

UNIVERSITÉ DU QUÉBEC À MONTRÉAL

ÉTUDE OCULOMÉTRIQUE SUR LA LECTURE DES CONSTRUCTIONS COMPARATIVES ET ADVERSATIVES

MÉMOIRE

PRÉSENTÉ

COMME EXIGENCE PARTIELLE

DE LA MAÎTRISE EN LINGUISTIQUE

PAR

GHYSLAIN CANTIN-SAVOIE

DÉCEMBRE 2025

UNIVERSITÉ DU QUÉBEC À MONTRÉAL
Service des bibliothèques

Avertissement

La diffusion de ce mémoire se fait dans le respect des droits de son auteur, qui a signé le formulaire *Autorisation de reproduire et de diffuser un travail de recherche de cycles supérieurs* (SDU-522 – Rév.12-2023). Cette autorisation stipule que «conformément à l'article 11 du Règlement no 8 des études de cycles supérieurs, [l'auteur] concède à l'Université du Québec à Montréal une licence non exclusive d'utilisation et de publication de la totalité ou d'une partie importante de [son] travail de recherche pour des fins pédagogiques et non commerciales. Plus précisément, [l'auteur] autorise l'Université du Québec à Montréal à reproduire, diffuser, prêter, distribuer ou vendre des copies de [son] travail de recherche à des fins non commerciales sur quelque support que ce soit, y compris l'Internet. Cette licence et cette autorisation n'entraînent pas une renonciation de [la] part [de l'auteur] à [ses] droits moraux ni à [ses] droits de propriété intellectuelle. Sauf entente contraire, [l'auteur] conserve la liberté de diffuser et de commercialiser ou non ce travail dont [il] possède un exemplaire.»

REMERCIEMENTS

Je tiens à remercier mes parents, Hélène Cantin et Patrice Savoie, pour leur support moral et financier pendant mes études, ainsi que pour la vie.

Je tiens à remercier ma douce et bien-aimée blonde et conjointe, Florence Marcotte, pour son soutien et légendaire patience durant ma maîtrise.

Je tiens à remercier Herman Goulet-Ouellet, dont la vive flamme intellectuelle m'a chauffé les fesses et a embrasé mon retour aux études.

Je tiens à remercier Christine Despaties, conseillère au SASESH de l'UQAM, qui avec toutes nos rencontres a amoindri le choc de mon retour aux études et m'a appris que quand on prend le temps d'en contempler les morceaux, une avalanche, c'est juste un casse-tête.

Je tiens à remercier Yoann Léveillé qui m'a aidé à cultiver - et conserver - mon intérêt pour la recherche et la linguistique ; l'appel à la compréhension de la parlure que l'on parle et le langage que l'on langue, cell phone.

Je tiens à remercier Charlie Beaulieu dont l'amitié fut un gnomon alors que la Terre ne tournait plus rond, non seulement pendant la réponse à la co-vid, mais encore aujourd'hui.

Je tiens à remercier Francis-Alexandre Lepage qui, se tenant droit face aux souffles de la vie durant son propre retour aux études, a remis de l'essence dans mon moteur et m'a ainsi aidé à me souvenir d'où vont les choses : par en avant.

Je tiens à remercier Alexandra Elbakian, récipiendaire d'un prix Nobel qui sera reconnu rétroactivement, pour l'accès à sa précieuse bibliothèque.

Je tiens à remercier Richard Compton, qui m'a pris sous son aile comme assistant de recherche en syntaxe lorsque j'en étais seulement à ma troisième session de baccalauréat, et m'a ainsi permis d'entrer directement dans le monde de la recherche.

Dernièrement mais non moindrement, je tiens à remercier mes directeurs de recherche, Denis Foucambert et Grégoire Winterstein, sans qui ce projet n'aurait pas eu lieu. Spécifiquement :

- Je tiens à remercier Denis Foucambert pour m'avoir introduit à l'univers de l'étude de la lecture et de l'oculométrie, ainsi que pour m'avoir donné la responsabilité et l'honneur de mettre sur place le laboratoire d'oculométrie du département de linguistique de l'UQAM : l'oculocal.
- Je tiens à remercier Grégoire Winterstein pour m'avoir compté parmi - et confié - les premiers balbutiements du SLIC (Sémantique et Linguistique Informatique et Computationnelle). Je tiens aussi à remercier Grégoire pour le pied-de-biche qu'il a utilisé sur la porte de mon esprit, principalement à l'aide de ses cours. Mes capacités de programmation dont il a guidé le développement me permettront assurément de générer des cordes (ou spaghettis, dépendant des points de vue 😊) à ajouter à mon arc.

TABLE DES MATIÈRES

TABLE DES FIGURES	vii
LISTE DES TABLEAUX	ix
RÉSUMÉ	xiv
CHAPITRE 1 CADRE THÉORIQUE	1
1.1 Connecteur adversatif <i>mais</i>	1
1.1.1 Approche contrastive formelle.....	2
1.1.2 Approche inférentielles	6
1.1.3 Les adjectifs gradables	10
1.1.4 Approches d'ambiguïté.....	17
1.1.5 Récapitulatif	18
1.2 Oculométrie	19
1.2.1 Mouvements oculaires	20
1.2.2 Immobilité oculaire : La fixation	21
1.2.3 Particularités de la lecture	23
CHAPITRE 2 MÉTHODE	27
2.1 Participants	27
2.1.1 Recrutement	27
2.2 Items	28
2.2.1 Contraintes globales.....	28
2.2.2 Les contextes	29
2.2.3 Les phrases.....	33
2.2.4 Questions	35

2.3	Mémoire de travail	36
2.3.1	Matériel pour test de mémoire de travail	37
2.3.2	Matériel	38
CHAPITRE 3	ANALYSES	40
3.1	Transformation et élimination des données	40
3.2	Résultats au niveau phrastique	41
3.2.1	Définition des variables au niveau phrastique	42
3.2.2	Considération des facteurs aléatoires	43
3.2.3	Sélection du modèle au niveau phrastique.....	45
3.2.4	La probabilité de régression	45
3.2.5	La durée de fixation totale	47
3.3	Second niveau d'analyse : Les mots	49
3.3.1	Durée du parcours régressif.....	51
3.3.2	La durée de la première fixation	55
3.4	Durée du premier regard.....	58
3.4.1	Probabilité de régressions sortantes	60
3.4.2	Probabilité de régression entrante	63
CONCLUSION		66
3.5	Description de la différence entre <i>mais moins</i> et <i>mais plus</i>	66
3.6	Limites	70
ANNEXE A	ITEMS EXPÉRIMENTAUX POUR LE TEST DE MÉMOIRE DE TRAVAIL	72
ANNEXE B	ITEMS EXPÉRIMENTAUX DE L'EXPÉRIENCE PRINCIPALE	79
ANNEXE C	TABLEAUX D'ANALYSES	96

C.1	Tableaux pour la durée du parcours régressif	96
C.1.1	Par zones	96
C.1.2	Par mots de la zone critique	99
C.2	Tableaux pour la durée de la première fixation	103
C.2.1	Par zones	103
C.2.2	Par mots de la zone critique	106
C.3	Tableaux pour la durée du premier regard	110
C.3.1	Par zones	110
C.3.2	Par mots de la zone critique	113
C.4	Tableaux pour la probabilité de fixation	117
C.4.1	Par zones	117
C.4.2	Par mots de la zone critique	120
C.5	Tableaux pour probabilité de régression reçue	124
C.5.1	Par zones	124
C.5.2	Par mots de la zone critique	127
C.6	Tableaux pour la probabilité de régression causée	131
C.6.1	Par zones	131
C.6.2	Par mots de la zone critique	134
	BIBLIOGRAPHIE	139

TABLE DES FIGURES

Figure 1.1	Échelle sur la dimension <i>taille</i> , allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d0	11
Figure 1.2	Échelle sur la dimension <i>taille</i> , allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d0, en comparaison avec le but à balise minimale, indiqué par l'espace en jaune.	13
Figure 1.3	Échelle sur la dimension <i>taille</i> , allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d0, en comparaison avec le but à balise maximale, indiqué par l'espace en jaune.	15
Figure 1.4	Échelle sur la dimension <i>taille</i> , allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d0	16
Figure 1.5	Représentation de l'axe visuel lors de 2 fixations distinctes, puis de l'angle entre les deux axes comme mesure de distance	22
Figure 1.6	Schématisation du champs perceptuel, montrant son asymétrie adaptée au système d'écriture.....	23
Figure 3.1	Probabilités de régression, séparées par conditions	46
Figure 3.2	Durée de fixation en millisecondes, séparée selon les variables Valence et Contexte	48
Figure 3.3	Durée du parcours régressif en millisecondes sur tous les mots	52
Figure 3.4	Moyenne de la durée du parcours régressif en millisecondes pour chaque zone, un fond blanc signifiant un effet significatif de la Valence	54
Figure 3.5	Durée du parcours régressif sur chaque mot de la zone critique, un fond blanc signifiant un effet significatif de la Valence	55
Figure 3.6	Durée de la première fixation en millisecondes pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables Valence et Contexte	56
Figure 3.7	Durée du premier regard en millisecondes pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables Valence et Contexte	58
Figure 3.8	Durée du premier regard en millisecondes pour chaque mot de la zone critique, un fond blanc signifiant un effet significatif de la Valence	60

Figure 3.9	Probabilité d'être le point de départ d'une régression oculaire pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables Valence et Contexte	61
Figure 3.10	Moyenne de la probabilité d'être le point de départ d'une régression oculaire pour chaque zone, un fond blanc signifiant un effet significatif de la Valence	63
Figure 3.11	Probabilité d'être le point d'arrivée d'une régression oculaire pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables Valence et Contexte	64
Figure 3.12	Lieu des effets de durées sur un exemple de phrase expérimentale	67

LISTE DES TABLEAUX

Table 1.1	Prédictions de la félicité de phrases avec <i>mais</i> selon les approches de (Sæbø, 2003) (pS) et de (Umbach, 2005) (pU) comparées aux résultats obtenus dans l'expérience de (Winterstein et al., 2014) (pW), la couleur verte indiquant une corrélation avec les résultats, la couleur rouge indiquant une différence	4
Table 1.2	Prédictions de la félicité de phrases avec <i>mais</i> selon les approches corrigées de (Sæbø, 2003) (pS) et de (Umbach, 2005) (pU), comparées aux résultats obtenus dans l'expérience de (Winterstein et al., 2014) (pW), la couleur verte indiquant une corrélation avec les résultats, la couleur rouge indiquant une différence	5
Table 1.3	Prédictions de la félicité de phrases avec <i>mais moins</i> ou <i>mais plus</i> selon le type de buts argumentatifs.	17
Table 2.1	Répartition des mots par zone	34
Table 2.2	Répartition des conditions selon la version	39
Table 3.1	Résultats du modèle pour la probabilité de régression	47
Table 3.2	Résultats du modèle pour le logarithme de la durée de fixation en millisecondes pour toute la phrase	49
Table 3.3		50
Table 3.4	Résultats du modèle pour la durée du parcours régressif sur tous les mots.	53
Table 3.5	Résultats du modèle pour le logarithme de la durée de première fixation par mot, considérant tous les mots de l'échantillon	57
Table 3.6	Résultats du modèle pour la durée du premier regard sur tous les mots.	59
Table 3.7	Résultats du modèle pour le logarithme de la probabilité d'être le point de départ d'une régression oculaire	62
Table 3.8	Résultats du modèle pour le logarithme de la probabilité d'être le point d'arrivée d'une régression oculaire	65
Table C.1	Durée du parcours régressif sur les mots de la zone -1	96
Table C.2	Durée du parcours régressif sur les mots de la zone 0	97

Table C.3	Durée du parcours régressif sur les mots de la zone 1	97
Table C.4	Durée du parcours régressif sur les mots de la zone 2	98
Table C.5	Durée du parcours régressif sur les mots de la zone 3	98
Table C.6	Durée du parcours régressif sur les mots de la zone 4	99
Table C.7	Durée du parcours régressif sur le mot à l'index (mais) -3	99
Table C.8	Durée du parcours régressif sur le mot à l'index (mais) -2	100
Table C.9	Durée du parcours régressif sur le mot à l'index (mais) -1	100
Table C.10	Durée du parcours régressif sur le mot à l'index (mais) 0	101
Table C.11	Durée du parcours régressif sur le mot à l'index (mais) 1	101
Table C.12	Durée du parcours régressif sur le mot à l'index (mais) 2	102
Table C.13	Durée du parcours régressif sur le mot à l'index (mais) 3	102
Table C.14	Durée du parcours régressif sur le mot à l'index (mais) 4	103
Table C.15	Durée de la première fixation sur les mots de la zone -1	103
Table C.16	Durée de la première fixation sur les mots de la zone 0	104
Table C.17	Durée de la première fixation sur les mots de la zone 1	104
Table C.18	Durée de la première fixation sur les mots de la zone 2	105
Table C.19	Durée de la première fixation sur les mots de la zone 3	105
Table C.20	Durée de la première fixation sur les mots de la zone 4	106
Table C.21	Durée de la première fixation sur le mot à l'index (mais) -3	106
Table C.22	Durée de la première fixation sur le mot à l'index (mais) -2	107
Table C.23	Durée de la première fixation sur le mot à l'index (mais) -1	107

Table C.24	Durée de la première fixation sur le mot à l'index (mais) 0	108
Table C.25	Durée de la première fixation sur le mot à l'index (mais) 1	108
Table C.26	Durée de la première fixation sur le mot à l'index (mais) 2	109
Table C.27	Durée de la première fixation sur le mot à l'index (mais) 3	109
Table C.28	Durée de la première fixation sur le mot à l'index (mais) 4	110
Table C.29	Durée du premier regard sur les mots de la zone -1	110
Table C.30	Durée du premier regard sur les mots de la zone 0	111
Table C.31	Durée du premier regard sur les mots de la zone 1	111
Table C.32	Durée du premier regard sur les mots de la zone 2	112
Table C.33	Durée du premier regard sur les mots de la zone 3	112
Table C.34	Durée du premier regard sur les mots de la zone 4	113
Table C.35	Durée du premier regard sur le mot à l'index (mais) -3	113
Table C.36	Durée du premier regard sur le mot à l'index (mais) -2	114
Table C.37	Durée du premier regard sur le mot à l'index (mais) -1	114
Table C.38	Durée du premier regard sur le mot à l'index (mais) 0	115
Table C.39	Durée du premier regard sur le mot à l'index (mais) 1	115
Table C.40	Durée du premier regard sur le mot à l'index (mais) 2	116
Table C.41	Durée du premier regard sur le mot à l'index (mais) 3	116
Table C.42	Durée du premier regard sur le mot à l'index (mais) 4	117
Table C.43	Probabilité de fixation sur les mots de la zone -1	117
Table C.44	Probabilité de fixation sur les mots de la zone 0	118

Table C.45	Probabilité de fixation sur les mots de la zone 1.....	118
Table C.46	Probabilité de fixation sur les mots de la zone 2	119
Table C.47	Probabilité de fixation sur les mots de la zone 3	119
Table C.48	Probabilité de fixation sur les mots de la zone 4	120
Table C.49	Probabilité de fixation sur le mot à l'index (mais) -3	120
Table C.50	Probabilité de fixation sur le mot à l'index (mais) -2	121
Table C.51	Probabilité de fixation sur le mot à l'index (mais) -1	121
Table C.52	Probabilité de fixation sur le mot à l'index (mais) 0.....	122
Table C.53	Probabilité de fixation sur le mot à l'index (mais) 1	122
Table C.54	Probabilité de fixation sur le mot à l'index (mais) 2.....	123
Table C.55	Probabilité de fixation sur le mot à l'index (mais) 3.....	123
Table C.56	Probabilité de fixation sur le mot à l'index (mais) 4.....	124
Table C.57	Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone -1....	124
Table C.58	Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 0	125
Table C.59	Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 1.....	125
Table C.60	Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 2	126
Table C.61	Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 3	126
Table C.62	Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 4	127
Table C.63	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais -3..	127
Table C.64	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais -2..	128
Table C.65	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais -1 ..	128

Table C.66	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 0 ..	129
Table C.67	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 1 ...	129
Table C.68	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 2 ..	130
Table C.69	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 3 ..	130
Table C.70	Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 4 ..	131
Table C.71	Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone -1...	131
Table C.72	Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 0 ...	132
Table C.73	Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 1...	132
Table C.74	Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 2 ...	133
Table C.75	Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 3 ...	133
Table C.76	Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 4 ...	134
Table C.77	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais -3	134
Table C.78	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais -2	135
Table C.79	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais -1.	135
Table C.80	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 0 .	136
Table C.81	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 1..	136
Table C.82	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 2 .	137
Table C.83	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 3 .	137
Table C.84	Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 4 .	138

RÉSUMÉ

Ce mémoire vise à explorer la sémantique du connecteur adversatif *mais* en observant les processus de lecture via une méthode oculométrique. 28 participants ont lu des phrases sur un moniteur équipé d'un oculomètre Tobii Pro Fusion 250 Hz. Outre une multitude de leurres, les phrases lues étaient des constructions comparatives avec *mais* (ex. : Paul est grand mais moins grand que Pierre) dont la valence était négative (ex. : mais moins grand) ou positive (ex. : mais plus grand) précédées d'un contexte qui pouvait introduire une condition d'égalité (ex. : Paul doit absolument avoir la même taille que Pierre) ou non. Les effets de la valence et du type de contexte furent explorés sur la durée de fixation et la probabilité de régression oculaire pour chaque phrase, ainsi que sur la durée de première fixation, la durée du premier regard, la durée du parcours régressif (*regression path*), la probabilité de régression entrante et la probabilité de régression sortante sur chaque mots ainsi que sur chaque zone, définies comme groupement de mots séparés de façon à considérer *mais* comme faisant partie de la zone centrale (ex. : [Paul est grand], [mais moins grand], [que Pierre]). Le niveau de scolarité et le résultat d'un test de mémoire de travail pour chaque individus, le nombre de caractères et la fréquence lexicale des mots utilisés ont également été prélevées à des fins de contrôle. Les résultats montre un effet de la valence, indiquant que *mais plus* entraîne plus de difficultés de lectures que *mais moins*, et aucun effet significatif du type de contexte.

CHAPITRE 1

CADRE THÉORIQUE

Dans cette section se trouvent les outils et considérations théoriques nécessaires à l'exécution et à la compréhension de ce projet. Il s'y trouve premièrement une description du connecteur adversatif *mais*, selon différentes approches sémantiques. Ensuite, cette section contient les outils théoriques utilisés en lien avec les techniques oculométriques qui informeront les mesures prélevées lors de notre expérience.

1.1 Connecteur adversatif *mais*

Mais est considéré comme une conjonction de coordination. Tel que décrit dans l'introduction en ??, elle lie ensemble deux propositions qui doivent être opposées. Plusieurs types d'usages sont possibles, qui varient notamment selon le nombre d'entités et le nombre de propriétés que l'on oppose. L'usage qui nous intéresse dans le cadre de ce mémoire est l'usage comparatif, tel qu'exemplifié en (1).

(1) Paul est grand, mais moins grand que Pierre.

Ici, le degré d'atteinte d'une propriété gradable (*grand* dans l'exemple (1)) est comparé vis-à-vis de deux entités distinctes, *Paul* et *Pierre*.

Ci-dessous, on y décrit les deux types d'approches principales de la sémantique de *mais* : les approches contrastives formelles et les approches inférentielles. Pour chacune, le fonctionnement global de l'approche est exposé, suivi d'une tentative de l'utiliser pour expliquer la sémantique des phrases comparatives adversatives, comme celle en (1)

Ensuite, les caractéristiques intrinsèques des adjectifs gradables sont approfondies, et l'intégration des usages comparatifs aux deux approches précédentes est de nouveau faite, cette fois-ci avec plus de succès.

Finalement, une récapitulation des éléments vus dans cette section est proposée, suivie de quelques prédictions sur la nature des résultats escomptés, informées par nos différentes approches théoriques.

1.1.1 Approche contrastive formelle

La première chose qui caractérise l'approche contrastive formelle est que l'usage contrastif de *mais* est considéré comme son usage prototypique (Winterstein, 2012). Les usages contrastifs de *mais* sélectionnent deux conjoints selon des contraintes spécifiques. Le premier énoncé doit contenir deux unités sémantiques de type différent, comme l'entité et le **prédicat** qu'on peut voir dans le premier conjoint de (2-a). Le deuxième énoncé conjoint a les mêmes contraintes, formant ainsi deux paires d'unités sémantiques; dans le cas de (2-a), il s'agit de la paire d'entités <Lemmy, Ritchie> et de la paire de prédicats <**jouer de la basse**, **jouer de la guitare**>.

- (2) a. Lemmy **joue de la basse**, mais Ritchie **joue de la guitare**.
b. Ritchie **joue de la guitare**, mais Lemmy **joue de la basse**.
c. Lemmy **joue de la basse** et Ritchie **joue de la guitare**.

L'usage contrastif est dit "symétrique", c'est-à-dire qu'inverser l'ordre des propositions conjointes, comme en (2-b) affecte peu ou pas du tout le sens véhiculé par la phrase. De même, remplacer *mais* par *et* comme en (2-c) affecte en apparence très peu le sens.

La deuxième particularité des approches contrastives formelles est que l'opposition nécessaire pour la félicité de l'utilisation de *mais* est comprise comme émergeant de la sémantique des propriétés lexicales des mots contenus dans les propositions conjointes (Winterstein, 2012).

Les notions de *focus* et de *topic*, implicites à la structure informationnelle de la phrase, peuvent être utilisées pour expliquer l'opposition entre les deux propositions conjointes par *mais* (Umbach, 2005). En remarquant que certains mots en anglais sont prononcés avec une certaine mise en relief, plusieurs chercheurs se sont penchés sur l'apport sémantique de ce qu'ils appellent *focus*. En se basant sur les travaux de Jackendoff (1972), Rooth (1985) propose qu'un focus invoque un ensemble (**P-set**) contenant la propriété mise en relief ainsi qu'une alternative par rapport à une question à laquelle la proposition répond. Par exemple, imaginons une situation contenant 4 personnes, John, Bill, Sue et Marie, et contemplons l'exemple en anglais suivant.

- (3) John introduced Bill to Sue_F.

- a. {John introduced Bill to Sue, John introduced Bill to Mary}
- b. Who did John introduce Bill to ?

La mise en relief de *Sue* en (3) signifie que John n'a pas présenté Bill à Marie, mais bien à Sue. Une façon de formaliser cet apport sémantique particulier est de considérer les deux possibilités comme faisant partie d'un ensemble comme en (3-a), dont un seul objet est sélectionné. Les propositions alternatives présentes dans l'ensemble sont reliées aux questions auxquelles la phrase initiale répond, et donc au sens de la phrase (Von Steutterheim et Klein, 1989; Büring, 2019). Certains connecteurs, comme *always* et *only*, sont sensibles aux propriétés des différents types de focus (Beaver et Clark, 2003). Si on considère que *mais* l'est tout autant, l'opposition entre les deux propositions conjointes peut être expliquée en observant le contenu des ensembles d'alternatives (Alt-set), dont on peut voir un exemple en (4).

(4) Paul est grand, mais Bill est petit.

- a. **Question** : QUI est grand ?
- b. **Alt-set** : {Paul est grand, Bill est grand}

Umbach (2005) et Sæbø (2003) utilisent ces ensembles pour caractériser les conditions d'emploi de *mais*. Leurs méthodes pour arriver à leurs conditions de félicité sont différentes, mais le résultat est très similaire, comme on peut voir en (5).

(5) Paul est grand, mais Bill est petit.

- a. **Alt-Set** {Paul est grand, Bill est grand}
- b. **Condition Sæbø** : *Bill n'est pas grand.* doit être vrai.
- c. **Condition Umbach** : *Bill est grand.* doit être faux.

Dans le cas de (5), si Bill est petit, alors il est vrai que Bill n'est pas grand, ainsi la condition pour l'acceptabilité de *mais* proposée par Sæbø, en (5-b) est remplie, prédisant une utilisation correcte de la phrase en (5), ce qui est le cas. Similairement, si Bill est petit, il est faux que Bill est grand, ainsi la condition d'Umbach, exposée en (5-c) prédit tout autant correctement une utilisation acceptable de *mais*. Puisqu'il s'agit d'une

distinction sans différences, nous n'utiliserons pour les exemples suivants que la condition d'évaluation de Sæbø. Observons maintenant les prédictions de ces approches vis-à-vis du type de phrases avec *mais* qui nous intéresse, les comparatives, en (6).

- (6) Paul est grand mais plus/moins grand que Pierre.
- a. **Alt-Set** {Paul est grand, Pierre est grand}
 - b. **Condition Sæbø** : *Pierre n'est pas grand*. doit être vrai.

Selon les résultats obtenus dans l'expérience de (Winterstein *et al.*, 2014), les prédictions correctes sont que lorsque la phrase comparative est construite avec *mais moins*, elle sera acceptée alors qu'elle ne le sera pas lorsqu'elle est construite avec *mais plus*. Les méthodes de (Sæbø, 2003) et (Umbach, 2005) demandent toutes deux que *Pierre n'est pas grand* soit vraie dans le contexte d'une phrase comme (6), ce qui n'est le cas ni avec *moins* ni avec *plus*, tel qu'exposé dans le tableau 1.1 qui contient les prédictions des différentes méthodes.

Phrase	pS	pU	pW
Paul est grand, mais moins grand que Pierre	non	non	oui
Paul est grand, mais plus grand que Pierre	non	non	non

Table 1.1 Prédications de la félicité de phrases avec *mais* selon les approches de (Sæbø, 2003) (pS) et de (Umbach, 2005) (pU) comparées aux résultats obtenus dans l'expérience de (Winterstein *et al.*, 2014) (pW), la couleur verte indiquant une corrélation avec les résultats, la couleur rouge indiquant une différence

Ces prédictions reposent sur l'hypothèse que l'ensemble d'alternatives émerge nécessairement de la question "QUI est grand?" et prédisent uniformément l'inacceptabilité des comparatives, ce qui ne rend pas compte des résultats de l'expérience de (Winterstein *et al.*, 2014), qui constate une utilisation correcte de *mais* dans le cas des phrases avec *mais moins* et une utilisation incorrecte lorsqu'il y a *mais plus*, pour des phrases hors contexte.

Pour corriger les divers problèmes que leurs méthodes respectives rencontrent, (Sæbø, 2003) et (Umbach, 2005) montrent une ouverture à l'idée que la paire d'alternatives soit générée autour d'un autre aspect de la sémantique de la phrase. Dans le cas de (Sæbø, 2003), il pourrait s'agir de la propriété *grand*, formant

l'ensemble d'alternatives {*Paul est grand*, *Paul est non-grand*} qui sont des réponses à la question implicite "Paul est QUOI?". Dans le cas de (Umbach, 2005), la négation est implicite à *mais* et enclenche donc, en se basant sur les travaux de (Givón, 1987), l'ensemble d'alternatives contenant ladite négation ainsi que sa forme positive, donc l'ensemble {*Paul est grand*, *mais plus grand que Pierre*; *Paul est grand et aussi plus grand que Pierre*}. Permettons-nous donc de réévaluer leurs prédictions selon leurs corrections.

- (7) Paul est grand, mais moins/plus grand que Pierre.
- a. **Alt-Set Sæbø** : {Paul est grand, Paul est non-grand}
 - b. **Condition Sæbø** : *Paul n'est pas non-grand* (Paul est grand) doit être vraie.
 - c. **Alt-set Umbach** : {Paul est grand, mais moins/plus grand que Pierre; Paul est grand, et aussi plus/moins grand que Pierre}
 - d. **Condition Umbach** : *Paul est grand, et aussi plus grand que Pierre* doit être faux.

Comme il peut être observé dans le tableau 1.2, la correction de (Sæbø, 2003) est trop permissive et ne prédit donc pas bien l'inacceptabilité des phrases avec *mais plus*, alors que la correction de (Umbach, 2005) reste trop contraignante, ne permettant pas de rendre compte de l'acceptabilité des phrases comparatives avec *mais moins*.

Phrase	pS	pU	pW
Paul est grand, mais moins grand que Pierre	oui	non	oui
Paul est grand, mais plus grand que Pierre	oui	non	non

Table 1.2 Prédications de la félicité de phrases avec *mais* selon les approches corrigées de (Sæbø, 2003) (pS) et de (Umbach, 2005) (pU), comparées aux résultats obtenus dans l'expérience de (Winterstein *et al.*, 2014) (pW), la couleur verte indiquant une corrélation avec les résultats, la couleur rouge indiquant une différence

Les questions implicites "QUI est grand?", "Paul est QUOI?" et la négation implicite à l'intérieur de *mais* ne sont pas à même de générer l'ensemble d'alternatives permettant d'expliquer l'opposition présente dans les phrases comparatives avec *mais*. En 1.1.3, je décris des propriétés intrinsèques des adjectifs gradables, comme *grand*, qui peuvent donner lieu à la génération des ensembles d'alternatives permettant à l'approche contrastive formelle de bien expliquer l'opposition présente dans les usages comparatifs de *mais*.

1.1.2 Approche inférentielles

Les approches inférentielles se définissent par deux caractéristiques : (1) l'usage canonique théorisé est celui de **déni d'attente**, et (2) l'opposition ne se produit pas entre les caractéristiques sémantiques lexicales des mots contenus dans les deux propositions conjointes, mais bien en lien avec une troisième proposition qui doit être inférée du contexte ou d'ailleurs.

Contrairement à l'usage contrastif, l'usage de déni d'attente n'a pas de contraintes structurelles strictes vis-à-vis du nombre de propriétés ou de prédicats qu'il contient. Plutôt, il s'agit d'une contrainte sémantique : le premier conjoint doit introduire une attente que le deuxième conjoint nie, comme on peut voir en (8-a), avec l'entité <Édith> et les propriétés <être végétarienne, manger du bœuf>.

- (8) a. Édith est végétarienne, mais elle mange du bœuf.
- b. Une végétarienne ne mange normalement pas de bœuf.

On peut s'attendre d'une végétarienne, normalement, qu'elle ne mange pas de bœuf. Cette attente est niée par le deuxième énoncé conjoint par *mais* dans la phrase en (8-a). Puisque l'attente en (8-b) est générée directement par la définition de ce qu'est une végétarienne, qui se trouve directement dans la phrase en (8-a), on parle ici de déni d'attente direct. Certaines attentes "directes" ne sont pas si normales cependant, comme celle en (9-b) entraînée par la phrase en (9-a).

- (9) a. **Phrase** : Il est gros, mais il est intelligent.
- b. **Attente** : Les gens gros ne sont pas intelligents.

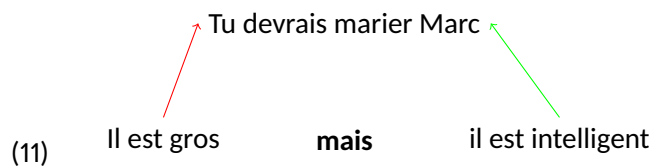
Quelqu'un disant la phrase en (9-a), hors contexte, sera perçu comme prenant l'attente en (9-b) comme acquis, et donc comme quelqu'un de quelque peu méchant¹. Or, cette même phrase peut aussi s'agir d'un usage de déni d'attente **indirecte** (aussi nommé **argumentatif**) dans le bon contexte, comme celui en (10-a).

1. Cette attente est forcée par la nécessité pour les conjoints de *mais* de partager une dimension d'échelle, en vertu du processus décrit un peu plus tard en 1.1.3.

- (10) a. **Contexte** : Marie considère ses éventuelles épousailles avec Marc. Tentant de l'aider à prendre une décision éclairée, son amie Gertrude s'exclame :
- b. **Phrase** : Il est gros, mais il est intelligent.
- c. **Attente** : Marc n'est pas bon à marier, car il est gros.

Tout comme (9-a), l'attente en (10-c) expose certains préjugés de la part de Gertrude. Cependant, contrairement aux phrases en (9-a) et en (8-a), le deuxième conjoint de la phrase (10-b) n'est pas la négation de l'attente en (10-c) ².

Une façon plus claire de percevoir ce qui se passe ici, permise par les approches inférentielles, est de considérer une proposition tierce inférée à l'aide du contexte, comme en (11).



Tel qu'indiqué par la flèche rouge, *il est gros* est en opposition à l'énoncé inféré *Tu devrais marier Marc* que l'on désigne l'énoncé **pivot**, alors que *il est intelligent*, tel qu'indiqué par la flèche verte, l'appuie. Même si les approches inférentielles s'accordent par rapport à cette vision globale de la sémantique de *mais*, c'est dans deux détails que les différences entre les approches émergent : (1) La nature du lien d'opposition (ou d'affirmation) entre les conjoints et le pivot ainsi que (2) le processus d'accès à l'énoncé pivot.

La théorie de la pertinence (*Relevance Theory*) (Blakemore, 1989; Blakemore, 2002; Sperber et Wilson, 1986), tel que cité dans (Winterstein, 2012), considère que *mais* introduit une proposition qui doit contredire et éliminer une assomption rendue disponible par la proposition précédente. Une des faiblesses de cette

2. Notons ici que l'approche contrastive formelle vue précédemment rencontre elle aussi un problème substantiel lorsque confronté à ce type d'usage. Pour que l'opposition en (8-a) soit traité comme un déni d'alternative, il faudrait imaginer que le premier conjoint, constituant intrinsèquement une réponse à la question **Quelles sont les raisons de ne pas épouser Marc ?**, génère l'ensemble d'alternatives {Il est gros, il est stupide} duquel la deuxième entité sera niée par le deuxième conjoint de la phrase. Pour justifier cette approche cependant, il faudrait démontrer l'impact de cette question implicite dans d'autres situations, ce que je me sens incapable de faire, et n'essaierai donc pas.

stratégie est que certaines assomptions accessibles peuvent être niées avec *mais*, mais ne résultent pas nécessairement en une phrase acceptable, comme on peut voir en (12).

- (12) a. **Premier conjoint** : Jean a résolu quelques problèmes.
b. **Assomption accessible** : Jean ne les a pas tous résolus.
c. **Phrase avec *mais*** : *Jean a résolu quelques problèmes, mais tous.

L'implicature de quantité en (12-b) est rendue accessible (Carston et Uchida, 1998) au deuxième conjoint par le premier en (12-a), mais sa contradiction ne résulte pas en une phrase acceptable, comme on peut voir en (12-c). Ainsi, le processus de sélection du pivot se doit d'être un peu plus restrictif, ce que l'approche proposée par la théorie de l'argumentation dans la langue (AdL) (Anscombe et Ducrot, 1983) permet.

Selon cette théorie, les gens parlent avec une intention, et les locuteurs au moment de l'énonciation ont un but. Par exemple, l'énoncé en (13-a) a plus souvent le sens d'un des buts en (13-b) que celui d'une description vériconditionnelle des faits comme en (13-c).

- (13) a. Il commence à faire noir.
b. {Allumes la lumière s'il-te-plaît, Nous devrions rentrer à la maison sous peu}
c. La quantité de lux caractérisant la luminosité de cet endroit à cet instant précis est inférieure à la quantité de lux pour le même endroit lors d'un instant antérieur.

Ainsi, une phrase comprise est une phrase dont un but argumentatif a été inféré. Plutôt que de référer à la description en (13-c), la phrase en (13-a) évoquera un ensemble de buts argumentatifs potentiels. Selon cette vision, comprendre une phrase, c'est inférer de l'ensemble des buts possibles celui qui était envisagé par le locuteur. Ici, *mais* est interprété comme un connecteur sensible au but argumentatif qui sélectionne ses conjoints de façon à ce que le premier argumente en faveur d'une proposition pivot, et le deuxième argumente contre (Anscombe et Ducrot, 1977).

- (14) La voiture est belle, mais elle est chère.

- a. Je vais acheter la voiture.

Le premier conjoint de la phrase en (14) argumente en faveur du pivot en (14-a), alors que le deuxième conjoint argumente contre, nous dirons donc que les deux énoncés joints par *mais* doivent être **contre-orientés**. La nature précise de cette opposition est cependant ouverte à l'interprétation et nécessite l'application d'un formalisme. Par exemple, si nous faisons preuve d'imagination en utilisant celui de (Sæbø, 2003) et (Umbach, 2005) en considérant le but argumentatif comme étant assez intrinsèque à la sémantique de l'énoncé pour constituer une alternative dans l'ensemble d'alternatives, il est possible d'obtenir une prédiction favorable, tel que démontré ci-dessous en (15).

(15) La voiture est belle, mais elle est chère

- a. Je vais acheter la voiture
- b. **Alt-set** : {la voiture est belle, je vais acheter la voiture}
- c. **Évaluation (Umbach)** : *Je vais acheter la voiture* doit être faux.

La phrase en (15) a un but argumentatif, en (15-a), autour duquel ses conjoints s'opposent. Selon l'interprétation d'Umbach en (15-c), le deuxième élément de l'ensemble d'alternatives en (15-b) doit être faux pour que la phrase en (15) soit jugée acceptable. Puisque *je vais acheter la voiture* n'est pas nécessairement faux, cette méthode prédirait incorrectement une utilisation inacceptable de *mais*. Cependant, soulevons ici qu'une certitude vis-à-vis de *je vais acheter la voiture*, exposée en (16-b) serait en effet fausse, menant à une prédiction correcte de la méthode contrastive formelle.

- (16) a. S'il y a une chose dont je suis certain dans la vie, c'est que je n'achèterai jamais cette voiture. Elle est belle, mais elle est chère.
- b. ? S'il y a une chose dont je suis certain dans la vie, c'est que j'achèterai cette voiture. Elle est belle, mais elle est chère.

La méthode contrastive formelle, se servant du but argumentatif comme alternative, ne rend pas compte convenablement de l'incertitude évoquée par la proposition "*Elle est belle, mais elle est chère*" lorsque

celle-ci est prononcée hors contexte.

Or, en voyant l'argumentation sous un regard probabiliste, on peut concevoir qu'argumenter pour une proposition, c'est augmenter sa probabilité alors qu'argumenter contre, c'est diminuer sa probabilité (Merin, 1999). En effet, *elle est belle* augmente la probabilité que quelqu'un achète la voiture, alors que *elle est chère* la diminue. Cette vision permet de rendre compte du caractère incertain de la phrase soulevé en (15). De plus, puisque le deuxième conjoint de la phrase en (15), *elle est chère*, diminue la probabilité du but argumentatif, cette probabilité ne peut plus être 1, ce qui exclue la certitude et donc explique le caractère étrange de la séquence en (16-b).

Ainsi, une phrase acceptable avec *mais* en est une permettant un pivot tel que (1) le premier conjoint en augmente la probabilité et (2) le deuxième conjoint en diminue la probabilité. Dans la section 1.1.3 qui suit se trouve une précision des caractéristiques des adjectifs gradables qui permettent d'expliquer qu'une moins grande quantité de ces pivots sont possibles avec *mais plus grand*, ce qui justifie la différence au niveau de l'acceptabilité et du temps de traitement.

1.1.3 Les adjectifs gradables

Trois caractéristiques sont attribuées aux adjectifs gradables : (1) une dimension, (2) un degré et (3) une relation d'ordre (Kennedy et McNally, 2005). Prenons l'exemple de *Paul est grand*, en comparant avec *moins grand que Pierre* et *plus grand que Pierre* en 1.1.

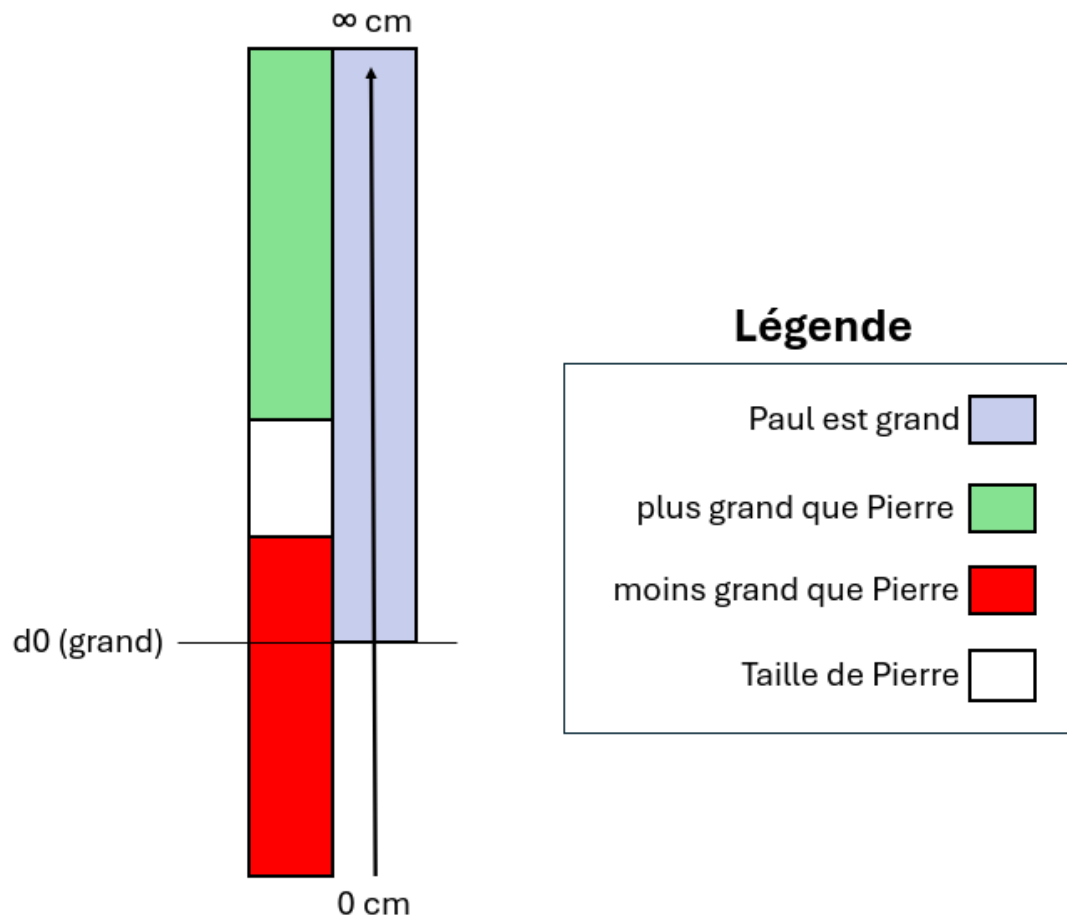


Figure 1.1 Échelle sur la dimension *taille*, allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d0

L'attribution à *Paul* de l'adjectif *grand* est par le fait même l'attribution à Paul d'un degré sur une échelle de taille tel que sa taille est plus grande que d0, soit le degré en dessous duquel une chose n'est pas grande et au-dessus duquel une chose est grande. Ce degré varie selon les situations et le contexte, par exemple, il n'est pas le même si l'on parle de grandes fourmis ou de grandes baleines. Quoiqu'il en soit, *grand* signifie "au-dessus de d0", tel que montré par le rectangle bleu dans la figure 1.1.

Aussi, cette conceptualisation des adjectifs gradables permet d'observer que l'espace que la taille potentielle de Paul occupe (en bleu dans la figure 1.1) et l'espace occupé par *plus grand que Pierre* (en vert) partagent une infinité de degrés sur l'échelle de taille. Ainsi, la probabilité que la taille de Paul insinuée par *Paul*

est grand soit dans l'intervalle impliqué par *plus grand que Pierre* est a priori infiniment plus grande que celle que la taille de Paul soit hors de cette intervalle, et donc inférieure à celle de Pierre. Puisque *Paul est grand* et *plus grand que Pierre* partagent autant d'espace sur l'échelle des tailles auxquelles ces propositions peuvent faire référence, une opposition entre ces deux propositions semble d'emblée moins plausible.

De plus, cette façon de voir les choses permet à l'approche contrastive formelle, vue en 1.1.1 de bien prédire l'acceptabilité des comparatives avec *mais* en utilisant l'aspect intrinsèquement comparatif des adjectifs gradables pour générer l'ensemble {Paul est plus grand que d0, Paul est plus grand que Pierre}, qui seraient des réponses alternatives à la question "plus grand que QUOI?". Cependant, cette méthode n'explique pas pourquoi, dans certains contextes, comme en ??, rappelés ci-dessous en (17), les comparatives avec *mais plus* sont acceptables; puisqu'elle utilise seulement les morceaux de la phrase, elle ne peut prédire que l'acceptabilité de la phrase seule, alors qu'une approche inférentielle n'a pas cette limite.

- (17) **Contexte** : On cherche un cascadeur pour doubler Pierre, un acteur, dans des scènes d'action. Il faut que le cascadeur ait exactement la même taille que Pierre pour qu'on ne remarque pas la doublure à l'écran. C'est difficile parce que Pierre est de grande taille.
- a. **Phrase** : Au niveau de la carrure, Paul est grand, mais plus grand que Pierre et il n'est pas facile de prendre une décision.
 - b. **Opposition impliquée** : Paul est peut-être le cascadeur recherché, mais probablement pas.

Pour **AdL**, une approche inférentielle vue en 1.1.2, les énoncés évoquent un ensemble de buts potentiels dont l'inférence permet la compréhension du sens de la phrase. Utilisant encore notre exemple de *Paul est grand*, *mais moins*/**plus grand que Pierre*, il est possible d'expliquer l'inacceptabilité des phrases avec *mais plus* en stipulant qu'il y a moins de buts argumentatifs pouvant servir de pivot pour ces phrases que pour celles avec *mais moins que*. Compter ces pivots serait trop long face à des possibilités infinies, permettons-nous donc de diviser 3 types de buts argumentatifs évoqués par *Paul est grand*.

- (18) Paul est grand
- a. **But à balise minimale** : Il doit se placer à l'arrière durant la prise de photo.
 - b. **But à balise maximale** : Il peut passer par la trappe à chat.

c. **But à 2 balises** : Il peut faire le manège.

L'exemple en (18-a) est dit à balise minimale, puisqu'à partir d'une certaine taille, ce but est atteint. En dessous de cette taille, il n'est pas vrai que Paul doit se placer à l'arrière pour la prise de photo. De plus, il n'y a pas de balise maximale à ce type de but. Ainsi, si Paul est juste un peu plus grand que le minimum requis pour se placer à l'arrière, il devra se placer à l'arrière. S'il dépasse intensément la balise minimale, par exemple s'il est aussi grand qu'un gratte-ciel, il devra aussi se placer à l'arrière pendant la prise de photo. Observons comment notre phrase comparative *Paul est grand, mais moins/*plus grand que Pierre* interagit avec ce type de buts avec l'aide de la figure 1.2.

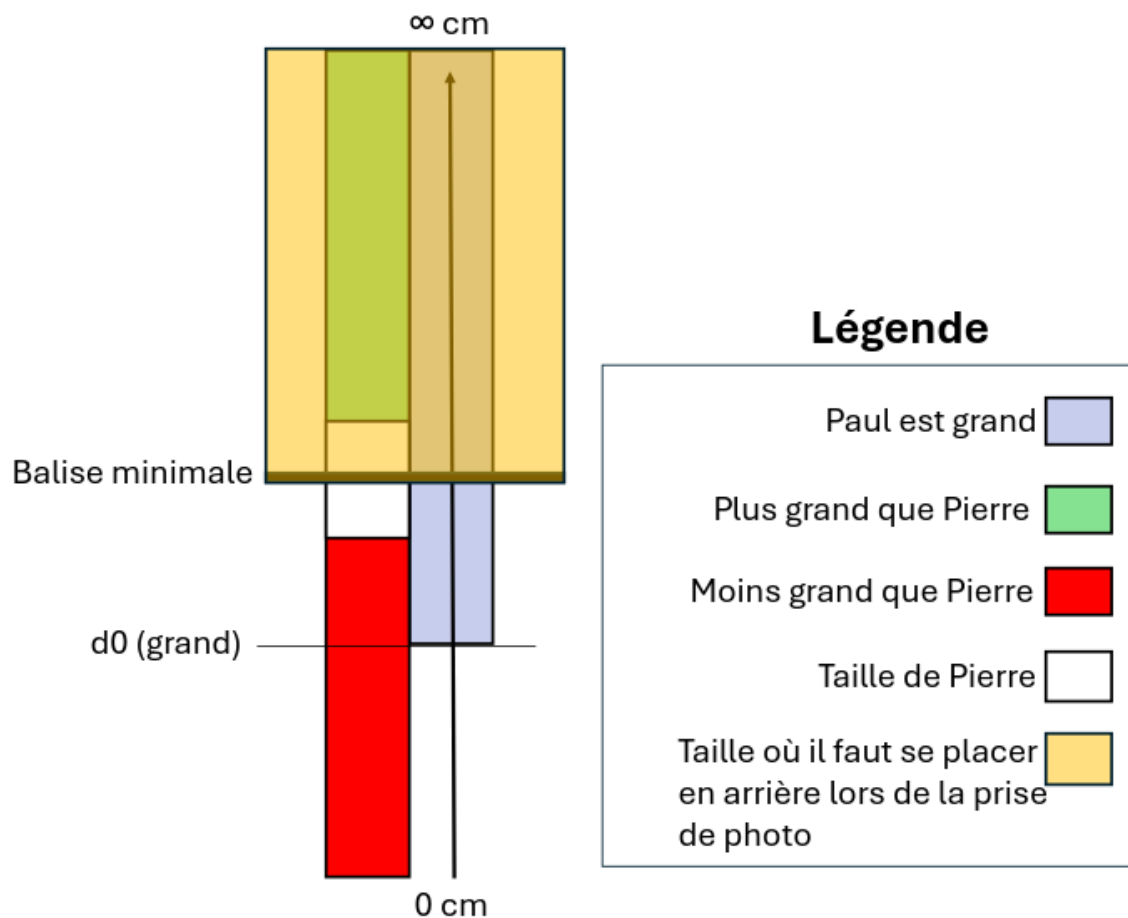


Figure 1.2 Échelle sur la dimension *taille*, allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d_0 , en comparaison avec le but à balise minimale, indiqué par l'espace en jaune.

Le rectangle bleu en 1.2 (représentant le lieu potentiel évoqué par *Paul est grand* sur l'échelle de taille) inclus l'espace jaune (représentant le lieu potentiel du degré évoqué par le but à balise minimale) de manière à ce qu'il y a plus de degrés partagés entre ces deux zones que de degrés qui ne le sont pas³. Ainsi, *Paul est grand* augmente la probabilité que Paul aie la taille requise pour devoir se placer en arrière pour la prise de photo, et est donc un argument en faveur du but argumentatif à balise minimale.

Limiter la taille de Paul à celle représentée par le rectangle rouge (*moins grand que Pierre*) a pour effet de réduire la probabilité que la taille de Paul se situe dans l'espace jaune représentant le but. Ainsi *moins grand que Pierre* est un contre-argument vis-à-vis de ce type de but, ce qui nous permet de prédire que la phrase *Paul est grand, mais moins grand que Pierre* serait correcte.

Le rectangle vert, soit l'espace représenté par *plus grand que Pierre*, partage plus de points avec l'espace jaune représentant le but, et est donc un argument en faveur du but. Ainsi, la phrase *Paul est grand, mais plus grand que Paul* serait prédite incorrecte avec ce type de but.

La phrase en (18-b), *Paul peut passer par la trappe à chat* peut constituer un but à balise maximale. C'est-à-dire que la taille de Paul ne peut dépasser cette balise. Une fois en dessous de la balise, la limite inférieure est celle de l'échelle, soit dans ce cas-ci 0 cm.

3. Cette assertion est possible en utilisant l'infinité de l'adjectif *grand*. Si c'est vraiment ce qui explique la différence d'acceptabilité entre les comparatives avec *mais plus* et *mais moins*, nous pourrions prédire que la différence serait moins remarquable sur des adjectifs à échelle fermée, comme *complété*, ce qui serait intéressant de tester dans une étude à cet effet.

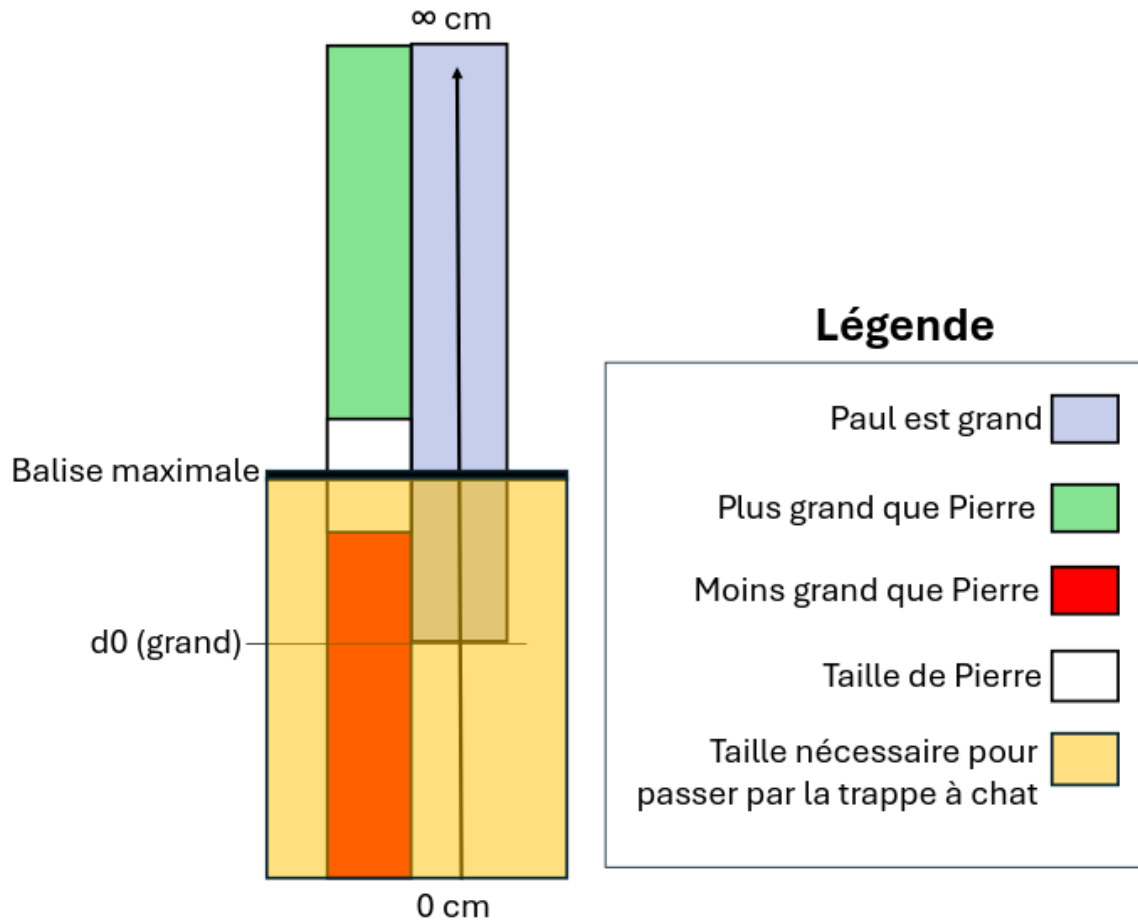


Figure 1.3 Échelle sur la dimension *taille*, allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d0, en comparaison avec le but à balise maximale, indiqué par l'espace en jaune.

Pour ce type de but, c'est l'inverse. *Paul est grand* diminue la probabilité d'atteindre le but, alors que *moins grand que Pierre* l'augmente. Cependant, puisque ces deux énoncés sont tout de même contre-orientés, on prédit encore une fois l'acceptabilité de la comparative avec *mais moins*. Similairement, *plus grand que Pierre*, tout comme *Paul est grand*, diminue la probabilité que Paul puisse passer par la trappe à chat, ce qui prédit que la comparative avec *mais plus* serait incorrecte avec ce type de buts argumentatifs.

Finalement, le troisième type de buts désignant la taille sont ceux à balise double, comme l'exemple en (18-c), *Paul peut faire le manège*. En effet, si la taille de Paul est inférieure à un certain seuil, il sera trop petit pour pouvoir entrer dans un manège. De plus, si sa taille est supérieure à un autre seuil, il sera trop

grand pour entrer sécuritairement dans le wagon du manège.

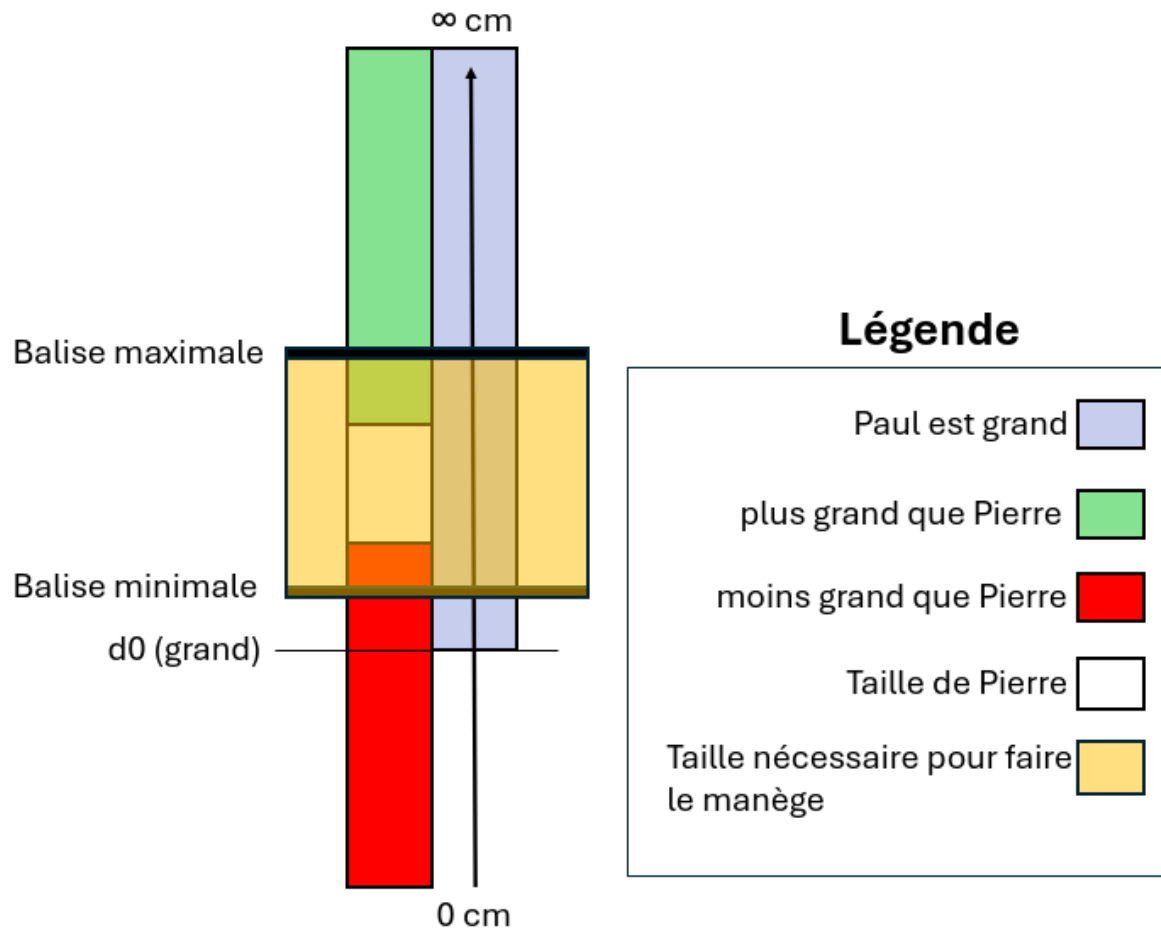


Figure 1.4 Échelle sur la dimension *taille*, allant de 0 à ∞ , avec en bleu la taille potentielle de Paul, qui dépasse la balise contextuelle d_0

Ici, *Paul est grand* augmente la probabilité de dépasser la balise minimale, et ainsi argumente en faveur du but argumentatif. *moins grand que Pierre*, en diminuant la probabilité de dépasser la balise minimale, sert de contre-argument, ce qui prédit l'acceptabilité des phrases comparatives avec *mais moins*. De plus, *plus grand que Pierre*, en augmentant la probabilité que la taille de Paul dépasse la balise maximale, est aussi un contre-argument. Ainsi, dans le cas des buts à balises doubles, les comparatives avec *mais plus* serait prédites correctes par une vision probabiliste de **AdL** informée par les caractéristiques des propriétés gradables, comme on peut voir dans le tableau 1.3.

Type de but	mais moins	mais plus
Balise minimale	oui	non
Balise maximale	oui	non
Balises doubles	oui	oui

Table 1.3 Prédications de la félicité de phrases avec *mais moins* ou *mais plus* selon le type de buts argumentatifs.

Les conditions d'égalités insérées dans les contextes comme en ?? occasionnent un but à balise double, puisque, par exemple, la taille de Paul doit être ni supérieure, ni inférieure à celle de Pierre. L'approche exposée ci-dessus n'est pas précurseur à l'invention des items, mais explique plutôt bien pourquoi l'insertion d'une condition d'égalité à l'aide d'un contexte résulte en des phrases qui seront jugées acceptables.

1.1.4 Approches d'ambiguïté

De par la diversité des usages potentiels de l'adversatif *mais*, certains auteurs, comme Izutsu (2008) ont proposé que les connecteurs adversatifs tels que *mais* soient en fait ambigus. En effet, *mais* se traduit en un mot différent en espagnol si son usage est correctif (Anscombe et Ducrot, 1977), et *mais* se traduit en un mot différent en russe si son usage est contrastif (Jasinskaja *et al.*, 2008). Cette diversité porte à croire que *mais*, en français, est un mot ambigu, dont les contextes sémantiques et syntaxiques détermineront la signification utilisée. Cependant, la délimitation entre différents types d'usages de *mais* est elle aussi souvent ambiguë, comme on peut voir dans l'exemple suivant.

- (19) a. Pensez-vous que McGregor a une chance de gagner ?
b. McGregor est puissant, mais Mayweather est endurant. Je dirais donc non.

Dans l'exemple ci-dessus, la phrase en (19-b) est un exemple classique de phrase contrastive. Il y a en effet 2 individus différents, 2 propriétés différentes, et le fait que McGregor soit puissant n'enlève rien à l'endurance de Mayweather, et vice versa. Cependant, dans le contexte de (19), "McGregor est puissant" nous encourage à penser que la réponse à la question posée en (19-a) est *oui* alors que "Mayweather est endurant" nous encourage à penser que la réponse à la question en (19-a) est *non*, exactement comme dans un usage

argumentatif de *mais*.

Le fait de dire que *mais* est ambigu et que le contexte détermine la signification visée serait plus facile à considérer si la distinction entre les significations possibles n'était pas elle-même ambiguë. Plutôt, nous pencherons dans ce texte vers des approches qui considèrent *mais* comme étant une seule unité sémantique.

1.1.5 Récapitulatif

Le connecteur adversatif *mais* peut être utilisé de différentes façons et reçoit un éventail varié de descriptions sémantiques. Cette section a fait un balayage des différentes utilisations du connecteur en précisant de quels types d'utilisation de *mais* il est question dans ce travail, soit les constructions **comparatives** de *mais* et a exploré différentes approches sémantiques qui cherchent à décrire le comportement de *mais*.

Nous décidons tout d'abord d'exclure les approches ambiguës, préférant dresser un portrait unique d'un mot unique. Deux types d'approches ont été décrites en plus de profondeur, soit les approches contrastives formelles et les approches inférentielles.

Les approches contrastives formelles décrivent plutôt bien l'opposition requise entre les propositions conjointes par *mais* en utilisant les propriétés lexicales des éléments qu'elles contiennent. Pour les usages comparatifs de *mais*, il suffit d'utiliser l'aspect intrinsèquement comparatif des adjectifs gradables en s'inspirant du degré 0 (d0) de (Kennedy et McNally, 2005) pour permettre la formation de l'ensemble de réponses alternatives à la question "plus grand que QUOI?". Le premier énoncé conjoint avec *mais* doit sélectionner le premier élément de l'ensemble, alors que le deuxième doit le nier, ce qui est le cas avec *moins grand que ...*, mais pas avec *mais plus que* Selon ce type d'approche, *mais plus grand que* serait lexicalement incorrect, et le resterait, peu importe le contexte. Cette prédiction va à l'encontre des résultats de (Winterstein et al., 2014) sur les jugements d'acceptabilités, mais corrobore la persistance de la différence dans les temps de traitement qui fut découverte au cours de la même expérience. Dans le cadre de mon mémoire, des résultats montrant que le contexte a un effet significatif sur la différence de comportements oculaires entre les comparatives avec *mais moins* et celles avec *mais plus* seraient incompatibles avec ce type d'approche, puisque la phrase et ses propriétés lexicales sont à elles seules suffisantes pour expliquer la différence de traitement.

Les approches inférentielles, qui utilisent un énoncé pivot autour duquel l'opposition se produit, brillent lorsque ce pivot est visiblement disponible dans le contexte. Pour certaines de ces approches, les phrases hors contexte posent problème, puisqu'on doit alors justifier la présence de la proposition tierce par un phénomène sémantique ou pragmatique quelconque, par exemple des implicatures conventionnelles ou de quantité. Cependant, ces implicatures ne constituant pas toujours des pivots permettant de bonnes prédictions par rapport à l'acceptabilité des phrases avec *mais*, leur accessibilité seule n'est pas un critère suffisant, même si c'est ce que la théorie de la pertinence (RT) propose. L'**AdL** (argumentation dans la langue) évite le problème de justifier l'émergence de propositions tierces selon certains phénomènes sémantiques en stipulant qu'à chaque énoncé est lié son but argumentatif, qui peut être représenté sous la forme d'un énoncé. De plus, ce cadre propose une façon de trier quelles propositions pivots sont permises : Le premier énoncé conjoint par *mais* doit argumenter en faveur du pivot, soit augmenter sa probabilité, alors que le deuxième énoncé doit argumenter contre, en baissant sa probabilité. L'interprétation probabiliste de **AdL** par (Merin, 1999) et la prise en compte des caractéristiques intrinsèques des adjectifs gradables de (Kennedy et McNally, 2005) permettent d'expliquer la différence d'acceptabilité entre *mais plus* et *mais moins* : *mais moins* permet tous les types de buts argumentatifs, alors que *mais plus* permet seulement certains types de buts. Dans le cadre de mon mémoire, une absence d'effet du contexte soulignerait la nécessité de modifier notre version de **RT** pour justifier l'absence d'effet du contexte, alors qu'une absence d'effet de la différence entre *plus* et *moins* (instanciée dans mon mémoire sous la variable **Valence**) serait incompatible avec l'aspect plus lexicaliste que l'on prête ici à **AdL**.

1.2 Oculométrie

L'oculométrie est loin d'être athéorique. Ainsi, cette section exposera les connaissances oculométriques préalables à la compréhension de mon mémoire. Il y sont décrits les différents types de mouvements oculaires, ainsi que la façon dont chacun est pris en compte, que ce soit par prise de mesure ou contrôle de l'environnement. Ensuite, nous décrivons en détail l'absence de mouvement qui nous intéresse, soit la fixation oculaire. Finalement, nous abordons trois particularités spécifiques à l'étude de la lecture en oculométrie, soit (1) le champs perceptuel asymétrique, (2) le lien indirect entre la durée de fixation et la durée de traitement, et (3) l'absence de corrélation entre la durée de fixation et l'amplitude saccadique, qui contraint quelque peu la nature des processus cognitifs dont nous pouvons conclure l'existence à partir de nos données.

1.2.1 Mouvements oculaires

Il y a 4 types de mouvements oculaires (Purves *et al.*, 2001).

La **poursuite visuelle** est un mouvement lent qui se produit lorsqu'un individu suit des yeux un objet se déplaçant lentement dans l'espace (Purves *et al.*, 2001). Puisque nos stimuli sont des textes immobiles, nous ne parlerons pas de poursuites visuelles dans ce mémoire.

Les mouvements **vestibulo-oculaires** sont des réflexes dont l'utilité est de continuer de regarder dans une direction donnée malgré des mouvements du corps et de la tête qui détournerait en temps normal le regard du point ciblé (Purves *et al.*, 2001). Puisque nos sujets ont la tête appuyée sur un appuiement tout au long de l'expérience, nous n'extrayons pas de mesures liées à ce type de mouvement, et ne les traitons donc pas dans ce mémoire.

Les mouvements de **vergences** (*convergence* ou *divergence*) sont ceux pour lesquels les yeux s'orientent chacun dans une direction, de façon asymétrique. Ce type de mouvements se produit pour s'adapter à des objets de distances inégales ou variables (Purves *et al.*, 2001). Puisque l'écran est situé à la même distance pour tous les participants tout au long de l'expérience, il ne sera pas question de ce type de mouvement.

Finalement, la **saccade** est un mouvement balistique rapide qui vise à déplacer le point de fixation d'un endroit à un autre (Purves *et al.*, 2001). De par la vitesse des saccades, l'information recueillie durant une saccade est si pauvre qu'elle échappe à notre conscience, ce qu'on appelle la cécité saccadique (Bridgeman *et al.*, 1975). Or, puisque ce qui nous intéresse dans le cadre de ce mémoire est la façon dont l'information est extraite dans la lecture de nos phrases, les mesures propres aux saccades, comme la latence ou l'amplitude, ne seront pas considérées. Toutefois, puisqu'une saccade est une transition entre deux fixations, elle peut être considérée comme la réaction au changement de focus attentionnel durant la lecture. Ainsi, nous les avons compté et avons pris la peine de les qualifier de progressive (de gauche à droite) ou régressive (de droite à gauche) pour mieux comprendre les mesures prélevées de nos fixations. Par exemple, une fixation faite sur un mot suite à une saccade régressive pourrait signifier que l'information véhiculée par ce mot était manquante et suffisamment nécessaire pour justifier un retour vers l'arrière.

1.2.1.1 La dilatation de la pupille

Une pupille normale se contracte lorsqu'elle est illuminée ou lorsque l'autre pupille est illuminée, et se dilate à la noirceur, de façon à permettre à une plus grande quantité de lumière d'atteindre les différents neuro-récepteurs servant à la capture de l'information visuelle (Spector, 1990). La dilatation de la pupille est liée à toutes sortes de phénomènes, comme l'effort cognitif, la valence émotionnelle, le stress, ou la douleur (Wang, 2011). Même si ces phénomènes ne sont pas notre objet d'études, certains, comme l'effet de l'effort cognitif, auraient pu servir de mesures intéressantes, ne serait-ce qu'à des fins de contrôle. Cependant, la prise en compte de ces informations aurait requis (1) une modification considérable des items expérimentaux pour éviter toute confusion entre les différents facteurs corrélant avec la dilatation de la pupille ainsi que (2) une modification de la présentation des items pour contrôler l'effet d'illumination, qui consiste à ajouter un délai entre chaque stimuli pour laisser le temps à la pupille de revenir à une taille normale, ce qui aurait allongé péniblement la passation de l'expérience. Ainsi, même si le logiciel utilisé capture d'emblée la taille de la pupille, elle n'est pas utilisée dans ce mémoire.

1.2.2 Immobilité oculaire : La fixation

Une fixation, c'est lorsque les yeux ralentissent suffisamment pour permettre la capture d'informations nouvelles (McConkie *et al.*, 1985; Ishida et Ikeda, 1989). Comme tout phénomène dynamique, sa capture complète n'est pas réalisable, et dans ce cas-ci pas très éthique. En observant à la place une multiplicité de moments statiques (250 par secondes dans le cas de notre oculomètre), nous passons par deux intermédiaires pour qualifier les fixations ayant lieu au cours de l'expérience : (1) leur durée et (2) leur lieu.

La **durée** d'une fixation est simplement le temps total où le mouvement de l'oeil est jugé comme étant "assez lent" pour être caractérisé de fixation, jusqu'à ce qu'il ne le soit plus. Ce qu'on entend par "assez lent" dépend des chercheurs, de leur question de recherche, du matériel utilisé, et même parfois de certaines caractéristiques cognitives des participants à l'expérience (Holmqvist *et al.*, 2011). Ainsi, différents filtres de données oculométriques ne produiront pas des fixations de la même durée. Les configurations spécifiques du filtre que nous utilisons pour extraire les fixations de nos données sont décrites en de plus amples détails dans le chapitre 3.

Le ralentissement qui caractérise une fixation est aussi accompagné d'une orientation spécifique de l'oeil, de façon à ce que la lumière émanant d'un objet spécifique rencontre le moins de résistance en chemin

vers la fovée. La fovée, ou *fovea centralis*, est une zone située à l'arrière de l'oeil, près du nerf optique, qui contient une grande quantité de cellules coniques, qui permettent de discriminer les couleurs, les détails, ainsi que les mouvements rapides (Purves *et al.*, 2001). Pour atteindre cette orientation, l'oeil bouge de façon à s'orienter sur l'*axe visuel*, soit une ligne imaginaire qui se trace entre la fovée et l'objet fixée.

L'oculomètre ne permettant pas de repérer spécifiquement la fovée qui se trouve à l'arrière de l'oeil, l'*axe visuel* ne peut pas être observée. Cependant, l'*axe optique*, qui est la ligne imaginaire de la rétine à l'objet fixé placée de façon à être perpendiculaire au centre de la pupille, peut être calculée en observant la position de ladite pupille avec l'oculomètre. Du positionnement de l'*axe optique* peut être calculé celui de l'*axe visuel*, résultant en un ensemble de coordonnées sur l'écran correspondant au lieu de fixation.

Un changement d'attention de la part du lecteur se caractérisera par un changement de lieu de fixation, ainsi une rotation de l'*axe visuel*. L'angle à la jonction des deux axes nous permet donc d'évaluer la distance entre deux fixations, tel que démontré dans la visualisation aptement dessinée en 1.5.

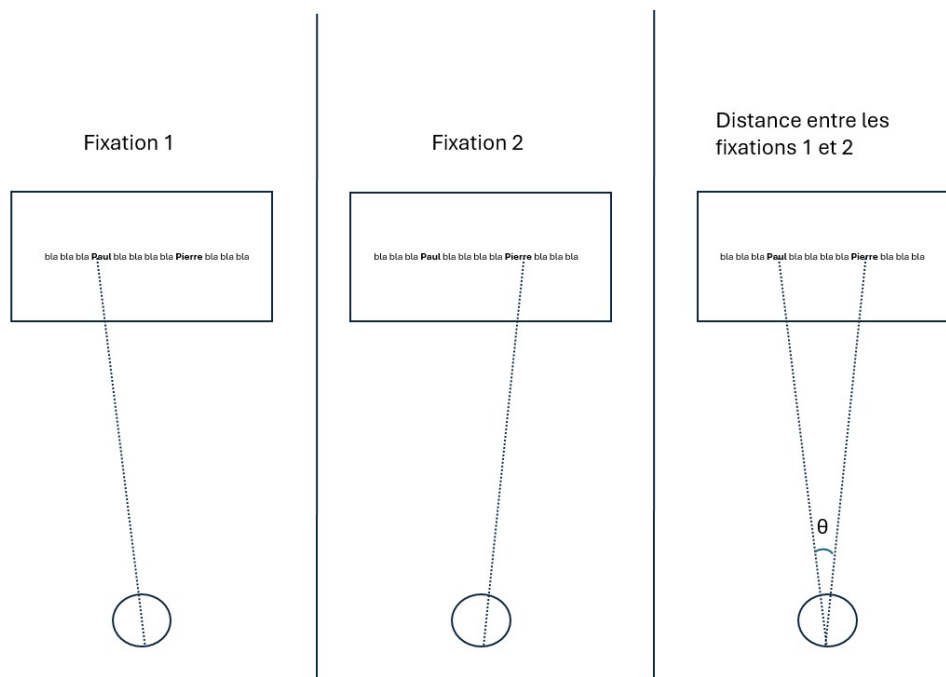


Figure 1.5 Représentation de l'axe visuel lors de 2 fixations distinctes, puis de l'angle entre les deux axes comme mesure de distance

Ainsi, le calcul de l'*axe visuel* nous permet à la fois de retrouver le lieu de la fixation et de calculer les distances entre les différentes fixations, qui seront indiqués pour la suite de ce mémoire en degrés d'angle

visuel.⁴

1.2.3 Particularités de la lecture

Ce qui suit ne s'applique plus de façon générale à la vision humaine et à l'oculométrie, mais bien spécifiquement à ce qui touche le phénomène de la lecture. On y voit trois particularités, soit l'asymétrie du champs perceptuel, l'indirectitude du lien entre les durées de fixation et le traitement, et la description sommaire de deux types de processus décisionnels qui sépare la durée des fixations de son emplacement.

1.2.3.1 Champs perceptuel

Si le lieu de fixation, décrit plus haut, représente ce à quoi on porte une attention particulière, il ne s'agit pas de tout ce que l'on voit. Le champs de vision se trouvant à proximité du lieu de fixation se nomme *vision parafovéale*. En lecture, l'individu n'est pas sensible à tout ce qui se trouve en vision parafovéale, mais plutôt à un sous ensemble, que nous nommons *champs perceptuel*. Ce champs perceptuel a la particularité, en lecture, d'être asymétrique, comme on peut voir dans l'image en 1.6, tiré de (Zihl, 1993).

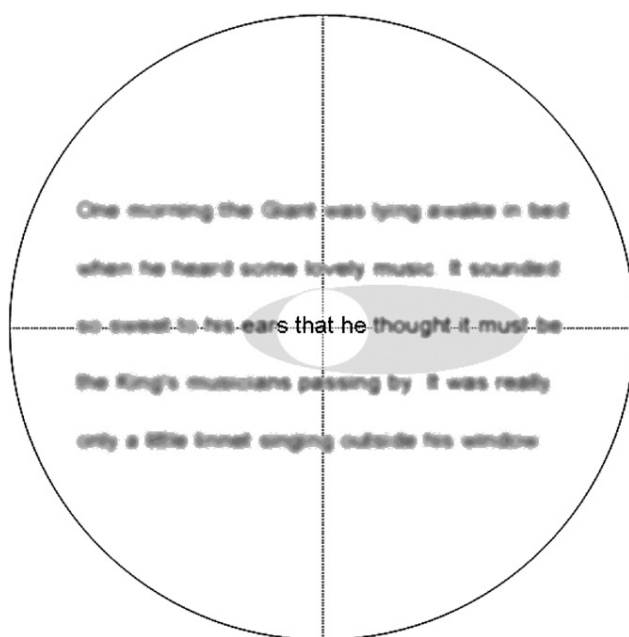


Figure 1.6 Schématisation du champs perceptuel, montrant son asymétrie adaptée au système d'écriture.

4. Il est aussi possible de mesurer la distance en pixels, mais ça ne prend pas en compte la distance de l'oeil à l'écran, ni le fait que cette dite distance n'est pas la même pour le centre ou pour l'extrémité du moniteur.

Pour les systèmes d'écritures qui se lisent de gauche à droite, comme celui du français utilisé dans notre expérience, les sujets sont sensibles à un champs perceptuel qui commence de 3 à 4 caractères à gauche du lieu de fixation, sans dépasser le début du mot fixé (Rayner *et al.*, 1980), et qui fini environ 14 ou 15 caractère à droite, peu importe les frontières des mots (Rayner *et al.*, 1982). La quantité de caractères perçu vers la droite peut augmenter avec l'expérience en lecture, mais pas celle à gauche (Rayner, 1986).

La vision parafovéale permet de traiter complètement les petits mots comme les mots fonctionnels, pour ensuite ne pas avoir à les fixer, tout en facilitant la lecture des mots suivant la fixation. (Blanchard *et al.*, 1989). Quand un mot suivant celui fixé est accessible en vision parafovéale et est sémantiquement non-plausible, le temps de la fixation augmente.

1.2.3.2 Durées de fixation et traitement

Puisque les informations par rapport au texte lu sont capturées pendant les fixations, il est tentant de simplement additionner les durées de fixation pour chaque mot ; le mot fixé le plus longtemps est celui qui a été le plus amplement traité par l'individu lors de sa lecture. Hélas, la relation entre le temps de traitement et le temps de fixation n'est pas si directe (Rayner, 1998). En effet, plusieurs mots, souvent ceux très petits et fréquents, ne sont jamais fixés, même s'ils ont été traités par le lecteur (Fisher et Shebilske, 1985). Aussi, la fixation faite avant de passer tout droit (*skip*) sur un mot est plus longue (Pollatsek *et al.*, 1986), indiquant qu'une certaine partie du traitement d'un mot se fait avant de l'avoir fixé, probablement avec l'information perçue grâce à la vision parafovéale. De plus, après avoir fixé un mot particulièrement compliqué, les fixations suivantes seront plus longues, même si elles ne sont pas sur ce même mot. Ce phénomène est souvent appelé *Spillover effect*, et nous le nommons **effet de débordement** pour le reste de ce mémoire. Finalement, plusieurs mesures liées aux temps de fixation augmentent considérablement en fin de phrase pendant l'intégration des différentes informations extraites lors de la lecture, ce qu'on appelle l'effet *wrap-up* (Just et Carpenter, 1980), et que je nommerai **effet de fermeture** pour le restant de ce mémoire.

Malgré ces difficultés, il y a une façon de déduire des informations sur les processus de lecture avec l'oculométrie. Il s'agit de prendre une zone d'intérêt, dans notre cas le lieu de la paire minimale (*plus, moins*), et prendre toutes sortes de mesures par rapport aux zones sur et entourant cette zone d'intérêt (Schmauder, 1992). Si une différence significative au niveau d'une seule mesure peut être causée par tout le bruit statistique intrinsèque à la lecture, plusieurs différences significatives sur la même zone sur différentes mesures nous permet avec plus d'assurance d'en conclure qu'il y a une différence dans les processus sous-jacents.

C'est pourquoi nous prélevons les mesures listées ci-dessous.

- La durée de la première fixation pour chaque mot.
- La durée du premier regard pour chaque mot.
- La durée totale des fixations pour toute la phrase
- La probabilité de retourner en arrière pour toute la phrase.
- La durée du parcours régressif (*regression-path duration*)
- La probabilité d'être le point de départ d'une régression oculaire pour chaque mot
- La probabilité d'être le point d'arrivée d'une régression oculaire pour chaque mot.

Pour chacune de ces mesures, une définition ainsi qu'un aperçu des processus cognitifs qui y sont reliés se retrouvent dans le chapitre 3.

1.2.3.3 Où et Quand

Dans les mouvements visuels qui ne font pas partie d'une lecture, il y a une corrélation positive entre la durée d'une fixation et la distance parcouru pendant la saccade qui la précède (Nattkemper et Prinz, 1987). En lecture, ce n'est pas le cas (Rayner et Fischer, 1996). Ce manque de corrélation laisse à penser que deux types de décisions sont prises indépendamment l'une de l'autre : (1) **quand** changer l'attention, ce qui influence la durée des fixations, et (2) **où** fixer l'attention de la prochaine fixation, ce qui influence l'amplitude de la saccade résultante.

La question d'**où** arrêter son regard est principalement entraînée par des facteurs de configuration visuo-spatiales, comme l'espacement entre les caractères et la longueur du mot fixé ainsi que celle du mot qui suit. Ces facteurs peuvent être mesurés en tenant compte de la longueur d'une saccade et du lieu de première fixation d'un mot, c'est-à-dire le caractère précis sur lequel le regard atterrit pour la première fois lorsqu'il observe un mot. Ce sont des mesures sur lesquelles nous ne nous attarderons qu'à des fins de contrôles, pour nous assurer que la configuration visuo-spatiale soit la même pour chaque participant, mais sans plus.

La question de **quand** arrêter une fixation est quant à elle influencée par des facteurs cognitifs qui peuvent être impactés, entre autres, par la fréquence (Rayner et Fischer, 1996) ou la probabilité d'apparition (Binder *et al.*, 1999) du mot observé. Autrement dit, on s'attarderait plus longtemps sur des mots plus surprenants et moins fréquents. Ainsi, ces types de facteurs sur le traitement cognitif de la phrase lue peuvent être

observés avec des mesures comme le temps de fixation de chaque mot et le temps de fixation total, que nous prélèverons dans notre expérience.

CHAPITRE 2

MÉTHODE

Ce chapitre décrit les actions entreprises dans le but de préparer, monter, et passer l'expérience. On y voit tout d'abord une description globale des participants et de leur processus de recrutement. Ensuite, est donné un portrait détaillé du processus de créations des items expérimentaux, suivi d'un bref aperçu de la méthode de prélèvement de la mémoire de travail, puis d'une liste de l'étendu du matériel physique utilisé. Finalement, la croix de fixation est décrite, ainsi que les conséquences d'une erreur dans son instanciation.

2.1 Participants

Pour assurer un équilibre entre puissance statistique et faisabilité, nous visions 30 participants. Suite à une erreur remarquée lors de la 27^e passation, nous avons décidé d'en passer 30 autres à partir de ce point, ce qui nous donne 57 participants au total. Aucune restriction par rapport à l'âge n'est appliquée, de façon à satisfaire les inquiétudes du comité d'éthique vis-à-vis de l'âgisme et à permettre à une plus grande quantité de sujets de participer. L'âge est cependant noté, nous permettant de le prendre en compte dans les analyses statistiques, puisque l'âge affecte considérablement la vitesse de lecture (Rayner *et al.*, 2006; Gordon *et al.*, 2015). Les sujets ont, préalablement à la passation de l'expérience, dû affirmer (1) ne pas être dyslexiques, (2) avoir une vision normale ou corrigée et (3) être locuteurs natifs du français québécois.

2.1.1 Recrutement

Des messages ont été publiés sur Facebook contenant un appel aux participants, disponible en annexe, ici. Ce message contenait des indications sur les détails de l'expérimentation, les critères de sélection, et la promesse d'une compensation monétaire de 15\$ pour les sujets de notre expérience, puisque celle-ci est plutôt longue, soit d'une durée d'environ 1h.

Puisque cela attirait peu de personnes, des messages privés ont été envoyés à mes connaissances Facebook, les référant aux formulaires de recrutement précédemment mentionné. Des visites ont aussi été faites dans des classes de premier cycle en linguistique pour les inviter directement à participer et à prendre rendez-vous. Ces deux dernières méthodes se sont avérées plus efficaces que l'attente passive d'une réponse à une annonce globale.

2.2 Items

Les textes qui sont montrés aux participants sont de 3 natures : (1) les contextes, (2) les phrases et (3) les questions. Ceux-ci, tous présents en annexe B, sont soumis aux mêmes contraintes globales, décrites dans la section 2.2.1, en plus d'être soumis à des contraintes spécifiques à leur type, ce qui est décrit dans les sections 2.2.2, 2.2.3 et 2.2.4.

2.2.1 Contraintes globales

Puisque la fréquence d'un item lexical peut avoir un effet sur les temps de premières fixations, les temps de fixations totales, et le nombre de refixations (Rayner et Fischer, 1996), nous nous sommes assuré que la variation dans la fréquence des mots lexicaux ne soit pas trop grande. Tout d'abord, les *stop words* ont été ignorés, utilisant une liste de *stop words* en français (Govignon, 2021), auquel nous avons ajouté le mot *y*. Un *stop word* est un mot dont la probabilité de présence dans un document quelconque est similaire ou égale à sa probabilité d'apparition dans d'autres documents qui ne traitent pas nécessairement du même sujet (Wilbur et Sirotkin, 1992). Cette liste contient donc des mots (ex : *de*, *à*, *le*) qui apparaissent dans la majorité des textes, sont très fréquents, et sont principalement fonctionnels.

Ensuite, nous avons attribués à tous les mots n'étant pas des *stop words* une valeur de fréquence représentant la moyenne de l'apparition de son mot-forme (token) dans le corpus Frantext, soit une collection de 218 romans datés de 1950 à 2000, en utilisant la base de données lexicale francophone Lexique 3 (New et al., 2004). Pour chaque texte affiché, peu importe sa nature, la moyenne des fréquences des mots qui le remplissent fut calculée, puis comparée à des balises minimales et maximales résultantes de la formule de différence interquartile ci-dessous en (1).

(1) **balise minimale** : moyenne - (différence interquartile * 1.5)

balise maximale : moyenne + (différence interquartile * 1.5)

La fréquence lexicale moyenne par texte était de 283.89, avec un écart interquartile de 302.63. Selon notre calcul, la balise minimale était de -170.06, et la balise maximale de 737.84. Le fait que la balise minimale était négative a été interprété comme une indication qu'il n'y avait pas de mots trop peu fréquents dans les items, et la balise minimale fut donc établie à 0. Toutes les phrases ayant une moyenne de fréquence

supérieure à 737.84 furent inspectées, et plusieurs mots ont été changés pour des mots moins fréquents. Par exemple, le verbe *vouloir*, qui est très fréquent, fut changé pour *souhaiter* ou *désirer* à maintes reprises.

Nous nous sommes assuré que la taille des textes soit similaire à celle des autres textes du même type pour éviter que la différence dans leur longueur n'ajoute une variable additionnelle. Cependant, nous n'avons pas utilisé la méthode de l'écart interquartile pour délimiter les balises minimales et maximales pour la quantité de caractères des items, puisque l'intervalle qui en résultait était trop petit. En effet, les contraintes syntaxiques et sémantiques étaient prioritaires, de par leur apport à la question de recherche, et les contraintes de fréquence étaient aussi cruciales quand on prend en compte l'impact de la fréquence sur les temps de lecture (Rayner et Fischer, 1996). Nous avons donc décidé de fixer un écart de 20 caractères. Les contextes avaient entre 222 et 242 caractères, les phrases entre 128 et 108 et pouvaient ainsi être affichées sur une seule ligne sans que la police soit trop petite, et les questions entre 48 et 68.

Dix items ainsi que 30 leurres ont été recueillis directement d'une expérience préalable traitant d'un sujet très similaire (Winterstein *et al.*, 2014), puis modifiés pour être plus adaptés à l'expérience de ce projet. Notamment, puisque nos participants étaient locuteurs natifs du français québécois, quelques expressions ont été changées. Par exemple, l'expression *tableau pendu au mur* avait une connotation quelque peu fatale, et a donc été changée pour *tableau accroché au mur*. 10 items et 10 leurres additionnels ont été créés de toutes pièces, résultant en 40 leurres et 20 items expérimentaux au total.

2.2.2 Les contextes

Les contextes remplissent deux rôles dans le cadre de mon expérience, et ainsi sont sujets à deux types de contraintes. Le premier rôle en est un **général**, ainsi les contextes sont créés et modifiés en prenant en considération des contraintes générales qui ne sont pas spécifiques à notre expérience, et qui sont décrites en 2.2.2.1. Ensuite, dans le cadre spécifique de notre expérience, les contextes se divisent en deux catégories, **facilitateur** et **neutre**, et servent ainsi à introduire une variable indépendante expérimentale liée à notre question de recherche. Les décisions prises par rapport à cet aspect de nos contextes sont présentées en 2.2.2.2.

2.2.2.1 Utilités et contraintes générales des contextes

Le premier rôle que les contextes remplissent dans le cadre de ce mémoire est d'introduire les personnages, lieux, et événements nécessaires pour la compréhension des phrases qui les suivent, dont on peut voir un exemple en (2).

- (2) Le gâteau d'Alice est moelleux, mais moins moelleux que celui fait par sa grand-mère pour chacun de ses anniversaires.

Cette phrase en isolation soulève en effet beaucoup de questions : Qui est Alice ? Pourquoi parle-t-on de son gâteau ? Pourquoi parle-t-on de l'aspect moelleux de ce gâteau ? Pourquoi compare-t-on son gâteau à celle de sa grand-mère ? Le contexte (3) a l'utilité première de répondre à toutes ces questions préalablement à l'affichage de la phrase expérimentale.

- (3) Alice essaie de reproduire le gâteau que cuisinait sa grand-mère pour chacun de ses anniversaires. Pour impressionner sa famille avec ses aptitudes en cuisine, elle veut notamment reproduire avec fidélité la texture moelleuse du gâteau.

Ces contextes agissent aussi comme un retour à la case départ et donc comme un tampon entre les différents univers sémantiques auxquels les diverses phrases font allusion. En effet, la présentation du contexte assure, autant que possible, que c'est l'aspect moelleux du gâteau qui sera au centre de l'attention, et non pas sa potentielle couleur rouge sombre, même s'il est possible qu'un item expérimental qui traite de la couleur rouge, comme en (4), soit présenté préalablement.

- (4) Sur sa palette, Léonie a créé un rouge sombre, mais moins sombre que celui utilisé sur le tableau dans sa chambre.

Les contextes ont aussi l'utilité inhérente de contenir des mots. En effet, plusieurs méthodes d'analyses statistiques impliquent de prendre des mesures sur une grande quantité de mots pour établir des temps de

lecture résiduels grâce à des modèles prédictifs (Winterstein *et al.*, 2014), ou bien un point de référence à partir duquel comparer les différentes mesures obtenues (Toivo et Scheepers, 2019).

Les contextes introduisaient souvent les personnages à l'aide de prénoms. Pour assurer qu'une différence notable dans la quantité des personnages masculins et féminins ne biaise pas la lecture des items, il y a une quantité égale de personnages féminins et de personnages masculins. Pour faciliter cette égalité, aucun prénom graphiquement épiciène, tel que *Claude* ou *Alex* n'est utilisé. Plutôt, des versions où le genre majoritaire des individus désignés par le prénom était plus évident ont été utilisées, comme *Claudette* ou *Alexandre*.

Les thèmes présents dans les items, et donc introduits par les contextes, ont été filtrés et modifiés de façon à ne pas faire référence explicitement à des éléments culturels, dont leurs connaissances pourraient être variables chez les différents participants. Par exemple, un excellent item traitant de Bella et de son triangle amoureux avec Edward et Jacob a été modifié de façon à ne pas faire apparaître la référence à *Twilight* (Meyer, 2008), ce qui aurait pu entraîner une variation dans la reconnaissance de la référence culturelle, et ainsi un biais statistique évitable.

2.2.2.2 Les contextes facilitateurs et neutres

Chaque contexte qui introduit les items expérimentaux est créé en deux versions : une dite **facilitatrice** et une dite **neutre**. Les contextes **facilitateurs** introduisent l'idée que les deux entités mentionnées dans l'item expérimental se doivent d'être exactement égales par rapport à l'atteinte de la propriété gradable autour de laquelle on les compare. Si on observe la phrase (5-a), le contexte qui l'introduit invite les lecteurs à être attentifs à la possibilité que les deux personnages aient exactement le même âge, et donc atteint de façon équivalente la propriété *vieux*.

(5) **Contexte** : Jeanne Doe est amnésique et ne se rappelle plus qui elle est. Les enquêteurs pensent qu'elle est la sœur jumelle de Marc Leclerc, un vieil avocat disparu il y a 10 ans. Pour s'en assurer, ils font passer à Jeanne une batterie de tests.

- a. Selon les tests, Jeanne est vieille, mais moins vieille que Marc, même s'ils ont tous les deux les cheveux blancs.

Le contexte neutre (6), quant à lui, ne nous invite pas à vérifier dans l'item subséquent que Jeanne a le même âge que Pierre. Ainsi, le contexte neutre ne donne accès à aucun matériel permettant d'inférer un pivot tel que décrit en 1.1.2, contrairement aux contextes facilitateurs. Il est à noter ici que le contexte neutre est un exemple de modification minimale d'un contexte **facilitateur** à un contexte **neutre**. Le changement de *jumelle* à *petite*, ne retire qu'un seul caractère et conserve une relation de fraternité entre les deux personnages, sans changer la propriété gradable (*vieux*) qui y est introduite, tout en évacuant l'assomption d'équivalence par rapport à l'âge des personnages.

- (6) Jeanne Doe est amnésique et ne se rappelle plus qui elle est. Les enquêteurs pensent qu'elle est la petite sœur de Marc Leclerc, un vieil avocat disparu il y a 10 ans. Pour s'en assurer, ils font passer à Jeanne une batterie de tests.

Malheureusement, tous les contextes neutres n'éliminent pas l'assomption d'équivalence associée aux contextes **facilitateurs** avec des modifications aussi minimales, comme on peut voir dans la paire de contextes (7), pour laquelle une certaine assomption d'équivalence est conservée, mais la propriété pour laquelle on assume une équivalence est substituée (de *grand* à *agile*).

- (7) **Contexte facilitateur** : On cherche un cascadeur pour doubler Pierre, un acteur, dans des scènes d'action. Il faut que le cascadeur ait exactement la même taille que Pierre pour qu'on ne remarque pas la doublure à l'écran. C'est difficile parce que Pierre est de grande taille.

Contexte neutre : On cherche un cascadeur pour doubler Pierre, un acteur très grand, dans des scènes d'action. Puisque le personnage de Pierre est agile, il faut que le cascadeur soit plus agile que Pierre. C'est difficile parce que Pierre est plutôt agile.

Phrase item : Au niveau de la carrure, Paul est grand, mais moins grand que Pierre et il n'est pas facile de prendre une décision.

L'exercice de substitution n'implique pas seulement de retirer l'assomption d'équivalence, il faut aussi (1) garder la mention explicite de la propriété *grand*, pour ne pas entraîner de possibles différences vis-à-vis de l'amorçage ou l'emprunte rétinienne et (2) garder le nombre de caractères ainsi que (3) la fréquence lexicale à l'intérieur de balises minimales et maximales déterminées dans la section 2.2.1. Ainsi, la catégorie

neutre ne fait pas que retirer le pivot **facilitateur**, elle peut entraîner divers effets sémantiques qui ne sont pas tous prévus et qui ne sont pas uniformes à travers tous les stimuli. La prise en compte des items comme facteurs aléatoires dans les modèles statistiques vise à atténuer l'importance, entre autres, de cet effet imprévisible. Il serait toutefois éclairant d'investiguer l'effet des différents types de substitutions entre les contextes **neutres** et **facilitateurs** sur les mesures recueillies lors de la lecture des phrases qui suivent ces contextes.

2.2.3 Les phrases

Les phrases expérimentales sont des usages comparatifs de l'adversatif *mais*, tel qu'exemplifié ci-dessous en (8).

- (8) Au niveau de la carrure, Paul est grand *mais* **moins/plus** grand que Pierre et il n'est pas facile de prendre une décision.

Comme vu dans la section 2.2.2, les entités (*Paul* et *Pierre*) ainsi que la propriété (*grand*) sont déjà explicitement mentionnées dans le contexte, ce qui résulte en une certaine redondance lexicale entre les phrases et les contextes qui les précèdent. Cette répétition a pour but de diminuer le risque d'ambiguïté dans la résolution des anaphores et a été répliquée dans les leurres, tel qu'exemplifié en (9), pour éviter qu'ils soient identifiables par leur absence de répétition.

- (9) a. **Marc-André** est passionné d'aménagement intérieur. Après avoir changé de **divan**, il cherche une **table** de salon pour donner à la pièce une allure plus complète. Pour gagner de l'espace, la **table** doit être moins longue que le divan.
- b. Dans un magasin **Marc-André** a vu une **table** de salon longue et moins longue que son **divan** dont il est très content.

Les phrases recueillies dans l'expérience de (Winterstein *et al.*, 2014) ont été créées en étant divisées par zones, pour s'adapter au protocole d'auto-présentation segmentée. Pour adapter mes nouvelles phrases à celles de l'expérience précédente, elles ont elles aussi été divisées par zones, qui sont définies en 2.2.3.1.

2.2.3.1 Division des phrases par zones

Les zones furent définies comme légères modifications des zones établies dans l'expérience de (Winterstein *et al.*, 2014). Le tableau 2.1 présente l'item expérimental relié à l'exemple initial de ce mémoire, avec chacun de ses mots identifiés relativement à la zone à laquelle ils appartiennent.

Au niveau de la carrure,	Paul est grand,	mais moins grand	que Pierre	et il n'est pas facile	de prendre une décision.
-1	0	1	2	3	4

Table 2.1 Répartition des mots par zone

Notons ici que les zones n'ont pas toutes des index positifs. Ce choix est fait pour rester près de la numérotation effectuée dans (Winterstein *et al.*, 2014), de façon à pouvoir souligner les résultats similaires, s'il y a lieu.

La zone -1 (*Au niveau de la carrure,*) contient un groupe complément. Dans le protocole d'auto-présentation segmenté duquel la moitié des items émerge, il n'y avait pas nécessairement de contexte précédant la phrase. Ainsi, ce petit groupe complément servait d'entrée en matière. De plus, ce syntagme permet aux entités et propriétés se situant à gauche dans la phrase (*Paul* et *grand*) de ne pas être trop proches de l'extrémité gauche de l'écran, où d'autres effets oculaires ont potentiellement lieu. Cette zone n'étant pas centrale à l'étude, la quantité de mots qu'elle contient est plus variable, pour permettre d'accommoder les zones plus importantes en respectant les balises globales pour le nombre de caractères décrites dans la section 2.2.1. La quantité de mots qu'elle contient varie donc entre 0 et 5, avec une moyenne de 3.1 et une déviation standard de 1.65.

La zone 0 (*Paul est grand*) contient entre 3 et 7 mots avec une moyenne de 4.6 et un écart-type de 1.33. Elle contient le nom d'une des entités comparées (*Paul*), la copule *est*, et la propriété gradable autour de laquelle les entités sont comparées (*grand*).

La zone 1 (*mais plus grand*) contient toujours 3 mots : (1) *mais*, (2) *moins* ou *plus*, et (3) une réitération de la propriété gradable introduite dans le contexte et mentionnée dans la zone 0 (*grand*). Le choix du deuxième mot de cette zone (*plus* ou *moins*) est un des objets de cette étude et détermine la variable **Valence**.

La zone 2 (*que Pierre*) contient entre 2 et 4 mots, avec une moyenne de 2.45 et une déviation standard de 0.59. Elle commence toujours par le mot *que*, suivi par la mention de la deuxième entité comparée (*Pierre*).

Les zones 3 et 4 (*et il n'est pas facile | de prendre une décision*) contiennent deux syntagmes permettant de conclure la phrase sans que celle-ci soit trop abrupte, en laissant un tampon spatial permettant à nos mots clés, soit *mais*, *plus*, *moins*, nos entités et nos propriétés gradables, de ne pas se retrouver à l'extrémité droite de l'écran.

Plus précisément, l'effet de fermeture, vu en 1.2.3.2, est escompté au niveau de la zone 4, qui est la dernière zone et contient donc les derniers mots. S'il y a un effet causé par des mots se trouvant dans la zone 1 ou 2, il se peut qu'il soit apparent dans la zone 3 par effet de débordement, comme vu précédemment en 1.2.3.2. La zone 3 a donc pour objectif de séparer, d'autant que faire se peut, de potentiels effets de débordement l'effet de fermeture escompté. La zone 3 contient entre 1 et 6 mots avec une moyenne de 3.7 et une déviation standard de 1.32. La zone 4 contient entre 0 et 7 mots avec une moyenne de 3.25 et une déviation standard de 1.9.

2.2.4 Questions

Le troisième type de texte et la dernière partie de mes items consiste en une question polaire, c'est-à-dire qui se répond par *oui* ou par *non* comme en (10).

(10) Pierre est-il plus grand que Paul ?

Elles servaient premièrement à détourner l'attention des participants de l'objet d'étude. Dans la partie *de-briefing* de l'expérience, je leur posais la question suivante.

(11) D'après toi, c'est quoi le sujet spécifique de mon mémoire, et de l'expérience que tu viens de passer ?

La majorité des participants orientaient leur réponse à cette question autour de l'impression que leurs réponses aux questions polaires étaient centrales à l'expérience, et seulement 3 personnes ont mentionné

qu'il y avait quelque chose d'étrange dans la façon dont les phrases étaient construites, indiquant ainsi que l'aspect de détournement de l'attention des questions fut un succès.

Deuxièmement, les questions servaient à assurer que les participants lisent effectivement les phrases dans le but de les comprendre, et ne font pas que les parcourir du regard.

83% des questions furent créées de façon à ce qu'il y ait une bonne réponse de prédite. Dans le reste des cas, nous trouvions intéressant de voir ce que les gens allaient répondre. Par exemple, certains leurres contiennent le mot *ou* et la réponse à la question de compréhension diffère selon l'interprétation inclusive ou exclusive de *ou*. Ces réponses ne sont pas considérées dans l'étape d'exclusion des participants. Nous nous sommes assuré qu'il y ait une quantité égale de questions dont la bonne réponse est *oui* que de questions dont la bonne réponse est *non*, et qu'elles aient rapport avec la phrase et le contexte qui les accompagnent. Pour faciliter l'atteinte d'une réponse polaire, les mots atteignant la modalité épistémique, comme "pense" ou "probable" ont été évités le plus possible. Les questions ont été construites et modifiées de façon à questionner spécifiquement un élément présent dans la phrase, et non dans le contexte, ce qui s'avérait parfois trop difficile à se remémorer.

2.3 Mémoire de travail

La mémoire de travail, qui est l'habilité d'emmagasiner de l'information temporairement en vue de l'utiliser (Baddeley, 2012), fut prélevée chez nos participants de façon à assurer que les variations dans les résultats soient bel et bien dues à nos variables indépendantes plutôt qu'à une variation des capacités individuelles de chaque participant.

Une mémoire de travail plus haute entraîne une diminution de la durée totale de fixation (Traxler *et al.*, 2012) ainsi qu'une augmentation de la probabilité de refixation et une tendance vers la droite d'environ 0.5 degré quant à la position d'atterrissage post-saccadique au sein d'un mot (Kuperman et Van Dyke, 2011). Il est possible que ces relations ne soient pas directement causales, mais bien dérivatives (Traxler *et al.*, 2012), c'est-à-dire que la mémoire de travail affecterait certaines capacités, comme celle de lecture de mot ou de compréhension à l'écoute (Kim *et al.*, 2021), qui elles-mêmes auraient un impact sur les différentes mesures oculométriques. La méthode décrite ci-dessous pour évaluer la mémoire de travail, le *Reading Span Test*, en est une au cœur de cette ambiguïté ; il n'est pas certain que le score qui en résulte reflète seulement la mémoire de travail. Cependant, que ce soit une limite cognitive intrinsèque ou l'indice d'un réseau complexe

de compétences personnelles, l'entité mesurée est dans tous les cas un facteur de variation individuelle, et nous est utile en tant que tel.

La méthode d'évaluation est celle proposée par Daneman et Carpenter en 1980 (Daneman et Carpenter, 1980). Le sujet lit à voix haute une série de phrases consécutives, sans pauses entre chaque phrase. Une fois la série de phrases terminée, il doit dire le dernier mot de chacune des phrases contenues dans la série. Le sujet est exposé à des séries contenant de plus en plus de phrases (de 2 à 6 phrases), jusqu'à ce qu'une erreur de rappel soit commise. Le nombre de phrases contenu dans la série où l'erreur se trouve est noté (ex. : 4), et l'exercice est répété deux autres fois avec des séries et phrases différentes. La cueillette de données complètes résulte donc en un trio d'entiers (ex. : 4,5,5) qui représente la capacité de mémoire de travail linguistique écrite de chaque participant. Il y a plusieurs façons de compiler ces entiers pour en faire un score unique, nous avons décidé d'en faire une moyenne (ex. 4.67).

Si la méthode utilisée a très peu changé depuis 1980, les items que nous avons affichés sont substantiellement différents. Le matériel utilisé a été créé par Denis Foucambert et al. en 2014, fortement inspiré des observations de Desmettes et al. (Desmette *et al.*, 1995). En 2.3.1 se trouve un résumé des contraintes liées au processus de créations des items, qui eux-mêmes peuvent être retrouvés en Annexe A.

2.3.1 Matériel pour test de mémoire de travail

Les mots de 2 syllabes avec la plus grande fréquence lexicale (New *et al.*, 2004; Ferrand *et al.*, 2001) ont été extraits de corpus de mots concrets (Ferrand et Alario, 1998) et abstraits (Ferrand, 2001) de façon à en obtenir 30 de chaque. Ensuite, 60 phrases précédant ces mots ont été créées en respectant les contraintes listées ci-dessous :

- Les phrases contiennent entre 13 et 15 mots.
- Les phrases contiennent entre 22 et 25 syllabes.
- Les phrases ne contiennent pas de virgules.
- Les constituants syntaxiques sont dans leur ordre canonique.
- Les noms et adjectifs se trouvent tous une seule fois dans le matériel.
- Les phrases ne contiennent pas de mots ayant subi de correction dans la réforme orthographique de 1990.
- Les phrases ne contiennent pas de temps de verbes composés.

- Les phrases ne contiennent pas de participe passé.
- Aucune phrase ne se trouve dans sa forme passive.
- Les homophones **tous/tout**, **ses/ces** et **leur/leurs** sont absents.

En plus de respecter ces diverses contraintes, les phrases furent construites avec le souci d'éviter des accords complexes inaudibles ainsi qu'avec celui de sélectionner des mots d'un vocabulaire non soutenu dont l'orthographe est facile.

Finalement, les phrases ont été rassemblées en séries de façon à éviter des liens sémantiques entre les phrases et entre les derniers mots de celles-ci, tout en assurant une répartition équilibrée des mots abstraits et concrets au sein des séries.

2.3.2 Matériel

L'oculomètre utilisé est un Tobii Pro Fusion 250Hz¹, manufacturé par Tobii AB, de la compagnie Tobii Pro. Ce produit fonctionne par détection de la réflexion cornéenne et vient équipé de deux caméras, pour une détection de l'œil plus précise et pour la création d'un modèle 3D. Il vient aussi équipé de deux systèmes d'illumination de la pupille, pour éviter des complications connues par rapport aux différences d'ajustement requis en lien avec la couleur de l'iris. L'oculomètre change automatiquement de l'un à l'autre selon la situation. Selon le mode d'évaluation, le taux d'échantillonnage diffère, soit 250Hz pour le "bright mode", et 120 Hz pour le "Dark mode". Il est capable d'obtenir une précision de 0.04 degré, et une exactitude de 0.3 degré. Le moniteur est un Asus TUF Gaming VG259Q, avec une taille de 24.5", une profondeur de couleur de 8-bits, une résolution de 1920 par 1080, et une fréquence de rafraîchissement de 59.94 Hz.

Nous avons utilisé Tobii Pro Lab pour présenter l'expérience et extraire les données oculométriques. Pour l'instant, la randomisation dans la présentation des items sur Tobii Pro Lab ne permet pas d'insérer des

1. Pour s'assurer que 250 Hz représente une résolution temporelle suffisante, nous utilisons un heuristique sur la base d'un raisonnement inspiré du théorème Nyquist-Shannon (Nyquist, 1928; Shannon, 2006) selon lequel la fréquence d'échantillonnage doit être au moins deux fois supérieure à la fréquence maximale présente dans le signal observé. Ce principe est couramment adapté en oculométrie pour motiver l'idée qu'il faut disposer d'au moins 2 échantillons à l'intérieur d'un événement de durée D pour pouvoir le détecter et en estimer la temporalité (Holmqvist *et al.*, 2011; Andersson *et al.*, 2010). À 250 Hz (4 ms entre chaque échantillon), même des fixations exceptionnellement courtes d'une durée de 50 ms sont représentées par environ 12 échantillons, ce qui dépasse largement le minimum requis pour une estimation fiable de la durée.

conditions, ce qui permettrait d'instancier un carré latin. Ainsi, à la place d'utiliser un système de carré latin permettant de minimiser la probabilité que deux participants soient exposés aux mêmes items dans les mêmes exactes conditions, quatre versions (timeline selon le lexique de Tobii Pro Lab) de la même expérience ont du être créées, pour répartir les 4 conditions expérimentales comme on peut voir dans le tableau 2.2. Cela a eu l'effet d'alourdir le projet et d'augmenter sévèrement le coût computationnel de l'extraction des données.

Items expérimentaux	Version 1	Version 2	Version 3	Version 4
1-5	facil. moins	neutre moins	facil. plus	neutre plus
6-10	neutre plus	facil. moins	neutre moins	facil. plus
11-15	facil. plus	neutre plus	facil. moins	neutre moins
16-20	neutre moins	facil. plus	neutre plus	facil. moins

Table 2.2 Répartition des conditions selon la version

Les caractères utilisés pour l'instanciation des items étaient de police *Times New Roman*, de taille 36, et de couleur noir (0, 0, 0) sur un fond gris (211, 211, 211).

CHAPITRE 3

ANALYSES

Ce chapitre présente les processus appliqués lors des analyses statistiques, ainsi que les résultats obtenus. Il ne s'agit ici que d'une exposition, la discussion permettant d'expliquer ce que les résultats veulent dire. Tout d'abord, le processus de transformation des données brutes et d'élimination des données est exposé. L'analyse statistique est ensuite divisée en deux niveaux distincts, chacun présenté dans sa propre section de ce chapitre. Premièrement, les statistiques sont présentées en prenant les phrases comme unité d'analyse. Ensuite, les mesures sont prélevées au niveau des mots, résultant en plusieurs analyses : sur tous les mots du corpus expérimental, sur les mots des zones (telles que décrite en 2.2.3.1 en se basant sur les travaux de (Winterstein *et al.*, 2014)), puis finalement sur chacun des mots se trouvant à proximité de *mais* dans ce que nous appelons la **zone critique**, décrite en 3.3.

3.1 Transformation et élimination des données

Les résultats bruts obtenus par l'oculomètre sont des coordonnées représentant la portion de l'écran vers laquelle chaque œil est orienté, et ce à chaque deux cent cinquantième de seconde (0.004 seconde). Il arrive parfois que les coordonnées du point de regard d'un œil ne soient pas disponibles. Cela peut arriver lorsque quelqu'un cligne des yeux, qu'il y ait une poussière devant l'appareil, ou dans toutes sortes d'autres circonstances. À partir du moment où les données ne sont disponibles pour aucun des deux yeux, ce point de donnée est considéré comme invalide et ne paraît pas dans les jeux de données servant à nos analyses. Autrement, la moyenne des coordonnées des deux yeux est calculée et est utilisée pour les calculs subséquents. Les items expérimentaux pour lesquels le taux de points de données invalides dépassaient 33% ont été exclus.

Nous avons appliqué le filtre de fixation Tobii I-VT (TobiiAB, 2023). Celui-ci prend les coordonnées correspondant à la moyenne des points de fixation des deux yeux et applique une formule de réduction de bruit de moyenne mouvante avec une fenêtre de trois échantillons. À l'intérieur de trois de ces fenêtres, donc de 20ms, la vélocité doit être inférieure à 30° par seconde pour que les points de données soient considérés comme faisant partie d'une seule fixation. Sinon, ils font partie d'une saccade. Les fixations à moins de 0.5° de distance et séparées par moins de 75ms furent considérées comme faisant partie d'une même fixation. Les fixations de moins de 60ms furent rejetées. L'ensemble de données après l'application des filtres

contenait 14 470 fixations.

Notre 19e item expérimental, nommé *PI19* ou *text* (448), dépendant de sa condition, était plus long que les autres. Même si sa quantité de caractères était à l'intérieur des balises établies dans la section 2.2.1, il était assez long pour être affiché sur deux lignes à l'écran. Pour ne pas biaiser inutilement les données, cet item fut retiré de l'analyse. La quantité de phrases à analyser a donc été réduite de 945 à 896 et la quantité de mots de 11 700 à 10 902.

La croix de fixation est un outil méthodologique permettant d'amener le regard des sujets en dehors des stimuli expérimentaux. Or, les 27 premiers participants se voyaient présenter, entre chaque texte, une croix de fixation se situant au centre de l'écran, qui est le paramètre par défaut du logiciel utilisé. Soupçonnant un effet d'amorçage causé par le fait que la première fixation se trouve au milieu de la phrase, où des mots clés comme *mais*, *moins* et *plus* se situent, nous avons fait passer 28 autres participants, pour lesquels la croix de fixation se situait à gauche de la phrase, évitant ainsi le plus possible l'effet d'amorçage soupçonné. La croix de fixation changeant significativement l'effet des variables qui nous intéressent, **valence** et **contexte**, nous avons décidé d'ignorer les résultats des participants pour qui la croix se situait au centre. Le retrait des participants pour qui la croix était au centre réduit la quantité de mots à analyser de 10902 à 5718 et le nombre de phrases de 896 à 466. Puisqu'il s'agit d'une perte de donnée considérable et que des effets significatifs semblent survenir lorsque la croix est au centre, une analyse plus poussée d'un potentiel amorçage entraîné par la croix serait intéressante, mais ne sera pas sujette à investigation dans ce présent mémoire.

5 phrases contenaient un nombre suspicieusement bas de saccades, et ainsi une très petite durée totale de fixation. Ainsi, les phrases contenant 4 saccades et moins furent retirées, réduisant la quantité de phrases observées de 466 à 461, et le nombre de mots fixés observés de 5718 à 5703.

3.2 Résultats au niveau phrastique

Après l'exclusion des données jugées invalides décrite dans la section 3.1, il restait 461 observations de phrases dans notre jeu de données. À ce niveau d'analyse, 2 mesures furent observées, soit la probabilité de régression et la durée de fixation, qui peuvent être considérées comme mesures globales de la difficulté de lecture (Goldberg et Kotval, 1999). La formule `lmer()` du paquet R `lme4` (Bates *et al.*, 2015) nous permet d'analyser ces mesures du niveau phrastique en les prenant tour à tour comme variable dépendante. Ci-

dessous se trouve une définition de chaque variable indépendante insérée dans le modèle (3.2.1) ainsi que des précisions liées à notre interprétation des facteurs groupes (3.2.2) et à notre sélection finale du modèle (3.2.3). Ensuite, il s'y trouve une description des résultats obtenus suite à l'insertion dans ces modèles des variables dépendantes représentant les mesures de la durée de fixation par caractère (3.2.5) et de la probabilité de régression (3.2.4).

3.2.1 Définition des variables au niveau phrastique

En (1) se trouve un aperçu du modèle dont la présente section traite. Nous y décrivons ensuite individuellement chaque variable indépendante choisie comme effet fixe dans notre modèle.

$$(1) \quad VD \sim \text{Valence} + \text{Contexte} + \text{ScolCal} + \text{Empan} + (1|\text{Participant}) + (1|\text{Media})$$

VD, ou variable dépendante, y représente une de nos deux mesures extraites au niveau phrastique, qui seront détaillées dans les sections 3.2.4 et 3.2.5.

La **Valence** est catégorielle et indique si *mais* est suivi de *moins* ou *plus* dans l'item expérimental.

Le **Contexte** représente le type de court texte précédant l'affichage de la phrase expérimentale et varie entre deux catégories, soit **facilitateur** ou **neutre**. Pour une explication des différences entre ces deux types de contexte, référez-vous à la section 2.2.2.

ScolCal véhicule l'information relative au calibre scolaire de nos participants. Il s'agit d'un nombre entier variant entre 1 et 4 qui représente le dernier niveau de scolarité achevé de chaque participant ; 1 étant le secondaire, 2 étant le cégep¹, 3 étant le premier cycle universitaire et 4 étant les cycles supérieurs. Cette variable ainsi définie a la particularité de présumer une augmentation plus grande entre le secondaire et le cégep ($3 - 2 = 1$) qu'entre la maîtrise et le doctorat ($4 - 3 = 1$). Cette vision de la scolarité comme étant numériquement ordonnée implique aussi que chaque niveau supérieur vient après l'autre et qu'il s'agit d'une progression linéaire, ce qui n'est pas nécessairement le cas, puisqu'il est en effet possible d'atteindre

1. Le Cégep (Collège d'Enseignement Général Et Supérieur), est un niveau de scolarité unique au Québec pour lequel l'âge minimal d'entrée tourne autour de 17 à 18 ans, sans âge maximal, créé dans le but de rendre l'éducation supérieure accessible et abordable.

un palier en en éludant d'autres. Malgré ses défauts, cette variable ainsi définie est cruciale pour rendre compte de la prédiction que les comportements de lectures changent par rapport à l'augmentation du niveau de scolarité (BENSOLTANA,). Puisqu'elle fait office de facteur contrôle, un effet significatif de cette variable ne sera pas abordé dans ce mémoire, car il ne s'agit ni d'une surprise ni d'un questionnement préalablement établi.

L'**Empan** représente la moyenne des résultats de 3 séries d'un test de mémoire de travail linguistique dont les processus de création et de mise en place sont décrits dans la section 2.3. Cette variable est incluse comme facteur de contrôle dans notre modèle, car une forte influence de la mémoire de travail sur les processus de lecture est attendue (Traxler *et al.*, 2012).

3.2.2 Considération des facteurs aléatoires

MediaType est une variable catégorielle contenant un chiffre de 1 à 20, qui représente le "type" de phrase affiché. Par exemple, la phrase suivante est de type 1, qu'elle contienne le mot "plus" ou "moins".

- (2) Au niveau de la carrure, Paul est grand, mais moins/plus grand que Pierre et il n'est pas facile de prendre une décision.

Malgré les mesures prises pour contrôler la variation entre les énoncés lors de la création des items, décrites en 2.2.3, il est certain qu'une variabilité inconnue et donc non contrôlée perdure. C'est pourquoi **MediaType** est utilisée ici comme facteur aléatoire.

Participant contient l'identifiant du participant. Ceux qui sont considérés dans cette analyse sont nommés de *Vrai28* à *Vrai55*, dont 13 mâles et 15 femelles âgés de 19 à 68 ans, avec une moyenne d'âge de 31.5, une médiane de 29.5 et un écart type de 10.56. Cette variable renferme une grande quantité de variations dont on ignore les sources exactes, chaque individu étant aussi idiosyncratique et précieusement unique qu'un flocon de neige, et est donc considéré comme facteur aléatoire.

Pour les deux facteurs aléatoires, **Media** et **Participant**, 4 syntaxes différentes, affichées en (3) ont été considérées, ayant toutes leur propre ensemble d'impacts sur la régression et ainsi nos analyses.

- (3) a. Effets fixes : ... + **Media** + **Participant**
 b. Intercepts aléatoires : ... + (1|**Media**) + (1|**Participant**)
 c. Pentes aléatoires : ... + (nbCarac | **Media**) + (empan | **Participant**)
 d. Pentes et intercepts aléatoires : ... + (1 + nbCarac | **Media**) + (1 + empan | **Participant**)

Le fait de considérer les facteurs groupes **Media** et **Participant** comme des effets fixes, illustré en (3-a), aurait le désavantage de ne pas prendre en compte que pour une variable dépendante donnée, disons la durée totale de fixation, chaque participant et chaque item aura sa propre moyenne. Cela réduirait considérablement le potentiel de généralisation de l'analyse, ainsi cette syntaxe fut rejetée, car trop simpliste.

Le fait d'instancier les facteurs **Media** et **Participant** avec des pentes aléatoires, comme illustré en (3-c) et (3-d), aurait l'avantage et inconvénient de permettre une analyse plus complexe de la variation entre les groupes. Cela permettrait d'un côté de rendre compte du fait que l'effet de la **mémoire de travail**, par exemple, n'aura pas le même effet sur la durée totale de fixation pour tous les participants, permettant une analyse plus précise et fidèle au phénomène observé. Cependant, la quantité de données requise pour permettre la généralisation de cette fidélité augmenterait substantiellement, en même temps que la complexité du modèle et, avec celle-ci, la probabilité que le modèle ne converge pas ou soit reconnu coupable de surapprentissage. De plus, pour interpréter les résultats avec des facteurs aléatoires qui ont une syntaxe comme **prédicteur|groupe**, il serait nécessaire de générer des hypothèses connexes outre celles centrales à mon mémoire dans le but de justifier l'impact du prédicteur sur la variation du groupe. Pour ces raisons, les syntaxes en (3-c) et (3-d) furent rejetées.

Par élimination surgit notre grand vainqueur, soit la vision des intercepts comme étant aléatoires, mais pas les pentes, tel qu'illustré en (3-b). Cette syntaxe permet un équilibre acceptable entre la complexité et la puissance du modèle, c'est donc comme cela que nos facteurs groupes seront instanciés dans nos modèles tout au long de l'analyse. Notons que ce choix à pour effet que la fonction lmer() du paquet lme4() interprète le modèle comme étant à effets aléatoires croisés. Puisque tous les participants ont vu tous les items, cette interprétation est jugée satisfaisante. Notons cependant qu'un ICC unique est généré pour le croisement des deux facteurs aléatoires.

3.2.3 Sélection du modèle au niveau phrastique

Pour nous assurer que les modèles soient adéquatement ajustés, nous avons, pour chaque variable dépendante, fait tourner quatre modèles, dans lesquels **VD** représente la variable dépendante, **FC** les facteurs contrôles et **FA** les facteurs aléatoires.

(4) **Modèle 1** : $VD \sim \text{Valence} + FC + FA$

Modèle 2 : $VD \sim \text{Contexte} + FC + FA$

Modèle 3 : $VD \sim \text{Valence} + \text{Contexte} + FC + FA$

Modèle 4 : $VD \sim \text{Valence} * \text{Contexte} + FC + FA$

En comparant ces différents modèles avec la fonction R `anova(mx, my)`, il a été découvert que le modèle 2, ne contenant que **Contexte**, est ajusté aux données de manière significativement différente du modèle 3, indiquant un poids appréciable de la variable **Valence** dans l'explicabilité de la variation présente dans les différentes variables dépendantes. Puisqu'il n'y avait pas de différences significatives entre les ajustements des modèles 3 et 4, il a été décidé d'ignorer les interactions pour prioriser, à ajustements égaux, le modèle le plus simple qui, dans ce cas est le modèle 3, dont la forme est identique au modèle montré en (1), répété ici en (5).

(5) $VD \sim \text{Valence} + \text{Contexte} + \text{ScolCal} + \text{Empan} + (1|\text{Participant}) + (1|\text{Media})$

3.2.4 La probabilité de régression

La probabilité de régression est calculée comme le ratio du nombre de régressions oculaires divisé par le nombre de saccades. Elle représente ainsi la probabilité, pour une saccade, que celle-ci soit vers la gauche, donc régressive. Dans la figure 3.1, on peut voir un aperçu de la probabilité de régression, divisé selon nos conditions expérimentales. On peut y voir que les boîtes rouges et jaunes, qui représentent les phrases avec "plus", sont placées marginalement plus hautes que les autres, indiquant peut-être que la probabilité de régression est plus haute lorsque la phrase contient "plus".

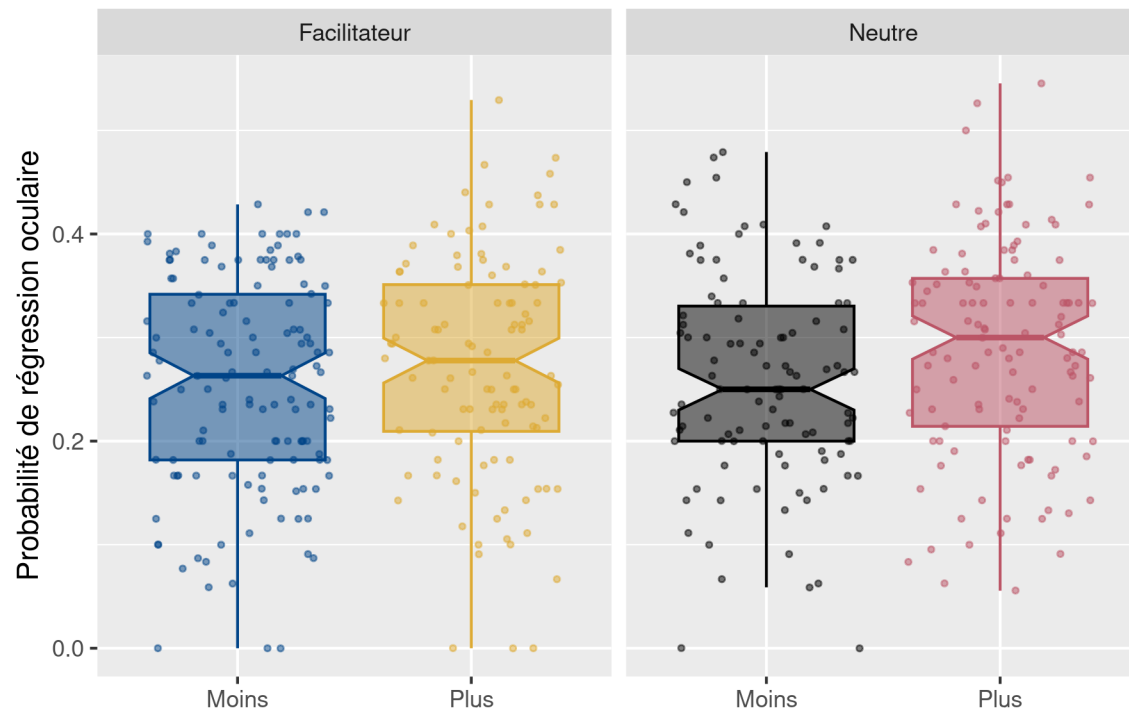


Figure 3.1 Probabilités de régression, séparées par conditions

Pour vérifier que cet effet n'est pas qu'une illusion, nous vérifions avec notre modèle de régression linéaire multiniveau, qui nous confirme que le fait de changer "moins" pour "plus" à l'effet significatif de faire augmenter la probabilité de régression de 3% en moyenne.

<i>Predictors</i>	probReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	0.31	-0.17	0.18 – 0.44	-0.41 – 0.08	<0.001
Valence [Plus]	0.02	0.22	0.01 – 0.04	0.06 – 0.39	0.007
Contexte [Neutre]	0.01	0.12	-0.00 – 0.03	-0.04 – 0.29	0.142
ScolCal	-0.00	-0.02	-0.03 – 0.02	-0.18 – 0.14	0.789
Empan	-0.01	-0.06	-0.05 – 0.02	-0.20 – 0.08	0.430
Random Effects					
σ^2	0.01				
T00 Participant	0.00				
T00 MediaType	0.00				
ICC	0.28				
N MediaType	19				
N Participant	28				
Observations	461				
Marginal R ² / Conditional R ²	0.020 / 0.292				

Table 3.1 Résultats du modèle pour la probabilité de régression

3.2.5 La durée de fixation totale

La durée de fixation est la somme en millisecondes de toutes les durées des fixations observées au cours de la lecture de chaque item.² Comme on peut voir dans la figure 3.2 par le fait que les boîtes jaunes et rouges sont plus élevées que les deux autres, *mais plus* semble entraîner une durée de fixation totale plus grande.

2. Cette mesure fut prise simplement plutôt qu'en divisant par le nombre de caractères, ce qui aurait causé des problèmes vu la relation non linéaire entre la durée de fixation et le nombre de caractères (Trueswell *et al.*, 1994).

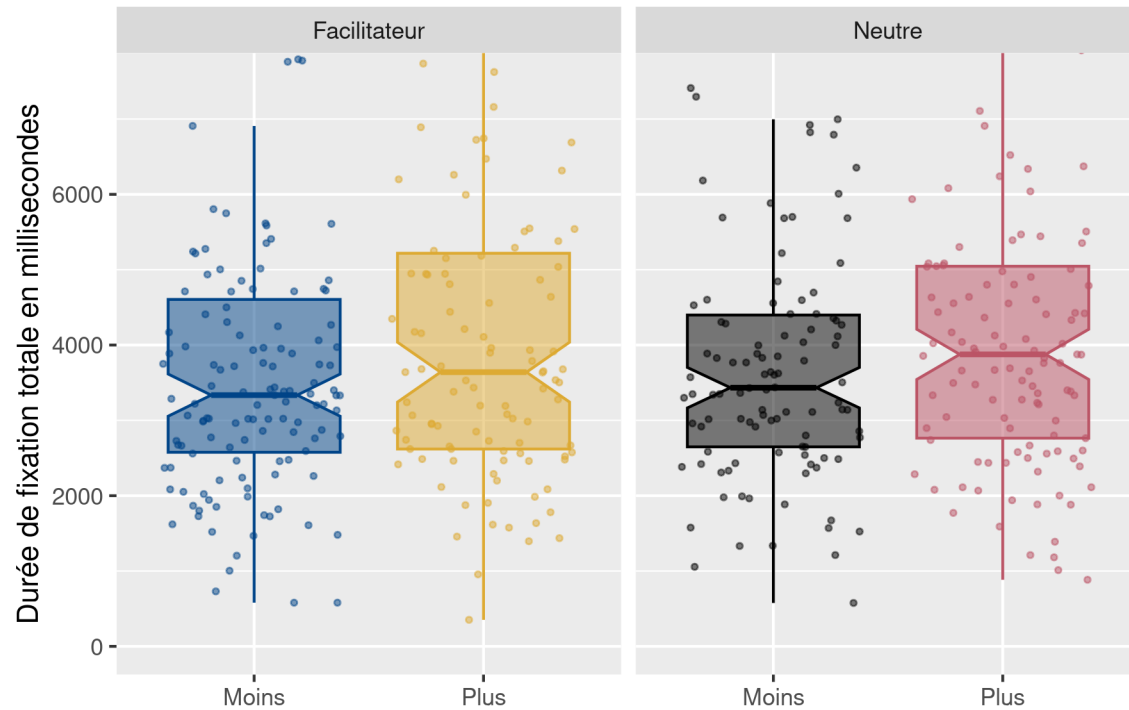


Figure 3.2 Durée de fixation en millisecondes, séparée selon les variables **Valence** et **Contexte**

La durée de fixation fut considérée comme variable dépendante de notre modèle. Cependant, lorsque cette mesure était insérée telle quelle, la distribution des résidus du modèle n'était pas normale, avec un kurtosis de 13.27157 et une asymétrie de 4.374365. Pour pallier à ce problème, la variable dépendante fut transformée avec la fonction $\log()$, résultant en la variable **logDurFix**. Lorsque celle-ci était insérée comme variable dépendante, la distribution des résidus du modèle montrait un kurtosis de 4.374365 et une asymétrie de -0.103389, ce qui a été jugé suffisamment normal.

<i>Predictors</i>	logdurFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	8.21	-0.16	7.56 – 8.87	-0.44 – 0.12	<0.001
Valence [Plus]	0.12	0.22	0.04 – 0.19	0.08 – 0.37	0.003
Contexte [Neutre]	0.05	0.09	-0.03 – 0.12	-0.06 – 0.23	0.233
ScolCal	0.03	0.04	-0.09 – 0.14	-0.13 – 0.20	0.675
Empan	-0.05	-0.04	-0.22 – 0.12	-0.18 – 0.10	0.587
Random Effects					
σ^2	0.16				
T00 Participant	0.10				
T00 MediaType	0.02				
ICC	0.43				
N MediaType	19				
N Participant	28				
Observations	461				
Marginal R ² / Conditional R ²	0.017 / 0.438				

Table 3.2 Résultats du modèle pour le logarithme de la durée de fixation en millisecondes pour toute la phrase

L'impression véhiculée par la figure 3.2 est confirmée par les résultats du modèle. En effet, le fait de changer *moins* pour *plus* a l'effet significatif d'augmenter le logarithme des durées de fixations de 0.22 écart type en moyenne, comme on peut voir dans le tableau en 3.2, ce qui constitue un effet faible, mais non-négligeable.

3.3 Second niveau d'analyse : Les mots

Pendant la cueillette de données, dès qu'un mot était fixé au moins une fois, il gagnait sa place en tant que rangée dans notre jeu de données des mots, résultant en un total de 11 700 rangées au total. Après l'élimination de données décrite en 3.1, il en restait 5943.

Chaque mot se voit attribuer deux index pour identifier sa position dans la phrase. La variable **phrIndex** attribue au premier mot de la phrase l'index 1, puis augmente de 1 pour chaque mot subséquent. **ZeroMaisIndex** est un index bâti en prenant le mot *mais* comme étant le point 0. Le mot avant aura donc l'index -1, le mot suivant directement *mais* aura l'index 1, et ainsi de suite. Contraints par les diverses considérations

lors de la création des items (voir 2.2.1), les phrases et les zones qui les remplissent ne contiennent pas toutes le même nombre de mots. Ainsi, s'il y a certains index où tous les mots seront de types, de classes grammaticales et de fonction dans la phrase similaire (comme le mot 0 selon **ZeroMaisIndex**, qui est toujours le mot *mais*), d'autres sont associés à des mots sans liens concrets entre eux (comme le mot 6 selon **ZeroMaisIndex**, qui est parfois le dernier mot de la zone 2, parfois le premier de la zone 3, et qui peut faire partie de n'importe quelle catégorie grammaticale.)

Pour pallier à ce problème, j'ai créée *post-hoc* la **Zone Critique**.³ Cette zone comporte les mots de -3 à 4 selon **ZeroMaisIndex**, qui sont les mots avec le moins de variation, comme on peut voir dans le gabarit présenté en 3.3.

ZeroMaisIndex	-3	-2	-1	0	1	2	3	4
Type de mot	Entité1	est	adjectif	mais	plus/moins	adjectif	que	Entité2
Exemple	Paul	est	grand	mais	plus	grand	que	Pierre

Table 3.3

Le modèle utilisé pour l'analyse des mots est le même que celui décrit en 3.2, à l'exception de 2 facteurs fixes additionnels.

nbCarac représente le nombre de caractères au sein de chaque mot, ce qui pourrait influencer maintes mesures, dont la durée de première fixation, et la quantité de refixations à l'intérieur du premier regard (Rayner, 1998).

sFreqLex désigne la fréquence lexicale de chaque mot selon Lexique3, prélevée du corpus Frantext selon le processus décrit en 2.2.1, mise à l'échelle à l'aide de la fonction R `scale()` (Becker *et al.*, 1988).

Pour chaque mesure prélevée, nous visualiserons tout d'abord l'effet de nos variables **Contexte** et **Valence** sur la mesure pour l'ensemble des mots de l'échantillon, nous donnant une impression globale des résultats. Ensuite, nous montrerons, s'il y a lieu, pour quelles zones la **Valence** et le **Contexte** influencent significati-

3. Le fait d'annoter, pour chaque mot, sa catégorie grammaticale ainsi que sa fonction dans notre expérience nous permettrait de voir l'influence, par exemple, de la deuxième entité nommée peu importe où elle se trouve dans la phrase. Cependant, puisque cette idée ne m'est venue à l'esprit que trop tard, il n'en sera pas question dans ce mémoire.

vement la mesure observée. Finalement, nous plongerons mot par mot dans la **Zone Critique**. Les tableaux d'analyses ne se trouvant pas dans cette section se trouvent en annexe.

3.3.1 Durée du parcours régressif

La durée du parcours régressif est une traduction personnalisée de *regression path*, parfois aussi appelée *go-past duration*. Il s'agit de la somme de toutes les durées de fixation, peu importe sur quel mot, à partir du moment où le mot est fixé pour la première fois jusqu'à ce qu'un mot à sa droite soit fixé (Konieczny *et al.*, 1997). Cela calcule ainsi le temps passé sur le mot en soi, et le temps passé à retourner en arrière.

Lorsque prélevée sur tous les mots de notre jeu de données, cette mesure entraînait une distribution quasi bimodale des résidus de notre modèle. Pour pallier à ce problème, les mots pour lesquels la durée du parcours régressif est 0 furent retirés pour l'analyse de cette mesure, qui fut donc faite sur 5684 mots plutôt que 5943. Pour corriger un problème de distribution anormale qui persistait, une transformation $\log()$ fut aussi appliquée à la variable, changeant le kurtosis des résidus du modèle de 63.80232 à 6.594594 et l'asymétrie de 6.45065 à 0.5409211.

La visualisation fournie en 3.3 ne permet pas de détecter une grande différence, mis à part le fait que la boîte noire semble être un peu plus basse que les trois autres, suggérant peut-être un petit effet de la **Valence** ou du **Contexte**.

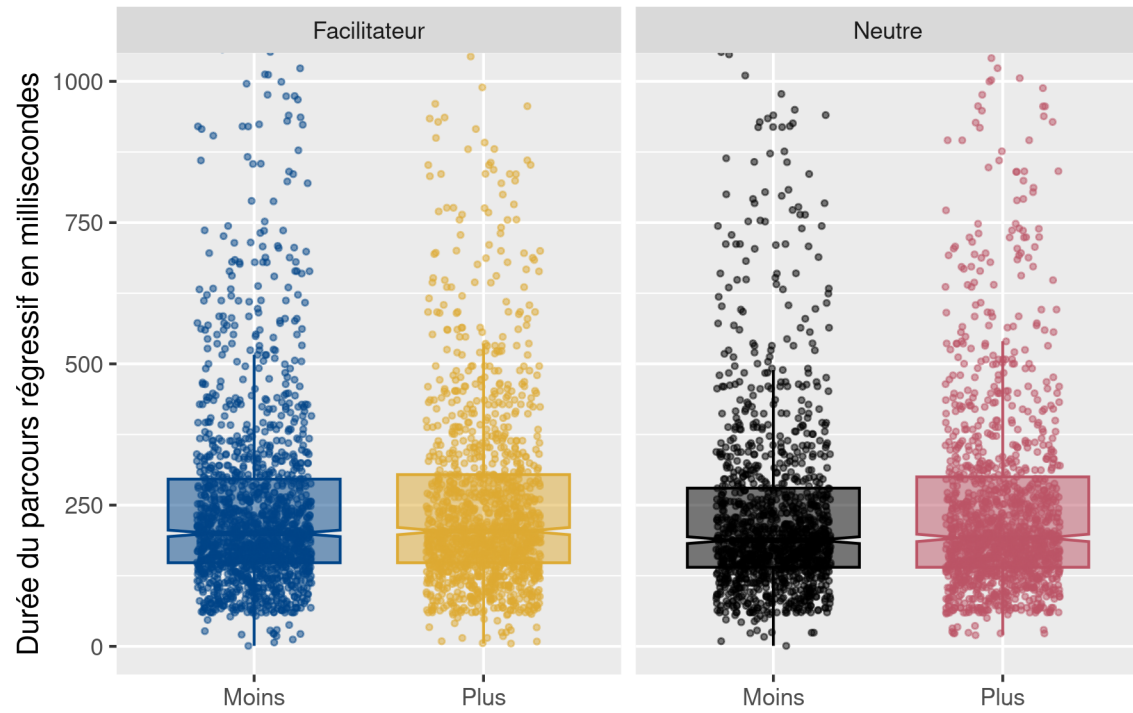


Figure 3.3 Durée du parcours régressif en millisecondes sur tous les mots

Les résultats du modèle en 3.4 montrent qu'il n'y a aucun effet significatif de la **Valence** ou du **Contexte** sur la durée du parcours régressif.

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.16	-0.02	4.67 – 5.65	-0.12 – 0.07	<0.001
Valence [Plus]	0.03	0.04	-0.01 – 0.07	-0.01 – 0.09	0.110
Contexte [Neutre]	-0.01	-0.01	-0.05 – 0.03	-0.07 – 0.04	0.634
ScolCal	0.05	0.05	-0.03 – 0.12	-0.03 – 0.12	0.209
Empan	-0.01	-0.00	-0.13 – 0.12	-0.08 – 0.07	0.911
sFreqLex	0.12	0.17	0.10 – 0.15	0.14 – 0.20	<0.001
nbCarac	0.02	0.08	0.01 – 0.03	0.05 – 0.11	<0.001
Random Effects					
σ^2	0.51				
T00 Participant	0.02				
T00 MediaType	0.01				
ICC	0.05				
N Participant	28				
N MediaType	19				
Observations	5684				
Marginal R ² / Conditional R ²	0.023 / 0.071				

Table 3.4 Résultats du modèle pour la durée du parcours régressif sur tous les mots.

En groupant les mots par zones et en se restreignant à celles-ci pour l'analyse de nos mesures, on peut voir que la **Valence** influence significativement la durée du parcours régressif sur les mots compris dans la zone 3, comme il est possible de voir dans la zone blanche de la figure 3.4.

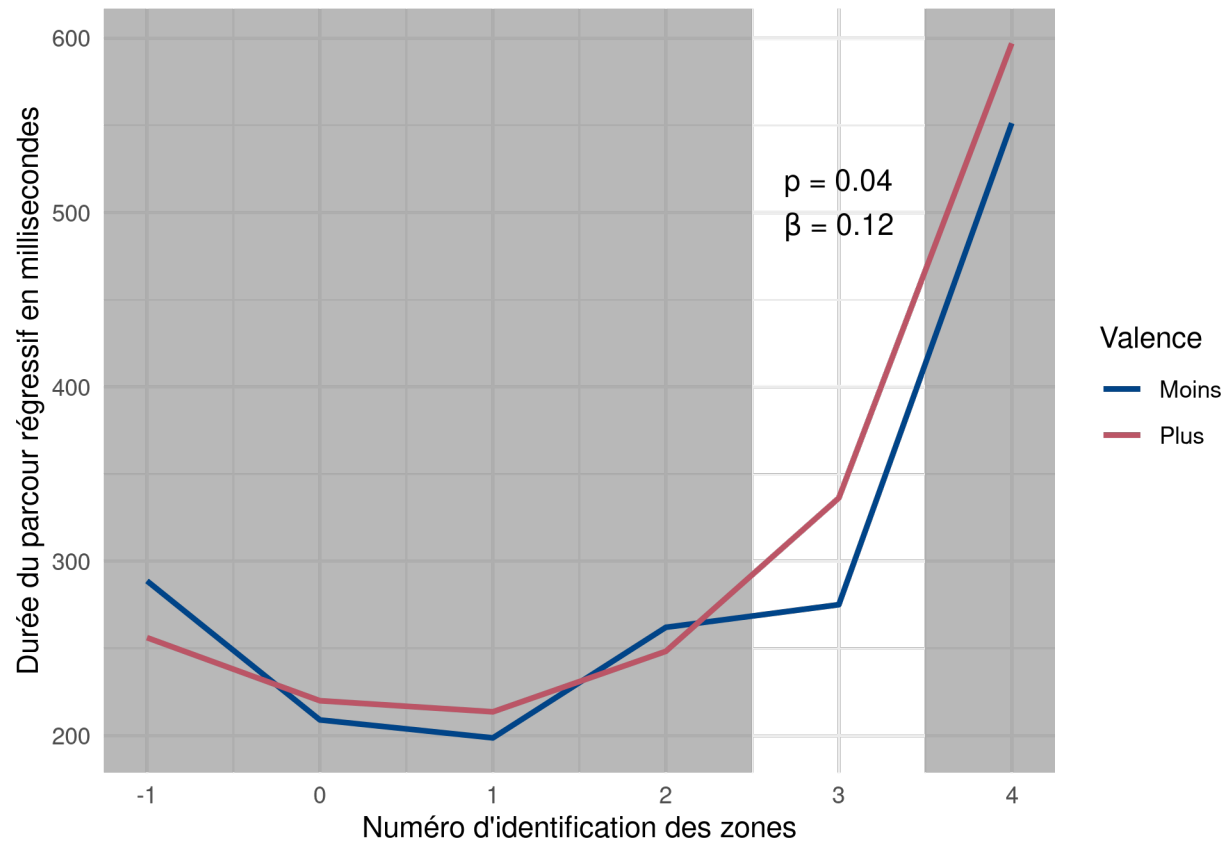


Figure 3.4 Moyenne de la durée du parcours régressif en millisecondes pour chaque zone, un fond blanc signifiant un effet significatif de la **Valence**

Le β positif ci-dessus montre que les phrases avec *mais plus* ont des plus grandes durées de parcours régressif sur les mots contenus dans la zone 3.

En observant le comportement de la durée du parcours régressif sur les mots de la zone critique (définie en 3.3, ex : *Paul est grand, mais plus grand que Pierre*), on découvre un effet significatif positif de la **Valence** au niveau de l'indexe 2 (prenant *mais* comme étant le point 0), soit celui où la réitération de la propriété gradable est juchée.

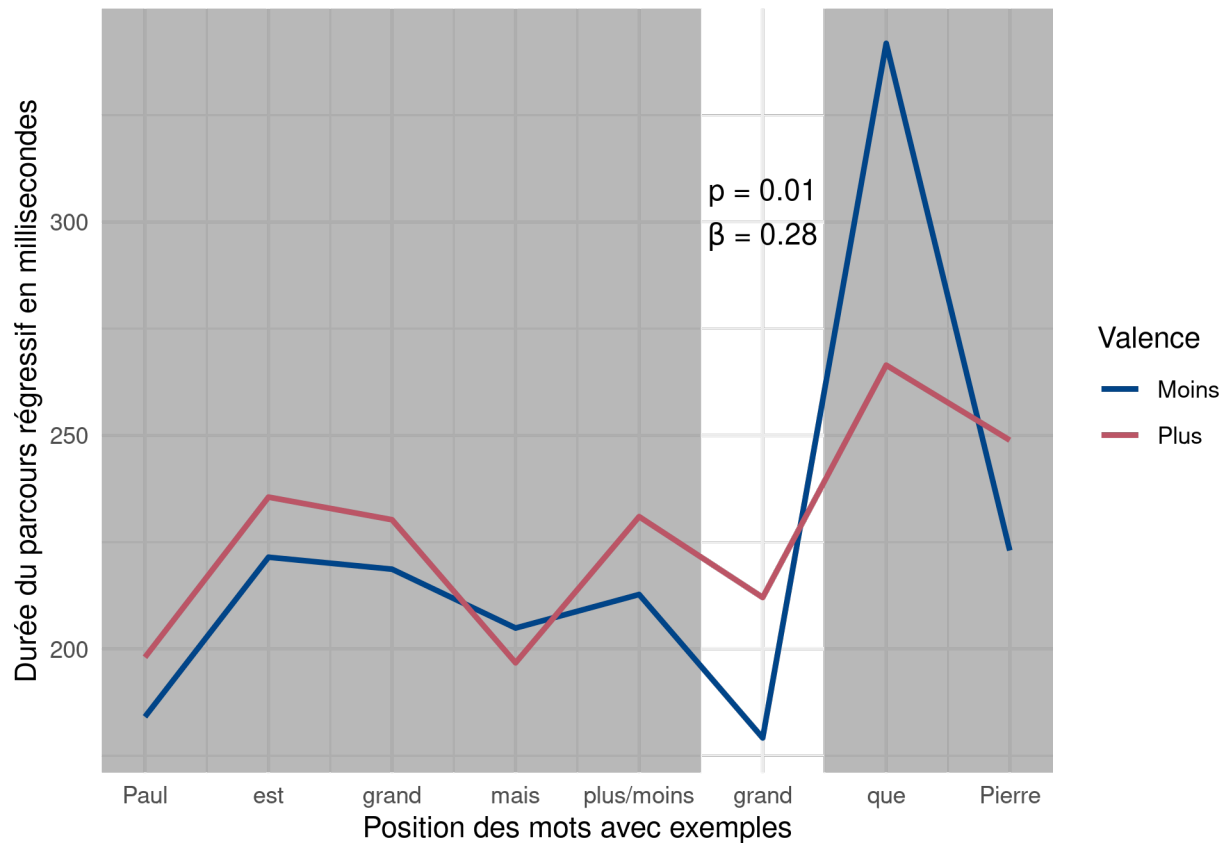


Figure 3.5 Durée du parcours régressif sur chaque mot de la zone critique, un fond blanc signifiant un effet significatif de la **Valence**

Le β positif en 3.5 montre que la durée du parcours régressif est plus haute au niveau de l'indexe 2 dans les phrases avec *mais plus* comparé aux phrases avec *mais moins*. Les tableaux montrant les résultats, significatifs ou non, des analyses des différentes zones ainsi que de chaque mot de la zone critique se trouvent ici, dans l'annexe C.

3.3.2 La durée de la première fixation

Lorsqu'un mot est fixé pour la première fois, la durée de cette fixation en millisecondes est simplement capturée et rangée dans la variable **durPremFix**. La durée de la première fixation est influencée par des processus très rapides, comme la reconnaissance des caractères (Rayner, 1998).

En prenant la phrase entière comme fenêtre d'analyse, la considération de la durée de la première fixation

par mots comme variable dépendante de notre modèle résultait en une distribution anormale des résidus de ce dernier. La variable **logdurPremFix** fut donc créée avec la transformation $\log()$, changeant le kurtosis des résidus de 80.46873 à 5.161747 et l'asymétrie de 5.155506 à -0.2210302. Cette variable ainsi transformée montrant des résidus appréciablement normaux pour les modèles des analyses subséquentes, elle est utilisée tout au long des analyses pour la durée de la première fixation.

Selon la représentation graphique de la distribution de cette mesure, aucun patron clair ne semble émerger, comme on peut voir dans la figure 3.6.

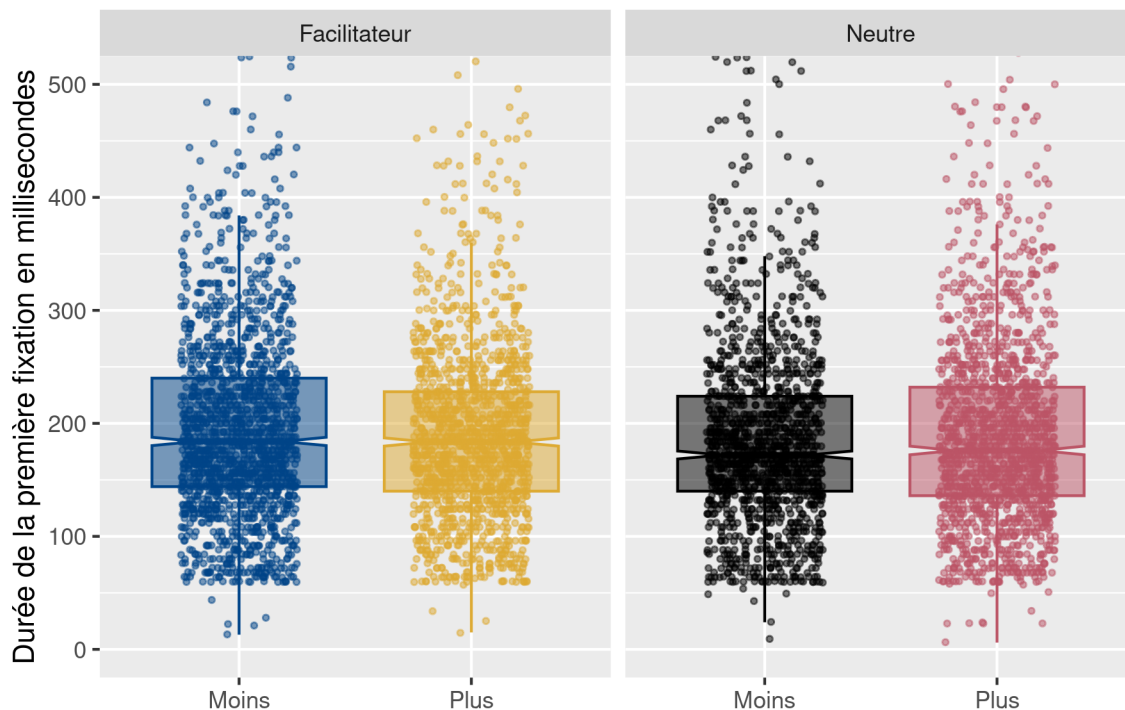


Figure 3.6 Durée de la première fixation en millisecondes pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables **Valence** et **Contexte**

Une très petite proportion de la variation de cette mesure est expliquée par notre modèle, soit seulement 3% sans et 12% avec les facteurs aléatoires. La **Valence** et le **Contexte** ne présentent aucun effet significatif, comme on peut voir ci-dessous en 3.5.

logdurPremFix			
<i>Predictors</i>	<i>Estimates</i>	<i>CI</i>	<i>p</i>
(Intercept)	4.94	4.47 – 5.41	<0.001
Valence [Plus]	-0.01	-0.03 – 0.01	0.446
Contexte [Neutre]	-0.01	-0.03 – 0.01	0.390
ZeroMaisIndex	-0.00	-0.01 – 0.00	0.211
phrIndex	0.01	0.01 – 0.02	<0.001
ScolCal	0.07	0.00 – 0.15	0.042
Empan	-0.03	-0.15 – 0.10	0.681
sFreqLex	-0.02	-0.03 – -0.00	0.026
nbCarac	0.00	-0.00 – 0.01	0.270
Random Effects			
σ^2	0.20		
T00 Participant	0.02		
ICC	0.09		
N Participant	28		
Observations	5943		
Marginal R ² / Conditional R ²	0.032 / 0.120		

Table 3.5 Résultats du modèle pour le logarithme de la durée de première fixation par mot, considérant tous les mots de l'échantillon

En séparant la phrase en ses zones ou bien en plongeant mot par mot dans la **zone critique**, aucun effet significatif de la durée de la première fixation ne fut trouvée à l'aide de nos modèle. Les tableaux véhiculant cette information se trouve dans l'annexe C, ici.

3.4 Durée du premier regard

La durée du premier regard consiste en la somme des durées de toutes les fixations à partir du moment où le mot est regardé pour la première fois jusqu'à la première fixation ayant lieu à l'extérieur du mot. Dans 87% des cas, cette mesure est équivalente à la durée de la première fixation, puisqu'il arrive souvent qu'une seule fixation ne survienne avant l'exécution d'une saccade sortante. Ces mesures étant très similaires, elles réagiraient de la même façon vis-à-vis les différents processus cognitifs dans la majorité des cas (Rayner, 1998). Cette mesure est cependant dite être affectée par la probabilité d'occurrence dans le contexte du mot observée, plutôt que sa simple probabilité d'occurrence dans la langue (Inhoff, 1984).

En observant la figure en 3.7 qui représente la durée du premier regard sur tous les mots, on peut voir que les boîtes bleues et rouges semblent être un peu plus hautes que les deux autres.⁴

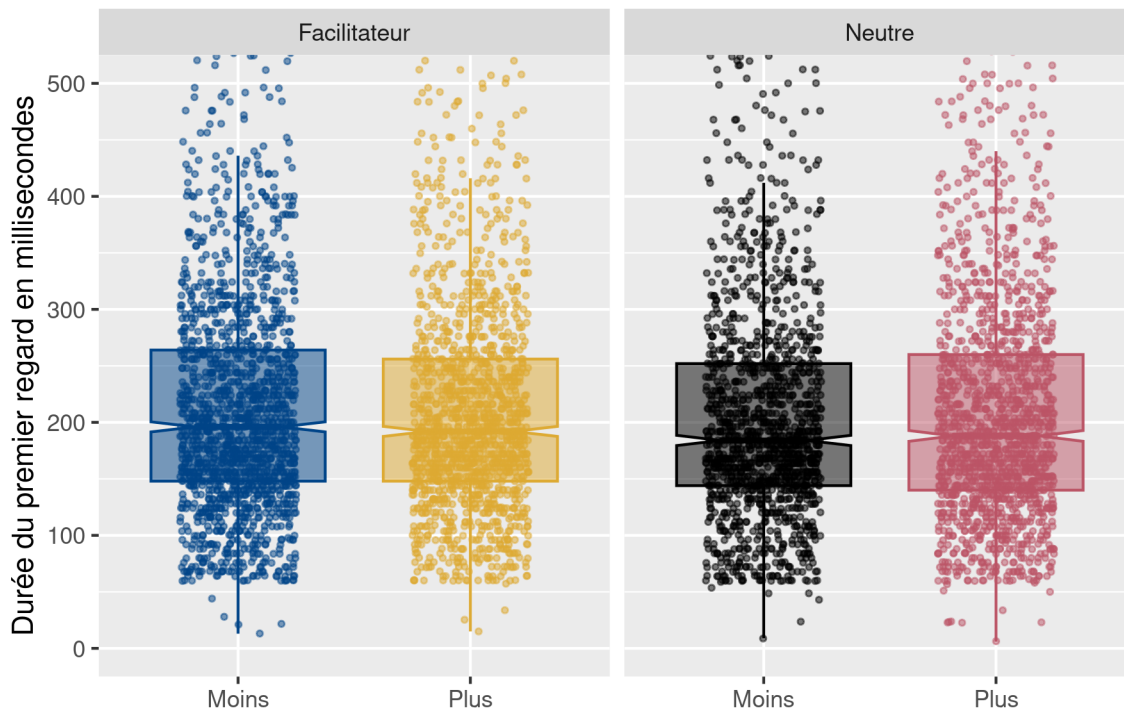


Figure 3.7 Durée du premier regard en millisecondes pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables **Valence** et **Contexte**

4. Cela ne pointe pas vers un effet clair d'une des deux variables, mais bien vers une possible interaction entre les deux, si ce n'est que très petite. L'analyse avec modèle de cette possible interaction c'est avérée inconcluante, voir annexe C.

En essayant la durée du premier regard comme variable dépendante de notre modèle, les résidus de celui-ci n'étaient pas distribués normalement. Ainsi, une transformation $\log()$ fut exécutée sur **durPremRegard** pour la transformer en **logdurPremRegard**, changeant le kurtosis des résidus du modèle de 216.9406 à 5.260592 et l'asymétrie de 2.605572 à 0.2498745. Les analyses subséquentes avec cette variable ainsi transformée ne nous posant plus de problème de distribution des résidus, nous utilisons **logdurPremRegard** pour toutes les analyses présentées dans cette section, ainsi que pour celles qui ne sont pas présentées et qui se trouvent en annexe.

<i>Predictors</i>	logdurPremRegard				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.04	-0.01	4.59 – 5.48	-0.11 – 0.09	<0.001
Valence [Plus]	0.01	0.02	-0.02 – 0.04	-0.03 – 0.07	0.454
Contexte [Neutre]	-0.01	-0.02	-0.04 – 0.02	-0.07 – 0.03	0.484
ScolCal	0.10	0.13	0.03 – 0.16	0.04 – 0.23	0.006
Empan	-0.04	-0.03	-0.15 – 0.08	-0.12 – 0.06	0.506
sFreqLex	-0.00	-0.00	-0.02 – 0.01	-0.03 – 0.03	0.891
nbCarac	0.03	0.15	0.02 – 0.03	0.12 – 0.18	<0.001
Random Effects					
σ^2	0.26				
T ₀₀ Participant	0.02				
T ₀₀ MediaType	0.00				
ICC	0.06				
N Participant	28				
N MediaType	19				
Observations	5943				
Marginal R ² / Conditional R ²	0.038 / 0.098				

Table 3.6 Résultats du modèle pour la durée du premier regard sur tous les mots.

En plus de n'expliquer qu'une petite proportion de la variation au sein de cette mesure (3% sans et 10% avec les facteurs aléatoires), aucun effet significatif de la **Valence** ou du **Contexte** n'a été détecté lorsque tout l'échantillon est considéré.

Similairement à la durée de première fixation, aucun effet significatif de la **Valence** ou du **Contexte** n'a été trouvé sur les analyses exécutées sur les mots contenus dans chacune des zones.

Cependant, une investigation mot par mot de la zone critique nous permet de voir que la **Valence** a un effet significatif positif sur la durée du premier regard pour la première apparition de la propriété gradable (ex : *grand*, $\beta = 0.21$, $p = 0.036$) ainsi que pour l'apparition de la deuxième entité (ex : *Pierre*, $\beta = 0.21$, $p = 0.038$).

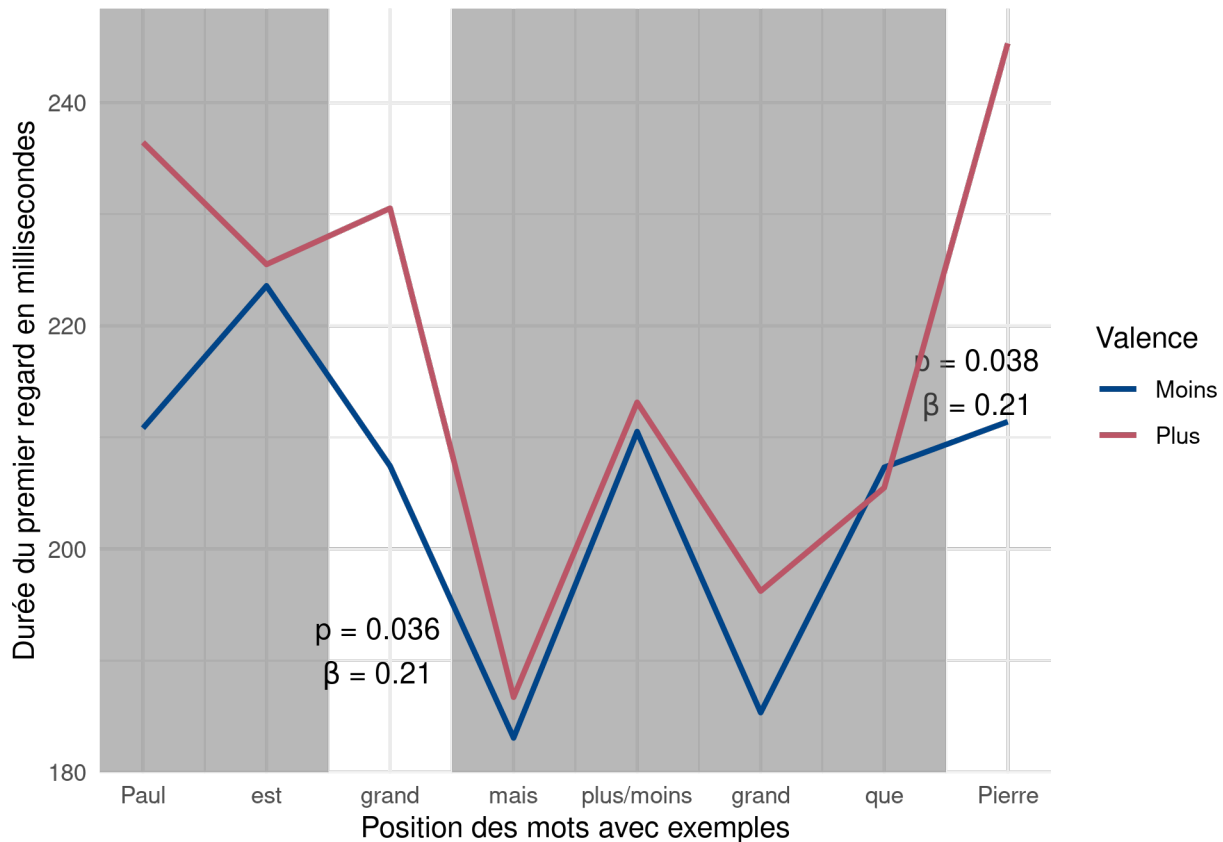


Figure 3.8 Durée du premier regard en millisecondes pour chaque mot de la zone critique, un fond blanc signifiant un effet significatif de la **Valence**

3.4.1 Probabilité de régressions sortantes

Cette mesure est une façon d'observer les mots qui sont à l'origine de mouvements oculaires régressifs. La probabilité de régressions sortantes est le nombre de régressions émanant de chaque mot, divisé par le nombre total de régressions oculaires durant la lecture de la phrase expérimentale. Cette mesure est inspirée de *regressions out* (Cook et Wei, 2019) qui permet de quantifier la propension d'une zone ou d'un mot à provoquer des régressions oculaires.

Plusieurs mots dans l'échantillon ne déclenchent pas de régressions oculaires, résultant en une quantité de 0 suffisante pour provoquer une distribution bimodale de la variable ainsi que des résidus du modèle. Pour contourner ce problème, nous avons retiré 3586 mots de l'échantillon. L'analyse de cette variable aura donc lieu sur les 2431 restants. Après l'insertion dans notre modèle de cette variable, la distribution n'était plus bimodale, mais était quand même anormale. Nous avons donc effectué une transformation $\log()$ de façon à changer le kurtosis de 20.04671 à 4.223336 et l'asymétrie de 3.053863 à 1.095626.

On peut voir en 3.9 que les boîtes bleu et noir représentant les probabilités de provoquer une régression oculaire pour les mots de phrases avec *mais moins* sont plus haute que les autres, indiquant un potentiel effet négatif de la **Valence**.

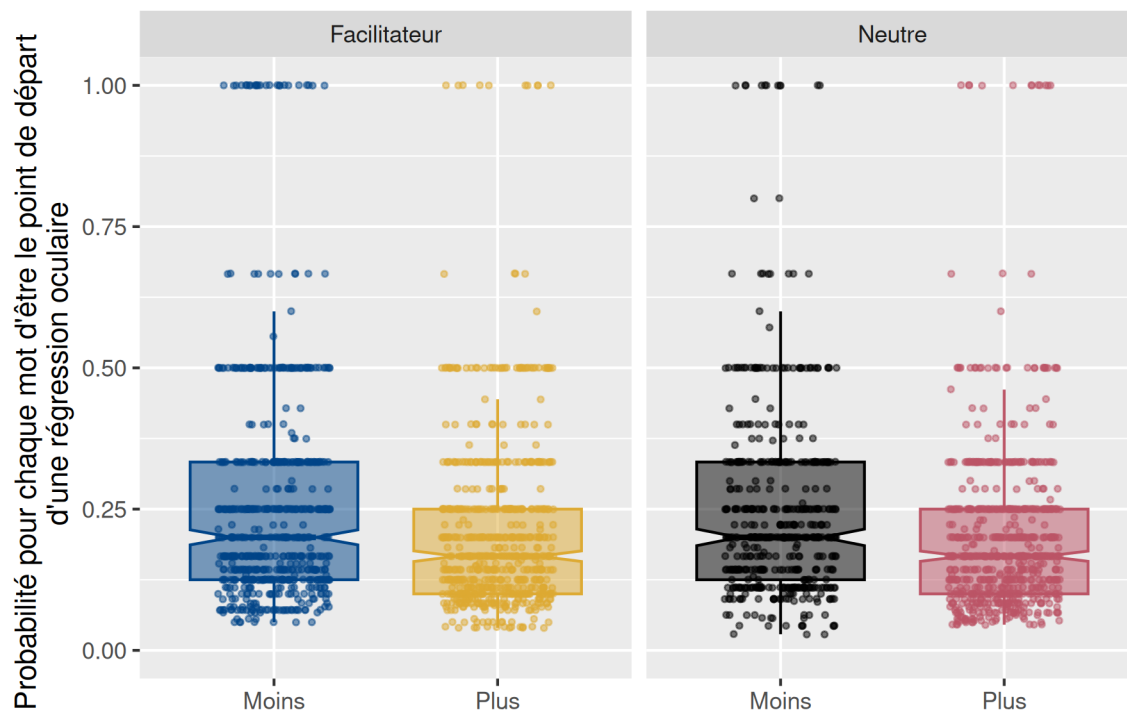


Figure 3.9 Probabilité d'être le point de départ d'une régression oculaire pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables **Valence** et **Contexte**

Rappelons-nous cependant que la quantité totale de régressions oculaires dans les phrases avec *mais plus* est plus haute. Ainsi, la probabilité par mot d'être le point de départ d'une régression oculaire sera divisée par un dénominateur plus petit pour les phrases avec *mais moins*, résultant en des chiffres plus gros. Pour empêcher que cette mesure ne soit que fonction de la quantité totale de régressions au sein de la phrase,

le nombre total de régressions dans la phrase, **nbRegPhrase**, est utilisé dans les modèles subséquents à des fins de contrôle. Corroborant notre intuition par rapport à la représentation graphique, les modèles montrent un effet significatif de la **Valence** ($\beta = -0.33$, $p < 0.001$) sur la probabilité d'être le point de départ d'une saccade régressive.

<i>Predictors</i>	logprobRegSort				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.67	0.34	-3.52 – -1.83	0.13 – 0.56	<0.001
Valence [Plus]	-0.20	-0.33	-0.24 – -0.16	-0.40 – -0.26	<0.001
Contexte [Neutre]	0.01	0.02	-0.03 – 0.05	-0.05 – 0.09	0.625
ScolCal	-0.05	-0.06	-0.18 – 0.08	-0.23 – 0.10	0.455
Empan	0.04	0.02	-0.18 – 0.25	-0.12 – 0.17	0.744
sFreqLex	0.05	0.09	0.03 – 0.07	0.05 – 0.13	<0.001
nbCarac	0.04	0.18	0.03 – 0.05	0.14 – 0.23	<0.001
Random Effects					
σ^2	0.24				
T ₀₀ Participant	0.06				
T ₀₀ MediaType	0.03				
ICC	0.28				
N Participant	28				
N MediaType	19				
Observations	2351				
Marginal R ² / Conditional R ²	0.058 / 0.326				

Table 3.7 Résultats du modèle pour le logarithme de la probabilité d'être le point de départ d'une régression oculaire

En prenant les mots de chaque zone individuellement comme échantillons pour notre analyse, on découvre tel qu'affiché en 3.10 que la pénultième zone montrent un effet significatif de la **Valence**. Dans le cas présent, cela signifie que les chances de retourner en arrière en rencontrant un mot augmentent pour les mots de la zone 3, affichée en blanc, pour les phrases avec *mais moins* par rapport aux phrases avec *mais plus*.

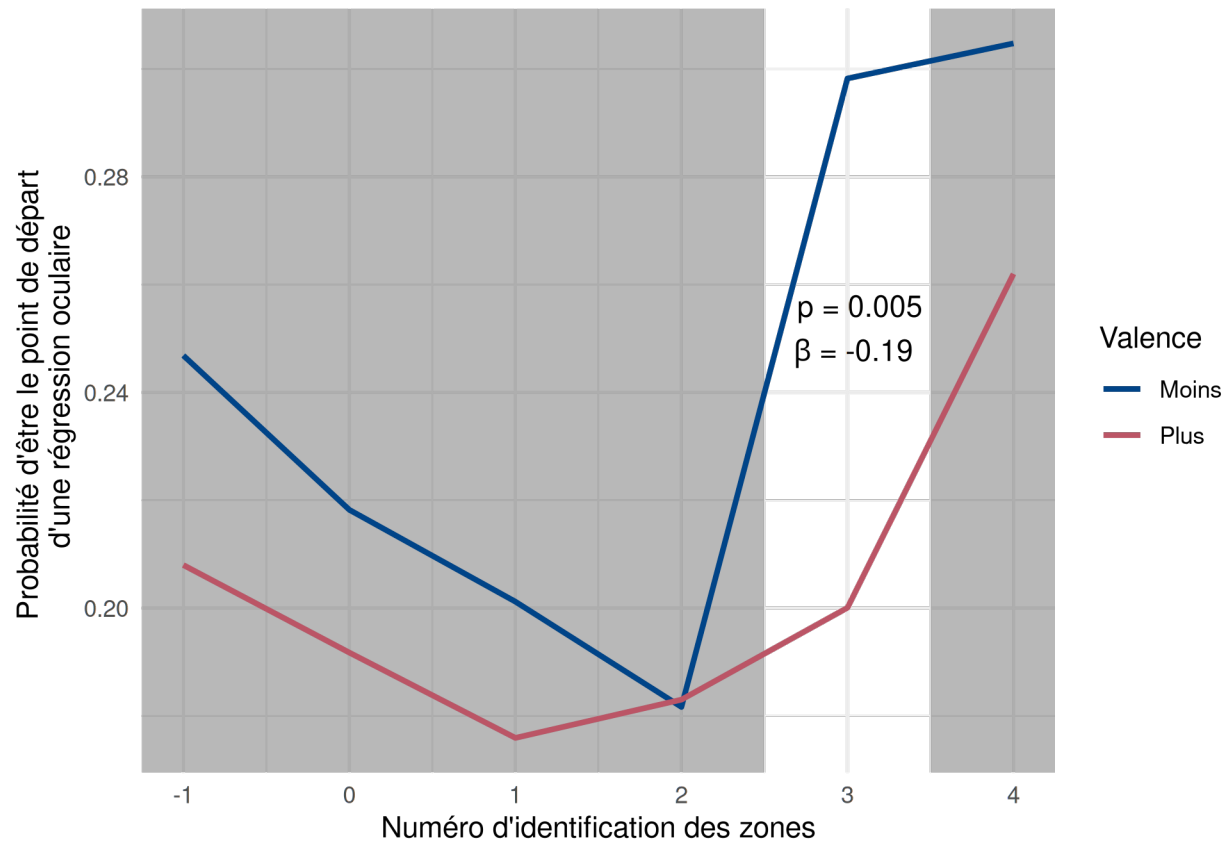


Figure 3.10 Moyenne de la probabilité d'être le point de départ d'une régression oculaire pour chaque zone, un fond blanc signifiant un effet significatif de la **Valence**

Aucun effet significatif du Contexte ne fut trouvé en séparant par zones. De plus, aucun effet de la Valence ou du Contexte ne fut trouvé sur en cherchant sur chacun des mots de la zone critique.

3.4.2 Probabilité de régression entrante

La probabilité de régression entrante est un ratio du nombre de régressions finissant sur chaque mot, divisé par le nombre total de régressions oculaires pour le visionnement de l'item entier. Cette mesure est l'autre côté de la probabilité de régressions sortante, et permet de voir quel mot se voit être la cible privilégiée de régressions oculaires, et si la **Valence** et le **Contexte** influence ce phénomène.

Puisqu'une quantité substantielle de mots ne sont pas la cible de régressions oculaires, une grande quantité de 0 peuplait cette variable, résultant en une distribution bimodale. Pour ne pas avoir à concocter un autre

modèle qui survit à cette distribution, nous avons opté pour le retrait des mots pour lesquels la probabilité de recevoir une fixation était 0. Ainsi, au lieu d'effectuer les analyses sur 5937 mots, elle a été faite sur 2597 mots. En insérant **probRegEnt** dans nos modèles, la distribution des résidus n'était plus bimodale, mais tout de même anormale. Nous avons donc effectué une transformation $\log()$, changeant le kurtosis des résidus du modèle de 15.28585 à 3.955388 et son asymétrie de 2.70051 à 1.063931.

En 3.11 on peut voir que la probabilité des régressions entrantes semble plus haute pour les phrases avec *mais moins*.

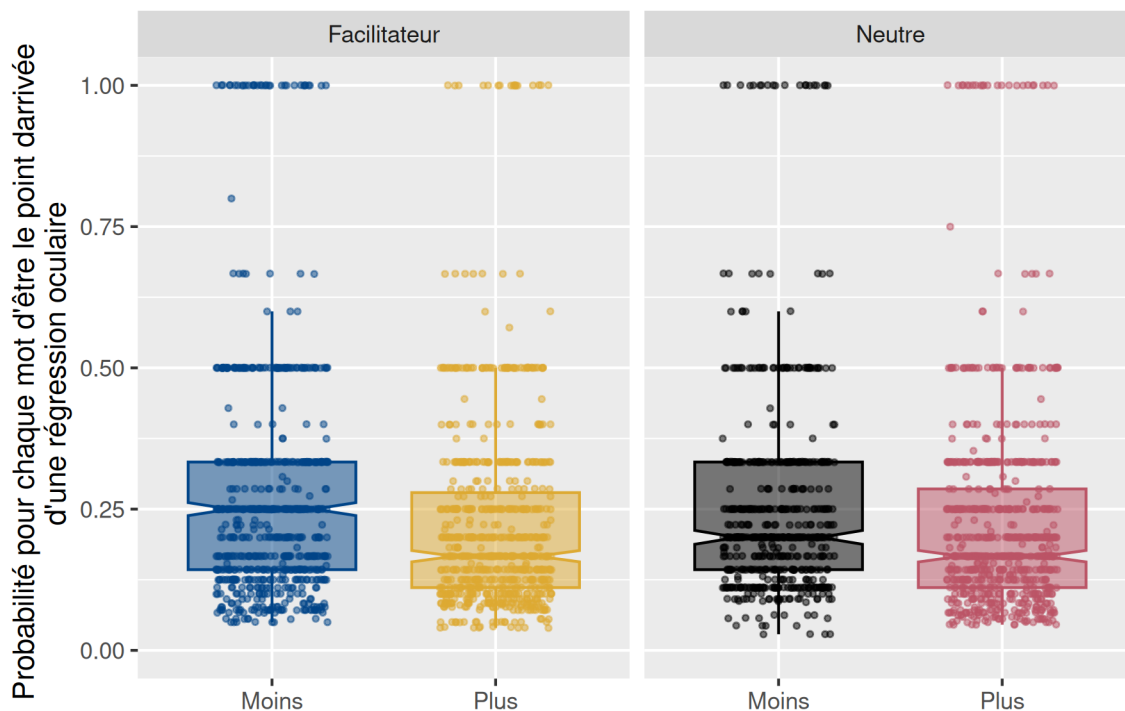


Figure 3.11 Probabilité d'être le point d'arrivée d'une régression oculaire pour chaque mot sur l'ensemble de l'échantillon, séparé selon les variables **Valence** et **Contexte**

La différence perçue au niveau de la Valence en observant les boîtes noire et bleu n'était pas une illusion. Cependant, si l'on se fie à notre modèle en contrôlant pour la quantité totale de régression dans la phrase, l'on y trouve aucun effet significatif de la Valence ou du Contexte, que ce soit global, au niveau des zones, ou au niveau des mots de la zone critique. Les résultats globaux se trouvent ci-dessous en 3.8, le reste se trouve en annexe.

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-0.78	0.05	-1.17 – -0.39	-0.03 – 0.13	<0.001
Valence [Plus]	-0.02	-0.03	-0.06 – 0.01	-0.09 – 0.02	0.195
Contexte [Neutre]	-0.02	-0.03	-0.05 – 0.02	-0.08 – 0.02	0.309
ScolCal	0.01	0.01	-0.05 – 0.07	-0.05 – 0.08	0.711
Empan	-0.05	-0.03	-0.15 – 0.05	-0.09 – 0.03	0.352
sFreqLex	0.01	0.01	-0.01 – 0.03	-0.02 – 0.04	0.347
nbCarac	0.02	0.06	0.01 – 0.02	0.03 – 0.09	<0.001
nbRegPhrase	-0.10	-0.74	-0.10 – -0.09	-0.77 – -0.71	<0.001
Random Effects					
σ^2	0.18				
τ_{00} Participant	0.01				
τ_{00} MediaType	0.00				
ICC	0.07				
$N_{\text{Participant}}$	28				
$N_{\text{MediaType}}$	19				
Observations	2597				
Marginal R^2 / Conditional R^2	0.561 / 0.594				

Table 3.8 Résultats du modèle pour le logarithme de la probabilité d'être le point d'arrivée d'une régression oculaire

CONCLUSION

Simplement, la réponse à notre première question de recherche est : OUI, il y a une différence de traitement entre les comparatives adversatives avec *mais plus* et celles avec *mais moins*, indiquée dans nos résultats par de multiples effets de la **Valence** à plusieurs niveaux d'analyse, corroborant les résultats d'auto-présentation segmentée de l'expérience de (Winterstein *et al.*, 2014). Le fait que la durée totale de fixation et que la probabilité de régression sur la phrase complète sont significativement plus élevées pour les phrases avec *mais plus* indique que celles-ci sont globalement plus difficiles à lire.

La réponse à notre deuxième question de recherche est NON : Le contexte "neutre" ou "facilitateur" n'affecte pas la lecture de la phrase, puisqu'aucun effet significatif du **Contexte** n'a été détecté pour aucune mesure, sur aucun de nos multiples niveaux d'analyse. Ainsi, des théories supposant un effet du contexte, comme **RT**, sont incompatibles avec nos résultats.

Le reste de ce chapitre nous permet d'élaborer, dans la section 3.5, et d'aborder les différents outils théoriques présentés dans le chapitre 1 à la lumière des résultats exposés dans le chapitre 3. Finalement, la section 3.6 contient une énonciation des limites rencontrées au courant de cette étude ainsi que diverses pistes d'approfondissement pour des recherches subséquentes.

3.5 Description de la différence entre *mais moins* et *mais plus*

Commençons par visualiser, à l'aide de la figure en 3.12 l'endroit où les effets significatifs de la **Valence** liées à la durée, c'est-à-dire à la décision de **Quand** exécuter une saccade, surgissent.

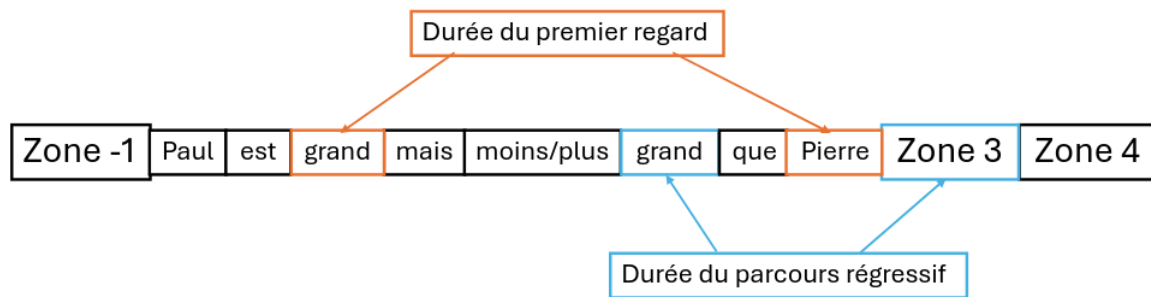


Figure 3.12 Lieu des effets de durées sur un exemple de phrase expérimentale

L'effet de la **Valence** survenant le plus tard dans la phrase en est un de la durée du parcours régressif au niveau de la zone 3. Si cet effet était dû à la difficulté d'un mot compris dans cette zone, on pourrait s'attendre à un effet de débordement (*spillover effect*) au niveau de la zone 4⁵. Plutôt, puisque la durée du premier regard est significativement plus élevée dans les phrases avec *mais plus* au niveau du mot à l'index 4 (prenant *mais* comme le point 0), qui représente le dernier mot de la zone 2, je suspecte que l'effet observé sur la zone 3 est en fait un effet de débordement suite à la complexité rencontrée dans les zones précédentes.

Si cette interprétation des effets sur la zone 3 est la bonne, cela signifie qu'aucun effet significatif de la **Valence** n'est provoqué à l'extérieur de la zone critique, dont un exemple est souligné dans l'exemple d'item ci-dessous en (6).

- (6) Au niveau de la carrure, Paul est grand, mais *moins/plus* grand que Pierre et il n'est pas facile de prendre une décision.

Cette observation est cohérente avec les explications se servant d'une échelle de comparaison comme celle en 1.1.3 : seulement et toute la partie soulignée en (6) est nécessaire pour la génération de l'échelle. Par exemple, un mot avant *Pierre*, il est encore possible pour l'échelle de comparaison de ne pas traiter de la

5. Le fait que le nombre de mots est plus variable dans les zone 3 et 4 jette un doute sur cette affirmation. Une solution est proposée en 3.6.

dimension évoquée par l'adjectif gradable (*taille*, dans l'exemple qui nous concerne), comme on peut voir en (7-b).

- (7) a. Paul est grand, mais *plus/moins* grand que Pierre...
- b. Paul est grand, mais *plus/moins* grand que rocambolesque...

En effet, l'échelle de comparaison invoquée par la phrase en (7-b) ne traite plus de la dimension *taille*, et n'a donc plus les propriétés sémantiques décrites plus tôt en 1.1.3 nous permettant de générer des prédictions. Or, c'est au mot *Pierre*, en (7-a), que l'échelle de comparaison se voit confirmer la dimension *Taille*, et c'est aussi sur le mot *Pierre* que la toute dernière différence de traitement entraînée par la **Valence** se situe.

Notamment, la division par zones permet de constater qu'il n'y a pas de différence causée par la **Valence** à la fin de la phrase, où se produit l'effet de fermeture, soit une hausse attendue du temps de lecture. Ce résultat va à l'encontre d'une théorie qui expliquerait la différence entre *mais moins* et *mais plus* en mentionnant l'intégration du contexte. En d'autres mots, c'est seulement autour des mots en (7-a) que la différence est observée, alors ces mots seuls devraient être nécessaires pour l'explication sémantique du phénomène.

La préférence pour une explication sémantique de nature plus lexicaliste est de surcroît appuyée par la vitesse des effets observés sur la **Valence**. Il est rassurant que ces derniers ne soient pas trop rapides, comme l'indique l'absence d'effet significatif de la **Valence** sur la durée de la première fixation. Une telle significativité aurait indiqué une différence notable au niveau du décodage graphique ; les phrases avec *mais plus* seraient plus difficiles à traiter dû à la différence de nombre de lettres et de fréquences lexicales entre *moins* et *plus*. L'absence de cet effet démontre l'efficacité des facteurs contrôles **sfreqlex** (la fréquence lexicale mise à l'échelle) et **nbCarac** (le nombre de caractères dans le mot).

Il s'agit tout de même d'effets rapides qui se perçoivent au niveau de la durée du premier regard. Notamment, l'effet significatif de la **Valence** sur la durée du premier regard au niveau de la première itération de l'adjectif gradable, qui se situe avant *plus/moins*, semble indiquer que la **Valence** exerce son influence en vision parafovéale. Un tel effet serait compatible avec une théorie nécessitant seulement les propriétés des items lexicaux, sans intégration contextuelle, ainsi qu'avec un modèle de la lecture permettant à des aspects

sémantiques d'être perçus en vision parafovéale ⁶.

Cependant, de par la nature de la mesure prescrite, un tel effet n'est pas certain. La durée du premier regard commence à partir du moment où le mot est fixé pour la première fois, mais n'exclue pas la possibilité d'avoir passé par-dessus un mot, avoir fixé la zone que l'on suspecte déclencher l'effet de **Valence** (*plus/moins*), puis être revenu fixer la première itération de l'adjectif gradable pour une première fois. Pour s'assurer que ce n'est pas le cas, il serait important d'analyser les mesures sur les mots de manière plus chronologique. Une éventuelle solution serait de considérer comme facteur contrôle l'index du mot le plus à droite ayant déjà été fixé au moment de la prescription de la durée du premier regard.

S'il y a un effet parafovéal de la **Valence**, il pourrait être dû à la perception de *plus* ou *moins* ou par celle de *plus grand* ou *moins grand*. La deuxième option a l'avantage d'être congruente avec notre explication sémantique du phénomène. En effet, *plus grand* partage presque tout son espace sémantique avec *grand*, ce qui n'est pas le cas de *plus*, qui n'a pas encore de dimension scalaire. De plus, avant d'avoir repéré *grand*, il est encore possible que la phrase traite d'une échelle de comparaison ne permettant pas la prédiction d'un effet de la **Valence** de manière similaire à (7-b), montré ci-dessous en (8).

(8) Paul est grand, mais *plus/moins* cool que Pierre...

Le champ perceptuel requis pour être affecté par *plus* ou *moins* tourne autour de 9 lettres vers la droite (dépendant du point de fixation), ce qui est réaliste pour la majorité des lecteurs. Pour être affecté par *grand*, il faut un champ perceptuel asymétrique s'étendant environ à 14 caractères vers la droite, ce qui indique de bonnes capacités de lectures. Ainsi, il serait possible de percevoir des différences interindividuelles face à l'effet parafovéal de la **Valence** selon les capacités de lectures des individus. Dans le cas de notre protocole expérimental, ces capacités sont inférées à l'aide des variables **Empan** (mémoire de travail) et **ScolCal** (calibre scolaire), qui ne montrent aucun effet significatif sur la durée du premier regard au mot à l'index -1. Cela est peut-être indicatif que *moins* ou *plus* sont nécessaires pour provoquer cet effet, sans nécessité de l'adjectif qui suit. Cependant, pour en être certain, il faudrait vérifier avec un protocole de manipulation du

6. Il y a débat dans la littérature : un côté disant que ce n'est pas possible (Inhoff et Rayner, 1980), l'autre disant que ce l'est (Yan et al., 2009). Le potentiel effet parafovéal découvert dans ma recherche n'est qu'ancillaire à mon objet d'étude, et nécessite plus ample investigation avant d'appuyer un camp parmi ce débat théorique.

champ perceptuel, comme le *boundary paradigm* de (Rayner, 1975).

3.6 Limites

La limite la plus notable liée à mon expérience est sans doute l'erreur effectuée dans l'instanciation des items au niveau de la croix de fixation. En me fiant de trop près aux tutoriels du logiciel et en ne changeant pas les paramètres par défaut, la croix de fixation apparaissait, pour les 27 premiers participants, directement au milieu et ainsi sur plusieurs mots et zones clés, ce qui m'a encouragé à retirer la moitié de mes participants de mes analyses statistiques. Une analyse comparative des deux ensembles de données (un avec la croix à gauche de la phrase, l'autre dans le milieu) serait intéressante pour examiner de potentiels effets d'amorçage, mais puisque l'expérience n'a pas été conçue à cet effet, les découvertes encourues auraient des capacités prédictives limitées.

Durant la passation de l'expérience, il m'arrivait de dire des phrases semblables à : "Assure-toi d'être le plus confortable possible." Il est possible que cette mention de *plus* sans contre-balance avec *moins* aie contribué à un effet d'amorçage ou d'habituation biaisant les résultats.

Au cours de la création de mon jeu de données, plusieurs informations ont été attachées à chaque mot, comme le nombre de caractères, ou la fréquence lexicale du mot. Avec le recul, plusieurs autres informations auraient été utiles. Par exemple, une indication de la chronologie de chaque fixation aurait permis d'élucider certaines questions, notamment la présence d'effet parafovéal de la **Valence** ainsi que les lieux d'atterrissages des régressions oculaires après la fixation de différents mots. Aussi, des informations sur les catégories grammaticales et sémantiques de chaque mot auraient permis de découper les zones d'intérêts non pas par positions dans la phrase, mais bien selon des critères linguistiques. Dans l'entreprise d'une telle réanalyse, je suggère de commencer avec un logiciel permettant un accès moins restreint aux données brutes que celui permis par TobiiProLab, de façon à générer le meilleur fichier de données possible dès le départ plutôt que de devoir fusionner divers jeux de données moins bien adaptés.

Un contrôle plus sophistiqué de la fréquence lexicale, spécifiquement du mot *plus* ainsi que des cooccurrences de *mais plus* et *mais moins* seraient essentielles. En effet, le contrôle s'est fait à l'aide de Lexique 3, dans lequel l'entrée correspondant à la graphie *plus* ne contient qu'un seul mot. Or la distinction entre le mot *plus* (+) et le mot *plus* (adverbe de négation, parfois prononcé *pu* en français québécois) permettrait de vérifier si les locutions *mais plus* et *mais moins* sont réellement de fréquences similaires, et à quel point.

L'instanciation de l'expérience dans Tobii Pro Lab (en 2022, des mises à jour à cet effet ont peut-être eu lieu depuis) a rendu trop difficile l'exécution d'un carré latin, qui a pour but d'assurer qu'aucun participant n'a été exposé aux mêmes items dans les mêmes conditions que les autres participants. La solution choisie, soit de faire quatre versions de l'expérience, a entraîné des problèmes vis-à-vis des noms des fichiers et du coût computationnel de l'extraction des données.

Au moment d'extraire les fréquences lexicales pour chaque mot n'étant pas un *stop-word*, j'ai pris la fréquence de l'unique adverbe *plus* présent dans Lexique3. Cependant, une ambiguïté réside en ce mot lorsque considéré seulement dans sa forme graphique entre *plus* (+) et l'adverbe négatif *plus* (pouvant se prononcer *pu* en français québécois). Avec l'accès au corpus FranText sur lequel Lexique3 est construit, il serait possible d'utiliser Spacy (Honnibal et Montani, 2017), qui avec un trait de "polarité" permet de distinguer les deux formes de *plus*, ce qui permettrait de calculer la fréquence lexicale du bon mot, au lieu de celle d'un amalgame homographique.

L'idée même de l'expérience de (Winterstein *et al.*, 2014) peut être retracée à un exemple dans un article de (Winterstein, 2012), traduit ci-dessous.

- (9) a. *Lemmy est grand, mais il est plus grand que son frère.
b. Lemmy est grand, et il est plus grand que son frère.
c. Lemmy est grand, il est plus grand que son frère.

Puisque (9-a) a un caractère dégradé, contrairement à (9-b) et (9-c), il est indéniable que le problème est causé par *mais*, ce qui démarre le long chemin sur lequel ce mémoire trotte. Vu l'aspect fondateur de ces exemples vis-à-vis de mon projet, il serait pertinent dans une étude subséquente de tester, plutôt qu'assumer, l'acceptabilité de comparatives asyndétiques comme en (9-c) ainsi que celles jointes avec *et*, comme en (9-b).

ANNEXE A
ITEMS EXPÉRIMENTAUX POUR LE TEST DE MÉMOIRE DE TRAVAIL

Série	Phrase	ITEM	CONCRET	Frequence	Pied_Mot	PIED_PH	MOTS_PH
1,00	Les rares secondes passées avec lui me donnent un aperçu très précis du bonheur.	bonheur	0,00	156,35	2,00	24,00	14,00
1,00	Le dresseur ira demander à son employé de lui choisir le meilleur cheval.	cheval	1,00	110,27	2,00	22,00	13,00
1,00	Les mauvaises herbes sont nuisibles et les jardiniers adroits les détruisent sans pitié.	pitié	0,00	57,91	2,00	24,00	13,00
1,00	Mon fils me demandera la permission avant de percer le lobe de son oreille.	oreille	1,00	103,45	2,00	23,00	14,00
1,00	Le dernier élève sortant de la salle de classe ferme les rideaux et la fenêtre.	fenêtre	1,00	199,39	2,00	24,00	15,00
1,00	Son ami et collègue ne pouvait endurer une minute de plus son humeur.	humeur	0,00	52,57	2,00	23,00	13,00
1,00	L'animateur lui mentionnera plusieurs noms et elle devra les garder à l'esprit.	esprit	0,00	182,84	2,00	23,00	14,00
1,00	Sa gardienne refuse de lui appliquer sa crème et de lui brosser les cheveux.	cheveux	1,00	263,18	2,00	22,00	14,00
1,00	Elle déposera des miettes de pain sur le gazon afin d'attirer un oiseau.	oiseau	1,00	47,97	2,00	23,00	14,00

1,00	Nous avons l'impression que l'examen de demain nous demandera un énorme effort.	effort	0,00	98,18	2,00	23,00	14,00
1,00	Le brave soldat pensait souvent au fait qu'il habitait loin de sa petite famille.	famille	0,00	241,69	2,00	23,00	15,00
1,00	Il commençait à ressentir pour la riche étrangère une profonde et furieuse passion.	passion	0,00	68,72	2,00	23,00	13,00
1,00	La pente du sentier était si raide qu'elle faisait surchauffer le moteur du camion.	camion	1,00	30,27	2,00	23,00	15,00
1,00	La secrétaire aux pantalons colorés vole des feuilles et les cache dans son bureau.	bureau	1,00	130,07	2,00	24,00	14,00
1,00	Le parent moderne élève son garçon comme sa fille dans une grande douceur.	douceur	0,00	66,08	2,00	23,00	13,00
1,00	Les spécialistes de son domaine se doivent bien de saluer son véritable génie.	génie	0,00	47,43	2,00	24,00	13,00
1,00	Les sportifs mettent beaucoup d'énergie pour arriver à remporter une seule victoire.	victoire	0,00	57,23	2,00	24,00	13,00
1,00	La voisine aime se baigner dans sa piscine creusée et se faire dorser au soleil.	soleil	1,00	328,78	2,00	24,00	15,00
1,00	La demoiselle avec le foulard et la tuque prenait son temps pour enlever son manteau.	manteau	1,00	58,99	2,00	24,00	15,00

1,00	Le verrier nettoie son sable et fait chauffer son feu pour se fabriquer une bouteille.	bouteille	1,00	70,41	2,00	22,00	15,00
2,00	Nous achetons ce produit de nettoyage parce qu'il résiste bien à la chaleur.	chaleur	0,00	112,23	2,00	23,00	14,00
2,00	Le terrible accident de voiture fera l'objet d'un reportage dans le journal.	journal	1,00	124,32	2,00	23,00	14,00
2,00	Plonger dans l'eau froide pour aller sauver un nageur demande beaucoup de courage.	courage	0,00	69,80	2,00	22,00	14,00
2,00	Il faut orienter le feuillage des plantes vertes vers une forte source de lumière.	lumière	0,00	238,65	2,00	24,00	14,00
2,00	Les vêtements de ma tante ne sont pas propres et elle doit refaire sa valise.	valise	1,00	47,43	2,00	23,00	15,00
2,00	Mon jeune cousin se cherche un endroit agréable et paisible pour finir son travail.	travail	0,00	223,99	2,00	23,00	14,00
2,00	La compétition sportive aura lieu dans une semaine et elle perturbe son sommeil.	sommeil	0,00	112,03	2,00	24,00	13,00
2,00	Il faut toujours lui demander plusieurs fois de terminer le contenu de son assiette.	assiette	1,00	36,28	2,00	23,00	14,00
2,00	Je ne sais pas pendant combien de jours nous devons marcher pour gravir cette montagne.	montagne	1,00	49,80	2,00	22,00	15,00

2,00	La situation complexe décrite par cet auteur à succès est vraiment sans espoir.	espoir	0,00	90,74	2,00	24,00	13,00
2,00	Mon ancien copain imaginait souvent que nous vivions une très belle histoire d'amour.	amour	0,00	373,58	2,00	24,00	14,00
2,00	Le vent de la tempête de neige est tellement puissant qu'il soulève notre casquette.	casquette	1,00	34,39	2,00	23,00	15,00
2,00	Il devra remettre sa rencontre avec les diplomates s'il manque le prochain avion.	avion	1,00	46,82	2,00	24,00	14,00
2,00	Nous pouvons nous rendre dans ce conduit aux odeurs désagréables par une courte échelle	échelle	1,00	28,04	2,00	24,00	14,00
2,00	Ce peintre est célèbre dans le monde entier pour son utilisation de la couleur.	couleur	0,00	118,65	2,00	22,00	14,00
2,00	Cuisiner à l'huile d'olive vierge est un geste bénéfique pour la santé.	santé	0,00	52,43	2,00	22,00	14,00
2,00	Des obstacles nombreux et répétés risquent de le faire basculer dans la folie.	folie	0,00	52,43	2,00	23,00	13,00
2,00	La femme de ménage repasse deux robes roses et les range dans l'armoire.	armoire	1,00	38,58	2,00	22,00	14,00
2,00	Le navigateur expérimenté fait toujours très attention à l'état de son bateau.	bateau	1,00	61,22	2,00	24,00	13,00

2,00	Ton oncle rêve de s'acheter une motocyclette et de se construire une maison.	maison	1,00	461,55	2,00	24,00	14,00
3,00	Il est clair que sauter en parachute lui donne une vive sensation de plaisir.	plaisir	0,00	208,78	2,00	24,00	14,00
3,00	Ce dessinateur bien connu se creuse la tête longuement avant de choisir un crayon.	crayon	1,00	25,47	2,00	25,00	14,00
3,00	Je trouve que ce politicien est détestable parce qu'il change toujours d'idée.	idée	0,00	241,08	2,00	23,00	14,00
3,00	Mon frère le plus âgé sera bien content que ma mère repasse sa chemise.	chemise	1,00	74,59	2,00	22,00	14,00
3,00	Il pique le morceau de viande avec une fourchette et le coupe avec un couteau.	couteau	1,00	44,26	2,00	22,00	15,00
3,00	Un enfant trop fatigué éprouve de la difficulté à se lever tôt le matin.	matin	0,00	376,89	2,00	24,00	14,00
3,00	Les proches des victimes de ce dangereux criminel poursuivront le coupable en justice.	justice	0,00	46,22	2,00	24,00	13,00
3,00	Je me demande comment ils réussiront à lui faire un lavage de cerveau.	cerveau	1,00	28,92	2,00	22,00	13,00
3,00	Il parvient à sécher son corps en entier seulement avec le coin de sa serviette.	serviette	1,00	26,62	2,00	22,00	15,00

3,00	Cet homme ayant confiance en lui ne voulait quand même pas suivre mon judicieux conseil	conseil	0,00	58,18	2,00	23,00	15,00
3,00	Plusieurs personnes sont aptes à prendre les bonnes décisions pour le bien du pays.	pays	0,00	241,55	2,00	23,00	14,00
3,00	Le docteur dira à ce malade qu'il perdra d'ici peu complètement la mémoire.	mémoire	0,00	105,74	2,00	23,00	15,00
3,00	La compagne du pêcheur offre des rabais importants afin de vendre son poisson.	poisson	1,00	30,14	2,00	23,00	13,00
3,00	Tu dépenseras tes économies pour t'acheter un ordinateur portable et des lunettes.	lunettes	1,00	67,84	2,00	25,00	13,00
3,00	Un voyageur sachant toujours se tirer d'affaire n'a rien à redouter du destin.	destin	0,00	62,77	2,00	23,00	15,00
3,00	Les médicaments prescrits par le médecin ne permettent pas de soulager sa douleur.	douleur	0,00	77,84	2,00	24,00	13,00
3,00	Cette vendeuse considère que les bottes de pluie sont à la mode du moment.	moment	0,00	611,62	2,00	23,00	14,00
3,00	Il est assez fréquent pour un policier d'assister à la naissance d'un bébé.	bébé	1,00	36,22	2,00	22,00	15,00
3,00	L'agent de police doit se calmer et se concentrer avant de tourner la poignée.	poignée	1,00	36,35	2,00	23,00	15,00

3,00	Le froid intense nous obligeait à revêtir encore une fois notre immense chapeau.	chapeau	1,00	72,91	2,00	24,00	13,00
------	--	---------	------	-------	------	-------	-------

ANNEXE B
ITEMS EXPÉRIMENTAUX DE L'EXPÉRIENCE PRINCIPALE

phrase	nbchar	nbmots	catégorie
On cherche un cascadeur pour doubler Pierre, un acteur, dans des scènes d'action. Il faut que le cascadeur ait exactement la même taille que Pierre pour qu'on ne remarque pas la doublure à l'écran. C'est difficile parce que Pierre est de grande taille.	252	47	Cont_Facil
Lorsqu'un pays est en forte croissance, les exportations doivent augmenter rapidement et exactement au même rythme que le PIB, pour éviter un déséquilibre économique.	166	25	Cont_Facil
Une enseignante a choisi un problème de maths assez difficile pour son examen de jeudi. Vendredi, elle a un autre examen avec un groupe de même niveau et veut trouver un problème de difficulté rigoureusement équivalente.	220	36	Cont_Facil
Une mécanicienne doit ajouter de l'huile dans un moteur de Formule 1. Elle a besoin d'une huile qui ait les exactes mêmes propriétés de viscosité que l'huile d'origine du moteur, connue pour sa forte viscosité.	210	39	Cont_Facil
Le système d'arrosage manuel d'un jardin public a été endommagé par la foudre. Pour le remplacer on a installé un système automatique récent. Il est essentiel que le nouveau système conserve absolument la même fréquence d'arrosage que le précédent, pour éviter d'abîmer certaines plantes délicates.	298	49	Cont_Facil
Michel veut refaire une partie de sa grande demeure et a fait réaliser plusieurs devis par différents architectes. Michel souhaite conserver ses plafonds, tous assez hauts, exactement à la même hauteur pour préserver le style des lieux.	236	37	Cont_Facil
La mairie oblige Anabelle à faire creuser une tranchée pour la mise en place d'une évacuation des eaux usées. La tranchée doit être identique à celle de la maison voisine d'Annabelle et être bien large pour permettre une bonne évacuation des eaux.	247	44	Cont_Facil

Elodie a hérité d'une voiture en panne de son parrain et veut réparer ce véhicule en trouvant une épave du même modèle que cette voiture pour récupérer des pièces détachées. L'épave doit avoir exactement le même âge que cette voiture qui était récente pour s'assurer que les pièces seront compatibles.	301	53	Cont_Facil
Simone est en train de restaurer un tableau accroché au mur de sa chambre. Elle a besoin de fabriquer un rouge sombre qui ait exactement le même degré de lumière que l'original.	177	33	Cont_Facil
Alice essaie de reproduire le gâteau que faisait sa grand-mère pour chacun de ses anniversaires. Elle veut notamment reproduire avec fidélité la texture moëlleuse de ce gâteau.	177	28	Cont_Facil
Maurice veut accéder au coffre-fort de son frère jumeau, Ivan, qui est protégé par un système de reconnaissance faciale. Il espère ressembler assez à son frère pour pouvoir déverrouiller la serrure.	198	32	Cont_Facil
Stéphanie, une criminelle, veut copier et vendre une reproduction d'une oeuvre d'art d'une valeur inestimable. Il ne s'agit pas que de reproduire le dessin, l'oeuvre doit montrer autant de traces d'usures que la version originale, pour tromper les experts qui évalueront son authenticité.	288	49	Cont_Facil
Jeanne Doe est amnésique et ne sais plus qui elle est. Les enquêteurs pensent qu'elle est peut-être la soeur jumelle de Marc Leclerc, un vieil homme portée disparu il y a 10 ans de cela. Pour s'en assurer, ils font passer à Jeanne une batterie de tests pour découvrir son âge.	276	54	Cont_Facil
Valérie est agente et cherche un boxeur pour s'affronter à son client, Maxime, qui est reconnu comme étant très dangeureux. Elle sait que si les deux combattants sont tout aussi dangereux, l'évènement sera intéressant et attirera beaucoup de téléspectateur.	259	41	Cont_Facil
Yvette et Diane veulent échanger leurs maisons respectives. Selon les conseils de leur agents immobiliers, si leurs maisons n'ont pas la même valeur après une évaluation du marché, il devra y avoir un transfert d'argent en plus de l'échange d'habitations pour que les deux parties en ressortent gagnants.	304	52	Cont_Facil

Une entreprise de démolition veut faire s'affaisser un bâtiment solidement érigé sur deux murs de fondations principaux, en utilisant de la dynamite. Avant l'explosion, les deux murs doivent avoir exactement la même solidité, sans quoi le bâtiment ne tombera pas droit et risque d'endommager les infrastructures environnantes.	326	50	Cont_Facil
Dans un bar, Georges s'apprête à affronter Serge, qui est fort, dans un bras de fer. La patronne du bar chuchote à ses employés que si le match est ex aequo, elle donnera une consommation gratuite à tout le monde présent.	223	42	Cont_Facil
Dans un zoo, on veut connecter deux vivariums, un petit et un gros, à l'aide d'un tunnel. L'air présent dans le gros vivarium y est humide, pour accommoder les serpents tropicaux qui s'y trouvent. Pour ne pas dérégler leur écosystème, l'air contenu dans les deux vivariums doivent avoir exactement le même taux d'humidité.	322	59	Cont_Facil
Julien, qui a 16 ans et la peau très pâle, se fait faire une fausse carte d'identité pour pouvoir entrer dans les bars. La photo présente sur la fausse carte doit lui ressembler le plus possible si il veut que sa stratégie fonctionne.	234	44	Cont_Facil
Elon, qui est très intelligent, veut exécuter le premier transfert de conscience en changeant de corps avec un autre individu, mais il ne sait pas encore qui. L'autre personne devra être tout aussi intelligente que lui s'il veut que le transfert de corps se passe sans problème.	278	49	Cont_Facil
On cherche un cascadeur pour doubler Pierre, un acteur, dans des scènes d'action. Pour donner l'impression que le personnage de Pierre est très agile, il faut que le cascadeur soit plus agile que Pierre. C'est difficile parce que Pierre est plutôt agile.	254	45	Cont_Neutre
Lorsqu'un pays est en forte croissance, les exportations et le PIB doivent être en augmentation constante, pour éviter un déséquilibre économique.	146	22	Cont_Neutre
Une enseignante a choisi un problème de maths assez difficile pour son examen de jeudi. Vendredi, elle a un autre examen avec un groupe contenant beaucoup plus d'étudiants et veut trouver un problème plus court, pour que la tâche de correction ne soit pas trop longue.	268	47	Cont_Neutre

Un mécanicien doit ajouter de l'huile dans un moteur de Formule 1. Suite aux nouvelles mesures de sécurité, il ne peut utiliser l'huile d'origine du moteur, qui est reconnue pour sa tendance à prendre en feu.	208	39	Cont_Neutre
Le système d'arrosage manuel d'un jardin public a été endommagé par la foudre. Pour le remplacer on a installé un système automatique récent. Il est essentiel que le nouveau système conserve absolument la même couleur que le précédent, pour éviter d'enlaidir l'aménagement paysager.	283	47	Cont_Neutre
Michel veut refaire une partie de sa grande demeure et a fait réaliser plusieurs devis par différents architectes. Michel souhaite reproduire la même forme pour les moulures, pour conserver l'apparence élégante des lieux.	221	34	Cont_Neutre
La mairie oblige Anabelle à faire creuser une tranchée pour la mise en place d'une évacuation des eaux usées. La tranchée ne doit pas empiéter sur la maison voisine d'Annabelle et être bien large pour permettre une bonne évacuation des eaux.	242	43	Cont_Neutre
Elodie a hérité d'une voiture en panne de son parrain et veut réparer ce véhicule en trouvant une épave du même modèle que cette voiture pour récupérer des pièces détachées. L'épave doit avoir exactement la même taille que cette voiture qui était énorme pour s'assurer que les pièces soient compatibles.	303	53	Cont_Neutre
Simone est en train de restaurer un tableau accroché au mur sa chambre. Elle en profite pour changer les teintes de de rouge par rapport à l'original, qui la rendaient mal à l'aise.	181	35	Cont_Neutre
Alice essaie de reproduire le gâteau que faisait sa grand-mère pour chacun de ses anniversaires. Elle veut explorer différentes textures dans le but de modifier la recette, et en faire une qui soit la sienne.	208	36	Cont_Neutre
Maurice veut accéder au coffre-fort de son frère jumeau Ivan, qui est protégé par un système de reconnaissance automatique de la rétine. Il espère avoir des yeux assez similaires à ceux de son frère pour pouvoir débloquent la serrure.	233	40	Cont_Neutre
Stéphanie, une artiste, veut s'inspirer d'un oeuvre très connu pour en faire une reproduction personnalisée. Après avoir laissé sa version dans le grenier humide pendant quelques mois, elle montre certaines traces d'usure.	223	35	Cont_Neutre

Jeanne Doe est amnésique et ne sais plus qui elle est. Les enquêteurs pensent qu'elle est la petite soeur de Marc Leclerc, un vieil homme portée disparu il y a 10 ans de cela. Pour s'en assurer, ils font passer à Jeanne une batterie de tests.	242	48	Cont_Neutre
Valérie est agente et cherche un boxeur pour s'affronter à son client, Maxime, qui est reconnu comme étant très dangeureux. Elle sait que si les adversaires sont ont le même poids, le combat sera autorisé et elle pourra en faire la publicité.	242	43	Cont_Neutre
Yvette et Diane veulent rénover leurs maisons respectives, pour en augmenter la valeur. Pour sauver de l'argent, ils engagent le même contracteur général, qui commence par regarder l'état et la valeur initiale de leurs habitations.	231	37	Cont_Neutre
Une entreprise de construction veux ajouter des étages à un immeuble solide-ment érigé sur deux murs de fondations principaux. Pour que le ne s'effondre pas, les deux murs doivent être suffisamment solides.	205	33	Cont_Neutre
Dans un bar, Georges s'apprête à affronter Serge, qui est fort, dans un bras de fer. La patronne du bar s'exclame à tous qu'elle affrontera elle-même le gagnant de ce match.	173	35	Cont_Neutre
Dans un zoo, on veut connecter un petit vivarium avec un gros, dans lesquels vivent plusieurs serpents tropicaux, qui ont besoin que l'air y soit humide. Pour ne pas dérègler leur écosystème, l'air contenu dans les deux vivariums doivent avoir exactement la même température	273	46	Cont_Neutre
Julien, qui a 16 ans et la peau très pâle, se fait faire une fausse carte d'identité pour pouvoir entrer dans les bars. Par erreur, la fausse carte prétend qu'il a 48 ans, alors qu'il a l'air d'en avoir 13.	206	45	Cont_Neutre
Elon, qui est très intelligent, veut exécuter le premier transfert de conscience en changeant de corps avec un autre individu, mais il ne sait pas encore qui. Le crâne de l'autre personne devra avoir la même forme que le sien s'il veut que le transfert de corps se passe sans problème.	285	53	Cont_Neutre
Au niveau de la carrure, Paul est grand mais moins grand que Pierre et il n'est pas facile de prendre une décision	114	23	Items_moins
Dans certains pays émergents, les exportations augmentent vite, mais moins vite que le PIB ce qui peut laisser certains experts perplexes.	138	21	Items_moins

Pour son examen de vendredi, l'enseignante a trouvé un problème difficile, mais moins difficile que celui qu'elle avait choisi pour son examen de jeudi.	152	26	Items_moins
Pour la formule 1 qu'elle entretient, la mécanicienne a trouvé une huile visqueuse, mais moins visqueuse que celle utilisée auparavant et qui avait été importée.	161	26	Items_moins
Avec le nouveau système automatique, l'arrosage des jardins publics est fréquent, mais moins fréquent que celui de l'ancien système que la foudre a détruit.	156	26	Items_moins
Pour les travaux du salon, un projet prévoit des plafonds hauts mais moins hauts que ceux dans l'antichambre situé au nord-est du salon.	136	25	Items_moins
Les ouvriers ont creusé une tranchée large, mais moins large que celle de la maison voisine qui a déjà été rénovée.	115	21	Items_moins
Un garage propose une carcasse de voiture récente, mais moins récente que celle qu'Élodie vient de récupérer de son parrain.	124	21	Items_moins
Sur sa palette, Simone a réussi à créer un rouge sombre, mais moins sombre que celui utilisé sur le tableau qu'elle a dans sa chambre.	134	26	Items_moins
Le gâteau d'Alice est moelleux mais moins moelleux que celui fait par sa grand-mère pour son anniversaire.	106	19	Items_moins
Le nez de Maurice est croche mais moins croche que celui de son frère Ivan depuis son accident.	95	18	Items_moins
La copie est endommagée, mais moins endommagée que la version originale, et cela est particulièrement visible dans les coins.	125	19	Items_moins
Selon les tests, Jeanne est vieille mais moins vieille que Marc même s'ils ont tous les deux les cheveux blancs.	112	21	Items_moins
Le boxeur Benjamin est dangereux, mais moins dangereux que Maxime et les téléspectateurs veulent un match serré.	112	17	Items_moins
La maison de Diane vaut cher mais moins cher que celle d'Yvette selon une évaluation rigoureuse du marché immobilier.	117	20	Items_moins
Le mur du côté de la rue est solide, mais moins solide que celui du côté de la cour arrière, près de la ruelle.	111	24	Items_moins
Avec ses gros bras, Georges est fort mais moins fort que Serge et la compétition s'annonce être féroce.	104	19	Items_moins

L'air dans le petit vivarium est humide, mais moins humide que l'air contenu dans le gros vivarium.	99	19	Items_moins
Au niveau du QI, Albert est intelligent, mais moins intelligent qu'Elon selon les tests attestés les plus reconnus.	115	19	Items_moins
Au niveau de la carrure, Paul est grand mais plus grand que Pierre et il n'est pas facile de prendre une décision	113	23	Items_plus
Dans certains pays émergents, les exportations augmentent vite, mais plus vite que le PIB ce qui peut laisser certains experts perplexes.	137	21	Items_plus
Pour son examen de vendredi, l'enseignante a trouvé un problème difficile, mais plus difficile que celui qu'elle avait choisi pour son examen de jeudi.	151	26	Items_plus
Pour la formule 1 qu'elle entretient, la mécanicienne a trouvé une huile visqueuse, mais plus visqueuse que celle utilisée auparavant et qui avait été importée.	160	26	Items_plus
Avec le nouveau système automatique, l'arrosage des jardins publics est fréquent, mais plus fréquent que celui de l'ancien système que la foudre a détruit.	155	26	Items_plus
Pour les travaux du salon, un projet prévoit des plafonds hauts mais plus hauts que ceux dans l'antichambre situé au nord-est du salon.	135	25	Items_plus
Les ouvriers ont creusé une tranchée large, mais plus large que celle de la maison voisine qui a déjà été rénovée.	114	21	Items_plus
Un garage propose une carcasse de voiture récente, mais plus récente que celle qu'Élodie vient de récupérer de son parrain.	123	21	Items_plus
Sur sa palette, Simone a réussi à créer un rouge sombre, mais plus sombre que celui utilisé sur le tableau qu'elle a dans sa chambre.	133	26	Items_plus
Le gâteau d'Alice est moelleux mais plus moelleux que celui fait par sa grand-mère pour son anniversaire.	105	19	Items_plus
Le nez de Maurice est croche mais plus croche que celui de son frère Ivan depuis son accident.	94	18	Items_plus
La copie est endommagée, mais plus endommagée que la version originale, et cela est particulièrement visible dans les coins.	124	19	Items_plus
Selon les tests, Jeanne est vieille mais plus vieille que Marc même s'ils ont tous les deux les cheveux blancs.	111	21	Items_plus

Le boxeur Benjamin est dangereux, mais plus dangereux que Maxime et les téléspectateurs veulent un match serré.	111	17	Items_plus
La maison de Diane vaut cher mais plus cher que celle d'Yvette selon une évaluation rigoureuse du marché immobilier.	116	20	Items_plus
Le mur du côté de la rue est solide, mais plus solide que celui du côté de la cour arrière, près de la ruelle.	110	24	Items_plus
Avec ses gros bras, Georges est fort mais plus fort que Serge et la compétition s'annonce être féroce.	103	19	Items_plus
L'air dans le petit vivarium est humide, mais plus humide que l'air contenu dans le gros vivarium.	98	19	Items_plus
Sur la fausse carte de Julien, sa peau est pâle mais plus pâle qu'en réalité car il sort très rarement au soleil.	113	23	Items_plus
Au niveau du QI, Albert est intelligent, mais plus intelligent qu'Elon selon les tests attestés les plus reconnus.	114	19	Items_plus
Va-t-on embaucher Paul comme doublure ?	38	7	Question_faci_item
Est-il probable que ces pays connaissent une bonne stabilité économique ?	72	11	Question_faci_item
Les examens de jeudi et vendredi vont-ils être équilibrés ?	58	10	Question_faci_item
Est-ce que la nouvelle huile convient au moteur ?	48	9	Question_faci_item
Le nouveau système est-il adapté aux plantes fragiles ?	54	9	Question_faci_item
Est-il probable que Michel demande une révision du projet ?	58	10	Question_faci_item
La mairie risque-t-elle d'objecter à la façon dont la tranchée a été creusée ?	77	16	Question_faci_item
Est-ce qu'Elodie va rechercher une autre épave pour réparer sa voiture ?	71	13	Question_faci_item
Verra-t-on une différence sur le tableau si Simone utilise le rouge qu'il a créé ?	81	17	Question_faci_item
Est-ce qu'Alice doit s'améliorer si elle veut reproduire parfaitement le gâteau ?	80	14	Question_faci_item
Maurice pourra-t-il accéder au coffre-fort de son frère jumeau ?	63	12	Question_faci_item
La copie de la peinture pourra-t-elle tromper les experts ?	58	11	Question_faci_item
Est-il possible que Jeanne Doe soit la soeur jumelle de Marc Leclerc ?	68	13	Question_faci_item
Est-ce que les téléspectateurs seront intéressés si Benjamin affronte Maxime ?	76	11	Question_faci_item
Est-ce que leur échange nécessitera un montant d'argent ?	56	10	Question_faci_item
Est-ce que le bâtiment s'effondrera en ligne droite ?	52	10	Question_faci_item

Est-ce que la patronne du bar paiera la tournée à tout le monde ?	64	14	Question_faci_item
Est-ce que l'écosystème des serpents sera dérèglé ?	50	9	Question_faci_item
Les portiers qui examinent sa fausse carte risquent-ils de percevoir une différence ?	84	13	Question_faci_item
Est-ce que Elon pourra changer de corps avec Albert sans problème ?	66	12	Question_faci_item
Est-ce que Paul est plus grand que Pierre ?	42	9	Question_neutre_item
Est-il probable que ces pays connaissent une bonne stabilité économique ?	72	11	Question_neutre_item
Est-ce que l'examen de vendredi sera plus difficile que celui de jeudi ?	71	14	Question_neutre_item
Est-ce que la nouvelle huile est visqueuse ?	43	8	Question_neutre_item
L'arrosage de l'ancien système est-il plus fréquent que celui du nouveau ?	73	14	Question_neutre_item
Est-ce que les murs de l'antichambre sont haut ?	47	10	Question_neutre_item
La tranchée creuser par les ouvriers est-elle moins large que celle de la maison voisine ?	89	16	Question_neutre_item
Est-ce que Élodie possède assez d'informations pour prendre sa décision ?	72	12	Question_neutre_item
Les tableaux seront-ils identiques si Simone utilise le rouge qu'elle a créé ?	77	14	Question_neutre_item
Est-ce qu'Alice a reproduit fidèlement la recette de sa grand-mère ?	67	13	Question_neutre_item
Le nez de Ivan est-il droit ?	28	7	Question_neutre_item
Est-ce que le grenier a endommagé la peinture ?	46	9	Question_neutre_item
Est-il possible que Jeanne Doe soit la soeur de Marc Leclerc ?	61	12	Question_neutre_item
Est-ce que Maxime est dangereux ?	33	6	Question_neutre_item
Est-ce que la maison d'Yvette est cher ?	39	9	Question_neutre_item
Est-il certain que le bâtiment sera assez solide ?	49	9	Question_neutre_item
Est-ce que la patronne du bar affrontera Serge ?	47	9	Question_neutre_item
Est-ce que l'air est sec dans le petit vivarium ?	48	11	Question_neutre_item
Est-ce que la couleur de sa peau est aussi sombre que sur la photo ?	67	15	Question_neutre_item
Est-ce que Elon est intelligent ?	32	6	Question_neutre_item
Yannick prépare son déménagement. Il fait ses boîtes et veut éviter d'en faire des trop lourdes. Il regarde avec quoi combler le carton qu'il est en train de remplir.	166	31	Cont_Fill

Marc-André est passionné d'aménagement intérieur. Après avoir changé de canapé, il cherche une table de salon qui soit un peu moins longue que le canapé pour gagner de l'espace.	177	32	Cont_Fill
En montant la côte pour rentrer chez elle, Camille a cassé sa chaîne de vélo. Elle est allée la faire changer dans un magasin spécialisé.	137	25	Cont_Fill
Une metteuse en scène fait un casting pour son prochain spectacle. Elle auditionne un danseur qui propose une chorégraphie sur une musique très calme.	151	24	Cont_Fill
Guillaume cherche un vêtement pour la mi-saison. Il veut quelque-chose qui cache ses petites jambes. La vendeuse lui propose un imperméable d'une nouvelle marque.	163	27	Cont_Fill
Le ministère a demandé qu'un de ses rapports d'audit soit plus court à lire en n'omettant que les détails triviaux.	117	24	Cont_Fill
Marion et Romain quittent Montréal pour trois mois. Marion a demandé à sa soeur Arianne de s'occuper de leurs plantes durant cette période. Romain n'est pas persuadé que ce soit une bonne idée parce qu'aucune des plantes d'Arianne n'a survécu ces dernières années.	264	48	Cont_Fill
Les parents de Sandra lui rendent visite ce week-end. Ils ont fait un long voyage car ils habitent à 500km de chez elle. Sandra a voulu aller les chercher à la gare, mais en chemin elle a rencontré son vieil ami François.	221	43	Cont_Fill
Madame Dupont est en train d'attendre le plombier chez elle. De façon inattendue une voisine sonne à la porte pour demander un œuf. Pendant que madame Dupont cherche l'œuf, le plombier arrive et la voisine le laisse rentrer.	224	40	Cont_Fill
A un feu rouge, un taxi rentre dans une autre voiture, conduite par une dame âgée. Par chance, personne n'est blessé et le chauffeur de taxi appelle immédiatement la police. Quand ils arrivent, ils demandent en premier à la dame d'expliquer ce qui s'est passé.	260	48	Cont_Fill
Un professeur de math est mécontent de ses étudiants. Non seulement ils n'ont pas fait leurs devoirs, mais ils passent leur temps à discuter en cours. Il demande à la professeur d'anglais si elle rencontre les mêmes problèmes. Apparemment les étudiants ne sont concentrés que lorsqu'elle est très stricte avec eux et qu'elle menace de leur rajouter du travail.	360	63	Cont_Fill

Louise et Claire vont magasiner ensemble. Claire cherche un cadeau pour son mari alors que Louise a besoin d'une nouvelle paire de chaussures. Pendant qu'elle essaie une paire de sandales bleues, Claire aperçoit une écharpe qui pourrait plaire à son mari. Comme elle n'est pas certaine de la couleur, elle demande conseil à Louise.	331	57	Cont_Fill
Une journaliste du Monde essaie d'obtenir un entretien avec le ministre des affaires étrangères islandais. Bien qu'il y ait passé beaucoup de temps, jusqu'ici elle a seulement pu parler à son assistant. Il a promis de le rappeler dès que le ministre aurait un moment de libre dans son agenda, mais la journaliste a des doutes à ce sujet.	337	62	Cont_Fill
Une agence touristique montréalaise teste une nouvelle offre : pour chaque quartier ils embauchent un guide qui y a grandi afin d'offrir l'expérience la plus authentique possible à leurs clients. Le premier groupe de touristes est interrogé après leur visite afin d'évaluer cette nouvelle stratégie. Certains se plaignent de la guide de Rosemont parce qu'elle parlait trop rapidement, alors que le guide du quartier latin est généralement apprécié.	447	71	Cont_Fill
Un libraire vient de recevoir une livraison de plusieurs ouvrages, dont certains très rares et chers. Il demande à son employée de trier les livres en fonction de leur valeur, et de protéger les plus rares dans une vitrine sécurisée. Après plusieurs heures de travail, elle hésite sur le classement d'un ouvrage qui paraît ancien.	330	56	Cont_Fill
Sarah et Robert visitent un zoo. En passant devant l'enclos des tatous, Sarah remarque qu'ils sont endormis alors que le panneau affirme que ce sont des animaux diurnes. Robert répond qu'elle se trompe probablement et que les tatous sont en fait des animaux nocturnes.	268	47	Cont_Fill
Gérard est avec sa petite fille Véronique. Ils font leurs courses dans un supermarché. Elle lui demande s'il peut lui acheter une friandise. Gérard pense qu'elle est mauvaise pour sa santé, et lui répond qu'il préférerait qu'elle mange une pomme.	246	44	Cont_Fill

Aujourd'hui c'est le printemps et il fait beau. Tous les couples se rendent au parc pour pique-niquer. Boris et Catherine ont ramené une bouteille de vin, une baguette de pain et un bout de fromage. Une fois sur place, ils réalisent qu'il est moisi et ils se demandent ce qu'ils vont bien pouvoir manger.	305	59	Cont_Fill
Un employé discute avec sa représentante syndicale de la dernière réunion avec le service des ressources humaines. Il s'intéresse notamment à une rumeur de plan de licenciement.	177	28	Cont_Fill
Amélie et Marie, deux collègues, discutent de Richard, un de leurs superviseurs. Amélie a entendu que Richard était régulièrement absent. Marie lui répond qu'il n'est pas faux que Richard va souvent au Brésil.	209	35	Cont_Fill
Bertrand interroge son avocate sur un procès en cours. Il veut notamment savoir comment a réagi la partie adverse.	114	19	Cont_Fill
Un chercheur interroge la directrice de son labo. Il veut savoir si les résultats de la commission d'évaluation ont déjà été divulgués.	135	23	Cont_Fill
Une audience au tribunal de commerce touche à sa fin. Avant de clore la séance, la juge demande si l'avocat de la défense à quelque chose à ajouter.	148	29	Cont_Fill
Une mère interroge un professeur à propos de sa fille Isabelle. Elle est surtout inquiète de savoir si sa fille va devoir redoubler.	132	23	Cont_Fill
Un journaliste interroge une experte à la commission européenne. Il lui demande quel est son pronostic économique.	114	17	Cont_Fill
Une envoyée spéciale interroge un responsable des pompiers suite à de fortes pluies dans la région. L'envoyée veut connaître la situation sur le terrain.	153	25	Cont_Fill
Une auxiliaire de laboratoire interroge une chercheuse dans un labo pharmaceutique pour savoir ce que pense le patron de la dernière expérience menée.	150	23	Cont_Fill
Un professeur interroge la directrice de département qui revient de la commission de validation des examens. Il veut connaître la décision qui a été prise au sujet de Geneviève, une étudiante.	191	31	Cont_Fill
Au cours d'un conseil de classe à propos d'un élève de dernière année, un professeur demande pourquoi Cédric, un élève, a décidé de changer d'orientation.	154	28	Cont_Fill
La construction d'un incinérateur est en discussion dans un conseil municipal. La mairesse veut savoir quelles sont les réactions des habitants du village.	155	24	Cont_Fill

Monsieur Tremblay est à l'épicerie après un rendez-vous chez sa nutritionniste. Tentant de sélectionner un jus, il se rappelle qu'elle lui a dit qu'il fallait qu'il consomme des oranges ou des pamplemousses pour faire le plein de vitamine D.	241	44	Cont_Fill
Un groupe d'amis veut organiser un party, et ils discutent pour décider chez qui il aura lieu. La personne choisie ne peut avoir un chien ou un chat, car Bob est allergique.	173	33	Cont_Fill
Pour un numéro d'un cirque, un dresseur danse avec un jaguar. Avant de commencer, il avise l'audience que pour des raisons de sécurité on peut prendre des photos avant ou après le numéro, mais pas pendant.	205	38	Cont_Fill
Pour s'inscrire à une université au Québec, il faut être âgé de plus de 20 ans ou avoir un diplôme d'étude collégiale.	117	24	Cont_Fill
Il y a longtemps, Isabelle a rencontré une voyante qui lui a prédit qu'elle finirait avec un comptable millionnaire ou un entraîneur de 6 pieds.	144	26	Cont_Fill
Sophie vient de finir de travailler, après quoi elle doit se rendre à un rendez-vous chez le dentiste. Puisqu'elle est en retard, elle a le temps de se changer ou d'arriver à l'heure.	183	37	Cont_Fill
Condamné à mort, Jack peut, selon la loi, choisir d'être électrocuté ou pendu. Un sourire dément au lèvres, il dit qu'il voudrait être pendu et électrocuté en même temps.	169	31	Cont_Fill
Fanny commande un café avec un, deux, ou trois sucres, en précisant qu'elle n'a pas de préférences. Le commis lui remets un café avec 6 sucres.	144	28	Cont_Fill
Au courant de l'année, Bella a été en couple avec Jacob ou Édouard, dépendant des mois.	87	17	Cont_Fill
Christine veut appeler son amie Sylvie, dont il ne trouve plus le numéro de téléphone.	87	15	Cont_Fill
Alexandre est un amateur de bon vin. Dans sa cave il a conservé plusieurs bouteilles de grands crus de sa région natale. Pour l'anniversaire de sa fille il décide d'ouvrir une des meilleures bouteilles. Au moment de l'ouvrir il réalise que le vin a tourné au vinaigre.	268	50	Cont_Fill

Sur l'étagère, les bottins téléphoniques sont lourds et plus lourds que les dossiers qui se trouvent à l'autre côté.	117	21	Phrase_Fill
Dans un magasin Marc-André a vu une table de salon longue et moins longue que son canapé donc il est très content.	116	23	Phrase_Fill
La nouvelle chaîne du vélo de Camille est solide et plus solide que l'ancienne, elle en est très satisfaite.	108	20	Phrase_Fill
Les mouvements du danseur sont lents mais aussi lents que la musique.	69	12	Phrase_Fill
L'imperméable proposé par la vendeuse est court mais aussi court que le genou de Guillaume.	91	16	Phrase_Fill
La nouvelle version du rapport est détaillée mais aussi détaillée que l'ancienne donc il n'est pas raccourci.	109	19	Phrase_Fill
Au final, Romain décide de lui faire confiance quand même.	58	10	Phrase_Fill
Les parents de Sandra ont attendu longtemps à la gare, mais au final ils étaient content de le voir.	100	19	Phrase_Fill
Comme elle est en retard, elle part immédiatement après avoir reçu l'œuf.	73	13	Phrase_Fill
Elle l'accuse de ne pas avoir fait attention autour de lui.	59	12	Phrase_Fill
Il décide de ne pas appliquer sa méthode.	41	8	Phrase_Fill
Elle suggère d'acheter l'écharpe et d'aller prendre un café ensuite.	68	13	Phrase_Fill
Le lendemain elle est surprise de recevoir un appel du ministre en personne.	76	13	Phrase_Fill
Toutefois l'agence décide de ne pas le garder parce qu'il est systématiquement en retard.	89	16	Phrase_Fill
Elle va demander au libraire s'il peut la conseiller sur le meilleur endroit où le ranger.	90	17	Phrase_Fill
Après en avoir débattu, ils décident de demander à un gardien.	62	11	Phrase_Fill
Véronique n'est pas d'accord avec son père et décide de poser quand même la friandise dans le chariot.	102	20	Phrase_Fill
Heureusement ils rencontrent Isaac et Audrey qui ont des provisions avec eux.	77	12	Phrase_Fill
La représentante syndicale lui répond que la DRH n'a aucun doute que le licenciement est bien dans les intentions de la direction.	130	23	Phrase_Fill
Marie lui répond qu'il n'est pas faux que Richard va souvent au Brésil.	71	15	Phrase_Fill

Son avocate lui répond que de façon surprenante, l'avocat adverse ne conteste pas que Bertrand apporte une preuve inattaquable.	127	20	Phrase_Fill
La directrice dit qu'ils ne sont pas encore publics, mais que la commission n'a pas nié que le projet fait des progrès.	119	24	Phrase_Fill
L'avocat répond que sa cliente ne disconvient pas que le commerçant a pu faire une erreur.	90	17	Phrase_Fill
Le professeur lui répond qu'il ne nie pas qu'Isabelle a fait de grands efforts.	79	16	Phrase_Fill
L'experte lui répond que personne ne doute que l'Europe doive faire une cure d'austérité.	89	17	Phrase_Fill
Le pompier lui répond que les mesures de précaution n'ont pas empêché que la crue envahisse les zones habitées.	111	20	Phrase_Fill
La chercheuse lui dit que le patron n'exige pas que l'expérience soit recommencée.	82	15	Phrase_Fill
La directrice lui apprend qu'ils n'ont pas obtenu que Geneviève puisse repasser son examen.	91	16	Phrase_Fill
La déléguée des élèves lui répond que les parents de Cédric ne veulent pas qu'il devienne médecin.	98	18	Phrase_Fill
Son conseiller lui répond que personne ne désire que l'incinérateur soit installé aussi près du village.	104	17	Phrase_Fill
Il choisit un jus d'orange et pamplemousse.	42	8	Phrase_Fill
Sandrine a un chat.	20	4	Phrase_Fill
Hélène a pris une photo avant le numéro et une photo après.	59	12	Phrase_Fill
Marie-Pierre a 28 ans et vient d'obtenir son diplôme d'études collégiales en sciences pures.	91	17	Phrase_Fill
Elle est présentement mariée à un comptable de 6 pieds, qui est aussi entraîneur et millionnaire	95	16	Phrase_Fill
Puisqu'elle est en retard, elle a le temps de se changer ou d'arriver à l'heure.	80	18	Phrase_Fill
Un sourire dément au lèvres, il dit qu'il voudrait être pendu et électrocuté en même temps.	90	17	Phrase_Fill
Le commis lui remet un café avec 6 sucres.	43	9	Phrase_Fill
Jacob et Édouard l'aiment beaucoup et sont dévoués et fidèles.	63	11	Phrase_Fill

Il se rappelle vividement que les derniers chiffres sont 13 ou 16.	66	12	Phrase_Fill
Alexandre dit à sa fille qu'il va descendre pour aller chercher une autre bouteille.	84	15	Phrase_Fill
Les bottins téléphoniques sont-ils lourds ?	42	6	Question_Neutre_Fill
La table de salon est-elle longue ?	34	7	Question_Neutre_Fill
La nouvelle chaîne de vélo est-elle fragile ?	44	8	Question_Neutre_Fill
Le danseur bouge-t-il lentement ?	32	6	Question_Neutre_Fill
L'imperméable proposé par la vendeuse est-il long ?	51	9	Question_Neutre_Fill
Le nouveau rapport est-il plus court que le précédent ?	54	10	Question_Neutre_Fill
Est-ce que Romain fait confiance à Arianne ?	43	8	Question_Neutre_Fill
Est-ce que les parents de Sandra ont rencontré François ?	55	10	Question_Neutre_Fill
Est-ce que le plombier est en retard ?	37	8	Question_Neutre_Fill
Est-ce que la dame âgée a appelé la police ?	43	10	Question_Neutre_Fill
Est-ce que la professeur d'anglais est stricte avec ses étudiants ?	66	12	Question_Neutre_Fill
Est-ce que Louise est intéressée par les sandales bleues ?	57	10	Question_Neutre_Fill
Est-ce que l'assistant a transmis le message du journaliste au ministre ?	72	13	Question_Neutre_Fill
Est-ce que le guide Rosemont est souvent en retard ?	51	10	Question_Neutre_Fill
Est-ce que tous les livres sont rangés derrière une vitrine ?	60	11	Question_Neutre_Fill
Est-ce que Sarah affirme que les tatous sont des animaux nocturnes ?	67	12	Question_Neutre_Fill
Est-ce que Véronique a envie de manger une pomme ?	49	10	Question_Neutre_Fill
Est-ce que Boris et Catherine vont pouvoir boire leur vin ?	58	11	Question_Neutre_Fill
Est-ce que la DRH pense que la direction va licencier des employés ?	68	13	Question_Neutre_Fill
Est-ce que Richard est déjà allé au Brésil ?	44	9	Question_Neutre_Fill
Est-ce que l'avocat adverse reconnaît les arguments de Bertrand ?	65	11	Question_Neutre_Fill
Est-ce que le projet fait des progrès ?	39	8	Question_Neutre_Fill
Est-ce que la cliente pense qu'il est possible que le commerçant se soit trompé ?	81	16	Question_Neutre_Fill
Est-ce que le professeur reconnaît qu'Isabelle a fait des efforts ?	66	12	Question_Neutre_Fill
L'experte pense-t-elle que l'austérité soit nécessaire ?	56	10	Question_Neutre_Fill
Les zones habitées ont-elles été inondées ?	43	7	Question_Neutre_Fill
Le patron est-il satisfait de l'expérience ?	44	8	Question_Neutre_Fill
Geneviève va-t-elle repasser son examen ?	41	7	Question_Neutre_Fill

Est-ce que les parents de Cédric veulent qu'il devienne médecin ?	65	12	Question_Neutre_Fill
Les habitants du village veulent-ils l'incinérateur dans le village ?	69	11	Question_Neutre_Fill
Est-ce que le jus que monsieur Tremblay a choisi est conforme au conseil de sa nutritionniste ?	94	17	Question_Neutre_Fill
Est-ce que le party peut avoir lieu chez Sandrine ?	50	10	Question_Neutre_Fill
Est-ce que Hélène a respecté les instructions du dresseur ?	58	10	Question_Neutre_Fill
Est-ce que Marie-Pierre peut s'inscrire à une université au Québec ?	67	13	Question_Neutre_Fill
Est-ce que la prédiction de la voyante s'est réalisée ?	54	11	Question_Neutre_Fill
Est-ce que Sophie arrivera à l'heure à son rendez-vous avec des nouveaux vêtements ?	83	16	Question_Neutre_Fill
Est-ce que ce mode d'exécution est légal ?	41	9	Question_Neutre_Fill
Le commis a-t-il donné à Fanny le café qu'il a commandé ?	56	14	Question_Neutre_Fill
Est-ce que Édouard a été en couple avec Jacob ?	46	10	Question_Neutre_Fill
Se peut-il que le numéro de Sylvie soit 514-795-1136 ?	53	12	Question_Neutre_Fill
Est-ce que c'est la fille d'Alexandre qui va aller chercher une autre bouteille ?	80	16	Question_Neutre_Fill

ANNEXE C

TABLEAUX D'ANALYSES

C.1 Tableaux pour la durée du parcours régressif

C.1.1 Par zones

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.06	-0.05	4.29 – 5.83	-0.23 – 0.13	<0.001
Valence [Plus]	-0.01	-0.01	-0.10 – 0.09	-0.15 – 0.13	0.915
Contexte [Neutre]	0.07	0.09	-0.04 – 0.17	-0.05 – 0.24	0.204
ScolCal	0.04	0.04	-0.08 – 0.16	-0.08 – 0.17	0.509
Empan	-0.01	-0.01	-0.21 – 0.19	-0.13 – 0.11	0.915
sFreqLex	0.11	0.19	0.06 – 0.17	0.10 – 0.28	<0.001
nbCarac	0.02	0.06	-0.01 – 0.05	-0.03 – 0.16	0.206
Random Effects					
σ^2	0.44				
T00 Participant	0.04				
T00 MediaType	0.02				
ICC	0.11				
N Participant	28				
N MediaType	16				
Observations	778				
Marginal R ² / Conditional R ²	0.029 / 0.137				

Table C.1 Durée du parcours régressif sur les mots de la zone -1

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.21	-0.03	4.64 – 5.78	-0.16 – 0.10	<0.001
Valence [Plus]	0.03	0.05	-0.04 – 0.10	-0.06 – 0.16	0.409
Contexte [Neutre]	0.00	0.01	-0.06 – 0.07	-0.10 – 0.12	0.903
ScolCal	0.02	0.03	-0.06 – 0.11	-0.08 – 0.13	0.603
Empan	0.01	0.01	-0.13 – 0.16	-0.09 – 0.11	0.862
sFreqLex	0.02	0.02	-0.04 – 0.07	-0.05 – 0.09	0.554
nbCarac	-0.02	-0.09	-0.03 – -0.01	-0.16 – -0.03	0.007
Random Effects					
σ^2	0.35				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.05				
N Participant	28				
N MediaType	19				
Observations	1299				
Marginal R ² / Conditional R ²	0.012 / 0.065				

Table C.2 Durée du parcours régressif sur les mots de la zone 0

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.10	-0.09	4.51 – 5.69	-0.23 – 0.06	<0.001
Valence [Plus]	0.05	0.08	-0.03 – 0.12	-0.05 – 0.22	0.235
Contexte [Neutre]	0.02	0.04	-0.05 – 0.09	-0.09 – 0.17	0.532
ScolCal	0.02	0.03	-0.06 – 0.11	-0.08 – 0.14	0.599
Empan	-0.01	-0.00	-0.15 – 0.14	-0.11 – 0.11	0.944
sFreqLex	-0.03	-0.01	-0.23 – 0.18	-0.10 – 0.08	0.797
nbCarac	-0.01	-0.03	-0.03 – 0.02	-0.11 – 0.06	0.556
Random Effects					
σ^2	0.30				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.06				
N Participant	28				
N MediaType	19				
Observations	946				
Marginal R ² / Conditional R ²	0.003 / 0.067				

Table C.3 Durée du parcours régressif sur les mots de la zone 1

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.51	-0.03	4.73 – 6.29	-0.20 – 0.15	<0.001
Valence [Plus]	0.02	0.03	-0.09 – 0.12	-0.13 – 0.19	0.746
Contexte [Neutre]	-0.01	-0.02	-0.11 – 0.09	-0.18 – 0.14	0.847
ScolCal	0.07	0.08	-0.05 – 0.19	-0.05 – 0.21	0.243
Empan	-0.07	-0.04	-0.27 – 0.13	-0.17 – 0.08	0.501
sFreqLex	0.00	0.01	-0.08 – 0.09	-0.09 – 0.10	0.913
nbCarac	-0.03	-0.08	-0.06 – 0.01	-0.18 – 0.02	0.100
Random Effects					
σ^2	0.38				
T00 Participant	0.03				
T00 MediaType	0.00				
ICC	0.08				
N Participant	28				
N MediaType	19				
Observations	599				
Marginal R ² / Conditional R ²	0.014 / 0.097				

Table C.4 Durée du parcours régressif sur les mots de la zone 2

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.42	-0.01	4.75 – 6.09	-0.23 – 0.21	<0.001
Valence [Plus]	0.08	0.12	0.00 – 0.16	0.01 – 0.24	0.040
Contexte [Neutre]	-0.01	-0.02	-0.10 – 0.07	-0.14 – 0.10	0.749
ScolCal	0.07	0.08	-0.03 – 0.17	-0.03 – 0.18	0.172
Empan	-0.05	-0.03	-0.22 – 0.12	-0.13 – 0.07	0.579
sFreqLex	0.04	0.06	-0.01 – 0.09	-0.02 – 0.13	0.149
nbCarac	0.00	0.02	-0.01 – 0.02	-0.06 – 0.10	0.574
Random Effects					
σ^2	0.39				
T00 Participant	0.03				
T00 MediaType	0.07				
ICC	0.20				
N Participant	28				
N MediaType	19				
Observations	1021				
Marginal R ² / Conditional R ²	0.011 / 0.207				

Table C.5 Durée du parcours régressif sur les mots de la zone 3

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.17	0.04	4.33 – 6.01	-0.13 – 0.21	<0.001
Valence [Plus]	0.03	0.03	-0.08 – 0.13	-0.09 – 0.15	0.592
Contexte [Neutre]	-0.06	-0.06	-0.16 – 0.05	-0.18 – 0.06	0.307
ScolCal	0.11	0.09	-0.03 – 0.24	-0.02 – 0.20	0.116
Empan	0.01	0.00	-0.21 – 0.22	-0.11 – 0.11	0.955
sFreqLex	0.12	0.20	0.07 – 0.17	0.12 – 0.29	<0.001
nbCarac	0.09	0.28	0.06 – 0.11	0.20 – 0.35	<0.001
Random Effects					
σ^2	0.70				
T00 Participant	0.05				
T00 MediaType	0.03				
ICC	0.10				
N Participant	28				
N MediaType	16				
Observations	1041				
Marginal R ² / Conditional R ²	0.056 / 0.149				

Table C.6 Durée du parcours régressif sur les mots de la zone 4

C.1.2 Par mots de la zone critique

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	4.60	-0.08	3.67 – 5.54	-0.39 – 0.22	<0.001
Valence [Plus]	0.09	0.16	-0.06 – 0.23	-0.12 – 0.44	0.249
Contexte [Neutre]	0.00	0.01	-0.14 – 0.15	-0.28 – 0.29	0.948
ScolCal	-0.04	-0.06	-0.18 – 0.09	-0.26 – 0.13	0.535
Empan	0.17	0.14	-0.07 – 0.41	-0.05 – 0.33	0.160
sFreqLex	0.00	0.01	-0.12 – 0.13	-0.21 – 0.22	0.948
nbCarac	-0.01	-0.08	-0.05 – 0.03	-0.33 – 0.17	0.546
Random Effects					
σ^2	0.23				
T00 Participant	0.04				
T00 MediaType	0.02				
ICC	0.19				
N Participant	28				
N MediaType	13				
Observations	190				
Marginal R ² / Conditional R ²	0.033 / 0.215				

Table C.7 Durée du parcours régressif sur le mot à l'index (mais) -3

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.66	-0.08	4.75 – 6.57	-0.32 – 0.17	<0.001
Valence [Plus]	0.01	0.01	-0.12 – 0.13	-0.21 – 0.24	0.903
Contexte [Neutre]	0.07	0.12	-0.06 – 0.19	-0.10 – 0.35	0.279
ScolCal	0.03	0.04	-0.10 – 0.15	-0.13 – 0.20	0.682
Empan	-0.08	-0.06	-0.30 – 0.14	-0.22 – 0.10	0.492
sFreqLex	-0.08	-0.05	-0.48 – 0.31	-0.31 – 0.20	0.684
nbCarac	-0.04	-0.18	-0.10 – 0.02	-0.44 – 0.08	0.166
Random Effects					
σ^2	0.26				
T00 Participant	0.04				
T00 MediaType	0.01				
ICC	0.16				
N Participant	28				
N MediaType	19				
Observations	299				
Marginal R ² / Conditional R ²	0.028 / 0.180				

Table C.8 Durée du parcours régressif sur le mot à l'index (mais) -2

<i>Predictors</i>	logdurParcReg				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	4.42	-0.06	1.95 – 6.89	-0.25 – 0.13	<0.001
Valence [Plus]	0.08	0.13	-0.04 – 0.21	-0.07 – 0.34	0.201
Contexte [Neutre]	-0.00	-0.01	-0.13 – 0.13	-0.21 – 0.20	0.950
ScolCal	0.01	0.01	-0.11 – 0.13	-0.12 – 0.15	0.864
Empan	0.02	0.01	-0.18 – 0.22	-0.12 – 0.15	0.843
sFreqLex	-1.79	-0.04	-6.65 – 3.07	-0.15 – 0.07	0.470
nbCarac	-0.03	-0.09	-0.06 – 0.01	-0.20 – 0.02	0.116
Random Effects					
σ^2	0.38				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.05				
N Participant	28				
N MediaType	19				
Observations	371				
Marginal R ² / Conditional R ²	0.011 / 0.064				

Table C.9 Durée du parcours régressif sur le mot à l'index (mais) -1

logdurParcReg					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.47	-0.04	4.76 – 6.18	-0.25 – 0.17	<0.001
Valence [Plus]	-0.02	-0.03	-0.15 – 0.12	-0.26 – 0.20	0.804
Contexte [Neutre]	0.06	0.10	-0.07 – 0.19	-0.13 – 0.34	0.379
ScolCal	0.05	0.07	-0.04 – 0.15	-0.06 – 0.19	0.286
Empan	-0.13	-0.09	-0.30 – 0.05	-0.21 – 0.03	0.160
Random Effects					
σ^2	0.33				
T00 Participant	0.00				
ICC	0.01				
N Participant	28				
Observations	284				
Marginal R ² / Conditional R ²	0.014 / 0.021				

Table C.10 Durée du parcours régressif sur le mot à l'index (mais) 0

logdurParcReg					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.26	0.08	4.55 – 5.97	-0.24 – 0.39	<0.001
Valence [Plus]	-0.15	-0.26	-0.48 – 0.17	-0.82 – 0.30	0.357
Contexte [Neutre]	0.04	0.07	-0.09 – 0.17	-0.15 – 0.29	0.529
ScolCal	-0.02	-0.03	-0.13 – 0.08	-0.16 – 0.10	0.671
Empan	0.02	0.02	-0.15 – 0.20	-0.11 – 0.15	0.793
sFreqLex	0.36	0.15	-0.31 – 1.03	-0.13 – 0.43	0.298
Random Effects					
σ^2	0.33				
T00 Participant	0.01				
ICC	0.04				
N Participant	28				
Observations	336				
Marginal R ² / Conditional R ²	0.006 / 0.043				

Table C.11 Durée du parcours régressif sur le mot à l'index (mais) 1

logdurParcReg					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	3.79	-0.13	1.71 – 5.88	-0.34 – 0.08	<0.001
Valence [Plus]	0.15	0.28	0.04 – 0.26	0.07 – 0.50	0.010
Contexte [Neutre]	-0.03	-0.05	-0.14 – 0.08	-0.26 – 0.16	0.645
ScolCal	0.01	0.02	-0.10 – 0.13	-0.14 – 0.18	0.829
Empan	0.06	0.05	-0.13 – 0.25	-0.11 – 0.20	0.545
sFreqLex	-2.11	-0.06	-6.11 – 1.89	-0.17 – 0.05	0.301
nbCarac	-0.01	-0.03	-0.03 – 0.02	-0.14 – 0.09	0.664
Random Effects					
σ^2	0.25				
T00 Participant	0.03				
ICC	0.09				
N Participant	28				
Observations	326				
Marginal R ² / Conditional R ²	0.024 / 0.117				

Table C.12 Durée du parcours régressif sur le mot à l'index (mais) 2

logdurParcReg					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	4.77	-0.06	3.12 – 6.42	-0.37 – 0.25	<0.001
Valence [Plus]	-0.02	-0.03	-0.23 – 0.19	-0.32 – 0.25	0.829
Contexte [Neutre]	0.03	0.03	-0.18 – 0.23	-0.24 – 0.31	0.804
ScolCal	0.07	0.07	-0.19 – 0.33	-0.18 – 0.32	0.601
Empan	0.10	0.05	-0.34 – 0.54	-0.17 – 0.27	0.656
Random Effects					
σ^2	0.42				
T00 Participant	0.14				
ICC	0.25				
N Participant	27				
Observations	169				
Marginal R ² / Conditional R ²	0.009 / 0.261				

Table C.13 Durée du parcours régressif sur le mot à l'index (mais) 3

logdurParcReg					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.80	-0.07	4.78 – 6.81	-0.28 – 0.14	<0.001
Valence [Plus]	0.08	0.12	-0.06 – 0.21	-0.10 – 0.34	0.270
Contexte [Neutre]	0.01	0.01	-0.13 – 0.14	-0.21 – 0.23	0.932
ScolCal	0.06	0.08	-0.06 – 0.19	-0.07 – 0.23	0.317
Empan	-0.19	-0.13	-0.41 – 0.02	-0.28 – 0.02	0.082
sFreqLex	-0.05	-0.05	-0.25 – 0.16	-0.24 – 0.15	0.645
nbCarac	0.00	0.00	-0.13 – 0.13	-0.19 – 0.19	0.998
Random Effects					
σ^2	0.35				
T00 Participant	0.03				
ICC	0.08				
N Participant	28				
Observations	318				
Marginal R^2 / Conditional R^2	0.024 / 0.097				

Table C.14 Durée du parcours régressif sur le mot à l'index (mais) 4

C.2 Tableaux pour la durée de la première fixation

C.2.1 Par zones

logdurPremFix					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	4.98	0.02	4.34 – 5.62	-0.16 – 0.20	<0.001
Valence [Plus]	-0.04	-0.08	-0.10 – 0.02	-0.22 – 0.05	0.227
Contexte [Neutre]	0.01	0.03	-0.05 – 0.08	-0.11 – 0.17	0.667
ScolCal	0.04	0.07	-0.06 – 0.14	-0.09 – 0.23	0.394
Empan	0.00	0.00	-0.16 – 0.17	-0.15 – 0.16	0.960
sFreqLex	-0.02	-0.06	-0.06 – 0.01	-0.14 – 0.03	0.187
nbCarac	0.00	0.02	-0.01 – 0.02	-0.07 – 0.10	0.727
Random Effects					
σ^2	0.20				
T00 Participant	0.03				
T00 MediaType	0.00				
ICC	0.13				
N Participant	28				
N MediaType	16				
Observations	786				
Marginal R^2 / Conditional R^2	0.010 / 0.144				

Table C.15 Durée de la première fixation sur les mots de la zone -1

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.01	-0.02	4.44 – 5.58	-0.19 – 0.14	<0.001
Valence [Plus]	0.03	0.07	-0.02 – 0.08	-0.03 – 0.18	0.180
Contexte [Neutre]	-0.01	-0.03	-0.06 – 0.04	-0.14 – 0.08	0.631
ScolCal	0.07	0.11	-0.02 – 0.16	-0.03 – 0.25	0.117
Empan	-0.01	-0.01	-0.15 – 0.14	-0.14 – 0.12	0.912
sFreqLex	-0.01	-0.01	-0.05 – 0.03	-0.08 – 0.05	0.655
nbCarac	0.00	0.02	-0.01 – 0.01	-0.05 – 0.08	0.667
Random Effects					
σ^2	0.20				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.13				
N Participant	28				
N MediaType	19				
Observations	1325				
Marginal R ² / Conditional R ²	0.014 / 0.141				

Table C.16 Durée de la première fixation sur les mots de la zone 0

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.13	0.01	4.48 – 5.79	-0.18 – 0.21	<0.001
Valence [Plus]	0.01	0.01	-0.05 – 0.06	-0.12 – 0.14	0.863
Contexte [Neutre]	-0.04	-0.10	-0.10 – 0.01	-0.22 – 0.02	0.105
ScolCal	0.07	0.12	-0.02 – 0.17	-0.04 – 0.29	0.139
Empan	-0.04	-0.04	-0.21 – 0.12	-0.20 – 0.12	0.615
sFreqLex	0.00	0.00	-0.15 – 0.16	-0.08 – 0.09	0.963
nbCarac	-0.01	-0.04	-0.03 – 0.01	-0.12 – 0.05	0.397
Random Effects					
σ^2	0.16				
T00 Participant	0.03				
T00 MediaType	0.00				
ICC	0.18				
N Participant	28				
N MediaType	19				
Observations	988				
Marginal R ² / Conditional R ²	0.019 / 0.196				

Table C.17 Durée de la première fixation sur les mots de la zone 1

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.07	-0.05	4.34 – 5.80	-0.27 – 0.17	<0.001
Valence [Plus]	0.02	0.05	-0.04 – 0.09	-0.10 – 0.20	0.479
Contexte [Neutre]	-0.01	-0.01	-0.08 – 0.06	-0.16 – 0.13	0.845
ScolCal	0.06	0.09	-0.05 – 0.17	-0.08 – 0.27	0.298
Empan	-0.03	-0.02	-0.21 – 0.16	-0.19 – 0.14	0.777
sFreqLex	-0.00	-0.01	-0.07 – 0.06	-0.10 – 0.09	0.887
nbCarac	0.00	0.01	-0.02 – 0.03	-0.08 – 0.11	0.771
Random Effects					
σ^2	0.17				
T00 Participant	0.04				
T00 MediaType	0.01				
ICC	0.22				
N Participant	28				
N MediaType	19				
Observations	635				
Marginal R ² / Conditional R ²	0.009 / 0.226				

Table C.18 Durée de la première fixation sur les mots de la zone 2

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.19	0.04	4.68 – 5.69	-0.12 – 0.20	<0.001
Valence [Plus]	-0.04	-0.09	-0.09 – 0.01	-0.21 – 0.02	0.116
Contexte [Neutre]	0.01	0.03	-0.04 – 0.07	-0.09 – 0.15	0.641
ScolCal	0.07	0.12	-0.01 – 0.15	-0.01 – 0.24	0.072
Empan	-0.04	-0.04	-0.17 – 0.09	-0.16 – 0.08	0.522
sFreqLex	-0.03	-0.05	-0.06 – 0.01	-0.13 – 0.02	0.129
nbCarac	0.01	0.05	-0.00 – 0.02	-0.02 – 0.12	0.186
Random Effects					
σ^2	0.17				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.11				
N Participant	28				
N MediaType	19				
Observations	1086				
Marginal R ² / Conditional R ²	0.025 / 0.128				

Table C.19 Durée de la première fixation sur les mots de la zone 3

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.14	0.03	4.54 – 5.75	-0.14 – 0.21	<0.001
Valence [Plus]	-0.04	-0.09	-0.10 – 0.01	-0.20 – 0.02	0.127
Contexte [Neutre]	-0.03	-0.06	-0.09 – 0.03	-0.18 – 0.05	0.277
ScolCal	0.12	0.18	0.03 – 0.21	0.04 – 0.32	0.012
Empan	-0.03	-0.02	-0.18 – 0.13	-0.17 – 0.12	0.734
sFreqLex	-0.04	-0.13	-0.07 – -0.02	-0.21 – -0.06	0.001
nbCarac	0.00	0.00	-0.01 – 0.01	-0.07 – 0.08	0.951
Random Effects					
σ^2	0.21				
T00 Participant	0.03				
T00 MediaType	0.01				
ICC	0.14				
N Participant	28				
N MediaType	16				
Observations	1123				
Marginal R ² / Conditional R ²	0.050 / 0.184				

Table C.20 Durée de la première fixation sur les mots de la zone 4

C.2.2 Par mots de la zone critique

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	4.31	-0.12	3.56 – 5.06	-0.39 – 0.15	<0.001
Valence [Plus]	0.07	0.16	-0.05 – 0.19	-0.12 – 0.44	0.272
Contexte [Neutre]	0.04	0.10	-0.08 – 0.17	-0.19 – 0.38	0.503
ScolCal	0.03	0.06	-0.08 – 0.14	-0.13 – 0.25	0.544
Empan	0.13	0.13	-0.06 – 0.33	-0.06 – 0.32	0.170
sFreqLex	0.06	0.12	-0.03 – 0.14	-0.05 – 0.29	0.174
nbCarac	0.03	0.20	0.00 – 0.05	0.03 – 0.38	0.019
Random Effects					
σ^2	0.17				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.12				
N Participant	28				
N MediaType	13				
Observations	192				
Marginal R ² / Conditional R ²	0.061 / 0.173				

Table C.21 Durée de la première fixation sur le mot à l'index (mais) -3

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.36	0.01	4.66 – 6.06	-0.21 – 0.23	<0.001
Valence [Plus]	-0.03	-0.06	-0.13 – 0.08	-0.29 – 0.17	0.618
Contexte [Neutre]	0.01	0.02	-0.10 – 0.11	-0.21 – 0.24	0.883
ScolCal	0.03	0.04	-0.07 – 0.13	-0.12 – 0.20	0.605
Empan	-0.03	-0.03	-0.20 – 0.14	-0.18 – 0.12	0.715
sFreqLex	-0.16	-0.13	-0.42 – 0.09	-0.32 – 0.07	0.205
nbCarac	-0.02	-0.13	-0.06 – 0.01	-0.32 – 0.07	0.200
Random Effects					
σ^2	0.20				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.09				
N Participant	28				
N MediaType	19				
Observations	309				
Marginal R ² / Conditional R ²	0.009 / 0.100				

Table C.22 Durée de la première fixation sur le mot à l'index (mais) -2

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	4.94	-0.09	3.02 – 6.86	-0.30 – 0.12	<0.001
Valence [Plus]	0.09	0.19	-0.00 – 0.18	-0.00 – 0.38	0.055
Contexte [Neutre]	-0.00	-0.01	-0.10 – 0.09	-0.21 – 0.19	0.924
ScolCal	0.06	0.10	-0.03 – 0.16	-0.05 – 0.26	0.199
Empan	-0.08	-0.08	-0.25 – 0.08	-0.23 – 0.08	0.324
sFreqLex	-0.84	-0.03	-4.59 – 2.90	-0.14 – 0.09	0.659
nbCarac	-0.01	-0.05	-0.04 – 0.02	-0.17 – 0.07	0.418
Random Effects					
σ^2	0.20				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.12				
N Participant	28				
N MediaType	19				
Observations	383				
Marginal R ² / Conditional R ²	0.024 / 0.138				

Table C.23 Durée de la première fixation sur le mot à l'index (mais) -1

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.46	0.04	4.77 – 6.14	-0.20 – 0.29	<0.001
Valence [Plus]	0.02	0.05	-0.07 – 0.11	-0.17 – 0.27	0.668
Contexte [Neutre]	-0.08	-0.20	-0.17 – 0.01	-0.42 – 0.02	0.081
ScolCal	0.03	0.05	-0.07 – 0.13	-0.12 – 0.23	0.559
Empan	-0.11	-0.10	-0.28 – 0.07	-0.26 – 0.06	0.226
Random Effects					
σ^2	0.15				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.15				
N Participant	28				
N MediaType	19				
Observations	295				
Marginal R ² / Conditional R ²	0.022 / 0.164				

Table C.24 Durée de la première fixation sur le mot à l'index (mais) 0

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.14	-0.05	4.31 – 5.98	-0.39 – 0.29	<0.001
Valence [Plus]	0.04	0.09	-0.21 – 0.30	-0.45 – 0.64	0.743
Contexte [Neutre]	-0.03	-0.07	-0.13 – 0.06	-0.27 – 0.13	0.486
ScolCal	0.05	0.09	-0.07 – 0.18	-0.11 – 0.28	0.399
Empan	-0.05	-0.05	-0.26 – 0.16	-0.25 – 0.15	0.636
sFreqLex	-0.10	-0.05	-0.61 – 0.42	-0.32 – 0.22	0.708
Random Effects					
σ^2	0.18				
T00 Participant	0.05				
T00 MediaType	0.00				
ICC	0.21				
N Participant	28				
N MediaType	19				
Observations	349				
Marginal R ² / Conditional R ²	0.009 / 0.218				

Table C.25 Durée de la première fixation sur le mot à l'index (mais) 1

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	4.89	-0.04	3.28 – 6.51	-0.27 – 0.18	<0.001
Valence [Plus]	0.04	0.08	-0.05 – 0.12	-0.12 – 0.28	0.418
Contexte [Neutre]	-0.01	-0.02	-0.10 – 0.08	-0.22 – 0.18	0.823
ScolCal	0.12	0.20	0.01 – 0.23	0.01 – 0.38	0.036
Empan	-0.01	-0.01	-0.20 – 0.17	-0.19 – 0.16	0.896
sFreqLex	0.18	0.01	-2.80 – 3.17	-0.10 – 0.11	0.904
nbCarac	0.00	0.01	-0.02 – 0.02	-0.10 – 0.12	0.859
Random Effects					
σ^2	0.16				
T00 Participant	0.03				
ICC	0.17				
N Participant	28				
Observations	344				
Marginal R^2 / Conditional R^2	0.038 / 0.202				

Table C.26 Durée de la première fixation sur le mot à l'index (mais) 2

<i>Predictors</i>	logdurPremFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	4.87	0.10	3.88 – 5.86	-0.19 – 0.39	<0.001
Valence [Plus]	-0.02	-0.04	-0.15 – 0.11	-0.32 – 0.23	0.745
Contexte [Neutre]	-0.11	-0.24	-0.24 – 0.01	-0.51 – 0.02	0.075
ScolCal	0.11	0.17	-0.04 – 0.26	-0.07 – 0.40	0.159
Empan	0.01	0.01	-0.25 – 0.27	-0.20 – 0.22	0.941
Random Effects					
σ^2	0.17				
T00 Participant	0.05				
ICC	0.24				
N Participant	28				
Observations	181				
Marginal R^2 / Conditional R^2	0.039 / 0.267				

Table C.27 Durée de la première fixation sur le mot à l'index (mais) 3

logdurPremFix					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	4.87	0.10	3.88 – 5.86	-0.19 – 0.39	<0.001
Valence [Plus]	-0.02	-0.04	-0.15 – 0.11	-0.32 – 0.23	0.745
Contexte [Neutre]	-0.11	-0.24	-0.24 – 0.01	-0.51 – 0.02	0.075
ScolCal	0.11	0.17	-0.04 – 0.26	-0.07 – 0.40	0.159
Empan	0.01	0.01	-0.25 – 0.27	-0.20 – 0.22	0.941
Random Effects					
σ^2	0.17				
T00 Participant	0.05				
ICC	0.24				
N Participant	28				
Observations	181				
Marginal R ² / Conditional R ²	0.039 / 0.267				

Table C.28 Durée de la première fixation sur le mot à l'index (mais) 4

C.3 Tableaux pour la durée du premier regard

C.3.1 Par zones

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.05	0.01	4.46 – 5.65	-0.15 – 0.18	<0.001
Valence [Plus]	-0.03	-0.06	-0.10 – 0.04	-0.20 – 0.08	0.386
Contexte [Neutre]	0.02	0.03	-0.06 – 0.09	-0.11 – 0.17	0.684
ScolCal	0.06	0.09	-0.03 – 0.15	-0.05 – 0.22	0.217
Empan	-0.02	-0.02	-0.17 – 0.13	-0.15 – 0.11	0.794
sFreqLex	-0.03	-0.06	-0.07 – 0.01	-0.15 – 0.02	0.145
nbCarac	0.01	0.05	-0.01 – 0.03	-0.04 – 0.14	0.253
Random Effects					
σ^2	0.23				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.10				
N Participant	28				
N MediaType	16				
Observations	786				
Marginal R ² / Conditional R ²	0.017 / 0.111				

Table C.29 Durée du premier regard sur les mots de la zone -1

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.02	-0.04	4.49 – 5.54	-0.19 – 0.11	<0.001
Valence [Plus]	0.06	0.11	-0.00 – 0.11	-0.00 – 0.22	0.051
Contexte [Neutre]	-0.01	-0.02	-0.07 – 0.04	-0.13 – 0.09	0.665
ScolCal	0.08	0.11	0.00 – 0.16	0.00 – 0.23	0.049
Empan	-0.02	-0.02	-0.15 – 0.11	-0.13 – 0.09	0.772
sFreqLex	-0.02	-0.03	-0.07 – 0.02	-0.10 – 0.03	0.355
nbCarac	0.02	0.10	0.01 – 0.03	0.03 – 0.17	0.004
Random Effects					
σ^2	0.24				
T00 Participant	0.02				
T00 MediaType	0.01				
ICC	0.09				
N Participant	28				
N MediaType	19				
Observations	1325				
Marginal R ² / Conditional R ²	0.031 / 0.119				

Table C.30 Durée du premier regard sur les mots de la zone 0

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.13	0.00	4.47 – 5.80	-0.18 – 0.18	<0.001
Valence [Plus]	0.01	0.03	-0.05 – 0.08	-0.10 – 0.16	0.641
Contexte [Neutre]	-0.05	-0.10	-0.11 – 0.01	-0.22 – 0.02	0.093
ScolCal	0.09	0.13	-0.02 – 0.19	-0.02 – 0.28	0.096
Empan	-0.05	-0.05	-0.22 – 0.12	-0.19 – 0.10	0.531
sFreqLex	0.05	0.03	-0.12 – 0.22	-0.06 – 0.11	0.535
nbCarac	0.01	0.04	-0.01 – 0.03	-0.05 – 0.12	0.371
Random Effects					
σ^2	0.20				
T00 Participant	0.03				
T00 MediaType	0.00				
ICC	0.15				
N Participant	28				
N MediaType	19				
Observations	988				
Marginal R ² / Conditional R ²	0.020 / 0.169				

Table C.31 Durée du premier regard sur les mots de la zone 1

logdurPremRegard					
<i>Predictors</i>	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	4.99	-0.03	4.25 – 5.73	-0.24 – 0.18	<0.001
Valence [Plus]	0.04	0.07	-0.04 – 0.11	-0.08 – 0.22	0.371
Contexte [Neutre]	-0.04	-0.07	-0.11 – 0.04	-0.22 – 0.08	0.359
ScolCal	0.08	0.12	-0.03 – 0.19	-0.04 – 0.27	0.151
Empan	-0.03	-0.02	-0.22 – 0.16	-0.17 – 0.12	0.761
sFreqLex	-0.00	-0.00	-0.07 – 0.07	-0.10 – 0.09	0.934
nbCarac	0.04	0.12	0.01 – 0.07	0.02 – 0.22	0.016
Random Effects					
σ^2	0.23				
T00 Participant	0.04				
T00 MediaType	0.01				
ICC	0.17				
N Participant	28				
N MediaType	19				
Observations	635				
Marginal R ² / Conditional R ²	0.028 / 0.197				

Table C.32 Durée du premier regard sur les mots de la zone 2

logdurPremRegard					
<i>Predictors</i>	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	5.28	-0.01	4.76 – 5.79	-0.18 – 0.16	<0.001
Valence [Plus]	-0.01	-0.03	-0.07 – 0.05	-0.14 – 0.09	0.646
Contexte [Neutre]	0.05	0.09	-0.01 – 0.11	-0.03 – 0.21	0.137
ScolCal	0.09	0.13	0.02 – 0.17	0.02 – 0.24	0.018
Empan	-0.09	-0.07	-0.22 – 0.04	-0.17 – 0.03	0.190
sFreqLex	-0.02	-0.03	-0.06 – 0.02	-0.10 – 0.04	0.384
nbCarac	0.03	0.19	0.02 – 0.04	0.12 – 0.27	<0.001
Random Effects					
σ^2	0.24				
T00 Participant	0.02				
T00 MediaType	0.02				
ICC	0.12				
N Participant	28				
N MediaType	19				
Observations	1086				
Marginal R ² / Conditional R ²	0.063 / 0.174				

Table C.33 Durée du premier regard sur les mots de la zone 3

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	4.98	-0.02	4.25 – 5.71	-0.18 – 0.14	<0.001
Valence [Plus]	-0.01	-0.01	-0.07 – 0.06	-0.12 – 0.10	0.819
Contexte [Neutre]	-0.03	-0.04	-0.09 – 0.04	-0.15 – 0.07	0.459
ScolCal	0.16	0.19	0.04 – 0.27	0.05 – 0.33	0.007
Empan	-0.03	-0.02	-0.22 – 0.16	-0.16 – 0.11	0.745
sFreqLex	-0.03	-0.07	-0.06 – -0.00	-0.15 – -0.00	0.042
nbCarac	0.04	0.17	0.02 – 0.05	0.10 – 0.24	<0.001
Random Effects					
σ^2	0.31				
T00 Participant	0.04				
T00 MediaType	0.00				
ICC	0.12				
N Participant	28				
N MediaType	16				
Observations	1123				
Marginal R ² / Conditional R ²	0.082 / 0.195				

Table C.34 Durée du premier regard sur les mots de la zone 4

C.3.2 Par mots de la zone critique

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	4.51	-0.18	3.77 – 5.24	-0.44 – 0.08	<0.001
Valence [Plus]	0.09	0.18	-0.05 – 0.24	-0.10 – 0.46	0.210
Contexte [Neutre]	0.10	0.18	-0.05 – 0.25	-0.10 – 0.46	0.210
ScolCal	0.04	0.06	-0.06 – 0.14	-0.09 – 0.21	0.439
Empan	0.07	0.06	-0.11 – 0.26	-0.09 – 0.20	0.453
sFreqLex	0.07	0.12	-0.04 – 0.18	-0.07 – 0.30	0.226
nbCarac	0.05	0.28	0.01 – 0.08	0.07 – 0.48	0.008
Random Effects					
σ^2	0.26				
T00 Participant	0.00				
T00 MediaType	0.01				
ICC	0.04				
N Participant	28				
N MediaType	13				
Observations	192				
Marginal R ² / Conditional R ²	0.079 / 0.117				

Table C.35 Durée du premier regard sur le mot à l'index (mais) -3

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.50	-0.02	4.83 – 6.17	-0.22 – 0.19	<0.001
Valence [Plus]	0.01	0.01	-0.11 – 0.12	-0.22 – 0.24	0.930
Contexte [Neutre]	0.01	0.01	-0.11 – 0.12	-0.21 – 0.24	0.902
ScolCal	0.01	0.02	-0.08 – 0.11	-0.11 – 0.16	0.755
Empan	-0.08	-0.07	-0.25 – 0.08	-0.20 – 0.06	0.306
sFreqLex	-0.13	-0.09	-0.42 – 0.17	-0.30 – 0.12	0.392
nbCarac	0.01	0.05	-0.03 – 0.05	-0.16 – 0.25	0.668
Random Effects					
σ^2	0.24				
T00 Participant	0.01				
T00 MediaType	0.00				
ICC	0.05				
N Participant	28				
N MediaType	19				
Observations	309				
Marginal R ² / Conditional R ²	0.022 / 0.071				

Table C.36 Durée du premier regard sur le mot à l'index (mais) -2

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	4.77	-0.07	2.62 – 6.93	-0.27 – 0.13	<0.001
Valence [Plus]	0.11	0.21	0.01 – 0.21	0.01 – 0.41	0.036
Contexte [Neutre]	-0.03	-0.06	-0.14 – 0.07	-0.26 – 0.14	0.535
ScolCal	0.09	0.13	-0.01 – 0.19	-0.01 – 0.27	0.068
Empan	-0.08	-0.07	-0.25 – 0.09	-0.21 – 0.07	0.345
sFreqLex	-1.12	-0.03	-5.38 – 3.13	-0.15 – 0.09	0.604
nbCarac	-0.00	-0.01	-0.03 – 0.03	-0.13 – 0.11	0.909
Random Effects					
σ^2	0.25				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.09				
N Participant	28				
N MediaType	19				
Observations	383				
Marginal R ² / Conditional R ²	0.031 / 0.118				

Table C.37 Durée du premier regard sur le mot à l'index (mais) -1

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.48	0.04	4.80 – 6.15	-0.21 – 0.28	<0.001
Valence [Plus]	0.01	0.03	-0.08 – 0.11	-0.19 – 0.25	0.792
Contexte [Neutre]	-0.07	-0.17	-0.17 – 0.02	-0.40 – 0.05	0.129
ScolCal	0.01	0.02	-0.08 – 0.11	-0.14 – 0.19	0.797
Empan	-0.09	-0.09	-0.27 – 0.08	-0.24 – 0.07	0.280
Random Effects					
σ^2	0.16				
T00 Participant	0.02				
T00 MediaType	0.00				
ICC	0.13				
N Participant	28				
N MediaType	19				
Observations	295				
Marginal R ² / Conditional R ²	0.015 / 0.142				

Table C.38 Durée du premier regard sur le mot à l'index (mais) 0

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.39	-0.11	4.49 – 6.29	-0.44 – 0.23	<0.001
Valence [Plus]	0.10	0.18	-0.19 – 0.39	-0.37 – 0.73	0.511
Contexte [Neutre]	-0.02	-0.03	-0.13 – 0.09	-0.24 – 0.17	0.738
ScolCal	0.09	0.12	-0.05 – 0.22	-0.06 – 0.31	0.196
Empan	-0.12	-0.10	-0.35 – 0.11	-0.29 – 0.09	0.296
sFreqLex	-0.16	-0.07	-0.75 – 0.43	-0.34 – 0.20	0.593
Random Effects					
σ^2	0.24				
T00 Participant	0.05				
ICC	0.17				
N Participant	28				
Observations	349				
Marginal R ² / Conditional R ²	0.022 / 0.192				

Table C.39 Durée du premier regard sur le mot à l'index (mais) 1

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	4.58	-0.03	2.84 – 6.32	-0.25 – 0.19	<0.001
Valence [Plus]	0.05	0.12	-0.04 – 0.15	-0.09 – 0.32	0.258
Contexte [Neutre]	-0.04	-0.09	-0.13 – 0.05	-0.29 – 0.11	0.391
ScolCal	0.12	0.19	0.01 – 0.23	0.01 – 0.36	0.034
Empan	0.01	0.01	-0.18 – 0.20	-0.16 – 0.17	0.916
sFreqLex	-0.14	-0.00	-3.38 – 3.11	-0.11 – 0.10	0.934
nbCarac	0.02	0.10	-0.00 – 0.05	-0.01 – 0.21	0.062
Random Effects					
σ^2	0.19				
T00 Participant	0.03				
ICC	0.15				
N Participant	28				
Observations	344				
Marginal R ² / Conditional R ²	0.047 / 0.187				

Table C.40 Durée du premier regard sur le mot à l'index (mais) 2

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.08	0.08	4.00 – 6.16	-0.22 – 0.38	<0.001
Valence [Plus]	-0.06	-0.12	-0.20 – 0.08	-0.39 – 0.16	0.407
Contexte [Neutre]	-0.09	-0.17	-0.23 – 0.05	-0.43 – 0.10	0.225
ScolCal	0.14	0.20	-0.02 – 0.30	-0.03 – 0.42	0.091
Empan	-0.05	-0.03	-0.33 – 0.24	-0.24 – 0.17	0.743
Random Effects					
σ^2	0.20				
T00 Participant	0.06				
T00 MediaType	0.01				
ICC	0.24				
N Participant	28				
N MediaType	19				
Observations	181				
Marginal R ² / Conditional R ²	0.041 / 0.270				

Table C.41 Durée du premier regard sur le mot à l'index (mais) 3

logdurPremRegard					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	5.08	-0.16	4.02 – 6.14	-0.42 – 0.10	<0.001
Valence [Plus]	0.10	0.21	0.01 – 0.20	0.01 – 0.40	0.038
Contexte [Neutre]	0.01	0.01	-0.09 – 0.10	-0.19 – 0.21	0.920
ScolCal	0.10	0.15	-0.04 – 0.24	-0.06 – 0.35	0.167
Empan	-0.08	-0.07	-0.31 – 0.16	-0.26 – 0.13	0.518
sFreqLex	-0.02	-0.03	-0.21 – 0.16	-0.25 – 0.19	0.812
nbCarac	0.04	0.08	-0.08 – 0.16	-0.15 – 0.31	0.485
Random Effects					
g ²	0.19				
T00 Participant	0.06				
T00 MediaType	0.01				
ICC	0.26				
N Participant	28				
N MediaType	19				
Observations	338				
Marginal R ² / Conditional R ²	0.040 / 0.294				

Table C.42 Durée du premier regard sur le mot à l'index (mais) 4

C.4 Tableaux pour la probabilité de fixation

C.4.1 Par zones

logprobFix					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-2.83	0.04	-3.63 – -2.04	-0.19 – 0.27	<0.001
Valence [Plus]	-0.03	-0.05	-0.09 – 0.04	-0.18 – 0.08	0.433
Contexte [Neutre]	-0.01	-0.02	-0.08 – 0.06	-0.15 – 0.12	0.783
ScolCal	-0.03	-0.04	-0.15 – 0.10	-0.22 – 0.14	0.686
Empan	-0.04	-0.03	-0.24 – 0.16	-0.21 – 0.14	0.699
sFreqLex	0.05	0.11	0.01 – 0.09	0.03 – 0.19	0.010
nbCarac	0.06	0.27	0.04 – 0.08	0.18 – 0.36	<0.001
Random Effects					
g ²	0.20				
T00 Participant	0.05				
T00 MediaType	0.02				
ICC	0.25				
N Participant	28				
N MediaType	16				
Observations	786				
Marginal R ² / Conditional R ²	0.053 / 0.292				

Table C.43 Probabilité de fixation sur les mots de la zone -1

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.68	0.06	-3.36 – -2.00	-0.12 – 0.24	<0.001
Valence [Plus]	-0.04	-0.08	-0.09 – 0.01	-0.18 – 0.02	0.111
Contexte [Neutre]	0.03	0.05	-0.02 – 0.07	-0.05 – 0.15	0.316
ScolCal	-0.00	-0.00	-0.11 – 0.10	-0.16 – 0.16	0.972
Empan	-0.05	-0.05	-0.23 – 0.12	-0.20 – 0.10	0.540
sFreqLex	-0.04	-0.06	-0.08 – -0.00	-0.12 – -0.00	0.045
nbCarac	0.04	0.24	0.03 – 0.05	0.18 – 0.31	<0.001
Random Effects					
σ^2	0.18				
T00 Participant	0.04				
T00 MediaType	0.00				
ICC	0.19				
N Participant	28				
N MediaType	19				
Observations	1325				
Marginal R ² / Conditional R ²	0.082 / 0.254				

Table C.44 Probabilité de fixation sur les mots de la zone 0

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.82	0.03	-3.53 – -2.10	-0.17 – 0.23	<0.001
Valence [Plus]	0.02	0.04	-0.04 – 0.08	-0.09 – 0.16	0.549
Contexte [Neutre]	-0.03	-0.06	-0.08 – 0.03	-0.17 – 0.06	0.334
ScolCal	0.00	0.00	-0.11 – 0.11	-0.17 – 0.17	0.983
Empan	-0.01	-0.01	-0.19 – 0.18	-0.17 – 0.16	0.951
sFreqLex	-0.02	-0.01	-0.18 – 0.13	-0.10 – 0.07	0.783
nbCarac	0.03	0.11	0.01 – 0.05	0.02 – 0.19	0.012
Random Effects					
σ^2	0.17				
T00 Participant	0.04				
T00 MediaType	0.00				
ICC	0.20				
N Participant	28				
N MediaType	19				
Observations	988				
Marginal R ² / Conditional R ²	0.013 / 0.213				

Table C.45 Probabilité de fixation sur les mots de la zone 1

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-3.33	0.06	-4.02 – -2.64	-0.14 – 0.27	<0.001
Valence [Plus]	-0.03	-0.05	-0.10 – 0.04	-0.20 – 0.09	0.463
Contexte [Neutre]	-0.00	-0.00	-0.07 – 0.07	-0.15 – 0.14	0.968
ScolCal	-0.02	-0.03	-0.13 – 0.08	-0.19 – 0.12	0.669
Empan	0.08	0.07	-0.09 – 0.26	-0.08 – 0.22	0.365
sFreqLex	0.07	0.10	0.01 – 0.13	0.01 – 0.20	0.029
nbCarac	0.08	0.31	0.06 – 0.11	0.21 – 0.40	<0.001
Random Effects					
σ^2	0.18				
T00 Participant	0.03				
T00 MediaType	0.01				
ICC	0.19				
N Participant	28				
N MediaType	19				
Observations	635				
Marginal R ² / Conditional R ²	0.077 / 0.253				

Table C.46 Probabilité de fixation sur les mots de la zone 2

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.59	0.11	-3.19 – -1.98	-0.09 – 0.31	<0.001
Valence [Plus]	-0.03	-0.06	-0.09 – 0.02	-0.17 – 0.04	0.248
Contexte [Neutre]	-0.04	-0.07	-0.10 – 0.03	-0.18 – 0.05	0.250
ScolCal	-0.05	-0.07	-0.14 – 0.04	-0.20 – 0.05	0.250
Empan	-0.03	-0.02	-0.18 – 0.13	-0.13 – 0.10	0.751
sFreqLex	-0.02	-0.04	-0.06 – 0.02	-0.11 – 0.03	0.246
nbCarac	0.03	0.21	0.02 – 0.05	0.14 – 0.28	<0.001
Random Effects					
σ^2	0.23				
T00 Participant	0.03				
T00 MediaType	0.03				
ICC	0.20				
N Participant	28				
N MediaType	19				
Observations	1086				
Marginal R ² / Conditional R ²	0.064 / 0.247				

Table C.47 Probabilité de fixation sur les mots de la zone 3

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.91	-0.03	-3.53 – -2.29	-0.23 – 0.16	<0.001
Valence [Plus]	-0.00	-0.00	-0.06 – 0.05	-0.11 – 0.11	0.966
Contexte [Neutre]	-0.03	-0.05	-0.09 – 0.03	-0.17 – 0.06	0.353
ScolCal	-0.04	-0.07	-0.14 – 0.05	-0.21 – 0.08	0.365
Empan	0.08	0.07	-0.08 – 0.24	-0.07 – 0.21	0.328
sFreqLex	0.05	0.15	0.02 – 0.08	0.07 – 0.22	<0.001
nbCarac	0.01	0.06	-0.00 – 0.02	-0.01 – 0.13	0.112
Random Effects					
σ^2	0.22				
T00 Participant	0.03				
T00 MediaType	0.01				
ICC	0.16				
N Participant	28				
N MediaType	16				
Observations	1123				
Marginal R ² / Conditional R ²	0.022 / 0.182				

Table C.48 Probabilité de fixation sur les mots de la zone 4

C.4.2 Par mots de la zone critique

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.14	-0.01	-3.02 – -1.27	-0.32 – 0.29	<0.001
Valence [Plus]	-0.05	-0.09	-0.18 – 0.08	-0.35 – 0.16	0.471
Contexte [Neutre]	0.05	0.10	-0.09 – 0.19	-0.16 – 0.36	0.461
ScolCal	0.07	0.11	-0.06 – 0.20	-0.08 – 0.30	0.270
Empan	-0.18	-0.15	-0.40 – 0.05	-0.33 – 0.04	0.123
sFreqLex	-0.14	-0.24	-0.25 – -0.04	-0.42 – -0.07	0.007
Random Effects					
σ^2	0.19				
T00 Participant	0.04				
T00 MediaType	0.02				
ICC	0.24				
N Participant	28				
N MediaType	13				
Observations	192				
Marginal R ² / Conditional R ²	0.087 / 0.310				

Table C.49 Probabilité de fixation sur le mot à l'index (mais) -3

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.37	0.04	-3.16 – -1.58	-0.20 – 0.27	<0.001
Valence [Plus]	0.01	0.02	-0.09 – 0.11	-0.19 – 0.23	0.873
Contexte [Neutre]	-0.00	-0.00	-0.10 – 0.10	-0.21 – 0.20	0.964
ScolCal	-0.07	-0.10	-0.18 – 0.05	-0.28 – 0.08	0.265
Empan	-0.06	-0.05	-0.26 – 0.14	-0.21 – 0.12	0.571
sFreqLex	-0.36	-0.26	-0.53 – -0.19	-0.39 – -0.14	<0.001
Random Effects					
σ^2	0.18				
T00 Participant	0.03				
T00 MediaType	0.01				
ICC	0.18				
N Participant	28				
N MediaType	19				
Observations	309				
Marginal R ² / Conditional R ²	0.084 / 0.249				

Table C.50 Probabilité de fixation sur le mot à l'index (mais) -2

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.45	-0.01	-4.91 – 0.00	-0.27 – 0.26	0.050
Valence [Plus]	-0.02	-0.05	-0.11 – 0.06	-0.23 – 0.13	0.580
Contexte [Neutre]	0.06	0.12	-0.03 – 0.14	-0.06 – 0.30	0.197
ScolCal	-0.00	-0.00	-0.13 – 0.13	-0.21 – 0.20	0.969
Empan	0.02	0.02	-0.20 – 0.24	-0.18 – 0.22	0.858
sFreqLex	0.49	0.02	-4.15 – 5.14	-0.13 – 0.16	0.834
Random Effects					
σ^2	0.16				
T00 Participant	0.05				
T00 MediaType	0.02				
ICC	0.31				
N Participant	28				
N MediaType	19				
Observations	383				
Marginal R ² / Conditional R ²	0.004 / 0.311				

Table C.51 Probabilité de fixation sur le mot à l'index (mais) -1

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.71	0.06	-3.47 – -1.96	-0.21 – 0.33	<0.001
Valence [Plus]	0.00	0.00	-0.09 – 0.09	-0.21 – 0.21	0.994
Contexte [Neutre]	-0.03	-0.07	-0.12 – 0.06	-0.28 – 0.15	0.543
ScolCal	-0.03	-0.06	-0.14 – 0.08	-0.26 – 0.14	0.570
Empan	-0.01	-0.01	-0.20 – 0.19	-0.19 – 0.18	0.942
Random Effects					
σ^2	0.13				
T00 Participant	0.03				
T00 MediaType	0.01				
ICC	0.23				
N Participant	28				
N MediaType	19				
Observations	295				
Marginal R ² / Conditional R ²	0.004 / 0.238				

Table C.52 Probabilité de fixation sur le mot à l'index (mais) 0

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.34	-0.12	-3.21 – -1.48	-0.46 – 0.21	<0.001
Valence [Plus]	0.14	0.28	-0.13 – 0.40	-0.26 – 0.82	0.304
Contexte [Neutre]	-0.00	-0.01	-0.10 – 0.09	-0.20 – 0.19	0.957
ScolCal	0.07	0.10	-0.06 – 0.20	-0.09 – 0.30	0.310
Empan	-0.14	-0.13	-0.36 – 0.08	-0.32 – 0.07	0.202
sFreqLex	-0.23	-0.11	-0.76 – 0.30	-0.38 – 0.15	0.395
Random Effects					
σ^2	0.19				
T00 Participant	0.05				
T00 MediaType	0.00				
ICC	0.22				
N Participant	28				
N MediaType	19				
Observations	349				
Marginal R ² / Conditional R ²	0.024 / 0.240				

Table C.53 Probabilité de fixation sur le mot à l'index (mais) 1

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-3.59	0.02	-5.86 – -1.32	-0.25 – 0.28	0.002
Valence [Plus]	-0.00	-0.01	-0.09 – 0.09	-0.20 – 0.19	0.952
Contexte [Neutre]	-0.02	-0.04	-0.11 – 0.07	-0.23 – 0.15	0.652
ScolCal	-0.05	-0.09	-0.18 – 0.08	-0.29 – 0.12	0.421
Empan	0.15	0.13	-0.07 – 0.37	-0.06 – 0.33	0.182
sFreqLex	-0.95	-0.03	-5.16 – 3.25	-0.17 – 0.11	0.656
Random Effects					
σ^2	0.15				
T00 Participant	0.05				
T00 MediaType	0.01				
ICC	0.29				
N Participant	28				
N MediaType	19				
Observations	344				
Marginal R ² / Conditional R ²	0.021 / 0.307				

Table C.54 Probabilité de fixation sur le mot à l'index (mais) 2

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.83	-0.02	-3.93 – -1.73	-0.34 – 0.29	<0.001
Valence [Plus]	0.04	0.09	-0.08 – 0.17	-0.17 – 0.36	0.484
Contexte [Neutre]	-0.01	-0.03	-0.14 – 0.11	-0.29 – 0.23	0.829
ScolCal	0.02	0.02	-0.15 – 0.18	-0.23 – 0.28	0.851
Empan	-0.05	-0.04	-0.34 – 0.24	-0.26 – 0.19	0.746
Random Effects					
σ^2	0.16				
T00 Participant	0.07				
T00 MediaType	0.00				
ICC	0.33				
N Participant	28				
N MediaType	19				
Observations	181				
Marginal R ² / Conditional R ²	0.004 / 0.335				

Table C.55 Probabilité de fixation sur le mot à l'index (mais) 3

<i>Predictors</i>	logprobFix				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.89	0.03	-3.69 – -2.09	-0.24 – 0.30	<0.001
Valence [Plus]	-0.04	-0.09	-0.13 – 0.05	-0.29 – 0.10	0.347
Contexte [Neutre]	0.03	0.06	-0.07 – 0.12	-0.14 – 0.25	0.585
ScolCal	-0.04	-0.07	-0.16 – 0.08	-0.26 – 0.13	0.507
Empan	0.08	0.07	-0.12 – 0.29	-0.11 – 0.26	0.442
sFreqLex	-0.08	-0.10	-0.19 – 0.03	-0.24 – 0.04	0.174
Random Effects					
σ^2	0.16				
T00 Participant	0.04				
T00 MediaType	0.02				
ICC	0.28				
N Participant	28				
N MediaType	19				
Observations	338				
Marginal R ² / Conditional R ²	0.021 / 0.293				

Table C.56 Probabilité de fixation sur le mot à l'index (mais) 4

C.5 Tableaux pour probabilité de régression reçue

C.5.1 Par zones

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.18	0.28	-3.27 – -1.09	0.03 – 0.53	<0.001
Valence [Plus]	-0.17	-0.26	-0.27 – -0.07	-0.41 – -0.11	0.001
Contexte [Neutre]	-0.11	-0.17	-0.21 – -0.01	-0.32 – -0.01	0.035
ScolCal	-0.03	-0.03	-0.20 – 0.14	-0.23 – 0.16	0.735
Empan	-0.04	-0.03	-0.32 – 0.23	-0.22 – 0.16	0.756
sFreqLex	-0.02	-0.03	-0.07 – 0.04	-0.13 – 0.07	0.568
nbCarac	0.03	0.11	0.00 – 0.07	0.01 – 0.21	0.036
Random Effects					
σ^2	0.31				
T00 Participant	0.09				
T00 MediaType	0.03				
ICC	0.28				
N Participant	28				
N MediaType	16				
Observations	571				
Marginal R ² / Conditional R ²	0.046 / 0.316				

Table C.57 Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone -1

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.47	0.12	-3.57 – -1.38	-0.13 – 0.38	<0.001
Valence [Plus]	-0.18	-0.30	-0.27 – -0.10	-0.44 – -0.15	<0.001
Contexte [Neutre]	0.10	0.15	0.01 – 0.19	0.01 – 0.30	0.037
ScolCal	-0.03	-0.03	-0.19 – 0.14	-0.23 – 0.17	0.748
Empan	-0.02	-0.01	-0.30 – 0.26	-0.19 – 0.16	0.880
sFreqLex	0.04	0.05	-0.03 – 0.11	-0.04 – 0.14	0.265
nbCarac	0.03	0.13	0.01 – 0.05	0.04 – 0.23	0.005
Random Effects					
σ^2	0.27				
T00 Participant	0.09				
T00 MediaType	0.03				
ICC	0.31				
N Participant	28				
N MediaType	19				
Observations	628				
Marginal R ² / Conditional R ²	0.038 / 0.340				

Table C.58 Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 0

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.73	0.13	-3.71 – -1.74	-0.15 – 0.41	<0.001
Valence [Plus]	-0.01	-0.02	-0.13 – 0.11	-0.22 – 0.18	0.849
Contexte [Neutre]	-0.05	-0.09	-0.16 – 0.06	-0.27 – 0.10	0.370
ScolCal	-0.03	-0.04	-0.17 – 0.11	-0.23 – 0.15	0.684
Empan	0.01	0.01	-0.23 – 0.26	-0.16 – 0.18	0.910
sFreqLex	0.11	0.05	-0.22 – 0.43	-0.09 – 0.18	0.519
nbCarac	0.04	0.11	-0.01 – 0.09	-0.02 – 0.25	0.103
Random Effects					
σ^2	0.25				
T00 Participant	0.06				
T00 MediaType	0.03				
ICC	0.27				
N Participant	28				
N MediaType	19				
Observations	396				
Marginal R ² / Conditional R ²	0.011 / 0.281				

Table C.59 Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 1

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.38	0.18	-3.44 – -1.33	-0.13 – 0.49	<0.001
Valence [Plus]	-0.17	-0.30	-0.30 – -0.04	-0.53 – -0.06	0.013
Contexte [Neutre]	0.10	0.18	-0.03 – 0.23	-0.05 – 0.41	0.129
ScolCal	-0.06	-0.07	-0.22 – 0.10	-0.28 – 0.13	0.482
Empan	-0.10	-0.07	-0.37 – 0.17	-0.27 – 0.12	0.467
sFreqLex	0.06	0.06	-0.08 – 0.20	-0.09 – 0.22	0.424
nbCarac	0.06	0.18	0.01 – 0.12	0.03 – 0.33	0.022
Random Effects					
σ^2	0.21				
T00 Participant	0.06				
T00 MediaType	0.04				
ICC	0.33				
N Participant	28				
N MediaType	19				
Observations	240				
Marginal R ² / Conditional R ²	0.061 / 0.371				

Table C.60 Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 2

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.14	0.43	-3.10 – -1.18	0.16 – 0.70	<0.001
Valence [Plus]	-0.29	-0.45	-0.40 – -0.17	-0.63 – -0.28	<0.001
Contexte [Neutre]	-0.05	-0.08	-0.17 – 0.07	-0.26 – 0.10	0.395
ScolCal	-0.08	-0.09	-0.22 – 0.06	-0.25 – 0.07	0.266
Empan	-0.01	-0.01	-0.26 – 0.23	-0.16 – 0.14	0.911
sFreqLex	-0.06	-0.09	-0.12 – 0.01	-0.20 – 0.02	0.115
nbCarac	0.00	0.00	-0.02 – 0.02	-0.11 – 0.11	0.952
Random Effects					
σ^2	0.27				
T00 Participant	0.05				
T00 MediaType	0.05				
ICC	0.27				
N Participant	28				
N MediaType	19				
Observations	397				
Marginal R ² / Conditional R ²	0.071 / 0.322				

Table C.61 Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 3

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.75	0.18	-3.69 – -1.81	-0.10 – 0.47	<0.001
Valence [Plus]	-0.16	-0.28	-0.27 – -0.06	-0.46 – -0.10	0.002
Contexte [Neutre]	0.05	0.08	-0.06 – 0.15	-0.10 – 0.26	0.387
ScolCal	-0.13	-0.16	-0.28 – 0.02	-0.36 – 0.03	0.093
Empan	0.16	0.12	-0.09 – 0.40	-0.06 – 0.30	0.205
sFreqLex	-0.04	-0.10	-0.09 – 0.01	-0.23 – 0.02	0.096
nbCarac	0.00	0.01	-0.03 – 0.03	-0.11 – 0.12	0.876
Random Effects					
σ^2	0.22				
T00 Participant	0.06				
T00 MediaType	0.04				
ICC	0.32				
N Participant	27				
N MediaType	16				
Observations	377				
Marginal R ² / Conditional R ²	0.063 / 0.359				

Table C.62 Probabilité d'être le point d'arrivée d'une régression oculaire sur les mots de la zone 4

C.5.2 Par mots de la zone critique

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.03	-0.06	-3.69 – -0.37	-0.44 – 0.31	0.017
Valence [Plus]	-0.18	-0.26	-0.42 – 0.06	-0.61 – 0.09	0.138
Contexte [Neutre]	0.26	0.38	0.02 – 0.51	0.03 – 0.73	0.036
ScolCal	0.05	0.05	-0.17 – 0.27	-0.19 – 0.29	0.673
Empan	-0.25	-0.13	-0.67 – 0.17	-0.35 – 0.09	0.241
sFreqLex	0.05	0.07	-0.12 – 0.22	-0.17 – 0.31	0.571
nbCarac	0.06	0.30	0.01 – 0.12	0.04 – 0.57	0.025
Random Effects					
σ^2	0.33				
T00 Participant	0.10				
T00 MediaType	0.02				
ICC	0.27				
N Participant	28				
N MediaType	13				
Observations	115				
Marginal R ² / Conditional R ²	0.126 / 0.360				

Table C.63 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais -3

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.32	0.13	-3.60 – -1.04	-0.20 – 0.45	<0.001
Valence [Plus]	-0.13	-0.22	-0.30 – 0.04	-0.51 – 0.07	0.131
Contexte [Neutre]	0.08	0.14	-0.09 – 0.26	-0.15 – 0.43	0.346
ScolCal	-0.11	-0.13	-0.28 – 0.07	-0.36 – 0.09	0.244
Empan	-0.05	-0.03	-0.36 – 0.27	-0.23 – 0.17	0.765
sFreqLex	-0.02	-0.01	-0.56 – 0.52	-0.33 – 0.31	0.941
nbCarac	0.05	0.17	-0.04 – 0.13	-0.14 – 0.48	0.283
Random Effects					
σ^2	0.27				
T00 Participant	0.07				
T00 MediaType	0.02				
ICC	0.25				
N Participant	28				
N MediaType	19				
Observations	162				
Marginal R ² / Conditional R ²	0.068 / 0.304				

Table C.64 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais -2

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-3.78	0.15	-7.01 – -0.56	-0.16 – 0.47	0.022
Valence [Plus]	-0.21	-0.36	-0.39 – -0.03	-0.66 – -0.06	0.020
Contexte [Neutre]	0.09	0.16	-0.09 – 0.28	-0.16 – 0.47	0.321
ScolCal	-0.10	-0.13	-0.30 – 0.10	-0.40 – 0.13	0.319
Empan	0.15	0.09	-0.24 – 0.55	-0.14 – 0.32	0.442
sFreqLex	-0.74	-0.02	-6.47 – 4.99	-0.18 – 0.14	0.799
nbCarac	0.07	0.26	0.02 – 0.11	0.10 – 0.42	0.002
Random Effects					
σ^2	0.22				
T00 Participant	0.10				
T00 MediaType	0.00				
ICC	0.31				
N Participant	27				
N MediaType	19				
Observations	132				
Marginal R ² / Conditional R ²	0.127 / 0.398				

Table C.65 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais -1

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-3.08	-0.01	-4.06 – -2.10	-0.37 – 0.35	<0.001
Valence [Plus]	0.01	0.02	-0.16 – 0.19	-0.29 – 0.33	0.898
Contexte [Neutre]	0.00	0.00	-0.18 – 0.18	-0.32 – 0.32	0.986
ScolCal	-0.02	-0.03	-0.15 – 0.11	-0.22 – 0.16	0.783
Empan	0.11	0.08	-0.15 – 0.36	-0.11 – 0.26	0.412
Random Effects					
σ^2	0.25				
T00 Participant	0.02				
T00 MediaType	0.05				
ICC	0.22				
N Participant	28				
N MediaType	19				
Observations	145				
Marginal R ² / Conditional R ²	0.006 / 0.224				

Table C.66 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 0

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.13	-0.09	-3.42 – -0.83	-0.75 – 0.57	0.001
Valence [Plus]	0.11	0.20	-0.43 – 0.65	-0.77 – 1.17	0.681
Contexte [Neutre]	-0.04	-0.07	-0.20 – 0.13	-0.37 – 0.23	0.660
ScolCal	-0.01	-0.01	-0.19 – 0.18	-0.24 – 0.23	0.946
Empan	-0.14	-0.10	-0.46 – 0.17	-0.32 – 0.12	0.373
sFreqLex	-0.53	-0.23	-1.61 – 0.56	-0.70 – 0.24	0.338
Random Effects					
σ^2	0.22				
T00 Participant	0.07				
T00 MediaType	0.03				
ICC	0.31				
N Participant	27				
N MediaType	19				
Observations	150				
Marginal R ² / Conditional R ²	0.028 / 0.325				

Table C.67 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 1

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.63	0.11	-6.14 – 0.87	-0.27 – 0.49	0.139
Valence [Plus]	0.08	0.14	-0.14 – 0.31	-0.23 – 0.50	0.457
Contexte [Neutre]	-0.15	-0.25	-0.38 – 0.07	-0.62 – 0.12	0.179
ScolCal	-0.09	-0.10	-0.30 – 0.13	-0.36 – 0.15	0.425
Empan	0.08	0.05	-0.29 – 0.44	-0.19 – 0.30	0.683
sFreqLex	1.59	0.05	-5.03 – 8.21	-0.14 – 0.24	0.635
nbCarac	0.11	0.39	0.05 – 0.16	0.19 – 0.58	<0.001
Random Effects					
σ^2	0.26				
T00 Participant	0.08				
T00 MediaType	0.00				
N Participant	26				
N MediaType	19				
Observations	101				
Marginal R ² / Conditional R ²	0.190 / NA				

Table C.68 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 2

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.85	0.18	-4.47 – -1.23	-0.30 – 0.66	0.001
Valence [Plus]	-0.11	-0.21	-0.33 – 0.10	-0.62 – 0.19	0.299
Contexte [Neutre]	0.02	0.03	-0.19 – 0.23	-0.35 – 0.42	0.864
ScolCal	-0.04	-0.06	-0.28 – 0.20	-0.39 – 0.28	0.727
Empan	0.04	0.03	-0.39 – 0.47	-0.25 – 0.30	0.850
Random Effects					
σ^2	0.16				
T00 Participant	0.10				
T00 MediaType	0.04				
ICC	0.47				
N Participant	25				
N MediaType	19				
Observations	81				
Marginal R ² / Conditional R ²	0.015 / 0.482				

Table C.69 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 3

<i>Predictors</i>	logprobRegEnt				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.06	0.16	-3.83 – -0.28	-0.22 – 0.55	0.023
Valence [Plus]	-0.19	-0.33	-0.40 – 0.01	-0.68 – 0.02	0.066
Contexte [Neutre]	0.04	0.07	-0.17 – 0.26	-0.29 – 0.43	0.705
ScolCal	-0.12	-0.15	-0.29 – 0.05	-0.36 – 0.07	0.174
Empan	-0.01	-0.01	-0.30 – 0.27	-0.22 – 0.20	0.929
sFreqLex	-0.02	-0.01	-0.47 – 0.44	-0.39 – 0.36	0.941
nbCarac	-0.05	-0.06	-0.35 – 0.26	-0.45 – 0.34	0.769
Random Effects					
σ^2	0.26				
T00 Participant	0.03				
T00 MediaType	0.07				
ICC	0.27				
N Participant	27				
N MediaType	18				
Observations	114				
Marginal R ² / Conditional R ²	0.050 / 0.303				

Table C.70 Probabilité d'être le point d'arrivée d'une régression oculaire sur le mot à l'index-mais 4

C.6 Tableaux pour la probabilité de régression causée

C.6.1 Par zones

<i>Predictors</i>	logprobRegCause				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.87	0.32	-4.06 – -1.68	0.02 – 0.62	<0.001
Valence [Plus]	-0.13	-0.23	-0.26 – -0.01	-0.44 – -0.02	0.036
Contexte [Neutre]	-0.14	-0.24	-0.27 – -0.01	-0.46 – -0.01	0.037
ScolCal	-0.11	-0.14	-0.29 – 0.07	-0.36 – 0.09	0.233
Empan	0.03	0.02	-0.27 – 0.33	-0.19 – 0.23	0.837
sFreqLex	0.09	0.21	0.03 – 0.15	0.07 – 0.35	0.004
nbCarac	0.08	0.32	0.04 – 0.11	0.18 – 0.47	<0.001
Random Effects					
σ^2	0.19				
T00 Participant	0.09				
T00 MediaType	0.02				
ICC	0.38				
N Participant	27				
N MediaType	15				
Observations	252				
Marginal R ² / Conditional R ²	0.124 / 0.453				

Table C.71 Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone -1

logprobRegCause					
<i>Predictors</i>	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.37	0.19	-3.28 – -1.47	-0.05 – 0.43	<0.001
Valence [Plus]	-0.19	-0.36	-0.29 – -0.10	-0.52 – -0.19	<0.001
Contexte [Neutre]	0.07	0.12	-0.03 – 0.16	-0.05 – 0.30	0.168
ScolCal	-0.08	-0.12	-0.21 – 0.05	-0.30 – 0.07	0.213
Empan	-0.07	-0.05	-0.30 – 0.16	-0.21 – 0.11	0.540
sFreqLex	0.04	0.05	-0.04 – 0.13	-0.05 – 0.15	0.333
nbCarac	0.04	0.22	0.02 – 0.06	0.11 – 0.33	<0.001
Random Effects					
σ^2	0.22				
T00 Participant	0.05				
T00 MediaType	0.01				
ICC	0.23				
N Participant	28				
N MediaType	19				
Observations	451				
Marginal R ² / Conditional R ²	0.094 / 0.305				

Table C.72 Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 0

logprobRegCause					
<i>Predictors</i>	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.73	0.23	-3.75 – -1.71	-0.06 – 0.51	<0.001
Valence [Plus]	-0.11	-0.21	-0.23 – 0.01	-0.42 – 0.01	0.061
Contexte [Neutre]	0.02	0.03	-0.09 – 0.13	-0.17 – 0.23	0.776
ScolCal	0.01	0.01	-0.14 – 0.15	-0.20 – 0.21	0.946
Empan	-0.03	-0.02	-0.29 – 0.23	-0.20 – 0.16	0.827
sFreqLex	-0.00	-0.00	-0.30 – 0.29	-0.14 – 0.13	0.982
nbCarac	0.05	0.15	0.00 – 0.09	0.01 – 0.30	0.038
Random Effects					
σ^2	0.22				
T00 Participant	0.06				
T00 MediaType	0.03				
ICC	0.30				
N Participant	28				
N MediaType	19				
Observations	336				
Marginal R ² / Conditional R ²	0.033 / 0.319				

Table C.73 Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 1

logprobRegCause					
<i>Predictors</i>	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-3.03	0.11	-4.10 – -1.97	-0.19 – 0.42	<0.001
Valence [Plus]	-0.12	-0.23	-0.26 – 0.01	-0.47 – 0.02	0.066
Contexte [Neutre]	0.13	0.23	-0.00 – 0.26	-0.00 – 0.47	0.052
ScolCal	-0.07	-0.09	-0.22 – 0.09	-0.30 – 0.12	0.413
Empan	0.01	0.01	-0.26 – 0.28	-0.18 – 0.20	0.925
sFreqLex	0.07	0.09	-0.05 – 0.19	-0.06 – 0.23	0.228
nbCarac	0.10	0.33	0.06 – 0.15	0.18 – 0.48	<0.001
Random Effects					
σ^2	0.20				
T00 Participant	0.06				
T00 MediaType	0.02				
ICC	0.29				
N Participant	28				
N MediaType	19				
Observations	231				
Marginal R ² / Conditional R ²	0.115 / 0.370				

Table C.74 Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 2

logprobRegCause					
<i>Predictors</i>	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-2.57	0.35	-3.43 – -1.72	0.11 – 0.59	<0.001
Valence [Plus]	-0.30	-0.47	-0.41 – -0.19	-0.64 – -0.29	<0.001
Contexte [Neutre]	0.03	0.05	-0.08 – 0.14	-0.13 – 0.23	0.589
ScolCal	-0.10	-0.12	-0.23 – 0.02	-0.26 – 0.03	0.105
Empan	0.07	0.04	-0.15 – 0.28	-0.09 – 0.18	0.548
sFreqLex	-0.06	-0.09	-0.13 – 0.01	-0.19 – 0.02	0.108
nbCarac	0.02	0.15	0.00 – 0.05	0.03 – 0.27	0.015
Random Effects					
σ^2	0.27				
T00 Participant	0.04				
T00 MediaType	0.04				
ICC	0.23				
N Participant	28				
N MediaType	19				
Observations	410				
Marginal R ² / Conditional R ²	0.124 / 0.325				

Table C.75 Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 3

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-2.85	0.24	-3.81 – -1.89	-0.03 – 0.51	<0.001
Valence [Plus]	-0.16	-0.26	-0.24 – -0.08	-0.39 – -0.14	<0.001
Contexte [Neutre]	-0.01	-0.02	-0.09 – 0.07	-0.14 – 0.11	0.813
ScolCal	-0.06	-0.07	-0.21 – 0.09	-0.26 – 0.11	0.421
Empan	0.12	0.09	-0.12 – 0.37	-0.09 – 0.27	0.328
sFreqLex	0.05	0.12	0.01 – 0.09	0.02 – 0.21	0.016
nbCarac	0.03	0.13	0.01 – 0.05	0.04 – 0.22	0.005
Random Effects					
σ^2	0.23				
T00 Participant	0.07				
T00 MediaType	0.05				
ICC	0.34				
N Participant	28				
N MediaType	16				
Observations	671				
Marginal R ² / Conditional R ²	0.038 / 0.369				

Table C.76 Probabilité d'être le point de départ d'une régression oculaire sur les mots de la zone 4

C.6.2 Par mots de la zone critique

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-1.68	0.05	-3.44 – 0.08	-0.44 – 0.54	0.061
Valence [Plus]	-0.26	-0.43	-0.55 – 0.03	-0.91 – 0.04	0.075
Contexte [Neutre]	0.26	0.43	-0.03 – 0.55	-0.05 – 0.91	0.079
ScolCal	0.04	0.04	-0.19 – 0.27	-0.23 – 0.32	0.752
Empan	-0.28	-0.17	-0.72 – 0.16	-0.44 – 0.10	0.204
sFreqLex	0.13	0.12	-0.19 – 0.45	-0.17 – 0.40	0.410
Random Effects					
σ^2	0.26				
T00 Participant	0.07				
T00 MediaType	0.03				
ICC	0.28				
N Participant	27				
N MediaType	11				
Observations	66				
Marginal R ² / Conditional R ²	0.115 / 0.359				

Table C.77 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais -3

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-1.75	0.05	-3.07 – -0.43	-0.31 – 0.41	0.010
Valence [Plus]	-0.06	-0.12	-0.26 – -0.13	-0.48 – 0.24	0.513
Contexte [Neutre]	0.02	0.03	-0.18 – 0.21	-0.33 – 0.39	0.866
ScolCal	-0.05	-0.08	-0.21 – 0.10	-0.31 – 0.15	0.495
Empan	-0.23	-0.14	-0.58 – 0.12	-0.35 – 0.07	0.191
sFreqLex	-0.23	-0.15	-0.51 – 0.04	-0.33 – 0.03	0.099
Random Effects					
σ^2	0.26				
T00 Participant	0.03				
T00 MediaType	0.00				
ICC	0.11				
N Participant	26				
N MediaType	19				
Observations	119				
Marginal R ² / Conditional R ²	0.060 / 0.160				

Table C.78 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais -2

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-0.40	0.08	-3.54 – 2.74	-0.20 – 0.37	0.800
Valence [Plus]	-0.20	-0.35	-0.36 – -0.04	-0.64 – -0.07	0.016
Contexte [Neutre]	0.13	0.24	-0.03 – 0.30	-0.06 – 0.54	0.117
ScolCal	-0.11	-0.15	-0.24 – 0.02	-0.33 – 0.03	0.108
Empan	0.01	0.01	-0.23 – 0.24	-0.17 – 0.18	0.955
sFreqLex	5.38	0.16	-0.73 – 11.50	-0.02 – 0.33	0.084
nbCarac	0.10	0.38	0.05 – 0.14	0.20 – 0.55	<0.001
Random Effects					
σ^2	0.23				
T00 Participant	0.02				
T00 MediaType	0.01				
ICC	0.12				
N Participant	28				
N MediaType	19				
Observations	155				
Marginal R ² / Conditional R ²	0.172 / 0.272				

Table C.79 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais -1

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-2.29	0.29	-4.78 – 0.21	-0.19 – 0.77	0.071
Valence [Plus]	-0.12	-0.22	-0.36 – 0.12	-0.66 – 0.21	0.305
Contexte [Neutre]	-0.15	-0.26	-0.40 – 0.11	-0.73 – 0.21	0.267
ScolCal	-0.01	-0.01	-0.25 – 0.24	-0.37 – 0.35	0.951
Empan	-0.10	-0.06	-0.73 – 0.52	-0.40 – 0.28	0.741
Random Effects					
σ^2	0.20				
T00 Participant	0.12				
T00 MediaType	0.01				
ICC	0.37				
N Participant	22				
N MediaType	19				
Observations	65				
Marginal R ² / Conditional R ²	0.033 / 0.396				

Table C.80 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 0

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-2.67	0.29	-3.81 – -1.53	-0.26 – 0.85	<0.001
Valence [Plus]	-0.21	-0.43	-0.66 – 0.23	-1.33 – 0.47	0.348
Contexte [Neutre]	-0.01	-0.02	-0.17 – 0.16	-0.35 – 0.31	0.917
ScolCal	0.04	0.06	-0.12 – 0.19	-0.19 – 0.30	0.631
Empan	0.02	0.01	-0.27 – 0.30	-0.22 – 0.24	0.911
sFreqLex	0.05	0.03	-0.86 – 0.96	-0.42 – 0.47	0.910
Random Effects					
σ^2	0.19				
T00 Participant	0.05				
T00 MediaType	0.02				
ICC	0.25				
N Participant	28				
N MediaType	19				
Observations	133				
Marginal R ² / Conditional R ²	0.035 / 0.273				

Table C.81 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 1

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-1.49	0.04	-5.64 – 2.66	-0.33 – 0.41	0.479
Valence [Plus]	-0.07	-0.13	-0.28 – 0.13	-0.47 – 0.22	0.473
Contexte [Neutre]	0.09	0.15	-0.10 – 0.28	-0.16 – 0.47	0.340
ScolCal	-0.03	-0.04	-0.23 – 0.16	-0.29 – 0.21	0.742
Empan	-0.07	-0.05	-0.40 – 0.26	-0.28 – 0.18	0.674
sFreqLex	2.49	0.06	-5.26 – 10.25	-0.13 – 0.26	0.526
nbCarac	0.07	0.24	0.01 – 0.12	0.03 – 0.44	0.022
Random Effects					
σ^2	0.26				
T00 Participant	0.08				
T00 MediaType	0.02				
ICC	0.26				
N Participant	26				
N MediaType	19				
Observations	138				
Marginal R ² / Conditional R ²	0.057 / 0.305				

Table C.82 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 2

logprobRegCause					
Predictors	Estimates	std. Beta	CI	standardized CI	p
(Intercept)	-2.68	0.18	-4.44 – -0.93	-0.35 – 0.72	0.003
Valence [Plus]	-0.16	-0.29	-0.44 – 0.12	-0.81 – 0.23	0.268
Contexte [Neutre]	0.08	0.15	-0.17 – 0.33	-0.32 – 0.61	0.531
ScolCal	0.01	0.02	-0.26 – 0.29	-0.33 – 0.36	0.916
Empan	-0.06	-0.03	-0.52 – 0.41	-0.33 – 0.26	0.815
Random Effects					
σ^2	0.18				
T00 Participant	0.10				
T00 MediaType	0.02				
ICC	0.40				
N Participant	23				
N MediaType	17				
Observations	64				
Marginal R ² / Conditional R ²	0.021 / 0.416				

Table C.83 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 3

<i>Predictors</i>	logprobRegCause				
	<i>Estimates</i>	<i>std. Beta</i>	<i>CI</i>	<i>standardized CI</i>	<i>p</i>
(Intercept)	-3.22	-0.12	-4.63 – -1.82	-0.51 – 0.27	<0.001
Valence [Plus]	-0.08	-0.16	-0.26 – 0.09	-0.51 – 0.18	0.346
Contexte [Neutre]	0.21	0.41	0.03 – 0.38	0.07 – 0.75	0.020
ScolCal	-0.09	-0.14	-0.25 – 0.07	-0.37 – 0.10	0.254
Empan	0.07	0.05	-0.23 – 0.36	-0.16 – 0.26	0.656
sFreqLex	0.08	0.09	-0.20 – 0.36	-0.23 – 0.41	0.576
nbCarac	0.10	0.18	-0.08 – 0.28	-0.16 – 0.52	0.288
Random Effects					
σ^2	0.20				
T00 Participant	0.04				
T00 MediaType	0.02				
ICC	0.23				
N Participant	28				
N MediaType	19				
Observations	122				
Marginal R ² / Conditional R ²	0.077 / 0.290				

Table C.84 Probabilité d'être le point de départ d'une régression oculaire sur le mot à l'index-mais 4

BIBLIOGRAPHIE

- Andersson, R., Nyström, M. et Holmqvist, K. (2010). Sampling frequency and eye-tracking measures : how speed affects durations, latencies, and more. *Journal of Eye Movement Research*, 3(3).
- Anscombre, J.-C. et Ducrot, O. (1977). Deux mais en français ? *Lingua*, 43(1), 23-40.
- Anscombre, J.-C. et Ducrot, O. (1983). *L'argumentation dans la langue*. Editions Mardaga.
- Baddeley, A. (2012). Working memory : Theories, models, and controversies. *Annual review of psychology*, 63(1), 1-29.
- Bates, D., Mächler, M., Bolker, B. et Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1), 1-48. <http://dx.doi.org/10.18637/jss.v067.i01>
- Beaver, D. et Clark, B. (2003). Always and only : Why not all focus-sensitive operators are alike. *Natural language semantics*, 11(4), 323-362.
- Becker, R., Chambers, J. et Wilks, A. (1988). The new s language. wadsworth & brooks/cole. *Computer Science Series*, Pacific Grove, CA.
- BENSOLTANA, D. Corrélation entre fonctionnement oculomoteur, processus visuo-cognitif et apprentissage. *DU COMPORTEMENT*, p. 50.
- Binder, K. S., Pollatsek, A. et Rayner, K. (1999). Extraction of information to the left of the fixated word in reading. *Journal of Experimental Psychology : Human Perception and Performance*, 25(4), 1162.
- Blakemore, D. (1989). Denial and contrast : A relevance theoretic analysis of " but ". *Linguistics and philosophy*, 15-37.
- Blakemore, D. (2002). *Relevance and linguistic meaning : The semantics and pragmatics of discourse markers*, volume 99. Cambridge university press.
- Blanchard, H. E., Pollatsek, A. et Rayner, K. (1989). The acquisition of parafoveal word information in reading. *Perception & Psychophysics*, 46(1), 85-94.
- Bridgeman, B., Hendry, D. et Stark, L. (1975). Failure to detect displacement of the visual world during saccadic eye movements. *Vision research*, 15(6), 719-722.
- Büring, D. (2019). Focus, questions and givenness. In *Questions in discourse* 6-44. Brill.
- Carston, R. et Uchida, S. (1998). *Relevance theory : Applications and implications*, volume 37. John Benjamins Publishing.
- Cook, A. E. et Wei, W. (2019). What can eye movements tell us about higher level comprehension ? *Vision*, 3(3), 45.
- Daneman, M. et Carpenter, P. A. (1980). Individual differences in working memory and reading. *Journal of verbal learning and verbal behavior*, 19(4), 450-466.

- Desmette, D., Hupet, M., Schelstraete, M.-A. et Van der Linden, M. (1995). Adaptation en langue française du «reading span test» de daneman et carpenter (1980). *L'année Psychologique*, 95(3), 459–482.
- Ferrand, L. (2001). Normes d'associations verbales pour 260 mots «abstraits». *L'Année psychologique*, 101(4), 683–721.
- Ferrand, L. et Alario, F.-X. (1998). Normes d'associations verbales pour 366 noms d'objets concrets. *L'Année psychologique*, 98(4), 659–709.
- Ferrand, L., Matos, R., New, B. et Pallier, C. (2001). Une base de données lexicales du français contemporain sur internet : Lexique. *L'Année Psychologique*, 101, 447–462.
- Fisher, D. F. et Shebilske, W. L. (1985). There is more that meets the eye than the eyemind assumption. *Eye movements and human information processing*, 149–158.
- Givón, T. (1987). Beyond foreground and background. *Coherence and grounding in discourse*, 11, 175–188.
- Goldberg, J. H. et Kotval, X. P. (1999). Computer interface evaluation using eye movements : methods and constructs. *International journal of industrial ergonomics*, 24(6), 631–645.
- Gordon, P., Lowder, M. et Hoedemaker, R. (2015). Cognitive-linguistic processes and aging.
- Govignon, G. (2021). stop words french.
https://gitlab.binets.fr/gabriel.govignon/projet-inf583/-/blob/91ca09bcee26a7dfd054875a6223d7d3e3e483fe/twitterapp/stop_words_french.txt.
 Accédé le 14 Avril 2022.
- Holmqvist, K., Nyström, M., Andersson, R., Dewhurst, R., Jarodzka, H. et Van de Weijer, J. (2011). *Eye tracking : A comprehensive guide to methods and measures*. oup Oxford.
- Honnibal, M. et Montani, I. (2017). spaCy 2 : Natural language understanding with Bloom embeddings, convolutional neural networks and incremental parsing. To appear.
- Inhoff, A. W. (1984). Two stages of word processing during eye fixations in the reading of prose. *Journal of verbal learning and verbal behavior*, 23(5), 612–624.
- Inhoff, A. W. et Rayner, K. (1980). Parafoveal word perception : A case against semantic preprocessing. *Perception & Psychophysics*, 27(5), 457–464.
- Ishida, T. et Ikeda, M. (1989). Temporal properties of information extraction in reading studied by a text-mask replacement technique. *JOSA A*, 6(10), 1624–1632.
- Jasinskaja, K., Zeevat, H. et al. (2008). Explaining additive, adversative and contrast marking in russian and english. *Revue de Sémantique et Pragmatique*, 24(1), 65–91.
- Just, M. A. et Carpenter, P. A. (1980). A theory of reading : from eye fixations to comprehension. *Psychological review*, 87(4), 329.
- Kennedy, C. et McNally, L. (2005). Scale structure, degree modification, and the semantics of gradable predicates. *Language*, 345–381.

- Kim, Y.-S. G., Petscher, Y. et Vorstius, C. (2021). The relations of online reading processes (eye movements) with working memory, emergent literacy skills, and reading proficiency. *Scientific Studies of Reading*, 25(4), 351–369.
- Konieczny, L., Hemforth, B., Scheepers, C. et Strube, G. (1997). The role of lexical heads in parsing : Evidence from german. *Language and cognitive processes*, 12(2-3), 307–348.
- Kuperman, V. et Van Dyke, J. A. (2011). Effects of individual differences in verbal skills on eye-movement patterns during sentence reading. *Journal of memory and language*, 65(1), 42–73.
- McConkie, G. W., Underwood, N. R., Zola, D. et Wolverton, G. (1985). Some temporal characteristics of processing during reading. *Journal of Experimental Psychology : Human Perception and Performance*, 11(2), 168.
- Merin, A. (1999). Information, relevance, and social decisionmaking : Some principles and results of decision-theoretic semantics. *Logic, Language, and computation*, 2, 179–221.
- Meyer, S. (2008). *Twilight*. Twilight Saga. Little, Brown Book Group Limited. Récupéré de <https://books.google.ca/books?id=3WVTPwAACAAJ>
- Nattkemper, D. et Prinz, W. (1987). Saccade amplitude determines fixation duration : Evidence from continuous search. In *Eye Movements from Physiology to Cognition* 285–292. Elsevier.
- New, B., Pallier, C., Brysbaert, M. et Ferrand, L. (2004). Lexique 2 : A new french lexical database. *Behavior Research Methods, Instruments, & Computers*, 36(3), 516–524.
- Nyquist, H. (1928). Thermal agitation of electric charge in conductors. *Physical review*, 32(1), 110.
- Pollatsek, A., Rayner, K. et Balota, D. A. (1986). Inferences about eye movement control from the perceptual span in reading. *Perception & Psychophysics*, 40, 123–130.
- Purves, D., Augustine, G. J., Fitzpatrick, D., Katz, L. C., LaMantia, A.-S., McNamara, J. O., Williams, S. M. et al. (2001). Types of eye movements and their functions. *Neuroscience*, 20, 361–390.
- Rayner, K. (1975). The perceptual span and peripheral cues in reading. *Cognitive psychology*, 7(1), 65–81.
- Rayner, K. (1986). Eye movements and the perceptual span in beginning and skilled readers. *Journal of experimental child psychology*, 41(2), 211–236.
- Rayner, K. (1998). Eye movements in reading and information processing : 20 years of research. *Psychological bulletin*, 124(3), 372.
- Rayner, K. et Fischer, M. H. (1996). Mindless reading revisited : Eye movements during reading and scanning are different. *Perception & psychophysics*, 58(5), 734–747.
- Rayner, K., Reichle, E. D., Stroud, M. J., Williams, C. C. et Pollatsek, A. (2006). The effect of word frequency, word predictability, and font difficulty on the eye movements of young and older readers. *Psychology and aging*, 21(3), 448.
- Rayner, K., Well, A. D. et Pollatsek, A. (1980). Asymmetry of the effective visual field in reading. *Perception*

- & *Psychophysics*, 27(6), 537–544.
- Rayner, K., Well, A. D., Pollatsek, A. et Bertera, J. H. (1982). The availability of useful information to the right of fixation in reading. *Perception & Psychophysics*, 31(6), 537–550.
- Sæbø, K. J. (2003). Presupposition and contrast : German 'aber' as a topic particle. Dans *Proceedings of Sinn und Bedeutung*, volume 7, 257–271.
- Schmauder, A. R. (1992). Eye movements and reading processes. In *Eye movements and visual cognition : Scene perception and reading* 369–378. Springer.
- Shannon, C. E. (2006). Communication in the presence of noise. *Proceedings of the IRE*, 37(1), 10–21.
- Spector, R. H. (1990). The pupils. *Clinical Methods : The History, Physical, and Laboratory Examinations*. 3rd edition.
- Sperber, D. et Wilson, D. (1986). *Relevance : Communication and cognition*, volume 142. Harvard University Press Cambridge, MA.
- TobiiAB. (2023). *Tobii Pro Lab*. Danderyd, Stockholm. Récupéré de <http://www.tobii.com/>
- Toivo, W. et Scheepers, C. (2019). Pupillary responses to affective words in bilinguals' first versus second language. *Plos one*, 14(4), e0210450.
- Traxler, M. J., Long, D. L., Tooley, K. M., Johns, C. L., Zirnstein, M. et Jonathan, E. (2012). Individual differences in eye-movements during reading : Working memory and speed-of-processing effects. *Journal of eye movement research*, 5(1).
- Trueswell, J. C., Tanenhaus, M. K. et Garnsey, S. M. (1994). Semantic influences on parsing : Use of thematic role information in syntactic ambiguity resolution. *Journal of memory and language*, 33(3), 285–318.
- Umbach, C. (2005). Contrast and information structure : A focus-based analysis of but.
- Von Stutterheim, C. et Klein, W. (1989). Referential movement in descriptive and narrative discourse. In *North-Holland Linguistic Series : Linguistic Variations*, volume 54 39–76. Elsevier.
- Wang, J. T.-y. (2011). Pupil dilation and eye tracking. *A handbook of process tracing methods for decision research : A critical review and user's guide*, 185–204.
- Wilbur, W. J. et Sirotkin, K. (1992). The automatic identification of stop words. *Journal of information science*, 18(1), 45–55.
- Winterstein, G. (2012). What but-sentences argue for : An argumentative analysis of but. *Lingua*, 122(15), 1864–1885.
- Winterstein, G., Ellsiepen, E., Jayes, J. et Hemforth, B. (2014). Effects of context on the processing of adversative and comparative constructions. Dans *CUNY*.
- Yan, M., Richter, E. M., Shu, H. et Kliegl, R. (2009). Readers of chinese extract semantic information from

parafoveal words. *Psychonomic bulletin & review*, 16(3), 561-566.

Zihl, J. (1993). Reading disorders, nonaphasic, and their treatment. *Neuroscience Year*, 143-144 (1993).