Musiek, F. E. (1985). Sequencing and temporal ordering in the auditory system. In F. E. Musiek & M. L. Pinheiro (Eds.), *Assessment of central auditory dysfunction: Foundations and clinical correlates* (pp. 219-238). Baltimore: Williams & Wilkins.
- Polka, L., & Rvachew, S. (2005). The impact of otitis media with effusion on infant phonetic perception. *Infancy*, 8(2), 101-117.
- Ponton, C. W., Eggermont, J. J., Coupland, S. G., & Winkelaar, R. (1992). Frequency-specific maturation of the eighth nerve and brain-stem auditory pathway: evidence from derived auditory brain-stem responses (ABRs). *J Acoust Soc Am*, 91(3), 1576-1586.
- Ponton, C. W., Eggermont, J. J., Kwong, B., & Don, M. (2000). Maturation of human central auditory system activity: evidence from multi-channel evoked potentials. *Clin Neurophysiol*, 111(2), 220-236.
- Posner, M. I., & Snyder, C. R. R. (1975). Attention and cognitive control. In R. L. Solso (Ed.), *Information processing and cognition: The Loyola symposium*. Hillsdale, NJ: Erlbaum.
- Quillian, M. R. (1967). Word concepts: a theory and simulation of some basic semantic capabilities. *Behav Sci*, 12(5), 410-430.

- Rabbitt, P. M. (1968). Channel-capacity, intelligibility and immediate memory. *Q J Exp Psychol*, 20(3), 241-248. doi: 10.1080/14640746808400158
- Radeau, M. (1983). Semantic priming between spoken words in adults and children. *Canadian journal of psychology*, 37(4), 547-556.
- Radeau, M., Besson, M., Fonteneau, E., & Castro, S. L. (1998). Semantic, repetition and rime priming between spoken words: behavioral and electrophysiological evidence. *Biol Psychol*, 48(2), 183-204.
- Ramus, F. (2001). Dyslexia. Talk of two theories. *Nature*, 412(6845), 393-395. doi: 10.1038/35086683
- Ramus, F. (2003). Developmental dyslexia: specific phonological deficit or general sensorimotor dysfunction? *Curr Opin Neurobiol*, 13(2), 212-218.
- Ramus, F., & Ahissar, M. (2012). Developmental dyslexia: the difficulties of interpreting poor performance, and the importance of normal performance. *Cogn Neuropsychol*, 29(1-2), 104-122. doi: 10.1080/02643294.2012.677420
- Ramus, F., Rosen, S., Dakin, S. C., Day, B. L., Castellote, J. M., White, S., & Frith, U. (2003). Theories of developmental dyslexia: insights from a multiple case study of dyslexic adults. *Brain*, 126(Pt 4), 841-865.
- Ramus, F., & Szenkovits, G. (2008). What phonological deficit? *Q J Exp Psychol (Hove)*, 61(1), 129-141. doi: 10.1080/17470210701508822
- Ratcliff, R., & McKoon, G. (1988). A retrieval theory of priming in memory. *Psychol Rev*, 95(3), 385-408.
- Ratcliff, R., & McKoon, G. (1994). Retrieving information from memory: spreading-activation theories versus compound-cue theories. *Psychol Rev*, 101(1), 177-184; discussion 185-177.
- Raven, J. (1938). *Progressive matrices: A perceptual test of intelligence*. London: HK Lewis.
- Rayner, K., Pacht, J., & Duffy, S. A. (1994). Effects of prior encounter and global discourse bias on the processing of lexically ambiguous words: evidence from eye fixations. *Journal of Memory and Language*, 33, 527-544.
- Reynvoet, B., Gevers, W., & Caessens, B. (2005). Unconscious primes activate motor codes through semantics. *J Exp Psychol Learn Mem Cogn*, 31(5), 991-1000. doi: 10.1037/0278-7393.31.5.991
- Richardson, U., Thomson, J. M., Scott, S. K., & Goswami, U. (2004). Auditory processing skills and phonological representation in dyslexic children. *Dyslexia*, 10(3), 215-233. doi: 10.1002/dys.276
- Rickard, N. A., Smales, C. J., & Rickard, K. L. (2013). A computer-based auditory sequential pattern test for school-aged children. *Int J Pediatr Otorhinolaryngol*, 77(5), 838-842. doi: 10.1016/j.ijporl.2013.02.024
- Rodd, J. (2004). The effect of semantic ambiguity on reading aloud: a twist in the tale. *Psychon Bull Rev*, 11(3), 440-445.
- Rodd, J., Gaskell, G., & Marslen-Wilson, W. (2002). Making sense of semantic ambiguity: Semantic competition in lexical access. *Journal of Memory and Language*, 46, 245-266.
- Rodd, J., Lopez Cutrin, B., Kirsch, H., Millar, A., & Davis, M. (2013). Long-term priming of the meaning of ambiguous words. *Journal of Memory and Language*, 68, 180-198.

- Rohm, D., Klimesch, W., Haider, H., & Doppelmayr, M. (2001). The role of theta and alpha oscillations for language comprehension in the human electroencephalogram. *Neurosci Lett*, 310(2-3), 137-140.
- Romei, L., Wambacq, I. J., Besing, J., Koehnke, J., & Jerger, J. (2011). Neural indices of spoken word processing in background multi-talker babble. *Int J Audiol*, 50(5), 321-333. doi: 10.3109/14992027.2010.547875
- Rosen, S. (2003). Auditory processing in dyslexia and specific language impairment: is there a deficit? What is its nature? Does it explain anything? *Journal of Phonetics*, 31, 509-527.
- Rosen, S. (2005). "A riddle wrapped in a mystery inside an enigma": defining central auditory processing disorder. [Comment]. *Am J Audiol*, 14(2), 139-142; discussion 143-150. doi: 10.1044/1059-0889(2005/015)
- Rosen, S., & Manganari, E. (2001). Is there a relationship between speech and nonspeech auditory processing in children with dyslexia? *J Speech Lang Hear Res*, 44(4), 720-736.
- Rugg, M. D. (1985). The effects of semantic priming and word repetition on event-related potentials. *Psychophysiology*, 22(6), 642-647.
- Russeler, J., Becker, P., Johannes, S., & Munte, T. F. (2007). Semantic, syntactic, and phonological processing of written words in adult developmental dyslexic readers: an event-related brain potential study. *BMC Neurosci*, 8, 52. doi: 10.1186/1471-2202-8-52
- Rutter, M., Caspi, A., Fergusson, D., Horwood, J., Goodman, R. J., Maughan, B., . . . Carroll, J. (2004). Sex differences in developmental reading disability. *Journal of American Medical Association*, 291(16), 2007-2012.
- Sabisch, B., Hahne, A., Glass, E., von Suchodoletz, W., & Friederici, A. D. (2006). Auditory language comprehension in children with developmental dyslexia: evidence from event-related brain potentials. *J Cogn Neurosci*, 18(10), 1676-1695. doi: 10.1162/jocn.2006.18.10.1676
- Salthouse, T. A. (1996). Constraints on theories of cognitive aging. *Psychon Bull Rev*, 3(3), 287-299. doi: 10.3758/BF03210753
- Schacter, D. L., & Church, B. A. (1992). Auditory priming: implicit and explicit memory for words and voices. *J Exp Psychol Learn Mem Cogn*, 18(5), 915-930.
- Schneider, B. A., Trehub, S. E., Morrongiello, B. A., & Thorpe, L. A. (1989). Developmental changes in masked thresholds. *J Acoust Soc Am*, 86(5), 1733-1742.
- Schoof, T., & Rosen, S. (2014). The role of auditory and cognitive factors in understanding speech in noise by normal-hearing older listeners. *Front Aging Neurosci*, 6, 307. doi: 10.3389/fnagi.2014.00307
- Schulte-Körne, G., Deimel, W., Bartling, J., & Remschmidt, H. (1998). Auditory processing and dyslexia: evidence for a specific speech processing deficit. *Neuroreport*, 9(2), 337-340.
- Sereno, S. C., Rayner, K., & Posner, M. I. (1998). Establishing a time-line of word recognition: evidence from eye movements and event-related potentials. *Neuroreport*, 9(10), 2195-2200.
- Serniclaes, W., Van Heghe, S., Mousty, P., Carre, R., & Sprenger-Charolles, L. (2004). Allophonic mode of speech perception in dyslexia. *J Exp Child Psychol*, 87(4), 336-361. doi: 10.1016/j.jecp.2004.02.001

- Shafer, V. L., Yu, Y. H., & Wagner, M. (2014). Maturation of cortical auditory evoked potentials (CAEPs) to speech recorded from frontocentral and temporal sites: Three months to eight years of age. *Int J Psychophysiol*. doi: 10.1016/j.ijpsycho.2014.08.1390
- Share, D., Jorm, A., Maclean, R., & Matthews, R. (2002). Temporal processing and reading disability. *Reading and Writing : An Interdisciplinary Journal*, 15, 151-171.
- Sharma, A., Kraus, N., McGee, T. J., & Nicol, T. G. (1997). Developmental changes in P1 and N1 central auditory responses elicited by consonant-vowel syllables. *Electroencephalogr Clin Neurophysiol*, 104(6), 540-545.
- Sharma, M., Purdy, S. C., & Kelly, A. S. (2009). Comorbidity of auditory processing, language, and reading disorders. *J Speech Lang Hear Res*, 52(3), 706-722. doi: 10.1044/1092-4388(2008/07-0226)
- Sidis, B. (1898). The Psychology of Suggestion. *Science*, 8(188), 162-163. doi: 10.1126/science.8.188.162
- Simpson, S. A., & Cooke, M. (2005). Consonant identification in N-talker babble is a nonmonotonic function of N. *J Acoust Soc Am*, 118(5), 2775-2778.
- Smith, N. A., & Trainor, L. J. (2011). Auditory Stream Segregation Improves Infants' Selective Attention to Target Tones Amid Distractors. *Infancy*, 16(6), 655-668. doi: 10.1111/j.1532-7078.2011.00067.x
- Snowling, M. J. (2000). *Dyslexia*. Oxford: Blackwells.
- Sperling, A. J., Lu, Z. L., Manis, F. R., & Seidenberg, M. S. (2005). Deficits in perceptual noise exclusion in developmental dyslexia. *Nat Neurosci*, 8(7), 862-863. doi: 10.1038/nn1474
- Spetner, N. B., & Olsho, L. W. (1990). Auditory frequency resolution in human infancy. *Child Dev*, 61(3), 632-652.
- Sprenger-Charolles, L., Cole, P., Lacert, P., & Serniclaes, W. (2000). On subtypes of developmental dyslexia: evidence from processing time and accuracy scores. *Can J Exp Psychol*, 54(2), 87-104.
- Spruyt, A., De Houwer, J., Everaert, T., & Hermans, D. (2012). Unconscious semantic activation depends on feature-specific attention allocation. *Cognition*, 122(1), 91-95. doi: 10.1016/j.cognition.2011.08.017
- Stanovich, K. E. (1986a). Matthew effects in reading: Some consequences of individual differences in the acquisition of literacy. *Reading Research Quarterly*, 21, 360-406.
- Stanovich, K. E. (1986b). Matthew effects in reading: Some consequences of individual differences in the acquisition of literacy. *Reading Research Quarterly*, 21, 360-407.
- Stanovich, K. E. (1988). Explaining the differences between the dyslexic and the garden-variety poor reader: the phonological-core variable-difference model. *J Learn Disabil*, 21(10), 590-604.
- Stanovich, K. E., Siegel, L. S., & Gottardo, A. (1997). Converging evidence for phonological and surface subtypes of reading disability. *Journal of Educational Psychology*, 89(1), 114-127.
- Stein, J., & Walsh, V. (1997). To see but not to read; the magnocellular theory of dyslexia. *Trends Neurosci*, 20(4), 147-152.
- Strauss, A., Kotz, S. A., & Obleser, J. (2013). Narrowed expectancies under degraded speech: revisiting the N400. *J Cogn Neurosci*, 25(8), 1383-1395. doi: 10.1162/jocn_a_00389

- Strauss, A., Wostmann, M., & Obleser, J. (2014). Cortical alpha oscillations as a tool for auditory selective inhibition. *Front Hum Neurosci*, 8, 350. doi: 10.3389/fnhum.2014.00350
- Surprenant, A. (1999). The effect of noise on memory for spoken syllables. *International journal of psychology*, 34(5/6), 328-333.
- Sussman, E., Steinschneider, M., Gumenyuk, V., Grushko, J., & Lawson, K. (2008). The maturation of human evoked brain potentials to sounds presented at different stimulus rates. *Hear Res*, 236(1-2), 61-79. doi: 10.1016/j.heares.2007.12.001
- Sussman, E., Wong, R., Horvath, J., Winkler, I., & Wang, W. (2007). The development of the perceptual organization of sound by frequency separation in 5-11-year-old children. *Hear Res*, 225(1-2), 117-127. doi: 10.1016/j.heares.2006.12.013
- Sutcliffe, P., & Bishop, D. (2005). Psychophysical design influences frequency discrimination performance in young children. *J Exp Child Psychol*, 91(3), 249-270. doi: 10.1016/j.jecp.2005.03.004
- Swaab, T., Brown, C., & Hagoort, P. (2003). Understanding words in sentence contexts: the time course of ambiguity resolution. *Brain Lang*, 86(2), 326-343.
- Tallal, P. (1980). Auditory temporal perception, phonics, and reading disabilities in children. *Brain Lang*, 9(2), 182-198.
- Tallal, P., & Gaab, N. (2006). Dynamic auditory processing, musical experience and language development. *Trends Neurosci*, 29(7), 382-390. doi: 10.1016/j.tins.2006.06.003
- Tallal, P., & Piercy, M. (1973). Defects of non-verbal auditory perception in children with developmental aphasia. *Nature*, 241(5390), 468-469.
- Thompson-Schill, S. L., Kurtz, K. J., & Gabrieli, J. D. E. (1998). Effects of semantic and associative relatedness on automatic priming. *Journal of Memory and Language*, 38, 440-458.
- Trainor, L., McFadden, M., Hodgson, L., Darragh, L., Barlow, J., Matsos, L., & Sonnadara, R. (2003). Changes in auditory cortex and the development of mismatch negativity between 2 and 6 months of age. *Int J Psychophysiol*, 51(1), 5-15.
- Trehub, S. E., Schneider, B. A., & Endman, M. (1980). Developmental changes in infants' sensitivity to octave-band noises. *J Exp Child Psychol*, 29(2), 282-293.
- Trehub, S. E., Schneider, B. A., & Henderson, J. L. (1995). Gap detection in infants, children, and adults. *J Acoust Soc Am*, 98(5 Pt 1), 2532-2541.
- Tun, P. A., McCoy, S., & Wingfield, A. (2009). Aging, hearing acuity, and the attentional costs of effortful listening. *Psychol Aging*, 24(3), 761-766. doi: 10.1037/a0014802
- Van den Bussche, E., Hughes, G., Humbeeck, N. V., & Reynvoet, B. (2010). The relation between consciousness and attention: an empirical study using the priming paradigm. *Conscious Cogn*, 19(1), 86-97. doi: 10.1016/j.concog.2009.12.019
- Van den Bussche, E., Van den Noortgate, W., & Reynvoet, B. (2009). Mechanisms of masked priming: a meta-analysis. *Psychol Bull*, 135(3), 452-477. doi: 10.1037/a0015329
- Van Deun, L., van Wieringen, A., Van den Bogaert, T., Scherf, F., Offeciers, F. E., Van de Heyning, P. H., . . . Wouters, J. (2009). Sound localization, sound lateralization, and binaural masking level differences in young children with normal hearing. *Ear Hear*, 30(2), 178-190. doi: 10.1097/AUD.0b013e318194256b

- Van Engen, K., & Bradlow, A. (2007). Sentence recognition in native- and foreign-language multi-talker background noise. *Journal of the Acoustical Society of America*, 121(1), 519-526.
- Van Petten, C. (2014). Selective Attention, Processing Load, and Semantics: Insights from Human Electrophysiology. In G. R. Mangun (Ed.), *Cognitive Electrophysiology of Attention* (pp. 236-253): Hardbound.
- Van Petten, C., & Kutas, M. (1987). Ambiguous words in context: An event-related potential analysis of the time course of meaning activation. *Journal of Memory and Language*, 26, 188-208.
- Vu, H., Kellas, G., & Paul, S. T. (1998). Sources of sentence constraint on lexical ambiguity resolution. *Mem Cognit*, 26(5), 979-1001.
- Werner, L. A., Fay, R., & Popper, A. (2011). *Human auditory development* (Vol. 42): Springer Science & Business Media.
- Werner, L. A., & Marean, G. C. (1996). *Human auditory development*. Boulder, CO: Westview Press.
- West, W. C., & Holcomb, P. J. (2002). Event-related potentials during discourse-level semantic integration of complex pictures. *Brain Res Cogn Brain Res*, 13(3), 363-375.
- White-Schwoch, T., Woodruff Carr, K., Thompson, E. C., Anderson, S., Nicol, T., Bradlow, A. R., . . . Kraus, N. (2015). Auditory Processing in Noise: A Preschool Biomarker for Literacy. *PLoS Biol*, 13(7), e1002196. doi: 10.1371/journal.pbio.1002196
- White, S., Milne, E., Rosen, S., Hansen, P., Swettenham, J., Frith, U., & Ramus, F. (2006). The role of sensorimotor impairments in dyslexia: a multiple case study of dyslexic children. *Dev Sci*, 9(3), 237-255; discussion 265-239. doi: 10.1111/j.1467-7687.2006.00483.x
- Wiley, J., & Rayner, K. (2000). Effects of titles on the processing of text and lexically ambiguous words: evidence from eye movements. *Mem Cognit*, 28(6), 1011-1021.
- Willems, R. M., Oostenveld, R., & Hagoort, P. (2008). Early decreases in alpha and gamma band power distinguish linguistic from visual information during spoken sentence comprehension. *Brain Res*, 1219, 78-90. doi: 10.1016/j.brainres.2008.04.065
- Wingfield, A., Tun, P. A., & McCoy, S. L. (2005). Hearing loss in Older Adulthood. What it is and howit interacts with cognitive performances. *Current directions in psychological science*, 14(3), 144-148.
- Wolf, M., & Bowers, P. G. (2000). Naming-speed processes and developmental reading disabilities: an introduction to the special issue on the double-deficit hypothesis. *J Learn Disabil*, 33(4), 322-324.
- Wood, N. L., Stadler, M. A., & Cowan, N. (1997). Is there implicit memory without attention? A reexamination of task demands in Eich's (1984) procedure. *Mem Cognit*, 25(6), 772-779.
- Wright, B. A., & Zecker, S. G. (2004). Learning problems, delayed development, and puberty. *Proc Natl Acad Sci U S A*, 101(26), 9942-9946. doi: 10.1073/pnas.0401825101
- Wunderlich, J. L., Cone-Wesson, B. K., & Shepherd, R. (2006). Maturation of the cortical auditory evoked potential in infants and young children. *Hear Res*, 212(1-2), 185-202. doi: 10.1016/j.heares.2005.11.010

- Ziegler, J. C., Pech-Georgel, C., George, F., & Lorenzi, C. (2009). Speech-perception-in-noise deficits in dyslexia. *Dev Sci*, 12(5), 732-745. doi: 10.1111/j.1467-7687.2009.00817.x
- Zwitserslood, P. (1989). The locus of the effects of sentential-semantic context in spoken-word processing. *Cognition*, 32(1), 25-64.

Annexes

Curriculum Vitae

Marie DEKERLE
17-12-1987

Groupe de Recherche ALP
Laboratoire L2C2
Institut des Sciences Cognitives
67 Bvd Pinel
69675 Bron Cedex - FRANCE
+33 6 89 14 04 19
marie.dekerle@isc.cnrs.fr

EDUCATION

- Since 2012** PhD Student in Neurosciences at Laboratory on Language, Brain and Cognition under the supervision of Dr. F. Meunier. *The impact of signal degradation on speech, from its representation to its comprehension*. Université Lyon 1.
- 2011** National Research Master 2 of Neuropsychology, with honors. Université Lyon 2.
- 2010** Master 1 of Cognitive Psychology and Neuropsychology. Université Lyon 2.
- 2009** Bachelor of Psychology. Université Catholique de Lyon.
- 2005** Scientific Baccalauréat (High School Diploma) Biology specialty, with honors.

INTERNSHIPS

- 2010-2011** Research internship at Lyon Neurosciences Research Center (INSERM, CNRS) under supervision of Dr. F. Meunier and Dr. V. Boulenger. *Masked semantic priming in cocktail party situation: electrophysiological and behavioral studies*.
- 2009-2010** Research internship at Laboratoire Dynamique du Langage (CNRS) under supervision of Dr. F. Meunier and Dr. V. Boulenger. *Semantic priming in speech in speech comprehension, contribution of cocktail party situation*.

PROFESSIONAL EXPERIENCES

- September 2011- August 2012** Temporary position at Lyon Neurosciences Research Center, CNRS, Lyon.
- June-July 2010** Temporary position at Laboratoire Dynamique du Langage, CNRS, Lyon.

SKILLS

- Languages** French (native language); English ; German (Scholar level)
- Computer** Programming software (E-Prime, Praat, Matlab, Presentation), statistical and data analyses (Statistica, Besa).

COMMUNICATIONS

Oral Communication

- Dekerle, M., Meunier, F., Donnadieu, S.** (2013). Central Auditory Processes and their links with verbal skills in typically developing children. Interspeech 25th-29th August. Lyon France

Posters

- Dekerle, M.**, Donnadieu, S. Meunier, F. (2014). Liens entre les traitements auditifs centraux et les compétences langagières chez les adultes normo-lecteurs et dyslexiques. 22^{ième} Journée Scientifique du CERNEC, 20 et 21 mars, Montréal, Québec, Canada.
- Dekerle, M.**, Hoen, M., Grataloup, C., Meunier, F. (2013). Speech-in-noise Comprehension and Poor Reading. Third Oxford-Kobe Symposium, 11th – 13th April, Oxford, UK.
- Dekerle, M.**, Boulenger, V., Hoen, M., Meunier, F. (2012). Semantic Priming at the Cocktail Party : One can hear a word without accessing its meaning. Mental Lexicon, 24th-26th October, Montreal, Canada.
- Dekerle, M.**, Boulenger, V., Hoen, M., Meunier, F. (2012). L’amorçage sémantique masqué en situation de cocktail party. 29^{ème} Journées d’études sur la Parole, du 4 au 8 juin, Grenoble, France.
- Dekerle, M.**, Boulenger, V., Meunier, F. (2012). Semantic Priming at the Cocktail Party : low intelligibility reveals automatic processes. 4th Workshop on Intelligibility and Quality of Speech in Noise, 5th-6th January, Cardiff, UK.
- Dekerle, M.**, Boulenger, V., Meunier, F. (2011). Semantic Priming at the Cocktail Party : automatic and strategic processes. . 17th ESCOP Meeting, 29th September – 2nd October, San Sebastian, Spain.
- Dekerle, M.**, Boulenger, V., Meunier, F. (2011). Semantic Priming at the Cocktail Party. 3rd Workshop on Intelligibility and Quality of Speech in Noise, Institut des Sciences de l’Homme, 6th -7th January, Lyon, France.

GRANTS

- 2014** Grant EXPLORA DOC for a 6 months international collaboration with University of Montreal, provided by Région Rhône-Alpes (4320€)

TEACHING

- 2013-2015** *Neuroscience*, (Lectures – 4h) Psychology Bachelor, Université Catholique de Lyon.
- 2014** *Neuroscience*, (Seminars – 44h) Psychology Bachelor, Université Catholique de Lyon.
- 2012** *Children language disorder* (Seminars – 17h), Language sciences Bachelor, Université Lyon 2.

OTHERS

- 2012- 2015** Representative of the PhD students of L2C2.
- 2014** Member of the organization committee for Pint of Science festival in Lyon.
- 2013** Scientific animation of the developmental disorder stand at the Brain’s Week (week for popularization of brain research).
- 2013** Volunteer for the association “Lyon Science – Academie” – proposing short internship for high school students from disadvantaged area interested in pursuing science studies.

Liste de stimuli Etude 1

Liste de mots prononcés par les Speaker 1 et Speaker 2 (voix Sémantiquement Consistantes) sur les essais associés à des mots cibles.

	PRIME	TARGET		PRIME	TARGET
S1	corbeau voler nid hirondelle merle	pigeon	S1	baigner sable pelle seau plage	serviette
S2	cage moineau oiseau roucouler volière		S2	bronzer transat solaire touriste vacances	
S1	informatique clavier ordinateur enceinte programme	écran	S1	vin célébrer trinquer alcool cave	champagne
S2	cliquer fichier connexion imprimante réseau		S2	boire vodka sommelier millésime cocktail	
S1	vache élever poule chèvre tracteur	cochon	S1	sculpteur métier enseignant maçon peintre	poète
S2	âne lapin fermier traire étable		S2	exercer cinéaste acteur plombier médecin	

	PRIME	TARGET		PRIME	TARGET
S1	tourelle armure roi défendre régner	château	S1	planer chouette serre falaise aigle	hibou
S2	donjon palefrenier pont douve cachot		S2	fauconnier rapace condor vautour nocturne	
S1	souche écorce sève feuille tronc	racine	S1	dompter jaguar sauvage rapide léopard	morsure
S2	branche élaguer nœud cime résine		S2	puma traquer proie panthère lynx	
S1	guichet train gare composter cheminot	billet	S1	créer pinceau aquarelle chevalet toile	artiste
S2	rail contrôleur wagon terminus locomotive		S2	peinture gouache tableau crayonner atelier	

	PRIME	TARGET		PRIME	TARGET
S1	pull chemise vêtir débardeur chandail	gilet	S1	minute trimestre heure chronométrer lundi	seconde
S2	chaussette jupon bermuda habiller pantalon		S2	sablier attendre date année calendrier	
S1	détenu arrêter tribunal parloir juger	prison	S1	pêcher ligne canne ferrer bouchon	filet
S2	cellule réclusion police avocat criminel		S2	hameçon vers épuisette chalutier moulinet	
S1	bonnet chapeau béret cagoule sombrero	foulard	S1	cil iris mascara lunette myopie	pupille
S2	chapelier couvrir casque visière képi		S2	loucher cerne sourcil maquiller lentille	

	PRIME	TARGET		PRIME	TARGET
S1	casier coffre trier commode placard	buffet	S1	curry épicer piment oriental muscade	cannelle
S2	meuble plier vêtement armoire bureau		S2	gingembre poivre coriandre arôme relever	
S1	robinet lavabo évier tremper shampooing	baignoire	S1	verre manger nappe attabler couvert	couteau
S2	laver douche Savon peignoir mouiller		S2	cuiller fourchette servir table repas	
S1	carreau ouverture aérer lucarne baie	rideau	S1	chêne pommier hêtre merisier scier	platane
S2	vitre fenêtre balcon persienne store		S2	érable oranger fleurir frêne verger	

	PRIME	TARGET		PRIME	TARGET
S1	carreler lino dalle faïence artisan	parquet	S1	marmite sauteuse poêlon casserole cuisson	cocotte
S2	joint mosaïque plancher béton moquette		S2	mijoter ébullition cuisiner cuire gazinière	
S1	bus transiter commun horaire arrêt	métro	S1	potager légume arrosage fertilisant limace	carotte
S2	navette terminus tramway voyager car		S2	terreau jardinage bêcher céleri navet	
S1	basket pantoufle botte soulier chausse-pied	lacet	S1	caniche basset chihuahua labrador bichon	berger
S2	escarpin sabot chausson marcher ballerine		S2	chien lévrier pédigrée épagneul Cocker	

	PRIME	TARGET		PRIME	TARGET
S1	doigt ongle attraper manucure tenir	poignet	S1	guitare mandoline contrebasse mélodie orchestre	violon
S2	pouce manipuler phalange bague main		S2	banjo harpe trompette symphonie interpréter	
S1	nez joues front visagiste grimacer	figure	S1	cuisiner potée dégustation viande gourmet	gigot
S2	farder oreille pommette bouche menton		S2	barbecue rôti saucisse fumé grillade	
S1	huitre moule bulot décortiquer palourde	crevette	S1	grandir cuisse coude hanche bras	épaule
S2	praire homard langoustine crabe gambas		S2	genou pied cheville poignet jambe	

	PRIME	TARGET		PRIME	TARGET
S1	corail masque plonger épave sable	trésor	S1	prêter découvert banquier dépenser épargne	salaire
S2	algue galet tuba récif palme		S2	crédit emprunter banque débit hypothèque	
S1	plomb fer cuivre souder chaudronnier	acier	S1	terrier furet ronger souris écureuil	raton
S2	bronze forgeron métal zinc étain		S2	hamster mulot cobaye gerbille loir	
S1	bague chaîne montre parer médaille	collier	S1	astronaute télescope constellation fusée orbite	planète
S2	pendentif bijou chevalière orner diadème		S2	étoile comète galaxie astre alunir	

	PRIME	TARGET		PRIME	TARGET
S1	ramper crocodile serpent caïman reptile	vipère	S1	volant capot banquette conduire circulation	essence
S2	boa iguane python lézards tortue		S2	caravane rouler voiture autoroute frein	
S1	cheminée poêle foyer cendre ramoner	foyer	S1	aiguille fil coudre ourlet patron	tissu
S2	allumette barbecue flamber suie étincelle		S2	broderie enfiler tricoter bouton dé	
S1	teindre peigne brosse pince bigoudis	chignon	S1	croix chapelle crypte église messe	prière
S2	cheveux tresse coiffer rincer démêler		S2	curé clocher autel bible évangile	

Liste de stimuli Etude 2

Liste des phrases utilisées en condition Ambigüe.

Les mots typiquement Québécois ou n'ayant pas le même sens en France et au Québec sont expliqués en bas de page.

Ambigüe	En chemin, il a appelé sa tante, il sera prêt pour le camping
Ambigüe	Le policier m'a mis une amende, heureusement il y a aussi les noix
Ambigüe	Cet immeuble a une bonne hauteur, en fait c'est une grande romancière
Ambigüe	Le sac est vraiment plein, et j'aurai préféré qu'il soit envié
Ambigüe	L'adolescent a fait une fugue, mais il n'est pas un bon compositeur
Ambigüe	Les poissons nagent dans les mares, et il y en a plein le café
Ambigüe	Dans l'armoire j'ai trouvé des vers, il y avait même des asticots
Ambigüe	Sur cette radio il y a beaucoup de rock, mais j'aime la géologie
Ambigüe	Cette ampoule avait beaucoup de watt, je n'ai jamais vu autant de coton
Ambigüe	La vieille dame a une belle canne, mais n'a pas d'autres animaux
Ambigüe	Il est enfermé dans une cellule, et ne peut sortir de ce globule
Ambigüe	Il leur a fait une farce, et il l'a mise dans la dinde
Ambigüe	Les policiers ont retrouvés les faux, ils les ont donnés aux jardiniers
Ambigüe	Elle a fait un accroc à son pull, il a un admirateur
Ambigüe	Sur le bord de la route il m'a fait un signe, puis il m'a fait un canard
Ambigüe	Elle est spécialiste des pies, car elle beaucoup étudié la trigonométrie
Ambigüe	Tu lui as fait une bonne blague, espérons qu'il ne perdra plus son tabac
Ambigüe	Chez moi, il y a un grand compteur, il a de nouvelles histoires
Ambigüe	Au château habitent les reines, alors qu'il n'y a rien pour les wapitis
Ambigüe	La randonnée est pleine de côtes, il n'y a pas d'autres fossiles
Ambigüe	Ce garçon a beaucoup de vaine, et donc beaucoup d'artères
Ambigüe	Ce n'est pas à cause d'eux, c'est le bacon
Ambigüe	Le coureur a pris son élan, et a fui en laissant son caribou
Ambigüe	Non, il n'est pas en taule, il est en acier
Ambigüe	Après les rectos, on fait les versos, mais il y a trop de sagittaire
Ambigüe	Il y avait là toute une bande, mais c'était une grande blessure
Ambigüe	Au lac, il a attrapé un thon mais je ne mange pas d'insectes ¹
Ambigüe	La phrase se termine par un point, parfois c'est tout le bras
Ambigüe	Les bateaux ont vu le phare, alors ils ont suivi le maquillage
Ambigüe	Le plombier a vu la fuite, mais n'a pas pris part à la poursuite
Ambigüe	Lorsque je l'ai vu je l'ai tout de suite reconnu, c'était mon erreur
Ambigüe	Ça c'est du beau boulot, on le placera dans la cour ²
Ambigüe	Nous sommes allés voir des canons, c'était de beaux chants
Ambigüe	Derrière les montagnes il y a les vallées, et il y aussi d'autres serviteurs
Ambigüe	Au japon, ils paient en yen, c'est plus pratique que les éléphants

¹ Taon se prononce de la même manière que thon

² Le terme cour est l'équivalent du jardin, alors que jardin sera employé dans le sens de potager

Ambigüe	Le sommelier a trouvé un vin, en fait c'est un très bon quarante
Ambigüe	Je ne sais pas quoi en faire, ou alors c'est de l'aluminium
Ambigüe	A la piscine j'ai un casier, mais je n'ai pas eu de procès
Ambigüe	Il a versé du vin dans sa coupe, elle en avait plein les cheveux
Ambigüe	Toute ces aventures l'avait crevé, il a dû changer le pneu
Ambigüe	Au restaurant nous avons de bonnes recettes, nous ne serons pas en déficit
Ambigüe	La voiture a 4 roues, mais il faut changer les blonds
Ambigüe	Sur le calendrier il n'y avait plus de dates, donc nous avons pris des pruneaux
Ambigüe	Le mannequin n'aime pas son teint, elle préfère la ciboulette
Ambigüe	Ils sont vivant ces êtres, les morts sont des érables
Ambigüe	L'instituteur n'a pas de classe, il n'est pas distingué
Ambigüe	Ce soir de tempête il y avait beaucoup d'éclairs, c'était un sacré dessert
Ambigüe	La ménagère a monté son balai, il y avait beaucoup de danseurs
Ambigüe	Chez le dentiste, il a parlé de son palais, l'autre avait des maisonnettes
Ambigüe	A la ferme il y avait un porc, je n'ai jamais vu autant de bateaux
Ambigüe	Ils n'ont pas respecté la police, il ne ressemble à rien ce document
Ambigüe	Pour y aller il a pris un car, il aurait pu prendre un huitième
Ambigüe	Il a déposé sa plainte, mais ça n'a pas soutenu le mur
Ambigüe	Sur le chantier j'ai vu une immense grue, mais je ne m'y connais pas en oiseaux
Ambigüe	Il faut que j'aile prendre l'air, j'ai besoin aussi du périmètre
Ambigüe	Si le cuisinier change tout ça en mets, alors nous aurons un bon juin
Ambigüe	Ce voyou ne respecte pas les règles, du coup il ne peut pas prendre de mesure
Ambigüe	J'avais besoin de basilic, il m'a acheté une cathédrale
Ambigüe	A la fin de la guerre nous attendions la paix, mais il n'y avait plus d'argent
Ambigüe	Dans le frigo il n'y a plus de lait, alors essaie de trouver des mignons
Ambigüe	Il faut y mettre un terme, on pourra prendre des bains
Ambigüe	Les scientifiques ont construit un vaisseau, ils veulent retourner sur la veine
Ambigüe	Avec autant de montagne et de lac, je ne peux oublier cette coiffure
Ambigüe	C'était une grosse entorse, il aurait dû faire attention au règlement
Ambigüe	La vieille dame était gelée ³ , mais elle n'avait pas pris de drogue.
Ambigüe	Près du tribunal j'ai vu des avocats, mais je n'aime pas le guacamole
Ambigüe	Cet écrivain avait tous les mots, mais il ne savait faire des remèdes
Ambigüe	Chez la couturière il y avait beaucoup de files, les clients devaient avoir de la patience
Ambigüe	Sur la table il y avait 52 cartes, mais nous n'étions que 2 pilotes
Ambigüe	La princesse a fait refaire sa couronne, elle n'a plus mal à la dent
Ambigüe	Cette voiture a réellement foncé, à présent elle est noire
Ambigüe	Il ne m'a pas parlé de toi, mais j'ai besoin de nouvelles tuiles
Ambigüe	Lorsqu'ils sont arrivés ils ont jeté l'ancre, ils étaient pleins de tâches

³ Etre gelé signifie être sous l'emprise de drogues

Ambigüe	C'est une horloge sans heures, elle est comme une neuve
Ambigüe	Le singe a beaucoup de poils, c'est pratique pour la cuisine
Ambigüe	La banque lui a accordé un prêt, ainsi il va gagner beaucoup d'herbes
Ambigüe	Il joue avec la balle, elle est sortie du revolver
Ambigüe	J'ai donné du fromage à ma souris, ça a cassé l'ordinateur
Ambigüe	J'ai écouté attentivement son comte, puis c'était au tour du baron
Ambigüe	Il avait vraiment trouvé l'inspiration, mais n'a jamais eu l'expiration
Ambigüe	Au soleil, on était sous les rayons, ils étaient pleins de biscuits
Ambigüe	Le chevalier n'avait qu'une mule parce qu'il avait perdu ses bottes
Ambigüe	Les enfants ont vu leur mère, elle remplace l'océan
Ambigüe	Il voulait l'emmener à l'hôtel, mais n'a pas pu entrer dans l'église
Ambigüe	Le mime s'était bien blanchi, plus personne n'avait de soupçons
Ambigüe	L'aventurier était dans les bois, coincé sur l'orignal ⁴
Ambigüe	Il est dans un autre état, en fait c'est un liquide
Ambigüe	Il va souvent à la pêche, elle préfère la nectarine
Ambigüe	Elle est folle de scotch, en fait elle aime les collages
Ambigüe	Le maçon a posé les murs, donc on verra bientôt notre confiture

⁴ L'orignal est un animal sauvage, très représentatif du Québec, il porte de très grands bois.