

Résumé

L'enchaînement, et plus particulièrement la liaison, représente un phénomène fréquent en français oral et peut mener certaines fois à des ambiguïtés lexicales. Les modèles d'accès au lexique actuels sont basés sur l'anglais et ne cherchent donc pas à rendre compte de tels phénomènes. L'étude exploratoire que nous avons entreprise avait pour objectif de déterminer le "taux d'ambiguïté" de suites de mots enchaînés avec liaison (ex. "petit ami" qui peut être interprété comme "petit tamis") et sans liaison (ex. "chaque ours" qui peut être interprété comme "chaque course"). Les résultats obtenus confirment que l'enchaînement, et en particulier la liaison, peut créer une réelle ambiguïté pendant la perception de la parole.

1. Introduction

La parole continue est affectée par des phénomènes d'enchaînement qui se traduisent par des processus de sandhi externe ou ajustements phonologiques entre les mots de la chaîne parlée. Les syllabes ouvertes prédominent en français et les frontières syllabiques ont presque toujours la priorité sur les frontières lexicales (ex. "cette arme" : /sɛ-tarm/ et non /set-arm/). L'enchaînement, qui est l'un de ces processus phonologiques, comprend à la fois l'enchaînement avec liaison et l'enchaînement sans liaison. Ces deux types d'enchaînement consistent à lier une consonne en finale de mot avec la voyelle ou la consonne du début du mot suivant. La différence entre les deux types d'enchaînement est la suivante : dans le cas de l'enchaînement sans liaison, la consonne en finale de mot est toujours prononcée, même devant une autre consonne ou en isolation, alors que dans le cas de l'enchaînement avec liaison, la consonne en finale de mot n'est prononcée que devant une voyelle ou un h "muet".

A notre connaissance, seules deux études ont examiné l'effet de l'enchaînement sur la reconnaissance des mots. Toutes deux ont été

[◊] Cette étude a pu être entreprise et menée à bien grâce à un subside du FNRS (12-33582.92).

menées afin de rendre compte de la difficulté que les apprenants rencontrent face à l'apprentissage du français. Matter (1986) a tenté de répondre à la question suivante : est-ce que des phénomènes de sandhi comme la liaison, l'enchaînement et l'élision retardent la reconnaissance de la parole en français à la fois chez les natifs et les non-natifs ? Plusieurs expériences ont été conduites sur l'enchaînement et la liaison à l'aide d'une tâche de détection de phonèmes. Les résultats significatifs obtenus sont les suivants : seule la liaison potentielle (ex. "grand talent" par rapport à "vrai talent") retarde la détection d'un mot commençant par /t/. La liaison ("grand arbre") et l'enchaînement ("grande amie") ne retardent pas la détection d'un mot commençant par /a/. En outre, dans toutes les expériences, les sujets non natifs ont mis plus de temps pour faire la tâche de détection que les sujets natifs du français. Enfin, l'auteur met en doute la tâche utilisée, et de ce fait, ne peut pas tirer de conclusions très probantes. Bradley & Dejean de la Bâtie (1990), quant à eux, ont utilisé le même matériel que Matter et ont étudié comment deux groupes d'auditeurs (sujets de langue maternelle française et étudiants du français langue étrangère) établissaient des frontières de mots en français. Les sujets devaient détecter le phonème /t/ en début de mot. Les auteurs ont obtenu des résultats, pour les deux groupes réunis, qui sont semblables à ceux de Matter, à savoir des temps de réaction plus longs pour "grand talent" que pour "vrai talent". Ils n'en concluent rien, cependant, sur la reconnaissance des mots enchaînés.

A l'exception de ces deux études qui portent surtout sur la perception de la parole chez les apprenants, les chercheurs ne se sont guère penchés sur l'effet de l'enchaînement lors de l'accès au lexique en parole continue, que ce soit en français ou dans une autre langue. Dès lors, il nous a paru important d'explorer cette voie dans le cadre de la recherche sur l'accès au lexique. Non seulement nous voulions nous assurer que la reconnaissance de mots enchaînés peut poser un problème au niveau de la perception mais nous désirions également savoir si les difficultés sont différentes selon le type d'enchaînement. Pour ce faire, nous nous sommes intéressés à trois types d'enchaînement. Un enchaînement avec liaison du type "son œuf" (que nous représenterons dorénavant par le sigle E/L-cv) et deux types d'enchaînement sans liaison : l'enchaînement entre la consonne finale d'un mot et la voyelle du

début du mot suivant comme dans "chaque ours" (sigle : E-cv), et l'enchaînement entre la consonne finale d'un mot et la consonne initiale du mot suivant comme dans "neuf lames" (sigle : E-cc). Il est à noter que dans notre étude chaque type d'enchaînement peut conduire à une ambiguïté lexicale complète. En effet, dans les exemples ci-dessus, les suites peuvent être interprétées de deux manières différentes : "son œuf" / "son neuf", "chaque ours" / "chaque course", "neuf lames" / "neuf flammes". Ce type d'ambiguïté, où le deuxième élément de la suite est un mot possible, avec ou sans la consonne d'enchaînement, est plus rare que l'ambiguïté momentanée créée en début du deuxième mot mais ensuite résolue avant la fin de celui-ci (ex. "nain" dans "un instrument", "tas" dans "était arrivé", etc.).

Dans cette étude exploratoire qui prélude à un programme de recherche plus conséquent (voir Besson & Grosjean, 1994), nous avons mené une expérience de discrimination dans laquelle nous avons présenté à des sujets des segments de phrases contenant des suites de mots enchaînés (ex. "Il s'agit de son œuf") et leur avons demandé de nous indiquer la provenance de ces segments. Ils avaient le choix entre les deux interprétations possibles ("Il s'agit de son œuf" et "Il s'agit de son neuf"). Nous avons fait l'hypothèse que les suites de la catégorie E/L-cv seraient discriminées avec difficulté. Quant aux suites du type E-cc, elles seraient discriminées assez facilement. Par contre, nous ne savions que dire des suites E-cv, car celles-ci partagent avec les suites E/L-cv le fait qu'il s'agit d'un enchaînement entre la consonne finale d'un mot et la voyelle du mot suivant (ex. "chaque ours"). Mais elles sont également proches des suites E-cc, car il s'agit d'un enchaînement et non d'une liaison. Par conséquent, les suites E-cv pourraient se comporter d'une manière similaire aux suites E/L-cv ou aux suites E-cc ou encore représenter une catégorie intermédiaire.

2. Méthode

Sujets : Seize sujets monolingues, dont la langue première est le français, ont pris part à l'expérience.

Matériel : Vingt-quatre suites de deux mots ont été réparties en trois groupes de huit, chaque groupe correspondant aux trois catégories d'enchaînement indiquées ci-dessus : enchaînement E/L-cv, ex. "son œuf" ; enchaînement E-cv ; ex. "chaque ours" ; enchaînement E-cc, ex. "neuf lames". Dans chaque groupe nous avons apparié chaque suite avec enchaînement à la suite correspondante sans enchaînement, ex. "son œuf"

(suite AE) appariée à "son neuf" (suite SE), "chaque ours" (suite AE) à "chaque course" (suite SE), etc. Avant d'intégrer ces suites dans un contexte plus large, nous nous sommes assurés que les trois groupes de suites ne différaient pas les uns des autres au niveau des variables suivantes : la fréquence du deuxième mot de la suite (ex. "œuf", "neuf", "ours", "course") ; le point d'unicité du deuxième mot ; et la cohérence sémantique entre le premier et le deuxième mot de la suite (ex. "son œuf", "son neuf", "chaque ours", "chaque course", etc.). Les suites ont ensuite été intégrées dans un énoncé de deux phrases. La première était une phrase introductive du type, "Cela a éveillé notre intérêt", et la deuxième était une phrase qui commençait avec "Il s'agit de", qui continuait avec la suite en question et qui se terminait avec un syntagme prépositionnel. (Ce syntagme était ajouté afin de s'assurer que la suite serait dite avec une prosodie de continuation). Ainsi, la suite "son œuf" se retrouvait dans l'énoncé suivant : "Cela a éveillé notre intérêt. Il s'agit de son œuf de pigeon probablement", et "son neuf" était inséré dans : "Cela a éveillé notre intérêt. Il s'agit de son neuf de carreau probablement".

Dans ce qui suit, nous ferons référence aux parties suivantes des énoncés :

- Segment AE (segment stimulus avec enchaînement) ; ex. "Il s'agit de son œuf"
- Segment SE (segment stimulus sans enchaînement) ; ex. "Il s'agit de son neuf"
- Suite AE (suite avec enchaînement) ; ex. "son œuf"
- Suite SE (suite sans enchaînement) ; ex. "son neuf"

Une lectrice française a lu les 48 énoncés (24 avec enchaînement et 24 sans) à un débit normal. Aucune instruction n'a été donnée quant à la prononciation de la consonne d'enchaînement afin de refléter le plus possible une situation naturelle. Les énoncés ont ensuite été digitalisés à l'aide du logiciel MacAdios et les segments stimuli en ont été extraits (ex. "Il s'agit de son œuf.", "Il s'agit de son neuf.", "Il s'agit de chaque ours.", "Il s'agit de chaque course.", etc.). Ceux-ci ont ensuite été enregistrés sur deux bandes différentes, chaque bande comportant 24 segments, douze segments AE et 12 segments SE. Les éléments d'une paire (segment AE et segment SE) se trouvaient toujours sur des bandes différentes. Bien que nous n'étions intéressés que par les segments AE (qui sont à l'origine de notre étude), nous avons utilisé tous les segments (AE et SE) car il fallait que les sujets puissent entendre chaque élément d'une paire et faire un choix entre l'une ou l'autre des interprétations.

Procédure : Les sujets ont écouté les deux bandes à l'aide d'un magnétophone à cassette et d'écouteurs, à 24 heures d'intervalle. Ils devaient indiquer si le segment présenté correspondait à l'une ou à l'autre des deux versions présentées par écrit. Ils devaient également donner un degré de confiance pour chaque réponse sur une échelle allant de 1 (très peu sûr) à 10 (très sûr).

Analyse des données : Nous n'avons analysé que les résultats des segments AE pour lesquels nous avons obtenu deux mesures :

1) **Taux de discrimination :** Étant donné que pour un segment AE, le sujet avait le choix entre la version écrite avec enchaînement (réponse AE) et la version écrite sans enchaînement (réponse SE), nous avons calculé un taux de discrimination pour chaque segment AE en prenant le nombre de réponses AE et en y soustrayant le nombre de réponses SE. Par exemple, pour le segment AE, "Il s'agit d'un ancien hectare", 6 sujets ont choisi la réponse AE et 10 sujets ont choisi la réponse SE ; le taux de discrimination est donc de $6 - 10 = -4$. Pour le segment AE, "Il s'agit d'une grande anse", 16 sujets ont choisi la réponse AE et aucun n'a choisi la réponse SE ; le taux de discrimination est donc de $16 - 0 = 16$. Une valeur positive et élevée signifie que la phrase a été discriminée facilement et que les sujets ont choisi la réponse qui correspond au segment AE. Une valeur négative et élevée signifie que la phrase a également été discriminée facilement, mais que les sujets ont choisi la réponse qui correspond au segment SE. Une valeur proche de 0 signifie que la phrase a été entendue autant de fois comme AE que

comme SE, et donc qu'elle est plus ambiguë. La valeur maximale est de 16 et la valeur minimale de -16.

2) **Degré de confiance :** Pour chaque réponse proposée, le sujet a donné un degré de confiance, que nous avons analysé tel quel, sans tenir compte de son choix de segment (AE ou SE). Nous avons donc simplement calculé la moyenne de chaque segment sur 16 sujets. Plus la valeur est haute, plus les sujets sont confiants dans leur choix. La valeur maximale est de 10 et la valeur minimale de 1.

3. Résultats et discussion

La Figure 1 (ci-dessous) montre le taux de discrimination des segments AE en fonction de la catégorie d'enchaînement : enchaînement avec liaison (E/L-cv), enchaînement (sans liaison) entre la consonne finale d'un mot et la voyelle du début du mot suivant (E-cv), et enchaînement (sans liaison) entre la consonne finale d'un mot et la consonne initiale du mot suivant (E-cc).

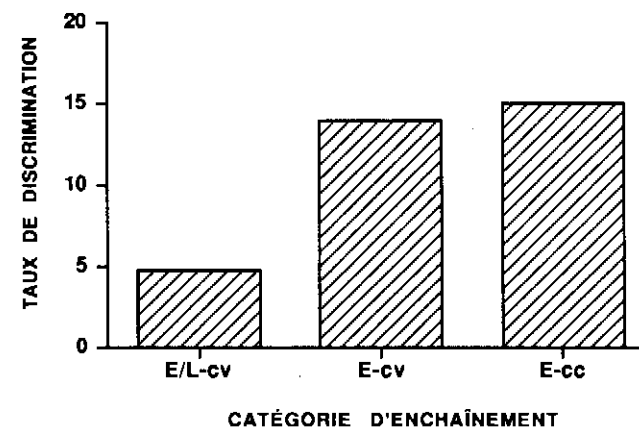


Figure 1. Taux de discrimination des segments AE en fonction de trois catégories d'enchaînement : E/L-cv, E-cv, E-cc.

Nous remarquons que les segments de la catégorie E/L-cv sont moins bien discriminés, en moyenne, que les segments des catégories E-cv et E-cc (moyennes de 4.75, 14 et 15, respectivement) et que les deux dernières catégories donnent des résultats semblables. Une analyse de variance par mot et par sujet confirme cette observation (par mot : $F(2, 21) = 26.90$,

$p < 0.001$; par sujet : $F(2, 30) = 42.04$, $p < 0.001$). Une analyse post hoc (par sujet) montre une différence significative entre la catégorie E/L-cv et la catégorie E-cv et entre la catégorie E/L-cv et la catégorie E-cc, mais pas de différence entre les catégories E-cv et E-cc.

Nous pouvons conclure que les segments qui contiennent des enchaînements avec liaison sont moins bien discriminés, et par conséquent, sont plus ambigus que les segments qui contiennent des enchaînements sans liaison.

La Figure 2 (ci-dessous) montre le degré de confiance dans les réponses données en fonction des trois catégories d'enchaînement : E/L-cv, E-cv et E-cc.

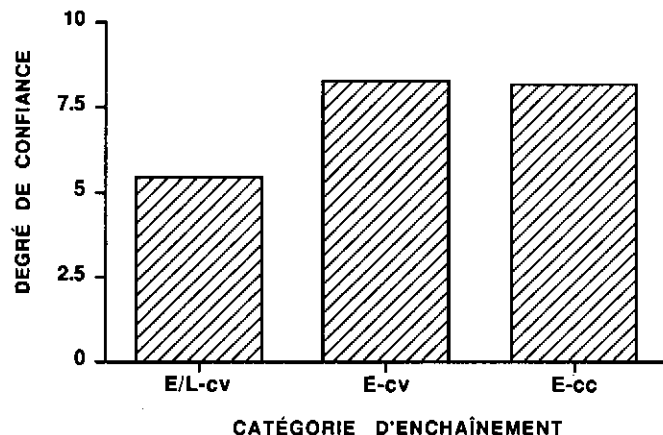


Figure 2. Degrés de confiance dans les réponses données en fonction de trois catégories d'enchaînement : E/L-cv, E-cv, E-cc.

Nous observons des effets similaires à ceux obtenus pour le taux de discrimination. Premièrement, les segments de la catégorie E/L-cv ont un degré de confiance plus bas, en moyenne, que les segments des catégories E-cv et E-cc (moyennes de 5.41, 8.29 et 8.17, respectivement) et deuxièmement les segments des deux catégories d'enchaînement sans liaison présentent des résultats presque identiques.

Une analyse de variance confirme ces observations (par mot : $F(2, 21) = 26.72$, $p < 0.001$; par sujet : $F(2, 30) = 60.03$, $p < 0.001$). Une analyse post hoc (par sujet) montre une différence significative entre la catégorie E/L-cv et la catégorie E-cv et entre la catégorie E/L-cv et la catégorie E-cc, mais pas entre la catégorie E-cv et la catégorie E-cc. Cette deuxième mesure (degré de confiance) confirme donc les résultats obtenus avec le taux de discrimination.

Nous retenons de cette étude que l'enchaînement avec liaison (E/L-cv) est potentiellement plus ambigu que l'enchaînement sans liaison, que ce dernier soit l'enchaînement entre la consonne finale d'un mot et la voyelle du début du mot suivant (E-cv) ou celui entre la consonne finale d'un mot et la consonne initiale du mot suivant (E-cc). Etant donné ces différences dans l'ambiguïté des suites, nous pouvons émettre l'hypothèse que l'accès au lexique se fera avec plus de difficulté, et donc plus lentement, pour les suites ambiguës du type E/L-cv que pour les suites des deux autres catégories.

4. Discussion générale

Bien qu'il faudra confirmer ces résultats avec une tâche de reconnaissance des mots en temps réel (voir Besson & Grosjean, 1994), nous pouvons déjà envisager la manière dont les modèles actuels d'accès au lexique pourraient rendre compte de la reconnaissance des mots enchaînés. Le modèle de la cohorte (Marslen-Wilson, 1987) aurait quelques difficultés à le faire car l'activation des candidats se fait à partir du début du mot (même si cela est moins explicite que dans la première version du modèle, Marslen-Wilson & Welsh, 1978) et la reconnaissance de plusieurs mots demeure une opération séquentielle. Or, ce qui pourrait se passer au début du deuxième mot dans une suite comme /pətitami/ n'est pas clair. Nous voyons deux possibilités qui semblent s'exclure mutuellement. Dans la première, la suite de sons /pəti/ est d'abord reconnue comme étant le mot "petit" et les sons qui suivent (ex. /ta/) activent tous les mots qui commencent de cette façon. Mais cette manière de faire empêchera (ou retardera) la reconnaissance du mot "ami". Dans la deuxième possibilité, la suite de sons /pətit/ est d'abord reconnue comme étant soit "petite" soit "petit" (il faudrait dans ce cas inclure une disposition spéciale pour les consonnes de liaison), et les sons suivants /am/ activent alors les mots qui commencent ainsi. Or, le problème

maintenant est que la reconnaissance de "tamis" sera soit empêchée, soit retardée. Par conséquent, nous voyons mal comment le modèle de la cohorte peut rendre compte de la reconnaissance de mots enchaînés.

Nous avons besoin d'un modèle qui soit moins séquentiel (qui ne stipule pas que la reconnaissance se fait mot par mot, un mot après l'autre) et qui permette l'activation de candidats à partir de plusieurs endroits dans la chaîne sonore. De plus, le modèle doit accepter la reconnaissance de plusieurs mots à la fois. Malgré certaines critiques qui lui sont faites (organisation temporelle, absence de liens avec les niveaux supérieurs, etc.), le modèle TRACE (McClelland & Elman, 1986) semble capable de rendre compte du traitement d'une suite ambiguë qui découle de la présence d'un enchaînement (avec quelques légères modifications éventuellement).

Prenons à nouveau l'exemple de /pətitami/. L'information du début de la suite active des candidats qui correspondent aux éléments acoustico-phonétiques perçus. Comme nous pouvons le voir sur la Figure 3 (page suivante), lorsque nous arrivons à la fin de /pəti/ (flèche no. 1 sur le diagramme), un certain nombre de candidats sont déjà actifs : "petit", "petite", "petite-fille", "petit-beurre", etc. Bien que TRACE ne tienne pas compte de la fréquence, nous avons arrangé les candidats selon leur fréquence dans l'exemple, les plus fréquents en haut de la liste, les moins fréquents en bas. "Petit-gris" sera donc moins actif que "petit" ou "petite". Au fur et à mesure que l'information phonétique continue d'arriver, de nouveaux candidats sont activés alors que d'autres perdent de leur activation (ils sont barrés dans le diagramme). Quand on arrive à la fin de la séquence /pətita/ (flèches 2a, 2b et 2c), "petit" et "petite" sont clairement actifs (flèche 2a) et "tamis" et "ami(e)" entrent en compétition (flèches 2b et 2c). Si ces deux concurrents étaient de même fréquence, seul le contexte pourrait les départager, mais dans notre exemple, "ami" étant beaucoup plus fréquent que "tamis", il sera plus actif et donc inhibera probablement "tamis". (Pour que "tamis" soit reconnu à la longue, il lui faudra l'aide du contexte). Enfin, lorsqu'on atteint la fin de la séquence /pətitami/ (flèches 3a et 3b), les candidats qui demeurent sont ceux qui correspondent directement à celle-ci ("petit" et "petite", "ami", "amie" et "tamis") ainsi que ceux qui seraient éventuellement possibles avec une continuation adéquate ("tas" si les mots suivants étaient

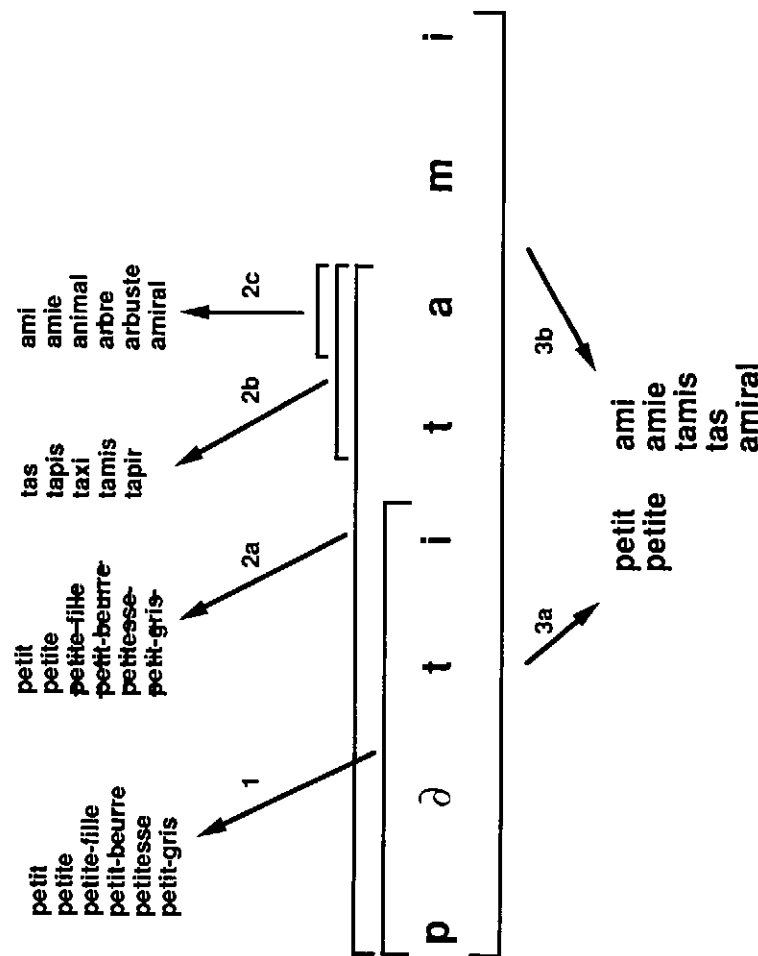


Figure 3. Le déroulement de la reconnaissance de la suite /pətitami/

"miniature", "misérable", "mis", etc. ; "amiral" si la suite se poursuivait avec /ral/, etc.). Ce qui est intéressant dans l'exemple donné est que "ami" et "amie" sont les premiers candidats (pour des raisons de fréquence) mais également que "tamis" reste candidat, certes moins actif, tant que d'autres informations (ascendantes ou descendantes) ne viennent les départager.

Par conséquent, plus l'information acoustico-phonétique permet de distinguer entre les candidats, moins il y aura de compétition entre eux et plus rapidement un seul candidat sortira du lot. C'est probablement ce qui se passe avec les suites qui appartiennent aux catégories E-cv et E-cc. Dans l'exemple de /ʃakurs/ ("chaque ours" ou "chaque course"), si l'indice acoustique au niveau du /k/ est en faveur de "chaque ours", alors "ours" sera plus activé, "course" le sera moins, la compétition sera moins intense entre les deux et "ours" sera reconnu plus rapidement. La rapidité de l'accès de mots enchaînés dépendra donc, en dehors d'informations descendantes, du degré d'ambiguïté de l'information acoustico-phonétique et, dans une moindre mesure, de la fréquence des concurrents en question.

Notons enfin que cette manière de considérer la reconnaissance des mots avec enchaînement permet d'éviter la présence de consonnes d'enchaînement au niveau de la représentation lexicale. En effet, on pourrait imaginer que l'entrée d'un mot qui commence avec une voyelle soit représentée par l'ensemble des suites phonétiques possibles dont celles avec liaison (ex. l'entrée "ami" correspondrait non seulement à "ami" mais également à "zami", "tami", "rami" etc.). Mais cela accroîtrait la taille du lexique de manière considérable (il faudrait que toutes les possibilités y soient présentes) et intégrerait au niveau de la représentation un phénomène avant tout de production. Nous préférons donc l'explication présentée ci-dessus dans laquelle les éléments enchaînés sont reconnus lorsque les informations ascendantes et descendantes ont permis l'activation suffisamment forte d'une suite de deux représentations lexicales.

5. Bibliographie

- BESSON, C. & F. GROSJEAN (1994): *L'effet de l'enchaînement sur la reconnaissance des mots dans la parole continue*, Manuscrit, Laboratoire de traitement du langage et de la parole, Université de Neuchâtel.
- BRADLEY, D. & B. DEJEAN DE LA BÂTIE (1990): *Resolving Word Boundaries in Spoken French*, Manuscrit non publié.
- MARSLEN-WILSON, W. (1987): "Functional parallelism in spoken word-recognition", *Cognition*, 25, 71-102.
- MARSLEN-WILSON, W. & A. WELSH (1978): "Processing interactions during word-recognition", *Cognitive Psychology*, 10, 29-63.
- MATTER, J. (1986): *A la recherche des frontières perdues*, Amsterdam, De Werelt.
- McCLELLAND, J. & J. ELMAN (1986): "The TRACE model of speech perception", *Cognitive Psychology*, 18, 1-86.